

Modelling the Bilingual Lexicon as a Multiplex Phonological Network

Eva Maria Luef

Institute of English and American Studies, University of Hamburg
Department of English and ELT Methodology, Charles University

Phonological word form networks of the mental lexicon are of psycholinguistic relevance, offering insights into the efficiency of lexical access. While much research has concentrated on first languages, there is growing evidence suggesting that phonological networks of second languages are equally significant for lexical processes. Bilingual language processing is proposed to involve the integration of first and second languages, with lexical activation spreading between similar word forms in both languages. Multiplex networks provide a framework to combine different phonological networks, allowing for the analysis of the integrated lexical system's behaviour during lexical processing. In the context of the present study, which focusses on German learners of English as a second language, a multiplex network analysis was constructed to model the interactive complexity of the bilingual mental lexicon. The study tested cross-linguistic effects in a word recognition task using English stimuli. Results revealed that during lexical processing in their second language English, German speakers also activate phonological neighbours from German. In addition, the bilinguals are attuned to the interconnectedness (i.e., clustering) of the German and English neighbours with one another in the phonological neighbourhood of the English target words. These findings can contribute to the ongoing debate on the degree of integration in the bilingual mental lexicon and shed light on the role that phonological networks can play in modelling bilingual lexical processing.

Public Significance Statement

The present study investigated the question of how much influence a first language exerts on lexical processing of a second language. Results showed that German words that are phonologically similar to English words delay word recognition in German learners of English. These findings indicate that different (but closely related) languages are cognitively stored and processed together in bilingual language users.

Keywords: multiplex network, German, English, phonological network, word recognition

Supplemental materials: <https://doi.org/10.1037/cep0000351.supp>

Cross-language activation in speech processing refers to the phenomenon in which multiple languages known by a language user can interfere and activate one another during lexical retrieval and production (Frances et al., 2021; Hoversten & Traxler, 2020; Marian et al., 2008; Sadat et al., 2016; Shook & Marian, 2013). When exposed to stimuli in one language, multilingual speakers also activate similar word forms in their other language(s) (Brysbaert et al., 1999; Dijkstra et al., 1999; Van Wijnendaele & Brysbaert, 2002). Studies have demonstrated this type of cross-linguistic interference (or “bilingual processing disadvantage”) between known languages in a bilingual speaker in auditory word recognition (e.g., Blumenfeld & Marian, 2007; Canseco-Gonzalez et al., 2010;

Ju & Luce, 2004; Marian & Spivey, 2003a, 2003b; Shook & Marian, 2012; Weber & Cutler, 2004). This has implications for the structure of the bilingual mental lexicon as well as the mechanisms involved in lexical processing in speakers of more than one language (Shook & Marian, 2013).

Theoretical models of spoken word recognition are centred on the idea that similar word forms spread lexical coactivation among one another (Luce & Pisoni, 1998; Marslen-Wilson & Zwitserlood, 1989; McClelland & Elman, 1986). Phonologically similar word forms cluster together in phonological neighbourhoods in a language and vie for lexical activation during word recognition (Vitevitch & Luce, 1998; Yates & Dickinson, 2023). Words with greater phonological overlap with numerous others in a language cause delays in lexical recognition. The competition effect arising from multiple neighbours must be resolved before successfully recognising a word (Norris et al., 1995). These phonological neighbourhood effects are well-established in monolingual studies (e.g., Chan & Vitevitch, 2009; Chen & Mirman, 2012; Vitevitch & Luce, 2016; Ziegler et al., 2003).

Cross-linguistic phonological neighbourhoods include word forms of various languages, and the dynamics of lexical competition seem to be similar as in the monolingual neighbourhoods, with neighbours in competition with one another (Marian et al., 2008,

Eva Maria Luef  <https://orcid.org/0000-0002-2362-2422>

The data underlying this study can be accessed at <https://osf.io/yzhxj/>. The author thanks Le Vy Dao, Marcel Frömberg, and Alice Eddyshaw for their assistance with German data processing.

Correspondence concerning this article should be addressed to Eva Maria Luef, Institute of English and American Studies, University of Hamburg, Von-Melle-Park 6, 20146 Hamburg, Germany. Email: eva.luef@uni-hamburg.de

2012; Vitevitch, 2012). Substantial evidence supports the idea that bilingual lexical access is nonselective, implying that language selection mechanisms do not restrict activation to representations belonging to the target language. Instead, languages are activated in parallel in a speaker's mind (Bartolotti & Marian, 2012; Costa et al., 2017; Costa & Santesteban, 2004; de Groot et al., 2000; Dijkstra et al., 1999; Jared & Kroll, 2001; Lemhöfer et al., 2004; Marian & Spivey, 2003a, 2003b; van Hell & Dijkstra, 2002). Lexical competition is suggested to be more pronounced in a second language (L2) than in a first language (L1) due to frequent spurious lexical competition, including coactivation of lexical competitors based on erroneous phonological perception (Hanulíková & Weber, 2012). For instance, for German and Dutch learners of English, the distinction between the vowels in the English words “pan” and “pen” might be opaque, resulting in activation spreading based on accented perception. Despite the generally smaller second language vocabulary, L2 listening may involve stronger lexical competition among words (Weber & Cutler, 2004).

Theoretical perspectives on bilingual lexical processing differ in terms of the dynamics of cross-linguistic activation spreading (Costa et al., 2017; see Schwartz & Kroll, 2007, for an overview). An influential model of bilingual word recognition is the BLINCS model (or “Bilingual Language Interaction Network for Comprehension of Speech”). It includes phonological, phonolexical, ortholexical, and semantic levels of representation and assumes interactive activation flow between the layers (Shook & Marian, 2013). The model associates similar words both within and across languages and forms clusters for each language and an integrated cluster where similar word forms of both languages are grouped together. If languages share more similar word forms, a clear separation of the languages will be difficult to achieve (Persici et al., 2019). The extensive cross-linguistic interactions that are predicted by the model are similar to the Bilingual Interactive Activation Model of Lexical Access, or “BIMOLA” (Lewy & Grosjean, 2008), which allows for bottom-up as well as top-down activation spreading. Both models, BLINCS and BIMOLA, describe bilingual activation as a result of interactive mapping within and between the language levels. However, they differ in terms of whether phonetic features (BIMOLA) or whether phonological features (BLINCS) are activated first (McDonald & Kaushanskaya, 2020).

The BLINCS model assumes that spoken word recognition begins with phonemic overlap at the phonological level and results in coactivation of words that share phonemes. Connections between words are constantly updated through linguistic experience; thus, the BLINCS model can capture different activation patterns in relation to language proficiency of the listeners (Shook & Marian, 2013). It is generally assumed that a less dominant language is accessed more slowly, thus exerting weaker lexical competition. In late bilinguals (i.e., second language learners), the second language is typically the less dominant and/or less proficient language (Blumenfeld & Marian, 2007; Canseco-Gonzalez et al., 2010; Weber & Cutler, 2004). Language proficiency is a crucial factor in L1–L2 cross-linguistic activation, with high-proficiency users exhibiting more pronounced spreading effects across their known languages (Helms-Park & Dronjic, 2016; Lemhöfer et al., 2008; Silverberg & Samuel, 2004; van Hell & Dijkstra, 2002; though there are contrasting views, as seen in Hameau et al., 2021). Additionally, the more dominant first language is known to exert greater influence

over the second language than vice versa (Spivey & Marian, 1999; Weber & Cutler, 2004).

While the majority of bilingual lexical competition studies have focused on cognates (e.g., English “arm” and German “Arm”) and cross-language homophones (e.g., English “kin” and German “Kinn”/chin; see Wang et al., 2020, for an overview), phonological contact points in the mental lexicon also include cross-linguistic phonological neighbours. Word forms such as English “lock” and German “Lack” ([lak], Engl. *lacquer*) share the majority of their phonology and would be expected to be found in an integrated phonological neighbourhood in speakers of English and German. Few studies to date have explicitly addressed cross-language neighbourhoods or the impact of “foreign neighbours” on word recognition, and the focus on distantly related languages (such as English and Spanish; see Vitevitch, 2012; also see Marian et al., 2012) complicates predictions for how closely related languages behave in interaction. The present study was designed to test cross-linguistic phonological neighbourhood effects and spoken word recognition in bilingual speakers of the two Germanic languages, English and German, concentrating on the question of whether cross-language phonological neighbourhoods are fully integrated in the bilingual mental lexicon. Exploring the interactive nature of the bilingual lexicon depends on the methodological tools to construct cognitive models that can combine different languages and inform about their degree of integration. Network science has been a useful method to analyse phonological neighbourhood statistics of single languages (Siew et al., 2019; Vitevitch, 2020, 2022), and it is well suited to be applied to study interactive effects between multiple languages (in the form of “multiplex networks”; Stella et al., 2024).

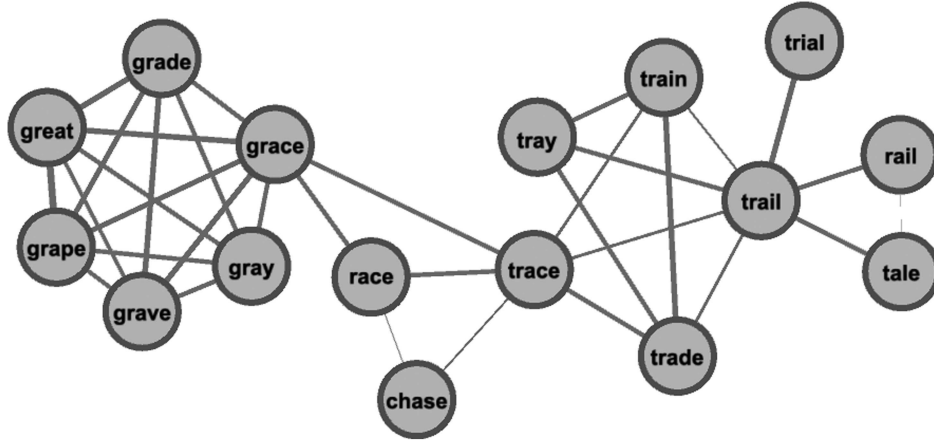
Phonological Networks

Network science examines complex systems by representing their components as “nodes” and the links between them as “edges” (Barabási, 2016). In phonological networks, the components are the word forms in a given lexicon, and edges are the links that connect them based on phonological similarity. While phonological relationships between words can be based on different metrics for network construction, for instance percentage of phonological similarity (Kapatsinski, 2006) or perceived similarity (Castro & Vitevitch, 2023; Siew & Castro, 2023), the majority of phonological networks have been based on the Levenshtein distance that defines words as similar if they differ by one phoneme via substitution (e.g., *cat-hat*), addition (e.g., *cat-catch*), or deletion (e.g., *cat-at*; see Levenshtein, 1966). The Levenshtein metric is commonly used in psycholinguistic studies of word neighbourhoods (Vitevitch & Luce, 2016); thus, its use in phonological networks increases comparability of findings. Figure 1 shows a partial phonological network of one-segment Levenshtein neighbours.

Phonological networks have been proven to be psycholinguistically meaningful as they influence lexical production and perception (e.g., Castro et al., 2020; Castro & Siew, 2020; Vitevitch, Chan, & Roodenrys, 2012; Vitevitch et al., 2014). Specifically, speech recognition has been shown to be impacted by the location of a word in the phonological network (Siew & Vitevitch, 2016). Densely linked network parts (i.e., the “giant component”) lead to recognition delays as activation can spread more extensively among phonological neighbours and their neighbours. These findings suggest that the proportion of densely linked versus loosely linked parts of a

Figure 1

A Partial Network of English-as-a-Second-Language at Advanced (Common European Framework of Reference for Languages C1) Proficiency



Note. Nodes are represented as phonological word forms; edges represent the links between them. Edge thickness indicates phonological similarity.

phonological network is crucial for word recognition. Phonological network science has also shown words in key positions in a phonological network (i.e., those that bind various parts of a network together) are endowed with a distinct processing advantage (Vitevitch & Goldstein, 2014). Phonological networks play a role in the probability of a word being learned in first and second language acquisition (Luef, 2022; Siew & Vitevitch, 2020a; Vitevitch & Goldstein, 2014), with words in more well-connected neighbourhoods having an advantage in word learning at initial stages of language acquisition.

Network science can take into account additional details on the links in phonological neighbourhoods, for instance, an analysis of how neighbours of target words are linked to one another. When inspecting Figure 1, the nodes “grace” and “trace” possess a similar number of neighbours (7 and 6), but the interlinking of the neighbours is different. This can be measured with the clustering coefficient, a network statistic based on the formula:

$$CC = 2e_n / (k_n(k_n - 1)), \quad (1)$$

with k_n being the number of neighbours per node (n) and e_n the number of connected pairs between all neighbours (Watts & Strogatz, 1998). Hence, $CC = 0$ if none of the neighbours of a node are linked, and $CC = 1$ if all are. The clustering coefficient has been shown to influence word recognition (Chan & Vitevitch, 2009, 2010; Yates, 2013). Words in highly clustered neighbourhoods (i.e., the majority of phonological neighbours of a target word are also neighbours of one another) are characterised by delayed lexical recognition (Chan & Vitevitch, 2010). The linking patterns of the neighbours lead to certain patterns of activation spreading within the neighbourhood, which in turn determine how much initial activation the target word will receive. Such investigations of specific phonological neighbourhood characteristics are made possible by a network approach to the phonological lexicon. In the above example, “trace” has a CC of 0.3, while “grace” has a CC of 0.5. From these network statistics, new predictions about lexical coactivation and the factors influencing the delay or facilitation

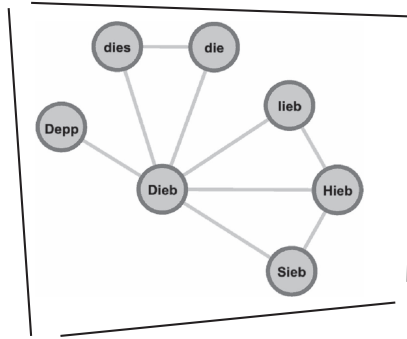
of lexical retrieval depending on a word’s location within the network can be derived (Castro et al., 2020; Siew, 2019; Vitevitch et al., 2011).

With network science, the links between words can be modified to show different connection strengths (i.e., “weighted edges” in network terms). Levenshtein neighbours may vary in terms of their phonological similarity, with words like “pat” and “bat” differing only in the voicing features of the word-initial segments, whereas neighbouring pairs such as “pat” and “sat” differ in terms of manner and place of articulation in the word-initial segments. As evident in Figure 1, the phonological similarity between “grade” and “great” is high, indicated by the thicker edge strength. The phonological similarity between “rail” and “tail” is low, resulting in very thin edges in the network visualisation. It is known that phonologically close neighbours spread more lexical activation and thus heighten lexical competition (Goldrick et al., 2013; Munson & Solomon, 2004; Scarborough, 2013). Including a phonological similarity measure in a phonological network based on Levenshtein neighbours can improve predictions about the diffusion of lexical activation (see, e.g., Luef, 2024b).

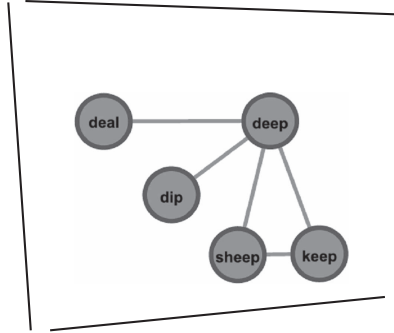
As network science has proven a valuable tool to chart lexical links in the mental lexicon of single-language users, it makes sense to combine multiple languages in a more complex, combinatorial network to test the effects of lexical links across languages. While simple networks depict nodes and their links in one language, multiplex networks can be constructed that link two languages together by placing links between neighbouring nodes across two network layers (Castro et al., 2020; Pio-Lopez et al., 2021; Stella, 2020; Stella et al., 2017). This approach provides an integrated analysis of word form links within and across languages, such as German and English, as depicted in Figure 2. The term “multiplex network” is used to describe combined networks that feature the same set of nodes replicated across multiple layers, whereas “multilayer networks” feature different sets of words linked across multiple layers (Stella et al., 2017, 2024). The distinction can be blurred when multiple network layers of phonological networks

Figure 2**Multiplex Network**

(a)

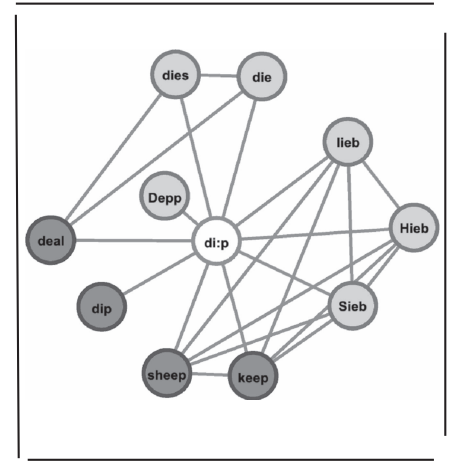


L1 German
“Dieb” [di:p]



L2 English
“deep” [di:p]

(b)



Note. Single language networks focusing on the target word form [di:p] (German “thief”) in L1 (first language) German and L2 (second language) English in (a). A combined multiplex network of German and English phonological neighbours is shown in (b).

show some common word forms. The term “multiplex network” will be used here to refer to phonological networks that represent largely distinct sets of words (see Figure 2a) but that show some degree of overlapping (identical) nodes across the layers (see Figure 2b).

Multiplex networks combining two phonological lexica are most meaningful when dealing with typologically related languages (such as English and German), where lexicon and phonotactics exhibit sufficient similarities to construct multiplex neighbourhood links. Given the assumption that activation spreads between similar word forms, multiplex phonological networks provide a tool to study bilingual lexical competition. Rather than focus on target words and their immediate phonological neighbourhoods (as in Figure 2a), the cognitive links across languages become analysable through various network-derived metrics (as in Figure 2b). The present study was designed to test whether lexical competition effects between German and English in late bilingual speakers (i.e., second language learners) are influenced by the structure of a multiplex cognitive model of bilingual speech processing. Such an approach is designed to account for lexical neighbourhood characteristics that have been shown to be meaningful in single-language phonological networks and that may be able to account for lexical processing in bilingual speech.

Bilingual Network Statistics

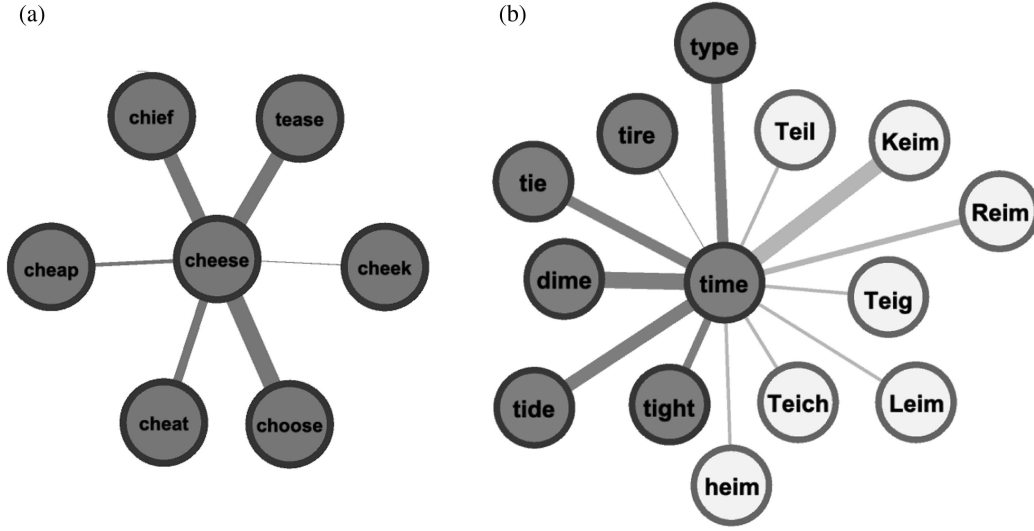
Several network statistics are relevant for understanding activation spreading in word recognition in bilingual cognition. The network degree centrality (or simply “degree”) quantifies the number of direct links a node possesses (Barabási, 2016); it is essentially a measure of phonological neighbourhood density. In multiplex networks, the metric *multiplex degree centrality* (or multiplex degree) considers the degree of a node, accounting for all links to a node from any layer of the network (Bianconi, 2018). This allows multiplex networks to assess whether phonological

neighbours from the target language alone or from both languages known by bilingual speakers can explain neighbourhood density effects in word recognition. Figures 3a and 3b provide examples of phonological neighbourhoods in an L2 English lexicon (Figure 3a) and in a German-English integrated lexicon (3b). Both neighbourhoods have the same number of English neighbours, but the neighbourhood in 3b includes additional German neighbours. The first research question addressed by the present study examines whether German neighbours are equally influential during English lexical recognition, potentially leading to increased competition in the neighbourhood and resulting in recognition delays of English target words such as “time.”

There is rarely perfect phonological agreement of segments across different languages, and English and German are no exception. For instance, German and English diphthongs such as [ai] in German “Keim” [kam] (Engl. *sprout*) and English “time” (see Figure 3b) differ slightly in terms of their segmental durations. Such phonetic details can add immense complexity to phonological networks. By focussing on phonology rather than phonetics, certain generalisations are introduced to a phonological network, but interpretability is heightened in service of a reasonable simplification.

A related note concerns the average relative phonological similarity between phonological neighbours across various languages. For instance, “time” and “dime” in Figure 3b are more similar than “time” and German “Teich” [taiç] (Engl. *pond*). Phonological similarities between neighbouring words can be estimated with the ALINE algorithm that compares word pairs in a number of phonological features, such as place and manner of articulation and vowel height (Kondrak, 2000), and calculates a similarity score between them. A phonological similarity score can be added to lexical links in a phonological network to provide details about phonological closeness of dyadic neighbour relationships (see Luef, 2024b). It can be assumed that cross-language neighbours are characterised by greater phonological distance to an English target word than the English

Figure 3
English and German Phonological Neighborhoods



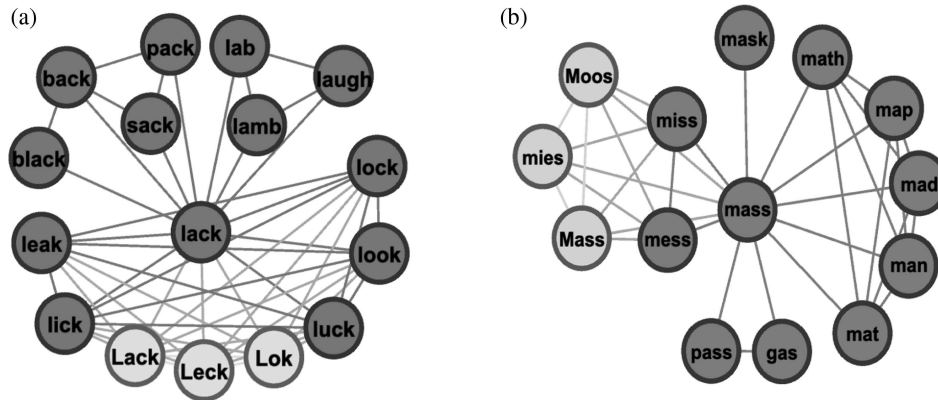
Note. (a and b) Two exemplary English target words with an equal number of English phonological neighbours ($k = 6$) but differing in the number of German phonological neighbours.

neighbours of that word; in fact, this is shown in the edge weights in [Figure 3b](#). When phonological similarity scores are incorporated in a bilingual multiplex network, weighted edges result, meaning that the links between words are based on phonological similarities (and displayed by thick or thin edges linking words in a network graph). This adds information concerning predictions about lexical diffusion and competition arising from the coactivation of cross-linguistic neighbours, especially in the context of bilingual comparisons when both languages are not equal in terms of phonological similarity of the Levenshtein neighbours.

Multiplex phonological neighbourhoods may exhibit varying degrees of interconnectedness among their constituents, a characteristic captured by the *multiplex clustering coefficient*. While the clustering

coefficient of a simple (i.e., single layer) network describes the linking of the neighbourhood beyond the target word, the multiplex clustering coefficient (or “multiplex CC”) incorporates all links across multiple layers in a multiplex network. In the present case, this means that the interlinking of all English and German neighbours of a target word is measured (see [Figure 4a and 4b](#)). In addition to degree influencing the speed and accuracy of word identification, clustering contributes to lexical competition and causes recognition delays in words where numerous neighbours are interconnected ([Chan & Vitevitch, 2009](#)). In [Figure 4a](#), the presence of German neighbours significantly enhances clustering in the neighbourhood, elevating the CC from 0.26 in the simple L2 English networks to 0.33 in the multiplex network. [Figure 4b](#) illustrates a neighbourhood where the German neighbours do not

Figure 4
Clustering in Phonological Neighborhoods



Note. (a and b) Two exemplary English target words and their clustered neighbourhoods. German neighbours raise clustering coefficient in 4a but not in 4b.

raise the overall clustering coefficient (English CC = 0.3, multiplex CC = 0.27). The second research question addressed by the present study investigates the impact of German clustering on English word recognition and whether the additional clustering introduced by German neighbours in a phonological neighbourhood leads to recognition delays for English target words.

Studies focussing on the phonological networks of individual languages have played a crucial role in advancing theories of speech processing (Vitevitch, 2020, 2022). The incorporation of multiple languages into a multiplex network can offer a model of the bilingual lexicon, showing how multilingual language users navigate and manage multiple languages within their cognitive systems. The shared processing of different languages emphasises the interconnectedness of these languages in a speaker's mind, highlighting the potential for leveraging existing linguistic knowledge to facilitate speech processing across different languages. This can show the ways in which multilingual individuals utilise their linguistic repertoire.

Method

Phonological networks of L1 German and L2 English were constructed before they were later combined in a multiplex network. Select words with specific network statistics from the multiplex network were then incorporated in a lexical decision task where participants' response accuracies and reaction times (RTs) were recorded to test whether the multiplex network constitutes a meaningful model of bilingual speech processing.

Network Data

Phonological networks of L1-German and L2-English were constructed based on data extracted from two corpora in the case of German, the *SUBTLEX-DE Corpus* (Brysaert et al., 2011) obtained from the Clearpond database (Marian et al., 2012), and in the case of English, the *English Vocabulary Profile* database, a subdatabase of the Cambridge Learner Corpus, the largest corpus of English as a second language (Capel, 2015). While corpora may be a suboptimal representation of language users' lexicon, often underestimating the actual vocabulary size (Gráf, 2017), they can be more realistic data than using dictionaries to construct lexical networks. In the case of second languages, corpora are a good method to model vocabulary knowledge as they reflect actual language use of the learners rather than theoretical word knowledge assumed by second language textbooks (see Vitevitch, Stamer, & Kieweg, 2012). For the multiplex network of German and English, it seemed most reasonable to combine similar types of data; hence, data from two language corpora were used.

L1 German Network

For the German network, a total of 27,751 words were obtained and further processed in the following ways: Proper nouns were removed (i.e., including names of people, regions, geographical structures, and landmarks were removed, e.g., *Anne*, *Paris*, *Toskana*/Engl. "Tuscany," *Mississippi Fluss*/"Mississippi River," *Eiffelturm*/"Eiffel Tower"), including nouns and adjectives (e.g., *italienisch*/"Italian"). In addition, recent and infrequent foreign or loan words were removed, for instance *Amigo*, *Empire*, *Mayday*, in addition to their Germanised morphological variants

(e.g., *bowlen*, *gebeamt*). Well-established loan words in German, typically characterised by Germanised pronunciation (e.g., *Apokalypse*, *Appartement*) and/or higher lexical frequency (e.g., *Shopping-Center*) and their morphological forms (*auschecken*/"check out," *einloggen*/"log in") were retained. All compound nouns merging English and German word forms, such as *Actionszene*/"action scene," were removed. Additional data removal encompassed specific medical terms (e.g., *Hepatitis*), abbreviations (e.g., *Abi* for *Abitur*/"high school diploma," *Dino* for *Dinosaurier*/"dinosaur"), informal swear words, numbers larger than 12 and not multiples of 10 (e.g., *dreißig*/"twenty-three") with the exception of frequent numeric phrases such as *eineinhalb*/"one-and-a-half" or *einmal*/"once" but not, for instance, *dreimal*/"three times" or *vierteilen*/"divide by four." All remaining words were lemmatised, and inflections were removed (e.g., *dem*, *diesen*, *habe*, *machst*, *glücklicher*), with the few exceptions including participle II forms, such as *gegessen*/"eaten," because they also function as adjectives. If present in the data set, both genders were retained for nouns, such as *Abgeordnete/Abgeordneter* ("female/male delegate") and *Arzt/Ärztin* ("male/female physician"). In case of minor phonological variation in one semantic word form, such as *Brüderschaft/Bruderschaft* ("brotherhood"), *Tür/Türe* ("door"), or *eklig/ekelig* ("disgusting"), the phonologically less marked one was retained (*Bruderschaft*, *Tür*, *eklig*). In addition, clitics were split into their constituent parts (e.g., *im* → *in*, *aufs* → *auf das*).

In the next step, all lemmata not contained in the German Clearpond database were removed, which amounted to 414 words, or 3.7% of the data. The final vocabulary on which the following analyses were performed was 14,546 words. Phonological transcriptions of all words were obtained from the Clearpond database (in the form of CP-SAMPA transcriptions), with the necessity to make certain adaptations and corrections to various transcriptions. First, dots in between symbols were removed and double-symbol transcription symbols were replaced by single symbols (e.g., "E3" was replaced with "e"). Furthermore, a number of German transcriptions were inaccurate, and problems included front vowels (/e/-i/, e.g., in *Besen*-/bizn/) and some instances of voiced fricatives in word-initial position (e.g., *Verb*-/fERp/). In addition, errors were found in loanword transcriptions (e.g., *circa*-/kIRka/), which were manually corrected. Since CP-SAMPA does not distinguish vowel length, word pairs like *Rahmen/rammen* and *Aal/All* became homophones, which is acceptable in a number of German varieties. Homophones were represented as one word form in the final network (e.g., *bis_Biss*), and their lexical frequency rates were averaged across the two words. As lexical frequency rates, the token frequencies obtained from the Clearpond database for German were used.

Phonological length of words was determined with the Excel function *len()* applied to the phonological transcriptions, resulting in a phoneme count of word length.

As a last step, all word forms in the German lexicon were compared to one another to determine their phonological distance, and words differing by one phonological segment via substitution, deletion, or addition (i.e., the Levenshtein distance; see Landauer & Streeter, 1973) were identified as phonological neighbours of one another. This meant that a link would be placed between them in the lexical network. Word form comparisons were conducted with an Oracle 12c Big-Data-Lite database (Bryla, 2015) and the function *edit_distance* of the package "utl_match."

The resulting data were analysed and visualised with the R package "igraph" (Csardi & Nepusz, 2006) and the network

software “Gephi” (Bastian et al., 2009), and various network statistics were calculated (see Table 1 for an overview).

L2 English Network

A total of 5,387 American English words of the proficiency level C1 (=advanced proficiency) were retrieved for network construction from the *English Vocabulary Profile*. Second language proficiency levels were defined according to the “Common European Framework of Reference for Languages” (see Council of Europe, 2018). The vocabulary list was processed in a similar way as for the German network, and phrases and clitics were split into their constituent parts (*think about—think, about; it’s—it, is*), and abbreviations and acronyms were removed (e.g., CD). Lemmatisation was applied to the vocabulary list, and it was decided to keep only a few cases of frequent inflected forms, including inflections of the verbs “be” and “have,” as well as inflections of the adjective “good” (*better, best*) but removing all other inflectional forms, including comparatives (e.g., *happy—happier*) and plurals (e.g., *cat—cats*). As lexical frequency counts, the token frequencies obtained from the Clearpond database for English were used.

The word list was then processed in the English database of Clearpond that contained 100% of the words. The final word count was 5,147. Phonological transcriptions of all words were obtained from the Clearpond database (in the form of CP-SAMPA transcriptions), and adaptations to the transcriptions included the removal of dots in between symbols and double-symbol transcription symbols (e.g., “r0” was replaced with “r”). Homophones were represented as one word form in the final network (e.g., *meet_meat*), with lexical frequency rates being averaged across the constituent words.

Phonological length of words was calculated in Excel with the function *len()*, and phoneme counts were used as measures of word length.

Phonological neighbours were calculated in the same way as in the German network, with all word forms of the English lexicon being compared to one another in an Oracle database. Words that differed by one segment via deletion, substitution, or addition were denoted as phonological neighbours of one another.

All network statistics of the English phonological network were computed with “igraph” and “Gephi.”

Multiplex Network: German and English

As a final step, an integrated vocabulary of German and English word forms was constructed that consisted of 19,599 word forms. Homophonous forms in German and English were merged as one word form, for instance in German *Dieb* / “thief” and English *deep* or in German *hell* / “bright” and English *hell*. Lexical frequency rates were averaged across both languages and log-transformed. The resulting vocabulary was processed in the Oracle database to determine all phonological neighbours in the combined dataset of German and English words. Subsequently, “igraph” and “Gephi” were used to compute network statistics (see Table 1).

The C1 English lexicon shares about 1.53% identical word forms with German. In German, many phonological neighbours arise from derivational morphology, for instance *morgens* / “in the mornings” or instances of participle I, such as *bewegend* / “moving,” which can potentially have implications for lexical processing of the language since the words share meaning as well as phonology (Vitevitch & Stamer, 2006).

The multiplex network constructed here is an unbalanced one in the sense that the number of German nodes and the number of English nodes differ significantly. This is a representation of a second language lexicon, with an L1-sized first language lexicon combined with a much smaller L2 lexicon. Network statistics are always influenced by network size, and network variables such as clustering coefficient or centrality measures are inevitably dependent on the size of the underlying data set. This makes comparisons across differently sized networks difficult, and comparative analyses across the German and English networks or across the single-language and multiplex networks have to be treated with caution.

Phonological Similarity

Phonological similarities between neighbouring words were calculated for English nodes with German neighbours (i.e., for all bilingual neighbourhoods that link German to English word forms). The R package “*alineR*” (Downey et al., 2017) computes the ALINE algorithm as proposed by Kondrak (2000) based on phonological transcriptions of word pairs. AlineR output values range from 0 (identical word forms) to 1 (maximal phonological distance). In order to obtain intuitive edge strengths for a network, those distance

Table 1

Overview of Network Statistics in the L1-German, L2-English, and the Integrated Multiplex (German-English) Networks

Variable	L1-German	L2-English	Multiplex network
Network size (number of words)	14,546	5,147	19,599
Connectivity (number of links)	8,553	5,294	15,301
Giant component words (%)	18.3%	33%	22.5%
Giant component links (%)	72.1%	90%	82.7%
Percentage of nodes in giant component of multiplex network	74%	26%	
German-English homophonous forms	0.4%		
Average degree	1.2	2.1	1.6
Average degree in GC	4.6	5.6	5.7
Average path length	7.9	6.3	7.5
Clustering coefficient	0.28	0.34	0.3
Ratio of edges to vertices/ β	0.59	1.03	0.8
Ratio of edges to vertices/ β (GC)	2.31	2.8	2.9

Note. L1 = (first language) German; L2 = (second language) English; GC= giant component.

scores were transformed to similarity scores by multiplying them times 100 and subtracting them from 100. This resulted in weighted edges that can be understood as percentage correspondences in terms of phonology; for instance, a score of 88 corresponds to 88% phonological similarity between two words. Table 2 lists the original alineR scores and the computed phonological similarity values for example words.

Lexical Decision Experiments

Participants and Experimental Designs

First-language speakers of German who are enrolled at a tertiary English language study programme at a German university and who know English at a proficiency level of C1 were recruited for two auditory lexical decision experiments. Participants either received credit towards the completion of a course requirement or were paid 10 Euros for their participation. None of them reported a history of speech or hearing disorders. The Research Ethics Committee of Charles University approved all experiments reported in this article (reference number UKFF/46874/2023).

Two research aims were pursued with the use of the auditory lexical decision tasks. First, it was examined whether German phonological neighbours influence the processing of English words. For this, two sets of English words (controlled for various factors; see descriptions below) were selected that showed no difference in English degree centrality but significant differences in multiplex degree centrality, with German neighbours densifying the English neighbourhoods. The second research question concerned the clustering coefficient of words in the multiplex network, and here two sets of English words (controlled for various factors; see descriptions below) were selected that only differed by multiplex clustering coefficient. In both experiments, participants were presented with either a word or a nonword over a set of headphones, and they had to decide whether or not what they heard is an actual English word. Reaction times were measured in order to obtain a time course of spoken word recognition and to see whether the variables *multiplex degree* and *multiplex clustering coefficient* had an effect on the temporal aspect of spoken word recognition. Given the assumption of an integrated bilingual lexicon with coactivation of German and English word forms, it was predicted that words of high multiplex degree centrality and high multiplex clustering coefficient would be responded to more slowly, reflecting detriments in processing in the reaction times. In addition, judgement accuracy of the responses was analysed.

Procedure

The two lexical decision tasks were created with PsychoPy 2023.2.3 (Peirce et al., 2019), which controlled the presentation of

the stimuli words (in random order) and the collection of responses (accuracy, reaction times). Audio recordings of real words and nonwords were obtained from the MALD database (Tucker et al., 2019). The experiments were put online on the Pavlovla platform between October 2023 and January 2024, and data were collected during that period. Each participant was tested individually at home at their personal computers and using their preferred headphones. Each trial started with a small cross appearing on the computer screen for 500 ms. Participants were then presented once with the target word. They were instructed to use the computer keys “a” and “k,” indicating whether a word was a nonword (pressing “a” with their left hand) or a real word (pressing “k” with their dominant hand) in the English language. They were instructed to respond as quickly and accurately as possible. The next trial only started once a response had been entered. Response times were measured from the onset of the stimulus to the onset of the key press response.

Prior to the start of the experiments, participants were given five practice trials to become familiar with the task and to adjust the volume of the auditory stimuli at their comfort level. The practice trials were not included in the data analyses. The experiment lasted about 10 min.

Only data of monolingual speakers of German who are learners of English as a second language at C1 level were included in the analyses.

Experiment 1: Degree Centrality

Forty participants were recruited for Experiment 1. Twenty English monosyllabic words were used as stimuli in the first experiment, in addition to 20 nonwords. The selected words were short, high frequency English words likely known by learners at lower proficiency levels than C1. While keeping English degree centrality similar, half of the words were characterised by high degree centrality in the multiplex network, and half had low multiplex degree centrality. This allowed the testing of the effect of German neighbours on the processing of the English words. Each target word had at least 50% German neighbours in their neighbourhood. Mann–Whitney *U* tests were employed to ascertain that the two groups of words—that is, high multiplex degree centrality group or “High mDeg group” and low multiplex degree centrality or “Low mDeg group”—did not differ in any of the control variables (see Table 3).

The variable “German-added degree” equalled the degree that the German neighbours add to the overall neighbourhood degree. It was interacted with the English degree variable to determine whether additional degree added by the German neighbours had a different effect on English words of low or high degree (see Experiment 1: Multiplex Degree section).

Table 2
Examples of AlineR Results and Computed Phonological Similarity Scores for Word Pairs

Word “cheese”	Neighbour	alineR distance score	Phonological similarity
“tʃiz”	“tʃuz” (choose)	0.12	88
“tʃiz”	“tiz” (tease)	0.16	84
“tʃiz”	“tʃik” (cheek)	0.35	65

Table 3
Variables Controlled in Experiment 1

Control variable	Data	Mann-Whitney <i>U</i> test, variable of interest across high mDeg and low mDeg group	High mDeg group, median	Low mDeg group, median
English degree centrality	Obtained from the English proficiency level C1 network	$U = 14.5, p = .008$	3 SEM = 0.33	6 SEM = 0.6
Phonemic length	Phoneme count per IPA transcribed word form	$U = 39, p = .47$	4 SEM = 0.15	4 SEM = 0.15
Lexical frequency rate	Obtained from Clearpond, log-transformed values ranged between 1 and 3	$U = 36, p = .32$	1.85 SEM = 0.21	1.43 SEM = 0.2
Phonotactic probability	Calculated with <i>Phonotactic Probability Calculator</i> (Vitevitch & Luce, 2004), biphone probabilities ranged between 0.001 and 0.6	$U = 27, p = .09$	0.009 SEM = 0.001	0.003 SEM = 0.001
Neighbourhood frequency	Obtained from Clearpond, values ranged from 16 to 379	$U = 28, p = .11$	47.2 SEM = 24.8	101.3 SEM = 35.9
English clustering coefficient	Obtained from the C1 network, values ranged between 0.16 and 1	$U = 47.5, p = .91$	0.33 SEM = 0.1	0.3 SEM = 0.08

Note. mDeg = multiplex degree; SEM = standard error of mean.

Furthermore, a variable termed “German neighborhood frequency” was calculated for each English target word neighbourhood to determine the degree of influence the German neighbours may exert (assuming that high-frequency German neighbours exert more influence in the neighbourhood in terms of activation; see Costa & Santesteban, 2004). Therefore, German lexical frequency rates of all German neighbours were obtained from Clearpond, and an average value was calculated, which was subsequently log-transformed. As German neighbourhood frequency was difficult to control, it was added as an independent variable to all statistical models.

Edge weight analysis was conducted on the multiplex degree network data. The phonological similarity scores of all English neighbours of a target word were added, and the variable name “weighted degree: English” was given to the variable. The phonological similarity scores of all German neighbours of a target word were also added, and “weighted degree: English” was subtracted from “weighted degree: German,” resulting in a variable that defines the edge weights that the German neighbours add to a neighbourhood (“German-added weighted degree”). Weighted degree tends to correlate with degree in networks, which precludes an inclusion of both variables in the same statistical model. Therefore, a separate weighted degree analysis was conducted on the multiplex weighted degree data.

The nonwords to supplement the real words in the lexical decision task were obtained from the MALD database of pseudowords and were chosen so that each target word had a corresponding nonword that differed by one or two phonological segments, preferably in word-final position. For instance, the word “fly” had the corresponding nonword [flam], while “loud” was matched with [laʊŋ] (see [Supplemental Material A1](#) for a list of the target words and nonwords). It was made sure that target and nonword had large phonological overlaps, including similarity in stress placement, length, and number of syllables. In addition, the chosen nonwords were not similar to any existing German words. Nonwords that could be mistaken for real words based on German-accented English phonology were avoided, such as [raʊŋ] which could possibly be perceived as “wrong” by German learners of English.

Experiment 2: Clustering Coefficient

Thirty-nine German students of English participated in the second experiment. Thirty-four English monosyllabic English words served as stimuli, in addition to 34 nonwords that were included in the experiment. The selected words were short, high frequency English words likely known by learners at lower proficiency levels than C1. The variable of interest in Experiment 2 was multiplex clustering coefficient, and half of the 34 real words were characterised by high multiplex clustering coefficients in the multiplex network (referred to as “High mCC group”), while the other half showed low multiplex clustering coefficients (referred to as “Low mCC group”). Wilcoxon signed rank tests were employed to exclude differences between the High mCC and the Low mCC groups of words concerning the control variables (see [Table 4](#)).

The variable “German-added CC” was calculated by subtracting the English CC from the multiplex CC of target words. In the case of the corresponding variable in Experiment 1, “German-added degree” is identical to German degree as calculated in the German phonological network. In the case of “German-added CC,” the calculation of how much CC the German neighbours add to an

Table 4
Variables Controlled in Experiment 2

Control variable	Data	Wilcoxon signed rank test, variable of interest across high mCC and low mCC group	High mCC group, <i>Mdn</i>	Low mCC group, median
English degree centrality	Obtained from the English proficiency level C1 network	$W = 32.5, p = .21$	13 SEM = 0.5	14 SEM = 0.6
Multiplex degree centrality	Obtained from the multiplex network	$W = 59, p = .98$	17 SEM = 0.7	17 SEM = 0.6
Proportion of German neighbours	Count of German phonological neighbours	$W = 34.5, p = .084$	20 SEM = 2.6	11.8 SEM = 1.8
Phonemic length	Phoneme count per IPA transcribed word form	$W = -0136, p = .58$	3 SEM = 0.1	3 SEM = 0.1
Lexical frequency rate	Obtained from Clearpond, log-transformed values ranged between 0.6 and 2	$W = 45, p = .23$	1.6 SEM = 0.1	1.26 SEM = 0.1
Phonotactic probability	Calculated with <i>Phonotactic Probability Calculator</i> (Vitevitch & Luce, 2004), biphone probabilities ranged between 0.00 to 0.05	$W = 43, p = .11$	0.01 SEM = 0.003	0.01 SEM = 0.004
Neighbourhood frequency	Obtained from Clearpond, values ranged from 20 to 918	$W = 58, p = .38$	158.9 SEM = 52.0	158.8 SEM = 56.7
English clustering coefficient	Obtained from the C1 network, values ranged between 0.18 and 0.33	$W = 53, p = .44$	0.27 SEM = 0.01	0.28 SEM = 0.01

Note. mCC = multiplex clustering coefficient; SEM = standard error of mean.

English neighbourhood differs from pure German CC. Since CC takes into account the ratio of neighbours and neighbourhood connections, the German neighbours in an English neighbourhood can change CC in ways that are unpredictable from the CC of the German words in the German network. In Experiment 2, “German-added CC” was interacted with English CC in order to investigate the possibility that additional CC may have a different effect on English words of various CC. As in Experiment 1, “German neighborhood frequency” was added to the statistical models as a control variable.

Nonword selection was conducted in an identical way to Experiment 1. Suitable words in terms of phoneme and syllable make-up in relation to the English target words and words that were not similar to or confusable with any German word were obtained from the MALD database (see [Supplemental Material A2](#) for a list of words and nonwords used in the second experiment).

Statistical Analyses

In a first step, all real words and their reaction times were retrieved. From the raw response data obtained in both experiments, all reaction times below 200 ms and above 4,000 ms were removed. In first-language experimentation, a value between 2,000 and 3,000 ms is typically chosen as the cutoff value; however, it makes sense to allow for more response time in second language learners as a large portion of the responses ranged between 2,000 and 4,000 ms (i.e., 17% of the data). Adult second language learners take more time in lexical judgements, and thus only latencies over 4,000 ms were considered outliers and excluded from analysis. Trials with latencies that were 3 standard deviations (*SDs*) above or below each participant’s mean reaction time (RT) were also considered outliers and removed. In Experiment 1, the data removal amounted to 3.8% of the original response data being removed, resulting in 96.2% of the trials being included in the analyses. Trials from one word were excluded due to very low overall rating accuracies: “thought” with an error rate of 56% in Experiment 1. In the second experiment, 2.3% of the data were removed, resulting in 97.7% of the trials being included in the analyses. All words from the second experiment were included in the analyses, as none of them had an error rate of over 30%.

Both experiments (multiplex degree, multiplex CC) were divided into accuracy and RT models, the former testing whether the accurate judgement of a word as a real or nonword was influenced by multiplex degree and the latter testing the effect of multiplex degree on the reaction times (in ms) to the correct responses in the experimental trials. Using the *lme4* package in R, generalized linear mixed effects (GLM) models were used to predict accuracy from the response data (Bates et al., 2014) when the outcome was a binary variable (correct/incorrect). Binomial regressions with “cloglog” links were specified in the GLMs since the data contained a large number of correct and a small number of incorrect responses (Experiment 1: 95 vs. 5%; Experiment 2: 92 vs. 8%). Linear mixed effects models (LMEs) were computed to determine the relationship between response times and the variables of interest. The models constructed for GLMs and LMEs were identical in their structure.

Experiment 1: Multiplex Degree. The multiplex degree models investigated in Experiment 1 included the following fixed effects: English degree and German-added degree and an interaction between the two to account for a possible polynomial effect where

German-added degree has a different effect at lower than at higher English degrees. The variables “English lexical frequency rate” and “German neighborhood frequency” were also added. Random slopes of “participant” and “word” were specified for the maximal number of fixed effects (before convergence became an issue) in all models. The LME models only included correct responses of the participants. In the LME models, degrees of freedom and significance values (p) can be obtained by approximations, with the Satterthwaite’s method being a popular method for this. The R package *lmerTest* (Kuznetsova et al., 2017) was used to calculate p values for all the LME models reported. The pseudocode of the models is shown in the following, with the dependent variable either accuracy of responses or reaction times:

$$\begin{aligned}
 \text{DV} \sim & \text{English Degree} \times \text{German-Added Degree} \\
 & + \text{English lexical frequency rate} \\
 & + \text{German neighbourhood frequency} \\
 & + (1 + \text{English degree}|\text{Participant}) \\
 & + (1 + \text{English degree}|\text{Word}) \\
 & + (1 + \text{German-added degree}|\text{Participant}) \\
 & + (1 + \text{German-added degree}|\text{Word}) \\
 & + (1 + \text{English lexical frequency rate}|\text{Participant}) \\
 & + (1 + \text{English lexical frequency rate}|\text{Word}) \\
 & + (1 + \text{German neighbourhood frequency}|\text{Participant}) \\
 & + (1 + \text{German neighbourhood frequency}|\text{Word}). \quad (2)
 \end{aligned}$$

Weighted Degree. The weighted degree model was structured in an identical way to the degree model and included as fixed effects “weighted degree: English” and “German-added weighted degree.” The two variables were interacted in order to see whether German-added weighted degree has different effects on English words with higher or low weighted degree. In addition, English lexical frequency rate, and German neighbourhood frequency were included. Random slopes could only be included for “weighted degree: English” before the model failed to converge. The following pseudocode was implemented:

$$\begin{aligned}
 \text{DV} \sim & \text{weighted degree: English} \times \text{German-added weighted degree} \\
 & + \text{English lexical frequency rate} \\
 & + \text{German neighborhood frequency} \\
 & + (1 + \text{weighted degree: English}|\text{Participant}) \\
 & + (1 + \text{weighted degree: English}|\text{Word}). \quad (3)
 \end{aligned}$$

Experiment 2: Multiplex Clustering Coefficient. The multiplex CC models calculated on data from Experiment 2 included the following variables: English CC, German-added CC, and an interaction between the two that can account for different effects of the German-added CC at various levels of English CC in the target words. In addition, English lexical frequency rate and German neighbourhood frequency were included as fixed effects. The maximal number of random slopes were included. For the binary regression involving accuracy of word judgement, a binomial regression was constructed. The response time data was processed with an LME model. The pseudocode of the models is shown below.

$$\begin{aligned}
 \text{DV} \sim & \text{English CC} \times \text{German-Added CC} \\
 & + \text{English lexical frequency rate} \\
 & + \text{German neighbourhood frequency} \\
 & + (1 + \text{English CC}|\text{Participant}) \\
 & + (1 + \text{English CC}|\text{Word}) \\
 & + (1 + \text{German-added CC}|\text{Participant}) \\
 & + (1 + \text{German-added CC}|\text{Word}) \\
 & + (1 + \text{English lexical frequency rate}|\text{Participant}) \\
 & + (1 + \text{English lexical frequency rate}|\text{Word}) \\
 & + (1 + \text{German neighbourhood frequency}|\text{Participant}) \\
 & + (1 + \text{German neighbourhood frequency}|\text{Word}). \quad (4)
 \end{aligned}$$

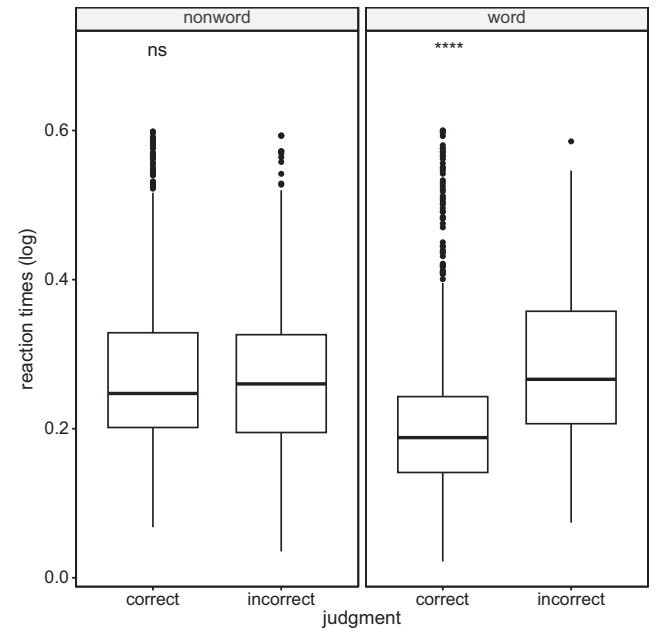
Results

Nonwords were responded to more slowly than real words in both experiments (see Figure 5).

Experiment 1: Multiplex Degree Accuracy Model

The sample size for this model was 749 items. Results showed no effect of multiplex degree on the accuracy of word judgments (see Table 5). The model showed no effects of any of the variables on judgement accuracy.

Figure 5
Nonwords Were Characterised by Longer Reaction Times



Note. Experiment 1: 696 correct and 53 incorrect word judgements; Experiment 2: 1,183 correct and 112 incorrect word judgements. *ns* = not statistically significant.
*** $p < 0.001$.

Table 5*Generalized Linear Mixed Effect Model Estimates for Fixed and Random Effects*

Variable	Variance	Estimate	SD	SE	p
Random effects					
Word					
English degree	8.930e – 01		9.450e – 0		
German-added degree	8.710e – 01		9.333e – 01		
English lexical frequency rate	7.833e – 01		8.850e – 01		
German neighbourhood frequency	2.385e – 08		1.544e – 04		
Participant					
English degree	8.746e – 01		9.352e – 01		
German-added degree	9.709e – 01		9.854e – 01		
English lexical frequency rate	7.750e – 02		2.784e – 01		
German neighbourhood frequency	4.264e – 10		2.065e – 05		
Fixed effects					
Intercept		2.274e + 01		1.962e + 01	.25
English degree		–3.13e + 01		2.901e + 01	.28
German-added degree		–2.27e + 01		1.996e + 01	.25
German neighbourhood frequency		–1.355e – 04		3.581e – 04	.71
English lexical frequency rate		5.457e – 01		1.290e + 00	.67
English Degree × German-Added Degree		3.269e + 01		2.964e + 01	.27

Note. Multiplex degree accuracy model. SE = standard error.

Experiment 1: Multiplex Degree Reaction Time Model

After removal of incorrect responses from the data ($N = 53$, =7% of data) the sample size of the reaction time model was 696 responses. Results show that reaction times to English words with German neighbors were not influenced by German-added degree, nor by any of the other variables entered into the model (see Table 6).

Weighted Edges: Phonological Similarity

The edge weight model included 696 correct responses to target words with English and German neighbours. Results showed that English weighted degree as well as German-added weighted degree affected reaction times (see Table 7).

Both English weighted degree and German-added weighted degree led to an increase in reaction times (see Figures 6 and 7).

Table 6*Linear Mixed Effect Model Estimates for Fixed and Random Effects*

Variable	Variance	Estimate	SD	SE	p
Random effects					
Word					
English degree	8.190e – 03		9.050e – 02		
German-added degree	8.902e – 03		9.435e – 02		
English lexical frequency rate	1.143e – 10		1.069e – 05		
German neighbourhood frequency	1.447e – 10		1.203e – 05		
Participant					
English degree	3.653e – 04		1.911e – 02		
German-added degree	3.091e – 03		5.560e – 02		
English lexical frequency rate	7.610e – 05		8.724e – 03		
German neighbourhood frequency	3.479e – 05		5.898e – 03		
Fixed effects					
Intercept		–0.052669		0.306726	.87
English degree		0.3547921		0.388149	.37
German-added degree		0.2441317		0.308624	.44
German neighbourhood frequency		–0.000678		0.008318	.94
English lexical frequency rate		–0.003782		0.011475	.75
English Degree × German-Added Degree		–0.320526		0.398141	.43

Note. Multiplex degree reaction time model. SE = standard error.

Table 7
Linear Mixed Effect Model Estimates for Fixed and Random Effects

Variable	Variance	Estimate	SD	SE	p
Random effects					
Word					
Weighted degree: English	0.0174614		0.13214		
Participant					
Weighted degree: English	0.0003709		0.01926		
Fixed effects					
Intercept		−0.694894		0.296337	.04*
Weighted degree: English		0.321654		0.102560	.009**
German-added weighted degree		0.361531		0.145134	.02*
German neighbourhood frequency		0.010237		0.007867	.21
English lexical frequency rate		0.017003		0.011159	.16
Weighted Degree: English x German-Added Weighted Degree		−0.138065		0.053731	.02*

Note. Weighted degree reaction time model; SE = standard error.

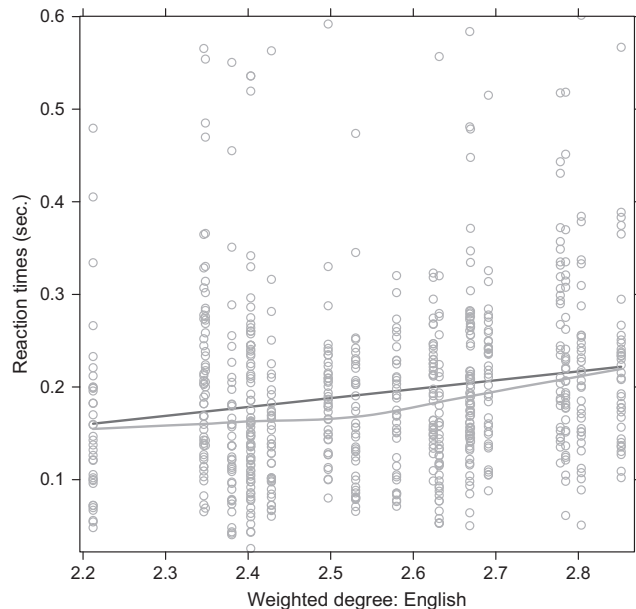
* $p < .05$. ** $p < .01$.

The interaction between English weighted degree and German-added weighted degree was such that the effect of additional German weighted degree had a greater effect on raising reaction times to target words of lower weighted English degree (see Figure 8).

Experiment 2: Multiplex CC Accuracy Model

The sample size of this model was 1,295 responses, and results showed no effects of German-added CC on the correct judgement of words (see Table 8).

Figure 6
Partial Residual Plot Showing That Higher English Weighted Degree Leads to Increased Reaction Times



Note. Partial residuals are displayed in lightgray, focal variable with corresponding regression line in darkgray.

Experiment 2: Multiplex CC Reaction Time Model

After the removal of incorrect responses ($N = 112$, =8.6%), the final sample size of the reaction time model was 1,183 responses. Results show that reaction times were influenced by English CC, German-added CC, and the interaction between the two (see Table 9 and Figure 9).

In addition, higher English CC was shown to lead to longer response latencies (see Figure 9).

The additional CC added by the German neighbours and English CC showed an interactional effect where German-added CC had the largest effect of raising reaction times in English words of lower CC values (see Figure 10).

The results of the experiments can be found on open science framework (see Luef, 2024a).

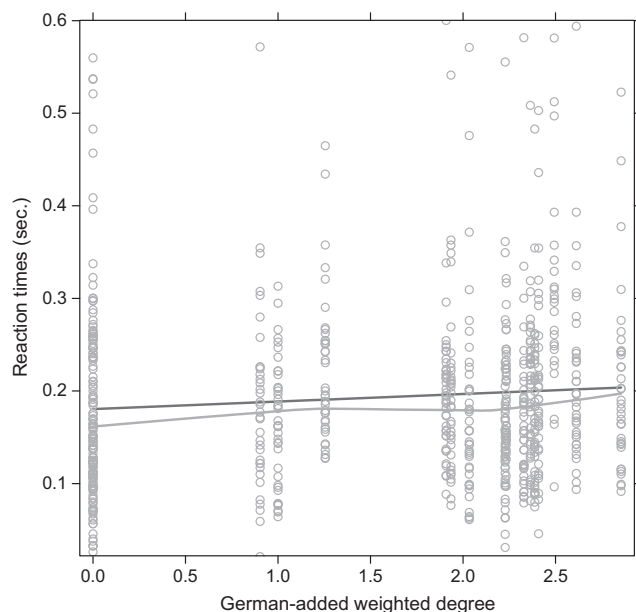
Discussion

The present study explored how first-language German words influence the recognition of English target words in German-English late bilinguals (German learners of English as a second language). Utilising network science to chart a multiplex network of L1 German and L2 English word forms, network-theoretical considerations of the German-English bilingual lexicon were tested in terms of their implications for spoken word recognition.

The results indicate that the networked multiplex German-English lexicon has predictive capabilities for English word recognition. While multiplex degree—a metric of cross-linguistic English-German phonological neighbourhood density—was not shown to influence reaction times to English words, weighted degree—a metric of phonological similarity between neighbouring words—had an effect on response latencies. Specifically, higher weighted degree (i.e., higher rates of phonological similarity) in cross-language German neighbours caused recognition delays in the English target words. This indicates that during the processing of English words, coactivation of German neighbours has an effect on lexical processing. The absence of an effect of multiplex degree in the present analyses could indicate that it is not the number of neighbours per se that impacts processing, but the phonological

Figure 7

Partial Residual Plot Showing That Weighted Degree Added by German Phonological Neighbours Leads to Increased Reaction Times in the English Target Words



Note. Partial residuals are displayed in lightgray, focal variable with corresponding regression line in darkgray.

relationships between neighbouring words are evaluated during lexical processing. L1 neighbours seem to exert influence on L2 target words primarily when the phonological relationships between the words is close, suggesting that one of the causes of the “bilingual processing disadvantage” can potentially be found on the segmental level. Moreover, the German-added processing effect was strongest in English words of lower weighted degrees (a discussion of the interaction effects can be found below).

The results of the experiment testing clustering coefficients of the bilingual neighbourhoods demonstrate an effect of English clustering in the monolingual English neighbourhoods. Higher clustering in the English-only neighbourhoods led to delays in reaction times to the English target words. These findings are in agreement with [Chan and Vitevitch \(2009\)](#), who found such an effect in monolingual L1-English neighbourhoods. More diffusion of the lexical activation in the clustered neighbourhoods likely leads to the longer latencies, as listeners take additional time to resolve the more wide-spread coactivation in the clustered neighbourhoods ([Vitevitch, 2022](#); [Vitevitch, Chan, & Roodenrys, 2012](#); [Vitevitch et al., 2011](#)). German-added CC had the largest effect on low-CC English words, where it significantly slowed down word recognition. The fact that especially low-English CC words were influenced by additional German CC indicates a special sensitivity to neighbourhood clustering in sparsely clustered neighbourhoods. In high-CC words that already show substantial levels of clustering, cross-linguistic clustering had less of an impact.

The cross-linguistic phonological neighbourhoods showed complex interactivity between the two involved lexica. It is noteworthy that the measured effects of German neighbours on English target

words (in terms of weighted degree and clustering coefficient) were more pronounced in cases where the network statistics were lower in the English words. The influence of the German phonological neighbours seems to unfold more fully when there is less lexical competition from the English neighbours. This could indicate that English neighbours are stronger competitors in a neighbourhood of an English target word, and when there are few of them, the German neighbours are able to exert more influence on lexical processing. What needs to be kept in mind is the fact that English and German neighbourhood densities were typically similar, and sparse English neighbourhoods tended to have sparse German neighbourhoods as well. This means that the weighted degrees and clustering coefficients of German neighbours in generally sparse neighbourhoods affected lexical processing the most.

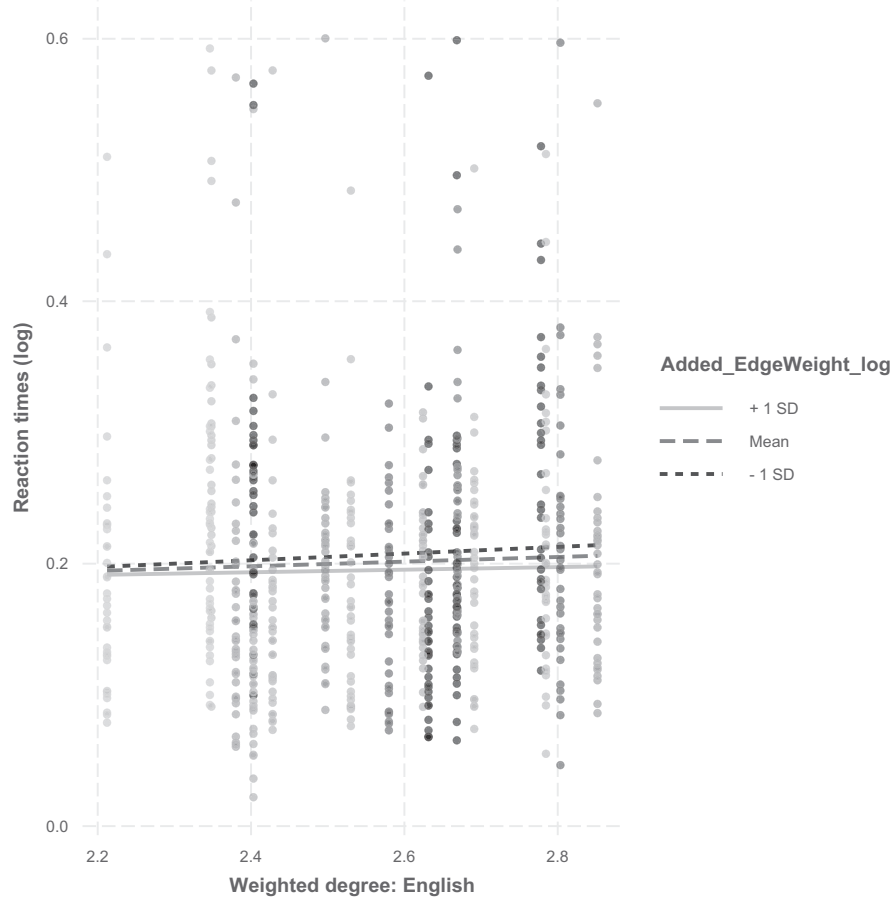
In sum, the findings of the present study suggest that word recognition mechanisms are not strictly limited to a target language that is in current use by a speaker but operate across languages in bilingual phonological neighbourhoods of second language learners. Both experiments showed measurable additional lexical processing demands on L2 target words that stemmed from the first-language vocabulary, and this is likely the result of lexical competition between word forms across both languages. Viewing the mental lexicon of late bilinguals as an accumulation of interconnected word forms across all known languages, with lexical competition among cross-linguistic phonological neighbours, can provide context for understanding the lexical processing delays observed in the German-English late bilingual speakers.

The integration of German and English word forms in bilinguals’ mental lexicon leads to the assumption that lexica of monolingual and multilingual language users are different in terms of lexical density and/or connectivity. The way the lexicon—and phonological neighbourhoods in it—are structured has implications for how quickly language users can retrieve words. Bilingual speakers seem to face more challenges as they have to navigate more word forms and resolve heightened lexical competition effects ([Marian & Spivey, 2003a, 2003b](#)).

The current results have to be evaluated in light of a number of limitations that apply to this study. First, participants were not tested on their knowledge of the target words and phonological neighbours. Even though they all reported proficiency level of C1 in English, it may well be the case that there are interpersonal differences in the L2 English and L1 German vocabularies, which could have influenced the results. Future studies should confirm vocabulary knowledge of all words that are presumed to be involved in bilingual lexical processing. A specific focus on pseudoneighbours can also make a significant contribution to bilingual lexical processing. By taking into account that German learners of English may misperceive “think” as either “sink” (Engl. “sink”) or “fink” (Engl. “finch”), the phonological neighbourhood of such a learner conceivably includes “sing,” “sin,” and “fin,” all of which would not be considered neighbours of “think” in the present experimental design. It would be ideal to incorporate a metric of pseudoneighbourhood that can capture the German (perceptual) accent in English in order to chart a more accurate model of the lexical competition that takes place in German-English late bilinguals. Last, the sample sizes of the experiments were different, which complicates comparisons. The absence of an effect of English/German degree centrality may have been influenced by the smaller sample size used for that analysis.

Figure 8

German-Added Weighted Degree Raised Reaction Times Especially in Target Words of Lower Weighted English Degree

**Table 8**

Generalized Linear Mixed Effect Model Estimates for Fixed and Random Effects

Variable	Variance	Estimate	SD	SE	p
Random effects					
Word					
German-added CC	1.649e + 01		4.061e + 00		
English CC	1.574e - 07		3.968e - 04		
English lexical frequency rate	4.780e - 01		6.914e - 01		
German neighbourhood frequency	4.856e - 10		2.204e - 05		
Participant					
German-added CC	1.424e - 03		3.774e - 02		
English CC	2.392e + 00		1.547e + 00		
English lexical frequency rate	1.297e - 03		3.602e - 02		
German neighbourhood frequency	2.992e - 02		1.730e - 01		
Fixed effects					
Intercept		-0.92190		1.19491	.44
English CC		-3.25716		2.01917	.11
German-added CC		9.67308		35.93398	.79
English lexical frequency rate		0.23748		0.27971	.4
German neighbourhood frequency		0.05277		0.11411	.64
English CC × German-Added CC		26.80637		61.02866	.66

Note. Multiplex CC accuracy model. *SE* = standard error; CC = clustering coefficient.

Table 9
Linear Mixed Effect Model Estimates for Fixed and Random Effects

Variable	Variance	Estimate	SD	SE	p
Random effects					
Word					
German-added CC	1.934e – 01		4.398e – 01		
English CC	1.774e – 01		4.212e – 01		
English lexical frequency rate	3.225e – 09		5.679e – 05		
German neighbourhood frequency	4.576e – 12		2.139e – 06		
Participant					
German-added CC	3.092e – 01		5.561e – 01		
English CC	2.939e – 07		5.421e – 04		
English lexical frequency rate	7.512e – 05		8.667e – 03		
German neighbourhood frequency	2.607e – 05		5.106e – 03		
Fixed effects					
Intercept		0.403758		0.071741	3.23e – 06***
English CC		0.331725		0.121013	.01*
German-added CC		–7.012010		3.406868	.05*
English lexical frequency rate		–0.004752		0.010297	.65
German neighbourhood frequency		0.000237		0.005395	.97
English CC × German-Added CC		–12.330141		5.861871	.048*

Note. Multiplex CC reaction time model. SE = standard error; CC = clustering coefficient.

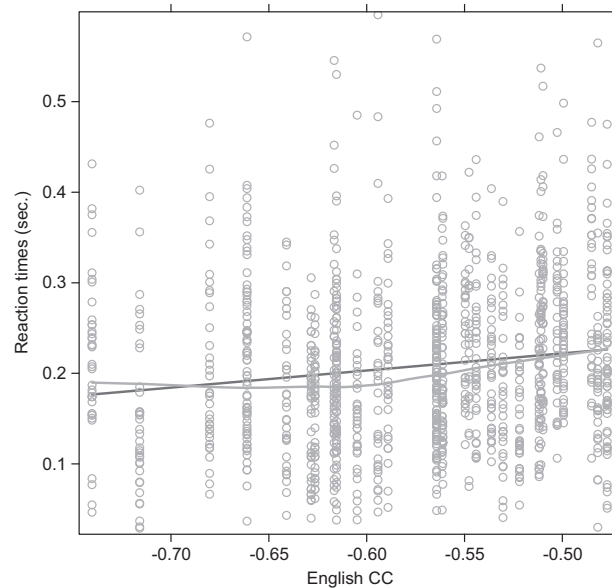
* $p < .05$. *** $p < .001$.

The field of network science has been heralded as “a new way to look at old questions in cognitive psychology” (Chan & Vitevitch, 2010, p. 1), and the application of multiplex networks to model the bilingual lexicon can serve as an example of that. The incorporation of a network-based approach to the mental lexicon of single-language users has contributed to our understanding of lexical competition, activation spreading, and word storage and

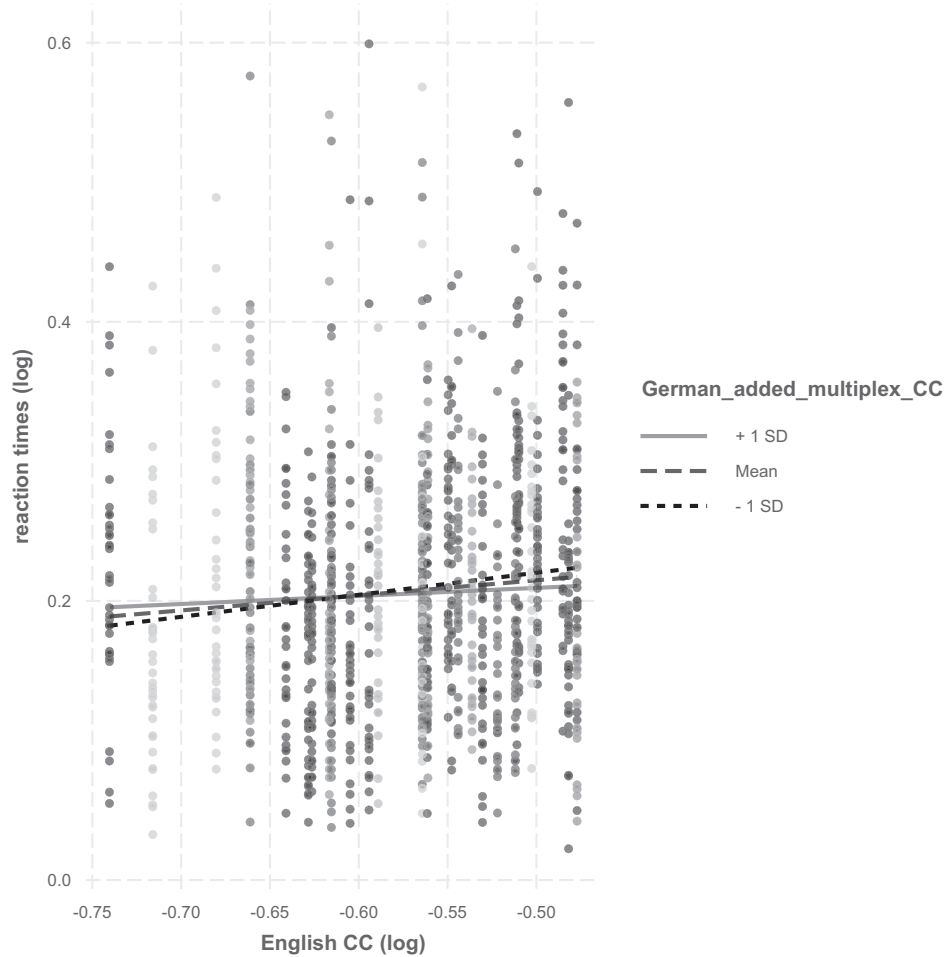
retrieval (Siew et al., 2019; Vitevitch, 2020, 2022). This not only affirms the validity of phonological networks as concrete concepts in psycholinguistics but also establishes them as essential tools for gaining a more nuanced insight into word relationships and phonological neighbourhoods. The inclusion of multiplex networks in the methodological toolkit of bi- and multilingualism can facilitate the conceptual integration of multiple language

Figure 9

Partial Residual Plot Showing That Reaction Times Were Longer for English Words in Neighbourhoods With Higher English CC



Note. Partial residuals are displayed in lightgray, focal variable with corresponding regression line in darkgray. CC = clustering coefficient.

Figure 10*German-Added CC Raised Reaction Times Especially in English Target Words of Lower English CC**Note.* CC = clustering coefficient.

cognition, offering a quantitative framework for multilingual modelling.

Emerging research opportunities in bilingual multiplex networks are diverse and encompass enquiries into cross-linguistic influence directions, such as the impact of a second language (L2) on the first language (L1) or the interactions between three or more languages. Multiplex networks are representations of the complex system *language* and can provide models of the versatile interactions between different languages in a speaker's mind, while at the same time preserving the unique characteristics of each contributing language.

Résumé

Les réseaux de formes de mots phonologiques du lexique mental présentent un intérêt psycholinguistique, car ils permettent de comprendre l'efficacité de l'accès au lexique. Bien que de nombreuses recherches se soient concentrées sur les premières langues, de plus en plus d'éléments suggèrent que les réseaux phonologiques des

secondes langues sont tout aussi importants pour les processus lexicaux. Il est proposé que le traitement linguistique bilingue implique l'intégration de la première et de la deuxième langue, l'activation lexicale se propageant entre les formes de mots similaires dans les deux langues. Les réseaux multiplexes fournissent un cadre pour combiner différents réseaux phonologiques, ce qui permet d'analyser le comportement du système lexical intégré pendant le traitement lexical. Dans le contexte de la présente étude, qui se concentre sur les apprenants allemands de l'anglais en tant que deuxième langue, une analyse de réseau multiplex a été construite pour modéliser la complexité interactive du lexique mental bilingue. L'étude a testé les effets interlinguistiques dans une tâche de reconnaissance de mots utilisant des stimuli anglais. Les résultats ont révélé que pendant le traitement lexical dans leur deuxième langue, l'anglais, les germanophones activent également les voisins phonologiques de l'allemand. En outre, les bilingues sont attentifs à l'interconnexion (c'est-à-dire au regroupement) des voisins allemands et anglais dans le voisinage phonologique des mots cibles anglais. Ces résultats peuvent contribuer au débat en cours sur le degré d'intégration du lexique mental bilingue et mettre

en lumière le rôle que les réseaux phonologiques peuvent jouer dans la modélisation du traitement lexical bilingue.

Mots-clés : réseau multiplex, allemand, anglais, réseau phonologique, reconnaissance des mots

References

- Barabási, A. L. (2016). *Network science*. Cambridge University Press.
- Bartolotti, J., & Marian, V. (2012). Bilingual memory: Structure, access, and processing. In J. Altarriba & L. Isurin (Eds.), *Memory, language and bilingualism: Theoretical and applied approaches* (pp. 7–47). Cambridge University Press. <https://doi.org/10.1017/CBO9781139035279.002>
- Bastian, M., Heymann, S., & Jacomy, M. (2009). *Gephi: An open source software for exploring and manipulating networks* [Paper presentation]. International AAAI Conference on Weblogs and Social Media, San Jose, CA, United States. <https://gephi.org/publications/gephi-bastian-feb09.pdf>
- Bates, D., Maechler, M., Bolker, B., & Walker, S. (2014). *Linear mixed-effects models using Eigen and S4* (R Package Version 1.1-7) [Computer software]. <https://cran.r-project.org/web/packages/lme4/lme4.pdf>
- Bianconi, G. (2018). Structural correlations of multiplex networks. In G. Bianconi (Ed.), *Multilayer networks: Structure and function* (pp. 129–145). Oxford University Press. <https://doi.org/10.1093/oso/9780198753919.003.0007>
- Blumenfeld, H. K., & Marian, V. (2007). Constraints on parallel activation in bilingual spoken language processing: Examining proficiency and lexical status using eye-tracking. *Language and Cognitive Processes*, 22(5), 633–660. <https://doi.org/10.1080/01690960601000746>
- Bryla, B. (2015). *Oracle Database 12c handbook: Manage a scalable, secure Oracle enterprise database environment*. McGraw Hill.
- Brysbaert, M., Buchmeier, M., Conrad, M., Jacobs, A. M., Bölte, J., & Böhl, A. (2011). The word frequency effect: A review of recent developments and implications for the choice of frequency estimates in German. *Experimental Psychology*, 58(5), 412–424. <https://doi.org/10.1027/1618-3169/a000123>
- Brysbaert, M., Van Dyck, G., & Van de Poel, M. (1999). Visual word recognition in bilinguals: Evidence from masked phonological priming. *Journal of Experimental Psychology: Human Perception and Performance*, 25(1), 137–148. <https://doi.org/10.1037/0096-1523.25.1.137>
- Canseco-Gonzalez, E., Brehm, L., Brick, C. A., Brown-Schmidt, S., Fischer, K., & Wagner, K. (2010). Carpet or carcel: The effect of age of acquisition and language mode on bilingual lexical access. *Language and Cognitive Processes*, 25(5), 669–705. <https://doi.org/10.1080/01690960903474912>
- Capel, A. (2015). The English vocabulary profile. In J. Harrison & F. Barker (Eds.), *English profile in practice* (pp. 9–27). Cambridge University Press.
- Castro, N., & Siew, C. S. Q. (2020). Contributions of modern network science to the cognitive sciences: Revisiting research spirals of representation and process. *Proceedings of the Royal Society: Mathematical, Physical and Engineering Sciences*, 476(2238), Article 20190825. <https://doi.org/10.1098/rspa.2019.0825>
- Castro, N., Stella, M., & Siew, C. S. Q. (2020). Quantifying the interplay of semantics and phonology during failures of word retrieval by people with aphasia using a multiplex lexical network. *Cognitive Science*, 44(9), Article e12881. <https://doi.org/10.1111/cogs.12881>
- Castro, N., & Vitevitch, M. S. (2023). Using network science and psycholinguistic megastudies to examine the dimensions of phonological similarity. *Language and Speech*, 66(1), 143–174. <https://doi.org/10.1177/00238309221095455>
- 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(6), 1934–1949. <https://doi.org/10.1037/a0016902>
- Chan, K. Y., & Vitevitch, M. S. (2010). Network structure influences speech production. *Cognitive Science*, 34(4), 685–697. <https://doi.org/10.1111/j.1551-6709.2010.01100.x>
- Chen, Q., & Mirman, D. (2012). Competition and cooperation among similar representations: Toward a unified account of facilitative and inhibitory effects of lexical neighbors. *Psychological Review*, 119(2), 417–430. <https://doi.org/10.1037/a0027175>
- Costa, A., Pannunzi, M., Deco, G., & Pickering, M. J. (2017). Do bilinguals automatically activate their native language when they are not using it? *Cognitive Science*, 41(6), 1629–1644. <https://doi.org/10.1111/cogs.12434>
- Costa, A., & Santesteban, M. (2004). Lexical access in bilingual speech production: Evidence from language switching in highly proficient bilinguals and L2 learners. *Journal of Memory and Language*, 50(4), 491–511. <https://doi.org/10.1016/j.jml.2004.02.002>
- Council of Europe. (2018). *Common European framework of reference for languages: Learning, teaching, assessment. Companion volume with new descriptors*. Cambridge University Press.
- Csardi, G., & Nepusz, T. (2006). The igraph software package for complex network research. *Frontiers in Immunology*, 1695, Article 862049. <https://doi.org/10.3389/fimmu.2022.862049>
- de Groot, A. M., Delmaar, P., & Lupker, S. J. (2000). The processing of interlexical homographs in a bilingual and monolingual task: Support for non-selective access to bilingual memory. *Quarterly Journal of Experimental Psychology*, 53(2), 397–428. <https://doi.org/10.1080/713755891>
- Dijkstra, T., Grainger, J., & Van Heuven, W. J. B. (1999). Recognition of cognates and interlingual homographs: The neglected role of phonology. *Journal of Memory and Language*, 41(4), 496–518. <https://doi.org/10.1006/jmla.1999.2654>
- Downey, S. S., Sun, G., & Norquest, P. (2017). alineR: An R package for optimizing feature-weighted alignments and linguistic distances. *The R Journal*, 9(1), 138–152. <https://doi.org/10.32614/RJ-2017-005>
- Frances, C., Navarra-Barindelli, E., & Martin, C. D. (2021). Inhibitory and facilitatory effects of phonological and orthographic similarity on L2 word recognition across modalities in bilinguals. *Scientific Reports*, 11, Article 12812. <https://doi.org/10.1038/s41598-021-92259-z>
- Goldrick, M., Vaughn, C., & Murphy, A. (2013). The effects of lexical neighbors on stop consonant articulation. *The Journal of the Acoustical Society of America*, 134(2), EL172–EL177. <https://doi.org/10.1121/1.4812821>
- Gráf, T. (2017). The story of the learner corpus LINDSEI_CZ. *Studies in Applied Linguistics*, 8(2), 22–35. https://dspace.cuni.cz/bitstream/handle/20.500.11956/97524/1541592_tomas_graf_22-35.pdf?sequence=1&isAflowed=y
- Hameau, S., Biedermann, B., & Nickels, L. (2021). Lexical activation in late bilinguals: Effects of phonological neighbourhood on spoken word production. *Language, Cognition and Neuroscience*, 36(4), 517–534. <https://doi.org/10.1080/23273798.2020.1863438>
- Hanulíková, A., & Weber, A. (2012). Sink positive: Linguistic experience with th substitutions influences nonnative word recognition. *Attention, Perception, & Psychophysics*, 74(3), 613–629. <https://doi.org/10.3758/s13414-011-0259-7>
- Helms-Park, R., & Dronjic, V. (2016). Crosslinguistic lexical influence: Cognate facilitation. In R. Alonso Alonso (Ed.), *Crosslinguistic influence in second language acquisition* (pp. 71–92). Multilingual Matters. <https://doi.org/10.21832/9781783094837-007>
- Hoversten, L. J., & Traxler, M. J. (2020). Zooming in on zooming out: Partial selectivity and dynamic tuning of bilingual language control during reading. *Cognition*, 195, Article 104118. <https://doi.org/10.1016/j.cognition.2019.104118>
- Jared, D., & Kroll, J. F. (2001). Do bilinguals activate phonological representations in one or both of their languages when naming words? *Journal of Memory and Language*, 44(1), 2–31. <https://doi.org/10.1006/jmla.2000.2747>

- Ju, M., & Luce, P. A. (2004). Falling on sensitive ears: Constraints on bilingual lexical activation. *Psychological Science*, 15(5), 314–318. <https://doi.org/10.1111/j.0956-7976.2004.00675.x>
- Kapatsinski, V. (2006). Sound similarity relations in the mental lexicon: Modeling the lexicon as a complex network. *Speech Research Lab Progress Report*, 27, 133–152. <https://citeseerx.ist.psu.edu/document?repid=rep1&type=pdf&doi=70c90848df98d4e7b61bb210dda317f972febece>
- Kondrak, G. (2000). *A new algorithm for the alignment of phonetic sequence* [Conference session]. 1st North American Chapter of the Association for Computational Linguistics Conference, Stroudsburg, PA, United States. <https://aclanthology.org/A00-2038.pdf>
- Kuznetsova, A., Brockhoff, P. B., & Christensen, R. H. (2017). lmerTest package: Tests in linear mixed effects models. *Journal of Statistical Software*, 82(3), 1–26. <https://doi.org/10.18637/jss.v082.i13>
- 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(2), 119–131. [https://doi.org/10.1016/S0022-5371\(73\)80001-5](https://doi.org/10.1016/S0022-5371(73)80001-5)
- Lemhöfer, K., Dijkstra, T., & Michel, M. (2004). Three languages, one ECHO: Cognate effects in trilingual word recognition. *Language and Cognitive Processes*, 19(5), 585–611. <https://doi.org/10.1080/01690960444000007>
- Lemhöfer, K., Dijkstra, T., Schriefers, H., Baayen, R. H., Grainger, J., & Zwitserlood, P. (2008). Native language influences on word recognition in a second language: A megastudy. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 34(1), 12–31. <https://doi.org/10.1037/0278-7393.34.1.12>
- Levenshtein, V. I. (1966). Binary codes capable of correcting deletions, insertions, and reversals. *Soviet Physics, Doklady*, 10(8), 707–710. <https://nymity.ch/sybilhunting/pdf/Levenshtein1966a.pdf>
- Lewy, N., & Grosjean, F. (2008). The Lewy and Grosjean BIMOLA model. In F. Grosjean (Ed.), *Studying bilinguals* (pp. 201–210). Oxford University Press.
- Luce, P. A., & Pisoni, D. B. (1998). Recognizing spoken words: The neighborhood activation model. *Ear and Hearing*, 19(1), 1–36. <https://doi.org/10.1097/00003446-199802000-00001>
- Luef, E. M. (2022). Growth algorithms in the phonological networks of second language learners: A replication of Siew & Vitevitch (2020a). *Journal of Experimental Psychology: General*, 151(12), e26–e44. <https://doi.org/10.1037/xge0001248>
- Luef, E. M. (2024a, May 16). *Late bilingual German-English phonological network*. <https://osf.io/yzhxj/>
- Luef, E. M. (2024b). Obsolescence effects in second language phonological networks. *Memory & Cognition*, 52, 771–792. <https://doi.org/10.3758/s13421-023-01500-9>
- Marian, V., Bartolotti, J., Chabal, S., & Shook, A. (2012). CLEARPOND: Cross-linguistic easy-access resource for phonological and orthographic neighborhood densities. *PLOS ONE*, 7(8), Article e43230. <https://doi.org/10.1371/journal.pone.0043230>
- Marian, V., Blumenfeld, H. K., & Boukrina, O. V. (2008). Sensitivity to phonological similarity within and across languages. *Journal of Psycholinguistic Research*, 37(3), 141–170. <https://doi.org/10.1007/s10936-007-9064-9>
- Marian, V., & Spivey, M. (2003a). Bilingual and monolingual processing of competing lexical items. *Applied Psycholinguistics*, 24(2), 173–193. <https://doi.org/10.1017/S01421716403000092>
- Marian, V., & Spivey, M. (2003b). Competing activation in bilingual language processing: Within- and between-language competition. *Bilingualism: Language and Cognition*, 6(2), 97–115. <https://doi.org/10.1017/S1366728903001068>
- Marslen-Wilson, W. D., & Zwitserlood, P. (1989). Accessing spoken words: The importance of word onsets. *Journal of Experimental Psychology: Human Perception and Performance*, 15(3), 576–585. <https://doi.org/10.1037/0096-1523.15.3.576>
- McClelland, J. L., & Elman, J. L. (1986). The TRACE model of speech perception. *Cognitive Psychology*, 18(1), 1–86. [https://doi.org/10.1016/0010-0285\(86\)90015-0](https://doi.org/10.1016/0010-0285(86)90015-0)
- McDonald, M., & Kaushanskaya, M. (2020). Factors modulating cross-linguistic co-activation in bilinguals. *Journal of Phonetics*, 81, Article 100981. <https://doi.org/10.1016/j.wocn.2020.100981>
- Munson, B., & Solomon, N. P. (2004). The effect of phonological neighborhood density on vowel articulation. *Journal of Speech, Language, and Hearing Research*, 47(5), 1048–1058. [https://doi.org/10.1044/1092-4388\(2004\)078](https://doi.org/10.1044/1092-4388(2004)078)
- Norris, D., McQueen, J. M., & Cutler, A. (1995). Competition and segmentation in spoken-word recognition. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 21(5), 1209–1228. <https://doi.org/10.1037/0278-7393.21.5.1209>
- Peirce, J., Gray, J. R., Simpson, S., MacAskill, M., Höchenberger, R., Sogo, H., Kastman, E., & Lindeløv, J. K. (2019). PsychoPy2: Experiments in behavior made easy. *Behavior Research Methods*, 51(1), 195–203. <https://doi.org/10.3758/s13428-018-01193-y>
- Persici, V., Vihman, M., Burro, R., & Majorano, M. (2019). Lexical access and competition in bilingual children: The role of proficiency and the lexical similarity of the two languages. *Journal of Experimental Child Psychology*, 179, 103–125. <https://doi.org/10.1016/j.jecp.2018.10.002>
- Pio-Lopez, L., Valdeolivas, A., Tichit, L., Remy, É., & Baudot, A. (2021). MultiVERSE: A multiplex and multiplex-heterogeneous network embedding approach. *Scientific Reports*, 11(1), Article 8794. <https://doi.org/10.1038/s41598-021-87987-1>
- Sadat, J., Martin, C. D., Magnuson, J. S., Alario, F.-X., & Costa, A. (2016). Breaking down the bilingual cost in speech production. *Cognitive Science*, 40(8), 1911–1940. <https://doi.org/10.1111/cogs.12315>
- Scarborough, R. (2013). Neighborhood-conditioned patterns in phonetic detail: Relating coarticulation and hyperarticulation. *Journal of Phonetics*, 41(6), 491–508. <https://doi.org/10.1016/j.wocn.2013.09.004>
- Schwartz, A. I., & Kroll, J. F. (2007). Language processing in bilingual speakers. In M. J. Traxler & M. A. Gernsbacher (Eds.), *Handbook of psycholinguistics* (pp. 963–995). Elsevier.
- Shook, A., & Marian, V. (2012). Bimodal bilinguals co-activate both languages during spoken comprehension. *Cognition*, 124(3), 314–324. <https://doi.org/10.1016/j.cognition.2012.05.014>
- Shook, A., & Marian, V. (2013). The bilingual language interaction network for comprehension of speech. *Bilingualism: Language and Cognition*, 16(2), 304–324. <https://doi.org/10.1017/S1366728912000466>
- Siew, C. S. Q. (2019). spreadr: An R package to simulate spreading activation in a network. *Behavioral Research Methods*, 51(2), 910–929. <https://doi.org/10.3758/s13428-018-1186-5>
- Siew, C. S. Q., & Castro, N. (2023). Phonological similarity judgments of word pairs reflect sensitivity to large-scale structure of the phonological lexicon. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 49(12), 1989–2002. <https://doi.org/10.1037/xlm0001271>
- Siew, C. S. Q., & Vitevitch, M. S. (2016). Spoken word recognition and serial recall of words from components in the phonological network. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 42(3), 394–410. <https://doi.org/10.1037/xlm0000139>
- Siew, C. S. Q., & Vitevitch, M. S. (2020a). An investigation of network growth principles in the phonological language network. *Journal of Experimental Psychology: General*, 149(12), 2376–2394. <https://doi.org/10.1037/xge0000876>
- Siew, C. S. Q., Wulff, D. U., Beckage, N. M., Kenett, Y. N., & Meštrović, A. (2019). Cognitive network science: A review of research on cognition through the lens of network representations, processes, and dynamics. *Complexity*, 2019, Article 2108423. <https://doi.org/10.1155/2019/2108423>
- Silverberg, S., & Samuel, A. G. (2004). The effect of age of second language acquisition on the representation and processing of second language

- words. *Journal of Memory and Language*, 51(3), 381–398. <https://doi.org/10.1016/j.jml.2004.05.003>
- Spivey, M., & Marian, V. (1999). Cross-talk between native and second languages: Partial activation of an irrelevant lexicon. *Psychological Science*, 10(3), 281–284. <https://doi.org/10.1111/1467-9280.00151>
- Stella, M. (2020). Multiplex networks quantify robustness of the mental lexicon to catastrophic concept failures, aphasic degradation and ageing. *Physica A: Statistical Mechanics and Its Applications*, 554, Article 124382. <https://doi.org/10.1016/j.physa.2020.124382>
- Stella, M., Beckage, N. M., & Brede, M. (2017). Multiplex lexical networks reveal patterns in early word acquisition in children. *Scientific Reports*, 7(1), Article 46730. <https://doi.org/10.1038/srep46730>
- Stella, M., Citraro, S., Rossetti, G., Marinazzo, D., Kenett, Y. N., & Vitevitch, M. S. (2024). Cognitive modelling of concepts in the mental lexicon with multilayer networks: Insights, advancements, and future challenges. *Psychonomic Bulletin & Review*. Advance online publication. <https://doi.org/10.3758/s13423-024-02473-9>
- Tucker, B. V., Brenner, D., Danielson, D. K., Kelley, M. C., Nenadić, F., & Sims, M. (2019). The massive auditory lexical decision (MALD) database. *Behavior Research Methods*, 51(3), 1187–1204. <https://doi.org/10.3758/s13428-018-1056-1>
- van Hell, J. G., & Dijkstra, T. (2002). Foreign language knowledge can influence native language performance in exclusively native contexts. *Psychonomic Bulletin & Review*, 9(4), 780–789. <https://doi.org/10.3758/BF03196335>
- Van Wijnendaele, I., & Brysbaert, M. (2002). Visual word recognition in bilinguals: Phonological priming from the second to the first language. *Journal of Experimental Psychology: Human Perception and Performance*, 28(3), 616–627. <https://doi.org/10.1037/0096-1523.28.3.616>
- Vitevitch, M. S. (2012). What do foreign neighbors say about the mental lexicon? *Bilingualism: Language and Cognition*, 15(1), 167–172. <https://doi.org/10.1017/S1366728911000149>
- Vitevitch, M. S. (2020). *Network science in cognitive psychology*. Routledge.
- Vitevitch, M. S. (2022). What can network science tell us about phonology and language processing. *Topics in Cognitive Science*, 14(1), 127–142. <https://doi.org/10.1111/tops.12532>
- Vitevitch, M. S., Chan, K. Y., & Goldstein, R. (2014). Insights into failed lexical retrieval from network science. *Cognitive Psychology*, 68, 1–32. <https://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(1), 30–44. <https://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, Article 369. <https://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. <https://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(4), 325–329. <https://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, Instruments, & Computers*, 36, 481–487. <https://doi.org/10.3758/BF03195594>
- Vitevitch, M. S., & Luce, P. A. (2016). Phonological neighborhood effects in spoken word perception and production. *Annual Review of Linguistics*, 2(1), 75–94. <https://doi.org/10.1146/annurev-linguistics-030514-124832>
- Vitevitch, M. S., & Stamer, M. K. (2006). The curious case of competition in Spanish speech production. *Language and Cognitive Processes*, 21(6), 760–770. <https://doi.org/10.1080/01690960500287196>
- Vitevitch, M. S., Stamer, M. K., & Kieweg, D. (2012). The Beginning Spanish Lexicon: A Web-based interface to calculate phonological similarity among Spanish words in adults learning Spanish as a foreign language. *Second Language Research*, 28(1), 103–112. <https://doi.org/10.1177/0267658311432199>
- Wang, X., Hui, B., & Chen, S. (2020). Language selective or non-selective in bilingual lexical access? It depends on lexical tones! *PLOS ONE*, 15(3), Article e0230412. <https://doi.org/10.1371/journal.pone.0230412>
- Watts, D. J., & Strogatz, D. H. (1998). Collective dynamics of ‘small-world’ networks. *Nature*, 393, 440–442. <https://doi.org/10.1038/30918>
- Weber, A., & Cutler, A. (2004). Lexical competition in non-native spoken-word recognition. *Journal of Memory and Language*, 50(1), 1–25. [https://doi.org/10.1016/S0749-596X\(03\)00105-0](https://doi.org/10.1016/S0749-596X(03)00105-0)
- Yates, M. (2013). How the clustering of phonological neighbours affects visual word recognition. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 39(5), 1649–1656. <https://doi.org/10.1037/a0032422>
- Yates, M., & Dickinson, D. (2023). How similarity influences word recognition. In J. Guendouzi, F. Loncke, & M. J. Williams (Eds.), *The Routledge international handbook of psycholinguistic and cognitive processes* (pp. 273–290). Routledge. <https://doi.org/10.4324/9781003204213-16>
- Ziegler, J. C., Muneaux, M., & Grainger, J. (2003). Neighborhood effects in auditory word recognition: Phonological competition and orthographic facilitation. *Journal of Memory and Language*, 48(4), 779–793. [https://doi.org/10.1016/S0749-596X\(03\)00006-8](https://doi.org/10.1016/S0749-596X(03)00006-8)

Received February 7, 2024

Revision received July 21, 2024

Accepted July 23, 2024 ■