

The Influence of Neighborhood Density on the Recognition of Spanish-Accented Words

Kit Ying Chan
James Madison University

Michael S. Vitevitch
University of Kansas

Foreign-accented speech is more difficult to recognize than the same words produced by a native speaker because the accented speech may activate many additional competitors, or it may strongly activate a single, but incorrect, word during lexical retrieval. Experiments 1 and 2 examined the recognition of native-produced and foreign-accented words varying in neighborhood density with auditory lexical decision and perceptual identification tasks, respectively. Experiment 1 found increased reaction times (RTs), especially for accented dense words. Analysis of misperceptions from Experiment 2 found that the mean number of phonologically distinct misperception tokens was higher for native than accented stimuli, suggesting that accented speech does not tend to activate more lexical candidates. Furthermore, a higher proportion of misperceptions in the accented condition (71%) compared with the native condition (58%) was accounted for by the most frequently reported misperception token, suggesting that accented speech instead tends to strongly activate 1 particular neighbor of the target word during lexical competition. Moreover, systematic phonemic substitutions in the misperceptions suggest that lawful acoustic-phonetic variations introduced by the accented speaker's L1 (native language) play a crucial role in determining which neighbor is activated as a strong competitor.

Keywords: spoken word recognition, foreign accent, neighborhood density

Foreign-accented speech deviates from speech *produced* by native speakers (NSs) of a language in a variety of ways (Flege, 1984), and contains greater within- and across-speaker variability compared with speech produced by NSs (Wade, Jongman, & Sereno, 2007). Previous research on the *recognition* of foreign-accented speech found that, compared with native speech, foreign-accented speech is generally less intelligible (Munro & Derwing, 1995a), requires more processing time (Munro & Derwing, 1995b), and is more vulnerable to the adverse effects of noise on its intelligibility (Lane, 1963; Munro & Derwing, 1998; van Wijngaarden, 2001).

A growing body of research on foreign-accented speech has focused on the perceptual learning that takes place in a native listener to recognize foreign-accented words. Native listeners use lexical knowledge to interpret ambiguous accented words and learn the lawful accent-general variations introduced by speakers' native phonological structure (Norris, McQueen, & Cutler, 2003). Native listeners can retune their phonetic categories to guide later

perception of speech spoken with the same foreign accent (Bradlow & Bent, 2008; Reinisch & Holt, 2013; Sidaras, Alexander, & Nygaard, 2009). Several studies also examine different factors affecting perceptual learning of accented speech, such as accent inconsistency (Witteman, Weber, & McQueen, 2014), acoustic variability (Wade et al., 2007), accent strength (Witteman, Weber, & McQueen, 2013), and native listener's familiarity with an accent (Witteman et al., 2013). These previous studies on the perceptual learning of foreign-accented speech considered what the native listeners learn, how adaptation works, and how much improvement can be achieved via perceptual learning. Given that the improvement after perceptual learning tends to be quite limited (Bradlow & Bent, 2008; Sidaras et al., 2009; Trude, Tremblay, & Brown-Schmidt, 2013; Wade et al., 2007), we chose to instead focus on the process of spoken word recognition (SWR) itself to more directly examine why foreign-accented words are more difficult to recognize in the first place. Currently, little is known about how foreign accents negatively impact the SWR process.

Influential models of SWR, including TRACE (McClelland & Elman, 1986), Shortlist (Norris, 1994), and Neighborhood Activation Model (NAM; Luce & Pisoni, 1998) all propose that incoming speech signals activate multiple representations of words stored in the lexicon, and that these activated word forms then compete for recognition. In the current study, we examined how the acoustic-phonetic deviations induced by foreign accents might affect lexical activation and competition. Bürki-Cohen, Miller, and Eimas (2001) found in a phoneme monitoring task that accent-related acoustic-phonetic deviations lead to a mismatch between the accented speech inputs and the native listener's segmental representations, thereby disrupting processing of the phonological segments. If a similar phonological mismatch occurs during SWR,

This article was published Online First December 8, 2014.

Kit Ying Chan, Department of Psychology, James Madison University; Michael S. Vitevitch, Department of Psychology, University of Kansas.

The experiments in this report partially fulfilled the requirements for a PhD degree in psychology awarded to Kit Ying Chan; Michael S. Vitevitch chaired the dissertation. We thank the members of the defense committee (Joan Sereno, Susan Kemper, Allard Jongman, and John Colombo) for their helpful comments and suggestions.

Correspondence concerning this article should be addressed to Kit Ying Chan, Department of Psychology, 91 E. Grace Street, MSC 7704, James Madison University, Harrisonburg, VA 22807. E-mail: chanky@jmu.edu

would (a) a larger set of lexical candidates be activated compared with the same word spoken by a NS, or (b) would a single, but incorrect, lexical candidate be more strongly activated than the target word? We sought in the present set of experiments to determine which of these scenarios occurs during foreign-accented word recognition.

For native word recognition, neighborhood density has been well-studied and shown to strongly influence lexical discrimination. Neighborhood density refers to the number of words that are phonologically similar (i.e., phonological neighbors) to a target word. A common metric used to assess phonological similarity between two words is the addition, deletion, or substitution of a phoneme in one word to form another word; such words are called phonological neighbors (Greenberg & Jenkins, 1967; Landauer & Streeter, 1973; Luce & Pisoni, 1998). For example, the word *cat* has as phonological neighbors the words: *_at*, *scat*, *rat*, *cut*, and *cap*. Words with many similar sounding neighbors are said to have dense neighborhoods, whereas words with few similar sounding neighbors are said to have sparse neighborhoods. In NAM (Luce & Pisoni, 1998), the probability of identifying a word depends not only on the incoming acoustic-phonetic information, but also higher-level information, including neighborhood density and the frequency of the target word and its neighbors.

Because of a large number of confusable competitors, the recognition of words (spoken by NSs) from dense neighborhoods tends to be less accurate and proceed more slowly than the recognition of words from sparse neighborhoods in native English listeners (Cluff & Luce, 1990; Goldinger, Luce, & Pisoni, 1989; Luce & Pisoni, 1998; Vitevitch, 2002, 2003; Vitevitch & Luce, 1999; Vitevitch, Stamer, & Sereno, 2008). NAM also states that identification of a word with less intelligible phonemes can still be relatively high if there are few neighbors that contain phonemes confusable with those of the target word. This implies that recognition of accented speech, which is generally less intelligible, might be influenced even more by neighborhood density than more intelligible, native speech. A previous study by Imai, Walley, and Flege (2005) found a significant neighborhood density effect on the recognition of Spanish-accented words by native English listeners. Participants were asked to identify native-produced and Spanish-accented words embedded in multitalker babbling noise. Spanish-accented sparse words were recognized more accurately than Spanish-accented dense words, whereas no such difference was observed for the native words. This finding shows that foreign accents affect the recognition of dense words more adversely than the recognition of sparse words. Consistent with NAM, Imai et al. (2005) suggested that finer-grained phonetic discriminations may be required to accurately recognize dense words given the large number of confusable competitors in dense neighborhoods.

Note that the study by Imai et al. (2005) used noise-degraded stimuli, and only recognition accuracy rate was measured. Note further that the recognition of foreign-accented speech is more adversely affected by noise than the recognition of native speech (Munro, 1998). Previous studies suggest that foreign accents also increase the processing time for recognizing words; are processing times (in addition to processing accuracy) differentially influenced in dense and sparse words spoken with foreign accents? Experiment 1 in the current study used a lexical decision task to examine this question by using noise-free stimuli and measuring reaction time (RT).

Furthermore, the findings from Imai et al. (2005) do not show precisely how foreign accents affect lexical activation and competition to make lexical discrimination more difficult in dense words. If the accented speakers' production of the L2 (second language) words deviates from that of NSs, it could sound like several different words. Apart from the target word, it is possible that words, which are less similar to the target words and normally not activated by native speech signals, might partially match the accented-speech signals and get activated (a situation we refer to as the *many-additional-competitor scenario*). Another possibility is that the accented-speech signal might sound like another word instead of the target word. That is, the accented production may sound more like one of the neighbors in the neighborhood than it does the target word (a situation we refer to as the *single-strong-competitor scenario*). The single-strong-competitor scenario may occur if foreign-accented speakers systematically mispronounce certain phonemes in the L2 as another phoneme. Experiment 2 in the current study aimed to examine which of the two scenarios contributes to the increased recognition difficulty of foreign-accented words. By analyzing the misperceptions from a perceptual identification task of native and foreign-accented words, we can gain insight into the set of lexical candidates that are most highly activated when listening to native and accented speech.

Experiment 1

In Experiment 1, participants were presented with either a word or a nonword without any noise. Participants were asked to decide as quickly and as accurately as possible whether the given stimulus was a real word in English or a nonsense word. RTs and accuracy rates were measured as dependent variables. Most of the previous studies used as stimuli sentences or stimulus words embedded in carrier sentences to study the processing costs in accented word recognition (Clarke & Garrett, 2004; Lane, 1963; Munro, 1998; Munro & Derwing, 1995a, 1995b; Schmid & Yeni-Komshian, 1999; van Wijngaarden, 2001). Without a sentence context, the participants in the present experiment could not use any semantic/syntactic cues for word recognition. Therefore, the impact of foreign accents on word recognition can be examined more directly in the present experiment.

Previous studies of SWR using native speech often find no difference in accuracy rates in the lexical decision task even when a significance difference is found in RTs (Vitevitch & Luce, 1999). Thus, accuracy rates were not expected to be different for the dense and sparse words in Experiment 1. This prediction differs somewhat from Imai et al. (2005), who found that Spanish-accented sparse words were recognized more accurately than Spanish-accented dense words. Recall, however, that Imai et al. used a perceptual identification task, which does not assess processing time like the lexical decision task used in the present experiment. Furthermore, they, somewhat curiously, did not observe an influence of neighborhood density for the native-produced words. Given the well-studied influence of neighborhood density on a wide variety of language processes (Cluff & Luce, 1990; Goldinger et al., 1989; Luce & Pisoni, 1998; Vitevitch, 2002, 2003; Vitevitch & Luce, 1999; Vitevitch et al., 2008), it was predicted that neighborhood density should influence recognition times in both the native- and accented-speech conditions. Given the greater processing cost found when processing foreign-

accented speech (Imai et al., 2005), we further expect a much bigger influence of neighborhood density in the foreign-accented condition than in the native-produced condition.

Method

Participants

Forty-eight NSs of American English were recruited from the pool of introductory psychology students enrolled at the James Madison University. The participants received partial credit toward the completion of the course for their participation. All participants were right-handed and reported no history of speech or hearing disorders. Participants in the native and accented conditions were similar in terms of their language backgrounds and experience. A majority of them learned Spanish as an L2 ($M_{\text{native}} = 84\%$ and $M_{\text{accented}} = 88\%$), but they do not speak Spanish or another language fluently ($M_{\text{native}} = 84\%$ and $M_{\text{accented}} = 88\%$). Most of them do not have family members or close friends with a Spanish accent ($M_{\text{native}} = 92\%$ and $M_{\text{accented}} = 84\%$).

Materials

All the 64 English monosyllabic stimulus words used in the experiment had a consonant-vowel-consonant (CVC) structure. Half of the stimuli had a dense neighborhood, $M = 27.4$, $SE = 0.351$, 95% CI [26.7, 28.2], and half had a sparse neighborhood, $M = 16.0$, $SE = 0.466$, 95% CI [15.0, 16.9], $F(1, 62) = 387$, $p < .001$. Although the words differed in neighborhood density, the words in the two conditions were similar in subjective familiarity, word frequency, neighborhood frequency, phonotactic probability, and distribution of phonemes. Subjective familiarity was measured on a 7-point scale (Nusbaum, Pisoni, & Davis, 1984). Word frequency refers to the average occurrence of a word in the language. Neighborhood frequency is defined as the mean word frequency of the neighbors of the target word. The phonotactic probability is measured by how often a certain segment occurs in a certain position in a word (positional segment frequency) and the segment-to-segment co-occurrence probability (biphone frequency; Vitevitch & Luce, 1998). The neighborhood density, mean positional segment frequency, and mean biphone frequency for each stimulus was obtained from the Web-based calculator described in Storkel and Hoover (2010). The means and standard errors for each of these lexical characteristics can be found in Table 1. These two groups of stimulus words and their lexical characteristics are listed in Appendix A and B.

Table 1
Means and Standard Errors (in Parentheses) for the Lexical Characteristics of the Dense and Sparse Conditions

Lexical characteristic	Dense	Sparse
Neighborhood density	27.4 (.351)	16.0 (.466)
Subjective familiarity	6.93 (.031)	6.87 (.044)
log word frequency	1.27 (.131)	1.30 (.123)
log neighborhood frequency	3.59 (1.57)	2.02 (.043)
Positional segment frequency	.152 (.005)	.146 (.007)
Biphone frequency	.007 (.0007)	.006 (.0009)

Distribution of phonemes. The distribution of phonemes in each of the phoneme positions in the words was balanced as much as possible across the dense and sparse neighborhood density conditions because certain English sounds or sequences of sounds are characteristically difficult for Spanish-accented speakers to produce. For example, for consonants, Spanish-accented speakers tend to produce /z/ as /s/ in the final position, /v/ as /b/ in the initial position, and /p, t, k/ in initial position with less aspiration (Magen, 1998; You, Alwan, Kazemzadeh, & Narayanan, 2005); for vowels, Spanish-accented speakers tend to have more difficulty producing vowels that exist in English but not in Spanish, including /ɪ, æ, ʌ/ (Sidas et al., 2009). Therefore, these English sounds or sequences of sounds whose production are characteristically difficult for Spanish-accented speakers were all matched among the stimuli in the two neighborhood density conditions.

The following phonemes were matched between the dense and sparse conditions, including onset consonants /p, t, b, d, f, s, ʃ, n, ɹ/, vowels /i, ɪ, ɜ, e, æ, ʌ, ɔ, o, u/, and final consonants /t, d, f, s, ʃ, z, v, ɹ/. The only unmatched onset consonants were an extra /g/ and /k/ in the sparse condition, and an extra /l/ and /w/ in the dense condition. The only unmatched vowels were two extra /au/ in the sparse condition and two extra /aɪ/ in the dense condition. The unmatched final consonants were those that are not characteristically difficult for Spanish-accented speakers to produce, including /p(4/3), k(3/6), b(2/0), g(2/1), n(5/7), m(3/1), and l(3/4)/, with their number of occurrence in the sparse and dense conditions in parentheses. The final consonants were categorized into different manners of articulations (stops, sibilant fricatives, nonsibilant fricatives, nasals, and liquids), and its distribution across the dense and sparse conditions was tested. A chi-square test for goodness-of-fit was not significant, $\chi^2(4, N = 64) = .128$, $p = .998$, suggesting no statistically significant difference in the distribution of different types of consonants in the final position across conditions. Overall, the distribution of constituent phonemes in the two conditions was similar; it is more likely that any difference observed in the lexical decision task is due to the difference in the independent variables (i.e., neighborhood density and accent), rather than differences in the phoneme distribution in the two conditions.

A list of 64 phonotactically legal nonwords with the same initial consonant, middle vowel, and phoneme length as the word stimuli were selected from the ARC nonword database as foils (Rastle, Harrington, & Coltheart, 2002). The phonological transcriptions of the nonwords are listed in Appendix C.

Speakers. A male non-native speaker (NNS) of English with Spanish as his native language was recruited through flyers sent through the University of Kansas international student association. This speaker was from Lima, Peru and had resided in the United States for a minimum of 1 year but less than 2 years. He was 35 years old and learned English when he was 23 years old. The speaker had learned English after puberty and was judged by 12 native listeners to have a heavy foreign accent in a pilot screening (details are further discussed in the next section). The speaker reported having no hearing or speech disorder and was paid \$10/hr for his participation. A male NS of American English from the Midwest was recruited from the University of Kansas to record the native version of the word stimuli under the same conditions as the NNS.

Each of the speakers read each word/nonword in a random order. The speech was recorded digitally at a 44.1-kHz sampling rate using a high-quality microphone and a solid-state recorder (Marantz PMD671) in an IAC (Industrial Acoustics Company) sound attenuated booth. Each stimulus was edited using Praat (Boersma & Weenink, 2009) into an individual sound file. The amplitude of the individual sound files was increased to their maximum without distorting the sound or changing the pitch of the words by Praat.

Degree of foreign-accentedness. The degree of foreign-accentedness of each speaker was determined by a foreign-accentedness rating task with 12 native English-speaking pilot listeners from the pool of introductory psychology students enrolled at the University of Kansas. These listeners were visually and auditorily presented a random sample of 16 stimulus words (8 dense and 8 sparse words) produced by each of the two speakers in a random order in a noise-free listening condition, and asked to rate each item for degree of accentedness using a 7-point scale ranging from 1 (*native-like*) to 7 (*strong foreign accent*).

The listeners' ratings ranged from 1 to 7, suggesting the use of the whole scale. The mean accentedness ratings (standard deviations are in parentheses) for the sparse items, dense items, and all items for the NS were 1.3(.7), 1.3(.4), and 1.3(.6), whereas those for the NNS were 5.7(1.0), 5.4(1.1), and 5.5(1.1). A 2 (accent: native vs. foreign-accented) $\times$ 2 (neighborhood density: dense vs. sparse) mixed-design analysis of variance (ANOVA), with accent as a within-words factor and neighborhood density as a between-words factor, shows that the foreign-accented English speaker (the NNS), $M = 5.52$, $SD = 1.08$, 95% CI [5.25, 5.79], was rated as having a stronger accent than was the native English speaker (the NS), $M = 1.32$, $SD = 0.560$, 95% CI [1.18, 1.46], $F(1, 62) = 775$, $p < .001$, partial $\eta^2 = .926$. There was no significant neighborhood density main effect, or interaction between accent and neighborhood density.

Stimulus duration. The duration of all the stimuli, including word and nonword, were submitted to a 2 (speaker) $\times$ 2 (neighborhood density) $\times$ 2 (lexicality) mixed-design ANOVA to check for any speaker effect, neighborhood density effect, lexicality effect, or interaction. The ANOVAs revealed no significant neighborhood density effect, $F(1, 124) = 0.094$, $p = .759$, nor lexicality effect, $F(1, 124) = 2.34$, $p = .126$. However, there was a significant speaker effect, $F(1, 124) = 24.28$, $p < .001$. The mean word durations (standard deviations are in parentheses) for the word and nonword stimuli were calculated for each speaker and neighborhood density condition, and are listed in Table 2. As seen in Table 2, the mean durations of words and nonwords were significantly different between the speakers.

A 2 $\times$ 2 mixed factorial design that includes *accent* as a between-subjects factor and *neighborhood density* as a within-subjects factor was adopted. Half (16 items) of the words were randomly selected from each of the two neighborhood density conditions to form list A, and the remaining half formed list B such that each list contained 16 dense and 16 sparse words. For counterbalancing purposes, half of the participants received list A first and then list B (designated by AB in the next paragraph), whereas the other half received list B first and then list A (designated by BA).

The speakers representing the native-produced and foreign-accented conditions were designated by NS and NNS. With *accent* as between-subjects factors, participants were randomly assigned to one of the four counterbalanced conditions, in which all the stimuli were produced only by one of the two speakers (with the order of list presentation indicated after the dash): NS-AB, NS-BA, NNS-AB, or NNS-BA. Thus, a given listener only heard stimuli spoken by one of the two speakers and each of the 64 stimulus words once—16 dense and 16 sparse words in the first block and another 16 dense and 16 sparse words in the second block. Items within blocks were presented in a different randomized order for each participant. Across participants, each stimulus word was presented in both native and foreign-accented form and evenly presented in the two blocks.

Procedure

Listeners were tested individually or in pairs. Each participant was seated in front of a Dell PC computer in an individual listening station separated by partitions. The presentation of stimuli and the collection of responses were controlled by Paradigm 2.0 (Perception Research Systems, 2007).

Each trial started with the word "Ready" appearing on the computer screen for 500 ms. Then the participants heard one of the randomly selected words or nonwords over a set of Sennheiser HD 25-SP II headphones at a comfortable listening level. Each stimulus was presented only once. The participants were instructed to respond as quickly and as accurately as possible whether the item they heard was a real English word or a nonword. If the item was a word, they were to press the button labeled "Word" with their right (dominant) hand. If the item was not a word, they were to press the button labeled "Nonword" with their left hand. RTs were measured from the onset of the stimulus to the onset of the button press response. After the participant pressed a response button, the next trial began. The experiment lasted about 15 minutes. Prior to the experi-

Table 2
Mean Word Durations in Milliseconds (ms) and Standard Deviations (in Parentheses) for the Word and Nonword Stimuli in Each Speaker and Neighborhood Density Condition

Speaker	Mean word duration/ms (SD)					
	Words			Nonwords		
	Dense	Sparse	All words	Dense	Sparse	All nonwords
Native Male	550 (96.9)	559 (102)	554 (98.7)	558 (94.0)	536 (80.9)	547 (87.7)
Accented Male	623 (146)	624 (139)	623 (142)	584 (98.9)	576 (104)	580 (101)

mental trials, each participant received 10 practice trials, which were not included in the data analyses.

Results

The current convention in psycholinguistic research is to perform analyses with participants as a random factor (subject analysis, F_1) and with items as a random factor (item analysis, F_2 ; however see Clark, 1973, for an alternative analysis). However, there is some debate about the proper use and interpretation of additional item analysis over subject analysis, especially when items are carefully matched or balanced across conditions on important variables correlated with the response measures (Raaijmakers, 2003; Raaijmakers, Schrijnemakers, & Gremmen, 1999). Because the stimulus items were well-controlled in the present study, additional item analyses are not appropriate or necessary (Raaijmakers et al., 1999). Just to be consistent with the conventions of the field, additional item analyses were reported in all of the experiments in the current study. However, the interpretation of the results was based on the subject analysis.

Only accurate responses for the word stimuli were included in the data analysis. RTs that were too rapid or too slow (i.e., below 500 ms or above 2,000 ms) were considered outliers. Using these cutoffs, a total of 1.6% of data, including .16% from the native sparse condition, .03% from the native dense condition, .68% from the accented sparse condition, and .72% from the accented dense condition, was excluded from the analysis. To factor out the effect of stimulus duration on the RTs, corrected RTs, which were obtained by subtracting the duration of each stimulus word from each subject's RT to that word, were used. It measures the amount of time it takes the participants to press the response button after the end of the utterance. With participants as the random variable, responses were pooled across stimulus items, yielding mean corrected RTs and accuracy rates in the dense and sparse conditions for each participant. These mean corrected RTs and accuracy rates were subjected to a 2×2 mixed-design ANOVA with neighborhood density as a within-subjects factor and accent as a between-subjects factor.

The ANOVA yielded a significant main effect of accent, $F_1(1, 46) = 10.37, p = .002$, partial $\eta^2 = .184$; $F_2(1, 61) = 50.8, p < .001$, partial $\eta^2 = .455$. The main effect of neighborhood density was also significant, $F_1(1, 46) = 37.0, p < .001$, partial $\eta^2 = .446$; $F_2(1, 61) = 0.305, p = .582$. The mean corrected RT for the foreign-accented condition, $M = 511$ ms, $SD = 119$, 95% CI [471, 553], was longer than the native-produced condition, $M = 418$ ms, $SD = 87.9$, 95% CI [378, 460]. The mean corrected RT for the dense words, $M = 487$ ms, $SD = 120$, 95% CI [456, 518], was longer than the sparse words, $M = 443$ ms, $SD = 105$, 95% CI [415, 472]. The interaction between accent and neighborhood density was also significant, $F_1(1, 46) = 5.83, p = .02$, partial $\eta^2 = .112$; $F_2(1, 61) = 0.276, p = .601$. Figure 1 shows the mean corrected RTs (standard deviations are in parentheses) as a function of accent (native, accented) and neighborhood density (dense, sparse). Post hoc tests with Bonferroni's correction ($p < .05$) revealed that dense words were responded to more slowly than sparse words in the accented condition, $F_1(1, 46) = 36.09, p < .001$, as well as in the native condition, $F_1(1, 46) = 6.72, p = .013$. Raw RTs were also analyzed with participants as the random variable for reference in Appendix D.

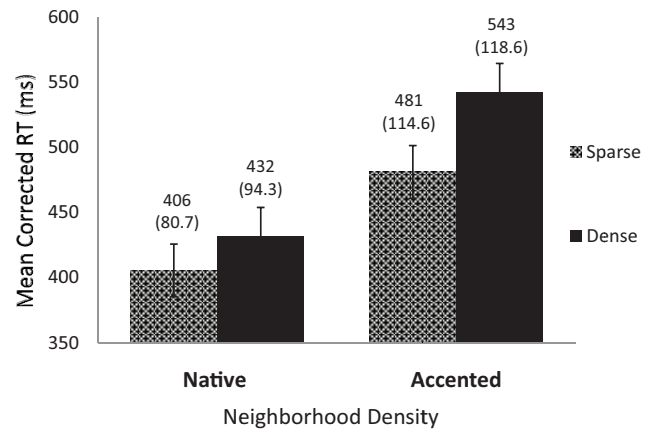


Figure 1. Mean corrected RTs and standard deviations (in parentheses) for the lexical decision task in Experiment 1 as a function of accent and neighborhood density (subject-analysis). Error bars indicate the 95% confidence intervals.

For accuracy rates, the main effects of accent was significant, $F_1(1, 46) = 114.7, p < .001$, partial $\eta^2 = .714$; $F_2(1, 62) = 22.1, p < .001$, partial $\eta^2 = .263$. The main effect of neighborhood density was also significant, $F_1(1, 46) = 20.5, p < .001$, partial $\eta^2 = .308$; $F_2(1, 62) = 1.48, p = .229$, as well as the interaction between accent and neighborhood density were significant, $F_1(1, 46) = 11.1, p = .002$, partial $\eta^2 = .195$; $F_2(1, 62) = 1.43, p = .235$. Participants responded to native-produced words, $M = 91.2\%$, $SD = 7.00$, 95% CI [88.8, 93.5] more accurately than foreign-accented words, $M = 73.6\%$, $SD = 7.90$, 95% CI [71.2, 75.9]. In contrast with initial predictions, participants responded to dense words, $M = 84.8\%$, $SD = 10.0$, 95% CI [82.7, 86.9], more accurately than sparse words, $M = 79.9\%$, $SD = 12.5$, 95% CI [78.0, 81.7]. Post hoc tests with Bonferroni's correction ($p < .05$) revealed that dense words were responded to more accurately than sparse words only in the accented condition, $F_1(1, 46) = 30.9, p < .001$. Figure 2 shows the mean accuracy rates (standard deviations in parentheses) as a function of accent (native, accented) and neighborhood density (dense, sparse).

Discussion

The results of Experiment 1 show that listeners took a longer time to respond to foreign-accented words than native words. This result is consistent with previous studies, which used sentence stimuli to demonstrate that foreign-accented speech takes a longer time to process than native speech (Clarke & Garrett, 2004; Munro & Derwing, 1995b). Using only individual words as stimuli, the listeners in the current study cannot use higher-level semantic and syntactic information from the sentence to help recognize the individual words. Thus, the increased processing time found in the current study directly reflects the additional time required to process a word when spoken with a foreign accent.

The significant main effect of neighborhood density showed that listeners took a longer time to respond to words from dense neighborhoods than from sparse neighborhoods. More importantly, the significant interaction between accent and neighborhood density indicated a markedly larger neighborhood density effect

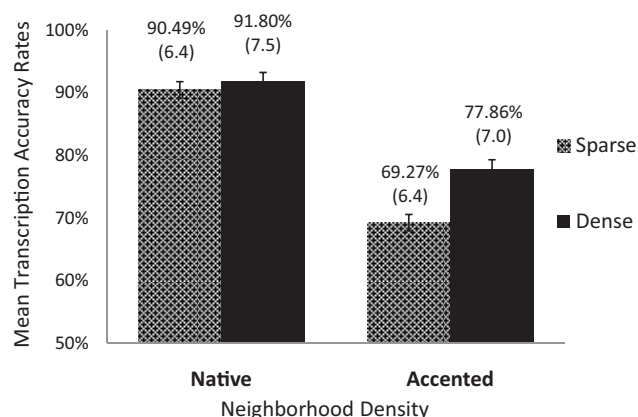


Figure 2. Mean transcription accuracy rates and standard deviations (in parentheses) for the lexical decision task in Experiment 1 as a function of accent and neighborhood density (subject-analysis). Error bars indicate the 95% confidence intervals.

for the accented stimuli, relative to the native stimuli. That is, the native listeners took more time to respond to words from dense neighborhoods than words from sparse neighborhoods in the foreign-accented condition. This suggests that lexical discrimination of words in a dense neighborhood becomes even more difficult in the presence of a foreign accent. This result is consistent with the previous results from Imai et al. (2005), which also showed an increased processing cost, in terms of lower transcription accuracy, for dense words than sparse words in foreign-accented condition using noise-degraded stimuli. Using a new set of well-balanced stimuli without noise degradation, the current experiment showed that the increased neighborhood density effect on foreign-accented word recognition can also be observed under more ideal listening conditions.

The accuracy rate result showed that participants responded to native-produced words more accurately than foreign-accented words. This result is consistent with previous studies showing that native listeners transcribed foreign-accented speech with more errors than native speech (Munro & Derwing, 1995a, 1995b). However, the significant interaction between accent and neighborhood density showed that participants responded to dense words more accurately than sparse words only in the foreign-accented condition, whereas participants responded to dense and sparse words with similar accuracy in the native condition. In short, foreign-accented dense words were responded to more slowly but more accurately than sparse words. This pattern of results may suggest a speed-accuracy trade-off when participants are responding to foreign-accented dense words, but this is probably not the best explanation. To further understand this result, it has to be considered in the context of the RTs and accuracy rates for the native condition. A different pattern of results was observed for the native condition: dense words were responded to more slowly but with equal accuracy compared with sparse words.

The different effect of neighborhood density on accuracy rate and RT as a function of accent can be explained by the accuracy assumption suggested by Luce and Pisoni (1998). The accuracy assumption suggests that when time restraints are imposed on a lexical decision by instructing participants to respond as quickly as

possible, accuracy rates will vary as a function of the amount of stimulus processing achieved before the response. Participants normally attempt to classify the stimulus before reaching a self-imposed RT deadline. If stimulus processing is completed before the self-imposed deadline, lexical decision responses will be accurate. But if stimulus processing is incomplete before the deadline, participants will execute a lexical decision response based on only partial information—the total lexical activity level of the decision system—resulting in numerous classification errors (Luce & Pisoni, 1998).

As foreign-accented stimuli probably require additional processing time that exceeds the response time deadline, lexical decisions for these stimuli are probably based on only partial information (i.e., the total lexical activity level of the decision system). For dense words with many word neighbors, the overall level of activity in the decision system tends to be higher. Therefore, dense words were more likely to be classified as words than sparse words only in the foreign-accented condition when longer processing time was required. Hence, the higher accuracy rate for the accented dense words was probably not due to more accurate lexical identification, but due to more lexical activity. No effect of neighborhood density was observed for accuracy rates for the native condition as decisions for native stimuli tended to be made before the response time deadline.

Another plausible explanation is that the RTs in the lexical decision do indicate accurate identification of the stimuli. As reflected by the longer RTs, dense words are particularly more difficult to recognize than sparse words in the presence of a foreign accent. However, as the lexical decision task only requires the participants to decide whether the stimulus item they heard is a real word or not, its accuracy rate might not really reflect the correct identification of the target words, depending on the characteristics of the stimuli. There is a higher chance for the dense word to be mistaken as another word due to the existence of many lexical neighbors. Substitution of a phoneme might result in another real word more often for dense words than for sparse words. Hence, misidentification of the target words as one of its similar sounding words might happen more easily for dense words, especially in the presence of foreign accents. As a result, in the current lexical decision task, accented dense words may have been classified as a word more “accurately” than accented sparse words, but participants might not have identified the right word when making this classification.

Both of these two plausible explanations suggest that the accuracy rate for accented speech in the lexical decision task might not be a good indicator of its identification accuracy. It is important to recognize this limitation of using a lexical decision task to examine accented speech processing and be cautious during data interpretation. To further examine whether the accented dense words were really more likely to be misidentified as one of their many similar sounding words, participants in Experiment 2 were asked to identify the same set of stimulus words, either native-produced or foreign-accented, in a perceptual identification task.

Experiment 2

The goal of Experiment 2 was to use a perceptual identification task to gain insight into the lexical competitors that are activated when native listeners hear accented speech, thereby exploring the

underlying cause of increased processing cost for accented speech. Is the lexical discrimination difficult for accented dense words because many additional competitors are activated (*the many-additional-competitors scenario*) or a single competitor is strongly activated (*single-strong-competitor scenario*)? Our use of the perceptual identification task contrasts with the way it was used in Imai et al. (2005), in which the perceptual identification task was used to assess how much accented speech impairs lexical processing via differences in accuracy rate. In the present experiment, participants were presented with a stimulus word, either native-produced or foreign-accented (against a background of white noise only for the native-produced condition), and were asked to identify it. By examining the misperceptions (perception errors) of the native and foreign-accented words, we can infer the set of lexical candidates that were highly activated and competing with the target word for lexical retrieval. It is important to note that the misperceptions are the erroneous responses the participants identified the target word as. Therefore, misperceptions can only allow us to infer the most activated competitor, but not the full set of lexical competitors being activated. Regardless of this limitation, a rich dataset of lexical competitors can still be collected by pooling misperceptions from all the participants. Analysis of this dataset allowed us to gain additional insight into how foreign accents influence the activation of lexical candidates during SWR.

To distinguish whether the many-additional-competitor scenario or the single-strong-competitor scenario increases the difficulty of recognizing foreign-accented words, we examined the distribution of misperceptions across all the misperception tokens. Each phonologically distinct misperception reported by one or more participants was considered as a misperception token. For example, the word *shear* was considered to have two misperception tokens if it was misidentified as *hear* (/hɪər/) by 3 participants, *here* (/hɪər/) by 2 participants, and *sear* (/sɪər/) by 1 participant.

If accented speech activates a broader range of lexical competitors than the same word spoken by a NS (the many-additional-competitor scenario), the total number of misperception tokens for each stimulus word in the accented condition is expected to be higher than that in the native condition. However, if a particular lexical competitor is dominantly activated (the single-strong-competitor scenario), the total number of misperception tokens for each stimulus word in the accented condition is expected to be equal or smaller than that in the native condition. Furthermore, it is likely that the dominantly activated misperception will be reported consistently across a majority of the listeners. In that case, for each accented word, a majority of its misperceptions is expected to be accounted for disproportionately by the dominantly activated misperception token (i.e., the first-most-reported token). Therefore, when we compare the distribution of misperceptions across the first-most-reported (dominantly activated) and the second-most-reported misperception tokens by ratio, we would expect the ratio to be larger for the accented condition than the native condition.

Method

Participants

Fifty-six NS of American English were recruited from the pool of introductory psychology students enrolled at the University of

Kansas. The participants received partial credit toward the completion of the course for their participation. All participants were right-handed and reported no history of speech or hearing disorders. None of the participants in the present experiment took part in the other experiment in this study.

Materials

The same 64-word stimuli from Experiment 1 were used in the present experiment. In an attempt to avoid both floor and ceiling performance in the identification of the words, we planned to add white noise of the same signal-to-noise ratio (S/N) to the sound files from both the native and foreign-accented conditions. After a series of pilot experiments using a variety of S/N ratios, we found that listeners experienced substantial difficulties (with an average accuracy rate of 50–60%) in identifying the foreign-accented words even without the addition of any noise. However, the native-produced words were identified almost perfectly (with an average recognition accuracy rate of 93–98%) when no noise was added to the sound files. Therefore, to minimize floor effects for the foreign-accented condition and ceiling effects for the native condition, we presented the foreign-accented speech with no noise but added a small amount of white noise to the native speech.

White noise with the same duration and relative amplitude as the sound file was added to the stimulus sound files from the NS of English using the GSU (Georgia State University) Praat tools (Owren, 2008) in Praat to produce stimuli with a +20 dB S/N (i.e., the mean amplitude of the resulting sound files were 20 dB more than that of the white noise). This S/N ratio resulted in a range of identification accuracy for the native condition (80–90%) that would yield a reasonable number of misperceptions for further analysis without unduly affecting normal word recognition.

With *accent* as a between-subjects factor and *neighborhood density* as a within-subjects factor, participants were randomly assigned to either the native-produced or foreign-accented condition. Thus, a given listener only heard stimuli spoken by one of the two speakers and each of the 32 dense and 32 sparse words once. Items were presented in a different randomized order for each participant. Across participants, each stimulus was presented in both native and foreign-accented form.

Procedure

Listeners were tested individually or in pairs. Each participant was seated in front of an iMac computer in an individual listening station separated by partitions. In the perceptual identification task, each trial begins with the word “Ready” appearing on the computer screen for 500 ms. Participants then heard one of the randomly selected stimulus words, either in quiet for the foreign-accented condition or embedded in white noise for the native-produced condition, through a set of Beyerdynamic DT 100 headphones at a comfortable listening level. Each stimulus was presented only once. Participants were instructed to use the computer keyboard to enter their response (or their best guess) for each word they heard over the headphones. They were instructed to type “?” if they were absolutely unable to identify the word. The participants were allowed as much time as they needed to type their response, and hit the Return key to initiate the next trial. Participants were able to see their responses on the computer screen when they were typing and could make cor-

rections to their responses before they hit the Return key. The experiment lasted about 10–15 min. Prior to the experiment, each participant received five practice trials to become familiar with the task. These practice trials were not included in the data analyses.

Results

Overall Accuracy

A response was scored as correct if the phonological transcription of the response and the stimulus was an exact match. Misspelling, transpositions of letters, and typographical errors that involve a single letter in the responses were scored as correct responses in certain conditions: (a) the omission of a letter in a word was scored as a correct response only if the response did not form another English word, and (b) the transposition or addition of a single letter in the word was scored as a correct response if the letter was within one key of the target letter on the keyboard. Responses that did not meet the above criteria were scored as incorrect. Items with an accuracy rate of three standard deviations or more below the mean were considered outliers, and were excluded from the data analysis. This resulted in five sparse words (*bib*, *beam*, *fad*, *lull*, and *psalm*) being excluded from data analysis in the native-produced condition. Based on this criterion, no dense words were excluded from data analysis in the native-produced condition, and no words of either type were excluded from data analysis in the foreign-accented condition.

With participants as the random variable, responses were pooled across stimulus items, yielding mean percent correct transcription scores in each of the two neighborhood density conditions for each participant. A 2 (accent: native vs. foreign-accented) $\times$ 2 (neighborhood density: dense vs. sparse) mixed-design ANOVA, with accent as a between-subjects factor and neighborhood density as a within-subjects factor, was conducted on the mean percent correct transcription scores. The ANOVA yielded a significant main effect of accent, $F_1(1, 54) = 660, p < .001$, partial $\eta^2 = .924$; $F_2(1, 57) = 59.7, p < .001$, partial $\eta^2 = .512$. Stimulus words spoken by foreign-accented speakers, $M = 52.1\%$, $SE = 1.08$, 95% CI [50.1, 54.2], were recognized less accurately than those spoken by NSs, $M = 89.5\%$, $SE = 1.00$, 95% CI [87.4, 91.5]. This result is especially striking because the foreign accented words were not mixed with any noise, whereas the native words were presented at a +20 dB S/N ratio. There was also a significant main effect of neighborhood density, $F_1(1, 54) = 11.2, p = .002$, partial $\eta^2 = .171$; $F_2(1, 57) = 1.36, p = .248$, such that dense words, $M = 68.9\%$, $SE = 1.00$, 95% CI [66.9, 70.9], were recognized less accurately than sparse words, $M = 72.7\%$, $SE = 0.900$, 95% CI [71.0, 74.4]. The interaction between neighborhood density and accent was not significant, $F_1(1, 54) = 0.002, p = .963$; $F_2(1, 57) = 0.203, p = .654$. The mean percent correct transcription scores and standard deviation for dense and sparse words were $M_{\text{dense}} = 87.5\%$, $SE = 1.40$, 95% CI [84.7, 90.3] versus $M_{\text{sparse}} = 91.4\%$, $SE = 1.20$, 95% CI [89.0, 93.9] for the native condition; and $M_{\text{dense}} = 50.2\%$, $SE = 1.40$, 95% CI [47.4, 53.0] versus $M_{\text{sparse}} = 54.0\%$, $SE = 1.20$, 95% CI [51.6, 56.4] for the accented condition.

Analysis of the Misperceptions

The misperceptions were analyzed to further understand how foreign accents might influence the activation of lexical candidates during SWR (i.e., the many-additional-competitor scenario vs. the single-strong-competitor scenario). By comparing the total number of misperception tokens collected for each stimulus in each accent condition, the first analysis examined whether a wide variety of competitors were activated (the many-additional-competitor scenario), or a small set of consistent competitors were activated (the single-strong-competitor scenario). To remove the influence of the raw number of misperceptions, the number of misperception token for each stimulus was normalized by dividing it by the total number of misperceptions. The normalized number of misperception tokens for each stimulus was then subjected to a 2 (accent: native vs. foreign-accented) $\times$ 2 (neighborhood density: dense vs. sparse) mixed-design ANOVA, with accent as a within-items factor and neighborhood density as a between-items factor. Nine sparse stimuli and eight dense stimuli had perfect identification accuracy and were not included in this analysis. The ANOVA yielded a significant main effect of accent, $F_2(1, 45) = 28.0, p < .001$, partial $\eta^2 = .384$. The mean normalized number of misperception tokens for native stimuli, $M = .673$, $SE = 0.037$, 95% CI [0.599, 0.747], was higher than that of the accented stimuli, $M = .370$, $SE = 0.040$, 95% CI [0.290, 0.450]. The main effect of neighborhood density was also significant, $F_2(1, 45) = 4.98, p = .030$, partial $\eta^2 = .100$. The mean normalized number of misperception token for sparse words, $M = .578$, $SE = .036$, 95% CI [0.505, 0.651], was higher than that of the dense words, $M = .465$, $SE = .035$, 95% CI [0.394, 0.536]. The interaction was not significant, $F_2(1, 45) = 0.460, p = .831$. Overall, rather than activating a wide variety of competitors, foreign accents tended to activate a smaller set of competitors.

Next, we examined how activation was distributed among the lexical competitors during foreign-accented word recognition. Does one particular lexical competitor tend to be dominantly activated, or are all the lexical competitors equivalently activated? If one particular lexical competitor is dominantly activated, we expected the listeners to reach consensus on the misperceptions (i.e., all of the listeners reported hearing the same wrong word). If all the lexical competitors are equivalently activated, then a range of misperceptions are likely to be reported across the listeners. To examine the distribution of misperceptions across all the misperception tokens, a relative frequency distribution was created for each stimulus in each accent condition by grouping all its identification responses into classes, including the correct response, and each of the misperception tokens reported (ordered from the most frequently reported one to the least frequently reported one). For example, if the word *shear* was misidentified as *hear* by 8 participants, *sear* by 2 participants, and *fear* by 1 participant, then *hear* was considered to be the most frequently reported misperception token, and *fear* was considered to be the least frequently reported misperception token. For easier comparison, the frequency distribution of responses (relative frequency in parentheses) across the correct responses and all the misperception tokens were pooled across stimuli in each neighborhood density and accent condition and are shown in Table 3.

A chi-square test of independence was used to determine whether misperceptions for the native and accented conditions

Table 3
The Frequency Distribution of Responses (Relative Frequency in Parentheses) Across the Correct Responses and All the Misperception Tokens in Each Neighborhood Density and Accent Condition

Accent	Neighborhood density	Correct responses	Most frequently reported misperception token	2nd most frequently reported misperception token	10 (1.1)	7 (0.8)	2 (0.2)	1 (0.1)	0 (0)	0 (0)	0 (0)	0 (0)	0 (0)	0 (0)	2nd least frequently reported misperception token	Least frequently reported misperception token
Native	Dense	785 (88)	69 (7.7)	22 (2.5)	10 (1.1)	7 (0.8)	2 (0.2)	1 (0.1)	0 (0)	0 (0)	0 (0)	0 (0)	0 (0)	0 (0)	0 (0)	0 (0)
	Sparse	750 (84)	80 (8.9)	25 (2.8)	13 (1.5)	10 (1.1)	7 (0.8)	5 (0.6)	3 (0.3)	1 (0.1)	1 (0.1)	1 (0.1)	0 (0)	0 (0)	0 (0)	0 (0)
	Total	1535 (86)	149 (8.3)	47 (2.7)	23 (1.3)	17 (1.0)	9 (0.5)	6 (0.3)	3 (0.2)	1 (0.1)	1 (0.1)	1 (0.1)	0 (0)	0 (0)	0 (0)	0 (0)
Accented	Dense	452 (50)	367 (41)	50 (5.6)	16 (1.8)	7 (0.8)	3 (0.3)	1 (0.1)	0 (0)	0 (0)	0 (0)	0 (0)	0 (0)	0 (0)	0 (0)	0 (0)
	Sparse	484 (54)	247 (28)	69 (7.7)	33 (3.7)	18 (2.0)	10 (1.1)	8 (0.9)	6 (0.7)	3 (0.3)	3 (0.3)	3 (0.3)	1 (0.1)	1 (0.1)	1 (0.1)	1 (0.1)
	Total	936 (52)	614 (35)	119 (6.6)	49 (2.7)	25 (1.4)	13 (0.7)	9 (0.5)	6 (0.3)	6 (0.3)	3 (0.2)	3 (0.2)	3 (0.2)	1 (0.1)	1 (0.1)	1 (0.1)

were distributed differently across the misperception tokens. To avoid low expected frequency in the contingency table, frequency counts for the fourth most frequently reported to the least frequently reported misperception tokens were pooled together for all of the following analyses. The chi-square test was significant, $\chi^2(3, N = 1,113) = 18.6, p < .001$, suggesting a statistically significant difference in the distribution of misperception responses for the native and accented conditions. The different misperception distribution pattern across the native and accented conditions was probably due to the higher tendency of listeners to consistently report the same misperception for the target words in the accented condition. The most frequently reported misperception token accounted for on average 71.7% of the total misperceptions for the accented condition, but only on average 58.0% of the total misperceptions for the native condition. The ratio of the first-most-reported and the second-most-reported misperception tokens is 5.15 for the accented condition and 3.17 for the native condition respectively, confirming that a majority of the misperceptions for accented stimuli was disproportionately accounted for by the dominantly activated misperception token.

Another chi-square test was used to analyze the distribution of misperceptions for the accented dense and sparse words. A statistically significant difference was found in the distribution of misperception responses for the accented dense and sparse words, $\chi^2(3, N = 856) = 67.8, p < .001$, indicating that listeners had a higher tendency to consistently report the same misperception for the dense target words in the presence of foreign accents. The most frequently reported misperception token accounted for on average 82.7% of the total misperceptions for the accented dense words, but only on average 60.0% of the total misperceptions for the accented sparse words. What were these most frequently reported misperception tokens? In the accented condition, they were phonological neighbors of the target words 93.0% of the time for dense words and 72.0% of the time for sparse words. In the native condition, they were phonological neighbors of the target words 87.5% of the time for dense words and 62.5% of the time for sparse words (see also Vitevitch, Chan, & Goldstein, 2014).

To further understand how listeners might systematically misperceive the foreign-accented words, the consonants and vowels that were consistently misperceived by a majority of the listeners was examined. It is important to note that the current set of stimuli does not contain all the consonants and vowels in English. Misperceptions that were consistently reported by at least seven or more listeners, which constitute 66.7% of the total misperceptions, were included in the following analysis. Table 4 shows the pattern of consonant and vowel confusions by the listeners and the percentage of time such phoneme substitutions occurred. All the listed consonant misperceptions occurred in the final position, except for the addition of /h/ in the onset position.

In Table 4, it can be seen that the majority of consonant misperceptions involved voiced target consonants being misperceived as their voiceless counterparts, including /b/ to /p/, /d/ to /t/, /g/ to /k/, /v/ to /f/, and /z/ to /s/. Also, most of the consonant misperceptions were from stops (including /b/, /d/, and /g/), fricatives (including /v/, /s/, and /z/), and nasals (including /n/, /ŋ/, and /m/). Furthermore, back vowels tended to be misperceived as each other, such as /ɔ/ being misperceived as /o/ and /u/, /o/ being misperceived as /u/, and /u/ being misperceived as /u/. Low vowels tended to be misperceived as each other, such as /æ/ being mis-

Table 4

Phonemic Substitutions for Consonants and Vowels (the Percentage of Time Such Substitution Occurred is Shown in Parentheses) Found in the Misperceptions Consistently Reported by Seven or More Listeners in the Accented Condition

Listeners' Responses (Percentage of Occurrence)				
Target Consonant				
b	p (87.5)			
d	l (9.82)	t (72.3)		
g	k (78.6)			
m	n (7.14)			
n	ŋ (14.3)	m (5.40)		
s	z (32.1)			
v	f (42.9)			
z	s (37.5)			
	h (25.0)			
	p (57.1)			
Target Vowel				
æ	ʌ (23.8)			
ʌ	æ (10.7)	ɔ (4.08)		aɪ (1.79)
a	ʌ (29.2)			
ɔ	ʌ (6.70)	o (23.2)		ʊ (4.02)
i	ɪ (12.5)			
o	ʊ (26.8)			
u	ʊ (6.25)			

perceived as /ʌ/ and vice versa, /ʌ/ being misperceived as /ɔ/ and vice versa, and /a/ being misperceived as /ʌ/. Also, /i/ was misperceived as /i/.

Discussion

As expected, participants identified foreign-accented words with lower accuracy than native words, and dense words with a lower accuracy than sparse words. The lack of an interaction between neighborhood density and accent contrasted with the findings of Imai et al. (2005), in which a significant neighborhood density effect was found during foreign-accented word recognition in native listeners, but (curiously) not during native word recognition. The current result is actually more consistent with previous findings of a robust neighborhood density effect in perceptual identification tasks with native-produced words (Luce & Pisoni, 1998; Vitevitch et al., 2008). The possible reason why the neighborhood density effect was not enhanced in the foreign-accented condition is that the neighborhood density effect was strong in the native condition in this experiment due to the noisy listening condition. Therefore, it might become comparable with the neighborhood density effect present in the foreign-accented condition.

A more detailed analysis of the misperceptions was conducted to examine the alternative word candidates that were activated during lexical competition. After normalizing for the raw number of misperceptions, the mean number of misperception tokens for native stimuli was higher than that of the accented stimuli; the mean number of misperception tokens for sparse words was higher than that of the dense words. These results suggest that accented speech does not tend to activate extra lexical candidates compared with native speech (i.e., the many-additional-competitor scenario); it actually activates a smaller set of lexical competitors. Moreover, regardless of accents,

dense words generally activate a smaller set of lexical competitors than sparse words.

The distribution of misperceptions for the native and accented words was also different. The most frequently reported misperception tokens accounted for a majority (71%) of the erroneous responses for accented words, but a lower proportion (58%) for native words. This suggests that listeners were pretty consistent in reporting the same misperception for the same accented target words, whereas listeners in the native condition misperceived the target words as being one of many possible words. That is, accented-speech signals tend to strongly activate one particular (but wrong) lexical candidate instead of weakly activating many different lexical candidates. This strongly activated word tends to be a phonological neighbor of the target word, and prevents the target word from winning the competition to be recognized, much like interlopers may keep one from resolving a tip of the tongue state (Brown & McNeill, 1966). That was particularly likely for accented dense words, in which 82.7% of the erroneous responses were accounted for by the most frequently reported misperception tokens, compared with only 60.0% for the accented sparse words.

These findings, along with the lower identification accuracy rates found for dense words than sparse words in the presence of foreign accents, indirectly verified that the higher accuracy rates for accented dense words than sparse words found in Experiment 1 was probably not due to more accurate identification of the dense words. Instead, it may be due to the nature of the lexical decision task, which only requires the participants to indicate whether the stimulus is a real word or not.

Listeners in the present experiment tended to systematically misperceive voiced stops and fricatives as their corresponding voiceless counterparts, and showed confusion between the nasals. These systematic consonant substitutions are consistent with the pronunciation variations found in Spanish-accented English spoken by young children, including /d/ to /t/, /ŋ/ to /n/, /v/ to /f/, and /z/ to /s/ (You et al., 2005). Listeners in the present experiment also tended to have difficulty distinguishing among the back vowels (/ɔ/, /o/, /u/, and /ʊ/), the low vowels (/æ/, /ʌ/, /ɔ/, and /a/), as well as the high front vowels /i/ and /i/. Previous studies examining Spanish-accented word identification also found /æ/, /ʌ/, /a/, and /i/ to be highly confusable for native listeners (Sidasas et al., 2009; Wade et al., 2007). Taken together, these commonly misperceived phonemes are those that exist in English but not in the Spanish phonetic inventory, such as /v/, /z/, /i/, /æ/, and /ʌ/ (Sidasas et al., 2009; Wade et al., 2007; You et al., 2005). This systematic phoneme substitution pattern reflects the systematic variations introduced by the speaker's native phonetic inventories and is probably specific to Spanish-accented speech only.

General Discussion

The goals of the current study were to examine (a) whether recognition of foreign-accented words are differentially influenced by neighborhood density, and (b) whether foreign accents activate many additional competitors (the many-additional-competitor scenario) or strongly activate one particular competitor (the single-strong-competitor scenario) such that lexical discrimination is more difficult in dense words. Using a lexical decision task with noise-free stimuli, Experiment 1 found an increased accent-induced processing cost, in terms of RT, especially for words with many similar sounding words, implying that word recognition is undermined for foreign-accented

dense words. Analyses of the misperceptions from a perceptual identification task in Experiment 2 showed that the mean number of misperception tokens was higher for native stimuli relative to the accented stimuli, as well as for sparse words relative to the dense words. The distribution of misperceptions for the native and accented words was also different. A majority (71.7%) of the erroneous responses in the accented condition was accounted for by the most frequently reported misperception tokens, which tend to be phonological neighbors of the target words. There were systematic phonemic substitutions in the misperceptions of Spanish-accented words, including nasals, voiced stops, and fricatives, as well as low, back, and high front vowels.

The current study advances previous findings of how foreign accents disrupt the recognition of words varying in neighborhood density by showing precisely how activation of lexical candidates is affected. Findings from the current study showed that mismatches induced by foreign accents do not weakly activate a large number of additional competitors (the many-additional-competitor scenario). Instead, foreign accents tend to strongly activate one particular word competitor (or a smaller set of words) that is phonologically similar to the target word during lexical retrieval. That alternative competitor might get more strongly activated than the target word and win the competition for recognition. More importantly, the present findings demonstrated that not all the neighbors containing phonemes confusable with those of the target word are equally activated as competitors. Consistent with NAM (Luce & Pisoni, 1998), the current results suggested that the accented acoustic-phonetic information plays a partial but important role in determining which competitors are strongly activated, influencing the probability of identification of the target words.

Like other previous studies that examined native listeners' identification of Spanish-accented words (Sidas et al., 2009; Wade et al., 2007; You et al., 2005), systematic phonemic substitutions were found in the misperceptions in this study. These reflect systematic misarticulations of particular L2 sounds by the foreign-accented speakers. This implies that foreign-accented speech contains lawful accent-general acoustic-phonetic variations that are introduced by the interaction of the phonetic inventories of the accented speakers' L1 (native language) and L2 (Norris et al., 2003). L2 learners' perception and production of L2 sounds is "filtered" through the phonological system of their native language (Polivanov, 1931). L2 speakers might misarticulate certain L2 sounds in a consistent way as they stick with their more familiar L1 articulatory habits when they produce L2 sounds. In addition to neighborhood density, these systematic acoustic-phonetic deviations induced by the foreign accent seem to play a crucial role in determining which neighbor of the target word is activated as a strong competitor. For example, if the Spanish-accented speaker tends to consistently misarticulate the final consonant /b/ in *cab* as /p/, only *cap* will be strongly activated as the competitor despite the existence of many other neighbors that are also confusable with the target word in the final phoneme, such as *cat*, *cad*, *cash*, and *calf*.

The systematic phonemic substitution patterns found in the current study largely conform to predictions from influential models of non-native speech perception—Perceptual Assimilation Model (PAM; Best, 1994) and Speech Learning Model (SLM; Flege, 1995, 2002). Both PAM and SLM predict that perception and production of a L2 phoneme will be hindered if the L2 learners classify that L2 phoneme to be an instance of an existing L1 phonemic category that is close in

phonetic space but different in articulatory realization. Inferred from the phoneme substitution patterns, the Spanish-accented speaker in this study tended to substitute English phonemes that do not exist in Spanish with similar phonemes in Spanish, such as /v/ with /f/, /z/ with /s/, and /l/ with /i/. This showed that our speaker "equated" or assimilated these unfamiliar non-native phonemes to the most similar counterparts from his L1. The voiced stops were also consistently misarticulated as their voiceless counterparts, including /b/ to /p/, /d/ to /t/, /g/ to /k/, suggesting that the speaker assimilated the English stops to the Spanish stops regardless of their differences in voice onset time and place of articulation.

The present findings also nicely complemented existing research on perceptual learning of accented speech. Dahan, Drucker, and Scarborough (2008) used a perceptual learning paradigm with a visual world eye-tracking task to examine whether listeners can use a newly learned feature of a regional accent to eliminate lexical competitors in subsequent speech recognition. Participants were exposed to NSs with a regional accent, in which the vowel /æ/ was raised to /ɛ/ before /g/ (e.g., *bag*), but not before /k/ (e.g., *back*). The participants' eye gaze on four written words (e.g., *bag*, *back*, *wig*, and *wick*) on the computer screen was monitored upon hearing the stimulus /bæ/ of *back*. Results showed that participants tended to fixate on *back* and ruled out *bag* as a competitor, demonstrating successful adaptation to the regional accent through dynamic adjustment of their representations. However, when Trude et al. (2013) adopted a similar paradigm to study perceptual learning of a particular vowel shift in French-accented English words, they found strikingly limited improvement. Listeners failed to learn the knowledge of the foreign accent and use it to rule out lexical competitors during subsequent speech recognition.

Results from these perceptual learning studies imply that when certain accent-related acoustic deviations result in a phonemic difference, the resulting phonemic mismatch could be disruptive to the lexical retrieval process by activating a strong lexical candidate competing with the target words. This finding is consistent with our findings that foreign-accented words are difficult to recognize because one (or a small set of) alternative word(s) strongly competes with the target word during lexical retrieval. Therefore, when phonemic categories are successfully recalibrated through perceptual learning of helpful accent distinction cues, as shown in the study by Dahan et al. (2008), the original competitor will no longer be strongly activated and recognition accuracy improves.¹

The current study along with the perceptual learning study by Trude et al. (2013) demonstrate that it is beneficial to examine the set of words competing with the target words during foreign-accented word recognition. The visual world eye-tracking paradigm used by Trude et al. (2013) is an online measure that tracks the listener's eye fixations to both the target and the competitors. It examines the time course of the lexical retrieval process during foreign-accented word recognition. In contrast, the current study relied on an offline measure—misperceptions in the perceptual identification task. This offline measure does not allow us to determine when the alternative competitors are strongly activated, or how strong their activation is over time relative to the target word. Therefore, it cannot tell us when the listeners arrive at that erroneous identification response. However,

¹ Note that in the current study, listeners might adapt to the Spanish accent. But their recognition accuracies were still low, as our Spanish-accented speaker did not strongly distinguish certain phonemes in his production.

misperception analysis seemed to be a better methodology for the current research questions—do foreign accents tend to activate a large or a small set of competitors compared with the same word in native speech? What are these competitors? The visual world eye-tracking paradigm typically tracks the participants' fixation across four words or pictures denoting the target word, the potential competitors, and filler words. Thus, the visual world paradigm would not be suitable under certain conditions like when the number and the identity of the competitors are unknown, as in the current study.

The present findings also provide valuable human behavioral data for computational modeling of recognition of foreign-accented speech. Currently, the major models of SWR, including TRACE (McClelland & Elman, 1986), Shortlist (Norris, 1994), and PARSYN (Luce, Goldinger, Auer, & Vitevitch, 2000), fail to account for the recognition of foreign-accented speech as they take in speech input that is abstract and idealized. Recently, Scharenborg, Wittenman, and Weber (2012) used Fine-Tracker, a computational model of human SWR that used real speech as input to simulate human recognition of foreign-accented words (Scharenborg, 2010). Fine-Tracker failed to simulate human listeners' performance on an item-by-item basis. However, it was in line with human listeners in recognizing words varying in degree of accentedness as well as showing slight improvement after brief exposure with the accent (Scharenborg et al., 2012). The current study examined the impact of foreign accent on the lexical retrieval process and provides insights into how the error recovery process operates in human listeners facing foreign accents. The rich misperception data from the current study could provide valuable human behavioral data for computational models of recognition of foreign-accented speech to simulate and compare with.

In sum, the current study suggested that foreign accents do not tend to activate a larger set of word candidates. Instead, one particular word (or a small set of words) that is phonologically similar to the target word is likely to be strongly activated during lexical retrieval and is mistaken as the target words. Also, certain pairs of phonemes were found to be highly confusable and systematically misperceived. The current study used only one speaker with a Spanish accent. Are these findings generalizable to other speakers or other foreign accents? All foreign-accented speech contains acoustic-phonetic deviations from native speech. However, foreign-accented speech spoken by speakers sharing the same native language contains systematic variations introduced by the speakers' native phonetic inventory (Sidas et al., 2009). Therefore, it is likely that the current findings are generalizable to different foreign accents, except for the specific phoneme substitution patterns in the misperceptions, which are likely to be specific to only the Spanish accent. Furthermore, the phoneme substitution patterns found in the current studies were consistent with previous studies also examining Spanish-accented speech (Sidas et al., 2009; Wade et al., 2007; You et al., 2005). Baese-Berk, Bradlow, and Wright (2013) found that listeners can even learn systematic variability in foreign-accent speech spoken by speakers with multiple language backgrounds. Hence, future research that uses other foreign-accented speech is required to examine whether all foreign accents impact the lexical activation and selection process in a similar way.

References

Baese-Berk, M. M., Bradlow, A. R., & Wright, B. A. (2013). Accent-independent adaptation to foreign accented speech. *The Journal of the*

Acoustical Society of America, 133, EL174–EL180. <http://dx.doi.org/10.1121/1.4789864>

Best, C. T. (1994). The emergence of native-language phonological influences in infants: A perceptual assimilation model. In J. C. Goodman & H. C. Nusbaum (Eds.), *The development of speech perception: The transition from speech sounds to spoken words* (pp. 167–224). Cambridge, MA: MIT Press.

Boersma, P., & Weenink, D. (2009). Boersma, P., & Weenink, D. Praat: Doing phonetics by computer (Version 5.1.05) [Computer program]. Amsterdam, The Netherlands Retrieved from <http://www.praat.org/>

Bradlow, A. R., & Bent, T. (2008). Perceptual adaptation to non-native speech. *Cognition*, 106, 707–729. <http://dx.doi.org/10.1016/j.cognition.2007.04.005>

Brown, R., & McNeill, D. N. (1966). The “tip of the tongue” phenomenon. *Journal of Verbal Learning & Verbal Behavior*, 5, 325–337. [http://dx.doi.org/10.1016/S0022-5371\(66\)80040-3](http://dx.doi.org/10.1016/S0022-5371(66)80040-3)

Bürki-Cohen, J., Miller, J. L., & Eimas, P. D. (2001). Perceiving non-native speech. *Language and Speech*, 44, 149–169. <http://dx.doi.org/10.1177/00238309010440020201>

Clark, H. H. (1973). The language-as-fixed-effect fallacy: A critique of language statistics in psychological research. *Journal of Verbal Learning & Verbal Behavior*, 12, 335–359. [http://dx.doi.org/10.1016/S0022-5371\(73\)80014-3](http://dx.doi.org/10.1016/S0022-5371(73)80014-3)

Clarke, C. M., & Garrett, M. F. (2004). Rapid adaptation to foreign-accented English. *The Journal of the Acoustical Society of America*, 116, 3647–3658. <http://dx.doi.org/10.1121/1.1815131>

Cluff, M. S., & Luce, P. A. (1990). Similarity neighborhoods of spoken two-syllable words: Retroactive effects on multiple activation. *Journal of Experimental Psychology: Human Perception and Performance*, 16, 551–563. <http://dx.doi.org/10.1037/0096-1523.16.3.551>

Dahan, D., Drucker, S. J., & Scarborough, R. A. (2008). Talker adaptation in speech perception: Adjusting the signal or the representations? *Cognition*, 108, 710–718. <http://dx.doi.org/10.1016/j.cognition.2008.06.003>

Flège, J. E. (1984). The detection of French accent by American listeners. *The Journal of the Acoustical Society of America*, 76, 692–707. <http://dx.doi.org/10.1121/1.391256>

Flège, J. E. (1995). Second-language speech learning: Theory, findings, and problems. In W. Strange (Ed.), *Speech perception and linguistic experience: Issues in cross-language research* (pp. 233–277). Timonium, MD: York Press.

Flège, J. E. (2002). Interactions between the native and second-language phonetic systems. In P. Burmeister, T. Piske & A. Rohde (Eds.), *An integrated view of language development: Papers in honor of Henning Wode* (pp. 217–244). Trier: Wissenschaftlicher Verlag.

Goldinger, S. D., Luce, P. A., & Pisoni, D. B. (1989). Priming lexical neighbors of spoken words: Effects of competition and inhibition. *Journal of Memory and Language*, 28, 501–518. [http://dx.doi.org/10.1016/0749-596X\(89\)90009-0](http://dx.doi.org/10.1016/0749-596X(89)90009-0)

Greenberg, J. H., & Jenkins, J. J. (1967). Studies in the psychological correlates of the sound system of American English. In L. A. Jakobovits & M. S. Miron (Eds.), *Readings in the psychology of language* (pp. 186–200). Englewood Cliffs, NJ: Prentice Hall.

Imai, S., Walley, A. C., & Flège, J. E. (2005). Lexical frequency and neighborhood density effects on the recognition of native and Spanish-accented words by native English and Spanish listeners. *The Journal of the Acoustical Society of America*, 117, 896–907. <http://dx.doi.org/10.1121/1.1823291>

Landauer, T. K., & Streeter, L. A. (1973). Structural differences between common and rare words: Failure of equivalence assumptions for theories of word recognition. *Journal of Verbal Learning & Verbal Behavior*, 12, 119–131. [http://dx.doi.org/10.1016/S0022-5371\(73\)80001-5](http://dx.doi.org/10.1016/S0022-5371(73)80001-5)

Lane, H. (1963). Foreign accent and speech distortion. *The Journal of the Acoustical Society of America*, 35, 451–453. <http://dx.doi.org/10.1121/1.1918501>

- Luce, P. A., Goldinger, S. D., Auer, E. T., Jr., & Vitevitch, M. S. (2000). Phonetic priming, neighborhood activation, and PARSYN. *Perception & Psychophysics*, 62, 615–625. <http://dx.doi.org/10.3758/BF03212113>
- Luce, P. A., & Pisoni, D. B. (1998). Recognizing spoken words: The neighborhood activation model. *Ear and Hearing*, 19, 1–36. <http://dx.doi.org/10.1097/00003446-199802000-00001>
- Magen, H. S. (1998). The perception of foreign-accented speech. *Journal of Phonetics*, 26, 381–400. <http://dx.doi.org/10.1006/jpho.1998.0081>
- McClelland, J. L., & Elman, J. L. (1986). The TRACE model of speech perception. *Cognitive Psychology*, 18, 1–86. [http://dx.doi.org/10.1016/0010-0285\(86\)90015-0](http://dx.doi.org/10.1016/0010-0285(86)90015-0)
- Munro, M. J. (1998). The effects of noise on the intelligibility of foreign-accented speech. *Studies in Second Language Acquisition*, 20, 139–154. <http://dx.doi.org/10.1017/S0272263198002022>
- Munro, M. J., & Derwing, T. M. (1995a). Foreign accent, comprehensibility, and intelligibility in the speech of second language learners. *Language Learning*, 45, 73–97. <http://dx.doi.org/10.1111/j.1467-1770.1995.tb00963.x>
- Munro, M. J., & Derwing, T. M. (1995b). Processing time, accent, and comprehensibility in the perception of native and foreign-accented speech. *Language and Speech*, 38, 289–306.
- Munro, M. J., & Derwing, T. M. (1998). The effects of speaking rate on listener evaluations of native and foreign-accented speech. *Language Learning*, 48, 159–182. <http://dx.doi.org/10.1111/1467-9922.00038>
- Norris, D. (1994). Shortlist: A connectionist model of continuous speech recognition. *Cognition*, 52, 189–234. [http://dx.doi.org/10.1016/0010-0277\(94\)90043-4](http://dx.doi.org/10.1016/0010-0277(94)90043-4)
- Norris, D., McQueen, J. M., & Cutler, A. (2003). Perceptual learning in speech. *Cognitive Psychology*, 47, 204–238. [http://dx.doi.org/10.1016/S0010-0285\(03\)00006-9](http://dx.doi.org/10.1016/S0010-0285(03)00006-9)
- Nusbaum, H. C., Pisoni, D. B., & Davis, C. K. (1984). Sizing up the Hoosier Mental Lexicon: Measuring the familiarity of 20,000 words. *Research on Speech Perception Progress Report*, 10, 357–376.
- Owren, M. J. (2008). GSU Praat Tools: Scripts for modifying and analyzing sounds using Praat acoustics software. *Behavior Research Methods*, 40, 822–829. <http://dx.doi.org/10.3758/BRM.40.3.822>
- Perception Research Systems. (2007). *Paradigm*. Retrieved from <http://www.paradigmexperiments.com>
- Polivanov, E. (1931). *La perception des sons d'une langue étrangère. Travaux du Cercle Linguistique de Prague 4*: 79–96. [English translation: The subjective nature of the perceptions of language sounds. In E.D. Polivanov (1974): *Selected works: articles on general linguistics*. The Hague: Mouton. 223–237].
- Raaijmakers, J. G. W. (2003). A further look at the “language-as-fixed-effect fallacy.” *Canadian Journal of Experimental Psychology/Revue canadienne de psychologie expérimentale*, 57, 141–151. <http://dx.doi.org/10.1037/h0087421>
- Raaijmakers, J. G. W., Schrijnemakers, J. M. C., & Gremmen, F. (1999). How to deal with “The language-as-fixed-effect fallacy”: Common misconceptions and alternative solutions. *Journal of Memory and Language*, 41, 416–426. <http://dx.doi.org/10.1006/jmla.1999.2650>
- Rastle, K., Harrington, J., & Coltheart, M. (2002). 358,534 nonwords: The ARC nonword database. *The Quarterly Journal of Experimental Psychology A: Human Experimental Psychology*, 55, 1339–1362. <http://dx.doi.org/10.1080/02724980244000099>
- Reinisch, E., & Holt, L. L. (2013). Lexically guided phonetic retuning of foreign-accented speech and its generalization. *Journal of Experimental Psychology: Human Perception and Performance*, 40, 539–555.
- Scharenborg, O. (2010). Modeling the use of durational information in human spoken-word recognition. *The Journal of the Acoustical Society of America*, 127, 3758–3770. <http://dx.doi.org/10.1121/1.3377050>
- Scharenborg, O., Witteman, M. J., & Weber, A. (2012). *Computational modelling of the recognition of foreign-accented speech*. Paper presented at Interspeech 2012.
- Schmid, P. M., & Yeni-Komshian, G. H. (1999). The effects of speaker accent and target predictability on perception of mispronunciations. *Journal of Speech, Language, and Hearing Research*, 42, 56–64. <http://dx.doi.org/10.1044/jslhr.4201.56>
- Sidaras, S. K., Alexander, J. E. D., & Nygaard, L. C. (2009). Perceptual learning of systematic variation in Spanish-accented speech. *The Journal of the Acoustical Society of America*, 125, 3306–3316. <http://dx.doi.org/10.1121/1.3101452>
- Storkel, H. L., & Hoover, J. R. (2010). An online calculator to compute phonotactic probability and neighborhood density on the basis of child corpora of spoken American English. *Behavior Research Methods*, 42, 497–506. <http://dx.doi.org/10.3758/BRM.42.2.497>
- Trude, A. M., Tremblay, A., & Brown-Schmidt, S. (2013). Limitations on adaptation to foreign accents. *Journal of Memory and Language*, 69, 349–367. <http://dx.doi.org/10.1016/j.jml.2013.05.002>
- van Wijngaarden, S. J. (2001). Intelligibility of native and non-native Dutch speech. *Speech Communication*, 35, 103–113. [http://dx.doi.org/10.1016/S0167-6393\(00\)00098-4](http://dx.doi.org/10.1016/S0167-6393(00)00098-4)
- Vitevitch, M. S. (2002). Naturalistic and experimental analyses of word frequency and neighborhood density effects in slips of the ear. *Language and Speech*, 45, 407–434. <http://dx.doi.org/10.1177/00238309020450040501>
- Vitevitch, M. S. (2003). The influence of sublexical and lexical representations on the processing of spoken words in English. *Clinical Linguistics & Phonetics*, 17, 487–499. <http://dx.doi.org/10.1080/0269920031000107541>
- Vitevitch, M. S., Chan, K. Y., & Goldstein, R. (2014). Insights into failed lexical retrieval from network science. *Cognitive Psychology*, 68, 1–32. <http://dx.doi.org/10.1016/j.cogpsych.2013.10.002>
- Vitevitch, M. S., & Luce, P. A. (1998). When words compete: Levels of processing in perception of spoken words. *Psychological Science*, 9, 325–329. <http://dx.doi.org/10.1111/1467-9280.00064>
- Vitevitch, M. S., & Luce, P. A. (1999). Probabilistic phonotactics and neighborhood activation in spoken word recognition. *Journal of Memory and Language*, 40, 374–408. <http://dx.doi.org/10.1006/jmla.1998.2618>
- Vitevitch, M. S., Stamer, M. K., & Sereno, J. A. (2008). Word length and lexical competition: Longer is the same as shorter. *Language and Speech*, 51, 361–383. <http://dx.doi.org/10.1177/0023830908099070>
- Wade, T., Jongman, A., & Sereno, J. (2007). Effects of acoustic variability in the perceptual learning of non-native-accented speech sounds. *Phonetica*, 64, 122–144. <http://dx.doi.org/10.1159/000107913>
- Witteman, M. J., Weber, A., & McQueen, J. M. (2013). Foreign accent strength and listener familiarity with an accent codetermine speed of perceptual adaptation. *Attention, Perception, & Psychophysics*, 75, 537–556. <http://dx.doi.org/10.3758/s13414-012-0404-y>
- Witteman, M. J., Weber, A., & McQueen, J. M. (2014). Tolerance for inconsistency in foreign-accented speech. *Psychonomic Bulletin & Review*, 21, 512–519.
- You, H., Alwan, A., Kazemzadeh, A., & Narayanan, S. (2005). *Pronunciation variations of Spanish-accented English spoken by young children*. Paper presented at the Interspeech 2005—Eurospeech, Lisbon, Portugal.

(Appendices follow)

Appendix A
Dense Neighborhood Words Used in Experiment 1

Stimulus word	Neighborhood density	Familiarity	log Word frequency	Pos. Seg. Freq.	Biphone Freq.	log NF
bug	26	7.00	0.60	0.1083	0.0047	1.80
buck	29	7.00	1.30	0.1439	0.0053	1.83
bought	25	7.00	1.75	0.1337	0.0025	2.34
duck	25	6.75	0.95	0.1445	0.0043	1.81
dumb	29	7.00	1.11	0.1404	0.0075	2.01
dune	27	7.00	0.00	0.1700	0.0043	1.96
dine	30	7.00	0.30	0.1822	0.0082	2.04
fun	25	7.00	1.64	0.1819	0.0067	2.06
fall	26	7.00	2.17	0.1368	0.0050	2.45
fine	28	6.92	2.21	0.1770	0.0065	2.12
cop	30	7.00	1.18	0.1903	0.0191	1.84
call	26	7.00	2.27	0.1829	0.0060	2.16
lash	26	6.17	0.78	0.1212	0.0065	1.63
lock	31	7.00	1.36	0.1481	0.0052	1.93
lead	31	7.00	2.42	0.1450	0.0077	2.10
lease	27	6.92	1.00	0.1447	0.0042	2.02
leave	26	7.00	2.31	0.0895	0.0038	1.76
kneel	27	7.00	0.70	0.1293	0.0044	1.74
nip	25	7.00	0.48	0.1571	0.0068	1.66
pop	29	7.00	0.90	0.1820	0.0103	1.69
rash	26	6.58	0.00	0.1372	0.0070	1.63
raise	30	7.00	1.72	0.0994	0.0042	1.86
son	26	7.00	2.44	0.2377	0.0116	2.43
seek	31	6.92	1.84	0.1877	0.0050	2.07
shear	26	7.00	1.74	0.1843	0.0062	2.19
shore	28	7.00	1.79	0.1374	0.0188	2.60
tuck	28	6.83	0.30	0.1372	0.0033	1.95
tall	27	7.00	1.74	0.1347	0.0044	2.25
tune	27	7.00	1.00	0.1627	0.0047	52.25
wed	25	7.00	0.30	0.1312	0.0061	2.37
wick	26	6.67	0.60	0.1700	0.0132	2.41
wine	30	7.00	1.86	0.1507	0.0064	1.98

Note. Pos. Seg. Freq. = position segment frequency (a measure of phonotactic probability); Biphone Freq. = biphone frequency (a measure of phonotactic probability); NF = neighborhood frequency.

(Appendices continue)

Appendix B

Sparse Neighborhood Words Used in Experiment 1

Stimulus word	Neighborhood density	Familiarity	log Word frequency	Pos. Seg. Freq.	Biphone Freq.	log NF
buzz	15	7.00	1.11	0.1105	0.0044	1.92
bib	13	6.83	0.30	0.1734	0.0064	2.25
beam	16	6.92	1.32	0.1324	0.0034	2.19
dash	15	6.92	1.04	0.1389	0.0039	1.52
dawn	19	7.00	1.45	0.1644	0.0022	2.10
deed	18	7.00	0.90	0.1216	0.0052	2.33
deep	18	7.00	2.04	0.1207	0.0045	1.97
fad	19	6.33	0.30	0.1640	0.0058	2.21
far	18	6.58	2.63	0.1855	0.0180	2.43
fish	13	7.00	1.54	0.1505	0.0059	1.87
gone	17	7.00	2.29	0.1386	0.0020	1.81
cup	18	7.00	1.65	0.1690	0.0055	2.07
calm	17	7.00	1.54	0.2026	0.0224	1.95
kiss	13	7.00	1.23	0.2677	0.0188	2.34
lull	15	6.25	0.30	0.1470	0.0064	1.66
love	11	6.67	2.37	0.0969	0.0030	1.91
lawn	19	7.00	1.18	0.1467	0.0030	2.41
league	19	7.00	1.84	0.0838	0.0030	1.86
null	17	6.17	1.11	0.1367	0.0060	1.51
neck	13	7.00	1.91	0.1502	0.0094	1.74
pool	18	7.00	2.05	0.1802	0.0018	1.99
wreck	18	7.00	0.90	0.1765	0.0156	1.91
robe	18	7.00	0.78	0.1254	0.0039	1.94
shun	19	6.33	0.00	0.1450	0.0062	2.24
psalm	11	6.92	0.60	0.2123	0.0080	2.27
chute	17	7.00	1.46	0.0978	0.0029	1.97
sour	10	6.92	0.48	0.1905	0.0009	1.76
tug	18	7.00	0.48	0.1016	0.0027	1.71
tape	16	7.00	1.54	0.1108	0.0029	2.08
town	14	7.00	2.33	0.1503	0.0045	2.19
walk	15	7.00	2.00	0.0903	0.0030	2.20
wipe	14	7.00	1.00	0.0917	0.0030	2.19

Note. Pos. Seg. Freq. = position segment frequency (a measure of phonotactic probability); Biphone Freq. = biphone frequency (a measure of phonotactic probability); NF = neighborhood frequency.

(Appendices continue)

Appendix C
International Phonetic Alphabet Transcriptions of the Nonword Foils Used in Experiment 1

bʌtʃ	bʌb
bʌp	bɪdʒ
bʊf	bil
dʌʃ	dæz
dʌt	duf
duθ	dik
dʌɪp	dit
fʌm	fæf
fʊθ	fɒn
fʌɪb	fɪd
kʊz	guð
kub	kʌk
læt	kɒŋ
lɒd	kɪg
lɛm	lʌt
lið	lʌθ
lim	lok
niv	lib
nɪθ	nʌs
pɒg	nɛd
ræd	puk
rem	rɛl
sʌv	rof
sig	ʃʌŋ
ʃɪk	sɒg
ʃof	ʃul
tʌdʒ	saʊt
tɒb	tʌl
tuv	ten
weg	taʊs
wɪd	wuk
wʌɪm	wʌɪb

(Appendices continue)

Appendix D

Analysis of Raw Reaction Times from Experiment 1 With Participants as the Random Variable Subjected to a 2×2 Mixed-Design ANOVA With Neighborhood Density as a Within-Subjects Factor and Accent as a Between-Subjects Factor

	Mean Reaction Time/ms (<i>SD</i>)	
	Neighborhood Density	
	Dense	Sparse
Accent		
Native	982 (94.3)	964 (80.7)
Accented	1165 (119)	1105 (121)
ANOVA's		
	<i>F</i> ₁ (1, 46)	<i>p</i> value
Effects		
Neighborhood Density	29.1	<.001
Accent	31.4	<.001
Neighborhood Density $\times$ Accent	8.48	0.006
Bonferroni Test		
	<i>F</i> ₁ (1, 46)	<i>p</i> value
Post hoc Effects for		
Neighborhood Density $\times$ Accent		
Dense versus Sparse in Native condition	3.07	0.08
Dense versus Sparse in Accented condition	34.5	<.001
Native versus Accented in Sparse condition	24.3	<.001
Native versus Accented in Dense condition	35.0	<.001

Received January 24, 2014

Revision received September 5, 2014

Accepted September 15, 2014 ■