

UC Merced

Proceedings of the Annual Meeting of the Cognitive Science Society

Title

The interaction between phonological and lexical variation in word recall in African American English

Permalink

<https://escholarship.org/uc/item/7gx404cj>

Journal

Proceedings of the Annual Meeting of the Cognitive Science Society, 40(0)

Authors

Mengesha, Zion

Zellou, Georgia

Publication Date

2018

The interaction between phonological and lexical variation in word recall in African American English

Zion Mengesha (zamengesha@stanford.edu)

Department of Linguistics, Margaret Jacks Hall, Bldg. 460
Stanford, CA 94305-2150 USA

Georgia Zellou (gzellou@ucdavis.edu)

Department of Linguistics, University of California, Davis, 469 Kerr Hall
Davis, CA 95616 USA

Abstract

Phonological characteristics of a voice, such as th-stopping (pronouncing *them* as “dem”) associated with African American English (AAE), provide indexical sociolinguistic information about the speaker. Word usage also signals this social dialect, i.e. usage of *crib* to mean *house*. The current study examines the effect of these sociolinguistic characteristics on word recall, as well as the interaction between the phonological and the lexical levels of variation. In a modified word recognition task, listeners displayed more accurate veridical word recall of AAE lexical items and voices. Furthermore, there was an interaction between phonological and lexical variation: listeners were even more accurate at recognizing AAE-specific lexical items heard in an AAE voice. This study adds to a growing body of work finding that sociolinguistic information influences word memory.

Keywords: memory, word recognition, lexical variation, dialect, African American English

Introduction

The factors that influence accurate recall of heard words inform models concerning the architecture of representations involved in language processing. For example, listeners falsely recall hearing the word *sleep* after hearing a word list containing several semantically related items, such as *tired*, *bed*, and *rest* (Roediger & McDermott, 1995), suggesting that lexical recognition activates semantic neighborhoods. Attention also mediates lexical encoding: in divided-attention conditions, lexical recall is more susceptible to the false memory of a non-presented item, such as *sleep*, relative to full-attention conditions (Otgaar et al., 2012; Perez-Mata et al., 2002). Furthermore, the likelihood of accurate word recall varies based on lexical categories and properties: recall is lower for emotional words (*joy*), relative to concrete (*socks*), or abstract (*old*) words (Bauer et al., 2009).

Yet, there are many other properties of words, and context, that might influence whether listeners accurately remember lexical items. Lexical (and phonological) variation can also carry social meaning: social dialects of American English, such as African American English (AAE), can be conveyed through the use of specific lexical items, such as *crib* for the Mainstream American English (MAE) *house*. Do dialect-specific lexical variants trigger different patterns of accurate word recall? Furthermore, a word’s pronunciation, e.g.,

producing *them* as “dem”, likewise is associated with AAE speakers. Thus, is there is an effect of speaker dialect on accurate lexical recall? Furthermore, do these two factors interact? There are many gaps in our understanding of how socially-conditioned lexical and phonological variation influences word recall.

Word Recall

There are two general types of models recruited to explain patterns of lexical recall. Single process models posit that when listeners retrieve a word, form and meaning are activated at the same time (e.g. Robinson & Roediger, 1997). For example, once a word is activated, both the word and its semantic neighborhoods influence false memories equally. False memories arise from continual activation of a semantic neighbor. When presenting 3, 6, 9, 12 or 15 items for a single semantic associate, Robinson and Roediger (1997) found the likelihood of veridical recall decrease as list length increases, but false memory increase as a function of list length.

Another type of account is a dual process model which posits that lexical processing leaves two kinds of memories, a verbatim form and a gist form, which are stored independently and, thus, can influence accurate recall independently. The verbatim form stores exactly what was said, while the gist form stores themes of the information being processed. Since they are independent, these forms influence memory in different ways. It has been shown that as verbatim memory gradually becomes weaker, gist extraction drives retrieval of a critical lure, occasioning false memories (Brainerd & Reyna, 2012).

A central question in the present study is which of these two models of processing allows for the mediation of social factors. A recent study by Sumner and Kataoka (2013) provides a clue. They played listeners word lists produced in either a British English, General American, or New York voice. They found that the social information cued by the talker’s voice influenced word recall: listeners were less accurate at recalling words produced by the NY talker. They interpret their result as supporting dual process models, where both word forms and semantic properties of words are activated differentially. An unresolved question is what role lexical variation plays in word recall.

A fuller understanding of the influence of social information in the simultaneous mapping of words and social categories is needed to inform our understanding of linguistic representations and speech comprehension. Thus, we look at how socially-conditioned lexical and phonological variation interact and influence word recall, in order to inform models of lexical memory.

Phonological and Lexical Variation

Prior work has shown that acoustic-phonetic properties of a voice convey social information that influence how words are perceived. For example, perceived gender and sexual orientation of the speaker influences fricative categorization (e.g., Strand & Johnson, 1996). Furthermore, phonetically cued social characteristics of a voice also influence how a word is encoded. For example, lexical activation of words with post-vocalic /r/ differs for a British English speaker and a New York English speaker (Sumner & Samuel, 2009). These findings indicate that perceived social characteristics of a speaker conveyed solely through their voice can influence how listeners encode words. Findings such as these have been vital in shaping and informing models of lexical representation. Exemplar-based theories provide a widely held account of the nature of lexical representations (e.g., Goldinger, 1998; Johnson, 2006; Pierrehumbert, 2002). Such theories posit that a lexical form is represented in memory as a cloud of traces, and every experienced token, with its particular phonetic details, leaves a unique memory trace. One approach to exemplar theory posits that heard tokens are directly encoded as they are experienced. This explains frequency effects on lexical encoding: an experienced low-frequency lexical item has greater representational weighting since there are fewer prior stored exemplars than a high frequency token (Goldinger, 1998). This account aligns with the single route model, insofar as the incoming token activates the word in form and meaning, regardless of attentional or social factors. Another perspective on exemplar theory allows for incoming tokens to be encoded more strongly if they receive greater attentional or social weighting (Pierrehumbert, 2002). Pierrehumbert posits that linguistic and social information are perceptually encoded together via rich exemplars and that socially idealized properties can weight exemplars more strongly, pulling the distributional space for lexical representations in one direction or another that can influence later perception of sounds and experiences. This perspective aligns more with a dual process account where attention is mediating how lexical recall is occurring. A proposed mechanism to account for the influence of social information on the sound-to-meaning mapping come from Sumner et al. (2013), who suggests that spoken word recognition operates in parallel with social representational meaning and interactivity between these processes can modulate linguistic comprehension.

These models of representation and processing tell us that words are stored with rich phonetic and social details from experiences and that influences how words are encoded. We

extend this to how words are remembered: does socio-linguistic variation affect accurate word recall?

Current Study

The view that listeners activate linguistic and social representations simultaneously is supported by evidence that congruent social and linguistic information has facilitating effects on lexical memory. For example, listeners better comprehend Chinese-accented English words when they co-occur with an image of an Asian face, versus a White face (McGowan, 2015). Similar socio-phonological congruency effects for AAE are seen: An AAE voice primes words pronounced with nonstandard (and, critically, specific to AAE) phonological variants, relative to standard variants (King & Sumner, 2014). Additionally, Casasanto (2008) found that listeners not only associate African American speakers with consonant cluster reduction, but recognized [mæs] as *mast* when they co-occurred with an image of an African American face, and as *mass* when it co-occurred with a White face, again demonstrating that knowledge about sociolinguistic variation influences perception. However, we do not fully understand how these phonological effects interact with the encoding and recall of dialect-specific lexical items. Hence, we ask whether AAE-specific lexical items are remembered better when produced with AAE phonological patterns, than with Mainstream American English (MAE) patterns.

As mentioned earlier, speakers' pronunciations and word choices co-vary in systematic ways. The current study focuses on the interaction between these two levels of grammatical variation—phonological and lexical—and how they influence memory for language. It is not well understood how these two levels of variation interact during lexical encoding. In the current study, we ask whether dialect-specific lexical variants trigger different patterns of accuracy in word recognition. Specifically, we ask whether there is an effect on accurate lexical recall of AAE words in an AAE voice pronunciation patterns associated with AAE, e.g. *th-stopping* (“dem” for *them*), *-ing/-in* variation (“walkin” for *walking*), *t/d* deletion, *th-fronting*: interdental fricative [θ] pronounced as a labiodental fricative [f] (“baf” for *bath*) (Green, 2002; Thomas, 2007), in addition to lexical variation such as *crib* instead of *house* (Smitherman, 2000).

However, a typical DRM-paradigm, where listeners hear words (e.g. *dream*, *bed* rest and awake) related to a lure (e.g. *sleep*), exposes participants to a word lists before testing lexical recall. This poses a problem for testing lexical variation of the kind here. Take the example of AAE *crib*, equivalent to MAE *house*. Out of context, it is not clear whether listeners will recruit the intended AAE usage of *crib* to mean *house*, or the MAE usage of *crib* to mean *baby's bed*. Therefore, we designed a modified word recognition task: target word items will be presented to listeners in a story, in order to provide sentential contexts in which each word can be perceived with the intended semantic meaning. Thus, we will test listeners' accuracy in recalling heard target words

(and correctly identifying semantic counterparts that were not heard) as a function of the socio-linguistic indexing of the speaker (AAE or MAE) and lexical usage (AAE or MAE).

Predictions

Previous research suggests that the phonological similarity between a token and previously experienced exemplars decreases the probability that listeners falsely recall non-presented words. For instance, using the DRM-paradigm Sumner & Kataoka (2013) find that listeners are less likely to accurately recall words for a New York English speaker (a less prestige variety) than for British English and General American speakers (more prestige varieties). Words spoken in General American are more frequent for General American participants than those uttered in various non-General American accents by any standard measure of global frequency. This leads to a more robust cluster of GA episodes. Though episodes of British English are sparse for General American participants, words uttered in this socially prestigious dialect have robust representations. The social prestige of BE has a more robust influence on processing than NY English, among General American and New York English participants. Therefore, we might predict that listeners will be less accurate at recalling words heard by an AAE speaker, since it is a less prestigious variety.

However, from an episodic approach, an interaction between lexical type and speaker accent is expected. We predict that for AAE-specific lexical items, exemplars should be stored with AAE phonetic detail. Upon hearing the AAE lexical item, previously experienced exemplars should be activated. This will increase encoding strength, which should lead to a benefit in accurate remembrance of AAE words. Participants exposed to the AAE talker should have higher rates of veridical word recognition. By this same cognitive mechanism, we also predict that the previously experienced exemplars of AAE lexical items should be dissimilar to the token produced by the Mainstream American English (MAE) talker. This dissimilarity will decrease the likelihood of activation for AAE lexical items, and weaken encoding strength. Participants exposed to the MAE voice should have lower rates of accurate recognition and higher rates of false recognition for AAE lexical items, but accurate remembrance of MAE words is probable.

Methods

Materials and Stimuli

As critical items, we selected 22 pairs of words that had the same semantic meaning and syntactic category. Each pair consisted of a AAE lexical item and a MAE lexical item counterpart. Items are provided in Table 1. Across the two lexical conditions, we balanced the items for number of syllables and morphological complexity, to the extent possible. Across word lists, similar numbers of monosyllabic target words (AAE=13 items, MAE=14 items) and bisyllabic (AAE=9 items, MAE=8 items) appeared.

A story was composed, containing slots for one word from each of the 22 pairs. Four versions of the story were created, where half of each of the 22 target words was realized as either an AAE or an MAE variant (e.g., “They go back to the [hood]/[city]”). In each version 11 AAE and 11 MAE target word variants occurred, word parings counterbalanced across the 4 versions. Due to an oversight, the word “dope” occurred twice in the AAE wordlist.

Two male native speakers of English, one a native AAE speaker and one a native MAE speaker, produced each story version. The speakers were recorded in a sound-attenuated booth with a head-mounted microphone at a sampling rate of 441kHz. Story productions were similar in length and speaking rate. Recordings were amplitude-normalized.

Table 1: MAE and AAE word pairs.

MAE word	AAE word
<i>car</i>	<i>whip</i>
<i>cash</i>	<i>moolah</i>
<i>city</i>	<i>hood</i>
<i>awesome</i>	<i>dope</i>
<i>cheating</i>	<i>trifling</i>
<i>crazy</i>	<i>cray</i>
<i>food</i>	<i>grub</i>
<i>friends</i>	<i>homies</i>
<i>on point</i>	<i>on fleek</i>
<i>shoes</i>	<i>kicks</i>
<i>jewelry</i>	<i>bling</i>
<i>weird</i>	<i>whack</i>
<i>fancy</i>	<i>banging</i>
<i>girlfriend</i>	<i>boo</i>
<i>group</i>	<i>crew</i>
<i>girl</i>	<i>shorty</i>
<i>leave</i>	<i>dip</i>
<i>working</i>	<i>grinding</i>
<i>house</i>	<i>crib</i>
<i>style</i>	<i>swag</i>
<i>call</i>	<i>holla</i>
<i>great</i>	<i>dope</i>

We aimed to quantify the amount of AAE specific variables. A phonetically trained coder tabulated the presence of specific phonological variables known to correlate with AAE including high rates of [-m] (77.77% occurrence in conditioning contexts), postvocalic r-dropping (37.5%), t/d deletion (55%), and th-stopping (68.75%); the MAE speaker has close to zero proportional usage of these variables.

Participants and Procedure

124 participants completed the study, native speakers of American English. These subjects, UC Davis undergraduate students who reported normal vision and hearing, received partial course credit for their participation.

First, participants heard one of the story recordings (65 heard one of the 4 story versions produced by the AAE

speaker; 59, heard one of the 4 story versions produced by the MAE speaker) in a sound-attenuated booth with headphones. Each participant was randomly assigned to a speaker and story version.

Following story exposure, participants completed a word recognition test. In the word recognition test, participants saw all 44 words (heard and unheard variants) randomized, one at a time, on a computer screen and indicated whether that word was heard in the recording or not via a button box response key press (using E-prime software and response box).

Results

Listener responses were coded for accuracy (1=correctly identified that they had heard or had not heard a lexical item, 0=incorrect). Since listeners had two response options, an aggregated pattern of chance performance would be 50%.

Recognition Accuracy

We first examined participants' accuracy for heard words (i.e., accuracy for heard words reflects "veridical" perception, following Roediger & McDermott, 1995). Responses to heard words were modeled using a logistic mixed-effects regression model (*lme4* package v.1.1-14; Bates et al., 2015) using the `glmer()` function in *lme4* using *R* (v.3.2.2). The model included two fixed-effect predictors. Word Type (AAE lexical item [baseline], MAE lexical item) tested whether listener recognition was influenced by the social dialect of the lexical item. Speaker (AAE speaker [baseline], MAE speaker) assessed whether the social dialect of the speaker influenced accurate lexical recognition of heard words. The two levels of each categorical variables were hand-assigned numerical weights, under the constraint that the sum of the weights equals to 0 (Davis, 2010). Thus, the model tests whether the observed difference in means between factors is within the sampling error of 0. In addition to these two main effects, the model also included the interaction between them to test the critical prediction that there will be differential recognition of AAE words in the AAE voice, relative to the MAE voice. The random effects structure of the model included random intercepts for participant and word as well as by-participant random slopes for Word Type (since Speaker was a between-subjects factor, random slopes for Speaker by participant were not permitted by the design).

Table 2: Output of logistic mixed effects model run on accuracy data for heard items.

	Est.	SE	<i>z</i>	<i>p</i>
(Intercept)	.48	.16	3	.002*
WordType	.33	.13	2.4	.01*
Speaker	.41	.11	3.5	.000*
WordType*Speaker	.18	.07	2.4	.01*

Table 2 provides the model output. Since the contrasts between levels within a predictor sum to zero, the intercept

can be interpreted as the estimated grand mean (mean of all conditions). Word Type significantly predicted accurate word recognition: overall, listeners recalled AAE words more accurately (61%) than MAE words (53%). Speaker was also a significant predictor: listeners who heard the AAE speaker remembered words better (64%) compared to listeners who heard the MAE speaker, who performed at chance (50%).

Crucially, a significant interaction between Speaker and Word Type, provided in Figure 1, reveals that listeners who heard the AAE voice were more accurate at recognizing AAE-specific lexical items (70%), relative to MAE lexical items (57%). Meanwhile, listeners who heard the MAE voice showed no difference between recalling AAE (51%) or MAE (48%) words, with both being at chance-level performance.

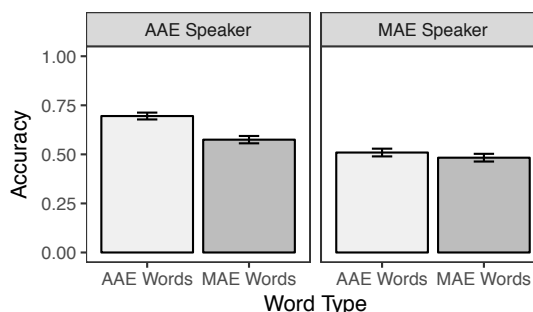


Figure 1: Mean recognition accuracy in recognizing heard words by Speaker and Word Type. Error bars depict standard errors of the mean.

False Recognition

Accuracy in rejecting non-heard lexical items was assessed, as well. We modeled responses to unheard lexical items using the procedure and model structure identical to that used in assessing heard word recognition. The output of the unheard data model is provided in Table 3. Figure 2 illustrates the performance for unheard words across conditions.

Table 3: Output of logistic mixed effects model run on accuracy data for unheard items.

	Est.	SE	<i>z</i>	<i>p</i>
(Intercept)	.22	.18	1.2	.2
WordType	.16	.17	.9	.3
Speaker	.27	.09	2.8	.004*
WordType*Speaker	.02	.07	.2	.7

The model assessing accuracy in correctly rejecting unheard lexical items revealed only a reliable simple main effect of Speaker: listeners were more accurate, overall, at correctly rejecting unheard lexical items in the AAE voice (57%), relative to at-chance performance for the MAE voice (48%).

Listeners who heard the AAE speaker were less likely to falsely recognize words, overall, than those who heard the MAE speaker.

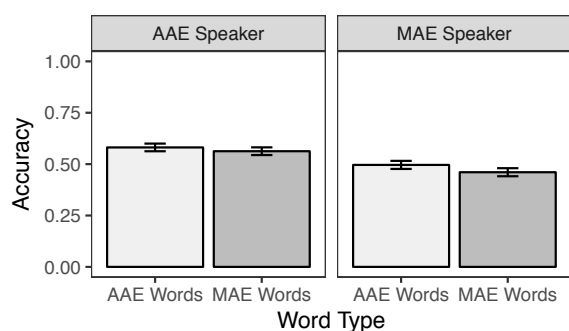


Figure 2: Mean recognition accuracy in identifying unheard words by Speaker and Word Type. Error bars depict standard errors of the mean.

Discussion

The current study was designed to investigate the effect of socially-conditioned phonological and lexical variation on word recall. We observed that listeners show more accurate veridical recall of words when they were pronounced with the constellation of phonological properties associated with African American English (AAE) varieties than when those words were pronounced with a Mainstream American English pronunciation. Additionally, we found an effect of socially-cued usage of lexical items on accurate word recall. AAE-specific lexical items were better recalled than MAE words. Furthermore, we found an interaction between these two main factors: AAE-specific lexical items were even better recalled when heard pronounced by an AAE voice. (Meanwhile, listeners who heard the MAE voice were at chance in recalling both AAE and MAE word types).

As outlined in the introduction, this study was strategically designed to address gaps in the literature on lexical recall. First, prior work has demonstrated that accurate word recall varies based on lexical and contextual factors. For example, words heard in greater-attention conditions are recalled more accurately (Otgaar et al., 2012). Here, we find that social factors influence accurate word recall. One interpretation of our finding of better word recall for the AAE voice is that some socio-linguistic conditions are more likely to elicit greater attention, leading to differential encoding and recall. This general observation supports prior work. For example, Sumner & Kataoka (2013) observed that social dialect influences word recall. However, the direction of the patterns between their study and the current one are different. They observed that words pronounced by British English and General American voices are similarly recalled, and that both elicit recall better than words pronounced by a New York English voice. They interpreted this difference as attention-based—listeners were more likely to have increased attention toward the BE voice (due to high social prestige) but divided attention toward the NY voice (due to negative social perceptions). However, we observe the opposite pattern:

words in the AAE voice was more accurately recalled, despite its lower prestige value, relative to MAE. Similar findings are reported by Szakay et al. (2016). In a study conducted with Maori L2 listeners, they find that hearing words in non-standard Maori English and Pakeha, a standard variant of New Zealand English, primed Maori, but that Maori English more robustly primed Maori words. Exploring which particular social factors are driving attention and how that mediates word recall are avenues for future work.

The present study also explored the interaction of lexical and phonological variation on word recall. AAE-lexical items were better recalled in the AAE voice, suggesting that linguistic recognition is stronger when socially-cued phonological and lexical variation are correlated. However, we observe that congruence between socially-conditioned phonological and lexical variation influences word recognition asymmetrically—there was not a parallel effect for MAE lexical items in the MAE voice. Greater attention to the AAE voice might explain this asymmetry; perhaps there was greater attention for AAE-specific words when pronounced in the AAE voice leading to more precise lexical recall. Meanwhile, there was seemingly less precise recall of the MAE voice overall. D’Onofrio’s (2016) discussion on confirmation bias, a concept from the field of social psychology, can also help explain the asymmetry. Confirmation bias is a view that listeners’ experiences trigger a schema that includes sociolinguistic expectancies for future experiences. Listeners pay more attention to experiences that fit this schema, while downplaying or ignoring those that contradict it. This schema modulates how we take in new episodes in addition to how we might incorrectly recognize others. Confirmation bias can be recruited to explain our observed interaction: hearing AAE lexical items that aligned with the AAE pronunciation lead to more precise recall and encoding of those words. The schema for the MAE talker is perhaps more general, thus no particular MAE or AAE word resonated within it. As mentioned earlier, though, the social mechanisms mediating attention, and how this might be influenced by social expectations, are areas for future work.

Furthermore, prior work has outlined two approaches to lexical memory: single or dual process lexical recall. Our finding socially-conditioned speaker- and lexical item-specific factors on word recall support a dual model. The gist trace in the dual-process stores themes of the information processed, and social information is one of the matrixes of thematic information used in the memory reconstruction process. Also, as mentioned in the introduction, models of lexical representation and encoding have demonstrated the role of experience with words and pronunciations in influencing production and perception. For example, listeners recognize a word heard in the same voice faster and more accurately than when the word is pronounced by a different speaker (Mullinex et al., 1989). However, it is not clear what the relationship between facts about representation and models of word recall is. While there is a large body of work about the detail stored in lexical representation (Pierrehumbert 2002; Johnson, 2006), how that interacts with

lexical recall needs to be further explored, as well. Whether the main finding that lexical items are better remembered if they are consistent with that accent is a broader linguistic pattern is beyond the scope of this paper. Further work using more speakers and other speaker groups would need to be done to show that this finding is robust and consistent.

Finally, our findings connect to issues above and beyond linguistic representations, memory, and speech comprehension. In educational, legal and judicial contexts, AAE is stigmatized, considered incomprehensible, and lacking intelligibility (Rickford & King, 2016). Our findings are seemingly contradictory to these ideologies. Not only do listeners correctly comprehend and recall AAE lexical items, but they do so *even better* when they are heard in an AAE voice. This raises larger issues about bias and social expectations and their role in language ideologies.

Acknowledgments

The authors would like to acknowledge John Rickford, Meredith Tamminga, Sharese King, Simon Todd, and participants of the 2018 LSA annual meeting in Salt Lake City for their feedback on this study.

References

- Bates, D., Mächler, M., Bolker, B., & Walker, S. (2014). Fitting linear mixed-effects models using lme4. *arXiv preprint arXiv:1406.5823*.
- Bauer, L. M., Olheiser, E. L., Altarriba, J., & Landi, N. (2009). Word type effects in false recall: Concrete, abstract, and emotion word critical lures. *The American journal of psychology*, 469-481.
- Brainerd, C. J., & Reyna, V. F. (2002). Fuzzy-trace theory and false memory. *Current Directions in Psychological Science*, 11(5), 164-169.
- Brainerd, C. J., & Reyna, V. F. (2012). Reliability of children's testimony in the era of developmental reversals. *Developmental Review*, 32(3), 224-267.
- Davis, M. J. (2010). Contrast coding in multiple regression analysis: Strengths, weaknesses, and utility of popular coding structures. *Journal of Data Science*, 8(1), 61-73.
- D'Onofrio, A. (2016). *Social Meaning in Linguistic Perception* (Doctoral dissertation, Stanford University).
- Casasanto, L. S. (2008, January). Does social information influence sentence processing? In *Proceedings of the Annual Meeting of the Cognitive Science Society* (Vol. 30, No. 30).
- Green, L. J. (2002). *African American English: a linguistic introduction*. Cambridge University Press.
- Johnson, K. (2006). Resonance in an exemplar-based lexicon: The emergence of social identity and phonology. *Journal of phonetics*, 34(4), 485-499.
- King, S., & Sumner, M. (2014, January). Voices and variants: Effects of voice on the form-based processing of words with different phonological variants. In *Proceedings of the Annual Meeting of the Cognitive Science Society* (Vol. 36, No. 36).
- McGowan, K. B. (2015). Social expectation improves speech perception in noise. *Language and Speech*, 58(4), 502-521.
- Mullennix, J. W., Pisoni, D. B., & Martin, C. S. (1989). Some effects of talker variability on spoken word recognition. *The Journal of the Acoustical Society of America*, 85(1), 365-378.
- Otgaar, H., Peters, M., & Howe, M. L. (2012). Dividing attention lowers children's but increases adults' false memories. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 38(1), 204.
- Pérez-Mata, M. N., Read, J. D., & Diges, M. (2002). Effects of divided attention and word concreteness on correct recall and false memory reports. *Memory*, 10(3), 161-177.
- Pierrehumbert, J. B. (2002). Word-specific phonetics. *Laboratory phonology* 7, 4(1), 101.
- Rickford, J. R., & King, S. (2016). Language and linguistics on trial: Hearing Rachel Jeantel (and other vernacular speakers) in the courtroom and beyond. *Language*, 92(4), 948-988.
- Robinson, K. J., & Roediger III, H. L. (1997). Associative processes in false recall and false recognition. *Psychological Science*, 8(3), 231-237.
- Roediger, H. L., & McDermott, K. B. (1995). Creating false memories: Remembering words not presented in lists. *Journal of experimental psychology: Learning, Memory, and Cognition*, 21(4), 803.
- Smitherman, G. (1977). Talking and testifying. *The Language of Black America*. Boston: Houghton.
- Strand, E. A., & Johnson, K. (1996, October). Gradient and visual speaker normalization in the perception of fricatives. In *KONVENS* (pp. 14-26).
- Sumner, M., & Samuel, A. G. (2009). The effect of experience on the perception and representation of dialect variants. *Journal of Memory and Language*, 60(4), 487-501.
- Sumner, M., & Kataoka, R. (2013). Effects of phonetically-cued talker variation on semantic encoding. *The Journal of the Acoustical Society of America*, 134(6), EL485-EL491.
- Sumner, M., Kim, S. K., King, E., & McGowan, K. B. (2014). The socially weighted encoding of spoken words: a dual-route approach to speech perception. *Frontiers in psychology*, 4, 1015.
- Szakay, A., Babel, M., & King, J. (2016). Social categories are shared across bilinguals' lexicons. *Journal of Phonetics*, 59, 92-109.
- Thomas, E. R. (2007). Phonological and phonetic characteristics of African American vernacular English. *Language and Linguistics Compass*, 1(5), 450-475.