



Using network science in the language sciences and clinic

Michael S. Vitevitch & Nichol Castro

To cite this article: Michael S. Vitevitch & Nichol Castro (2015) Using network science in the language sciences and clinic, International Journal of Speech-Language Pathology, 17:1, 13-25, DOI: [10.3109/17549507.2014.987819](https://doi.org/10.3109/17549507.2014.987819)

To link to this article: <http://dx.doi.org/10.3109/17549507.2014.987819>



View supplementary material 



Published online: 24 Dec 2014.



Submit your article to this journal 



Article views: 93



View related articles 



View Crossmark data 

Using network science in the language sciences and clinic

MICHAEL S. VITEVITCH & NICHOL CASTRO

Department of Psychology, University of Kansas, Lawrence, KS, USA

Abstract

A number of variables—word frequency, word length—have long been known to influence language processing. This study briefly reviews the effects in speech perception and production of two more recently examined variables: phonotactic probability and neighbourhood density. It then describes a new approach to study language, network science, which is an interdisciplinary field drawing from mathematics, computer science, physics and other disciplines. In this approach, nodes represent individual entities in a system (i.e. phonological word-forms in the lexicon), links between nodes represent relationships between nodes (i.e. phonological neighbours) and various measures enable researchers to assess the micro-level (i.e. the individual word), the macro-level (i.e. characteristics about the whole system) and the meso-level (i.e. how an individual fits into smaller sub-groups in the larger system). Although research on individual lexical characteristics such as word-frequency has increased understanding of language processing, these measures only assess the ‘micro-level’. Using network science, researchers can examine words at various levels in the system and how each word relates to the many other words stored in the lexicon. Several new findings using the network science approach are summarized to illustrate how this approach can be used to advance basic research as well as clinical practice.

Keywords: *Adults, children, psycholinguistic*

Introduction

Several decades ago, Cutler (1981) commented on the large number of factors that language scientists identified as influences on the production, recognition or acquisition of spoken words. These factors included, among others, semantic ambiguity of a word, number of meanings of a word, the length of the word, the stress-pattern of the word, concreteness of the word, the age at which the word was first learned (*Age of Acquisition*; AoA), the frequency with which the word occurs in the ambient language (word-frequency), morphological complexity of the word and the recognition point of a word (i.e. the point in a word that it becomes unique from all other words in the lexicon). The number of factors that researchers have identified as influences on the production, recognition or acquisition of words has only increased since that time and now includes *phonotactic probability* and *neighbourhood density*, in addition to many others.

Phonotactic probability refers to the frequency with which segments and sequences of segments occur in words (Vitevitch & Luce, 2005). A word, like *back*, or a non-word, like /fʌl/, which contain common segments (/f/, /ʌ/, /l/) and sequences of segments

(/fʌ/ and /ʌl/) that co-occur often in the language, is said to have high phonotactic probability, whereas a word, like *bag*, or a non-word, like /ʃʌtʃ/, which contain less common segments and sequences of segments that co-occur rarely in the language, is said to have low phonotactic probability.

Numerous studies have found that phonotactic probability influences various language processes: (1) *word segmentation* in infants (Mattys, Jusczyk, Luce, & Morgan, 1999), (2) the *production of spoken words* in adults (Goldrick & Larson, 2008; Vitevitch, Armbruster, & Chu, 2004), in children who stutter (Anderson & Byrd, 2008), in typically-developing children (Zamuner, Gerken, & Hammond, 2004) as well as the repetition accuracy in speakers with acquired output impairment after stroke (Lallini & Miller, 2011), (3) the *recognition of words* in adults with normal hearing and in adult users of cochlear implants (Vitevitch, Pisoni, Kirk, Hay-McCutcheon, & Yount, 2002), (4) *word learning* in typically-developing children (Storkel & Hoover, 2011), in children with Specific Language Impairment (SLI; Gray, Brinkley, & Svetina, 2012), in children with phonological delays (Storkel & Hoover, 2010b), in late talkers (MacRoy-Higgins, Schwartz, Shafer, &

Marton, 2013), and in adults (Storkel, Armbrüster, & Hogan, 2006) and (5) the *conjugation of verbs* in children with SLI (Leonard, Davis, & Deevy, 2007).

Phonotactic probability has also been implicated in *memory* for non-words (Gathercole, Frankish, Pickering, & Peaker, 1999; Messer, Leseman, Boom, & Mayo, 2010) and in *literacy-related skills*, as in children learning to spell (Apel, Wolter, & Masterson, 2006) and as evidenced by differences in processing stimuli that vary in phonotactic probability in children and adults with dyslexia (Bonte, Poelmans, & Blomert, 2007; Noordenbos, Segers, Mitterer, Serniclaes, & Verhoeven, 2013). Finally, both electro- and magneto-physiological components have been identified for the processing of stimuli that vary in phonotactic probability (Hunter, 2013; Pylkkänen, Stringfellow, & Marantz, 2002).

Neighbourhood density refers to the number of words that sound like a target word. A word is said to be a phonological neighbour of a target word if the substitution, addition or deletion of a single phoneme in any position in that word converts it to the target word (e.g. Greenberg & Jenkins, 1964; Landauer & Streeter, 1973; Luce & Pisoni, 1998). For example, the words *hat*, *cut*, *cap*, *scat* and *at* are considered neighbours of the word *cat* (*cat* has other words as neighbours, but only a few were listed for illustrative purposes). A word with many phonological neighbours is said to have a dense neighbourhood, whereas a word with few phonological neighbours is said to have a sparse neighbourhood.

Numerous studies have found that neighbourhood density influences various language processes: (1) the *acquisition of sounds* in children (Gierut, Morrisette, & Champion, 1999), (2) the *acquisition of words* in children (Storkel, 2004) and in second language learners (Stamer & Vitevitch, 2012; see also computational work in Vitevitch & Storkel, 2013), (3) *spoken word recognition* in young adults with no history of speech, language or hearing impairment in English and in Spanish (Luce & Pisoni, 1998; see also Vitevitch, 2002; Vitevitch & Luce 1998, 1999; Vitevitch & Rodriguez, 2005), in older adults with no history of speech, language or hearing impairment (e.g. Sommers, 1996) and in post-lingually deafened adults who had a cochlear implant (Kaiser, Kirk, Lachs, & Pisoni, 2003), as well as the recognition of accented speech (Chan & Vitevitch, in press; Imai, Walley, & Flege, 2005), (4) *spoken word production* in children who stutter (Arnold, Conture, & Ohde, 2005), in young adults with fluent speech in English and in Spanish (Munson & Solomon, 2004; Vitevitch, 1997, 2002; Vitevitch & Stamer, 2006), in older adults with fluent speech (Vitevitch & Sommers, 2003), in individuals with aphasia (Gordon, 2002) and even (5) *reading* by young adults with no history of speech, language or hearing impairment (Yates, Locker, & Simpson, 2004). For a more complete review of how neighbourhood density and pho-

notactic probability influence the perception and production of spoken words, see Vitevitch and Luce (2016).

In working with various colleagues investigating how phonotactic probability and neighbourhood density influence the production and recognition of spoken words we observed several interesting relationships. One observation was that words (or non-words) comprised of common segments and sequences of segments tended to be similar to many words in the language. That is, a word or non-word with high phonotactic probability tends to have a dense phonological neighbourhood. The relationship that we observed between phonotactic probability and neighbourhood density was easily quantified and captured in the statistically significant positive correlation between phonotactic probability and neighbourhood density (Vitevitch, Luce, Pisoni, & Auer, 1999); see Storkel and Lee (2011) for an attempt to dissociate the influence of these two variables.

Additional observations that were made, but at the time could not be quantified so easily, were that: (1) a phonological neighbor of one word was also a phonological neighbour of other words and (2) some words had phonological neighbours that tended to be neighbours just with that specific word, whereas other words had phonological neighbours that were also neighbours with other neighbours of that word. With the emergence of the field now known as *Network Science* (Newman, 2010), a suite of mathematical tools that could be used to quantify these and other relationships among phonological word-forms in the mental lexicon came to our attention. Several studies, briefly summarized in what follows, demonstrate that the way in which phonological word-forms are organized in the mental lexicon influences lexical processes such as the production, recognition and acquisition of spoken words.

The focus on the overall structure of the lexicon differs from a more mainstream psycholinguistic approach, which tends to focus on characteristics of an individual word—including phonotactic probability and neighbourhood density, as well as many of the variables described by Cutler (1981)—for an explanation of why some words are processed differently from others. What is most striking about the studies reviewed below is that these measures of individual words were controlled in the studies described below, thereby providing clear evidence that differences in the structural organization of words in the lexicon influence various lexical processes.

What is network science?

Network science draws on techniques used in mathematics, sociology, computer science, physics and a number of other fields to examine complex systems. What makes complex systems interesting to study is that the ‘whole’ is often ‘greater than the sum of its

parts', meaning that the interaction among entities in the system leads to system-wide behaviours that cannot be predicted from the local interaction that occurs between two adjacent entities. This holistic approach contrasts with the reductionist approach typically employed in contemporary psycholinguistics (and science in general) in which characteristics of individual entities—such as how frequent a word occurs in the ambient language—are believed to influence processing in the system. Therefore, these network science measures may capture important *interactions among words* that might be useful for understanding processing in the mental lexicon and for the development of language interventions.

One way to model a complex system is to use *nodes* (sometimes called *vertices*) to represent individual entities in the system and *connections* to represent relationships between entities in the system (the terms *edges*, *directed edges* or *arcs* are sometime used for connections that indicate a relationship in one direction, such as predator–prey relationships). When assembled, the nodes and connections of a system form a web-like structure or *network* (sometimes called a *graph*) to represent the entire system. As noted above, there are a number of disciplines that contribute to Network Science. One discipline tends to use one set of terms, whereas another discipline will use the other set of terms to refer to the same concepts. In the present case, we will use the terms: *node*, *connection* and *network* (see Supplementary Appendix to be found online at <http://informahealthcare.com/doi/abs/10.3109/17549507.2014.987819>—Key Terms for definitions of the terms that we introduce in what follows).

The network approach has been used to examine complex systems in economical, biological, social and technological domains (Barabási, 2009). An intuitive example of network analysis and its application is found in the social domain (i.e. a social network) in which nodes represent members of a street gang and connections are placed between gang members who participate together in gang-related activities, such as distributing illegal drugs. Law enforcement officers employing network analysis techniques could identify members of the gang who are the biggest suppliers of drugs to the other gang-members (e.g. which node directly connects to the most nodes in the system) and try to turn that gang member into a confidential informant, thereby providing law enforcement officers with important updates on the use and distribution of drugs. Alternatively, law enforcement officers could 'remove' that individual from the network by arresting that individual. Such an action could maximally disrupt the distribution of illegal drugs by the gang (until another gang member steps-in to fill that recently vacated role).

These same analysis techniques can be used to model the species in an eco-system, with connections indicating which species prey upon other

species (Montoya & Solé, 2002). Removal of a node in this case is equivalent to a species going extinct. Determining the broader and indirect effects on the ecosystem of one vs another species going extinct could provide invaluable information for individuals engaged in conservation efforts to decide where to best direct their limited resources.

More relevant to the language sciences, this approach has been used to examine connectivity in the brain (Sporns, 2010) and the cognitive processes and representations involved in semantic memory (Hills, Maouene, Maouene, Sheya, & Smith, 2009; Steyvers & Tenenbaum, 2005). In the research summarized here, we will focus on a network of phonological word-forms in the mental lexicon. We will also discuss other ways to use network science to study language and to provide clinical insight.

Network science and the phonological lexicon

Vitevitch (2008) applied the tools of network science to the mental lexicon by creating a network with ~ 20,000 English words represented as nodes¹ and connections placed between words that were phonologically similar. To operationally define 'phonologically similar', a commonly used metric was employed (i.e. the one-phoneme metric used in Luce & Pisoni, 1998). In the present case two words were connected if the addition, deletion or substitution of a phoneme in one word formed the other word. For example, the nodes for *cat* /kæt/ and *bat* /bæt/ would be connected (the underlined phonemes indicate where in the words the changed phoneme occurred). Figure 1 shows a small portion of this network.

Analysis of the whole network revealed several noteworthy characteristics about the structure of the mental lexicon. (The reader is encouraged to consult the sources cited herein, as well as other sources for more technical definitions of the various network measures that are described here.) Vitevitch (2008) found that the phonological network had a large group of nodes that were highly connected to each other (known as the giant component), as well as many smaller groups of words that were connected to each other, but not to the giant component. Such items are referred to in the network science literature as smaller components, but Vitevitch (2008) adopted the term 'lexical islands' to describe such groupings in the mental lexicon. An example of a lexical island is the component that Vitevitch (2008) referred to as the 'island of the shunned', because words in that component contained the sequence of segments /ʃʌn/, such as *faction*, *fiction* and *fission*. Vitevitch (2008) further observed that the lexical network contained many words that did not have any phonological neighbours. An individual node that is not connected to any other nodes is known as an isolate in the network science literature, but Vitevitch adopted the

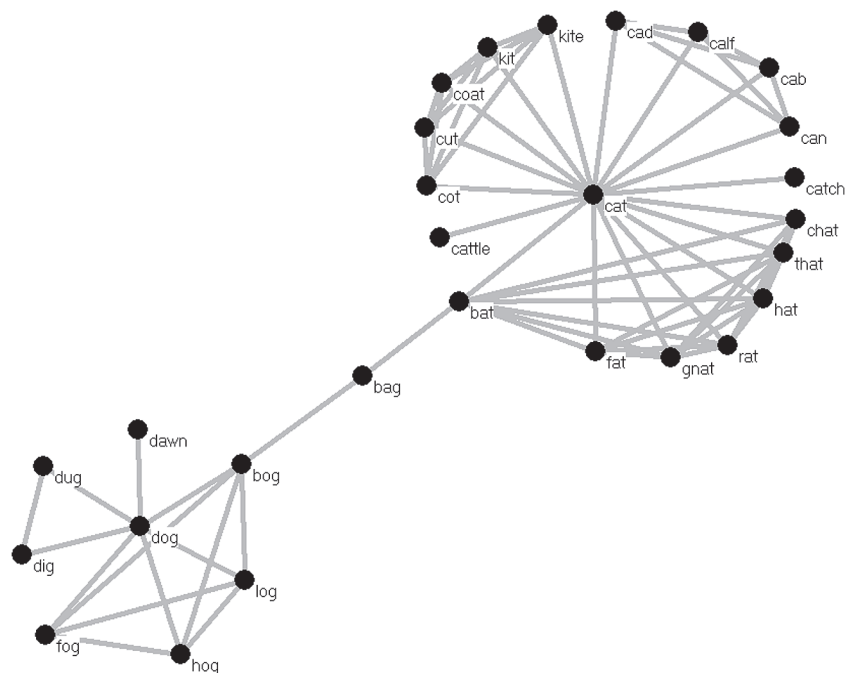


Figure 1. A small portion of the network examined in Vitevitch (2008). Each word is, of course, connected to other words in the lexicon, but only a few words are displayed here for the sake of image clarity. See additional discussion in the text for illustrations of certain network science concepts.

term ‘lexical hermits’ to describe such words in the mental lexicon.

Vitevitch (2008) further examined the giant component and found that it exhibited small-world characteristics (Watts & Strogatz, 1998). A network is said to exhibit small-world characteristics if it has a ‘short’ *average path length* (meaning that, on average, one can get from one randomly selected node to another randomly selected node in the network by traversing a small number of connections) and, relative to what would be expected, a high clustering coefficient. The *clustering coefficient* measures the extent to which the neighbours of a given node are also neighbours of each other (see Watts & Strogatz, 1998, for a more quantitative definition). Consider the word *dog* in Figure 1, which has as neighbours the words *dawn*, *dug*, *dig*, *log*, *fog*, *hog* and *bog*, many of which are also neighbours with each other (such as *log–fog*, *hog–bog*, etc.). In a small-world network, a node tends to have more neighbours being neighbours with each other than would be expected by chance, where ‘chance’ is determined by creating a network of similar size but with connections between nodes placed randomly rather than based on the relationships that occur in the system being examined.

Arbesman, Strogatz, and Vitevitch (2010) found similar structural features in phonological networks of Spanish, Mandarin, Hawaiian and Basque. Finding similar network features across these languages was somewhat surprising given the numerous differences among the languages that were sampled in characteristics like the typical length of a word, the

phoneme inventories, etc. and the different ‘families’ from which the languages were sampled. For example, English is from the Germanic branch of Indo-European languages, whereas Spanish is from the Romance branch of Indo-European languages. Mandarin, is not only a Sino-Tibetan language, but it further differs from English, Spanish, Hawaiian and Basque in that it uses tones to convey word meanings (note, tone was not represented in the phonological network, however). Hawaiian is an Austronesian language with a phoneme inventory that is smaller than the inventories found in English, Spanish, Mandarin and Basque. Finally, Basque is a linguistic isolate or not known to be related to any other language.

Observing the same characteristics in the phonological network of a number of different languages suggests that the network of phonological word-forms might be capturing important aspects of the structure of the mental lexicon or of language more generally. One of the fundamental assumptions of network science is that the structure of a network influences the dynamics of that system (Watts & Strogatz, 1998). That is, a certain process might operate very efficiently on a network that is structured in one way. However, in a network with the same number of nodes and the same number of connections—but with those nodes connected in a slightly different way—the same process might now be very inefficient. Given the fundamental assumption that the structure of a network influences the dynamics of that system, several colleagues and I (as described below) began to investigate how the structure among phonologi-

cal word-forms in the mental lexicon might influence various language-related processes.

How lexical structure influences lexical processing

One important measure of network structure is *degree*, which refers to the number of connections that a node has. In the network of phonological word-forms examined by Vitevitch (2008), degree corresponds to the more familiar term from psycholinguistics: *neighbourhood density*. Thus, a word with a dense neighbourhood would be represented in the network as a node (or phonological word-form) that is connected to many other nodes/phonological word-forms (such as the word *cat* in Figure 1), whereas a word with a sparse neighbourhood would be represented in the network as a node that is connected to few other nodes/phonological word-forms (such as the word *dog* in Figure 1). Given the extensive research on the influence of neighbourhood density on a variety of language-related processes (summarized above) we were encouraged to look for other measures of network structure that might also influence language processing in some way.

Our initial exploration began with the *clustering coefficient*, which measures the proportion of phonological neighbours of a word that are also phonological neighbours with each other. This network science measure provided us with a way to quantify one of the observations noted above: some words had phonological neighbours that tended to be neighbours just with that specific word, whereas other words had phonological neighbours that were also neighbours with other neighbours of a word. Thus, the clustering coefficient provided us with a long-sought-after way to measure more precisely the 'internal structure' of a phonological neighbourhood.

The clustering coefficient may appear conceptually similar to phonological neighbourhood density/degree, but it is important to note that they are, by definition, different measures. Furthermore, as shown in Chan and Vitevitch (2010), the clustering coefficient of over 6000 words in the lexicon (with two or more neighbours, which is the minimum number of neighbours required to compute clustering coefficient), was not significantly correlated with the neighbourhood density/degree of those words (see also Vitevitch, Chan, & Roodenrys, 2012).

Chan and Vitevitch (2009) found that words that were similar in neighbourhood density/degree, but varied in clustering coefficient, were responded to differentially in several conventional psycholinguistic tasks. That is, words with low clustering coefficient (the neighbours of a target word tended to be neighbours only with the target word and not with other neighbours) were responded to in a perceptual identification task and a lexical decision task more quickly and accurately than words with high clustering

coefficient (many of the neighbours of a target word were also neighbours with each other).

Furthermore, Chan and Vitevitch (2009) found that computer simulations of widely-accepted models of spoken word recognition were not able to account for the influence of clustering coefficient on spoken word recognition, suggesting that the structure of the lexicon, as measured by the clustering coefficient, may influence spoken word recognition. Although widely-accepted models of spoken word recognition were not able to account for the influence of clustering coefficient on spoken word recognition, Vitevitch, Ercal, and Adagarla (2011) demonstrated with computer simulations that the diffusion of activation across a network model could account for the influence of clustering coefficient on spoken word recognition. Subsequent research has found that the clustering coefficient also influences the production of spoken words (Chan & Vitevitch, 2010), certain aspects of short-term and long-term memory for words (Vitevitch et al., 2012) and the acquisition of novel word-forms (Goldstein & Vitevitch, 2014).

Additional studies have examined how other structures observed in the lexical network—including path-length, mixing pattern, the existence of communities and keyplayers in the network—might influence language-related processes. *Path-length* refers to the number of connections that must be traversed to get from one node/word to another. For example, in the network in Figure 1, the path-length between *dog* and *bog* is one connection, whereas the (shortest) path between the words *dog* and *cat* is four. Vitevitch, Goldstein, and Johnson (in press) analysed responses in a phonological associates task (the participant hears a word and says the first word that comes to mind that 'sounds like' that word) to examine how path length might influence what is perceived when a word is perceived erroneously.

Although participants were not given a precise definition of what it means for two words to 'sound like' each other, Vitevitch et al. found that over 80% of the responses differed from the target word by a single phoneme. For example, if the participant heard the word *dog*, they were likely to respond with *bog* or *dig*. More interesting, over 95% of the responses that differed from the target word by more than a single phoneme had a path that connected the target word to the more distant response, such as being presented with *dog* and responding with *bag* (see Figure 1). The existence of lexical intermediaries between the target word and more distant responses raises some concerns about measures of word-form similarity that ignore such items, such as the Orthographic Levenshtein Distance-20 (*OLD-20*; Yarkoni, Balota, & Yap, 2008) and the Phonological Levenshtein Distance-20 (*PLD-20*; Suárez, Tan, Yap, & Goh, 2011). Computations of *OLD-20/PLD-20* do not consider whether real-word intermediaries exist or not; the measure only considers the number of letter/phoneme changes, respectively, that are

required to turn one word into another. However, the findings of Vitevitch et al. show that distant phonological neighbours tend to be connected to a word via a path of real words, raising questions about the psychological validity of metrics such as *OLD-20* and *PLD-20* that do not take into account the existence of lexical intermediaries.

The *mixing pattern* of a network refers to a general tendency for how nodes in a network connect to each other (i.e. how entities in the system tend to mix together). In a social network, mixing might be defined based on the age of the individuals, resulting in the observation that people in the network tend to have as friends people that are comparable in age.

Mixing, however, can be defined on a variety of other characteristics, including the number of connections that a node has (i.e. degree). If nodes with many connections tend to connect to other nodes that also have many connections—there is an overall positive correlation in the degree of two connected nodes—it is said that the network exhibits assortative mixing by degree. If nodes with many connections tend to connect to other nodes that have few connections—there is an overall negative correlation in the degree of two connected nodes—it is said that the network exhibits disassortative mixing by degree. If there is no correlation in the degree of two connected nodes, then there is no observable mixing pattern.

In their analysis of the network structure of several different languages, Arbesman et al. (2010) found relatively high values of assortative mixing by degree (0.5–0.8 in the languages examined by Arbesman et al., whereas 0.1–0.3 is typically observed in social networks). The observation of assortative mixing by degree in phonological networks is interesting because mathematical simulations of networks with different mixing patterns suggest that the overall pattern of mixing exhibited in a network has implications for the ability of the system to maintain processing in the face of damage to the network. Newman (2002) found that removing nodes with a high-degree in networks with disassortative mixing by degree greatly disrupted the ability to traverse a path from one node to another node in the system (known as network connectivity). In contrast, network connectivity was not disrupted as much when high-degree nodes were removed from a network with assortative mixing by degree. In other words, networks with assortative mixing by degree are able to remain relatively connected in the face of targeted attacks to the system.

In addition to measuring the extent to which the phonological networks exhibited assortative mixing by degree, Arbesman et al. (2010) used the method employed by Newman (2002) to examine the resilience of the English network by targeting for removal either highly connected nodes or randomly selected nodes. What is typically observed in other domains is that the network remains relatively well connected when nodes are attacked at random, but when highly

connected nodes are targeted for removal the network falls apart, thereby disrupting processing in that system (e.g. Albert, Jeong, & Barabási, 2000; Newman, 2002). In contrast, Arbesman et al. observed similar and high levels of resilience in connectivity in the phonological network when either a random attack or an attack targeting highly connected nodes was carried out. The resilience of phonological networks with assortative mixing by degree observed in the computer simulations of Arbesman et al. (2010) made us wonder how network organization might contribute to the ‘resilience’ of language processing more broadly.

Vitevitch, Chan, and Goldstein (2014) reasoned that, if assortative mixing by degree contributes to the ‘resilience’ of language processing, then they should be able to find behavioural evidence for assortative mixing by degree in instances when lexical retrieval failed. Recall that in the phonological network degree corresponds to the psycholinguistic term, phonological neighbourhood density. Therefore, when lexical retrieval fails the neighbourhood density/degree of the word that is erroneously retrieved should be correlated to the neighbourhood density/degree of the word that was correctly produced. Vitevitch et al. analysed a corpus of slips-of-the-ear or speech errors in which the speaker produces an utterance correctly, but the listener ‘mishears’ what is said. They found a significant, positive correlation in the neighbourhood density/degree of the words that were produced and the neighbourhood density/degree of the words that were ‘misheard’ by the listener, indicating that assortative mixing by degree might have behavioural consequences for certain aspects of language processing.

To further examine how assortative mixing by degree might influence language processing, Vitevitch et al. (2014) simulated ‘failed’ lexical retrieval in a computer model of spoken word recognition and also in three psycholinguistics tasks that approximated, in a laboratory setting, failed lexical retrieval. Vitevitch et al. again found behavioural evidence for assortative mixing by degree; that is, a significant, positive correlation in the neighbourhood density/degree of the words that were presented to participants and the neighbourhood density/degree of the words that were given in response. The results of these studies on the network science metric known as assortative mixing by degree further suggest that the way word-forms are organized in the mental lexicon influences certain aspects of language processing.

Together these studies illustrate that network science consists of techniques and measurements that can be used to examine a system at multiple levels. At the *micro-level*, one can examine the characteristics of an individual node in the system and perhaps the nodes immediately connected to that individual in the system. Thus, the studies that examined the

influence of degree and clustering coefficient on lexical processing can be categorized as having examined the influence of the micro-structure of the lexicon on processing. At the *macro-level*, one examines the characteristics of the whole system by calculating the average value of a particular measure or by measuring a characteristic that describes a general tendency in the system. Thus, the studies that examined the influence of mixing (i.e. assortative mixing by degree) and path length on lexical processing can be categorized as having examined the influence of the macro-structure of the lexicon on processing.

In between the micro- and macro-levels lies the *meso-level*, where one can examine characteristics of groups or sub-sets of nodes that might be found in the system. A common technique used to examine the meso-level is *community detection* or attempting to find smaller sub-groups of nodes, called *communities*. Nodes within a given community tend to be more connected to each other than to nodes found in another community. Consider the neighbours of the word *cat* in Figure 1: *cot*, *cut*, *coat*, *kit* and *kite* might form one community, *cad*, *calf*, *cab* and *can* might form another community and *chat*, *that*, *hat*, *rat*, *gnat* and *fat* might form yet another community.

Using a common community detection algorithm, Siew (2013) found 17 communities of various sizes in the phonological network examined in Vitevitch (2008). Siew suggested that the presence of community structure in the lexical network might contribute to the rapidity of word recognition by restricting activation to a relatively small sub-set of lexical candidates (i.e. the words in the community) instead of allowing activation to spread rampantly to the entire lexicon.

Implications of network science for speech and language disorders

The studies reviewed above looked at a variety of network measures and their influence on lexical processing in typically-developing college-age adults. Despite the—at present—limited application of network science to the language sciences (for reviews see: Baronchelli, Ferrer i Cancho, Pastor-Satorras, Chater, & Christiansen, 2013; Cong & Liu, 2014), we believe there is much promise for the application of network science, especially for increasing our understanding of language development and our understanding of and ability to treat language disorders.

One network measure that might have fairly direct application to clinical practice is that of *keyplayers in the network*. Keyplayers are nodes in a network that, when removed, result in the network fracturing into several smaller components (see Borgatti (2006) for the algorithm used to find such nodes, as well as for information about software that will find such nodes in a network). In Figure 1, if the word *bag* (and its connections) were removed from the network, two

smaller components would be obtained: *dog* and its neighbours and *cat* and its neighbours, with no way to get from *dog* to *cat*. Vitevitch and Goldstein (2014) extracted a set of 25 words that held such ‘key’ positions in the larger phonological network (e.g. bring, fish, misty) and a set of 25 foil words (e.g. brief, firm, mystic) that were comparable to the ‘keywords’ on a number of lexical and network characteristics. They found that keywords were responded to more quickly and accurately than the foils in a perceptual identification task, an auditory naming task and an auditory lexical decision task, showing that the position of words in the phonological network plays an important role in the processing of those words.

Given the important role that ‘keywords’ in the network play in processing, one might attempt to introduce keywords at a developmentally appropriate time during the acquisition of a first or second language to accelerate or otherwise facilitate the acquisition of new words. Similarly, in individuals with acquired language disorders, including various types of aphasia, treatments that focus on the re-acquisition or rehabilitation of such keywords could facilitate language recovery. Additional analyses, simulations and empirical investigations are required, however, to verify what is at present optimistic speculation.

Work by Beckage, Smith, and Hills (2011) illustrates more directly how the principles and analytic techniques of network science can be used to increase our understanding of certain language disorders. Note that their analysis was of a lexical network in which the connections between words represented *semantic* similarity, instead of phonological similarity as in most of the studies described here. In their analysis, Beckage et al. (2011) made semantic networks for a group of typically-developing children and for a group of ‘late talking’ children, who had vocabularies (obtained from the Communicative Development Inventories (CDI), Dale & Fenson, 1996) that were smaller than most children at that age (15–36 months).

Beckage et al. (2011) observed that the networks for the typically-developing children had a higher clustering coefficient, a shorter path-length and greater average degree compared to the networks for the late-talking children (specifically for directed links coming ‘into’ the node, known as *in-degree*; see Supplementary Appendix to be found online at <http://informahealthcare.com/doi/abs/10.3109/17549507.2014.987819>—Key Terms for a definition of in-degree). In other words, the extent to which a child’s vocabulary resembled a small-world network was related to the child’s rate of vocabulary development, such that children who developed a vocabulary at the typical pace exhibited networks with small-world characteristics, whereas late-talkers showed the small-world characteristics to a lesser extent. The small-world network structure is known to contribute to rapid search (Kleinberg, 2000). It is perhaps not a coincidence that deviations from this

network structure were observed in the lexical network of late-talking individuals.

Beckage et al. (2011) suggested that the small-world-like structure observed in the semantic lexicon of typically-developing children likely arises from the biases in word acquisition identified (via network analyses) in Hills et al. (2009, 2010). However, late-talking children may instead show a preference for new words that are semantically novel compared to what is already known; they refer to such words as ‘oddballs’. For example, late-talking children may be more likely to acquire the word *telephone* than *dog* after learning the word *cat*, because *telephone* is less semantically similar to the already known word *cat*.²

The work by Beckage et al. (2011) and by Hills et al. (2009, 2010) on the acquisition of semantic information in the mental lexicon nicely illustrates how network analyses can be useful for shedding light on the processes that influence typical, disordered and delayed development (see also Kenett, Wechsler-Kashi, Kenett, Schwartz, Ben-Jacob, & Faust, 2013). In what follows we show how network analyses might also be useful for examining language processes at the other end of the developmental spectrum, namely in adults with either Broca’s or Wernicke’s aphasia.

Analysis of performance in the Philadelphia Naming Test

In this section we report the results of an analysis of data obtained from a database (freely available on-line) of various tests of cognitive function performed by individuals with various types of aphasia and by age-matched controls (Mirman, Strauss, Brecher, Walker, Sobel, Dell, et al., 2010). We examined some well-known linguistic variables in this analysis in order to replicate the results of previous studies of speech production. The replication of previous well-known results in this set of data would bolster our confidence in any novel results we might obtain as we explore how other network science measures that we have not examined previously might influence speech production.

We downloaded from the Moss Aphasia Psycholinguistics Project Database: <http://mrri.org/mappd> (accessed March 2014) the accuracy results of the 175 items in the Philadelphia Naming Test (PNT) for age-matched controls ($n = 20$), individuals with Broca’s Aphasia ($n = 58$) and individuals with Wernicke’s Aphasia ($n = 36$). A binomial multiple regression model was used to predict the odds of naming a picture correctly. Table I lists each variable that we examined, the beta coefficient (in log odds units) and the coefficient in odds units (calculated by taking the exponent of the beta coefficient). All analyses discussed below were significant at $p < 0.001$, indicating that a particular variable influences the likelihood of accurately naming a picture when all other variables are controlled.

Not surprising, the odds of accurately naming a picture in the PNT was influenced by the type of

Table I. Analysis of performance in the Philadelphia Naming Test.

Variable	Beta coefficient (log odds unit)	Coefficient (odds unit)
Intercept	4.382	79.99
Wernicke’s Aphasia	−4.183	0.015
Broca’s Aphasia	−3.727	0.024
Months Post-Onset	0.003	1.003
Word length	−0.195	0.823
Word Frequency	0.313	1.368
Degree/neighbourhood density	0.437	1.548
Closeness Centrality	−5.144	0.006
Component	0.720	2.054

Wernicke’s Aphasia is in comparison to age-matched controls. *Broca’s Aphasia* is in comparison to age-matched controls. *Word length* was centred at one phoneme. *Word frequency* counts came from Kučera and Francis (1967; we added 1 to each value and then performed a $\log_{10}$ transformation). For *degree/neighbourhood density*, we added 1 to each value and performed a $\log_{10}$ transformation, because, as observed in Vitevitch (2008), *degree/neighbourhood density* is not normally distributed. For the *Component*, words in the giant component were coded as 0, words outside of the giant component were coded as 1.

individual (i.e. Wernicke’s Aphasia, Broca’s Aphasia or age-matched control). The odds of correctly naming a picture for those with Wernicke’s aphasia was 0.015-times the odds for healthy controls and the odds of correctly naming a picture for those with Broca’s aphasia was 0.024-times the odds for healthy controls. In other words, healthy controls named roughly 67 pictures correctly for every one picture named correctly by an individual with Wernicke’s aphasia and healthy controls named roughly 42 pictures correctly for every one picture named correctly by an individual with Broca’s aphasia.

It is also not surprising that over time as individuals recover from the incident that led to either Wernicke’s or Broca’s Aphasia, that performance on cognitive tasks, like the PNT, improve (at least somewhat). In the present analysis we found that, for every month post-onset, the odds of correctly naming a picture was 1.003 or an increase in accuracy of 0.3% each month.

It has long been known that the length of a word (measured in various ways, including the number of syllables in the word, the number of letters in the word or the number of phonemes in the word) influences the likelihood of accurately producing a word (e.g. Hodgson & Ellis, 1998; Meyer, Roelofs, & Lev-elt, 2003; Santiago, MacKay, Palma, & Rho, 2000). For example, Bricker, Schuell, and Jenkins (1964) found that individuals with aphasia (type was not specified) made more errors spelling longer words than shorter words. In the present analysis we found that, as the length of the word (measured by the number of phonemes) increased, the odds of correctly naming a picture was 0.823 or a decrease in accuracy of roughly 18%, replicating the well-attested influence of word-length on speech production.

It is also well-known that the frequency with which a word occurs in the language influences how quickly

and accurately it is perceived or produced (e.g. Howes, 1957). Specifically, words that are common in the language tend to be produced and perceived more quickly and accurately than words that are used less often. In our analysis of data from the PNT, we found that, as the frequency of the word increases, the odds of correctly naming a picture was 1.368 or an increase of roughly 37%, replicating another well-attested influence on speech production.

Another effect that was replicated in the present set of data was the influence of neighbourhood density on speech production (e.g. Goldrick & Rapp, 2007; Harley & Bown, 1998; Vitevitch, 1997, 2002; Vitevitch & Sommers, 2003). Recall that neighbourhood density refers to the number of words that sound like a target word. In a network representation of phonological word-forms in the lexicon, the term degree refers to the number of connections incident to a node; in other words, how many words sound similar to that word. Because the psycholinguistic term neighbourhood density and the network science term degree both refer to the same concept we will use the combined term degree/neighbourhood density. As in previous studies of speech production we found in the present set of data that as the degree/neighbourhood density of a word increased, the odds of correctly naming a picture was 1.548 or an increase of roughly 55%.

Replicating several well-known and previously observed effects in the present set of data bolsters our confidence that any new effects of network variables that have not been extensively explored in psycholinguistics that we may presently observe are not spurious. One network variable we wish to explore in the present data is *closeness centrality*, which measures the distance from one node to all other nodes in the network (following the shortest path between any two nodes being considered). A node might, therefore, be considered ‘important’ if it is relatively close to all other nodes in the system. More precisely, closeness centrality is defined as:

$$C_v = \frac{1}{\sum_{u \in V} d(v, u)} \quad (1)$$

where $d(v, u)$ refers to the shortest distance (i.e. shortest path) between nodes v and u , $\sum$ refers to the sum of the path lengths from node v to all other nodes in the network.

As indicated in equation (1), closeness centrality is typically reported as the inverse of the distance from a node to every other node in the network. Therefore, a node that has high closeness centrality, a value close to 1, tends to be close to the other nodes in the network (meaning that one can get from that node to other nodes in the network by traversing relatively few connections). Conversely, a node that has low closeness centrality, a value close to 0, tends to be far away from the other nodes in the network (meaning that one must traverse many connections to get from that node to the other nodes in the network).

Iyengar, Madhavan, Zweig, and Natarajan (2012) demonstrated the influence of closeness centrality on language-related processing using a game called word-morph, in which participants were given a word and asked to form a disparate word by changing one letter at a time. For example, asked to ‘morph’ the word *bay* into the word *egg*, participants might have changed *bay* into *bad-bid-aid-add-ado-ago-ego* and finally into *egg*. (See Figure 1 for a way to morph the word *dog* into the word *cat*.) Once participants in this task identified certain ‘landmark’ words in the network of orthographically similar words—words that had high closeness centrality, like the word *aid* in the example above—the task of navigating from one word to another became trivial, enabling the participants to solve subsequent word-morph puzzles very quickly. The time it took to find a solution dropped from 10–18 minutes in the first 10 games, to ~ 2 minutes after playing 15 games, to ~ 30 seconds after playing 28 games, because participants would ‘morph’ the start-word (e.g. *bay*) into one of the landmark words that were high in closeness centrality (e.g. *aid*), then morph the landmark-word into the desired end-word (e.g. *egg*). Although this task is a contrived word-game rather than a conventional psycholinguistic task that assesses on-line lexical processing, the results of Iyengar et al. (2012) nevertheless illustrate how the tools of network science can be used to provide insights about linguistic representations and how the organization of those representations might influence processing.

In the present analysis of picture naming data from the PNT we observed that, as closeness centrality increased, the odds of correctly naming a picture was 0.006. That is, words that are close to all of the other words in the lexicon are named less accurately than words that are far away from all of the other words in the lexicon. Although this result may appear to contradict the findings reported by Iyengar et al. (2012), that is not the case. Recall that Iyengar et al. found that words with high closeness centrality proved to be very useful in the word-morph game. The demands of this off-line, language game are quite different from the demands of a confrontation naming-task (a.k.a. picture- or object-naming), such as that found in the PNT. In the word-morph game, being close to other words in the lexicon can help one quickly transform one word into another, leading to successful performance in the game. However, when the task is to retrieve from the lexicon a specific word, as in a picture-naming task, being close to all of the other words in the lexicon could lead to competition among candidate word-forms or to activation being diverted away from the target word-form to all of the other nearby word-forms, thereby decreasing the likelihood of successfully retrieving target words that have high closeness centrality. Given the different demands of the word-morph game and the picture-naming task, the present results do not necessarily contradict the findings from Iyengar et al. (2012).

The present findings regarding closeness centrality may also appear to contradict the well-attested findings regarding degree/neighbourhood density that were replicated in the present analysis: English words with many phonological neighbours are named more accurately than English words with few phonological neighbours (e.g. Goldrick & Rapp, 2007; Harley & Bown, 1998; Vitevitch, 1997, 2002; Vitevitch & Sommers, 2003). Computer simulations (e.g. Dell & Gordon, 2003; Gordon & Dell, 2001) further demonstrated that ‘phonological neighbours’ or words that differed from the target word by one phoneme played a facilitative role in retrieving a word-form from the lexicon, such that words with more neighbours were retrieved more accurately than words with few neighbours.

However, recent simulations looking at what might be described as ‘distant neighbours’ suggests that items that are less similar to a target word can exert a different influence on lexical retrieval than ‘near’ neighbours (Mirman & Magnuson, 2008). Therefore, it is not unreasonable for degree/neighbourhood density, which measures ‘near’ neighbours, to exert one type of influence on lexical retrieval and closeness centrality, which measures ‘distant’ neighbours, to exert a different kind of influence on lexical retrieval.

The replication of several well-studied effects in the present analysis boosts our confidence regarding the novel influence observed for closeness centrality on speech production. The work of Iyengar et al. (2012) as well as the present findings regarding closeness centrality point to a network characteristic that may warrant further investigation by language scientists and speech-language pathologists.

Another network variable that may warrant further investigation by language scientists is where in the network a word resides. Recall that Vitevitch (2008) found that the phonological network of English had a large group of nodes that were highly connected to each other (i.e. the giant component), as well as many smaller groups of words that were connected to each other, but not to the giant component (i.e. smaller components or ‘lexical islands’) and many words that did not have any phonological neighbours at all (i.e. isolates, ‘lexical hermits’). Further recall that Arbesman et al. (2010) found, across a handful of languages, that the giant component contained from 34% (English) to 66% (Mandarin) of the words in the lexicon, leaving a large proportion of words in either the lexical islands or as hermits. In most systems examined in network science, ~ 90% of the nodes are found in the giant component. The smaller proportion of nodes found in the giant component in phonological networks compared to other types of networks points to a characteristic that may warrant further investigation.

In the present analysis we examined whether a word located in the giant component (coded as 0 in the present analysis) might be processed differently than a word found outside of the giant component (coded as 1 in the present analysis; we did not dis-

tinguish between words in smaller components and isolates in this analysis). In the picture naming data from the PNT we observed that the odds of correctly naming a picture was 2.05-times greater for words not in the giant component. That is, words found outside of the giant component were named more accurately than words in the giant component (see also Siew & Vitevitch, 2014).

The present finding regarding more accurate naming of words found outside of the giant component is intriguing, because 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)—and the well-known influences of word frequency and of word-length on lexical retrieval (described above and replicated in this set of data), one might expect the words in the giant component to be named more accurately than words outside of the giant component. Recall, however, that in the binomial multiple regression technique used in the present analysis all other variables are controlled. Finding an influence of location in the phonological network when those other variables (e.g. word length, word frequency) are controlled points to another network characteristic that may warrant further investigation by language scientists (see Siew & Vitevitch, 2014 for a possible account of this finding).

Conclusion

In the present review we briefly summarized a large body of research looking at individual lexical characteristics, focusing specifically on *phonotactic probability* and *neighbourhood density*. We proceeded to introduce the emerging field of *network science* and illustrated how the computational tools of network science could be used to examine individual lexical characteristics (i.e. the *micro-level*), overall characteristics of a system (i.e. the *macro-level*), as well as characteristics of sub-sets of items (i.e. the *meso-level*) in the context of the phonological lexicon. The results of the studies that we summarized showed that more than just individual lexical characteristics influence processing. Rather, the structure of the lexicon at the micro-, meso- and macro-level influences various aspects of lexical processing.

Furthermore, we analysed data from an on-line database of the Philadelphia Naming Test to show the utility of network science for increasing our understanding of language processing at other points in the lifespan. A number of well-known findings were replicated and several novel influences of network science measures on lexical processing were also observed, pointing toward new avenues for research.

Without explicitly appealing to the overall structure of the lexicon—something that previous models of lexical processing have not done—it is difficult to see how current models of lexical retrieval could account for the novel findings that we observed.

The present manuscript focused primarily on phonological networks. By no means, however, is phonology the only aspect of the mental lexicon or of language more broadly that can be examined using network science; a bibliography of research using networks to study semantics, orthography and many other aspects of language can be found at: http://www.cs.upc.edu/~rferrericancho/linguistic_and_cognitive_networks.html.

A developing area of research in network science is on multiplex networks, in which two or more relationships between entities are represented in the network. In the case of a multiplex network of people links between people might represent different relationships, such as friendship ties, advice seeking/giving and people to whom you would loan \$100. This developing area of research holds much promise for the study of language to examine how semantic, phonological and orthographic factors contribute to and interact during language processing (see Mirman and Magnuson (2008) for a different approach to representing semantic and phonological information).

Although we are enthusiastic about the potential insights that network science can provide the language sciences, we recognize that networks are not suitable for all research questions. Researchers and clinicians who desire to use the theoretical framework and analytic tools of network science should think carefully about how well entities and relationships among those entities in a given domain map onto nodes and connections in a network representation (Butts, 2009; Valente, 2012). Similarly, not all measures employed in network science are appropriate for all domains (Borgatti, 2005). We urge more language scientists to consider how network science might lead to new questions and novel insights, but we also urge judicious and appropriate use of these techniques. See the Supplementary Appendix to be found online at <http://informa-healthcare.com/doi/abs/10.3109/17549507.2014.987819> for information on software that can be used to visualize a network and to compute the variables that are discussed.

Notes

1. The words used to make the network came from the Webster's Seventh Collegiate Dictionary (1967), which also forms the basis of several widely used databases in psycholinguistics (e.g. Storkel & Hoover, 2010a; Vitevitch & Luce, 2004). Although estimates of the size of the vocabulary of the average

person vary widely, this sample is believed to be sufficiently representative.

2. Interestingly, such an 'oddball' strategy may be advantageous for *triggering* the acquisition of novel phonological word-forms in typically-developing individuals (Storkel, 2011). That is, novel words that are phonologically less similar to other known words can be more easily identified as a novel word to which resources should be allocated in order to acquire it. Novel words that are phonologically similar to many known words may be erroneously identified as an already known word, thereby delaying the acquisition of that novel word (in typically-developing individuals)

Declaration of interest: The authors report no conflicts of interest. The authors alone are responsible for the content and writing of the paper.

This work was supported in part by NIDCD grant T32 – DC00052 Training Researchers in Language Impairments.

References

- Albert, R., Jeong, H., & Barabási, A. L. (2000). Error and attack tolerance of complex networks. *Nature*, 406, 378–382.
- Anderson, J. D., & Byrd, C. T. (2008). Phonotactic probability effects in children who stutter. *Journal of Speech, Language, and Hearing Research*, 51, 851–866.
- Apel, K., Wolter, J. A., & Masterson, J. J. (2006). Effects of Phonotactic and Orthotactic Probabilities During Fast Mapping on 5-Year-Olds' Learning to Spell. *Developmental Neuropsychology*, 29, 21–42.
- Arbesman, S., Strogatz, S. H., & Vitevitch, M. S. (2010). The Structure of Phonological Networks Across Multiple Languages. *International Journal of Bifurcation and Chaos*, 20, 679–685.
- Arnold, H. S., Conture, E. G., & Ohde, R. N. (2005). Phonological neighborhood density in the picture naming of young children who stutter: preliminary study. *Journal of Fluency Disorders*, 30, 125–148.
- Barabási, A. L. (2009). Scale-free networks: A decade and beyond. *Science*, 325, 412–413.
- Baronchelli, A., Ferrer i Cancho, R., Pastor-Satorras, R., Chater, N., & Christiansen, M. H. (2013). Networks in cognitive science. *Trends in Cognitive Science*, 17, 348–360.
- Bastian, M., Heymann, S., & Jacomy, M. (2009). Gephi: An open source software for exploring and manipulating networks. In Proceedings of the 3rd International AAAI Conference on Weblogs and Social Media; San Jose, CA, pp. 361–362.
- Batagelj, V., & Mrvar, A. (1988). Pajek – Program for large network analysis. *Connections*, 21, 47–57.
- Beckage, N., Smith, L., & Hills, T. (2011). Small Worlds and Semantic Network Growth in Typical and Late Talkers. *PLoS ONE*, 6, e19348.
- Bonte, M. L., Poelmans, H., & Blomert, L. (2007). Deviant neurophysiological responses to phonological regularities in speech in dyslexic children. *Neuropsychologia*, 45, 1427–1437.
- Borgatti, S. P. (2005). Centrality and network flow. *Social Networks*, 27, 55–71.
- Borgatti, S. P. (2006). Identifying sets of key players in a network. *Computational, Mathematical and Organizational Theory*, 12, 21–34.

- Bricker, A. L., Schuell, H., & Jenkins, J. J. (1964). Effect of word frequency and word length on aphasic spelling errors. *Journal of Speech & Hearing Research*, 7, 183–192.
- Butts, C. T. (2009). Revisiting the foundations of network analysis. *Science*, 325, 414–416.
- 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 & Performance*, 35, 1934–1949.
- Chan, K. Y., & Vitevitch, M. S. (2010). Network structure influences speech production. *Cognitive Science*, 34, 685–697.
- Chan, K. Y., & Vitevitch, M. S. (in press). The influence of neighborhood density on the recognition of Spanish-accented words. *Journal of Experimental Psychology: Human Perception and Performance*.
- Cong, J., & Liu, H. (2014). Approaching human language with complex networks. *Physics of Life Reviews*, 11, 598–618.
- Cutler, A. (1981). Making up materials is a confounded nuisance, or: Will we be able to run any psycholinguistic experiments at all in 1990? *Cognition*, 10, 65–70.
- Dale, P. S., & Fenson, L. (1996). Lexical development norms for young children. *Behavior Research Methods, Instruments, & Computers*, 28, 125–127.
- Dell, G. S., & Gordon, J. K. (2003). Neighbors in the lexicon: Friends or foes? In N. O. Schiller, & A. S. Meyer (Eds.), *Phonetics and phonology in language comprehension and production: Differences and similarities*. New York: Mouton.
- Gathercole, S. E., Frankish, C. R., Pickering, S. J., & Peaker, S. (1999). Phonotactic influences on short-term memory. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 25, 84–95.
- Gierut, J. A., Morrisette, M. L., & Champion, A. H. (1999). Lexical constraints in phonological acquisition. *Journal of Child Language*, 26, 261–294.
- Goldrick, M., & Larson, M. (2008). Phonotactic probability influences speech production. *Cognition*, 107, 1155–1164.
- Goldrick, M., & Rapp, B. (2007). Lexical and post-lexical phonological representations in spoken production. *Cognition*, 102, 219–260.
- Goldstein, R., & Vitevitch, M. S. (2014). The influence of clustering coefficient on word-learning: how groups of similar sounding words facilitate acquisition. *Frontiers in Language Sciences*, 5, 1307.
- Gordon, J. K. (2002). Phonological neighborhood effects in aphasic speech errors: Spontaneous and structured contexts. *Brain & Language*, 82, 113–145.
- Gordon, J. K., & Dell, G. S. (2001). Phonological neighborhood effects: Evidence from aphasia and connectionist models. *Brain and Language*, 79, 21–23.
- Gray, S., Brinkley, S., & Svetina, D. (2012). Word learning by preschoolers with SLI: Effect of phonotactic probability and object familiarity. *Journal of Speech, Language, and Hearing Research*, 55, 1289–1300.
- Greenberg, J. H., & Jenkins, J. J. (1964). Studies in the psychological correlates of the sounds system of American English. *Word*, 20, 157–177.
- Harley, T. A., & Bown, H. E. (1998). What causes a tip-of-the-tongue state? Evidence for lexical neighbourhood effects in speech production. *British Journal of Psychology*, 89, 151–174.
- Hills, T. T., Maouene, J., Riordan, B., & Smith, L. B. (2010). The Associative Structure of Language: Contextual Diversity in Early Word Learning. *Journal of Memory & Language*, 63, 259–273.
- Hills, T. T., Maouene, M., Maouene, J., Sheya, A., & Smith, L. B. (2009). Longitudinal analysis of early semantic networks: Preferential attachment or preferential acquisition? *Psychological Science*, 20, 729–739.
- Hodgson, C., & Ellis, A. W. (1998). Last in, first to go: Age of acquisition and naming in the elderly. *Brain and Language*, 64, 146–163.
- Howes, D. (1957). On the relation between the intelligibility and frequency of occurrence of English words. *Journal of the Acoustical Society of America*, 29, 296–305.
- Hunter, C. R. (2013). Early effects of neighborhood density and phonotactic probability of spoken words on event-related potentials. *Brain and Language*, 127, 463–474.
- Imai, S., Walley, A. C., & Flege, J. E. (2005). Lexical frequency and neighborhood density effects on the recognition of native and Spanish-accented words by native English and Spanish listeners. *Journal of the Acoustical Society of America*, 117, 896–907.
- Iyengar, S. R. S., Madhavan, C. E. V., Zweig, K. A., & Natarajan, A. (2012). Understanding human navigation using network analysis. *Topics in Cognitive Science*, 4, 121–134.
- Kaiser, A. R., Kirk, K. I., Lachs, L., & Pisoni, D. B. (2003). Talker and lexical effects on audiovisual word recognition by adults with cochlear implants. *Journal of Speech-Language-Hearing Research*, 46, 390–404.
- Kenett, Y. N., Wechsler-Kashi, D., Kenett, D. Y., Schwartz, R. G., Ben-Jacob, E., & Faust, M. (2013). Semantic organization in children with cochlear implants: Computational analysis of verbal fluency. *Frontiers in Psychology*, 4, 543.
- Kleinberg, J. M. (2000). Navigation in a small world. *Nature*, 406, 845.
- Kučera, H., & Francis, W. N. (1967). *Computational Analysis of Present Day American English*. Providence, RI: Brown University Press.
- Lallini, N., & Miller, N. (2011). Do phonological neighbourhood density and phonotactic probability influence speech output accuracy in acquired speech impairment? *Aphasiology*, 25, 176–190.
- 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.
- Leonard, L. B., Davis, J., Deevy, P. (2007). Phonotactic probability and past tense use by children with specific language impairment and their typically developing peers. *Clinical Linguistics & Phonetics*, 21, 747–758.
- Luce, P. A., & Pisoni, D. B. (1998). Recognizing spoken words: The neighborhood activation model. *Ear and Hearing*, 19, 1–36.
- MacRoy-Higgins, M., Schwartz, R. G., Shafer, V. L., & Marton, K. (2013). Influence of phonotactic probability/neighbourhood density on lexical learning in late talkers. *International Journal of Language & Communication Disorders*, 48, 188–199.
- Mattys, S. L., Jusczyk, P. W., Luce, P. A., & Morgan, J. L. (1999). Phonotactic and prosodic effects on word segmentation in infants. *Cognitive Psychology*, 38, 465–494.
- Messer, M. H., Leseman, P. P. M., Boom, J., & Mayo, A. Y. (2010). Phonotactic probability effect in nonword recall and its relationship with vocabulary in monolingual and bilingual preschoolers. *Journal of Experimental Child Psychology*, 105, 306–323.
- Meyer, A. S., Roelofs, A., & Levelt, W. J. M. (2003). Word length effects in object naming: The role of a response criterion. *Journal of Memory and Language*, 48, 131–147.
- Mirman, D., & Magnuson, J. S. (2008). Attractor dynamics and semantic neighborhood density: Processing is slowed by near neighbors and speeded by distant neighbors. *Journal of Experimental Psychology: Learning, Memory, & Cognition*, 34, 65–79.
- Mirman, D., Strauss, T. J., Brecher, A., Walker, G. M., Sobel, P., Dell, G. S., et al. (2010). A large, searchable, web-based database of aphasic performance on picture naming and other tests of cognitive function. *Cognitive Neuropsychology*, 27, 495–504.
- Montoya, J. M., & Solé, R. V. (2002). Small world patterns in food webs. *Journal of Theoretical Biology*, 214, 405–412.
- Munson, B., & Solomon, N. P. (2004). The Effect of Phonological Neighborhood Density on Vowel Articulation. *Journal of Speech, Language, and Hearing Research*, 47, 1048–1058.
- Newman, M. E. J. (2002). Assortative mixing in networks. *Physical Review Letters*, 89, 208701.
- Newman, M. E. J. (2010). *Networks: An Introduction*. Cambridge, MA: Oxford University Press.
- Noordenbos, M. W., Segers, E., Mitterer, H., Serniclaes, W., & Verhoeven, L. (2013). Deviant neural processing of phonotactic probabilities in adults with dyslexia. *NeuroReport: For Rapid Communication of Neuroscience Research*, 24, 746–750.

- Pylkkänen, L., Stringfellow, A., & Marantz, A. (2002). Neuromagnetic evidence for the timing of lexical activation: An MEG component sensitive to phonotactic probability but not to neighborhood density. *Brain and Language*, 81, 666–678.
- Santiago, J., MacKay, D. G., Palma, A., & Rho, C. (2000). Sequential activation processes in producing words and syllables: Evidence from picture naming. *Language and Cognitive Processes*, 15, 1–44.
- Siew, C. S. Q. (2013). Community structure in the phonological network. *Frontiers in Psychology*, 4, 553.
- Siew, C. S. Q., & Vitevitch, M. S. (2014). *Spoken word recognition and serial recall of words from components in the phonological network*. Manuscript submitted for publication.
- Sommers, M. S. (1996). The structural organization of the mental lexicon and its contributions to age-related declines in spoken word recognition. *Psychology and Aging*, 11, 333–341.
- Sporns, O. (2010). *Networks of the brain*. Cambridge, MA: MIT Press.
- Stamer, M. K., & Vitevitch, M. S. (2012). Phonological similarity influences word learning in adults learning Spanish as a foreign language. *Bilingualism: Language & Cognition*, 15, 490–502.
- Steyvers, M., & Tenenbaum, J. (2005). The Large-Scale Structure of Semantic Networks: Statistical Analyses and a Model of Semantic Growth. *Cognitive Science*, 29, 41–78.
- Storkel, H. L. (2004). Do children acquire dense neighborhoods? An investigation of similarity neighborhoods in lexical acquisition. *Applied Psycholinguistics*, 25, 201–221.
- Storkel, H. L. (2011). Differentiating word learning processes may yield new insights. *Journal of Child Language*, 38, 51–55.
- Storkel, H. L., & Hoover, J. R. (2010a). An on-line calculator to compute phonotactic probability and neighborhood density based on child corpora of spoken American English. *Behavior Research Methods*, 42, 497–506.
- Storkel, H. L., & Hoover, J. R. (2010b). Word learning by children with phonological delays: Differentiating effects of phonotactic probability and neighborhood density. *Journal of Communication Disorders*, 43, 105–119.
- Storkel, H. L., & Hoover, J. R. (2011). The influence of part-word phonotactic probability/neighborhood density on word learning by preschool children varying in expressive vocabulary. *Journal of Child Language*, 38, 628–643.
- Storkel, H. L., & Lee, S. Y. (2011). The independent effects of phonotactic probability and neighborhood density on lexical acquisition by preschool children. *Language and Cognitive Processes*, 26, 191–211.
- Storkel, H. L., Armbruster, J., & Hogan, T. P. (2006). Differentiating phonotactic probability and neighborhood density in adult word learning. *Journal of Speech, Language, and Hearing Research*, 49, 1175–1192.
- Suárez, L., Tan, S. H., Yap, M. J., & Goh, W. D. (2011). Observing neighborhood effects without neighbors. *Psychonomic Bulletin and Review*, 18, 605–611.
- Valente, T. W. (2012). Network interventions. *Science*, 337, 49–53.
- 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 Hearing Research*, 51, 408–422.
- Vitevitch, M. S., & Goldstein, R. (2014). Keywords in the mental lexicon. *Journal of Memory & Language*, 73, 131–147.
- Vitevitch, M. S., & Luce, P. A. (1998). When words compete: Levels of processing in spoken word perception. *Psychological Science*, 9, 325–329.
- Vitevitch, M., & Luce, P. A. (1999). Probabilistic phonotactics and neighborhood activation in spoken word recognition. *Journal of Memory & Language*, 40, 374–408.
- Vitevitch, M. S., & Luce, P. A. (2004). A web-based interface to calculate phonotactic probability for words and nonwords in English. *Behavior Research Methods, Instruments, and Computers*, 36, 481–487.
- Vitevitch, M. S., & Luce, P. A. (2005). Increases in phonotactic probability facilitate spoken nonword repetition. *Journal of Memory & Language*, 52, 193–204.
- Vitevitch, M. S., & Luce, P. A. (2016). *Phonological Neighborhood Effects in Spoken Word Perception and Production*. Manuscript submitted for publication.
- Vitevitch, M. S., & Rodríguez, E. (2005). Neighborhood density effects in spoken word recognition in Spanish. *Journal of Multilingual Communication Disorders*, 3, 64–73.
- Vitevitch, M. S., & Sommers, M. S. (2003). The facilitative influence of phonological similarity and neighborhood frequency in speech production in younger and older adults. *Memory & Cognition*, 31, 491–504.
- Vitevitch, M. S., & Stamer, M. K. (2006). The curious case of competition in Spanish speech production. *Language & Cognitive Processes*, 21, 760–770.
- Vitevitch, M. S., & Storkel, H. L. (2013). Examining the acquisition of phonological word forms with computational experiments. *Language & Speech*, 56, 491–527.
- Vitevitch, M. S., Armbruster, J., & Chu, S. (2004). Sublexical and lexical representations in speech production: Effects of phonotactic probability and onset density. *Journal of Experimental Psychology: Learning, Memory and Cognition*, 30, 514–529.
- Vitevitch, M. S., Chan, K. Y., & Goldstein, R. (2014). Insights into failed lexical retrieval from network science. *Cognitive Psychology*, 68, 1–32.
- Vitevitch, M. S., Chan, K. Y., & Roodenrys, S. (2012). Complex network structure influences processing in long-term and short-term memory. *Journal of Memory & Language*, 67, 30–44.
- Vitevitch, M. S., Ercal, G., & Adagarla, B. (2011). Simulating retrieval from a highly clustered network: Implications for spoken word recognition. *Frontiers in Language Sciences*, 2, 369.
- Vitevitch, M. S., Goldstein, R., & Johnson, E. (in press). In A. Mehler, P. Blanchard, B. Job, & S. Banish (Eds.), *Towards a Theoretical Framework for Analyzing Complex Linguistic Networks*. Springer (Understanding Complex Systems series).
- Vitevitch, M. S., Luce, P. A., Pisoni, D. B., & Auer, E. T. (1999). Phonotactics, neighborhood activation and lexical access for spoken words. *Brain and Language*, 68, 306–311.
- Vitevitch, M. S., Pisoni, D. B., Kirk, K. I., Hay-McCutcheon, M., & Yount, S. L. (2002). Effects of phonotactic probabilities on the processing of spoken words and nonwords by postlingually deafened adults with cochlear implants. *Volta Review*, 102, 283–302.
- Watts, D. J., & Strogatz, S. H. (1998). Collective dynamics of ‘small-world’ networks. *Nature*, 393, 409–410.
- Webster’s Seventh Collegiate Dictionary. (1967). Los Angeles: Library Reproduction Service.
- Yarkoni, T., Balota, D., & Yap, M. J. (2008). Moving beyond Coltheart’s N: A new measure of orthographic similarity. *Psychonomic Bulletin & Review*, 15, 971–979.
- Yates, M., Locker, L., Jr., & Simpson, G. B. (2004). The influence of phonological neighborhood on visual word perception. *Psychonomic Bulletin and Review*, 11, 452–457.
- Zamuner, T. S., Gerken, L., & Hammond, M. (2004). Phonotactic probabilities in young children’s speech production. *Journal of Child Language*, 31, 515–536.
- Zipf, G. K. (1935). *The Psycho-Biology of Language*. Boston, MA: Houghton-Mifflin.

Supplementary material available online

Supplementary Appendix.