

Journal of Experimental Psychology: Learning, Memory, and Cognition

Spoken Word Recognition and Serial Recall of Words From Components in the Phonological Network

Cynthia S. Q. Siew and Michael S. Vitevitch

Online First Publication, August 24, 2015. <http://dx.doi.org/10.1037/xlm0000139>

CITATION

Siew, C. S. Q., & Vitevitch, M. S. (2015, August 24). Spoken Word Recognition and Serial Recall of Words From Components in the Phonological Network. *Journal of Experimental Psychology: Learning, Memory, and Cognition*. Advance online publication. <http://dx.doi.org/10.1037/xlm0000139>

Spoken Word Recognition and Serial Recall of Words From Components in the Phonological Network

Cynthia S. Q. Siew and Michael S. Vitevitch
University of Kansas

Network science uses mathematical techniques to study complex systems such as the phonological lexicon (Vitevitch, 2008). The phonological network consists of a *giant component* (the largest connected component of the network) and *lexical islands* (smaller groups of words that are connected to each other, but not to the giant component). To determine if the component that a word resided in influenced lexical processing, language-related tasks (naming, lexical decision, and serial recall) were used to compare the processing of words from the giant component and from lexical islands. Results showed that words from lexical islands were recognized more quickly and recalled more accurately than words from the giant component. These findings can be accounted for via the diffusion of activation across a network. Implications for models of spoken word recognition and network science are also discussed.

Keywords: network science, word recognition, STM, network components, lexical islands

Network science is an emerging interdisciplinary field, which uses mathematical techniques to analyze a diverse array of complex systems in the biological, telecommunication, cognitive, and social domains (Barabási, 2009; Watts, 2004). In these complex networks, nodes represent entities such as people in a social group, Internet web pages, or words of a language, and connections typically represent relationships between any pair of entities; for instance, friendships among individuals, hypertext links between web pages, or phonological or semantic similarity between words. In recent years network science has been applied to the study of complex cognitive systems, in particular, the semantic and phonological relationships among words in the mental lexicon (Steyvers & Tenenbaum, 2005; Vitevitch, 2008).

In the language network examined in Vitevitch (2008), nodes represented phonological word forms and connections represented phonological similarity between words. Two words were considered phonologically similar if one word could be transformed to the other word via the substitution, addition, or deletion of one phoneme in any position (Landauer & Streeter, 1973; Luce & Pisoni, 1998). Vitevitch (2008) analyzed the phonological network using the mathematical and computational tools of network science and found that the network possessed a *giant component* (the largest connected component of the network), several *lexical islands* (smaller connected components of the network), and hermits (nodes that do not connect to any

nodes; known as “isolates” in the network science literature). The giant component of the phonological network exhibited characteristics of a small-world network: short average path lengths and high clustering coefficients relative to a network of comparable size, but with randomly placed connections.

Vitevitch (2008) also observed that the giant component of the phonological network consisted of 6,508 out of 19,340 words, about 33.7% of the entire network. In a comparative analysis of phonological networks of other languages including Spanish, Mandarin, Hawaiian, and Basque, the proportion of words residing in the giant components of these phonological networks ranged from 34% to 66% (Arbesman, Strogatz, & Vitevitch, 2010b). The proportion of nodes found in the giant components of these phonological networks is small relative to other real-world networks, where typically almost all nodes are connected to form a single connected component (Newman, 2001), making this an interesting aspect of the phonological network to examine further.

Previous studies of other aspects of the phonological network have demonstrated that the local structure of words (i.e., a word and the words that immediately surround that word) influences various aspects of spoken word recognition and production, word-learning, as well as short- and long-term memory processes (Chan & Vitevitch, 2009, 2010; Goldstein & Vitevitch, 2014; Vitevitch, Chan, & Roodenrys, 2012). Chan and Vitevitch (2009, 2010) showed that the clustering coefficient, or *C*, of a word has measurable effects on psycholinguistic tasks, such as perceptual identification, lexical decision, and picture naming. Clustering coefficient refers to the extent to which phonological neighbors of a word are also neighbors of each other (Watts & Strogatz, 1998). The phonological neighbors of high *C* words tend to be neighbors of each other, whereas the phonological neighbors of low *C* words do not tend to be neighbors of each other. Chan and Vitevitch (2009, 2010) found in various tasks that low *C* words were responded to more accurately and quickly than high *C* words.

Cynthia S. Q. Siew and Michael S. Vitevitch, Department of Psychology, University of Kansas.

The experiments reported here met in part the requirements for a Master's Thesis from the University of Kansas for the first author. We thank the members of the thesis committee: Susan Kemper and Joan Sereno.

Correspondence concerning this article should be addressed to Cynthia S. Q. Siew, Spoken Language Laboratory, Department of Psychology, 1415 Jayhawk Boulevard, University of Kansas, Lawrence, KS 66045. E-mail: cysnewsq@gmail.com

It has also been shown that the structure of the phonological network at the macro-level, which describes the overall structure of the network, influences various aspects of language processing (Vitevitch, Chan, & Goldstein, 2014). In particular, assortative mixing by degree, which refers to the tendency for highly connected nodes to be connected to other highly connected nodes in the network (Newman, 2002), represents one way to examine the macro-level structure of the phonological network. Vitevitch et al. (2014) investigated instances of failed lexical retrieval and found that errors made by participants reflected the presence of high assortative mixing by degree in the phonological network. This result suggests that phonological representations might indeed be organized as a complex network in the mental lexicon, not just at a micro-level as demonstrated by the clustering coefficient studies conducted by Chan and Vitevitch (2009, 2010), but also at a higher, macro-level as exemplified by assortative mixing by degree. The micro- and macro-level organization found in the mental lexicon could have important implications for understanding cognitive and lexical processes. For a computational analysis of the meso-level (or the structure of the network at a scale between the micro- and macro-level), see Siew (2013).

In the present study we used the network science approach to investigate another macro-level feature of the network: the various connected components that make up the phonological network. Specifically we asked whether the component a word is in (either the giant component or one of the lexical islands) influences processing in any way. If any processing differences exist between words that are found in either the giant component or lexical islands, this would further suggest that the phonological lexicon may be organized as an interconnected network, and further demonstrate that the overall macro-level structure of the lexicon has implications for lexical processing.

To date, little research has explicitly focused on examining words residing in lexical islands. One study by Arbesman, Strogatz, and Vitevitch (2010a) compared the words found in the lexical islands of English and Spanish, and found that Spanish words belonging to the same lexical island tended to be both phonologically and semantically related, whereas English words in the same island tended to be only phonologically related. This finding was offered as a possible explanation for the seemingly contradictory result regarding phonological neighborhoods in Spanish observed by Vitevitch and Rodríguez (2005). Vitevitch and Rodríguez found that Spanish words with many phonological neighbors were recognized more quickly than Spanish words with few phonological neighbors—the opposite of what is typically observed in English (Luce & Pisoni, 1998). In addition to being one of the few studies examining network components, the Arbesman et al. (2010a) study is also an example of how network analysis could provide additional insights into the lexical processes underlying spoken word recognition.

Other studies investigating the influence of various network science metrics, such as the clustering coefficient on language processing (e.g., Chan & Vitevitch, 2009, 2010; Goldstein & Vitevitch, 2014; Vitevitch et al., 2012), used words from the giant component as stimuli. Recall, however, that the giant component contained only 33.7% of the words from the entire mental lexicon. The majority of words in the lexicon are either lexical hermits, that is, words that do not have any phonological neighbors, or reside in smaller components, small groups of words that are connected to

each other, but disconnected from the giant component (referred to as *lexical islands*). In complex networks from other domains, an overwhelming majority of nodes are located within the giant component (Newman, 2001). Therefore, it is common practice to view smaller components and isolates as outliers, and exclude them from further analysis (e.g., Newman, 2001; Steyvers & Tenenbaum, 2005). In the phonological network, however, the majority of words are found outside of the giant component in the lexical islands and as hermits. It seems counterproductive to treat this large number of words as outliers and exclude them from further investigation. Rather, we argue that it is imperative that we extend our investigation to such words.

The present work is, to the best of our knowledge, the first experimental investigation of potential processing differences between words residing in the giant component and words residing in lexical islands (see also Vitevitch & Castro, 2015). For ease of exposition, the term “giant component words” will be used to refer to words residing in the giant component of the phonological network, and the term “lexical island words” will be used to refer to words residing in one of the lexical islands, which are disconnected from the giant component. In the present studies the giant component words were matched to the lexical island words on various psycholinguistic and network science characteristics that are known to influence processing. To examine how the location of a word in the lexical network might influence processing, we examined how participants responded to giant component words and to lexical island words in a variety of tasks commonly used in cognitive psychology, including word naming (Experiment 1), lexical decision (Experiment 2), and serial recall (Experiment 3).

Experiment 1

In the present experiment, an auditory naming task was used to examine how the location of a word in the lexical network might affect spoken word recognition. In the auditory naming task, a word is presented to participants over a set of headphones, and they must simply repeat the word as quickly and accurately as possible. We first consider a number of widely held assumptions about word recognition to make an initial prediction regarding how the location of a word in the lexical network might affect spoken word recognition.

In a network analysis of the phonological lexicon, Siew (2013) observed that giant component words tend to be short, monosyllabic words, whereas lexical island words tend to be long, multisyllabic words. Given the well-known relationship between word frequency and word length—commonly used words tend to be short words and less commonly used words tend to be longer words (Zipf, 1935)—we reasoned that the process of lexical retrieval might occur more often in the giant component than in the lexical islands.

We further reasoned that the common occurrence of lexical retrieval in the giant component may grant some sort of processing advantage to the giant component words over the lexical island words. For instance, activation of nearby neighbors of giant component words and lexical island words could spread to and partially activate giant component words and lexical island words, indirectly strengthening these partially activated items (Vitevitch & Goldstein, 2014; see also Nelson, McKinney, Gee, & Janczura, 1998). Given that lexical retrieval occurs more often within the

giant component than within lexical islands, it is reasonable to speculate that more frequent activation of nearby neighbors of giant component words might result in greater accumulation of “indirect” activation over time for these words compared with lexical island words. The strengthening of giant component words via accumulated indirect activation would therefore afford the giant component words a processing advantage over lexical island words.

From the earliest models (e.g., Morton, 1969) to more contemporary models of word recognition (e.g., Strauss, Harris, & Magnuson, 2007), the processing advantage for commonly occurring words over words that occur less often—the well-known word frequency effect in word recognition—is often represented via differences in activation thresholds for words that vary in frequency of occurrence. We adopt this commonly used approach in the present case and suggest that the strengthening of giant component words that occurs due to accumulated indirect activation may result in the giant component words having a lower activation threshold than lexical island words (even when word frequency is the same for both types of words, as is the case for the words used in the present experiments). The difference in activation threshold as a function of accumulated indirect activation could result in processing differences. Specifically, even though giant component words and lexical island words in the present experiments were matched on a number of relevant psycholinguistic characteristics, we expected that the more common experience of retrieving words from the giant component would result in a lower activation threshold for such words and therefore more rapid and accurate naming for giant component words compared with lexical island words.

Method

Participants. Twenty native English speakers were recruited from the introductory psychology subject pool at the University of Kansas. All participants had no previous history of speech or hearing disorders and received partial course credit for their participation.

Materials. Ninety-six English words were selected as stimuli for this experiment. Half of the stimuli were selected from the giant component of the phonological network and half were selected from various lexical islands in the phonological network. Of the 19,340 words in the phonological network, 6,508 words (33.65%) resided in the giant component, and 2,567 words (13.27%) were found in lexical islands of sizes ranging from 2 to 53 words. The lexical island words were selected such that they came from different islands of varying sizes (2 to 53). A male native speaker of American English (Michael S. Vitevitch) produced the stimuli by speaking at a normal speaking rate and volume into a high-quality microphone in an Industrial Acoustics Company (Bronx, NY) sound-attenuated booth. Individual sound files for each word were created from the digital recording and edited with SoundEdit 16 (Macromedia, Inc.; San Francisco, CA). All sound files were ensured to have the same amplitude by using the Normalization function in SoundEdit 16.

Table 1 shows the means and standard deviations of several psycholinguistic characteristics of giant component and lexical island words. A list of the stimuli is included in Appendix A.

Table 1

Characteristics of Giant Component and Lexical Island Words Used in Experiments 1–3

Lexical characteristic	Giant component	Lexical islands
	<i>M</i> (<i>SD</i>)	<i>M</i> (<i>SD</i>)
Number of phonemes	5.35 (0.53)	5.40 (0.64)
Subjective familiarity	6.60 (0.78)	6.77 (0.43)
Log frequency	0.93 (0.71)	1.11 (0.77)
Neighborhood density	2.73 (0.84)	2.83 (0.72)
Log neighborhood frequency	1.70 (0.53)	1.65 (0.47)
Mean positional probability	0.0533 (0.00852)	0.0542 (0.00760)
Mean biphone probability	0.00562 (0.00180)	0.00590 (0.00194)
Clustering coefficient	0.274 (0.353)	0.236 (0.311)
Onset duration (ms)	58 (3)	58 (4)
Stimuli duration (ms)	556 (92)	583 (70)
Overall file duration (ms)	675 (93)	700 (72)

Word length. Word length refers to the number of phonemes in a given word. Giant component words had a mean word length of 5.35 ($SD = 0.53$) and lexical island words had a mean word length of 5.40 ($SD = 0.64$), $F(1, 94) < 1$, $p = .73$.

Subjective familiarity. Subjective familiarity was measured on a 7-point scale (Nusbaum, Pisoni, & Davis, 1984). The rating scale ranged from 1 (*You have never seen the word before*) to 4 (*You recognize the word, but do not know the meaning*) to 7 (*You recognize the word and are confident that you know the meaning of the word*). Giant component words had a mean familiarity value of 6.60 ($SD = 0.78$) and lexical island words had a mean familiarity value of 6.77 ($SD = 0.43$), $F(1, 94) = 1.79$, $p = .19$. Therefore, both sets of words were considered highly familiar.

Word frequency. Word frequency refers to how often a given word occurs in a language. Log-base 10 of the raw frequency counts from Kučera and Francis (1967) were used. Giant component words had a mean word frequency of 0.93 ($SD = 0.71$) and lexical island words had a mean word frequency of 1.11 ($SD = 0.77$), $F(1, 94) = 1.45$, $p = .23$. Log-base 10 of the frequency counts from the more current SUBTLEX_{US} corpus (Brysbaert & New, 2009) were also obtained from the English Lexicon Project (Balota et al., 2007). Based on these frequency counts, giant component words had a mean word frequency of 2.39 ($SD = 0.86$) and lexical island words had a mean word frequency of 2.48 ($SD = 0.88$), $t(92) < 1$, $p = .63$.

Neighborhood density. Neighborhood density refers to the number of words that are phonologically similar to a given word (Luce & Pisoni, 1998). A word was considered to be phonologically similar to a target word if a single phoneme could be substituted, added, or deleted at any position of the target word to form that word. In the context of the phonological network in Vitevitch (2008), phonological neighborhood density corresponds to the network science term *degree*. Giant component words had a mean neighborhood density of 2.73 neighbors ($SD = 0.84$) and lexical island words had a mean neighborhood density of 2.83 neighbors ($SD = 0.72$), $F(1, 94) < 1$, $p = .52$.

Neighborhood frequency. Neighborhood frequency is the mean word frequency of the phonological neighbors of a word. Word frequency counts were obtained from Kučera and Francis (1967) and converted to log-base 10 values. Giant component words had a mean log neighborhood frequency of 1.70 ($SD =$

0.53) and lexical island words had a mean log neighborhood frequency of 1.65 ($SD = 0.47$), $F(1, 94) < 1$, $p = .60$.

Phonotactic probability. The phonotactic probability of a word refers to the probability that a segment occurs in a certain position of a word (positional segment probability), and the probability that two adjacent segments co-occur (biphone probability; Vitevitch & Luce, 2004). Giant component words had a mean positional segment probability of 0.0533 ($SD = 0.00852$) and lexical island words had a mean positional segment probability of 0.0542 ($SD = 0.00760$), $F(1, 94) < 1$, $p = .57$. Giant component words had a mean biphone probability of 0.00562 ($SD = 0.00180$) and lexical island words had a mean biphone probability of 0.00590 ($SD = 0.00194$), $F(1, 94) < 1$, $p = .48$.

Clustering coefficient. In the context of a phonological network, the clustering coefficient, C , refers to the extent to which the phonological neighbors of a word are also neighbors of each other. To calculate clustering coefficient, the number of connections between neighbors of a target word was counted and divided by the number of possible connections that could exist among the neighbors. Therefore, the clustering coefficient is the ratio of the actual number of connections existing among neighbors to the number of all possible connections among neighbors if every neighbor were connected (Batagelj & Mrvar, 1988). For a more precise definition of the clustering coefficient see Watts and Strogatz (1998), and for a definition of clustering coefficient in the context of a phonological network (and for how clustering coefficient differs from neighborhood density/degree) see Chan and Vitevitch (2009, 2010). C ranges from 0 to 1; when $C = 1$ all the neighbors of a word are neighbors of each other, and when $C = 0$ no neighbors of the word are neighbors of each other. Giant component words had a mean C of 0.274 ($SD = 0.353$) and lexical island words had a mean C of 0.236 ($SD = 0.311$), $F(1, 94) < 1$, $p = .58$.

Duration. The duration of the stimulus sound files was equivalent across both sets of words. The mean overall duration of sound files was 675 ms ($SD = 93$) for giant component words and 700 ms ($SD = 72$) for lexical island words, $F(1, 94) = 2.18$, $p = .14$. The mean onset duration, measured from the beginning of the sound file to the onset of the stimuli, was 58 ms ($SD = 3$) for giant component words and 58 ms ($SD = 4$) for lexical island words, $F(1, 94) < 1$, $p = .59$. The mean stimulus duration, measured from the onset to the offset of the word, was 556 ms ($SD = 92$) for giant component words and 583 ms ($SD = 70$) for lexical island words, $F(1, 94) = 2.64$, $p = .11$.

Uniqueness points. The uniqueness point of a word is the point, as measured from the beginning of the word, at which the word begins to diverge from all the other words in the lexicon (Marslen-Wilson, 1987). Studies have shown that uniqueness point plays a role in lexical decision and gating tasks (Marslen-Wilson, 1987; Tyler & Wessels, 1983). Giant component words had a mean uniqueness point of 4.13 phonemes ($SD = 0.94$) and lexical island words had a mean uniqueness point of 3.92 phonemes ($SD = 1.07$), $F(1, 94) = 1.02$, $p = .31$. The values of the uniqueness points here represent the position of the last overlapping phoneme for each word (see Luce, 1986).

Onset phoneme. To minimize acoustic and articulatory artifacts in the naming task (Treiman, Mullennix, Bijeljac-Babic, & Richmond-Welty, 1995), stimuli were selected such that the first phoneme was a consonant. Furthermore, chi-square analyses revealed that no particular phoneme was overrepresented in the

onsets of giant component words and lexical island words ($\chi^2 = 10.21$, $df = 14$, $p = .81$).

Procedure. Participants were tested individually. Each participant was seated in front of an iMac computer that was connected to a New Micros (Dallas, TX) response box. PsyScope 1.2.2 was used to randomize and present the stimuli via Beyerdynamic (Berlin, Germany) DT100 headphones at a comfortable listening level. A response box containing a dedicated timing board provided millisecond accuracy for the recording of response times.

In each trial, the word *READY* appeared on the screen for 500 ms. Participants heard one of the randomly selected stimuli and were instructed to repeat the word as quickly and accurately as possible. Reaction times were measured from stimulus onset to the onset of the participant's verbal response. Verbal responses were also recorded for offline scoring of accuracy. The next trial began 1 s after the participant's response was made. Prior to the experimental trials, each participant received five practice trials to become familiar with the task; these trials were not included in the subsequent analyses. There was a total of 96 trials and the experiment lasted approximately 10 min.

Results

Both reaction times (RTs) and accuracy rates were the dependent variables of interest. Accuracy was scored offline by Cynthia S. Q. Siew. Trials containing mispronunciations of the word or responses that triggered the voice-key prematurely (e.g., coughing, "uh") were coded as incorrect and excluded from the analyses. Trials with RTs that were less than 500 ms or more than 2,000 ms were considered to be outliers and also excluded. Excluded trials accounted for less than 2.92% of the data.

The convention in psycholinguistic research is to perform two types of analyses on participant and item means, treating participants and items as random factors in each of these analyses, respectively. However, there is a growing movement in the field to use alternative approaches such as multilevel modeling (Locker, Hoffman, & Bovaird, 2007) and hierarchical regression (e.g., Yap & Balota, 2009). In the following analyses we used hierarchical regression to assess the extent to which a variable (whether a word was found in the giant component or lexical islands) accounts for additional variability in the item means, over and above the variability that is already accounted for by other psycholinguistic and network characteristics (e.g., word frequency, neighborhood density). Therefore, item-level regression analyses were conducted on the mean RTs and accuracies for the stimuli.

A two-step hierarchical approach was used. Number of phonemes, familiarity, frequency, neighborhood density, neighborhood frequency, positional and biphone probabilities, C , and stimulus duration were entered in Step 1. Location, a dummy coded variable indicating whether a word resided in the giant component (coded as "0") or in a lexical island (coded as "1"), was entered in Step 2. Again, partitioning the regression analysis into two steps was done to determine if location of the word within the network accounted for additional variance over previously entered variables.

Reaction times. Table 2 presents the results of the regression analysis on naming RTs. In Step 1, frequency, positional probability, and stimulus duration significantly predicted naming RTs. Frequency was negatively correlated with RTs, standardized $\beta = -0.22$,

Table 2
Hierarchical Regression Results for Reaction Times From Experiment 1

Variable	β	SE	<i>t</i>	<i>p</i>	R^2	ΔR^2
Step 1						
Number of phonemes	.019	10.45	0.20	.84		
Subjective familiarity	-.061	9.14	-0.68	.50		
Log frequency	-.220	7.74	-2.44	.02*		
Neighborhood density	-.032	6.65	-0.39	.69		
Log neighborhood frequency	.148	10.65	1.76	.08†		
Positional probability	.344	922.30	2.96	.004**		
Biphone probability	-.208	4336.00	-1.65	.10		
Stimuli duration	.633	0.07	7.00	<.001***		
Clustering coefficient	-.061	15.56	-0.75	.46		
					.482***	
Step 2						
Location (dummy variable)	-.177	10.00	-2.24	.03*		
					.511***	.029*

† $p < .10$. * $p < .05$. ** $p < .01$. *** $p < .001$.

$t(86) = -2.44, p < .05$, replicating the well-known effect that high frequency words tend to be responded to more quickly than low frequency words. Positional probability was positively correlated with RTs, standardized $\beta = 0.344, t(86) = 2.96, p < .01$, such that words with high phonotactic probability were responded to less quickly than words with low phonotactic probability (replicating results reported in Vitevitch & Luce, 1998). It is not surprising that stimulus duration was positively correlated with RTs, standardized $\beta = 0.633, t(86) = 7.00, p < .001$, such that words of longer durations were responded to less quickly than words of shorter duration. Together, the variables entered at Step 1 explained 48.2% of the variance in naming RTs, accounting for a significant proportion of the variance in naming RTs, $R^2 = .482, F(9, 86) = 8.90, p < .001$.

In Step 2, location significantly predicted naming RTs, standardized $\beta = -0.177, t(85) = -2.24, p < .05$, such that lexical island words were responded to more quickly than giant component words. The influence of location accounted for an additional 2.9% of the variance, $\Delta R^2 = .029, F(1, 85) = 5.03, p < .05$. Together, the variables entered at both steps explained 51.1% of the variance in naming RTs, accounting for a significant proportion of variance in naming RTs, $R^2 = .511, F(10, 85) = 8.89, p < .001$.

Accuracy. Table 3 presents the results of the regression analysis on naming accuracies. In Step 1, only familiarity significantly predicted naming accuracies, standardized $\beta = 0.328, t(86) = 2.84, p < .01$, such that more familiar words were responded to more accurately than less familiar words. Together, the variables entered at Step 1 explained 14.8% of the variance in naming accuracies, which did not account for a significant proportion of variance in naming accuracies, $R^2 = .148, F(9, 86) = 1.66, p = .11$.

In Step 2, location did not significantly predict naming accuracy, standardized $\beta = 0.112, t(85) = 1.08, p = .28$, nor did it explain a significant proportion of variance, $\Delta R^2 = .012, F(1, 85) = 1.17, p = .28$. Together, the variables entered at both steps explained 16.0% of the variance in naming accuracies, which did not account for a significant proportion of variance in naming accuracies, $R^2 = .160, F(10, 85) = 1.62, p = .12$.

Table 4 shows the subject and item RT and accuracy means for the lexical island and giant component words. RTs for lexical island words ($M = 956$ ms, $SD = 55$) were faster than RTs for giant component words ($M = 968$ ms, $SD = 72$), and this was consistent across subject means as well.

Table 3
Hierarchical Regression Results for Accuracy Rates From Experiment 1

Variable	β	SE	<i>t</i>	<i>p</i>	R^2	ΔR^2
Step 1						
Number of phonemes	.080	0.89	0.66	.51		
Subjective familiarity	.328	0.78	2.84	.006**		
Log frequency	.053	0.66	0.46	.65		
Neighborhood density	-.021	0.57	-0.20	.84		
Log neighborhood frequency	.054	0.91	0.50	.62		
Positional probability	.084	79.13	0.56	.58		
Biphone probability	-.128	372.00	-0.79	.43		
Stimuli duration	.017	0.006	-0.15	.88		
Clustering coefficient	.081	1.34	0.78	.44		
					.148	
Step 2						
Location (dummy variable)	.112	0.88	1.08	.28		
					.16	.012

** $p < .01$.

Table 4
Subject and Item Means for Giant Component and Lexical
Island Words in Experiment 1

	Lexical island words	Giant component words
	<i>M</i> (<i>SD</i>)	<i>M</i> (<i>SD</i>)
Subject means		
Reaction times (ms)	955 (107)	967 (110)
Accuracy (%)	97.71 (3.09)	96.46 (4.06)
Item means		
Reaction times (ms)	956 (55)	968 (72)
Accuracy (%)	97.71 (3.41)	96.46 (4.94)

Accuracy rates were very high for both lexical island and giant component words, although slightly higher accuracy rates were observed for the lexical island words ($M = 97.71\%$, $SD = 3.41$) compared with giant component words ($M = 94.46\%$, $SD = 4.94$). This was consistent across subject means as well. The fact that accuracy rates are close to ceiling could explain why the location of word within the network did not significantly affect accuracy rates. This also suggests that there was no speed–accuracy trade-off in the performance of this task.

Discussion

Recall (a) the relationship between word length and word frequency (Zipf, 1935), and (b) the tendency for giant component words to be short and lexical island words to be long. Given these relationships, we reasoned that the process of lexical retrieval occurs more often in the giant component than in the lexical islands, bestowing a practice-effect-like processing advantage to the giant component words over the lexical island words. Contrary to our intuition, the results of the word naming task showed that lexical island words were processed more quickly than giant component words. Because every behavioral task used in laboratory settings has advantages and disadvantages, we sought to replicate this result in another task in Experiment 2 to increase our confidence that the observed effect was not spurious, or due to the characteristics of a particular task.

Experiment 2

Given the somewhat counterintuitive finding in Experiment 1 we sought to replicate the effect in a different psycholinguistic task, namely the lexical decision task. In this task, participants are presented with either a word or a nonword over a set of headphones. In this standard psycholinguistic task, participants are asked to decide as quickly and accurately as possible whether the given stimulus is a real word in English or a nonsense word. As in the auditory naming task used in Experiment 1, the lexical decision task allows for both accuracy and RT to be assessed.

Method

Participants. Twenty native English speakers were recruited from the introductory psychology subject pool as described in Experiment 1. All participants were right-handed and had no

previous history of speech or hearing disorders. None of the participants in the present experiment took part in Experiment 1.

Materials. The word stimuli for the present experiment consisted of the same 96 words used in Experiment 1. In addition, a list of 96 phonotactically legal nonwords was constructed by replacing a phoneme (at any position except the first and last positions) of the word stimuli with another phoneme. For instance, the nonword *porcel* (/pɔrɪsl/) was created by replacing /a/ in *parcel* (/pɑrɪsl/) with /o/. The phonological transcriptions of nonwords are listed in Appendix B.

The nonwords were recorded by the same male speaker in a similar manner as in Experiment 1. The same method for editing and digitizing the word stimuli was used to create individual sound files for each nonword.

Duration. The duration of the stimulus sound files was equivalent across both words and nonwords. The mean overall duration of sound files was 687 ms ($SD = 84$) for words and 665 ms ($SD = 75$) for nonwords, $F(1, 190) = 3.70$, $p = .06$. The mean onset duration, measured from the beginning of the sound file to the onset of the stimuli, was 58 ms ($SD = 3$) for words and 57 ms ($SD = 5$) for nonwords, $F(1, 190) < 1$, $p = .40$. The mean stimulus duration, measured from the onset to the offset of the word, was 569 ms ($SD = 82$) for words and 550 ms ($SD = 75$) for nonwords, $F(1, 190) = 2.84$, $p = .09$.

Procedure. Participants were tested in groups no larger than three. As in Experiment 1, each participant was seated in front of an iMac computer that was connected to a New Micros response box. PsyScope 1.2.2 was used to randomize and present the stimuli via BeyerDynamic DT100 headphones at a comfortable listening level. The response box contained a dedicated timing board, providing millisecond accuracy for the recording of response times.

In each trial, the word *READY* appeared on the screen for 500 ms. Participants heard one of the randomly selected stimuli and were instructed to decide, as quickly and accurately as possible, whether the item heard was a real English word or a nonword. If the item was a word, participants pressed the button labeled *WORD* with their right (dominant) index finger. If the item was a nonword, participants pressed the button labeled *NONWORD* with their left index finger. Reaction times were measured from stimulus onset to the onset of the participant's button press. The next trial began 1 s after the participant's response was made. Prior to the experimental trials, each participant received eight practice trials to become familiar with the task; these trials were not included in the subsequent analysis. There was a total of 192 trials and the experiment lasted approximately 15 min.

Results

Both RTs and accuracy rates were the dependent variables of interest. In lexical decision, only accurate responses for word stimuli were analyzed. Trials with RTs that were less than 500 ms or more than 2,000 ms were excluded. Excluded trials accounted for less than 9.53% of the data. As in Experiment 1, hierarchical regression analyses were conducted on the item means.

Reaction times. Table 5 presents the results of the regression analysis on the lexical decision RTs. In Step 1, familiarity, positional probability, and stimulus duration significantly predicted lexical decision RTs. Familiarity was negatively correlated with RTs, standardized $\beta = -0.475$, $t(86) = -5.12$, $p < .001$, such that more familiar

Table 5
Hierarchical Regression Results for Reaction Times From Experiment 2

Variable	β	SE	<i>t</i>	<i>p</i>	R^2	ΔR^2
Step 1						
Number of phonemes	.039	17.05	0.40	.69		
Subjective familiarity	-.475	14.91	-5.12	<.001***		
Log frequency	-.158	12.62	-1.70	.09†		
Neighborhood density	-.089	10.90	-1.06	.29		
Log neighborhood frequency	.027	17.37	0.31	.76		
Positional probability	.284	1,505.00	2.37	.02*		
Biphone probability	-.070	7,074.00	-0.53	.60		
Stimuli duration	.341	0.11	3.66	<.001***		
Clustering coefficient	-.113	25.38	-1.36	.18		
					.452***	
Step 2						
Location	-.203	16.20	-2.52	.01*		
					.490***	.038*

† $p < .10$. * $p < .05$. *** $p < .001$.

words were responded to more quickly than less familiar words. Positional probability was positively correlated with RTs, standardized $\beta = 0.284$, $t(86) = 2.37$, $p < .05$, such that words with high phonotactic probability were responded to less quickly than words with low phonotactic probability. Stimulus duration was positively correlated with RTs, standardized $\beta = 0.341$, $t(86) = 3.66$, $p < .001$, such that words of longer durations were responded to less quickly than words of shorter durations. Together, the variables entered at Step 1 explained 45.2% of the variance in lexical decision RTs, accounting for a significant proportion of the variance in lexical decision RTs, $R^2 = .452$, $F(9, 86) = 7.86$, $p < .001$.

In Step 2, location significantly predicted lexical decision RTs, standardized $\beta = -0.203$, $t(85) = -2.52$, $p = .01$, such that lexical island words were responded to more quickly than giant component words, and accounted for an additional 3.8% of the variance, $\Delta R^2 = .038$, $F(1, 85) = 6.35$, $p < .05$. Together, the variables entered at both steps explained 49.0% of the variance in lexical decision RTs, accounting for a significant proportion of variance in lexical decision RTs, $R^2 = .490$, $F(10, 85) = 8.15$, $p < .001$.

Accuracy. Table 6 presents the results of the regression analysis on the lexical decision accuracy rates. In Step 1, familiarity

and frequency significantly predicted lexical decision accuracy rates. Familiarity was positively correlated with accuracy, standardized $\beta = 0.616$, $t(86) = 7.83$, $p < .001$, such that more familiar words were responded to more accurately than less familiar words. Frequency was also positively correlated with accuracy rates, standardized $\beta = 0.192$, $t(86) = 2.44$, $p < .05$, such that high frequency words were responded to more accurately than low frequency words. Together, the variables entered at Step 1 explained 60.6% of the variance in lexical decision accuracy rates, accounting for a significant proportion of variance in lexical decision accuracies, $R^2 = .606$, $F(9, 86) = 14.69$, $p < .001$.

In Step 2, location did not significantly predict lexical decision accuracy, standardized $\beta = 0.058$, $t(85) = 0.82$, $p = .41$, nor did it explain a significant proportion of variance, $\Delta R^2 = .003$, $F(1, 85) = 0.62$, $p = .44$. Together, the variables entered at both steps explained 60.6% of the variance in lexical decision accuracy rates, accounting for a significant proportion of variance in lexical decision accuracies, $R^2 = .606$, $F(10, 85) = 13.24$, $p < .001$.

Table 7 shows the subject and item RT and accuracy means for the lexical island and giant component words. Reaction times for lexical island words ($M = 978$ ms, $SD = 83$) were

Table 6
Hierarchical Regression Results for Accuracy Rates From Experiment 2

Variable	β	SE	<i>t</i>	<i>p</i>	R^2	ΔR^2
Step 1						
Number of phonemes	-.013	2.03	-0.15	.88		
Subjective familiarity	.616	1.77	7.83	<.001***		
Log frequency	.192	1.50	2.44	.02*		
Neighborhood density	-.061	1.29	-0.86	.39		
Log neighborhood frequency	-.050	2.07	-0.68	.50		
Positional probability	.027	178.90	0.27	.79		
Biphone probability	-.032	841.30	-0.29	.77		
Stimuli duration	.157	0.01	1.99	.05†		
Clustering coefficient	.073	3.02	1.04	.30		
					.606***	
Step 2						
Location (dummy variable)	.058	1.99	0.82	.41		
					.609***	.003

† $p < .10$. * $p < .05$. *** $p < .001$.

Table 7
Subject and Item Means for Giant Component and Lexical Island Words in Experiment 2

	Lexical island words	Giant component words
	<i>M</i> (<i>SD</i>)	<i>M</i> (<i>SD</i>)
Subject means		
Reaction times (ms)	974 (71)	1,006 (92)
Accuracy (%)	93.02 (5.33)	87.92 (6.43)
Item means		
Reaction times (ms)	978 (83)	1,019 (114)
Accuracy (%)	93.02 (8.92)	87.92 (17.74)

faster than RTs for giant component words ($M = 1,019$ ms, $SD = 114$), and this was consistent across subject means as well.

Accuracy rates were very high across both lexical island and giant component conditions, although higher accuracy rates were observed for the lexical island words ($M = 93.02\%$, $SD = 8.92$) compared with giant component words ($M = 87.92\%$, $SD = 17.74$). This was consistent across subject means as well.

Discussion

The results of the lexical decision task replicated the somewhat counterintuitive results obtained in the word naming task in Experiment 1—lexical island words were responded to more quickly than giant component words. It is important to emphasize that the giant component and lexical island words were closely matched on a number of network and psycholinguistic variables that are known to influence lexical processing. Given the matching of relevant psycholinguistic variables, widely accepted models of spoken word recognition would not predict any difference in the responses to these two sets of words (e.g., Marslen-Wilson, 1987; McClelland & Elman, 1986; Norris & McQueen, 2008; Luce & Pisoni, 1998). Nevertheless, the results of Experiments 1 and 2 showed that lexical island words are processed more quickly than giant component words in both naming and lexical decision tasks. This result suggests that there may be some psychological validity to the idea that phonological word-forms in the mental lexicon are organized as a complex network, and, more important, that where a word is located in the network has an important influence on lexical processing.

Although (a) the relationship between word length and word frequency (Zipf, 1935), and (b) the tendency for giant component words to be short and lexical island words to be long led us to believe that the process of lexical retrieval might occur more often in the giant component than in the lexical islands, and bestow some sort of processing advantage to the giant component words over the lexical island words, the results of Experiments 1 and 2 are inconsistent with that intuition. Instead, lexical island words appear to have a processing advantage over comparable giant component words.

The assumptions of mainstream psycholinguistics led us to make a reasonable but incorrect prediction regarding how the location of a word in the lexical network might affect spoken word recognition. To account for the processing advantage of lexical island words over comparable giant component words we instead

appeal to the network framework described in Chan & Vitevitch (2009, 2010) and simulated with a simple diffusion mechanism in Vitevitch, Ercal, and Adagarla (2011).

In Vitevitch et al. (2011) *activation* was defined as a limited cognitive resource that spread unimpeded between connected nodes, and did not decay over time (compare the models of Anderson (1983); Collins and Loftus (1975); MacKay (1990), and others for very different assumptions about activation, how it spreads, decays, and so forth). In the very simple model used in Vitevitch et al. (2011), the target node received an initial burst of activation of 100 arbitrary units. A portion of the initial activation was retained by the target word, and the remaining amount of activation that was not retained in the target node was equally divided (i.e., spread) among the neighbors of the target word (referred to as *one-hop neighbors*).

Similar to the target node, each one-hop neighbor retained a portion of the activation it received from the target node, with the remaining activation being spread equally to the nodes to which it was connected (including the target, other one-hop neighbors, and nodes connected to the one-hop neighbors but not connected to the target word; such nodes are referred to as *two-hop neighbors* of the target word). This process of a node retaining a certain amount of activation and spreading the rest of the activation to nodes that it was connected to—resulting in activation being spread back and forth between the target, one-hop neighbors, and two-hop neighbors—occurred for 10 discrete time steps, at which point retrieval of the target items was said to occur.

The activation value in the target node was mapped inversely to response latency, such that higher activation values indicated that lexical retrieval occurred rapidly, and lower activation values indicate that lexical retrieval required more time to be completed, and directly to accuracy, such that higher activation values indicated a high probability of accurate retrieval, and lower activation values indicated a low probability of accurate retrieval. With this very simple diffusion model Vitevitch et al. (2011) were able to account for the psycholinguistic observations reported in Chan and Vitevitch (2009): Low *C* words were more quickly and accurately recognized than high *C* words.

In the simulations reported in Vitevitch et al. (2011) the size of the 2-hop networks was the same. The only difference in the networks examined in Vitevitch et al. (2011) was the number of connections among the 1-hop neighbors (i.e., the clustering coefficient, *C*). In the present study, the networks that surround each word extend beyond 2-hops, and vary in size. Using the basic framework described in Chan and Vitevitch (2009) and simulated in Vitevitch et al. (2011), we consider how the difference in the size of the network component that the words reside in (lexical island vs. giant component) might contribute to the results observed in the present set of experiments.

Recall that the giant component represented the largest connected component of the phonological network, whereas lexical islands constitute smaller, connected components of the phonological network. The giant component consisted of 6,508 words compared with the largest lexical island, which consisted of only 53 words. If we again assume that the target words in the giant component and in each island receive an initial burst of activation of 100 arbitrary units, then the diffusion of activation to every word connected to the target word (either directly or indirectly) results in each word in the giant component receiving a smaller

“share” of the initial burst of activation ($100/6,508 = 0.015$ units of activation) compared with the “share” of the activation that each word in even the largest of the lexical islands receives ($100/53 = 1.887$ units of activation). Using the mapping of activation values to response latency and accuracy rate as described in Vitevitch et al. (2011), the relative difference in the activation received by words in the giant component versus the lexical islands gives the words in the islands a clear processing advantage (despite being comparable on other psycholinguistic characteristics).

Even though conventional assumptions from psycholinguistics and mainstream models of word recognition cannot account for the results of Experiments 1 and 2, the network framework described in Chan and Vitevitch (2009, 2010) and simulated with a simple diffusion mechanism in Vitevitch et al. (2011) does provide an account. To further examine how activation in the giant component versus activation in the lexical islands affects retrieval of otherwise comparable words we used a serial recall task (Experiment 3). Using a serial recall task offers another way to compare the processing of giant component and lexical island words, especially as it has been argued that short-term memory (STM) involves processes that also occur in speech perception (Ellis, 1980; Schweickert, 1993) and speech production (Roodenrys, Hulme, Lethbridge, Hinton, & Nimmo, 2002). In addition, given that the dependent variable of interest in the serial recall task is recall accuracy, the serial recall task may reveal differences in accuracy rates between giant component and lexical island words that were not observed in Experiments 1 and 2. Based on the results of Experiments 1 and 2, we expect that lexical island words will be recalled more accurately than giant component words.

Experiment 3

In a serial recall task participants are presented with a sequence of items (e.g., words or numbers) and have to recall those items in the same order in which they were presented. The serial recall task is widely used by cognitive psychologists to examine STM and its underlying cognitive processes (Baddeley, Thomson, & Buchanan, 1975; Ebbinghaus, 1913; Hulme, Maughan, & Brown, 1991). However, there is evidence that suggests that long-term memory contributes to serial recall ability as well (Hulme et al., 1997; Tehan & Humphreys, 1988; Watkins, 1977), making it appropriate to use this task to further examine (in a complementary manner) that part of long-term memory known as the mental lexicon (see Vitevitch et al., 2012).

Method

Participants. Thirty-two native English speakers were recruited from the introductory psychology subject pool. All participants had no previous history of speech or hearing disorders and received partial course credit for their participation. These participants did not participate in Experiments 1 and 2.

Materials. The word stimuli for the present experiment consisted of the same 96 words used in Experiment 1. The words in each condition were pseudorandomly assigned to ensure that no phonological neighbors appeared in the same list. Eight lists consisting of six giant component words each and eight lists consisting of six lexical island words each were created. In addition, two different samples of the 16 lists (Versions A and B) were created to minimize order effects.

Procedure. Participants were tested individually. Each participant was randomly assigned to one of two versions of the word lists (A or B), with 16 participants being assigned to each version. As in the previous experiments, each participant was seated in front of an iMac computer. PsyScope 1.2.2 was used to randomize and present the stimuli via BeyerDynamic DT100 headphones at a comfortable listening level.

In each trial, the word *READY* appeared on the screen for 500 ms. Participants were presented with one of the 16 randomly selected lists over headphones, at a rate of approximately 1 word per second. At the end of each list, the prompt *RECALL* appeared on the screen and participants recalled out loud the list of words in the same order as they were presented. Participants were instructed to say “pass” if they could not recall the word in any particular position. Verbal responses were recorded for offline scoring of accuracy. The next trial began when participants finished recalling the words and pressed the spacebar. Prior to the experimental trials, each participant received four practice trials to become familiar with the task; these trials were not included in the subsequent analyses. There were a total of 16 trials and the experiment lasted approximately 15 min.

Results

In contrast to the previous two experiments, recall accuracy was the dependent variable of interest in this experiment. Accuracy was manually scored offline by Cynthia S. Q. Siew. Trials which contained mispronunciations of the word, or in which the participant said “pass” (or some indication of recall failure, e.g., “skip” or “don’t know”) were coded as incorrect trials.

A 2×6 two-way within-participants analysis of variance (ANOVA) was conducted. The independent variables were location (2; lexical island or giant component) and serial position (6; 1 through 6). The dependent variable was the mean accuracy rate in each condition. The Location $\times$ Position interaction was significant, $F(5, 155) = 3.32, p < .01$. To ensure that the significant interaction was not due to specific ordering effects of either Version A or Version B, list was included as a third independent variable in the ANOVA. The three-way interaction was not significant, indicating that the significant two-way interaction observed between location and position was consistent across both lists.

To further interpret the nature of the significant Location $\times$ Position interaction, tests of simple main effects of location were conducted at each level of position. At Position 1, the simple main effect of location was significant, $F(1, 31) = 9.83, p < .01$. At Positions 2 to 6, the simple main effect of location was not significant ($F_s < 1.70, p_s > .20$).

As shown in Figure 1, recall for words from lexical islands was significantly better than words from the giant component, but only for words in the first position of the serial recall curve. Recall for words from lexical islands or the giant component did not significantly differ across the other positions along the serial recall curve.

Discussion

As we predicted, serial recall of lexical island words was more accurate than recall of giant component words, although this was

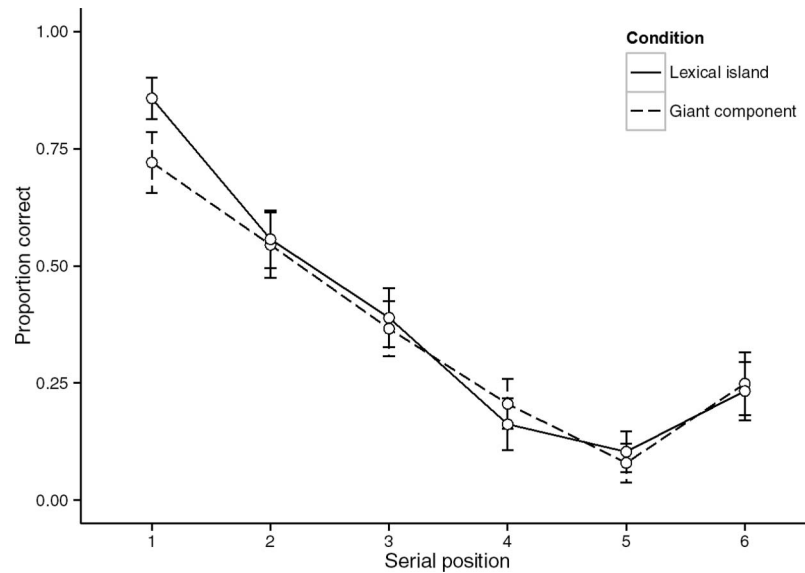


Figure 1. Serial recall curve showing proportion of accurate recall of giant component and lexical island words at each serial position. The error bars represent standard error.

only observed for words in the initial serial position. According to the memory literature, items presented in early serial positions are retrieved from long-term memory, whereas items presented in late serial positions are retrieved from STM (Craik, 1968; Watkins, 1977). Given that an advantage for lexical island words was only observed in the first serial position (i.e., a primacy effect), this would suggest that long-term memory plays a role in the superior serial recall performance of lexical island words in the first serial position. Given that the mental lexicon is part of long-term memory, it is perhaps not surprising that effects for where in the lexicon a word resides (i.e., islands vs. the giant component) were observed in the first position of the serial recall curve.

The advantage in serial recall of lexical island words over giant component words in the present experiment can again be accounted for with the diffusion of activation account discussed in the contexts of Experiments 1 and 2. According to this account, activation eventually diffuses to all of the other words within the component. For a word in the giant component, which contains more than 6,000 words, the “share” of activation that each word has is relatively less than the “share” of activation that each word has in even the largest lexical island (which contained about 50 words). The relative difference in the amount of activation for each lexical island word and for each giant component word gives lexical island words a processing advantage over giant component words, as observed in the present experiment.

However, the processing advantage for lexical island words over giant component words is not limited to the experiments reported here. We also examined the visual naming and visual lexical decision data from the English Lexicon Project (ELP; Balota et al., 2007) for evidence of such a processing advantage. The ELP database contains lexical characteristics (e.g., word frequency, number of orthographic and phonological neighbors) for more than 40,000 words, as well as behavioral data from 1,260 participants across six universities who responded to those words in a visual naming task and a visual lexical decision task. One important point

to note is that the behavioral data in the ELP was obtained by presenting stimuli visually rather than auditorily, as in the present experiments. It is nonetheless worthwhile to investigate whether the macro-structure of the phonological network might also influence visual word processing (especially in light of the results reported by Yates, 2013).

We analyzed in the ELP database the words used in the present studies (N.B., *cumber*, *remit*, *rollick*, and *scepter* were not found in the ELP database). Although none of the analyses reached statistical significance, the numerical trends observed in the visual naming and visual lexical decision data from the ELP were in the same direction as observed in the present studies.

We also attempted to use the ELP database to generalize to a different set of words (48 new words were selected from lexical islands such that compared with the original stimuli, equivalent numbers of lexical island words were obtained from lexical islands of the same sizes). Again, none of the analyses reached statistical significance, but the numerical trends observed in the visual naming and visual lexical decision data from the ELP for the new set of words were in the same direction as observed in the present studies. The trends observed in the ELP analyses further suggest that a processing advantage exists for lexical island words over giant component words.

More directly related to spoken language processing, Vitevitch and Castro (2015) analyzed data from individuals with Broca’s aphasia, with Wernicke’s aphasia, and from age-matched controls on the Philadelphia Naming Task (Mirman et al., 2010). Their analysis replicated a number of well-known findings, including well-studied influences on picture naming of word length, word frequency, and neighborhood density (Vitevitch, 1997, 2002). Relevant to the present set of experiments, Vitevitch and Castro found, using a binomial multiple regression model, that words found outside of the giant component were named about two times more accurately than words found in the giant component. Although Vitevitch and Castro did not distinguish between words in

smaller components (i.e., lexical island words) and isolates (i.e., hermits) in their analysis (i.e., they simply categorized words as being inside or outside of the giant component), their results clearly show the same processing disadvantage for words found in the giant component that was observed in the present set of experiments.

The analysis reported in Vitevitch and Castro (2015) is important because it provides a systematic replication of the processing disadvantage for words found in the giant component that was observed in the present set of experiments. Notably, Vitevitch and Castro found this effect in several different participant populations, in a different psycholinguistic task (picture-naming), in a different set of words (only the word *mountain* appeared in the present experiments and in the words in the Philadelphia Naming Test), and in a different dependent variable (i.e., naming accuracy) in data that were collected by an independent lab. Such replications demonstrate the internal and external validity of the present results.

General Discussion

Across three experiments (see also Vitevitch & Castro, 2015), a processing advantage was observed for lexical island words compared with giant component words. Lexical island words were produced more quickly in a naming task, recognized more quickly in a lexicon decision task, and recalled more accurately in a serial recall task than giant component words. As these two sets of words were matched on a number of relevant psycholinguistic and network characteristics known to influence lexical processing, the present findings show that where a word is located in the lexical network influences lexical retrieval.

Additional analyses of data from the English Lexicon Project revealed trends that were consistent with the results observed in the three psycholinguistic experiments reported above. These analyses were conducted for the stimuli used in the present experiments as well as for a different set of carefully matched stimuli. In both sets of words, lexical island words were responded to more quickly and accurately than giant component words. It is also worthwhile to point out that Vitevitch and Castro (2015) analyzed data from the Philadelphia Naming Test in the Moss Aphasia Psycholinguistics Project Database and found that individuals named pictures of words outside of the giant component (i.e., lexical island words and hermits) more accurately than pictures of giant component words. As these analyses explicitly controlled for various lexical characteristics known to influence the likelihood of correctly naming a picture, it is difficult for any explanation that does not consider the overall network structure of the mental lexicon to account for the findings of Vitevitch and Castro as well as the results from the present investigation. Taken together, it is striking to see that analyses of databases containing behavioral data from a variety of experimental designs and different participant populations provide converging evidence that lexical island words possess a lexical retrieval advantage over giant component words, strongly suggesting that our findings are not spurious results limited to a particular set of stimuli, nor were they a by-product of the experimental paradigms employed in the present study.

There are a number of widely accepted models of spoken word recognition including the cohort model (Marslen-Wilson, 1987; the latest adaptation being the distributed cohort model, Gaskell &

Marslen-Wilson, 1997); TRACE, the interactive-activation model proposed by McClelland and Elman (1986), Shortlist B (Norris & McQueen, 2008); neighborhood activation model (NAM; Luce & Pisoni, 1998); and PARSYN, the computational instantiation of NAM (Luce, Goldinger, Auer, & Vitevitch, 2000). Although these models have different computational architectures, they have successfully accounted for several well-established influences of word frequency, phonological neighborhood density, and phonotactic probability on spoken word recognition. It is not clear, however, how these well-known models of spoken word recognition would account for the present results. Recall that lexical island and giant component words selected for the present set of experiments were closely matched on a variety of lexical characteristics that are known to influence spoken word recognition, so current models of spoken word recognition would not predict any processing differences between the two sets of words. In contrast, we observed processing differences depending on whether words resided within the giant component or within lexical islands of the phonological network.

The results of the present experiments were also counter to our initial intuitions. We reasoned that (a) the inverse relationship between word frequency and word length, such that high frequency words tend to be short words, and low frequency words tend to be longer words (Zipf, 1935), and that (b) shorter, more frequently occurring words tend to be found in the giant component and longer, less frequently occurring words tend to be found among the lexical islands and hermits (Siew, 2013) would result in a great deal of lexical processing occurring within the giant component. We further reasoned that the accumulation of indirect activation (e.g., Nelson et al., 1998) in giant component words compared with lexical island words would result in more rapid and accurate responses to giant component words compared with lexical island words. Contrary to our intuition, we observed instead a processing advantage for lexical island words over giant component words.

To account for this counterintuitive finding, which cannot be accommodated by widely accepted models of spoken word recognition, we appealed to the network framework described in Chan and Vitevitch (2009) and simulated in Vitevitch et al. (2011). The diffusion of activation framework assumes that total activation remains constant over time (i.e., activation does not “decay” over time) and that higher activation levels are associated with faster retrieval times and more accurate retrieval (see Vitevitch et al., 2011 for more detailed descriptions of these assumptions). In this framework, a word-form (or node) was partially activated by the acoustic-phonetic input. Activation at that node would diffuse to other nodes connected to that target node, with activation continuing to diffuse to other connected nodes (including diffusing back to the target item). The more activation that accumulates in a node, the faster and more accurate a response would be for the word represented by that node.

In the present case, the giant component contained more than 6,000 words, and the largest lexical island contained about 50 words. If we again assume that the target words in each case receive an initial burst of activation of 100 arbitrary units, then the diffusion of activation to every word connected to the target word (either directly or indirectly) results in each word in the giant component receiving a smaller “share” of the initial burst of activation compared with the “share” of the activation that each word in even the largest of the lexical islands receives. The relative difference in the activation received by words in the giant component versus the lexical islands gives the words in the

islands a clear processing advantage (despite being comparable on other psycholinguistic characteristics).

The diffusion of activation across the network described in Chan and Vitevitch (2009) and simulated in Vitevitch et al. (2011) predicts gradient effects based on the size of the component. In practice, however, we expect the effects to appear more categorical in nature. We make what may appear to be contradictory predictions because of the distribution of component size in the lexical network. Consider that the giant component of English contains ~6,000 words (Vitevitch, 2008), and the next largest component (i.e., the largest lexical island) contains ~50 words; this is a difference of 2 orders of magnitude. The next largest component (i.e., the second largest lexical island) contains ~10 words. Arbesman et al. (2010a) examined a larger database of English words, but again found a difference in size between the giant component and lexical islands that was of several orders of magnitude, as shown in Figure 2 of that report. Given the small differences in size found among the lexical islands, and the very large difference in size when comparing the giant component to the lexical islands, any gradient effects that may exist may not be easily detected with conventional statistical methods, and may instead appear to be categorical in nature.

An alternative framework that may also account for the present findings comes from a memory retrieval model, which considers the discriminability of a particular item in the context of its neighbors within a psychological space (we thank an anonymous reviewer for suggesting this alternative approach). Brown, Neath, and Chater (2007) proposed the Scale-Independent Memory, Perception, and LEarning (SIMPLE) model to account for various memory phenomena, one of which—*isolation* (or *distinctiveness*) effects—may be of particular relevance to our present discussion. According to the SIMPLE model, successful memory retrieval occurs when an item is particularly distinctive and discriminable compared with its nearby neighbors on various dimensions. The temporal isolation effect—the finding whereby an item in the middle of a list that is preceded and followed by an exceptionally long pause during presentation is better recalled—is one example of an isolation effect. Generally, items that are more distinguishable on a temporal dimension are more accurately remembered because they are temporally distinctive—being “further” away from their competitors on a temporal dimension (Brown, Morin, & Lewandowsky, 2006).

By analogy, one could replace the temporal dimension in the SIMPLE model (Brown et al., 2006) with a spatial dimension representing phonological similarity-space. Previous studies demonstrate that the speed and accuracy of lexical retrieval depends on the ease of discriminating a target word from its neighbors in a phonological similarity-space. The neighborhood density effect is one example—words with several phonological neighbors are responded to more slowly than words with few phonological neighbors (Luce & Pisoni, 1998).

Our present finding—that lexical island words were responded to more quickly and accurately than giant component words—could be an example of the “isolation effect” described in the SIMPLE model (Brown et al., 2006). Both giant component and lexical island words were matched on neighborhood density/degree and clustering coefficients, so the average distance between giant component and lexical island words and their *local* neighbors are equivalent. Given that lexical islands are smaller connected components that do not connect to other network components, lexical island words are free of distant neighbors, which allow words in the lexical islands to “stand out” to

a greater extent in these sparsely populated regions of phonological similarity-space. On the other hand, giant component words are connected to many distant neighbors, reducing the discriminability of the target word, making it difficult for the giant component word to “stand out” and to be easily retrieved. Therefore, lexical island words may be more efficiently responded to than giant component words due to greater “phonological distinctiveness” from the rest of the lexicon. Our results suggest that in addition to nearby neighbors it is just as important to consider the influence of *distant* neighbors on the discriminability of the target word. The network science approach offers one way to quantify and study the influence of these distant neighbors on lexical processing by considering the overall structure of the phonological network.

Recent studies (e.g., Chan & Vitevitch, 2009, 2010; Vitevitch et al., 2012, 2014) as well as the present work show that applying the network science approach in psycholinguistics can lead to a more nuanced understanding of the processes and representations involved in spoken word recognition. It is important to note that the network science approach emphasizes how the structure (at various levels) of the network can influence processing in that system (Strogatz, 2001). Without explicitly appealing to the overall structure of the lexicon, it is difficult to see how current models of spoken word recognition would account for these findings.

The present finding—lexical island words have a processing advantage over giant component words—suggests that smaller connected components in the lexical network may play an important compensatory role in lexical retrieval, enabling words that, due to certain lexical characteristics, should be “at risk” for slow, laborious, and errorful retrieval to be retrieved with relative ease and efficiency. Not being connected to the giant component may limit, impair, or otherwise challenge certain processes found in other domains examined by network science, leading to a disadvantage for items in the smaller components. However, in the case of the mental lexicon, being “exiled” to a lexical island may be accompanied by certain advantages (e.g., fewer nodes to “share” activation with, a shorter network diameter, greater distinctiveness from the rest of the lexicon), resulting in a trade-off of sorts that makes for rapid and efficient language processing overall. That is, what appears to be an ineffective *structure*—a small “giant” component, and many lexical islands and hermits—may, in fact, contribute to an efficient *process* overall.

The peculiar structure of the lexical network might also enable the language system to maintain connectivity and resist succumbing to damage (i.e., it is robust), yet remain flexible enough to continue to evolve over time. In network science, the *robustness* of a network refers to the ability of a network to continue to function (as measured by changes in the average path length of the network which assesses the overall connectivity of the network) despite the removal of nodes from the network (Albert & Barabási, 2002). Highly connected nodes can be targeted for removal, or nodes can be removed via random selection. Typically a targeted attack on a network results in a large increase in the average path length—one needs several more steps in order to traverse the network, indicating that the functioning of the network is compromised (Albert & Barabási, 2002). It is interesting to note that Arbesman et al. (2010b) found that the phonological network was extremely resistant to *both* random and targeted attacks. As the vast majority of complex networks typically have a very large proportion of nodes residing in their giant components (Newman, 2001), this extraordinarily high amount of robustness of the phonological network could be due to the unique structure of the mental

lexicon—a small “giant” component, and many lexical islands and hermits.

Evolvability refers to the ability of the network to adapt over time (Payne & Wagner, 2014). Language is widely regarded as a highly evolvable system (Ke, Minett, Au, & Wang, 2002; Nowak & Krakauer, 1999; Solé, Corominas-Murtra, Valverde, & Steels, 2010). The unique structure of the phonological network might also contribute to the evolvability of language. For instance, new words may connect to words in the lexical islands or to hermits instead of to words in the giant component. Such a situation would help maintain stability in the overall function of the phonological network by keeping constant the structure and size of the giant component (where most lexical processing presumably occurs), while at the same time promoting growth in the mental lexicon by incorporating novel combinations of speech sounds (i.e., new words) into the smaller components of the phonological network without drastically influencing the overall functioning of the language network.

Although some researchers have noted a trade-off between robustness and evolvability of complex networks (Ciliberti, Martin, & Wagner, 2007), the phonological network may represent an exception to this general observation. In particular, the structure of the phonological network—a “small” giant component (in comparison with other real-world networks) and the presence of several smaller connected components and isolates—may afford the phonological network the best of both worlds: the ability to withstand random and targeted removal of nodes from the network (robustness), and the ability to adapt to environmental and sociocultural changes (evolvability).

The results of these experiments have shown that where in the network a word is located (the giant component vs. a lexical island) plays an important role in spoken word recognition and memory processes, over and beyond the lexical and network characteristics of individual words. It is important to note that the present work has shown that, in addition to the structure at the micro- and meso-levels of the phonological network, the *macro-level* structure of the phonological network has important implications for spoken word recognition processes, contributing to the growing body of research showing that the structure of the mental lexicon has measurable influences on cognitive processes (e.g., De Deyne, Navarro, Perfors, & Storms, 2012; Hills, Maouene, Maouene, Sheya, & Smith, 2009).

References

- Albert, R., & Barabási, A. L. (2002). Statistical mechanics of complex networks. *Reviews of Modern Physics*, 74, 47–97. <http://dx.doi.org/10.1103/RevModPhys.74.47>
- Anderson, J. R. (1983). A spreading activation theory of memory. *Journal of Verbal Learning and Verbal Behavior*, 22, 261–295. [http://dx.doi.org/10.1016/S0022-5371\(83\)90201-3](http://dx.doi.org/10.1016/S0022-5371(83)90201-3)
- Arbesman, S., Strogatz, S. H., & Vitevitch, M. S. (2010a). Comparative analysis of networks of phonologically similar words in English and Spanish. *Entropy*, 12, 327–337. <http://dx.doi.org/10.3390/e12030327>
- Arbesman, S., Strogatz, S. H., & Vitevitch, M. S. (2010b). The structure of phonological networks across multiple languages. *International Journal of Bifurcation and Chaos in Applied Sciences and Engineering*, 20, 679–685. <http://dx.doi.org/10.1142/S021812741002596X>
- Baddeley, A. D., Thomson, N., & Buchanan, M. (1975). Word length and the structure of short-term memory. *Journal of Verbal Learning and Verbal Behavior*, 14, 575–589. [http://dx.doi.org/10.1016/S0022-5371\(75\)80045-4](http://dx.doi.org/10.1016/S0022-5371(75)80045-4)
- Balota, D. A., Yap, M. J., Hutchison, K. A., Cortese, M. J., Kessler, B., Loftis, B., . . . Treiman, R. (2007). The English Lexicon Project. *Behavior Research Methods*, 39, 445–459. <http://dx.doi.org/10.3758/BF03193014>
- Barabási, A. L. (2009). Scale-free networks: A decade and beyond. *Science*, 325, 412–413. <http://dx.doi.org/10.1126/science.1173299>
- Batagelj, V., & Mrvar, A. (1988). Pajek—Program for large network analysis. *Connections*, 21, 47–57.
- Brown, G. D. A., Morin, C., & Lewandowsky, S. (2006). Evidence for time-based models of free recall. *Psychonomic Bulletin & Review*, 13, 717–723. <http://dx.doi.org/10.3758/BF03193986>
- Brown, G. D. A., Neath, I., & Chater, N. (2007). A temporal ratio model of memory. *Psychological Review*, 114, 539–576. <http://dx.doi.org/10.1037/0033-295X.114.3.539>
- Brysbaert, M., & New, B. (2009). Moving beyond Kučera and Francis: A critical evaluation of current word frequency norms and the introduction of a new and improved word frequency measure for American English. *Behavior Research Methods*, 41, 977–990. <http://dx.doi.org/10.3758/BRM.41.4.977>
- Chan, K. Y., & Vitevitch, M. S. (2009). The influence of the phonological neighborhood clustering coefficient on spoken word recognition. *Journal of Experimental Psychology: Human Perception and Performance*, 35, 1934–1949. <http://dx.doi.org/10.1037/a0016902>
- Chan, K. Y., & Vitevitch, M. S. (2010). Network structure influences speech production. *Cognitive Science*, 34, 685–697. <http://dx.doi.org/10.1111/j.1551-6709.2010.01100.x>
- Ciliberti, S., Martin, O. C., & Wagner, A. (2007). Innovation and robustness in complex regulatory gene networks. *Proceedings of the National Academy of Sciences of the United States of America*, 104, 13591–13596. <http://dx.doi.org/10.1073/pnas.0705396104>
- Collins, A. M., & Loftus, E. F. (1975). A spreading-activation theory of semantic processing. *Psychological Review*, 82, 407–428. <http://dx.doi.org/10.1037/0033-295X.82.6.407>
- Craik, F. I. M. (1968). Two components in free recall. *Journal of Verbal Learning and Verbal Behavior*, 7, 996–1004. [http://dx.doi.org/10.1016/S0022-5371\(68\)80058-1](http://dx.doi.org/10.1016/S0022-5371(68)80058-1)
- De Deyne, S., Navarro, D. J., Perfors, A., & Storms, G. (2012, August). *Strong structure in weak semantic similarity: A graph based account*. Paper presented at the Proceedings of the 34th Annual Conference of the Cognitive Science Society, Sapporo, Japan.
- Ebbinghaus, H. (1913). *Memory: A contribution to experimental psychology*. New York, NY: Teachers College, Columbia University. <http://dx.doi.org/10.1037/10011-000>
- Ellis, A. W. (1980). Errors in speech and short-term memory: The effects of phonemic similarity and syllable position. *Journal of Verbal Learning and Verbal Behavior*, 19, 624–634. [http://dx.doi.org/10.1016/S0022-5371\(80\)90672-6](http://dx.doi.org/10.1016/S0022-5371(80)90672-6)
- Gaskell, M. G., & Marslen-Wilson, W. D. (1997). Integrating form and meaning: A distributed model of speech perception. *Language and Cognitive Processes*, 12, 613–656. <http://dx.doi.org/10.1080/016909697386646>
- Goldstein, R., & Vitevitch, M. S. (2014). The influence of clustering coefficient on word-learning: How groups of similar sounding words facilitate acquisition. *Frontiers in Psychology*, 5, 1307. <http://dx.doi.org/10.3389/fpsyg.2014.01307>
- Hills, T. T., Maouene, M., Maouene, J., Sheya, A., & Smith, L. (2009). Longitudinal analysis of early semantic networks: Preferential attachment or preferential acquisition? *Psychological Science*, 20, 729–739. <http://dx.doi.org/10.1111/j.1467-9280.2009.02365.x>
- Hulme, C., Maughan, S., & Brown, G. D. A. (1991). Memory for familiar and unfamiliar words: Evidence for a long-term memory contribution to short-term memory span. *Journal of Memory and Language*, 30, 685–701. [http://dx.doi.org/10.1016/0749-596X\(91\)90032-F](http://dx.doi.org/10.1016/0749-596X(91)90032-F)

- Hulme, C., Roodenrys, S., Schweickert, R., Brown, G. D. A., Martin, M., & Stuart, G. (1997). Word-frequency effects on short-term memory tasks: Evidence for a redintegration process in immediate serial recall. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 23, 1217–1232. <http://dx.doi.org/10.1037/0278-7393.23.5.1217>
- Ke, J., Minett, J. W., Au, C.-P., & Wang, W. S.-Y. (2002). Self-organization and selection in the emergence of vocabulary. *Complexity*, 7, 41–54. <http://dx.doi.org/10.1002/cplx.10030>
- Kučera, H., & Francis, W. N. (1967). *Computational analysis of present-day American English*. Providence, RI: Brown University Press.
- 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 and 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)
- Locker, L., Jr., Hoffman, L., & Bovaird, J. A. (2007). On the use of multilevel modeling as an alternative to items analysis in psycholinguistic research. *Behavior Research Methods*, 39, 723–730. <http://dx.doi.org/10.3758/BF03192962>
- Luce, P. A. (1986). A computational analysis of uniqueness points in auditory word recognition. *Perception & Psychophysics*, 39, 155–158. <http://dx.doi.org/10.3758/BF03212485>
- 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>
- MacKay, D. G. (1990). Errors, ambiguity, and awareness in language perception and production. In B. Baars (Ed.), *The psychology of error: A window on the mind* (pp. 39–69). New York, NY: Plenum Press.
- Marslen-Wilson, W. D. (1987). Functional parallelism in spoken word-recognition. *Cognition*, 25, 71–102. [http://dx.doi.org/10.1016/0010-0277\(87\)90005-9](http://dx.doi.org/10.1016/0010-0277(87)90005-9)
- 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)
- Mirman, D., Strauss, T. J., Brecher, A., Walker, G. M., Sobel, P., Dell, G. S., & Schwartz, M. F. (2010). A large, searchable, web-based database of aphasic performance on picture naming and other tests of cognitive function. *Cognitive Neuropsychology*, 27, 495–504. <http://dx.doi.org/10.1080/02643294.2011.574112>
- Morton, J. (1969). Interaction of information in word recognition. *Psychological Review*, 76, 165–178. <http://dx.doi.org/10.1037/h0027366>
- Nelson, D. L., McKinney, V. M., Gee, N. R., & Janczura, G. A. (1998). Interpreting the influence of implicitly activated memories on recall and recognition. *Psychological Review*, 105, 299–324. <http://dx.doi.org/10.1037/0033-295X.105.2.299>
- Newman, M. E. J. (2001). The structure of scientific collaboration networks. *Proceedings of the National Academy of Sciences of the United States of America*, 98, 404–409. <http://dx.doi.org/10.1073/pnas.98.2.404>
- Newman, M. E. J. (2002). Assortative mixing in networks. *Physical Review Letters*, 89, 208701. <http://dx.doi.org/10.1103/PhysRevLett.89.208701>
- Norris, D., & McQueen, J. M. (2008). Shortlist B: A Bayesian model of continuous speech recognition. *Psychological Review*, 115, 357–395. <http://dx.doi.org/10.1037/0033-295X.115.2.357>
- Nowak, M. A., & Krakauer, D. C. (1999). The evolution of language. *Proceedings of the National Academy of Sciences of the United States of America*, 96, 8028–8033. <http://dx.doi.org/10.1073/pnas.96.14.8028>
- 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.
- Payne, J. L., & Wagner, A. (2014). The robustness and evolvability of transcription factor binding sites. *Science*, 343, 875–877. <http://dx.doi.org/10.1126/science.1249046>
- Roodenrys, S., Hulme, C., Lethbridge, A., Hinton, M., & Nimmo, L. M. (2002). Word-frequency and phonological-neighborhood effects on verbal short-term memory. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 28, 1019–1034. <http://dx.doi.org/10.1037/0278-7393.28.6.1019>
- Schweickert, R. (1993). A multinomial processing tree model for degradation and redintegration in immediate recall. *Memory & Cognition*, 21, 168–175. <http://dx.doi.org/10.3758/BF03202729>
- Siew, C. S. Q. (2013). Community structure in the phonological network. *Frontiers in Psychology*, 4, 553. <http://dx.doi.org/10.3389/fpsyg.2013.00553>
- Solé, R. V., Corominas-Murtra, B., Valverde, S., & Steels, L. (2010). Language networks: Their structure, function, and evolution. *Complexity*, 15, 20–26.
- Steyvers, M., & Tenenbaum, J. B. (2005). The large-scale structure of semantic networks: Statistical analyses and a model of semantic growth. *Cognitive Science*, 29, 41–78. http://dx.doi.org/10.1207/s15516709cog2901_3
- Strauss, T. J., Harris, H. D., & Magnuson, J. S. (2007). jTRACE: A reimplementation and extension of the TRACE model of speech perception and spoken word recognition. *Behavior Research Methods*, 39, 19–30. <http://dx.doi.org/10.3758/BF03192840>
- Strogatz, S. H. (2001, March 8). Exploring complex networks. *Nature*, 410, 268–276. <http://dx.doi.org/10.1038/35065725>
- Tehan, G., & Humphreys, M. S. (1988). Articulatory loop explanations of memory span and pronunciation rate correspondences: A cautionary note. *Bulletin of the Psychonomic Society*, 26, 293–296. <http://dx.doi.org/10.3758/BF03337662>
- Treiman, R., Mullennix, J., Bijeljac-Babic, R., & Richmond-Welty, E. D. (1995). The special role of rimes in the description, use, and acquisition of English orthography. *Journal of Experimental Psychology: General*, 124, 107–136. <http://dx.doi.org/10.1037/0096-3445.124.2.107>
- Tyler, L. K., & Wessels, J. (1983). Quantifying contextual contributions to word-recognition processes. *Perception & Psychophysics*, 34, 409–420. <http://dx.doi.org/10.3758/BF03203056>
- Vitevitch, M. S. (1997). The neighborhood characteristics of malapropisms. *Language and Speech*, 40, 211–228.
- Vitevitch, M. S. (2002). The influence of phonological similarity neighborhoods on speech production. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 28, 735–747.
- Vitevitch, M. S. (2008). What can graph theory tell us about word learning and lexical retrieval? *Journal of Speech, Language, and Hearing Research*, 51, 408–422. [http://dx.doi.org/10.1044/1092-4388\(2008\)030](http://dx.doi.org/10.1044/1092-4388(2008)030)
- Vitevitch, M. S., & Castro, N. (2015). Using network science in the language sciences and clinic. *International Journal of Speech-Language Pathology*, 17, 13–25. <http://dx.doi.org/10.3109/17549507.2014.987819>
- 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., Chan, K. Y., & Roodenrys, S. (2012). Complex network structure influences processing in long-term and short-term memory. *Journal of Memory and Language*, 67, 30–44. <http://dx.doi.org/10.1016/j.jml.2012.02.008>
- Vitevitch, M. S., Ercal, G., & Adagarla, B. (2011). Simulating retrieval from a highly clustered network: Implications for spoken word recognition. *Frontiers in Psychology*, 2, 369. <http://dx.doi.org/10.3389/fpsyg.2011.00369>
- Vitevitch, M. S., & Goldstein, R. (2014). Keywords in the mental lexicon. *Journal of Memory and Language*, 73, 131–147. <http://dx.doi.org/10.1016/j.jml.2014.03.005>
- 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. (2004). A web-based interface to calculate phonotactic probability for words and nonwords in English. *Behavior Research Methods*, 36, 481–487. <http://dx.doi.org/10.3758/BF03195594>
- Vitevitch, M. S., & Rodríguez, E. (2005). Neighborhood density effects in spoken word recognition in Spanish. *Journal of Multilingual Communication Disorders*, 3, 64–73. <http://dx.doi.org/10.1080/14769670400027332>
- Watkins, M. J. (1977). The intricacy of memory span. *Memory & Cognition*, 5, 529–534. <http://dx.doi.org/10.3758/BF03197396>
- Watts, D. J. (2004). The “new” science of networks. *Annual Review of Sociology*, 30, 243–270. <http://dx.doi.org/10.1146/annurev.soc.30.020404.104342>
- Watts, D. J., & Strogatz, S. H. (1998, June 4). Collective dynamics of “small-world” networks. *Nature*, 393, 440–442. <http://dx.doi.org/10.1038/30918>
- Yap, M. J., & Balota, D. A. (2009). Visual word recognition of multisyllabic words. *Journal of Memory and Language*, 60, 502–529. <http://dx.doi.org/10.1016/j.jml.2009.02.001>
- Yates, M. (2013). How the clustering of phonological neighbors affects visual word recognition. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 39, 1649–1656. <http://dx.doi.org/10.1037/a0032422>
- Zipf, G. K. (1935). *The psycho-biology of language*. Boston, MA: Houghton Mifflin.

Appendix A

List of Lexical Island and Giant Component Words and Mean Reaction Times (RT; in milliseconds) for Experiments 1 and 2

Word	Mean RT (Exp 1: Naming)	Mean RT (Exp 2: Lexical decision)
Lexical island words		
beckon	900	1,068
mission	885	811
portion	950	1,056
taken	869	886
concede	1,035	1,080
concern	1,062	1,033
confine	996	923
consign	1,045	1,206
coffin	954	900
deafen	916	1,060
siphon	1,027	1,053
soften	1,065	961
banish	955	955
furnish	1,059	1,063
manage	939	983
marriage	932	917
domain	926	988
regain	910	941
remain	963	1,004
retain	968	970
partition	987	993
permission	932	991
petition	958	928
position	1,021	979
central	1,009	1,064
locus	930	988
notice	914	915
report	926	1,011
lizard	918	933
nervous	944	930
service	1,059	1,062
warrant	922	1,045
happen	908	1,007
margin	914	934
peasant	883	842
revolve	902	948
gallop	871	869

(Appendix continues)

Appendix A (*continued*)

Word	Mean RT (Exp 1: Naming)	Mean RT (Exp 2: Lexical decision)
nominee	1,013	1,097
felon	994	921
village	982	948
cunning	956	930
retail	935	968
memory	886	924
trophy	937	877
treasure	908	875
solemn	1,017	1,151
radio	907	850
plaza	981	1,111
Giant component words		
parcel	962	1,027
ceiling	1,014	1,116
driven	865	908
temple	886	886
comic	923	994
century	1,005	1,018
panther	971	990
facet	975	1,234
cumber	916	1,033
stutter	1,019	940
rollick	973	1,417
scepter	1062	1,152
brittle	821	855
filing	1,049	1,119
scant	1,100	1,175
mountain	933	938
spiral	1,095	1,108
drench	976	1,072
repeat	884	874
grunt	897	883
coroner	1,034	1,146
remind	922	1,077
defend	955	989
mention	961	901
receive	989	1,021
limber	949	1,096
hardly	975	1,089
minute	907	935
squid	1,154	1,043
straighten	1,119	1,059
supposed	1,126	1,087
collect	932	956
mustard	956	949
cartridge	933	914
languor	984	1,058
dribble	908	928
magnet	979	1,072
parable	943	1,063
remit	953	1,199
reverse	864	954
knowledge	941	975
device	920	895
chapter	945	947
hamper	936	892
temporal	950	1,020
colleague	980	1,048
danger	874	800
salvage	1,037	1,066

(Appendices continue)

Appendix B

Nonwords Used in Experiment 2

Nonword	IPA	Nonword	IPA
beton	bɛtɪn	porcel	pɔ:sl
mizzion	mɪʒɪn	ceilong	sɪlɒŋ
ponfion	pɒnfɪn	druven	druvən
tapen	tɛpɪn	tample	tɪmpl
conzede	kənzɪd	cowic	kəwɪk
conbern	kənbɜ:n	cendury	sɛndə:ɪ
comfine	kəmfɪn	panker	pænkə:
conlign	kənlɪn	fapet	fæsɒt
cothin	kɔθɪn	cumler	kʌmlə:
deamen	dɛmən	stotter	stɔɪtə:
suphon	sufən	romick	rɒmɪk
saften	sʌfən	sepger	sɛpgə:
bonish	bɒnɪʃ	brimble	bɪmbl
fugish	fʊʒɪʃ	fileg	fɪlɛŋ
magage	mæɡɪdʒ	scaft	skæft
madiage	mædɪdʒ	moontain	mʊntɪn
dogain	dɒɡɛn	spooral	spʊəl
reshain	rɪʃɛn	drenth	dɹɛntʃ
redain	rɪdɛn	repout	rɪpaʊt
refain	rɪfɛn	glunt	ɡlʌnt
partution	pə:tu:ʃɪn	corofer	kɔ:ɹədə:
permission	pə:nlɪʃɪn	rehind	rɪhaɪnd
perition	pə:nlɪʃɪn	dejend	dɛjɛnd
polition	pə:lɪʃɪn	mendion	mɛndɪn
cendral	sɛndɹl	rebeive	rɪbɪv
lercus	lɜ:kɪs	lomber	lɒmbə:
nopice	nɒpɪs	harply	hɑ:plɪ
rekort	rɪkɔ:ɹt	minupe	mɪnɒt
lipard	lɪpə:ɹd	sqad	spwɪd
navous	nʌvəs	stroten	strɒtn
sernice	sɜ:nəs	summosed	səmozd
weerant	wɪɹənt	colluct	kəlukt
halen	hælən	musgard	mʌsgə:ɹd
mardin	mɑ:ɹdɪn	curtridge	kʊ:ɹɪdʒ
peasart	pɛzɑ:ɹt	langdor	læŋdə:
repolve	rɪpɒlv	driggle	dɹɪɡl
goillop	ɡɔ:ɪlɒp	magzet	mægzɪt
nomidee	nɒmɪdi	paradle	pæ:ɹədl
fenon	fɛnɪn	remut	rɪmʌt
vittage	vɪtɪdʒ	relerse	rɪlɜ:s
cuzzing	kʌʒɪŋ	knowpedge	nəpɪdʒ
rezail	rɪzɛɪl	degice	dɛɡaɪs
medory	mɛdɜ:ɪ	chapmer	tʃæpmə:
tromy	tɹɒmɪ	hamler	hæmlə:
trosure	tɹɔ:ʒə:	temtoral	tɛmtə:ɹl
somemn	səməm	comeague	kəmɪɡ
ravio	rɛvɪɔ	dadger	dɛddʒə:
plara	plæ:ɹə	salgage	sælɡɪdʒ

Note. IPA = International Phonetic Alphabet.

Received May 2, 2014

Revision received February 4, 2015

Accepted March 24, 2015 ■