

Central Lancashire Online Knowledge (CLoK)

Title	Acoustic analysis of like in North-West England shows context effects, but not function effects
Type	Article
URL	https://clok.uclan.ac.uk/id/eprint/54406/
DOI	https://doi.org/10.1017/S1360674325000115
Date	2025
Citation	Bürkle, Daniel Matthias (2025) Acoustic analysis of like in North-West England shows context effects, but not function effects. English Language and Linguistics.
Creators	Bürkle, Daniel Matthias

It is advisable to refer to the publisher's version if you intend to cite from the work.
<https://doi.org/10.1017/S1360674325000115>

For information about Research at UCLan please go to <http://www.uclan.ac.uk/research/>

All outputs in CLoK are protected by Intellectual Property Rights law, including Copyright law. Copyright, IPR and Moral Rights for the works on this site are retained by the individual authors and/or other copyright owners. Terms and conditions for use of this material are defined in the <http://clok.uclan.ac.uk/policies/>

RESEARCH ARTICLE

Acoustic analysis of *like* in North-West England shows context effects, but not function effects

Daniel Matthias Bürkle 

School of Psychology and Humanities, University of Central Lancashire, Preston PR1 2HE, UK
Email: DMBourke@uclan.ac.uk

(Received 17 July 2023; revised 30 January 2025; accepted 4 February 2025)

Abstract

If supposedly homophonous words were acoustically distinct despite sharing phonemic form, theories of mental storage may have to account for the consistent differences with separate storage for each homophone. Previous studies of the homophonous functions or word classes of the English word *like* showed such subphonemic differences between functions, though some studies also found effects of utterance context alongside these. Schlee & Turton (2018) argued that all these function effects reduce to context effects, since function is not independent of context – for example, quotative *like* typically occurs before a pause and thus is typically subject to lengthening because of its position, not due to a lexicalised acoustic distinction between functions. Testing this argument with new data from a different regional variety to those used by Schlee & Turton, we only find differences that can be explained by context, in line with their argument. This casts prior findings of acoustic distinctions between *like* functions in new light, and introduces the need for further research (especially including the frequency of different functions).

Keywords: *like*; Northern England English; sociophonetics; homophony

1. Background

Many linguistic theories assume, and much psycho- and sociolinguistic work predicts, that different functions of a word are distinct in the speaker's mind. The weakest version of this position is uncontroversial – if polysemous and homophonous words exist, then competent speakers must be distinguishing them in some way – but the nature of this distinction is debated: are the different functions stored as one lemma, or separately? Much of the research in this debate has focused on the useful theoretical distinction between polysemy (one lemma with different but related senses, e.g. *film* 'thin sheet of material' and *film* 'motion picture, originally recorded on thin strips of material') and homonymy (two lemmas with different and unrelated meanings, e.g. *bank* 'monetary institution' and *bank* 'side of a river'). Results are ambiguous as to whether this theoretical distinction is relevant for psycholinguistic descriptions of the mental lexicon (see Gries 2019: 32–5 for a summary).

One interesting strand in this research investigates language production. If different senses/meanings of a polysemous/homonymous word are stored distinctly, this would allow for (though not guarantee) systematic differences in speech production. Such systematic variation could also aid the listener, of course, as the systematic differences could be

mapped to the different senses/meanings and thus allow faster retrieval without the need for processing contextual information. It could even be argued that the existence of such systematic variation explains why polysemy and homonymy persist despite evident cognitive and communicative pressures against them: if a polysemous/homonymous word is not ambiguous to the listener due to systematic acoustic differences between senses/meanings, then there is less pressure to remove this supposed source of ambiguity (cf. Ferreira 2008; though see Trott & Bergen's (2022) argument that languages do in fact reduce this source of ambiguity).

These systematic differences could of course exist in any aspect of the linguistic form; for the purposes of this article, we focus on acoustic (arguably subphonemic) differences, as (following Jurafsky *et al.* 2002) several studies revealed interesting acoustic differences between supposedly homophonous words. For instance, Gahl (2008) and Lohmann (2018) showed that word duration, as a function of word/lemma frequency, can reliably distinguish supposedly homophonous English words: in spontaneous speech, the high-frequency word *time* is often shorter than its supposed homophone *thyme*, which is less frequent. Similarly, supposedly homophonous pairs that are distinguished by inflection are also distinguished acoustically, possibly due to stem frequency (Engemann & Plag 2021), relative frequencies of different word forms (Luef & Sun 2021), or inheritance of acoustic targets from the stem (Seyfarth *et al.* 2018).

The English word *like* is so frequent across contexts, registers and varieties in present-day usage that its frequency is often commented on by non-linguists, mostly negatively (see D'Arcy 2017). Unlike many other frequent or strongly socially indexed words, it is highly polysemous (we propose seventeen different senses/functions in this article; see section 2). For these reasons, it is the ideal subject for investigating the question of function-specific storage and production.

Prior evidence for separate storage of different *like* functions includes the work of Corrigan & Diskin (2019), Diskin-Holdaway (2021) and Pan & Aroonmanakun (2022), who found that second-language learners from various L1s and in various countries use the discourse marker *like*, but not some other functions of *like*, with the same frequency as comparable first-language speakers. As Zaykovskaya (2022) suggests, this may be conceptualised as over-use of discourse *like* (relative to the use of other functions of *like*) as a conspicuous sociolinguistic marker (similar to Davydova's (2021) analysis of L2 quotative *like*). Regardless of this intriguing possibility for future research, these findings all imply that different functions of *like* are stored and accessed separately.

Further evidence comes from acoustic differences between *like* functions. Drager (2009) found systematic differences in *like* tokens (realisation/quality and relative durations of the three segments making up a prototypical *like* token) produced by young female speakers of New Zealand English, arguing that these differences index the function of *like* as well as the speaker's social identity. Likewise, Podlubny *et al.* (2015) showed that duration and vowel quality differ between (some) functions of *like* in the speech of Canadian English speakers to such a degree that they may serve as cues to *like* function.

On the other hand, Schleef & Turton's (2018) study of *like* in London and Edinburgh, while documenting similar acoustic differences between functions, argued that these differences are due to the syntactic and prosodic context. That context is of course not the same for every token with the same function, but Schleef & Turton show that certain contexts are associated with certain functions quite reliably – for example, quotative *like* is rarely in a clitic group (Beckman & Hirschberg's (1993) ToBI boundary strength 0), and it is this clitic group environment that is most clearly associated with diphthongal vowels in *like*. Thus, they argue, the monophthongal vowels found in many quotative *likes* are not an acoustic cue to that function of *like*, but rather evidence of the fact that quotative *like* tends to be used in a context that happens to favour monophthongal realisations of vowels.

We must furthermore respect the fact that most of the functions of *like* are also distinguished by word class, because Hollmann (2020) showed that different word classes (even subclasses) systematically differ in some phonological dimensions. For instance, prototypically attributive adjectives (like *recent*) on average contain more syllables and are more likely to contain nasal consonants than prototypically predicative adjectives (like *better*). By the same token, it is possible that acoustic differences between functions of *like* are cues to word class, not a specific *like* function – the verb *like* may be produced with the acoustic information for the single lemma *like* plus the systematic/typical acoustic information for verbs, rather than there being a consistent but non-systematic set of acoustic information for the verb *like* or a reliable realisation effect from its typical context.

The present study further investigates acoustic differences between different functions of *like*, using rich data from another regional accent of England. Our primary aim, though, is to scrutinise Schlee & Turton's (2018) argument that apparently systematic acoustic differences in *like* may in fact stem from the most common positions and contexts of each function (and not to describe regional patterns in *like*). Parsimoniously, we expect any strong acoustic differences we find to be explained by context.

2. Functions of *like*

Different typologies of the different functions of *like* have been proposed (e.g. nine functions in D'Arcy 2007 and Drager 2009, eleven in Podlubny *et al.* 2015, and at least eleven in Dinkin 2016). We argue that there are seventeen functions in present-day British use, differentiated by word class, function, alternants and position. These are listed below, with examples from our data and explanations of their occurrence and function. It is impossible to reconcile this list of functions and terms with all prior research, because of inconsistent terms and at times less than perfectly precise definitions, but we use the same labels (or similar ones) as prior work wherever possible. We do not intend this typology to be the final word on *like*; its purpose is merely clear description and exploration of the present data.

1. the verb, as in *I don't really like the fantasy one*
2. the noun, as in *discrimination and the like*
3. the adverb, as in *your boss seems like a bit of an idiot*. This usually introduces a noun phrase complement and is usually obligatory, which may be why a few researchers (e.g. Dinkin 2016: 227) call it a preposition. We adopt the more widely used label 'adverb' here simply for ease of comparison.
4. the comparative adverb/connector, as in *this isn't anything like the last good one*. This type of *like* occurs between two noun phrases that are being compared (without a conjunction or copula).
5. the adjectival suffix, as in *sitting Buddha-like*. This type of *like* straightforwardly derives adjectives from nouns and other adjectives.
6. the conjunction introducing a non-obligatory clause of similarity, as in *[she's] flapping, like she always does*. As Blondeau & Nagy (2008) note, this *like* can usually be replaced by *as*, and is used for describing similarity to an existing entity/event.
7. the conjunction introducing a non-obligatory clause of comparison, as in *and then he'll look at you like you're crazy*. As López-Couso & Méndez-Naya (2012) note, this *like* can usually be replaced by *as if* or *as though*, but not by *as* alone, and is used for comparison to a hypothetical alternative.
8. the complementiser that, for expediency, we will call type 1 here, which introduces an obligatory clause (often following *feel* or *seem*), as in *you feel like you're making progress*. This *like* can be replaced by *as if*, *as though* and *that*, and can also be omitted.

9. the complementiser of type 2, which introduces an obligatory clause (often following *look* or *sound*), as in *that just sounds like they needed to add something on to the label*. This *like* can be replaced by *as if* or *as though*, but not by *that*, and can also not be omitted.
10. the quotative, as in *I'm like 'I couldn't tell you'*. This introduces complements, these being quotes or approximations of what was said/thought. Quotative *like* is often preceded by forms of *be* or *go*.
11. the approximator, as in *they're like 30, 35 quid*. This introduces a measure phrase with approximative meaning.
12. the exemplifier, as in *commute to get to like jobs and stuff*
13. the clause-/utterance-external discourse marker, as in *Like, he can do that*. This often alternates with phrases such as *you know*, *I guess* or *I mean*.
14. the clause-final discourse marker, as in *trashy taste in music like*. This does not have common alternants; it expresses hedging or softening, and has been described as a stereotypical feature of some regional varieties (Welsh English by McArthur 2003: 110; Irish English by Diskin 2017; and Newcastle English by Simpson 2022).
15. the clause- and utterance-internal discourse particle, as in *while I was like waiting for the train*. This is by far the most common in our data, and appears to be the function of *like* that has attracted lay attention to the supposed over-use of *like*. This *like* can be distinguished from many other functions by the fact that it clearly does not express similarity, comparison, equivalence or approximation. As D'Arcy (2007) notes, it does not have a common alternant.
16. the 'trailing off' discourse marker, as in *well like....* It may be argued that this occurs only in cases where speakers are interrupted (or interrupt themselves) and does thus not have a distinct and coherent function. Introspection and anecdotal evidence (of usage and intonation) suggest that it is used deliberately, with a distinct function of signalling that the utterance is incomplete but the speaker is ready to end their turn; therefore, we count it as a separate function here to allow any possible function-wise differences to surface in our analysis.
17. the 'social media' noun, as in *every like on this photo counts*. This may simply be one sense of the noun function above; due to its relative novelty, we counted it separately to allow any possible differences to surface in our analysis. Note that verbal use with the specific 'social media' meaning has also been attested (e.g. *Please like my page on Facebook*).

3. Data

Through paper notices on community noticeboards at the University of Central Lancashire and in the surrounding city of Preston, we recruited 11 volunteers (all speakers of North-West English by self-identification and experimenter judgment; aged 18–25; 5 self-identified females, 6 self-identified males) and recorded them in a phone conversation with a close friend or family member. By using a head-mounted microphone (Countryman E6 connected to a Zoom H4n battery-powered recorder) and placing this microphone on the side of the head where the speaker did not hold the telephone, it was possible to record speech from the participants only, but not the speech of their friends/family members. These conversations were held in a sound-dampened recording studio; participants were encouraged to place the call to their friend or family member using the landline telephone in this studio, but some preferred to use their mobile phones (thus it would not have been possible to use an acoustically and electromagnetically isolated recording room). Participants were asked to speak to their friend/family member for about 20 to 30 minutes on any topics they chose; the conversations as recorded ranged from 21 to 38 minutes. Following this conversation, participants were asked to read a list of 36 sentences including *likes* of 12 different functions

(see section A of the online supplementary material for the sentence list annotated with *like* functions; this sentence list did not contain distractors). Importantly, speakers were not aware of the content of the sentence list until the start of the sentence list recording; thus, while *like* was most likely salient to speakers once they saw it occur at least once in every sentence in the sentence list, *like* was not highlighted to speakers during the conversation recording.¹ After this sentence-list reading was complete, the recorder was stopped, and the session ended. Participants received GBP 10 as recompense for their time and travel costs. This procedure was approved by the ethics committee at the University of Central Lancashire (reference number BAHSS421).

The conversations were transcribed orthographically by hand. We used the Montreal Forced Aligner (MFA; McAuliffe *et al.* 2017) toolset to train a custom acoustic model for mapping transcribed words to pronunciation, as the existing models for the English language (apparently General American and Southern British) were not appropriate to our (North-West England) data. With this acoustic model, the MFA then aligned these transcripts to the recordings at word- and phone-level, producing Praat (Boersma & Weenink 2015) TextGrids.

The training process for the acoustic model was automatic (i.e. unsupervised); to check that the acoustic model was sound, we compared the MFA alignment against manual alignment for a continuous section of 30 seconds each (including pauses) from three speakers' conversation recordings. For each vowel in this alignment check sample, we compared the values obtained from automatic and manual alignment for six data points: the starting point (time-stamp), duration, and the first and second formants at 25 per cent and 75 per cent of that duration.

The starting points of what ought to be the same phone/segment in these two datasets are correlated very strongly ($r = 0.985$, $p < 0.001$). The mean difference between starting points is 12.4ms. This difference is less than 25 ms for 87.3 per cent of segments, and less than 100 ms for 98.2 per cent of segments. This compares favourably to McAuliffe *et al.*'s (2017: 500–1) differences (mean of 17.3 ms using Buckeye corpus data, 16.6 ms using Phonsay) and percentages under threshold (77 per cent less than 25 ms in Buckeye and 76 per cent less than 25ms in Phonsay; 98/95 per cent less than 100 ms; note that McAuliffe *et al.* used MFA 'out of the box', with the pre-existing acoustic model trained on the LibriSpeech corpus, while we used MFA with a custom acoustic model generated by unsupervised training on our dataset).

In line with previous research, we normalised the vowel formant measurements using Lobanov's method as implemented in the R package phonTools (version 0.2-2.1; Barreda 2015). Then, we calculated the Pillai score for each unique combination of speaker and vowel phoneme, as a measure of distance/difference between these two datasets for that speaker and vowel in five dimensions (F1 at 25%, F1 at 75%, F2 at 25%, F2 at 75%, and vowel duration). For details of the Pillai score, see Hay *et al.* (2006), Hall-Lew (2010) and Kelley & Tucker (2020); for present purposes, it is enough to know that the Pillai score ranges from 0 to 1, with lower scores indicating more overlap. In the present sample of vowel segment data, the Pillai scores for each speaker–vowel combination range from 0.006 to 0.897, with a mean of 0.236. Interpreting this dimensionless statistic is not straightforward, but by comparison with previous research, we interpret this as indicating strong overlap: For instance, Hay *et al.* (2006) report a Pillai score of 0.24 for the famously merged *beer/bare* pair in their two-dimensional New Zealand English data, and Kelly & Tucker's (2020: 143) target Pillai score for

¹ One speaker did discuss their use of *like* during their conversation, but this appears to have been prompted by their conversation partner (who was not involved with the design of this project). This one speaker's production of *like* may have been affected by this foregrounding; it is reasonable to assume that at least the conversation data from all other speakers is naturalistic.

their three-dimensional simulation of partially overlapping vowel categories is 0.66 (note that Kelly & Tucker report inverted Pillai scores, so their reported score of 0.34 is equivalent to $1 - 0.34 = 0.66$ in the values/terms we use in this article).²

Thus, we conclude (similarly to MacKenzie & Turton 2020 and Gnevsheva *et al.* 2020) that the alignment produced by the self-trained MFA is not identical to manual alignment, but similar enough for the purposes of our analysis. Therefore, we accepted the TextGrids produced by this model for the conversation data, and also used the same model to produce aligned TextGrids for the sentence list data.

For all 614 tokens of *like* in the conversation dataset and all 416 tokens of *like* in the sentence list dataset, we recorded the information listed in table 1. All duration and timing information uses MFA segment or word boundaries; ‘normalised’ indicates that this information was normalised using Gelman’s (2007) method of subtracting the mean and dividing the result by two standard deviations, to make easier the comparison of these numerical variables to each other as well as to the binary variable of gender. In line with prior research (e.g. Podlubny *et al.* 2015), we use the perceptual Bark scale (rather than raw frequency measurements in Hz) for formants to represent the perceptual prominence of any differences more accurately.

A higher break index on Beckman & Hirschberg’s (1993) scale means a perceptually stronger break – break index 0 indicates a clearly cliticised transition between words, break index 1 the more typical juncture between words, and break index 4 the end of an intonation phrase. While this is perceptual, it measures an important aspect of the phonological environment: clause-initial and clause-final *likes* are more likely to be followed by stronger breaks than quotatives, for instance, and the difference in break strength makes for different phonological environments with well-known effects (Schleef & Turton 2018: 47 and 52–8).

Table 2 shows the number of tokens from each function in this conversation data, using the function labels in section 2. For 35 tokens, it was not clear which function they had; these unclear tokens were excluded from all analysis in this article

It is evident from table 2 that some functions are much more common than others (in our data, but most likely this would also hold true for a larger-scale study of function frequency). For the present analyses, we combined the two conjunction functions with the two complementiser functions into one ‘conjunction/complementiser’ category, as these functions are syntactically rather similar.

Of the 36 sentences in the sentence list, 34 contained one *like* each; the other two sentences each contained two *likes*. Thus, the whole list of 36 sentences contained 38 *likes* (see supplementary material). Two sentence tokens with one *like* each (namely sentences number 6 and 7 in speaker number 5’s recording) were not recorded due to a technical fault and are thus not included in any analyses, leaving $(38 \times 11 - 2 =)$ 416 tokens in the sentence list dataset, for a total of $(579 + 416 =)$ 995 *like* tokens used in the analyses reported below.

To reduce the dimensionality of this data and eliminate multicollinearity between acoustic variables, which can cause problems in regression modelling, we conducted factor analysis for mixed data (FAMD) before fitting explanatory regression models (see section 4). FAMD calculates a lower-dimensionality representation of mixed (quantitative and categorical) data, just as principal components analysis does for quantitative-only data and multiple correspondence analysis does for categorical-only data (Pagès 2014). In fields ranging from grapevine diseases (Bertrand *et al.* 2007) to global conflicts (Vazquez *et al.* 2024), FAMD has

² Lobanov’s normalisation method is of course one option of several; to test for possible distortion, we re-calculated the Pillai scores after applying Barreda & Nearey’s (2018) method for normalising vowel formants and Gelman’s method for normalising segment durations. These scores range from 0.0006 to 0.64, with a mean 0.185. As this is fairly similar to the Lobanov-based Pillai scores, we accept the latter as sound for present purposes.

Table 1. Information calculated/extracted for each *like* token

random effect social variables	speaker ID	1–11 (anonymous)
	speaker gender	male or female (in open-ended self-identification, no speaker identified as gender other than these two)
	genre of recording	conversation or sentence list
variable of interest	function of <i>like</i> token	one of the 17 functions listed in section 2, or 'unclear' in ambiguous cases
basis of derived variables (below)	absolute duration of each of the three segments (/l/, vowel, and /k/)	0.03s – 0.3s, 0.03s – 0.35s, and 0.03s – 0.53s respectively
	first and second formants of the vowel, at time-points 25% and 75% of the vowel segment	2.79 – 12.97 for F1 at 25%; 1.46 – 10.51 for F1 at 75%; 7.18 – 14.47 for F2 at 25%; 8.05–15.87 for F2 at 75% (in Bark, calculated using Traunmüller's (1990) formula $\frac{26.81 \times f}{1960 + f} - 0.53$, where f is the raw measurement in Hz; time-points calculated from MFA segment boundaries)
dependent variables	normalised word duration of <i>like</i>	–0.84–2.41
	relative duration of the /k/ segment	0.08–0.75 (as proportion of the word duration)
	normalised ratio of /l/ duration to vowel duration	–0.24–1.09 (also used in Drager 2009, 2011 and Schlee & Turton 2018)
	normalised Euclidean distance between the two pairs of formants as a measure of how monophthongal or diphthongal the vowel token is	–1.07–2.56 (absolute difference between F1 at 25% and F1 at 75%, plus the absolute difference between F2 at 25% and F2 at 75%, as in Drager 2009, 2011 and Podlubny et al. 2015)
context variables	strength of the boundary immediately following <i>like</i>	0–4 (annotated manually using Beckman & Hirschberg's (1993) break indices, as in Schlee & Turton 2018)
	position of the <i>like</i> token in the sentence/utterance	initial, medial, or final (annotated manually, with <i>like</i> immediately before a restart/recast, long pause, or long disfluency tagged as final)
	normalised speech rate	–1.63–1.96 (vowels per second in a window extending up to 3 words before and after the <i>like</i> token, but not across pauses or unaligned words; calculated automatically from MFA boundaries)
	type of the segments immediately preceding and following <i>like</i>	plosive, nasal, fricative, affricate, liquid, glide, vowel, pause, or 'unknown' for unaligned words (determined automatically from TextGrid annotations)
	log-transformed bigram frequency of the preceding word + <i>like</i> as well as log-transformed bigram frequency of <i>like</i> + the following word	0–10.69 and 0–11.25 respectively (only if any word immediately preceded/followed the <i>like</i> token; calculated from van Heuven et al.'s (2014) SUBTLEX-UK frequencies, as in Schlee & Turton 2018)

Table 2. Tokens by function (conversation data only)

Function	Number of tokens
verb	22
noun	0
adverb	18
comparative adverb/connector	23
adjectival suffix	0
conjunction with optional clause of similarity	7
conjunction with optional clause of comparison	1
complementiser type 1 (<i>as if/as though/that/∅</i>)	3
complementiser type 2 (<i>as if/as though</i>)	4
quotative	68
approximator	53
exemplifier	33
clause-/utterance-external discourse marker	36
clause-final discourse marker	7
clause- and utterance-internal discourse particle	277
‘trailing off’ discourse marker	27
‘social media’ noun	0
Total	579

been used as an analysis method revealing correlations or as a dimension-reduction method prior to further analysis (as we use it here). The only prior use in linguistics that we are aware of is Onosson & Stewart’s (2021) description of Media Lengua vowels, though a paper that used FAMD in clustering the vocalisations of three Atlantic seals (Stansbury *et al.* 2015) may also be of interest.

For FAMD dimensionality-reduction on the present data, we used R version 4.0.5 (R Core Team 2021) and the FactoMineR package (version 2.10; Lê *et al.* 2008). This was done separately for each of the four intended dependent variables listed in table 1, using all of the context variables (to use the term from table 1) as well as the genre of recording and the three other variables that were not destined to be the dependent variable for the respective model (out of the word duration, /k/ duration, ratio of /l/ to vowel durations and Euclidean diphthongisation measure). This FAMD resulted in five dimensions of data, with each *like* token’s acoustic information being expressed by five coordinates within this space.

As we conducted FAMD separately for each of the four regression models, the FAMD dimensions are not comparable across models and do not represent the same variables across models, as they are composed of different variables for each model. This is unavoidable, but the alternative with comparable models (that is, a consistent single FAMD of context variables only, with the other variables added to each model individually) would reintroduce problems of multicollinearity.

4. Mixed-effects regression modelling

We used the lme4 package for R (version 1.1-27.1; Bates *et al.* 2015) to apply four separate linear mixed-effects regression models to this set of *like* tokens – one each for the four acoustic measures as dependent variable in prior research (diphthongisation of the vowel, ratio of /l/ segment duration to vowel segment duration, the duration of the /k/ segment expressed as a proportion of the word token duration, and the word token duration itself). We used restricted maximum-likelihood estimation (REML) and bound optimisation by quadratic approximation (BOBYQA) limited to a maximum of 100,000 function evaluations. The fixed independent variables in each model were the speaker's gender, the function of *like*, and the five dimensions calculated by FAMD; the models also included a random speaker-wise slope for each of these variables except the speaker's gender (as this does not vary within each speaker). For the purposes of analysis, we set the discourse particle *like* as the reference level for the variable of *like* function. The resulting models are discussed separately below. Note that all figures reported below are rounded to three decimal places. Due to this rounding, the *t* values reported here are not always identical to the fraction of reported coefficient estimate over reported standard error.

4.1. Model A: Diphthongisation measure

The model for the normalised diphthongisation measure (see table 3 for fixed-effect coefficients) shows that this measure of vowel quality is affected by four of the five FAMD dimensions calculated from acoustic variables and by two functions of *like*: the adjectival suffix tends to have a more monophthongal vowel than the discourse particle (and most other functions), and the noun *like* a more diphthongal vowel. Effects of other *like* functions are not strong enough to be distinguished from zero.

As table 2 shows, these two functions of *like* only occurred in our sentence-list data, not in conversation. Thus, it is not surprising that they show more consistent and extreme articulations (very diphthongal nouns, very monophthongal suffixes) than other categories: all the data for these two functions comes from the same syntactic and phonological environment.

The first FAMD dimension in this model largely incorporates the speech rate and information about what follows the *like* token. This connection is not surprising, given well-known effects like final lengthening, and neither is the fact that this dimension has a significant effect on the diphthongisation measure: shorter vowels are less likely to be diphthongal, after all. The second FAMD dimension represents the position of *like*. The third and fourth FAMD dimensions mostly represent the preceding and following segments and the genre of recording. Section B of the online supplementary material shows correlation circles and tables of variable contributions to these FAMD dimensions.

4.2. Model B: /l/-to-vowel duration ratio

The dependent variable for this model was calculated by dividing the absolute duration of a token's /l/ segment by the absolute duration of that token's vowel, thus expressing how long the /l/ is compared to the vowel. We have two reasons to use this ratio rather than the absolute duration of /l/ as a measure of /l/ reduction. Firstly, this reduces the impact of speech rate differences: the absolute /l/ duration is moderately correlated with speech rate (Spearman's $\rho = -0.35$, $p < 0.001$ in our data), but the /l/-to-vowel duration ratio is not (Spearman's $\rho = 0.07$, $p > 0.05$). Secondly, it makes our findings more easily comparable to

Table 3. Fixed-effect coefficients of model for vowel diphthongisation

Fixed effect	Coefficient estimate	Standard error	t
(intercept: discourse particle, female speaker)	0.026	0.068	0.383
adjectival suffix	−0.271	0.107	−2.540
adverb	0.031	0.083	0.372
approximator	−0.028	0.052	−0.535
conjunction/ complementiser	−0.044	0.125	−0.348
comparative adverb	0.102	0.137	0.744
clause-final discourse marker	0.407	0.555	0.733
clause-/utterance-external discourse marker	−0.060	0.136	−0.438
‘trailing off’ discourse marker	−0.094	0.078	−1.214
exemplifier	−0.030	0.208	−0.144
noun	0.206	0.076	2.715
quotative	0.020	0.097	0.202
‘social media’ <i>like</i>	−0.001	0.082	−0.011
verb	−0.014	0.097	−0.140
male speaker	0.014	0.062	−0.220
FAMD dimension 1	−0.102	0.018	−5.526
FAMD dimension 2	0.029	0.011	2.649
FAMD dimension 3	0.063	0.013	5.032
FAMD dimension 4	−0.045	0.021	−2.164
FAMD dimension 5	0.006	0.012	0.519

prior work which also used this ratio (for instance Drager 2009, 2011 and Schlee & Turton 2018).³ This duration ratio variable was normalised before modelling, as stated above.⁴

The mixed-effects model for this dependent variable (see table 4) shows that the /l/ segment is relatively short in the noun *like* and relatively long when produced by male speakers. Some of the acoustic and contextual information incorporated into the FAMD dimensions is also predictive (dimension 2, which largely comprises positional information in this model).

³ This second reason is also the reason why we do not use the ratio of /l/ duration to word duration. That said, there does not appear to be much difference between these two ratios: the vowel segment duration and whole word duration correlate strongly, of course (Spearman's $\rho = 0.75$, $p < 0.001$), and thus the /l/-to-word duration ratio is very similar to the /l/-to-vowel duration ratio (Spearman's $\rho = -0.78$, $p < 0.001$). Moreover, the coefficient estimates of a model using /l/-to-word duration ratio are effectively identical to the ones in model B.

⁴ Unsurprisingly, the correlations with speech rate are effectively unchanged by normalisation: In the normalised data, there is still a moderate negative correlation between absolute /l/ duration and speech rate (Spearman's $\rho = -0.34$, $p < 0.001$), but not between /l/-to-vowel duration ratio and speech rate (Spearman's $\rho = 0.08$, $p < 0.01$).

Table 4. Fixed-effect coefficients of model for /l/-to-vowel duration ratio

Fixed effect	Coefficient estimate	Standard error	t
(intercept: discourse particle, female speaker)	-0.050	0.035	-1.424
adjectival suffix	0.014	0.088	0.162
adverb	0.017	0.071	0.233
approximator	-0.023	0.061	-0.384
conjunction/ complementiser	-0.057	0.126	-0.449
comparative adverb	-0.011	0.124	-0.085
clause-final discourse marker	0.110	0.644	0.170
clause-/utterance-external discourse marker	-0.045	0.122	-0.370
'trailing off' discourse marker	-0.010	0.074	-0.141
exemplifier	0.043	0.174	0.247
noun	-0.190	0.091	-2.087
quotative	-0.018	0.104	-0.173
'social media' like	0.204	0.205	0.998
verb	0.001	0.076	0.019
male speaker	0.130	0.040	3.272
FAMD dimension 1	0.004	0.008	0.495
FAMD dimension 2	-0.049	0.023	-2.080
FAMD dimension 3	0.010	0.013	0.757
FAMD dimension 4	-0.023	0.028	-0.829
FAMD dimension 5	-0.001	0.013	-0.055

Following from the vowel quality effect described by model A above, this length ratio effect of the noun function is not surprising: the noun *like* in our data tended to have a diphthongal, long vowel, and as this vowel length is the denominator of the ratio used in model B, it naturally makes for a negative effect in model B just as it made for a positive effect in model A. The gender effect is unexpected, but robust even in descriptive statistics: In our data, male speakers produce longer /l/s and shorter vowels (mean /l/ duration 86.7 ms, mean vowel duration 75.7 ms) than female speakers do (mean /l/ duration 67.0 ms, mean vowel duration 102 ms).

4.3. Model C: /k/-to-word duration ratio

The dependent variable for this model was calculated by dividing the absolute duration of a token's /k/ segment by the absolute duration of the whole token (in other words, it expresses the length of the /k/ as a proportion of the length of the whole token). Previous work has studied /k/ by using variables that record expert judgments of /k/ realisation/

Table 5. Fixed-effect coefficients of model for /k/-to-word duration ratio

Fixed effect	Coefficient estimate	Standard error	t
(intercept: discourse particle, female speaker)	0.348	0.018	19.098
adjectival suffix	−0.011	0.026	−0.408
adverb	0.023	0.021	1.108
approximator	0.005	0.020	0.257
conjunction/ complementiser	0.054	0.037	1.461
comparative adverb	0.033	0.034	0.975
clause-final discourse marker	−0.066	0.142	−0.463
clause-/utterance-external discourse marker	−0.007	0.034	−0.195
‘trailing off’ discourse marker	−0.024	0.026	−0.937
exemplifier	−0.091	0.050	−1.834
noun	−0.026	0.017	−1.504
quotative	0.030	0.025	1.234
‘social media’ <i>like</i>	0.021	0.017	1.243
verb	< −0.001	0.024	−0.004
male speaker	−0.006	0.011	−0.565
FAMD dimension 1	−0.011	0.005	−2.081
FAMD dimension 2	0.005	0.003	1.482
FAMD dimension 3	0.006	0.005	1.367
FAMD dimension 4	−0.016	0.003	−5.538
FAMD dimension 5	0.007	0.004	2.042

quality; as this work by and large found /k/ reduction, we chose to use this duration ratio as a more straightforward quantity that measures /k/ reduction here.

The mixed-effects model for this dependent variable (shown in table 5) shows that it is affected by acoustic and contextual information (in FAMD dimensions) – there is no meaningful effect of *like* function.

The first FAMD dimension in this model can be interpreted as speech rate and related factors. The fourth dimension here is largely made up of information regarding the preceding and following segments, while the fifth dimension adds the /l/-to-vowel duration ratio. As with similar FAMD dimensions in model A above, it is not surprising that these positional and durational variables have effects on the /k/-to-vowel duration ratio in this model.

4.4. Model D: Word duration

Table 6 shows the fixed-effect coefficients of the model for the (normalised) duration of each *like* token. This direct phonological measure is strongly connected to the FAMD dimensions,

Table 6. Fixed-effect coefficients of model for word duration

Fixed effect	Coefficient estimate	Standard error	t
(intercept: discourse particle, female speaker)	-0.045	0.026	-1.714
adjectival suffix	-0.005	0.068	-0.076
adverb	0.002	0.045	0.039
approximator	-0.022	0.042	-0.529
conjunction/ complementiser	-0.168	0.087	-1.926
comparative adverb	-0.068	0.086	-0.798
clause-final discourse marker	0.094	0.399	0.234
clause-/utterance-external discourse marker	0.019	0.113	0.169
'trailing off' discourse marker	0.202	0.061	3.299
exemplifier	-0.109	0.149	-0.732
noun	0.026	0.078	0.329
quotative	-0.056	0.117	-0.482
'social media' <i>like</i>	0.059	0.044	1.349
verb	0.305	0.067	4.584
male speaker	-0.052	0.030	-1.731
FAMD dimension 1	-0.187	0.015	-12.816
FAMD dimension 2	0.012	0.008	1.501
FAMD dimension 3	0.072	0.014	5.106
FAMD dimension 4	0.021	0.013	1.622
FAMD dimension 5	0.005	0.009	0.556

which is unsurprising: it is widely known that word duration and other segmental and suprasegmental features can affect one another (as noted in other models above), and positional lengthening/shortening effects are also well-established. Moreover, one of the variables incorporated in these FAMD dimensions is speech rate, which naturally affects word duration; as speech rate contributed mostly to the first dimension in this model, it is again unsurprising to see a strong negative effect of that dimension on word duration (higher speech rate going together with lower word duration). The largest contributions to the third FAMD dimension here come from the vowel diphthongisation measure, genre of recording and positional information.

The model does show some effects of *like* functions: verbs and 'trailing off' marker *likes* are longer than other *likes*, and conjunctions/complementisers are shorter than other *likes*. These effects appear scattershot at first glance, but taking them in turn suggests that they are incidental to the usual positions of different *like* functions rather than inherent distinctions/cues to function directly: The verb *like* occurs 22 times in our conversation data and 33 times in the sentence-list data. As the sentence list data has a slower speech

rate overall, it is not surprising that we see longer word durations for this function of *like*, with relatively more of its data coming from that slower part of the dataset. Verbs are also more likely to be stress-bearing (though this was not measured/included in the models formally) and thus longer.

The ‘trailing off’ discourse marker occurs at the end of utterances or turns by definition and, due to its apparent function (see [section 2](#)), is naturally likely to be lengthened. Conversely, conjunctions and complementisers are usually utterance-internal and not followed by pauses (only 6 of 33 sentence-list and 1 of 15 conversation tokens with these functions are followed by pauses); thus, most of these conjunction and complementiser tokens were not subject to final or pre-pausal lengthening, which explains why they tend to be shorter. In other words, all effects of *like* function in model D can be explained by (typical) position and context: final, stressed and/or pre-pausal *like* tends to be longer.

5. Hierarchical cluster analysis

Linear mixed-effects regression modelling is widely used in linguistics, but has been criticised as it can lead to unsubstantiated conclusions (see e.g. Eager & Roy’s (2017) demonstration that mixed-effects regression models often fall short of their aim of accounting for random effects properly, especially with unbalanced and binary data). To alleviate this concern, we chose a second approach to answering our research question (are there acoustic differences between *like* functions that cannot be explained by token position/context?): hierarchical clustering.

After removing *like* tokens in the rare and phonologically particular functions of clause-final discourse marker (‘Irish *like*’) and ‘trailing off’ *like* as well as tokens whose function was unclear, we mean-centred the speech rate, word duration, relative /k/ duration, /l/-to-vowel duration ratio and diphthongisation measure, and scaled these to standard deviations within each variable. The genre of recording, boundary strength immediately following *like*, position of *like* in the sentence/utterance, speech rate, type of segments (if any) immediately preceding and following *like*, and the two bigram frequencies for each *like* token (all as described above) were also available to the clustering algorithm. Importantly, though, speaker ID and *like* token function were held out of this dataset when we used Ward’s (1963) minimum-variance method (as implemented in R) to calculate a best-fit clustering solution. Using the R package *dendextend* (version 1.15.1; Galili 2015), this clustering solution was plotted as a dendrogram.

This dendrogram contains a few small function-wise clusters for rare functions, but no larger speaker- or function-wise clusters. The truncated version shown in [figure 1](#) demonstrates that most functions occur across all clusters – each pie chart shows the proportion of functions within that branch of the dendrogram. Most noticeably, there are no sizeable clusters that do not contain some discourse particle *likes*. While this is not unexpected given the frequency of this type of *like*, and the fact that it can occur in different contexts, it is possible that the clustering algorithm was unable to find speaker- or function-wise clusters due to the existence of discourse particle *likes* across the variable space. Thus, the procedure was repeated without the discourse particle tokens (and a few other outliers, defined by being at least two standard deviations from the mean value of a variable) to investigate whether the unbalanced full dataset obscured interesting function-wise differences. However, the dendrogram for this reduced dataset does not show sizeable clusters for any single *like* function either (and is thus not shown here). Thus, we conclude that such similarities that do exist between different *like* tokens in our data are due to features other than the function of *like*, as the regression modelling suggested.

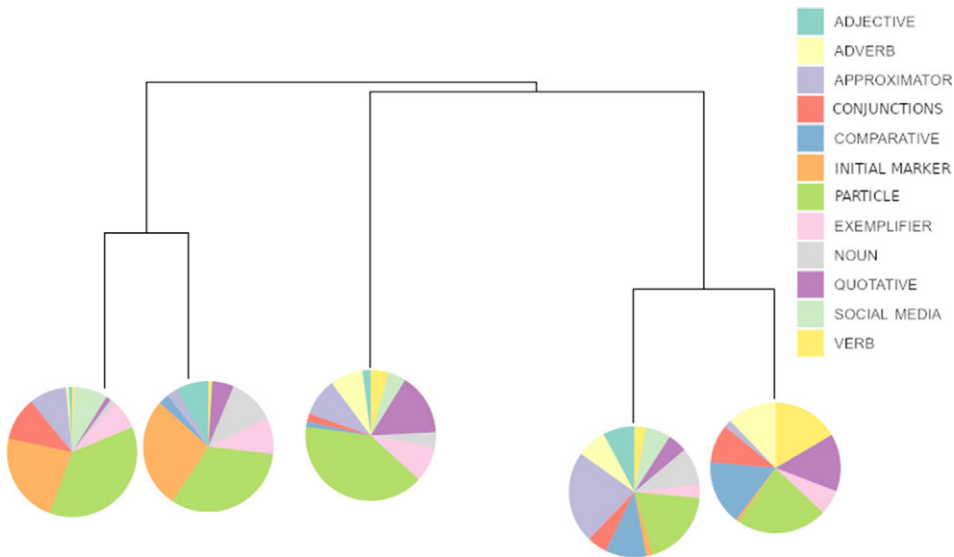


Figure 1. Dendrogram of clustering solution including discourse particle

6. General discussion

Three of the four linear mixed-effects regression models fitted to our data contain effects of *like* function. However, in all three models, these effects can be explained by context:

- nominal *like* tokens in our data have more diphthongal vowels (and, concurrently, a lower ratio between the durations of /l/ and that vowel) due to occurring in more consistent and extreme environments in our elicitation;
- likewise, the adjectival suffix *likes* tend to monophthongs in our data because of sentence-list intonation;
- verbs and 'trailing off' *likes* are longer because of stress and/or a following pause; and
- likewise, *like* tokens that function as conjunctions or complementisers are shorter than others because they rarely occur in a pre-pausal position and thus are more likely to be reduced.

In summary, these apparent effects of *like* function in models of acoustic measures can all be explained by lengthening or reduction due to context (or, in the case of the noun and adjectival suffix, type of recording/data).

It is interesting that these apparent function effects (which we argue are proxies of context effects) emerge in the models despite the presence of other variables that measure the context directly. We suggest two reasons why these apparent function effects do emerge (though the present study cannot support either of these suggestions strongly): the first possible explanation is that the context effects are so strong that they cover all other effects, but emerge even in imperfect proxy variables. This possible explanation is supported by the observation that most effects, and strongest effects, in models A, B, C and D are effects of the utterance context (incorporated in FAMD dimensions). The second possible explanation is that our model fitting and selection procedure was deliberately biased to include the variable that records the function of *like* tokens, by not incorporating it in dimension-reducing FAMD.

Hierarchical clustering also shows no sizeable clusters by *like* function when contextual information is available to the clustering algorithm alongside acoustic information. There are many acoustic differences between individual tokens of *like* in our data, but none that we can confidently ascribe to the function of *like* – in other words, different functions of *like* are not produced with systematic differences in our data. Thus, we see no reason to assume that these different functions are stored separately.

The present findings are further evidence for Schlee & Turton's (2018) argument that subphonemic differences in *like* can be explained by context. This argument appears to counter the earlier findings of Drager (2009, 2011) and Podlubny *et al.* (2015), who saw function-specific differences when some, but not all of the features of a *like* token's context were taken into account. For instance, Drager (2011) considers the type of the segments preceding and following each *like* token with great care, but not the fact that different functions of *like* are likely to occur in different prosodic settings – and it is this latter difference in typical setting that causes differences in *like* realisations, we argue. That said, there are two possible other explanations discussed in those earlier papers: frequency and social group membership/identity. We will address these possibilities in turn.

Word frequency has been shown to affect word duration. Gahl's seminal paper found a length difference between frequent and infrequent homophones of 28 ms (368 vs 396 ms in all data; Gahl 2008: 481) or 22 ms (374 vs 396 ms for just the first use in each text; Gahl 2008: 487). Of course, these figures are for many homophone pairs of different lengths and thus cannot be applied to *like* directly, but they serve as a point of comparison: in our conversation data, the longest reliable⁵ function-wise average word duration is 266 ms for verbs and utterance-/sentence-initial discourse markers, and the shortest is 182 ms for comparative adverbs/connectors. The most frequent function in our data, the clause- and utterance-internal discourse particle *like*, has a high average word duration when considered in this range, at 231 ms. The scale of the difference between averages in our data is larger than in Gahl (2008), and thus it is conceivable that frequency (of the different *like* functions, in this case) explains at least part of the difference. Of course, this frequency is likely connected to typical position – for instance, complementisers are surely rarer than discourse markers and particles not because speakers often choose to omit complementisers, but rather because speakers choose to use fewer constructions that require or allow complementisers in the first place. Frequency may well have an additional effect, and future studies would be well-advised to investigate it (with function-wise frequency data). That said, our length differences are clearly not due to frequency in the way Gahl's were: the most frequent *like* (the discourse particle) is not usually the shortest here; moreover, the shortest and longest *likes* have functions that, intuitively, do not differ in frequency. This discussion of frequency effects does of course not call into question Gahl (2008) or other frequency effect findings (for example, Lohmann's 2018 confirmation of Gahl's work, which found frequency effects while controlling for some features of position/context). However, it illustrates that our findings for *like* are unlikely to be frequency effects, as they do not match what would be predicted for frequency effects.

In addition to lemma frequency generally, there is speaker-specific frequency (or likelihood of use) for each function. As Drager (2011) argues convincingly, one speaker's production may be determined by how often they themselves use the lemma (for instance, forms that a particular speaker uses often are more likely to be reduced by *that speaker*). We

⁵ There are more extreme averages than these in our conversation data, but they have been ignored here either because there are too few tokens for the average to be reliable and meaningful (clause-final discourse marker, conjunctions and complementisers) or because they are greatly lengthened by definition (the 'trailing off' discourse marker).

were unable to include speaker-specific frequency in our analyses, as it would require much more data per speaker.

Likewise, a speaker's identity (or social group membership) may well affect their production of *like*. Drager's (2009, 2011) detailed study of *like* in one high school year group found differences in *like* between pupils according to where they tended to eat lunch (which indexed further social dimensions). This is intriguing, but unlikely to be the source of any systematic effect in our data, as our participants are arguably drawn from the 'general population' rather than one restricted setting. In such a setting, it is of course possible for acoustic differences to arise and serve as sociolinguistic markers; however, we see no reason to assume that *like* realisation (out of all the features of North-West England English) is indeed a sociolinguistic marker within that larger variety. Thus, while we do not rule out the possibility of function-wise differences as sociolinguistic markers, we believe it is much less likely to explain the present findings than the well-documented context effects are. Of course, function-specific subphonemic markers would be an indication of function-specific storage, as Drager argues; absent a study like Drager's that also accounts for the context effects we argue for, we see no convincing evidence to assume different functions of *like* are stored with subphonemic detail.

Thus, we assume that there is no psychologically real representation of different functions of *like*, at least not in this respect. Some of the information that distinguishes different functions of *like* (e.g. the quite systematic usage pattern, meaning and social evaluation of quotative *like*) is most likely psychologically real in some way, of course, as speakers could otherwise not use the different *likes* as consistently as they evidently do. However, we see no convincing evidence for subphonemic acoustic distinctions in the present study. In other words, we agree with Ferreira (2008): readily accessible information in an utterance (word frequency and semantics in Ferreira's data, syntax and prosody in ours) is often enough to disambiguate phonemically identical utterances in practice, and therefore speakers and listeners do not use other information (such as subphonemic distinctions) to further disambiguate them.

Our typology of seventeen functions of *like* is of course subject to further argument. Due to low numbers of tokens, we are unable to state conclusively that the 'social media' function or the 'trailing off' function are distinct from other functions. Similarly, the four *like* conjunction and complementiser functions appear to be distinct syntactically, but had to be combined in our analysis due to low numbers; there may be interesting distinctions between these functions or further argument on their similarity.

7. Conclusion

Having studied the subphonemic detail of *like* in use by young adult speakers of English from the North-West of England, we find no systematic differences between functions of *like* that cannot be accounted for by context (in line with Schleef & Turton 2018). This means we find no evidence for separate storage of different functions of *like*.

Future work could extend this analysis to data from prior studies (if available), to investigate whether the present argument also holds in these cases. It is, after all, conceivable that some regional varieties of English do differentiate *like* functions subphonemically while others do not, and reliable findings of such differentiation would show conclusively that these function-specific representations do indeed exist.

Another intriguing possibility for future research is to examine this phenomenon from the perspective of the listener, as done previously in Drager (2009) and Lohmann & Tucker (2021). If a listener cannot perceive or use supposed subphonemic differences, the discussion

of their psychological reality is of course not moot, but the listener-oriented argument for this reality would be falsified. It may not be fruitful to study discrimination of *like* tokens in isolation, if relative differences (to the context) are the cues listeners use.

Future research could also investigate the frequency of different functions of *like* in more detail, which would allow the use of relative per-function frequency in statistical analyses and thus the investigation of frequency effects as per Gahl (2008). That said, finding such frequency effects would not necessarily contradict the present findings, as these frequency effects have usually been described in the domain of duration (even for different word forms, as in Engemann & Plag 2021). Duration can, of course, affect other measures we used as dependent variables, but crucially this would be a frequency effect, which does not necessarily require function-specific mental representation (for example, in exemplar-theoretic models). Moreover, even these effects may be affected by context features: Lohmann & Conwell (2020) show that nouns tend to be followed by stronger boundaries than verbs (a function of English syntax), and that this (rather than separate storage) leads to nouns being longer than verbs in homophonous noun–verb pairs like *chat*. Future research investigating frequency effects would thus also be useful in further investigating whether these word-class effects are in fact contextual, just like the supposed *like* function effect we discussed here.

Acknowledgements. This article has benefitted greatly from the insight of Ryan Podlubny, Benjamin V. Tucker, Nick Palfreyman, the attendees of the 29th Manchester Phonology Meeting in 2022, a colleague who wishes to remain anonymous, and the reviewers and editors. The anonymous research participants who contributed speech samples are also acknowledged gratefully.

References

- Barreda, Santiago. 2015. *phonTools: Functions for phonetics in R*. R package version 0.2-2.1.
- Barreda, Santiago & Terrance M. Nearey. 2018. A regression approach to vowel normalization for missing and unbalanced data. *Journal of the Acoustical Society of America* 144(1), 500–20.
- Bates, Douglas, Martin Mächler, Ben Bolker & Steve Walker. 2015. Fitting linear mixed-effects models using lme4. *Journal of Statistical Software* 67(1), 1–48.
- Beckman, Mary E. & Julia Hirschberg. 1993. The ToBI annotation conventions. Unpublished manuscript. www1.cs.columbia.edu/~jjv/pubs/tobi_convent.pdf
- Bertrand, Frédéric, Myriam Maumy, Lionel Fussler, Nathalie Kobes, Serge Savary & Jacques Grosman. 2007. Using factor analyses to explore data generated by the National Grapevine Wood Diseases Survey. *Case Studies in Business, Industry and Government Statistics* 1(2), 183–202.
- Blondeau, Hélène & Naomi Nagy. 2008. Subordinate clause marking in Montreal Anglophone French and English. In Miriam Meyerhoff & Naomi Nagy (eds.), *Social lives in language: Sociolinguistics and multilingual speech communities*, 273–314. Amsterdam: John Benjamins.
- Boersma, Paul & David Weenink. 2015. *Praat: Doing phonetics by computer*. Version 5.4.21. www.praat.org
- Corrigan, Karen P. & Chloé Diskin. 2019. ‘Northmen, Southmen, comrades all?’ The adoption of discourse *like* by migrants north and south of the Irish border. *Language in Society* 49, 745–73.
- D’Arcy, Alexandra. 2007. *Like* and language ideology: Disentangling fact from fiction. *American Speech* 82(4), 386–419.
- D’Arcy, Alexandra. 2017. *Discourse-pragmatic variation in context: Eight hundred years of like*. Amsterdam: John Benjamins.
- Davydova, Julia. 2021. The role of sociocognitive salience in the acquisition of structured variation and linguistic diffusion: Evidence from quotative *be like*. *Language in Society* 50(2), 171–96.
- Dinkin, Aaron J. 2016. Variant-centered variation and the *like* conspiracy. *Linguistic Variation* 16(2), 221–46.
- Diskin, Chloé. 2017. The use of the discourse-pragmatic marker ‘*like*’ by native and non-native speakers of English in Ireland. *Journal of Pragmatics* 120, 144–57.
- Diskin-Holdaway, Chloé. 2021. *You know and like* among migrants in Ireland and Australia. *World Englishes* 40, 562–77.
- Drager, Katie. 2009. A sociophonetic ethnography of Selwyn Girls’ High. PhD dissertation, University of Canterbury. <http://hdl.handle.net/10092/4185>
- Drager, Katie. 2011. Sociophonetic variation and the lemma. *Journal of Phonetics* 39, 694–707.
- Eager, Christopher & Joseph Roy. 2017. Mixed effects models are sometimes terrible. <https://arxiv.org/abs/1701.04858v1>

- Engemann, Marie & Ingo Plag. 2021. Phonetic reduction and paradigm uniformity effects in spontaneous speech. *The Mental Lexicon* 16(1), 165–98.
- Ferreira, Victor S. 2008. Ambiguity, accessibility, and a division of labour for communicative success. *Psychology of Learning and Motivation* 49, 209–46.
- Gahl, Susanne. 2008. *Time and thyme* are not homophones: The effect of lemma frequency on word durations in spontaneous speech. *Language* 84(3), 474–96.
- Galili, Tal. 2015. dendextend: An R package for visualizing, adjusting, and comparing trees of hierarchical clustering. *Bioinformatics* 31(22), 3718–20.
- Gelman, Andrew. 2007. Scaling regression inputs by dividing by two standard deviations. *Statistics in Medicine* 27, 2865–73.
- Gnevshva, Ksenia, Simon Gonzalez & Robert Fromont. 2020. Australian English Bilingual Corpus: Automatic forced alignment accuracy in Russian and English. *Australian Journal of Linguistics* 40(2), 182–93.
- Gries, Stefan Th. 2019. Polysemy. In Ewa Dąbrowska & Dagmar Divjak. (eds.), *Cognitive linguistics: Key topics*, 24–44. Berlin: De Gruyter Mouton.
- Hall-Lew, Lauren. 2010. Improved representation of variance in measures of vowel merger. *Proceedings of Meetings in Acoustics* 9, 060002.
- Hay, Jen, Paul Warren & Katie Drager. 2006. Factors influencing speech perception in the context of a merger-in-progress. *Journal of Phonetics* 34(4), 458–84.
- Hollmann, Willem B. 2020. The ‘nouniness’ of attributive adjectives and ‘verbiness’ of predicative adjectives: Evidence from phonology. *English Language and Linguistics* 25(2), 257–79.
- Jurafsky, Daniel, Alan Bell & Cynthia Girand. 2002. The role of lemma in form variation. In Carlos Gussenhoven & Natasha Warner (eds.), *Papers in laboratory phonology VII*, 1–34. New York: Mouton de Gruyter.
- Kelley, Matthew C. & Benjamin V. Tucker. 2020. A comparison of four vowel overlap measures. *Journal of the Acoustical Society of America* 147, 137.
- Lê, Sebastien, Julie Josse & François Husson. 2008. FactoMineR: An R package for multivariate analysis. *Journal of Statistical Software* 25(1), 1–18.
- Lohmann, Arne. 2018. *Time and thyme* are NOT homophones: A closer look at Gahl’s work on the lemma-frequency effect, including a reanalysis. *Language* 94(2), e180–90.
- Lohmann, Arne & Erin Conwell, E. 2020. Phonetic effects of grammatical category: How category-specific prosodic phrasing and lexical frequency impact the duration of nouns and verbs. *Journal of Phonetics* 78, 100939.
- Lohmann, Arne & Benjamin V. Tucker. 2021. Testing the storage of prosody-induced phonetic detail via auditory lexical decision: A case study of noun/verb homophones. *The Mental Lexicon* 16(1), 133–64.
- López-Couso, María José & Belén Méndez-Naya. 2012. On the use of *as if*, *as though*, and *like* in Present-Day English complementation structures. *Journal of English Linguistics* 40(2), 172–95.
- Luef, Eva Maria & Jong-Seung Sun. 2021. Wordform-specific frequency effects cause acoustic variation in zero-inflected homophones. *Poznan Studies in Contemporary Linguistics* 56(4), 711–39.
- MacKenzie, Laurel & Danielle Turton. 2020. Assessing the accuracy of existing forced alignment software on varieties of British English. *Linguistics Vanguard* 6(s1), 20180061.
- McArthur, Tom. 2003. *The Oxford guide to World Englishes*. Oxford: Oxford University Press.
- McAuliffe, Michael, Michaela Socolof, Sarah Mihuc, Michael Wagner & Morgan Sonderegger. 2017. Montreal Forced Aligner: Trainable text-speech alignment using Kaldi. *Proceedings of the 18th Conference of the International Speech Communication Association*, 498–502.
- Onosson, Sky & Jesse Stewart. 2021. A multi-method approach to correlate identification in acoustic data: The case of Media Lengua. *Laboratory Phonology* 12(1), 13.
- Pagès, Jérôme. 2014. *Multiple factor analysis by example using R*. New York: CRC.
- Pan, Zhaoyi & Wirote Aroonmanakun. 2022. A corpus-based study on the use of spoken discourse markers by Thai EFL learners. *LEARN Journal: Language Education and Acquisition Research Network* 15(2), 187–213.
- Podlubny, Ryan, Kristina Geeraert & Benjamin V. Tucker. 2015. It’s all about, *like*, acoustics. *Proceedings of the 18th International Congress of Phonetic Sciences (ICPhS)*. www.internationalphoneticassociation.org/icphs-proceedings/ICPhS2015/proceedings.html
- R Core Team. 2021. *R: A language and environment for statistical computing*. Version 4.0.5. Vienna: R Foundation for Statistical Computing. www.R-project.org
- Schleef, Erik & Danielle Turton. 2018. Sociophonetic variation of *like* in British dialects: Effects of function, context and predictability. *English Language and Linguistics* 22(1), 35–75.
- Seyfarth, Scott, Marc Garellek, Gwendolyn Gillingham, Farrell Ackerman & Robert Malouf. 2018. Acoustic differences in morphologically-distinct homophones. *Language, Cognition and Neuroscience* 33(1), 32–49.
- Simpson, David. 2022. Geordie dictionary: I–L. <https://englandsnortheast.co.uk/geordie-dictionary-i-l/> (15 February 2024).

- Stansbury, Amanda L., Mafalda de Freitas, Gi-Mick Wu & Vincent M. Janik. 2015. Can a gray seal (*Halichoerus grypus*) generalize call classes? *Journal of Comparative Psychology* 129(4), 412–20.
- Trautmüller, Hartmut. 1990. Analytical expressions for the tonotopic sensory scale. *Journal of the Acoustical Society of America* 88(1), 137.
- Trott, Sean & Benjamin Bergen. 2022. Languages are efficient, but for whom? *Cognition* 225, 105094.
- van Heuven, Walter J. B., Pawel Mandera, Emmanuel Keuleers & Marc Brysbaert. 2014. SUBTLEX-UK: A new and improved word frequency database for British English. *Quarterly Journal of Experimental Psychology* 67(6), 1176–90.
- Vazquez, Kathleen, Jeffrey C. Johnson, David Griffith & Rachata Muneeppeerakul. 2024. Exploration of underlying patterns among conflict, socioeconomic and political factors. *PLoS ONE* 19(5), e0304580
- Ward, Joe H. 1963. Hierarchical grouping to optimize an objective function. *Journal of the American Statistical Association* 58, 236–44.
- Zaykovskaya, Irina. 2022. ‘When I come here I kinda blend in the language so that’s why I use *like*’: Multifunctional word *like* as a challenge for English learners. In Zihan Yin & Elaine Vine (eds.), *Multifunctionality in English: Corpora, language and academic literacy pedagogy*, 222–40. London: Routledge.