

Does top-down information about speaker age guise influence perceptual compensation for coarticulatory /u/-fronting?

Georgia Zellou (gzellou@ucdavis.edu)

Department of Linguistics, Phonetics Lab, UC Davis, 1 Shields Avenue
Davis, CA 95616 USA

Michelle Cohn (mdcohn@ucdavis.edu)

Department of Linguistics, Phonetics Lab, UC Davis, 1 Shields Avenue
Davis, CA 95616 USA

Aleese Block (asblock@ucdavis.edu)

Department of Linguistics, Phonetics Lab, UC Davis, 1 Shields Avenue
Davis, CA 95616 USA

Abstract

The current study explores whether the top-down influence of speaker age guise influences patterns of *compensation for coarticulation*. /u/-fronting variation in California is linked to both phonetic and social factors: /u/ in alveolar contexts is fronter than in bilabial contexts and /u/-fronting is more advanced in younger speakers. We investigate whether the *apparent age* of the speaker, via a guise depicting a 21-year-old woman or a 55-year-old woman, influences whether listeners compensate for coarticulation on /u/. Listeners performed a paired discrimination task of /u/ with a raised F2 (fronted) in an alveolar consonant context (/sut/), compared to non-fronted /u/ in a non-coronal context. Overall, discrimination was more veridical for the younger guise, than for the older guise, leading to the perception of more inherently fronted variants for the younger talker. Results indicate that apparent talker age may influence perception of /u/-fronting, but not only in coarticulatory contexts.

Keywords: speech perception; u-fronting; compensation for coarticulation; apparent speaker age

Introduction

Knowledge about the connection between social properties of speakers and their speech patterns interacts with the sound-to-meaning mapping (e.g., Sumner et al., 2013). Coarticulation, the overlapping of adjacent segments in speech, has been understudied with respect to this phenomenon. There is work in speech production demonstrating that amount of coarticulatory variation is connected to regional, social, diachronic, and stylistic speech patterns (Harrington et al., 2008; Kataoka, 2011; Scarborough & Zellou, 2013; Scarborough et al., 2015), suggesting that speakers can learn to use coarticulation in socially meaningful ways. In the current study, we ask whether listeners are sensitive to differences in expected coarticulatory patterns across social groups and use this information when compensating for coarticulation.

One method for exploring listeners' knowledge about the relationship between speaker groups and speech patterns is by presenting listeners with the same stimuli, but varying the speaker guise. For example, Niedzielski (1999) found that the apparent regional status of a talker can influence vowel categorization: participants' task was to match tokens with "Canadian raised" vowels (i.e., [aɪ] produced as [ʌɪ]) produced by a Detroit speaker to synthetic raised or canonical vowels. Half were told the speaker was from Detroit, while the other half were told speaker was from Canada. Listeners classified the apparent-Canadian as (accurately) producing the raised vowel variant; yet, they perceived the apparent-Michigander's vowels as reflecting the canonical standard American English pronunciation. In other words, the same sound was categorized differently depending on the social information provided about the talker: the Detroit participants reported awareness that vowel raising is a Canadian, not a Michigan, feature and this guided their low-level perceptual categorization of the same stimuli. Other types of explicit social information elicit this top-down influence on speech perception, such as photographs depicting speakers of various ages (Hay, Warren, & Drager, 2006) and stuffed animal toys referencing dialect regions (Hay & Drager, 2010).

Recent speech perception theories propose mechanisms to account for the influence of social information on the sound-to-meaning mapping. For example, Sumner et al. (2013) suggest that spoken word recognition operates in parallel with social representational mapping and interactivity between these processes can modulate linguistic comprehension. Pierrehumbert (2002) 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; this distribution then shapes subsequent perception of sounds and experiences. A fuller understanding of the role social information plays in the simultaneous mapping of sounds and

social categories is needed to inform our understanding of linguistic representations and speech comprehension.

Compensation for coarticulation

One area that remains understudied with respect to top-down effects of speaker social information is compensation for coarticulation, or segmental overlap of adjacent speech sounds. Compensation for coarticulation is a perceptual reduction, or elimination, of context-specific acoustic variation. For instance, a nasalized vowel in isolation is perceived as nasal, but when the same nasalized vowel occurs adjacent to nasal consonants (in a [m_m] frame) it is not perceived as nasal (Kawasaki, 1986). As soon as the putative source of coarticulation is heard, the acoustic variation is attributed to that source as the listener decides how the speech signal should be parsed into discrete underlying units. Coarticulatory compensation, thus, can be seen as a useful perceptual process since it provides listeners with ways of adjusting to systematic variation within the speech signal and identifying invariant units (Beddor et al., 2013; Zellou & Dahan, 2019). At the same time, compensation for coarticulation ascribes context-sensitive acoustic information to its source, making veridical acoustic perception more challenging (cf. Kawasaki, 1986). Furthermore, native-language experience influences the degree to which a listener compensates for coarticulatory detail present on a vowel: native Shona speakers, who produce less extensive coarticulation in Shona than in English, compensate less for English vowel-to-vowel coarticulation than English speakers (Beddor et al., 2002). In other words, listeners use their language-specific learned coarticulatory structures to parse the speech signal.

In the realm of speech production, coarticulation was traditionally viewed as the invariant, physiological connection between the speech signal and abstract phonemes and, therefore, having no relevance to the linguistic grammar. Yet, few studies have explored the *social* impact on compensation for coarticulation. In an eye tracking experiment, Coetzee et al. (2019) measured fixations toward visually presented CVC–CVN(C) pairs while participants heard the words produced by speakers of two varieties of Afrikaans: one variety that uses greater nasal coarticulation, and one that uses less. They found that listeners display learning for the Afrikaans dialect with greater anticipatory nasal coarticulation with earlier fixations. This work suggests that social knowledge about a given speaker's dialect can shape our ability to extract and learn patterns from their speech. Yet no prior work, to our knowledge, has examined whether manipulating top-down apparent social characteristics might impact listeners' compensation for coarticulation.

Examining these top-down influences can additionally inform theories of speech perception as to how listeners integrate social information, specifically during compensation for coarticulation. For example, evidence of perceptual compensation for non-linguistic stimuli and in non-human species has been used to argue for domain-

general auditory mechanisms in speech perception (Diehl et al., 2004); yet, aside from linguistic experience, this 'general approach' does not make explicit claims about how — or whether — other speaker-indexical information might be integrated. One possibility is that we might not expect listeners to show differences in perceptual compensation based on top-down social guise, given identical stimulus items, where listeners are relying on general auditory mapping mechanisms. Variationist accounts of coarticulation, on the other hand, have pointed to evidence that coarticulatory patterns are highly variable, both across and within speakers (e.g., Solé, 1992; Scarborough & Zellou, 2013; Zellou & Tamminga, 2014). That language-specific patterns of perceptual compensation are linked to produced coarticulatory extent means that phonetic representations can be rich enough to encode coarticulatory patterns via linguistic experience (Beddor & Krakow, 1999; Beddor et al., 2002). Furthermore, decades of empirical phonetic research establish that coarticulatory patterns can vary across languages (e.g., Beddor et al., 2002), across regional varieties of a language (Tamminga & Zellou, 2015), and across generations of speakers within one variety (Zellou & Tamminga, 2014; Harrington et al., 2008). In other words, there is evidence to support the view that coarticulatory structures are encoded in the grammatical system of a language and socially learned.

Vowel variation is linked to both social groups (Labov et al., 2006) and context-specific influences (Farnetani & Recasens, 1997). Variation in high back vowel fronting in California English, in particular, has been correlated with both social and phonetic factors. For one, the California Vowel Shift (CVS) is an ongoing sociolinguistic sound change; the most salient aspect of this shift includes fronting of the back rounded vowels. California /u/-fronting is most advanced in younger speakers' productions, suggesting that there are categorically different /u/ targets across generations (Hall-Lew, 2011). Also, /u/ before alveolar consonants is fronter than in bilabial contexts, due to coarticulation (Kataoka, 2011). Because /u/-fronting is linked to both social and coarticulatory factors, there can be ambiguity for listeners as to the source of a higher F2. While these two studies did not find an interaction between age and context on /u/-fronting in California speakers' productions (Hall-Lew, 2011; Kataoka, 2011), it is possible that listeners may still differentially apply social knowledge to these productions (here, resulting in differences in compensation for coarticulation). To test this question in the present study, we hold the acoustic information constant and vary the top-down guise (as a 'younger' or 'older' speaker).

Current study

In the current study, we explore whether the mapping of coarticulatory variation to linguistic structure is guided by a listener's *expectations* about a speaker's coarticulatory patterns based on apparent speaker age. In particular, we test whether a guise of an apparent older versus younger adult influences Californian listeners' compensation for

coarticulation for alveolar codas on /u/-fronting. As mentioned earlier, California /u/-fronting is most advanced in younger speakers, than older speakers (Hall-Lew, 2011). And, independently, /u/ before alveolar consonants is fronter due to coarticulation (Kataoka, 2011). It is plausible that listeners have formed social indices for different coarticulatory patterns based on experience with speaker age groups whose coarticulatory distributions vary (Harrington et al., 2008). Thus, we hypothesize that listeners use these representations to compensate differently depending on explicit social information provided about the talker. In particular, listeners might expect a younger speaker to have phonologized /u/-fronting, whereas they might expect older speakers' fronted /u/ to be the result of coarticulatory influences from an anterior consonant. Therefore, we predict listeners will be less likely to attribute a higher F2 to the articulatory effects of subsequent alveolar consonant (e.g. /sut/ "suit") when told the speaker is a younger adult, compared to an apparent older adult speaker.

To test compensation for coarticulation, we used a paired vowel discrimination paradigm, which has been used in prior work to examine perceptual compensation for coarticulation (Beddor & Krakow, 1999). In a paired discrimination task, listeners hear two pairs of words containing fronted or backed vowels in alveolar and bilabial consonant contexts. Participants' task is to indicate which pair contains vowels that sound the most different. We predict that the apparent age of the speaker will influence whether listeners compensate for coarticulation on /u/, leading to the veridical perception of more fronted variants for younger talkers. Another potential outcome is that social information does not influence the perception of coarticulation and there will be no difference in perceptual compensation behavior as a function of speaker age guise.

Methods

Participants

85 native English-speaking undergraduates (60 females, 25 males; mean age = 21.1 ± 2.3 years old) were recruited from the UC Davis Psychology subject pool. All participants indicated that they were from California. They participated in the experiment on online platform (Qualtrics) and received course credit for their participation. Sample size was calculated based on a power analysis in G*Power (Faul et al., 2007) assuming power of .95, 3 predictors, and a small-to-medium effect size.

Stimuli

Stimuli materials were created by eliciting 'soup' [sup] (non-coronal) and 'suit' [sut] (coronal) from a native Californian female. After vowels were extracted, F2 was manipulated ± 80 Hz to create *fronted* and *backed* versions of each vowel using the VocalToolkit package (Corretge, 2012) in Praat (Boersma & Weenink, 1996). For all vowels, f0 was controlled to have a smoothed falling f0 contour (from 225 Hz to 200 Hz, decreasing linearly from the start to the end of

the vowel). Vowel durations were also normalized to 150 ms, and stimuli were normalized in intensity (50 dB). The vowels were then spliced into the original word context.

Procedure

Participants completed a 4-interval forced-choice (4IAX) paired discrimination paradigm (Beddor & Krakow, 1999) to assess their ability to discriminate between fronted and backed versions of /u/ in coronal and non-coronal consonantal contexts. For each trial, two pairs of stimuli were presented to a listener: one pair contained acoustically identical vowels and the other pair contained acoustically different vowels (i.e., fronted vs. backed). Listeners were instructed to decide which pair contained *different* sounding vowels (two possible responses: first or second pair). Participants completed two types of trials (randomly presented), which varied the consonantal context: 'same context' (control) trials and 'different context' (test) trials.

Same context (control) trials contained identical consonant contexts across pairs, all non-coronal [sup] or coronal [sut]. In each set of stimuli, one pair contained vowels differing in backing, while the others were acoustically identical (both fronted or backed). For example: Pair 1 [sup_{Back}] [sup_{Back}] vs. Pair 2 [**sup_{Back}**] [**sup_{Front}**] (bolded pair contains acoustically distinct vowels). Order of differing vowels within and across pairs was counterbalanced across trials. We expect veridical acoustic vowel perception to be the highest in control trials, because when consonantal context is identical, compensation should equally attribute (or equally *not* attribute) acoustic variation to the source across vowels (Beddor et al., 2002; Zellou et al., 2020).

Different context (test) trials contained varying coda contexts: each pair of words included a non-coronal [sup] and coronal [sut] (order of [sut]/[sup] counterbalanced). One pair within the trial had acoustically identical vowels (e.g., both backed or both fronted); the other pair was had different vowels always in the direction of interest, with the fronted /u/ in a coronal context: [sut_{Front}], [sup_{Back}]. For example: Pair 1 [**sut_{Front}**] [**sup_{Back}**] vs. Pair 2 [sut_{Front}] [sup_{Front}] (bolded pair contains acoustically distinct vowels). Ordering within and across vowel pairs was counterbalanced across trials. If compensation occurs, listeners should hear the fronted vowel in the alveolar consonant context as the 'same' as the backed vowel, reflecting less veridical acoustic perception. Thus, the vowel in [sut_{Front}] might sound more similar to the vowel in [sup_{Back}] if compensation occurs.

Subjects heard the same set of stimuli across two blocks varying in apparent age of the speaker: one block presented a younger speaker guise, the other, an older-speaker guise (Figure 1). Participants were given both the apparent name and age of the speaker (e.g. "You will hear Linda, a 55-year-old woman, producing two pairs of words"). Participants' exposure to the older and younger guises was counterbalanced across subjects (e.g., "Linda" first, "Madison" second).

In total, participants completed 32 trials (randomly presented across test and control conditions). As the

experiments was conducted online, we included listening four comprehension questions interspersed throughout the experiment (e.g., “Who is older, Linda or Madison?”); participants’ data was retained only if they answered all four questions correctly ($n=85$).

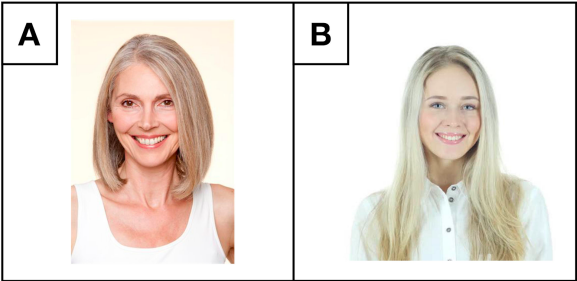


Figure 1: Stock images used for the apparent age guises. Older age guise: (panel A) 55-year-old “Linda” and Younger age guise: (panel B) 21-year-old “Madison”.

Predictions

Overall, we expect a **main effect of Context** on listeners’ accuracy in identifying acoustically distinct vowels. Listeners’ performance should be highest in the Same Context condition, relative to the Different Context conditions, where compensation for coarticulation from the differing coda consonant will make veridical vowel perception more challenging.

We also critically expect an **interaction between Context and Age Guise**: if listeners’ expectations about age-related differences in produced coarticulation guide their perceptual compensation, we predict different patterns of veridical perception in Different word context conditions. In particular, we expect less compensation, i.e., more veridical perception, for the younger age guise. Therefore, we predict that listeners will be more accurate in identifying acoustically distinct vowels in different-word pairs when given the younger age guise, relative to when these trials are presented with the older age guise.

Statistical analyses

We coded listeners’ responses as binomial data: if they selected the trial with acoustically different vowels (e.g., [sut_{Front}] [sut_{Back}], [sut_{Front}] [sup_{Back}], etc.) (=1) or not (=0). We analyzed these responses with a mixed effects logistic regression (*lme4* R package; Bates et al., 2015). Main effects included Age Guise (Older, Younger), Consonantal Context (2 levels: ‘Same context’ (i.e., control), ‘Different context’ (i.e., test)), and their interaction. Random effects included by-Subject random intercepts and by-Subject random slopes for Guise and Consonantal Context (and their interaction). Contrasts were sum coded. (*lmer* syntax: Response ~ Guise * Context + (1 + Guise*Context | Subject).)

Results

Figure 2 shows the proportion of trials where participants discriminated the acoustically distinct vowels. The model output for the logistic regression is provided in Table 1. As seen in Figure 2, higher values indicate more acoustic (i.e., more veridical) discrimination. First, we observed a main effect of Consonant Context: participants showed more veridical vowel perception when the differing vowels occurred in the same word context (e.g., Pair 1 [sup_{Back}] [sup_{Back}] vs. Pair 2 [sup_{Back}] [sup_{Front}]), relative to when they occurred in different contexts. For the different-context condition (right panel), mean values closer to chance performance (0.50) indicate greater compensation for coarticulation, while higher values indicate greater veridical vowel perception (indicating failure to compensate fully). There was also a main effect of Age Guise: vowel discrimination in the older guise (i.e., “Linda”) was lower than for the younger guise (i.e., “Madison”), indicating greater overall veridical vowel perception. No interaction between Consonant Context and Age Guise was observed.

Vowel discrimination by Guise

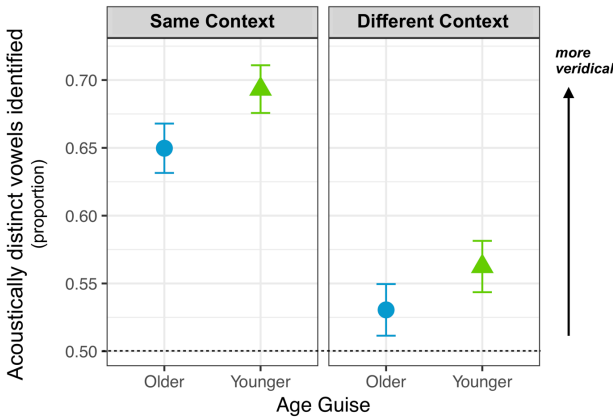


Figure 2: Mean proportion of acoustically distinct vowels identified (error bars show standard error) by Age Guise (Younger, Older) and Consonant Context (Same word or Different words). Chance performance (0.50) is indicated with a dotted line.

Table 1. Model Output.

	Coef.	SE	<i>z</i>	<i>p</i>
(Intercept)	0.46	0.05	8.59	<0.001***
Age Guise(older)	-0.09	0.04	-2.19	0.03 *
ConsContext(diff)	-0.28	0.04	-6.74	<0.001***
Age(older)*	0.02	0.04	0.57	0.57
ConsContext(diff)				
Num. observations = 2,752, Num. subs = 85				

Discussion

The current study examined top-down effects of apparent speaker age (either younger or older adult) on university-aged Californians’ social associations of fronted /u/ in the context

of alveolar codas via perceptual compensation for coarticulation. Listeners discriminated /u/ vowels in the word “suit” [sut] with differing F2 values (fronter or backer) in word pairs with same (all “suit”) or varying consonantal contexts (“suit” vs. “soup”). Overall, the consonantal context had an effect on patterns of perceptual compensation for coarticulation, whereby identical consonant contexts lead to more veridical vowel perception, relative to when vowels occurred in differing consonant contexts. Varying consonant context across word pairs in a trial makes veridical acoustic vowel perception more difficult due to coarticulatory compensation: a raised F2 on /u/ adjacent to an alveolar consonant might be factored out, making that vowel sound more similar to a backed /u/ in another context. Listeners’ above-chance performance in vowel discrimination for different-word contexts reflects partial compensation, in general, replicating the phenomenon of a failure to fully compensate for coarticulation that has been observed across numerous studies (Beddor & Krakow, 1999; Beddor et al., 2002; Zellou, 2017).

Additionally, we observed differences in overall vowel discrimination performance based on the apparent age guise of the speaker. While the stimuli were produced by a California native in her 20s, participants were given two different age guises: in one block, they were given an image of the speaker depicting a woman in her 20s, in the other block, they were given an image of the speaker depicting a woman in her 50s. We observed differences according to age guise: listeners displayed more veridical perception of acoustic differences in /u/ for the apparent young adult speaker guise, relative to the apparent older adult speaker guise. We can interpret our finding that listeners displayed less veridical vowel perception for the older speaker guise, than when they were given the younger speaker guise, as reflecting the effect of social information on distributions of experienced vowel patterns, which subsequently guide patterns of acoustic vowel perception (cf. Walker & Hay, 2011). The younger speaker guise ostensibly recruited activation of the pronunciation features associated with this accent, which included more advanced and phonologized /u/-fronting overall (Hall-Lew, 2011). The expectation for phonologization of /u/-fronting in this younger adult speaker guise would lead listeners to attribute more of the raised F2 (fronting) to the vowel. On the other hand, the older adult speaker guise lead to the expectation that /u/-fronting was less phonologized. So, listeners were less accurate in identifying the veridical signal, with a fronted /u/, given that top-down expectation. The current findings are in line with Hay, Warren, and Drager (2006), who found that apparent-speaker age and socio-economic status of speaker (signaled by photographs of people in various guises), influence how New Zealand listeners classify “near” and “square” vowels. Participants made more errors when viewing the younger guise, assuming it was a merged speaker, while fewer errors were made while viewing the older guise, indicating that they expected this speaker to be non-merged. In the current study, as in Hay et al., we observed that listeners given different

social information about acoustically identical stimuli displayed evidence of different phonetic and phonological interpretations of those sounds. In other words, we find further support that social knowledge can influence how listeners perceive variation.

However, we do not find that social knowledge influences patterns of compensation for coarticulation, as we predicted. Crucially, we did not observe an interaction between Context and Guise. While overall vowel perception was more accurate for the younger guise, listeners did not display show even higher performance in different-word contexts (relative to the same-word context) for this guise. In other words, listeners did not display differences in perceptual compensation based on age guise. There are several possibilities for this finding. For one, Hall-Lew (2011) did not find differences across ages in patterns of /u/-fronting based on consonant context. Therefore, it is possible that listeners *are* attending to variation in coarticulation across social groups, but this is not at play for these social groups for this acoustic feature. Future work looking at both production and perception in variation of coarticulation, *in tandem*, can identify social factors that might be relevant to listeners’ expectations about coarticulatory patterns. Additionally, all of our stimuli contained an initial coronal consonant (soup, suit), which may have triggered compensation across the board, weakening any potential interaction effects.

Furthermore, as previously mentioned, the stimuli were produced by a speaker in her 20s. Thus, the older speaker guise was mismatching in apparent and real voice age, whereas the younger speaker guise was not. The mismatch could have also led to the decrease in performance for the older speaker guise. Moreover, participants were also university-aged and may have more experience with the speech patterns of younger speakers, which could explain their higher performance for the younger guise. This is in line with Niedzielski (1999): listeners displayed more veridical perception of their social out-group, relative to their social in-group). Future work crossing an older voice with younger and older speaker guise could address these confounds.

Overall, the current results extend previous work on sociophonetic variation in perception by examining how social information influences listeners’ vowel perception (Niedzielski, 1999) and perceptual evaluation of coarticulation (cf. Coetzee et al., 2019). Perception of fine-grained acoustic vowel patterns is susceptible to the influence of apparent-speaker guise. It is still an open question of how coarticulatory detail is influenced by top-down knowledge. Nevertheless, future investigations and discussions of how social knowledge affects the perception of phonetic variation should explore coarticulatory and compensatory facts, as well.

Finally, we can speculate about the implications of the current findings for sound change. The influence of social factors on speech perception and how they might interact with linguistic change is beginning to be explored. In a set of studies, Warren (2005) investigated the ongoing merger of

the vowels in “near” and “square” in New Zealand English. For example, older speakers produce “square” raising only in contexts where the phonetic environment would make that natural, while younger speakers exhibit “square” raising in all contexts, suggesting that this vowel change has been fully reanalyzed for younger speakers (Warren, 2005).

These different interpretations of the ‘same sound’ (i.e., phoneme) mean that listeners might have differences in how the experience is encoded in linguistic memory based on the social information in that context (Pierrehumbert, 2002; Sumner et al., 2013). Both Hay et al. (2006) and findings from the current study suggest that exploring the role of social information in the perception of linguistic variation has important implications for understanding sound change. The phenomenon of different apparent social characteristics yielding different perceptual experiences can be a starting point for understanding the conditions under which sound change might occur. Ultimately, the current study indicates that looking at how the interaction of how coarticulation and sociolinguistic knowledge influences speech perception is a promising scientific area of research and opens many possible directions for future work.

References

- Bates, D., Maechler, M., Bolker, B., Walker, S., Christensen, R. H. B., Singmann, H., ... & Bolker, M. B. (2015). Package ‘lme4’. *Convergence*, 12(1), 2.
- Beddor, P. S., & Krakow, R. A. (1999). Perception of coarticulatory nasalization by speakers of English and Thai: Evidence for partial compensation. *The Journal of the Acoustical Society of America*, 106(5), 2868-2887.
- Beddor, P. S., Harnsberger, J. D., & Lindemann, S. (2002). Language-specific patterns of vowel-to-vowel coarticulation: Acoustic structures and their perceptual correlates. *Journal of Phonetics*, 30(4), 591-627.
- Beddor, P. S., McGowan, K. B., Boland, J. E., Coetzee, A. W., & Brasher, A. (2013). The time course of perception of coarticulation. *The Journal of the Acoustical Society of America*, 133(4), 2350-2366.
- Boersma, P. P. G., & Weenink, D. J. M. (1996). *Praat: Doing Phonetics by Computer: Version 3.4*. Instituut voor Fonetische Wetenschappen.
- Coetzee, A. W., Beddor, P. S., Styler, W., Tobin, S., Bekker, I., & Wissing, D. (2019). Producing and perceiving socially indexed coarticulation in Afrikaans. *Proceedings of ICPHS*.
- Corrette, R. (2012). *Praat vocal toolkit*. Barcelona, Spain: Praat. Retrieved from <http://praatvocaltoolkit.com>.
- Diehl, R. L., Lotto, A. J., & Holt, L. L. (2004). Speech perception. *Annu. Rev. Psychol.*, 55, 149-179.
- Elman, J. L., & McClelland, J. L. (1988). Cognitive penetration of the mechanisms of perception: Compensation for coarticulation of lexically restored phonemes. *Journal of Memory and Language*, 27(2), 143-165.
- Farnetani, E., & Recasens, D. (1997). Coarticulation and connected speech processes. *The handbook of phonetic sciences*, 371-404.
- Faul, F., Erdfelder, E., Lang, A. G., & Buchner, A. (2007). G* Power 3: A flexible statistical power analysis program for the social, behavioral, and biomedical sciences. *Behavior research methods*, 39(2), 175-191.
- Hall-Lew, L. (2011, August). The Completion of a Sound Change in California English. *Proceedings of ICPHS* (pp. 807-810).
- Harrington, J., Kleber, F., & Reubold, U. (2008). Compensation for coarticulation, /u/-fronting, and sound change in standard southern British: An acoustic and perceptual study. *The Journal of the Acoustical Society of America*, 123(5), 2825-2835.
- Hay, J., Warren, P., & Drager, K. (2006). Factors influencing speech perception in the context of a merger-in-progress. *Journal of Phonetics*, 34(4), 458-484.
- Hay, J. & Drager, K. (2010). Stuffed toys and speech perception. *Linguistics*, 48(4), pp. 865-892.
- Johnson, K., Strand, E. A., & D'Imperio, M. (1999). Auditory-visual integration of talker gender in vowel perception. *Journal of phonetics*, 27(4), 359-384.
- Kataoka, R. (2011). *Phonetic and cognitive bases of sound change* (Doctoral dissertation). UC, Berkeley.
- Kawasaki, H. (1986). Phonetic explanation for phonological universals: The case of distinctive vowel nasalization. In JJ Ohala, & Jaeger, JJ (Eds.), *Experimental phonology* (pp. 239-252).
- Labov, W., Ash, S., & Boberg, C. (2006). *Atlas of North American English: Phonetics. Phonology And Sound Change*, Mouton de Gruyter.
- Niedzielski, N. (1999). The effect of social information on the perception of sociolinguistic variables. *Journal of Language and Social Psychology*, 18(1), 62-85.
- Ohala, J. J. (1993). Coarticulation and phonology. *Language and speech*, 36(2-3), 155-170.
- Pierrehumbert, J. (2002). Word-specific phonetics. *Laboratory Phonology*, 7.
- Scarborough, R., & Zellou, G. (2013). Clarity in communication: “Clear” speech authenticity and lexical neighborhood density effects in speech production and perception. *The Journal of the Acoustical Society of America*, 134(5), 3793-3807.
- Scarborough, R., Zellou, G., Mirzayan, A., & Rood, D. S. (2015). Phonetic and phonological patterns of nasality in Lakota vowels. *Journal of the International Phonetic Association*, 45(3), 289-309.
- Solé, Maria-Josep. "Phonetic and phonological processes: The case of nasalization." *Language and speech* 35.1-2 (1992): 29-43.
- Strand, E. A. (1999). Uncovering the role of gender stereotypes in speech perception. *Journal of Language and Social Psychology*, 18(1), 86-100.
- 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.

- Tamminga, M., & Zellou, G. (2015). Cross-dialectal differences in nasal coarticulation in American English. In *ICPhS*.
- Walker, A., & Hay, J. (2011). Congruence between ‘word age’ and ‘voice age’ facilitates lexical access. *Laboratory Phonology*, 2(1), 219-237.
- Warren, P. (2005). Patterns of late rising in New Zealand English: Intonational variation or intonational change?. *Language Variation and Change*, 17(2), 209-230.
- Zellou, G. (2017). Individual differences in the production of nasal coarticulation and perceptual compensation. *Journal of Phonetics*, 61, 13-29.
- Zellou, G., & Tamminga, M. (2014). Nasal coarticulation changes over time in Philadelphia English. *Journal of Phonetics*, 47, 18-35.
- Zellou, G., & Dahan, D. (2019). Listeners maintain phonological uncertainty over time and across words: The case of vowel nasality in English. *Journal of Phonetics*, 76, 100910.
- Zellou, G., Barreda, S., & Ferenc Segedin, B. (2020). Partial perceptual compensation for nasal coarticulation is robust to fundamental frequency variation. *The Journal of the Acoustical Society of America*, 147(3), EL271-EL276.