

UC Merced

Proceedings of the Annual Meeting of the Cognitive Science Society

Title

Perception of Phonemically Ambiguous Spoken Sequences in French

Permalink

<https://escholarship.org/uc/item/423956tf>

Journal

Proceedings of the Annual Meeting of the Cognitive Science Society, 27(27)

ISSN

1069-7977

Authors

Schaeken, Walter
Schroyens, Walter J.

Publication Date

2005

Peer reviewed

Perception of Phonemically Ambiguous Spoken Sequences in French

Anne-Laure Schaegis (al.schaegis@infonie.fr)

Elsa Spinelli (elsa.spinelli@upmf-grenoble.fr)

Université Pierre Mendès France

Laboratoire de Psychologie et NeuroCognition

BP 48, 38040 Grenoble Cedex 9, FRANCE

Pauline Welby (welby@icp.inpg.fr)

Institut de la Communication Parlée, CNRS UMR 5009

Institut National Polytechnique de Grenoble, Université Stendhal

46, avenue Félix Viallet, 38031 Grenoble Cedex 1, FRANCE

Abstract

Because the speech signal is continuous, listeners must segment the speech stream in order to recognize words. Due to elision, some spoken utterances in French are phonemically ambiguous (e.g., *C'est l'affiche* 'It's the poster' vs. *C'est la fiche* 'It's the sheet', both [sɛlafɪʃ]), and correct segmentation is necessary for recognition and comprehension. The aim of this study was to assess if listeners discriminate and identify such phonemically ambiguous utterances. In Experiments 1 and 2, an ABX paradigm was used for a discrimination task. The observed accuracy shows that listeners succeeded in discriminating between the two ambiguous stimuli, with identical or different tokens of those stimuli. In Experiment 3, in a forced choice task, listeners were able to retrieve the correct segmentation and correctly identify such ambiguous stimuli. Acoustic analyses have identified some of the acoustic differences between members of the pairs (*l'affiche* vs. *la fiche*). These differences are likely to be used by listeners during word segmentation.

Keywords : Segmentation, Speech, Word Recognition.

Introduction

Unlike written language, where words are separated by blank spaces, there are no clear word boundaries in spoken language. This means that a given stretch of speech can be consistent with multiple lexical hypotheses, and that these hypotheses can begin at different points in the input. In processing the speech stream, the listener is therefore routinely confronted with temporary ambiguities. Thus in the French *son chat potelé* [sɔ̃ʃapɔtlɛ] 'his/her plump cat', the recognition system must select between competing hypotheses like *son chat* [sɔ̃ʃa]... 'his/her cat' and *son chapeau* [sɔ̃ʃapo]... 'his/her hat' which, to a first approximation, are equally supported by segmental information. Hence, in order to recognize words in spoken language, listeners must segment the speech stream into discrete word units.

How do listeners accomplish this task? Segmentation could be based on several sources of information contained in the speech signal. Listeners could

then exploit regularities associated with word beginnings and ends in segmenting the speech stream. One source of information could come from language-specific metrical structure. In English, for example, most content words start with a strong syllable. Cutler and Norris (1988) proposed the Metrical Segmentation Strategy (MSS), according to which listeners exploit such prosodic probabilities to segment speech. According to the MSS, lexical access is initiated at each strong syllable the listener encounters. In a word-spotting experiment, Cutler and Norris (1988) showed that CVCC words like *mint* [mɪnt] are easier to detect in Strong-Weak sequences (e.g., [mɪn.tʔf], where the first syllable is stressed) than in Strong-Strong sequences (e.g., [mɪn.tef], where both syllables are stressed). In the latter sequences, detection was hypothesized to be slowed by misalignment with the syllabic boundary before the second stressed syllable. Other prosodic cues have been shown to play a role in segmentation. For example, in French, the last syllable of a prosodic phrase is lengthened and has special prominence. Bacri and Banel (1994) showed that listeners could exploit this pattern in word segmentation. Given ambiguous sequences [ba.gaʒ] (*bagage* 'luggage' or *bas gage* 'low pledge'), listeners were more likely to hear one word (*bagage*) when the second syllable was lengthened and two words (*bas gage*) when the first was lengthened. This finding corresponds to the expectation that a phrase-final syllable will be lengthened (and that a phrase boundary will not occur in the middle of a word). More recently, Welby (2003a/b) showed that French listeners could use the presence of an optional rise in fundamental frequency (f_0) or even a simple "elbow" in the f_0 curve as a cue to a content word beginning. Listeners interpreted nonsense sequences like [me.la.mõ.din] as a single nonword *mélamondine* when the f_0 rise began at the first syllable ([me]), and two words when it began at the second syllable ([la]) *mes lamondines* 'my lamondines'.

Another source of information could come from the phonotactic rules of a given language. For example, if a certain phone sequence is not a possible syllable onset cluster in a given language (e.g., [mʁ] in French or [mr] in

Dutch), there is a good possibility that there is a word boundary between the two segments (McQueen, 1998).

Lower level analyses of the speech signal could also help listeners in segmentation. Nakatani and Dukes (1977) examined segmentally ambiguous English sequences such as *known ocean/no notion*, and found that there were acoustic cues for juncture at the beginning and (occasionally) at the end of words – glottal stops, laryngealization, aspiration of voiceless stops. Quené (1992), with a forced choice task, showed that listeners could exploit durational cues to detect word boundaries in pairs of Dutch words like *diep in* ‘deep in’/ *die pin* ‘that pin’. Dumay, Content and Frauenfelder (1999) examined whether such cues could be used for online segmentation. With a word-spotting task, they showed that subjects detected the French word *tante* ‘aunt’ [tāt] more rapidly in the nonsense sequence *tantrou* [tātʁu] when *tantrou* was extracted from the sequence *tante roublarde* ‘sly aunt’ ([tã.t#ʁu.blɑʁd]) in which the coda [t] of *tante* is resyllabified as an onset consonant through the process of enchainment, than when it was extracted from *temps troublant* ‘troubling time’ ([tã.#tʁu.blɑ̃], in which there is no resyllabification. Measurements showed phonetic differences between the conditions. For example, for critical sequences like [tātʁu] in the enchainment condition, the durations of both the pre-boundary vowel and the liquid were greater. The differences in reaction times show that listeners were sensitive to these differences and could use them in segmentation. Moreover, it seems that in many languages, word-initial consonants tend to be longer than consonants which are syllable- but not word-initial (Fougeron & Keating, 1997; Gow & Gordon, 1995; Oller, 1973, Fougeron, 2001). It has also been argued that durational differences in word-initial position (along with other acoustic cues) signal the fact that speakers strengthen their articulation of segments at the edges of prosodic domains (Cho & Keating, 2001; Fougeron, 2001; Fougeron & Keating, 1997). Evidence from Welby (2003a/b) suggests that these durational differences are sensitive to a word’s status as a content word or a function word. The onset consonant of the first syllable of a nonsense sequence like [me.la.mõ.din] was longer when the syllable was the first syllable of a pseudo-content word than when it was a function word syllable.

The aim of the current study is to assess if French listeners segment and disambiguate between phonemically ambiguous spoken phrases. One potential problem in creating phonemically ambiguous stimuli is that it is often accompanied by a change in the syntactic structure of the phrase (*diep in/ die pin; no notion/ known ocean*). In French, the vowel of the definite article *le* or *la* is elided before vowel-initial words (with the exception of those words beginning with so-called aspirated vowels). Thus it is easy to create ambiguous pairs of stimuli that do not differ in syntactic structure. In addition, the disambiguation of

segmentally ambiguous strings to which the process of elision can give rise and the acoustic correlates of this process have rarely been studied (in contrast to liaison and enchainment, which have received much attention in recent work). We present three experiments in which French listeners listened to ambiguous sequences composed of a definite article followed by a noun presented after a neutral context *c’est* (‘it is’) like those in (1). In Experiments 1 and 2, they performed a discrimination task and in Experiment 3, an identification task.

- (1) a. C’est la fiche. [sɛlafɪʃ] ‘It’s the sheet.’
b. C’est l’affiche. [sɛlafɪʃ] ‘It’s the poster.’

Experiment 1

Method

Participants. Twenty-eight undergraduate students from the University Pierre Mendès France, Grenoble II participated in this experiment for course credit. They were all native speakers of French and reported no hearing impairment. Participants for all experiments in the study were drawn from the same population, although none participated in more than one experiment.

Stimuli and Procedure. Thirty pairs of phonemically ambiguous sequences composed of a definite article followed by a noun (e.g., *la fiche / l’affiche*) were selected (ambiguous set). The noun in one member of the minimal pair was vowel-initial (e.g., *affiche*) and the noun in the other member of the pair was consonant-initial (e.g., *fiche*). The members of the pairs were matched in frequency (47.8 occurrences per million words for the consonant-initial set and 18.4 occurrences per million words for the vowel-initial set, $t(29)=1.24$, $p>.20$; frequency given by the French database *Lexique* (New, Pallier, Ferrand, & Matos, 2001).¹ As a control, 30 pairs of non-phonemically ambiguous sequences composed of a definite article followed by a noun such as *le bain* ‘the bath’ / *le pain* ‘the bread’ were also selected (unambiguous set). The two members of these pairs differed by only one phonetic feature (e.g., voicing). The 60 sequences were recorded by a naive speaker (21-year old female native speaker of French from Grenoble) in pseudo-cleft carrier sentences. An example is given in (2).

- (2) a. C’est la fiche qui me manque.
b. C’est l’affiche qui me manque.
‘It’s the poster/sheet that is missing.’

The experimental sentences were recorded twice in order to provide two tokens (two sound files, e.g. two productions of *la fiche*) of each sequence, and mixed with 320 filler sentences. The sentences were recorded onto a Korg D1200 digital recorder using a Shure SM10A headworn microphone at 44.1 kHz, then downsampled to 22.05 kHz. Experimental sequences composed of an article and a noun preceded by the neutral context *c’est* (e.g., *c’est la fiche*)

¹ The apparent difference in frequency between the vowel and consonant-initial sets is due the high frequency of one item (*vie*, 712 occurrences per million words).

were labeled, then excised from the carrier sentences using Praat software (Boersma & Weenink, 2004) and two Praat scripts written to semi-automate the process. The experimental sequences (ambiguous set and unambiguous set) were presented to the subjects for a discrimination task using the ABX paradigm. Stimuli were presented through Sennheiser headphones HD 212Pro at a comfortable listening level. Subjects were informed that they would hear two different sequences, then a third one that would either equal the first or the second one. They were informed that some sequences were ambiguous and would to a certain extent sound the same (they were given an example of ambiguity due to elision, i.e., *l'amer*, 'the bitter one', *la mère*, 'the mother'). They were asked to decide whether the third sequence was the same as the first one (i.e. contained the same "word") or the same as the second one, by pressing one of the two response buttons, and their responses were collected. Stimuli were counterbalanced in two experimental lists and two blocks. The target X in the ABX paradigm was the same token (token¹) as either A or B. Hence, for both ambiguous and unambiguous sets, subjects were presented to A¹B A¹, B A¹A¹ in one block (e.g., *la fiche¹*, *l'affiche¹*, *la fiche¹*; *la mie¹*, *l'amie¹*, *l'amie¹*) and A B¹B¹, B¹A B¹ in another block (e.g., *la fiche¹*, *l'affiche¹*, *l'affiche¹*; *la mie¹*, *l'amie¹*, *la mie¹*). The order of block presentation was counterbalanced, and the order of item presentation was randomized within each block. The experiment was controlled by E-Prime Software (E-prime Psychology Software Tools Inc.; Pittsburgh, USA). The experimental session lasted approximately 25 minutes.

Results and Discussion

Percentages of correct responses in the ambiguous and unambiguous sets are presented in Figure 1. Results showed that participants were above chance (50%) in discriminating between ambiguous stimuli ($t(27)=15.05$; $p<.001$) and between unambiguous stimuli ($t(27)=101.9$; $p<.001$). In addition, participants showed better performances at discriminating between unambiguous stimuli than ambiguous stimuli ($t(27)=12.1$; $p<.001$).

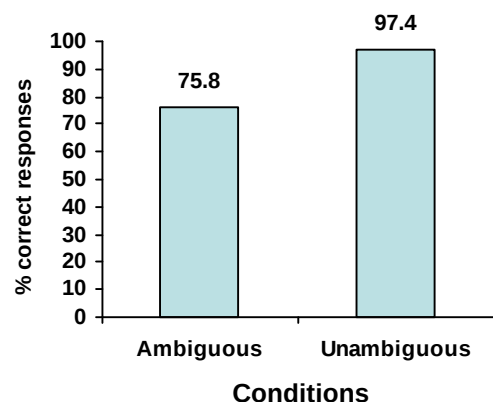


Figure 1: Percentages of correct responses in the ambiguous and unambiguous sets in Experiment 1.

The results of Experiment 1 show that listeners managed to discriminate between ambiguous sequences like *c'est la fiche* and *c'est l'affiche*. We hypothesize that although the two utterances are phonemically identical, there are acoustic differences between them so that vowel-initial and consonant-initial lexical candidates (*fiche*, *affiche*) can be differentially perceived depending on the intended segmentation. Note however that if subtle acoustic cues are associated with each segmentation, such cues are not powerful enough to totally disambiguate the sequences since listener performance in the ambiguous set does not meet that in the unambiguous set. In Experiment 1, the target X in the ABX paradigm was the same token (same sound file) as either A or B. Hence, it is possible that participants based their responses upon a non-linguistic comparison (e.g., non-relevant characteristics in the sound files). In Experiment 2, we used two different tokens. The target X was therefore a different token (token²) of either A or B. We hypothesized that if acoustic cues are used to discriminate between ambiguous sequences, such cues should be available from one production to another.

Experiment 2

Method

Participants. Twenty-eight participants took part in the experiment.

Stimuli and procedure. The same 30 pairs of phonemically ambiguous sequences (*C'est la fiche* / *l'affiche*, ambiguous set) and 30 pairs of non-phonemically ambiguous sequences (*C'est le bain* / *le pain*, unambiguous set) used in Experiment 1 were used in Experiment 2. The procedure paralleled that of Experiment 1 except that stimulus X in the ABX presentation was a different token of either type A or B. Hence, for both ambiguous and unambiguous sets, subjects were presented to A¹B A², B A¹A² in the first block (e.g., *la fiche¹*, *l'affiche¹*, *la fiche²*; *la mie¹*, *l'amie¹*, *l'amie²*) and A B¹B², B¹A B² in the second block (e.g., *la fiche¹*, *l'affiche¹*, *l'affiche²*; *la mie¹*, *l'amie¹*, *la mie²*).

Results and Discussion

Percentages of correct responses in the ambiguous and unambiguous sets are presented in Figure 2. As in Experiment 1, results showed that participants were above chance both in discriminating between ambiguous stimuli ($t(27)=8.6$; $p<.001$) and in discriminating between unambiguous stimuli ($t(27)=94.8$; $p<.001$). As in Experiment 1, participants showed better performance at discriminating between unambiguous stimuli than ambiguous stimuli ($t(27)=16.6$; $p<.001$).

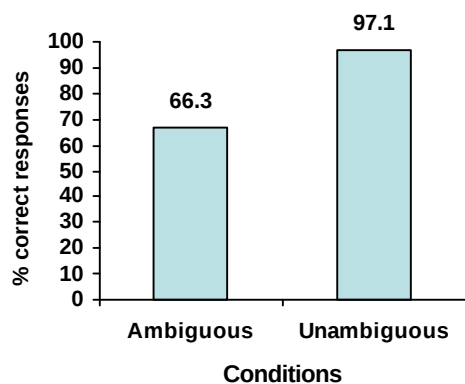


Figure 2: Percentages of correct responses in the ambiguous and unambiguous sets in Experiment 2.

The results of Experiment 2 show that despite the use of two different tokens for the comparison in the ABX paradigm, listeners still managed to discriminate between the two ambiguous sequences (*C'est la fiche* vs. *C'est l'affiche*). Hence, if acoustic cues can be used to differentiate vowel-initial segmentation from consonant-initial ones, such cues remain across different tokens of the same type. If such cues can be exploited to discriminate between the two sequences, one may ask if they are associated with each lexical entry. In other words, can listeners use them to identify the sequences? In Experiment 3, participants performed an identification task.

Experiment 3

Method

Participants. Twenty-eight participants took part in this experiment.

Stimuli and procedure. The same stimuli were used in this experiment as in Experiments 1 and 2 (30 pairs in the ambiguous set and 30 pairs in the unambiguous set). Participants were presented with one member of the pair (e.g., *C'est la fiche*) and were asked to identify the noun by making a forced choice between the two possible nouns (e.g., *fiche* vs. *affiche*). The experiment was controlled by E-Prime Software. Stimuli were presented through Sennheiser headphones HD 212Pro at a comfortable listening level. Stimuli were counterbalanced in two experimental lists so that each subject heard all conditions but only one member of each pair. Within each list, the order of item presentation was randomized. Participants were asked to respond by pressing one of the two response buttons, and their responses were collected.

Results and Discussion

Percentages of correct identification in the ambiguous and unambiguous sets are presented in Figure 3. Results showed that participants were above chance in identifying ambiguous stimuli ($t(27)=13.2$; $p<.001$) and unambiguous stimuli ($t(27)=174.6$; $p<.001$). In addition, they showed

better performance at identifying unambiguous stimuli than ambiguous stimuli ($t(27)=12.05$; $p<.001$).

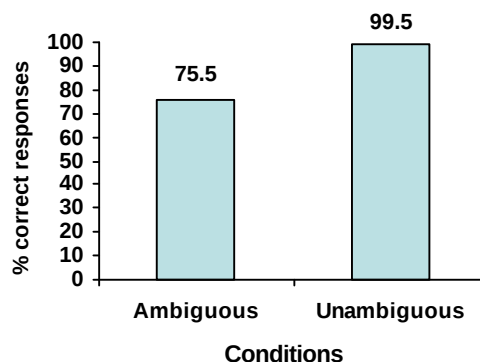


Figure 3: Percentages of correct responses to the identification task in Experiment 3.

Acoustic analyses

Acoustic analyses show significant durational differences between the two members of a pair (e.g., *la fiche*/ *l'affiche*). The mean duration of the first syllable of the phonemically ambiguous target sequences ([*la*] for all the experimental items) was greater for items like *l'affiche*, where the [*l*] of the definite article is resyllabified as the onset of the content word (151.23 ms), than for items like *la fiche*, where the [*l*] is not content word-initial (146.4 ms) ($t(59)=2.23$; $p<.05$). The [*l*] onset of the first syllable was also significantly longer when it was content word-initial. Onset durations were on average 69.2 ms for items like *l'affiche*, and 59.0 ms for items like *la fiche* ($t(59)=4.88$; $p<.001$). There was also a small, but statistically significant difference in duration of the first syllable vowel, with longer vowels in content word-initial syllables (87.4 ms vs. 82.0 ms, $t(59)=2.64$, $p<.05$). There were also durational differences in the second syllable of target sequences (e.g., [*fiʃ*] in [*la.fiʃ*]). This syllable was longer in content word-initial position (as in *la fiche*) than in non-initial position (as in *l'affiche*) (498.0 ms vs. 455.2 ms, $t(59)=6.62$; $p<.001$). Both the syllable onset and the vowel ([*f*] and [*i*] in [*la.fiʃ*]) were significantly longer in content word-initial syllables than in non-initial syllables (for onset consonants: 216.5 ms vs. 197.1 ms, $t(59)=6.94$; $p<.001$); for vowels: 229.4 ms vs. 211.7 ms, $t(59)=4.69$; $p<.001$).

Formant analyses also revealed differences. The F2 of the first syllable vowel (measured at the vowel midpoint) was significantly lower (1808 Hz) when the vowel was in a vowel-initial word (*l'affiche*) than when it was part of the determiner *la* (*la fiche*) (1834 Hz), indicating a more forward tongue position for the vowel of the determiner ($t(59)=2.92$; $p < 0.01$). There were no significant differences in F1. Unlike the vowel of the first syllable (always /a/), the vowel of the second syllable varied (e.g., /i/ in *la fiche*/ *l'affiche*, /u/ in *la toux* 'the cough')

l'atout 'the advantage'); we therefore separated the items into groups depending on vowel height (high/ mid-high and low/ mid-low) for the F1 analyses and vowel backness (front/back) for the F2 analyses. F1 was significantly higher (819 Hz) for low/mid-low vowels in consonant-initial words (e.g., [a] in *la marre*, 'the pond'), than for those in vowel-initial words (768 Hz, e.g., [a] in *l'amarre* 'the nautical rope'), indicating a lower tongue position for low/mid-low vowels in consonant-initial words ($t(16)=2.84$; $p<.05$). There was no difference in F1 for high/mid-high vowels. There were, however, significant differences in F2. F2 was significantly higher for front vowels (1961 Hz) in items like *la marre* than those like *l'amarre* (1921 Hz) ($t(35)=2.13$; $p<.05$), indicating that front vowels were articulated with a more forward tongue position for consonant-initial words. Similarly, back vowels are articulated with a more back tongue position in consonant-initial words (e.g. *la location* 'the rental') than in vowel-initial ones (e.g. *l'allocation* 'the grant'), as shown by a significant difference in F2 (1148 Hz vs. 1179 Hz, respectively, $t(14)=2.86$; $p<.05$).

Evidence of another potential difference comes from devoicing of the onset consonant of the second syllable. In this position, only five items (thus 10 tokens) had phonemically voiced obstruent onsets (e.g., *la joue* 'the cheek'/ *l'ajout* 'the addition', /la.ʒu/). The speaker showed some devoicing of these fricatives, most often in the condition in which they were word-medial (e.g. the /ʒ/of *l'ajout* 'the addition' is completely devoiced ([ʃ]) for both tokens). Of course, no firm conclusions can be drawn from so few tokens.

In summary, differences in formant structure and duration, when found, generally go in the direction of clearer, more canonical pronunciations of content word-initial syllables (*la location*, *l'allocation*). The one exception is the surprising result for F2 in the first vowel.

Finally, there are clear intonational differences between the two conditions. There is often a rise in fundamental frequency beginning at the left edge of the first content word syllable. This rise is particularly apparent in sequences with all sonorant segments, as shown in Figure 4. In both panels of the example, there is a rise beginning at the first content-word syllable. In the top panel, the rise begins at [la], the first syllable of *l'amie* 'the friend'; in the bottom panel, the rise begins a syllable later, at [mi], the first (and only) syllable of *mie* 'crumb'.

Fougeron, Bagou, Stefanuto & Frauenfelder (2002), Spinelli, McQueen & Cutler (2003) *inter alia* have found differences in duration and formant structure in cases of enchainment and liaison. It seems likely that the same types of cues would be used in disambiguating ambiguous sequences, no matter which process (enchainment, liaison, elision) gives rise to the ambiguity. However, the potential acoustic cues we have discussed here were observed in a single speaker. Given speaker-specific differences found by Fougeron et al. (2002), it is possible that other speakers may show somewhat different patterns. Nevertheless, we are

confident of the patterns found for this speaker; she was naive to the goals of the recording and subsequent experiments, and the critical ambiguous items constituted only 14% (60 out of 440) of the items recorded.

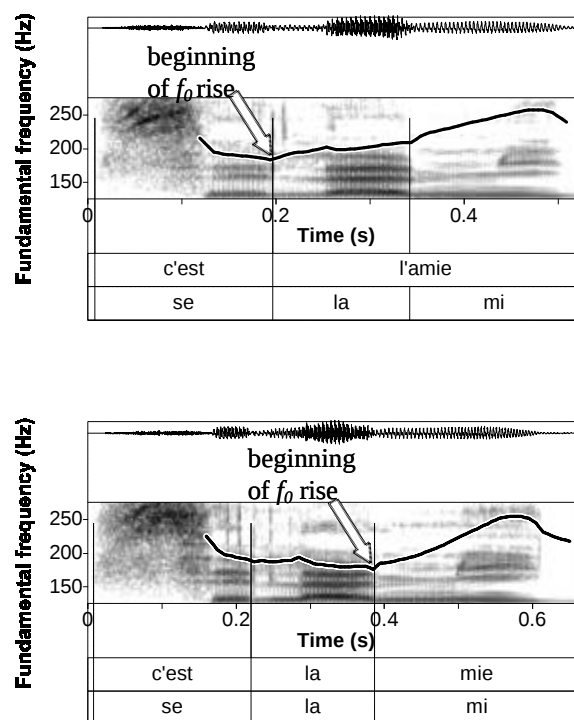


Figure 4: Waveform, spectrogram, and fundamental frequency curve of an experimental pair. (top) *C'est l'amie* 'It's the friend' (bottom) *C'est la mie*. 'It's the crumb'.

General discussion

In three experiments, we have found that listeners were above chance level in discriminating and identifying phonemically ambiguous sequences such as [selafiʃ]. Hence listeners managed to retrieve the correct segmentation when presented with such sequences. Our preliminary acoustic analyses showed that there are some cues associated with each intended segmentation that listeners could use in order to decode the speech stream. Further experiments are planned to examine whether the manipulation of these cues influence subject's segmentation performance. Our present results are in line with previous studies showing that listeners were sensitive to subtle acoustic differences and could use them in segmentation (Quené, 1992, *inter alia*). If specific cues are indeed present in the speech to signal intended segmentation, it remains to be seen whether these cues are used online to modulate lexical activation of competing candidates. Some research has shown that subphonemic differences can influence processing at the lexical level. Andruski, Blumstein and Burton (1994), for example, examined the effect of the alteration of VOT on

the activation of English words beginning with stop consonants. They found that words beginning with voiceless stops were more strongly activated when the input words had normal VOTs than when the VOTs had been shortened. In a similar vein, Marslen-Wilson and Warren (1994) showed that lexical and phonemic decisions were slower to words and nonwords which contained mismatching subphonemic information than to words and nonwords which did not contain mismatching information. For example, response latencies were longer to the word *job*, when the *jo* was spliced from a token of *jog* than when the *jo* was spliced from another token of *job*. These studies suggest that subphonemic information is passed up to the lexicon. Recently, Spinelli, McQueen and Cutler (2003) examined how lexical ambiguities in liaison contexts in French are processed on line. In French, the final /r/ of *dernier* 'last' is not pronounced when the following word begins with a consonant, but is pronounced when the following word begins with a vowel. Due to this process of liaison, some sequences become phonemically ambiguous (e.g., *dernier oignon* 'last onion' and *dernier rognon* 'last kidney'). In a cross-modal priming study, facilitation was found for both types of target (vowel-initial and consonant-initial) when they matched the speaker's intended segmentation, but this facilitation was weaker when they mismatched the intended segmentation. This suggests that in spite of the apparent homophony between the sequences, the appropriate lexical entry was activated. Moreover, acoustic analyses showed that consonants in the liaison environments were shorter than underlyingly word-initial consonants (e.g., [ʁ] in *dernier oignon* vs. *rognon*), suggesting that word recognition was influenced by subphonemic cues to the words that speakers intend. Taken together, these studies suggest that, during spoken word recognition, fine-grained differences in the speech signal influence processing at the lexical level and thus modulate lexical selection. Therefore, considering the subjects performance in segmenting ambiguous sequences like [selafiʃ], we would predict that the acoustic cues associated to the intended segmentation might be used on line to modulate the activation the two competitors (*fiche* and *affiche*). Further experiments with an online priming paradigm are planned to examine this specific question.

Acknowledgments

This research has been funded by an *Action Concertée* of the CNRS *Système Complexe en SHS*, *Pati papa ? modélisation de l'émergence d'un langage articulé dans une société d'agents sensori-moteurs en interaction*. We thank Laure Deguelle for recording our experimental stimuli.

References

Andruski, J.E., Blumstein, S.E., & Burton, M. (1994). The effect of subphonetic differences on lexical access. *Cognition*, 52, 163-187.

Banel, M.-H. and Bacri, N. (1994). On metrical patterns and lexical parsing in French. *Speech Comm*, 15, 115-126.

Boersma, P. & Weenink, D. (2004). Praat : doing phonetics by computer (Version 4.2.20) [Computer program]. Retrieved from <http://www.praat.org>

Cho, T., & Keating, P.A. (2001). Articulatory and acoustic studies on domain-initial strengthening in Korean. *Journal of Phonetics*, 29, 155-190.

Cutler, A. & Norris, D. (1988). The role of strong syllables in segmentation for lexical access. *Journal of Experimental Psychology: Human Perception and Performance*, 14, 113-121.

Dumay, N., Content, A. & Frauenfelder, U. H. (1999). Contribution de la structure syllabique de surface à la segmentation lexicale *Proceedings of the 2^{èmes} Journées d'Etude Linguistiques*, (Nantes), 105-109.

Fougeron, C. (2001). Articulatory properties of initial segments in several prosodic constituents in French. *Journal of Phonetics*, 29, 109-135.

Fougeron, C., Bagou, O., Stefanuto, M. & Frauenfelder, U. (2002). A la recherche d'indices de frontière lexicale dans la resyllabation. In the proceedings of XXIV *Journées d'Etude sur la Parole*, Nancy, 24-27 juin 2002.

Fougeron, C. & Keating, P.A. (1997). Articulatory strengthening at edges of prosodic domains. *JASA*, 101, 3728-3740.

Gow, D.W. Jr., & Gordon, P.C. (1995). Lexical and prelexical influences on word segmentation: Evidence from priming. *Journal of Experimental Psychology: Human, Perception and Performance*, 21, 344-359.

McQueen, J.M. (1998). Segmentation of continuous speech using phonotactics. *J. of Memory & Language*, 39, 21-46.

Marslen-Wilson, W., & Warren, P. (1994). Levels of perceptual representation and process in lexical access: Words, phonemes, and features. *Psychological Review*, 101, 653-675.

Nakatani L. H. & Dukes, K. D. (1977). Locus of segmental cues to word juncture. *JASA*, 62, 714-719.

New, B., Pallier, C., Ferrand, L., & Matos, R. (2001). Une base de données lexicales du français contemporain sur Internet: LEXIQUE. *L'Année Psychologique*, 101, 447-462.

Oller, D.K. (1973). The effect of position in utterance on speech segment duration in English. *JASA*, 54, 1235-1247.

Quené, H. (1992). Durational cues for word segmentation in Dutch. *J. of Phonetics*, 20, 331-350.

Spinelli, E., McQueen, J. & Cutler, A. (2003). Processing Resyllabified Words in French. *Journal of Memory and Language*, 48, 233-254.

Welby, P. (2003a). French Intonational Rises and their Role in Speech Segmentation. In *Proceedings of Eurospeech: The 8th Annual Conference on Speech Communication and Technology*, (Geneva), 2125-2128.

Welby, P. (2003b). *The slaying of Lady Mondegreen, being a study of French tonal association and alignment their role in speech segmentation*, PhD dissertation, The Ohio State University.