

New insights into American English V+/r/ sequences

María Riera & Joaquín Romero, Universitat Rovira i Virgili

Abstract

This paper presents an acoustic study of final V+/r/ sequences in American English stressed monosyllables. We provide experimental data to show the durational and spectral characteristics of the vowel, the consonant and the VC transition, we explain the presence of this transition in relation to the vowel and the consonant, and we examine the role of speaking rate. The results show the presence of a transitional vocalic element that varies significantly as a function of the vowel and speaking rate. They also show significant durational and spectral differences which can be interpreted as the result of VC coarticulation.

1. Introduction

1.1 Overview

The study presented in this paper forms part of a wider ongoing acoustic study that seeks a better understanding of the phonetic and phonological nature of final V+/l/ and V+/r/ sequences in American English stressed monosyllables by investigating the VC coarticulatory processes that take place in them. On the one hand, the present study expands on and replicates in part previous studies carried out by the authors (Riera & Romero 2006, 2007; Riera et al. 2009) in an attempt to gain new insights into the behavior of V+/r/ sequences in particular. On the other hand, the present study introduces innovative aspects related to participants, stimuli, segmentation procedures and measurements taken: the number of participants has been increased, the stimuli have been modified, a more objective method of segmentation and boundary identification has been applied and consonant (i.e., /r/) measurements have been included. In this study we provide experimental acoustic data to show the durational and spectral characteristics of the vowel, the consonant and the VC transition, we explain the presence of this transition in relation to the vowel and the consonant, and we examine the role of speaking rate.

1.2 Previous studies

Previous studies that have looked into V+/r/ sequences have focused on the schwa-like element that is often perceived in some of these sequences. Terms like *epenthetic schwa* (Warner et al. 2001), *excrecent schwa* (Gick & Wilson 2001, 2006) or *targetless schwa* (Browman & Goldstein 1992b) might be used to refer to this element. According to Gick & Wilson (2001, 2006), the perceptual presence of this element after high front vowels can be explained as the result of the tongue movement required in passing through a schwa-like configuration. Browman & Goldstein (1992b) make reference to the influence exerted by neighboring segments on what they call *targetless schwa*. Wells (2000) uses the term pre-r breaking¹ to refer to cases of schwa epenthesis in sequences containing high vowels, whereby monophthongs become diphthongs and diphthongs become triphthongs. Lavoie & Cohn (1999) state that monosyllables consisting of non-low tense pure vowels or diphthongs followed by a liquid can be pronounced with either one or two syllables. Hall (2003, 2006) distinguishes between schwa intrusion and schwa epenthesis/insertion. In her view, intrusive vowels are phonologically invisible, are inserted late in the phonological derivation, cannot act as syllable nuclei, do not add a syllable to the word and do not involve the addition of a vowel segment. Moreover, they are not likely to occur in the most marked types of CC clusters, tend to occur between heterorganic consonants, copy only over sonorants or gutturals and are either copy vowels or neutral and schwa-like in quality.

Riera & Romero (2006) provide an impressionistic analysis of V+/l/ and V+/r/ sequences by means of visual spectrographic observation in a preliminary descriptive study that relies on acoustic data from two speakers and considers the whole range of American English stressed vowels. The study acknowledges the presence of VC transitions in some of these sequences and of a variable schwa-like element which is not visually detectable to the same extent in all of the VC transitions. It also suggests a relationship between front versus back versus central vowels as well as between high and tense versus non-high and lax ones. No acoustic measurements are taken in this study. The role of speaking rate is evidenced only by the fact that VC transitions are more easily discernible in slow tokens than in fast ones. It is concluded that the presence of the transitional element is the result of a dynamic phonetic process of coarticulation rather than of a discrete phonological rule of epenthesis/insertion.

In experimental studies conducted by Riera & Romero (2007) and Riera et al. (2009), durational and spectral measurements reveal differences between the

¹ Wells (2000) also uses the term pre-/l/ breaking to refer to cases of schwa epenthesis in V+/l/ sequences.

schwa-like element and canonical schwa² as well as variability in the schwa-like element as a function of both the preceding vowel and speaking rate. The formant values of this element are significantly different from those of canonical schwa and tend to resemble more those of the preceding vowel the faster the speaking rate in both the V+/l/ (Riera & Romero 2007) and the V+/r/ (Riera et al. 2009) sequences. The phenomenon under analysis is regarded in these studies as a generalized process affecting all contexts (i.e., all stressed vowels + /l/ or /r/), rather than, for example, only high vowels, as has been implied by previous studies (Gick & Wilson 2001, 2006; Lavoie & Cohn 1999; Riera & Romero 2006; Wells 2000). As in Riera & Romero (2006), coarticulation, rather than epenthesis/insertion, is favored. The segmentation procedure in these studies is based solely on the observation of acoustic waveforms and spectrograms as well as on the auditory corroboration by the experimenters. These studies rely on acoustic data from only one (Riera & Romero 2007) or two (Riera et al. 2009) speakers. Durational, F1 and F2 measurements are obtained in both studies, but F3 measurements are obtained only for the V+/r/ sequences. Measurements for the vowel and the transitional element only are obtained in both studies; neither of them includes consonant (i.e., /l/ or /r/) measurements and thus the behavior of the transitional element is explained only in terms of its relationship with the preceding vowel. Speaking rate (i.e., slow vs. fast) differences are considered in both studies. Thus, the current study expands on these previous findings by offering a more reliable methodological approach to segmentation and by providing data for the consonants as well as for a larger pool of subjects. Also, it provides measurements of the different parts of the sequences taken at midpoint rather than mean measurements of them, which was the measurement procedure used in previous studies.

1.3 The present study: Objectives and hypotheses

As mentioned above, the overall main objective of this study is to investigate the VC coarticulatory processes that take place in final V+/r/ sequences in American English stressed monosyllables. In order to do this, (i) we provide experimental data to show the durational and spectral (i.e., F1, F2 & F3) characteristics of the vowel, the consonant and the transitional vocalic (i.e., schwa-like) element, (ii) we explain the presence of this element in relation to the vowel and the consonant, and (iii) we determine the role of speaking rate (iiia) by looking for durational, F1, F2 and F3 variability in the vowel, the transitional element

² Canonical schwa refers to a lexically-licensed vowel that shows relatively stable spectral characteristics and is not usually subject to significant contextual variability, as in the first syllable of the word *ahead*.

and the consonant, and (iiib) by comparing F1, F2 and F3 mean values in the different contexts (i.e., each of the V+/r/ sequences) and the different rates (i.e., slow and fast).

The results of this study are expected to provide evidence for the existence of coarticulatory processes and to make manifest the extent of VC coarticulation in the V+/r/ sequences under study. By looking into the behavior of the vowel, the transitional element and the consonant in the sequences, and by looking at the influence exerted by both the vowel on the transitional element and the transitional element on the consonant, the phonetic, rather than phonological, nature of the transitional element will be revealed.

We hypothesize (i) that there will be significant durational, F1, F2 and F3 variability in the vowel, the transitional element and the consonant, across contexts, and as a function of speaking rate, and (ii) that the F1, F2 and F3 mean values of the different contexts will tend to resemble each other more in the slow-rate productions than in the fast-rate ones. This is expected to be especially the case for the vowel and the transitional element but not so much for the consonant. The greatest differences are expected to be particularly noticeable for F1 and F2 but less so for F3.

The hypotheses presented here regarding the coarticulatory nature of the V+/r/ transitions are in accordance with the approach to speech production and gestural organization illustrated by the theory of Articulatory Phonology (Browman & Goldstein 1986, 1989, 1990a, 1990b, 1992a; Goldstein & Fowler 2003). Articulatory Phonology offers a view of phonological organization based on articulatory gestures as primitive units that are responsible for both phonological invariance and phonetic variability and thus bridges the gap between the two levels of description. A key aspect of the theory for our study is the fact that it contemplates time as an intrinsic part of the description of gestures, therefore providing a much more plausible explanation for the coarticulatory variability caused by rate differences, than would be given by a theory based on discrete underlying segments.

2. Method

2.1 Speakers

The subjects that participated in the experiment were six native speakers of American English. Four were male and two female. They all had rhotic accents. Three had a western accent (California, Utah, Wyoming), one a mid-western accent (Wisconsin) and one an upper-southern accent (Tennessee).

The last speaker reported having lived in different parts of the US but self-identified her accent as being mid-western. Their ages ranged from 24 to 40. Four of them were temporarily living in Spain for a period of at least one year at the time of the recording; two had been living in Spain for over five years. Only one speaker had some specialized phonetic training; the rest had none. All the speakers were unaware of the purposes of the experiment prior to being recorded. Sex, type of accent, age, place of residency, contact with or knowledge of the Spanish language, and specialized phonetic training were not considered relevant factors to affect the purposes of our experiment in any negative way.

2.2 Stimuli

The target words that were selected for the experiment reported in this paper were seven English monosyllables containing final V+/r/ sequences (i.e., *fear*, *fair*, *par*, *pore*, *poor*, *hire* and *power*). Fifteen English monosyllables containing final V+/l/ sequences (i.e., *feel*, *bill*, *pale*, *fell*, *pal*, *Poll*, *Paul*, *hole*, *pull*, *pool*, *hull*, *furl*, *pile*, *howl* and *boil*) were also included as target words to be separately analyzed as part of the wider ongoing study. Fifteen distracters, consisting of C₁VC₂, where C₂ was one of /t/ or /d/, were included as well. These were the words *heat*, *fit*, *hate*, *vet*, *fat*, *hot*, *fought*, *vote*, *hood*, *food*, *hut*, *heard*, *hide*, *void* and *vowed*. All the target words and distracters were inserted in the carrier sentence *Say ____ for me again*. In order to minimize unwanted coarticulatory effects, C₁ was a non-lingual (unlike /r/ and /l/) and oral (like /r/ and /l/) consonant in the target words, the distracters and the word *for*.

2.3 Data collection

The six speakers performed two readings each of ten randomized repetitions of the carrier sentence containing the target words and distracters reported in the previous subsection. The first reading was performed at a slow speaking rate; the second at a faster one. The speaking rate variable was controlled for by presenting the slow-rate readings at four-second intervals separated by a three-second break every 20 sentences and the fast-rate readings at one-second intervals with a three-second break every five sentences. The readings took place in two different sessions separated by a 30-minute period. Each of the sessions was preceded by an instruction period and a trial period of 20 tokens, which were not used for the analysis. After the second session, the speakers were informed of the purposes of the experiment and were asked to fill out a questionnaire to provide some very general personal information relevant only to the purposes of the experiment. The data were recorded at a 44,100 Hz sampling rate directly into a

laptop computer using an M-Audio Nova condenser microphone, an M-Audio Firewire Solo mobile interface, and the Praat speech analysis software (Boersma & Weenink 2010), which was also used for the subsequent data analysis.

2.4 Data analysis

2.4.1 Segmentation procedure

From a segmental point of view, the V+/r/ sequences under study are considered to be composed of two elements only (i.e., a vowel followed by a consonant). However, in order to identify the transitional element in them, the sequences had to be divided into three parts, corresponding to the vowel, the transitional element and the consonant. In the case of sequences containing diphthongs, they were divided into four parts and it was the second element of the diphthong that was taken into account for the analysis.

Given the dynamic nature of the transitional element, and therefore the difficulties in identifying and delimiting it, we applied a first differentiation algorithm to the F1, F2 and F3 traces as identified by an automatic formant tracking routine in order to obtain velocity curves for each of these spectral events. This allowed us to automatically identify inflexion points in the formant traces that corresponded with the boundaries between the three portions of the signal under study and thus made it possible to isolate the transitional element. A Praat script was written to obtain these first derivative traces and identify the peaks of formant change given by velocity maxima and minima. These peaks were then taken as reference points for boundary placement.

Figure 1 illustrates the segmentation procedure. The upper part of the figure shows the acoustic wave and the spectrogram for the slow version of the /ir/ sequence in the word *fear* as produced by one of the speakers. In the lower part there are a series of tiers which provide information about F1 and F2 first derivative peaks of formant change indicating velocity maxima and minima (first and third tiers). These maxima and minima correspond to inflection points in the velocity trace and can, therefore, be identified with the beginning and end of specific events. The second tier in this lower part of Figure 1 shows the segmentation of the sequence into three parts (i.e., vowel, transitional schwa-like element and consonant). The vertical lines in this second tier are determined by observing where the broken lines in the F1, F2 and, if necessary F3, tiers fall, and then by deciding which of these lines correspond to the beginning and end of the different parts of the sequence. In the case exemplified here, one peak provided by the F2 derivative was chosen to mark the beginning of the schwa-like element, whereas one peak in the F1 derivative was chosen to mark its end. Because it was not necessary to rely on F3 derivative peaks, these are not shown

in the figure.

As might be inferred from the information provided in Figure 1, problems related to determining boundary placement often arise. In such cases, the automated procedure needs to be complemented by visual observation of waveform and spectrographic cues as well as by auditory corroboration. This is particularly necessary in the case of fast tokens and sequences containing low back vowels (i.e., /ɑ/ and /ɔ/). The objective segmentation procedure can then be considered to be more reliable than the subjective one only to a certain extent, but nonetheless reliable enough to the point that it allows for consistency in the segmentation procedure.

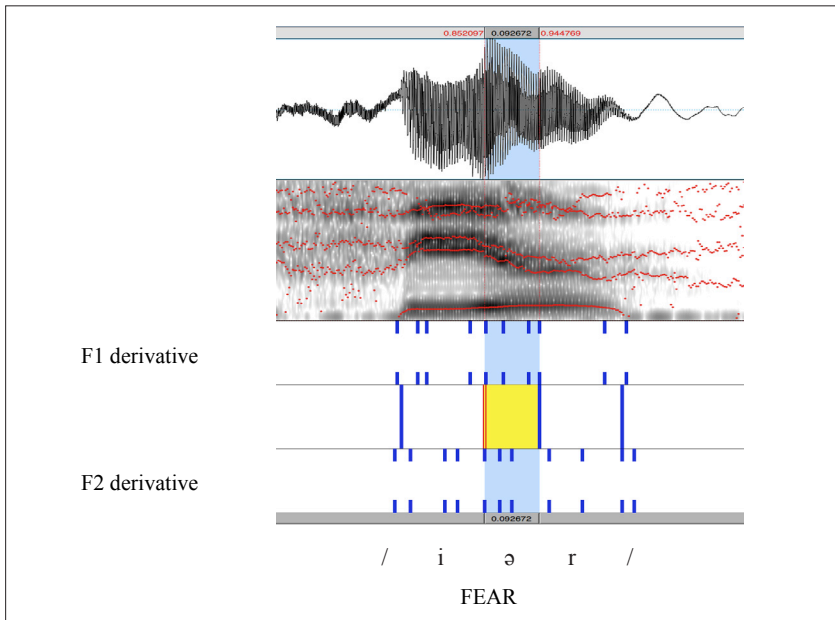


Figure 1 – Segmentation procedure for the /ir/ sequence corresponding to the slow version of the word FEAR as produced by one of the speakers. The vertical broken lines in the first and third textgrid tiers represent the F1 and F2 first derivative peaks of formant change indicating velocity maxima/minima. The second tier shows the /ir/ sequence segmented into three parts: vowel, transitional element and consonant.

2.4.2 Measurements

A Praat script was designed to extract midpoint duration, F1, F2 and F3 values for the vowel, the transitional element and the consonant. Mean values for each context (i.e., each of the V+/r/ sequences) and for each rate (i.e., slow and fast)

were then obtained and used for the statistical analyses. F1, F2 and F3 mean values were also used for comparisons between the slow and fast speaking rates.

2.4.3 Statistical analyses

Two-way factorial ANOVAs were performed to test for duration, F1, F2 and F3 overall variability in the vowel, the transitional element and the consonant. The independent variables were rate (i.e., slow and fast) and context (i.e., each of the V+/r/ sequences); the dependent variables were duration, F1, F2 and F3 mean values.

In the cases where interactions proved to be significant, independent one way ANOVAs for each of the two rates (i.e., slow and fast) were subsequently performed to confirm the variability shown by the two-way factorial ANOVAs, or to test for further variability, by looking at the two rates separately. The independent variable was context (i.e., each of the V+/r/ sequences); the dependent variables were duration, F1, F2 and F3 mean values.

3. Results

3.1 ANOVAs for variability

As mentioned above, the two-way factorial ANOVAs looked for duration, F1, F2 and F3 overall variability³ in the vowel, the transitional element and the consonant. Rate and context were the independent variables and duration, F1, F2 and F3 mean values the dependent variables. Significance level was set at $p < .01$. Significant differences were obtained in almost all cases. The results showing variability in the vowel were significant for all speakers, for rate, context and the interaction between rate and context, and for duration, F1, F2 and F3. The results showing variability in the transitional element and the consonant were significant for all speakers, for context, and for duration, F1, F2 and F3. They were non-significant in the following cases, which involve combinations of rate or the interaction between rate and context, and duration, F1, F2 or F3: Speaker 2 (rate, F2), Speaker 3 (rate*context, duration; rate, F3), Speaker 4 (rate, F3), and Speaker 6 (rate*context, duration; rate F1). The results showing variability in the consonant were non-significant in the following cases: Speaker

³ Here *variability in the vowel* does not refer to intra-token variability but rather to the comparison between the mean values for the vowels in the different contexts (i.e., *fear* vs. *hire* vs. *fair* vs. *par* vs. *pore* vs. *poor* vs. *power*). As expected, the results show that these vowels are indeed different. The reason why we have decided to make this seemingly obvious comparison is so that it can then be compared with the differences in the transitional element and thus show that the transitional element retains some of the variability of the vowel but is much more deeply affected by the lack of a specific articulatory target, as demonstrated by the significant differences across rates.

1 (rate*context, duration; rate, F1), Speaker 2 (rate, F2; rate*context, F3), Speaker 3 (rate*context, duration), Speaker 4 (rate*context, duration; rate*context, F1), Speaker 5 (rate*context, duration; rate, F2), and Speaker 6 (rate, F2).

The results of the separate one-way ANOVAs performed to confirm the variability shown by the two-way factorial ANOVAs, or to test for further variability, with context as the independent variable and duration, F1, F2 and F3 as the dependent variables, also yielded significant differences in almost all cases. Significance level was again set at $p < .01$. As with the two-way factorial ANOVAs, the results showing variability in the vowel were significant for all speakers, for duration, F1, F2 and F3, and for both rates. The results showing variability in the transitional element were significant for five speakers, for duration, F1, F2 and F3, and for both rates. The exception was Speaker 4, with non-significant differences for duration for both rates. The results showing variability in the consonant were significant for two speakers, for duration, F1, F2 and F3 for the slow rate. They were non-significant in the following cases: Speaker 1 (F2, slow), Speaker 2 (F3, slow), Speaker 4 (duration, slow), and Speaker 5 (duration, fast; F2, slow).

3.2 Means for comparisons between speaking rates and variability

Figures 2, 3 and 4 show scatter plots for mean F1, F2 and F3 values, respectively (i) for one speaker, (ii) for the seven V+/r/ contexts, (iii) for the vowel, the transitional element and the consonant, and (iv) for the slow-rate and fast-rate productions. Due to space constraints, the data from one speaker only will be used to exemplify what, as a general rule, applies to the other five speakers as well.

As can be observed in these scatter plots, formant values for the same target words show clear differences across speaking rates (i.e., compare slow and fast *fair* V F1, slow and fast *hire* T F2 or slow and fast *poor* C F3). This is especially noticeable for F1 and F2 and less so for F3. It is also particularly discernible in the case of the vowel and the transitional element, but not so much in the case of the consonant.

What can also be detected in these scatter plots is the fact that the difference between the mean values across contexts tends to be smaller in the slow-rate productions than in the fast-rate ones. In other words, there is greater dispersion between the mean values of the seven tokens in the fast rates than in the slow ones. Again, this can be easily seen in the case of F1 and F2 but is not easy to perceive in the case of F3. Likewise, it is more easily distinguishable in the vowel and the transitional element than in the consonant.

These observations provide the grounds to state that, albeit to a different extent,

there is variability in the vowel, the transitional element and the consonant as regards F1, F2 and F3 mean values in the vowel, the transitional element and, to a lesser extent, the consonant, across rates.

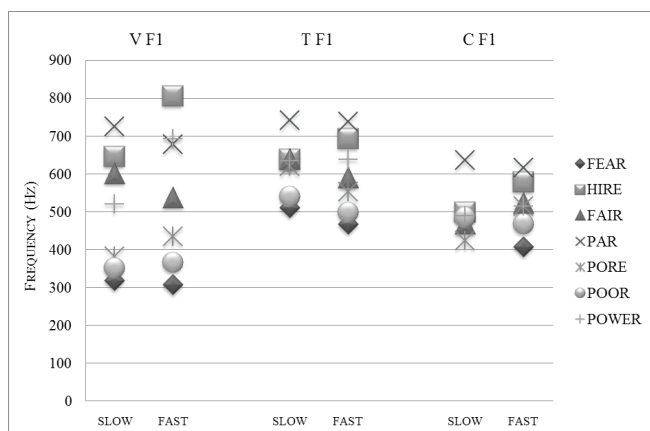


Figure 2 – Scatter plots for F1 values for one speaker by context. Each data point represents the mean for 10 tokens in each category. The vertical axis shows F1 frequency. The horizontal axis shows the values for the vowel (V), transition (T) and consonant (C) as well as the slow and fast rate values for each of these.

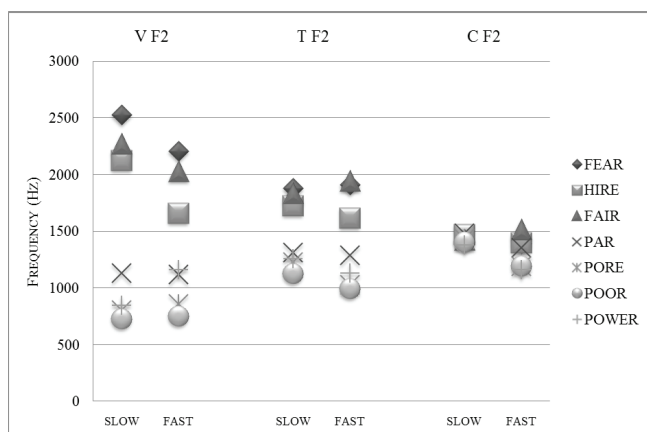


Figure 3 – Scatter plots for F2 values for one speaker by context. Each data point represents the mean for 10 tokens in each category. The vertical axis shows F2 frequency. The horizontal axis shows the values for the vowel (V), transition (T) and consonant (C) as well as the slow and fast rate values for each of these.

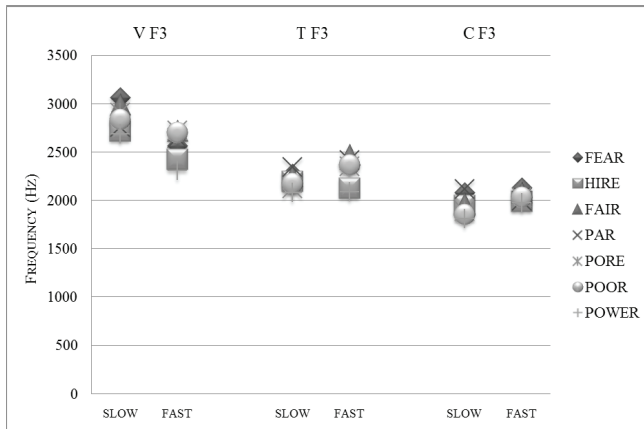


Figure 4 – Scatter plots for F3 values for one speaker by context. Each data point represents the mean for 10 tokens in each category. The vertical axis shows F3 frequency. The horizontal axis shows the values for the vowel (V), transition (T) and consonant (C) as well as the slow and fast rate values for each of these.

4. Discussion and conclusions

The purpose of this study has been to further investigate the VC coarticulatory processes that take place in final V+/r/ sequences in American English stressed monosyllables in order to contribute new insights into the behavior and nature of these sequences. These insights are meant to expand on the results obtained and conclusions reached in previous studies carried out by the same authors (Riera & Romero 2006, 2007; Riera et al. 2009). We have designed an experiment which replicates in part these previous studies but also introduces innovative aspects related to participants, stimuli, segmentation procedures and measurements taken. We have gathered acoustic data for the different constituent elements in the sequences that have allowed us to confirm already existing conclusions and reach new ones concerning the role played by speaking rate.

The first hypothesis (i.e., that there is significant durational, F1, F2 and F3 variability in the vowel, the transitional element and the consonant, across contexts, and as a function of speaking rate) has been confirmed by the results of the statistical analyses, as well as by the information presented in Figures 2, 3 and 4. This provides evidence for the existence of coarticulatory processes and shows the extent of VC coarticulation in the V+/r/ sequences which are the object of our study. Despite having mean duration, F1, F2 and F3 values similar

to those of a mid central vowel (i.e., schwa), the transitional element has been proven to be different in each of the different contexts (i.e., different vowels + /r/). This rules out the possibility of the transitional element being considered a segment and thus reveals its phonetic, rather than phonological, nature.

The second hypothesis (i.e., that the F1, F2 and F3 mean values of the different contexts tend to resemble each other more in the slow-rate productions than in the fast-rate ones) provides evidence for the dynamic nature of the sequences, in general, and of the transitional element, in particular. The information provided in Figures 2, 3 and 4 evidences that it takes longer for the vowel to attain the transitional element target and for this element to attain the consonant target in the slow productions than in the fast ones. It shows, therefore, how an increase in speech rate entails a decrease in time for the articulatory gestures to attain their targets. All in all, it proves that we are dealing with a process of coarticulation rather than epenthesis/insertion. This also complements the findings of previous studies (Riera & Romero 2007; Riera et al. 2009) that reveal how the coarticulatory influence of the vowels on their corresponding transitional elements is shown by the fact that the spectral values of these elements tend to resemble more those of the preceding vowels the faster the speaking rate.

Despite not varying much across V contexts, the acoustic characteristics of the /r/ in the different sequences show some variability, which can be taken as proof of the coarticulatory influence exerted by the vowel on the schwa-like element, by the schwa-like element on the consonant, and even by the vowel on the consonant. The fact that the variability is smaller in the /r/ than in the schwa-like element is explained by the fact that the /r/ is present underlyingly and, therefore, it is associated with clearly determined articulatory targets, whereas the schwa-like element does not correspond to any underlying segment and, therefore, has no specific articulatory targets.

The present study has not aimed at finding relationships between the sequences according to the phonological parameters for the classification of vowels (i.e., vowel height or frontness/backness). A possible further study would look into the role played by context (i.e., each of the seven different vowels in the V+/r/ sequences) as well as by examining vowel-transition and transition-consonant differences.

Finally, we believe that the limitations posed by an acoustic analysis of the type reported in this paper, based on segmentation as well as durational and spectral measurements, can only be overcome by an articulatory analysis of the type offered, for example, by the Electromagnetic Midsagittal Articulator (EMMA) technique. This type of study is meant to be considered for future research.

Acknowledgments

Research Group in Experimental Phonetics (Universitat Rovira i Virgili, Tarragona, Spain).

Research Groups 2005-SGR00864 and 2009-SGR003 (Generalitat de Catalunya, Spain: Institut d'Estudis Catalans and Universitat Autònoma de Barcelona).

Projects HUM2005-02746 and FFI2010-19206 (Ministerio de Educación y Ciencia, Spain: Universitat Autònoma de Barcelona).

The authors would also like to thank two anonymous reviewers for their insightful comments and suggestions.

References

- Boersma, Paul & David Weenink. 2010. *Praat: doing phonetics by computer: Version 5.2*.
- Browman, Catherine & Louis Goldstein. 1986. Towards an articulatory phonology. *Phonology Yearbook* 3. 219-252.
- Browman, Catherine & Louis Goldstein. 1989. Articulatory gestures as phonological units. *Phonology* 6(2). 201-251.
- Browman, Catherine & Louis Goldstein. 1990a. Gestural specification using dynamically-defined articulatory structures. *Journal of Phonetics* 18(3). 299-320.
- Browman, Catherine & Louis Goldstein. 1990b. Tiers in articulatory phonology, with some implications for casual speech. In John Kingston & Mary Beckman (eds.), *Papers in Laboratory Phonology I: between the grammar and physics of speech*, 341-376. Cambridge: Cambridge University Press.
- Browman, Catherine & Louis Goldstein. 1992a. Articulatory phonology: an overview. *Phonetica* 49(3-4). 155-180.
- Browman, Catherine & Louis Goldstein. 1992b. "Targetless" schwa: an articulatory analysis. In Gerry Docherty & Robert Ladd (eds.), *Papers in Laboratory Phonology II: gesture, segment, prosody*, 26-56. Cambridge: Cambridge University Press.
- Gick, Brian & Ian Wilson. 2001. Pre-liquid excrescent schwa: what happens when vocalic targets conflict. In Paul Dalsgaard, Børge Lindberg & Henrik Benner (eds.), *Proceedings of the 7th European Conference on Speech Communication and Technology (Eurospeech 2001)*, 273-276. Aalborg: Center for Personkommunikation.
- Gick, Brian & Ian Wilson. 2006. Excrescent schwa and vowel laxing: cross-linguistic responses to conflicting articulatory targets. In Louis Goldstein, Douglas Whalen & Catherine Best (eds.), *Laboratory Phonology 8*, 635-659. Berlin: Mouton de Gruyter.
- Goldstein, Louis & Carol Fowler. 2003. Articulatory phonology: a phonology for public

- language use. In Niels Schiller & Antje Meyer (eds.), *Phonetics and phonology in language comprehension and production: differences and similarities*, 159-207. Berlin: Mouton de Gruyter.
- Hall, Nancy. 2003. *Gestures and segments: vowel intrusion as overlap*. Amherst: University of Massachusetts, Amherst dissertation.
- Hall, Nancy. 2006. Cross-linguistic patterns of vowel intrusion. *Phonology* 23(3). 87-429.
- Lavoie, Lisa & Abigail Cohn. 1999. Sesquisyllables of English: the structure of vowel-liquid syllables. In John Ohala, Yoko Hasegawa, Manjari Ohala, Daniel Granville & Ashley Bailey (eds.), *Proceedings of the XIVth International Congress of Phonetic Sciences (ICPhS XIV)*, 109-112. Berkeley: University of California, Berkeley.
- Riera, María & Joaquín Romero. 2006. V+/l/ and V+/r/ sequences in American English: a preliminary descriptive acoustic study. In Alejandro Alcaraz Sintes, Concepción Soto Palomo & María de la Cinta Zunino Garrido (eds.), *Actas del XXIX Congreso Internacional de la Asociación Española de Estudios Anglo-Norteamericanos / Proceedings of the 29th AEDEAN International Conference*, 529-536. Jaén: Universidad de Jaén Servicio de Publicaciones.
- Riera, María & Joaquín Romero. 2007. Schwa in American English V+/l/ sequences: speaking rate effects. *Actes des/Proceedings of JEL'2007: Schwa(s) – 5^{èmes} Journées d'Études Linguistiques de Nantes*, 131-136. Nantes: Université de Nantes.
- Riera, María, Joaquín Romero & Ben Parrell. 2009. Schwa in American English V+/r/ sequences. In Marina Vigário, Sónia Frota & João Freitas (eds.), *Phonetics and phonology: Interactions and interrelations*, 15-34. Amsterdam: John Benjamins.
- Warner, Natasha, Allard Jongman, Ann Cutler & Doris Mücke. 2001. The phonological status of Dutch epenthetic schwa. *Phonology* 18(3). 387-420.
- Wells, John C. 2000. *Longman Pronunciation Dictionary*, 2nd edn. Harlow: Longman-Pearson Education.