

Article

Zeyu Li* and Ulrike Gut

The distribution of /w/ and /ɹ/ in Scottish Standard English

<https://doi.org/10.1515/cllt-2021-0052>

Received July 13, 2021; accepted January 17, 2022; published online February 8, 2022

Abstract: The Scottish English phoneme inventory is generally claimed to have a /ɹ/-/w/ contrast, although several studies have suggested that this historical contrast is weakening for Scottish English speakers in the urban areas of Glasgow, Edinburgh and Aberdeen. Little is known about whether the /ɹ/-/w/ contrast is maintained in supraregional Scottish Standard English (SSE). This study sets out to explore, based on the phonemically transcribed ICE-Scotland corpus, the distribution of [ɹ] and [w] in SSE, their acoustic properties and potentially influencing social and language-internal factors. A total of 1,241 <wh> tokens were extracted from the corpus, together with a matching number of <w> tokens, and the median of harmonicity was measured. The results show that [ɹ] and [w] produced for words beginning with <wh> are acoustically distinct from [w] produced for words beginning with <w>. [ɹ] is relatively frequent in SSE, but most speakers use both [ɹ] and [w] interchangeably for <wh> and some never use [ɹ]. The realisation of <wh> as [ɹ] is determined by preceding phonetic context and speaker gender.

Keywords: /ɹ/-/w/ contrast, harmonicity, ICE-Scotland, Scottish Standard English

1 Introduction

In Scottish English, like in some varieties of Irish, US American and Canadian English, the so-called *wine-whine* merger did not take place, in which historical Old and Middle English /hw/ was replaced by /w/ (Grant 1914: 38; Minkova 2012: 16). Consequently, Scottish English still has the voiceless labial-velar fricative /ɹ/ (Giegerich 1992: 36; Jones 2002: 27; Stuart-Smith 2008: 63; Wells 1982: 408),

*Corresponding author: Zeyu Li, College of Foreign Languages, Nankai University, Tianjin, China, E-mail: zeyu.li@nankai.edu.cn. <https://orcid.org/0000-0003-3980-6581>

Ulrike Gut, Department of English Linguistics, University of Münster, Münster, Germany, E-mail: gut@uni-muenster.de

which is used to pronounce the digraph <wh> in many words such as *which*, *when* and *what*.¹ This results in Scottish English, unlike most other varieties of English, having minimal pairs such as *which* versus *witch*, *where* versus *wear* and *whales* versus *Wales*. Minkova argues that the /ɹ/-/w/ contrast has become a recessive feature in Southern British English since late Old English but was maintained in the north and Scotland mostly due to external motivations such as “literacy, prestige, dialect borrowing, and word frequency” (2012: 35). Some southern speakers even redeveloped the contrast in the sixteenth and seventeenth century due to language contact with Scots pronunciation, which had enjoyed prestige at that time (Minkova 2012: 27–29).

However, it appears that the historical /w/-/ɹ/ contrast is weakening for an increasing number of Scottish English speakers. As early as in the 1980s, Macafee (1983: 32) noted that some speakers in Glasgow use either [ɹ] or [w] in words beginning with <wh->, which Stuart-Smith (2008: 63) confirms with Glasgow data from 1997. Likewise, Johnston (1997: 507) and Chirrey (1999: 227) reported that, at the end of the last century, only some speakers in Edinburgh were consistent in maintaining the /ɹ/-/w/ contrast, while others rather unpredictably varied between these two phonemes. By the beginning of the twenty-first century, Schützler found that /ɹ/ was realised as both [w] and [ɹ] by all Edinburgh speakers in his study (2010). Similarly, in Aberdeen, all the speakers analysed by Brato (2007, 2014) produce both [ɹ] and [w] in words beginning with <wh->.

Chirrey (1999) suggests that this eroding contrast “has a considerable time-depth”, as Edinburgh speakers as old as 73 did not fully preserve the contrast, and significant differences in the use of /ɹ/ and /w/ between older and younger speakers are consistently found across Scotland today: this holds true for the towns of Livingston (Robinson 2005), Glasgow (Timmins et al. 2004), Aberdeen (Brato 2014), and Edinburgh (Schützler 2010), where the younger speakers produce both /ɹ/ and /w/ for <wh->, while older speakers use fewer [w] for <wh-> than the younger ones. Interestingly, younger Glasgow speakers also occasionally produce [ɹ] for /w/ in words such as *wear*, *weather*, *wine* and *word* (Timmins et al. 2004).

The realisation of the /ɹ/-/w/ contrast in Scottish English is further influenced by the social factors class and gender. In the 27 middle-class Edinburgh speakers analysed by Schützler (2010), gender differences became apparent with a greater erosion of the /ɹ/ – /w/ contrast for male than for female speakers. In a project carried out in Glasgow, speech from 32 socially-stratified speakers, balanced for age and gender was collected and it was found that middle-class speakers, both adolescents and older speakers, use /ɹ/ more frequently than working-class

1 For historical reasons, *who* and its derivations are exceptions; Wells (1982: 409) claims that in South-East Scotland some <w> are also pronounced with /ɹ/ as in *weasel*.

speakers. Middle-class men in particular strongly favour the /ɹ/ variant, while working-class boys and girls show a very low use of /ɹ/ (Lawson and Stuart-Smith 1999; Stuart-Smith et al. 2007; Timmins et al. 2004). The same distribution across speakers of different classes was also found by Brato (2014) in his socially-stratified sample of 44 Aberdeen speakers.

In addition to sociolinguistic predictors for /ɹ/, Schützler (2010) and Brato (2014) also found language-internal factors influencing the realisation of the /ɹ/-/w/ contrast. Their findings show that the /ɹ/ variant favours stressed syllables and the postpausal position, which they explain by “the slightly more effortful /ɹ/” being more easily pronounced, least influenced by coarticulatory factors and receiving more attention with a preceding pause compared to an utterance-medial position (Brato 2014: 40; Schützler 2010: 15). More specifically, Schützler argued that the pre-aspiration effect of [ɹ] tends to be more pronounced in stressed syllables than in unstressed ones, and that the requirement of a pulmonic impulse in articulating [ɹ] favours a preceding pause, which allows speakers to produce it more easily.

While the speech of socially diverse Scottish speakers in the urban centres is well researched, including middle-class speakers from Edinburgh, Aberdeen and Glasgow, little is still known about whether the /ɹ/-/w/ contrast is maintained in the speech of other middle-class Scottish English speakers from other regions of Scotland as well. We would like to investigate middle-class speech from all over Scotland, which we refer to as Scottish Standard English (SSE) here. Nor does any prior research exist that explores whether the influencing factors, both social and linguistic, found for various types of urban Scottish English also determine the distribution of /w/ and /ɹ/ in the SSE. While many authors claim that the /ɹ/-/w/ contrast is present in contemporary SSE (Giegerich 1992: 36; Jones 2002: 27; Stuart-Smith 2008: 63), to the best of our knowledge, no empirical studies have been carried out so far to substantiate this. It is thus the first aim of this study to explore, based on a large corpus of SSE, whether the /ɹ/-/w/ contrast is maintained by SSE speakers. Some first evidence for a larger presence of /ɹ/ in the standard variety in Scotland might be deduced from those studies of urban Scottish English that compared two or more speaking styles differing in the level of formality. Some of them found that in the most formal speaking style, i.e. word lists, in which the speakers’ target pronunciation presumably is closest to the standard due to a high level of attention given to the pronunciation of each word, more [ɹ] are produced than in less formal styles such as passage reading and conversations (Brato 2007). Stuart-Smith et al. (2007) further found that the social factors of class and age as well as an age*gender interaction only significantly influenced the realisation of <wh> words in Glaswegian speech in the conversations but not in the word lists (where all speakers produced more [ɹ] than [w]), which points to an underlying representation of a prestigious accent of Scottish English that is shared by Scottish

speakers and that surfaces in formal contexts. Our first hypothesis to be tested therefore is that [ɱ] is realised frequently for <wh> in SSE and that many SSE speakers maintain a systematic contrast between /ɱ/ and /w/.

2 Acoustic properties of [ɱ]

In the literature, there is little consensus on the exact articulation of /ɱ/ in Scottish English: while Giegerich (1992: 36) refers to it as a voiceless bilabial fricative, Wells (1982: 408) classifies the consonant as a voiceless labial-velar fricative, but also suggests representing it as the diphone /hw/ or /xw/. By contrast, Robinson (2005: 186f.) defines it as a “voiceless lip-rounded consonant with audible friction at both velar and bilabial articulations”. Schützler (2010: 13), finally, describes the phonetic realisation of /ɱ/ as “a hybrid between an approximant and a fricative [that] can be interpreted as the combination of a voiced and a voiceless component, or at least as a partially devoiced approximant, thus: [xw]<[hw]<[w].”

All empirical studies on the realisation of the /ɱ/-/w/ contrast that have been carried out so far have observed ‘mixed’ or ‘in-between’ realisations of /ɱ/ in their auditory analyses (e.g. Brato 2007, 2014 for Aberdeen speakers²). Robinson (2005: 187), for instance, reports on realisations of /ɱ/ by Livingston speakers as a predominantly voiceless fricative with less bilabial articulation, which she referred to as [hɱ]. Timmins et al. (2004) and Stuart-Smith et al. (2007) found breathy-voiced labial velar approximants, which they labelled [wh], “a category of variants which sounded as if they were neither properly [w] nor [hw]” (Timmins et al. 2004: 16) and which also acoustically fell between the two sounds.

However, only one study so far has analysed the acoustic properties of the different realisations of /ɱ/ and /w/ in Scottish English. Lawson and Stuart-Smith (1999: 2543) found for Glaswegian teenagers aged 13–14 that [ɱ] often, but not always, showed a period of voiceless friction before the abrupt onset of the first (F1) and second (F2) formants. In those cases where no friction was visible in the spectrogram, sometimes the abrupt start of F1 and F2 was still found. The acoustic properties of [w] consisted of a low F1 with a low, weaker F2 and without any period of voiceless friction, while the ‘mixed’ categories that were observed in the auditory analyses did not seem to have any consistent acoustic correlates.

It is the second aim of this study to further explore the acoustic properties of the realisations of [ɱ] and [w] in SSE. We will test the hypothesis that in SSE a voiceless friction part will be consistently used for [ɱ], but not for [w]

² He also found [f/hw] for the Aberdeen speakers, where [f] is still a common dialectal variant of /ɱ/.

(Lawson and Stuart-Smith 1999: 2543). Moreover, we will hypothesise that the acoustic properties of [ʍ] and [w] are clearly distinct in SSE.

3 Aim and hypotheses

The aim of this study is to investigate the realisation of /ʍ/ and /w/ in SSE. While an ongoing /ʍ/-/w/ merger and its determining factors social class, age and gender as well as phonetic context have been found for various urban varieties of Scottish English, their actual realisation and distribution in contemporary SSE are still largely unknown. The present study will try to fill this research gap by testing the following hypotheses:

- /ʍ/ is largely present in SSE as claimed in Giegerich (1992: 36), Jones (2002: 27) and Stuart-Smith (2008: 63).
- There will be measurable voiceless friction for [ʍ], and the acoustic properties of [ʍ] and [w] will be distinct as suggested by Lawson and Stuart-Smith (1999: 2543).
- Both the social factors age and gender and language-internal phonetic factors influence the distribution of [ʍ] and [w] as found by Schützler (2010) and Brato (2014).

4 Method

The data were drawn from ICE-Scotland, the Scottish component of the International Corpus of English project (ICE; Greenbaum 1991) that aims to collect comparable corpora of all national varieties of English spoken around the world. ICE-Scotland is currently being compiled at the University of Münster (Schützler et al. 2017) and constitutes the first corpus of the ICE corpora family to contain time-aligned phonemic transcriptions. These were created in a two-step process: first, automatic phonemic annotations in SAMPA were created using WebMAUS (Schiel 2004) and were subsequently corrected manually by a team of five independent phonetically-trained transcribers including the two authors of this article. During the manual correction, misplaced phoneme boundaries were adjusted, missing or superfluous phonemes were inserted or removed respectively and incorrect SAMPA labels were corrected. Furthermore, each <wh> token was analysed auditorily for two rounds. A binary choice was made between the realisation of [w] and [ʍ] perceptually by a transcriber in the first round, and the second author checked all the transcriptions and corrected them whenever necessary. The sound quality was relatively good as only formal contexts were included for analysis.

As this study targets SSE and we expect it to be most likely to obtain speech that maintains the /m/ – /w/ contrast in formal contexts (see also Douglas 2020), only the following text categories of ICE-Scotland were searched: broadcast discussions, broadcast interviews, broadcast news and broadcast talks, legal presentations, demonstrations, non-broadcast talks, unscripted speeches as well as parliamentary debates.³ It is noteworthy that there is a likelihood that many speakers in this sub-corpus might have moderately to strongly anglicised speech in comparison with general middle-class Scottish population. After excluding all speakers who produced fewer than two <wh-> words and those for whom no information on their age is available, the final dataset comprised 138 speakers (aged 17–70, originating from all over Scotland, of whom 64 are female). They produced a total of 1,388 words containing initial <wh->.

The degree of acoustic periodicity of each token was measured by extracting the median of harmonicity (also referred to as harmonics-to-noise ratio, see e.g., Boersma 1993) using Praat (Boersma and Weenink 2017), which expresses the relationship between voiced (harmonic) and friction (noise) parts in a sound. Following Hamann and Sennema (2005), the median of harmonicity was measured with time steps of 0.01 s, a minimum pitch of 75 Hz, a silence threshold of 0.1 and 1 period per window. A value of 20 dB means that nearly all of the energy lies in the harmonic part, while for 0 dB the energy in the sound signal is equally distributed between voicing and friction. Twenty-three of the <wh-> tokens had to be excluded because of background noise or speaker overlap, which prevented a reliable acoustic analysis. A further 124 tokens had to be excluded because Praat yielded no values or values that had to be interpreted as measurement errors. Thus, the final token number of measured tokens of <wh-> words is 1,241 (see Appendix for a list of lexical items).

Furthermore, a matching number of tokens with initial <w-> (as e.g. in *water* and *went*) was extracted from the corpus. We attempted to include for each speaker an equal number of words beginning with <w->, matched in terms of preceding context (pausal or non-pausal) to the <wh-> words they had produced. However, this was not possible in all cases, as especially the necessary amount of post-pausal /w/ were not always available. The median of harmonicity of each <w-> token was also extracted using Praat, with measurement errors being excluded. There are 1,227 tokens of <w-> in the final dataset.

For the statistical analysis, linear mixed-effects regression models with MEDIAN OF HARMONICITY as dependent variable, and SPEAKER and WORD as random intercepts

³ Another practical reason for us to exclude informal categories was that, unlike the formal text categories, the acoustic quality of the informal ones was not always good enough for the demands of auditory and acoustic measurements.

were used. Fixed predictors including the social factors AGE and GENDER (m/f) and the internal factors preceding/following PHONETIC CONTEXT as well as the PRESENCE OF SCRIPT were tested. Preceding phonetic context was coded as either pause, voiced and voiceless; the following phonetic context as either the DRESS vowel as in the word *when*, the KIT vowel as in *which*, the LOT vowel as in *what* and the PRICE vowel as in *why*. Furthermore, the ICE text categories were divided into scripted (e.g., broadcast news) and unscripted (e.g., broadcast interviews). The models were fitted using the R-package {lme4} (Bates et al. 2015).

5 Results

Figure 1 compares the median of harmonicity of <wh-> and <w-> words observed in the data. Of the 1,241 <wh-> tokens, 37% had been transcribed in the corpus as [ʍ], while 63% had been transcribed as [w]. The pronunciation of all of the <w-> words extracted from the corpus had been transcribed as [w]. As Figure 1 shows, the realisation of [ʍ] for <wh-> displays the lowest harmonicity values ($mean = 5.48$, $sd = 3.62$, $n = 464$), which is followed by <wh-> tokens produced as [w] ($mean = 9.02$, $sd = 4.64$, $n = 785$), with <w-> tokens showing the highest median of harmonicity ($mean = 10.47$, $sd = 4.56$, $n = 1,227$). There were statistically significant differences between the mean values of the three groups as determined by a one-way ANOVA ($F(2, 2,473) = 214$, $p < 0.001$, $\omega^2 = 0.147$). Post hoc testing shows that the degree of median of harmonicity of [ʍ] is significantly lower than that of both <wh-> words transcribed in the corpus as [w] ($p < 0.001$, *Cohen's d* = -0.275) and <w-> words ($p < 0.001$, *Cohen's d* = -0.416). However, *Cohen's d* shows that the differences between the values of <wh->-[w] and <w-> tokens have a very small effect size ($p < 0.001$, *Cohen's d* = -0.143).

The /ʍ/-/w/ contrast in SSE was further analysed with a mixed-effects linear regression model in R using lme4 package in order to investigate potential language-internal and external factors affecting their distribution and acoustic properties. Table 1 illustrates the potential effects of LABEL (i.e. transcription in the corpus), AGE, GENDER, PRECEDING CONTEXT, FOLLOWING CONTEXT and PRESENCE OF SCRIPT in predicting the median of harmonicity of the tokens. The results of the full model show that three factors, both linguistic and social, exerted statistically significant effects: LABEL, GENDER and PRECEDING CONTEXT. The effects of AGE, FOLLOWING CONTEXT and PRESENCE OF SCRIPT, on the other hand, did not reach statistical significance. Figures 2–6 show the effects of these factors on the acoustic properties of <wh-> tokens predicted by the model.

Figure 2 shows that the acoustic properties measured for all <wh-> words are perfectly in line with their labels, i.e. the transcriptions in the corpus that had been

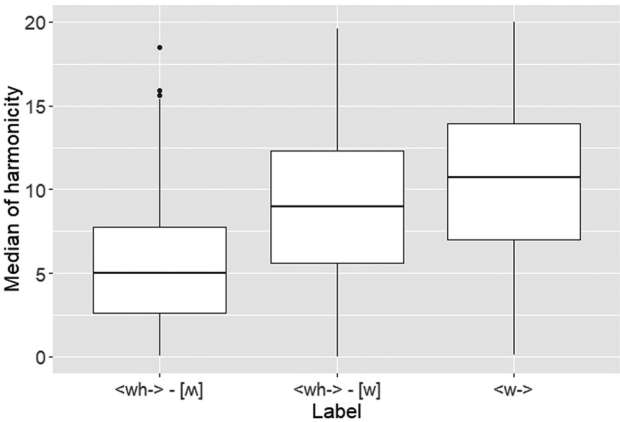


Figure 1: Median of harmonicity for <wh> words (transcribed in the corpus as [M] and [w] respectively) and <w> words.

Table 1: Linear mixed effect regression model of MEDIAN OF HARMONICITY as dependent variable, with SPEAKER and WORD as random intercepts.

Fixed effects	Levels	Estimate	Std. error	t-value	p <	
(Intercept)		7.276	0.923	7.887	0.000	***
LABEL	/w/	3.011	0.253	11.900	0.000	***
AGE		−0.031	0.016	−1.946	0.054	
GENDER	male	−1.086	0.399	−2.719	0.007	**
PRECEDING CONTEXT	voiced	1.781	0.247	7.217	0.000	***
	voiceless	0.628	0.341	1.841	0.066	
FOLLOWING CONTEXT	KIT	−1.197	0.738	−1.622	0.175	
	LOT	0.444	0.617	0.719	0.497	
	PRICE	−0.445	0.615	−0.725	0.485	
PRESENCE OF SCRIPT	unscripted	−0.248	0.382	−0.648	0.518	
Random effects	Type	Variance	Std. Dev			
SPEAKER	(Intercept)	2.5453	1.5954			
WORD	(Intercept)	0.4407	0.6639			
		Min.	Median	Max.		
Scaled residuals		−3.110	−0.001	3.233		

Levels of significance: $p < .05$ (*), $p < .01$ (**), $p < .001$ (***).

made based on an auditory analysis of the data. Thus, the <wh> tokens transcribed as [w] and [M] respectively have distinct degrees of median of harmonicity, with the values of [w] being significantly higher than those of [M] ($p < 0.001$).

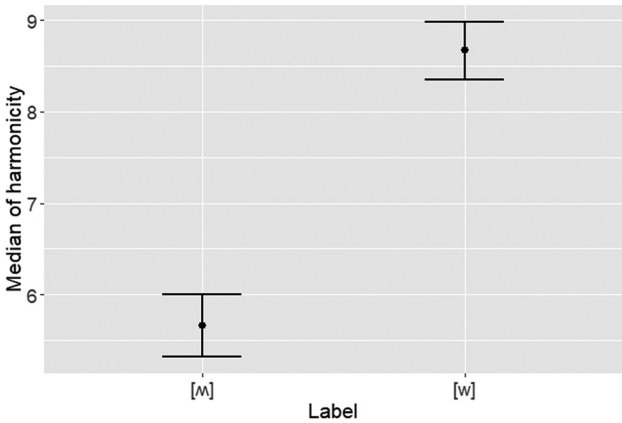


Figure 2: LABEL and the acoustic properties of <wh> tokens.

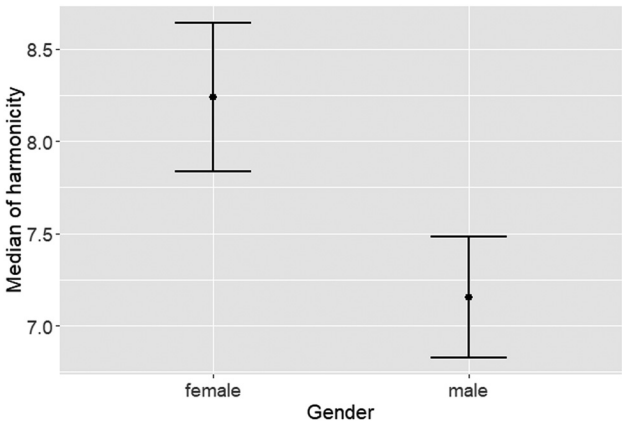


Figure 3: Effects of GENDER on the acoustic properties of <wh> tokens.

Figure 3 illustrates that the acoustic properties of the initial consonant of <wh> words are further constrained by speaker gender. More specifically, female speakers produced a significantly higher median of harmonicity, thus more [w]-like realisations containing less friction, than their male counterparts ($p = 0.007$).

The results of the mixed effects model show that the language-internal predictor PRECEDING CONTEXT also plays a statistically significant role in predicting the degree of median of harmonicity of the <wh> words produced by SSE speakers. Figure 4 shows that [ɰ] is most likely preceded by a pause, followed by the

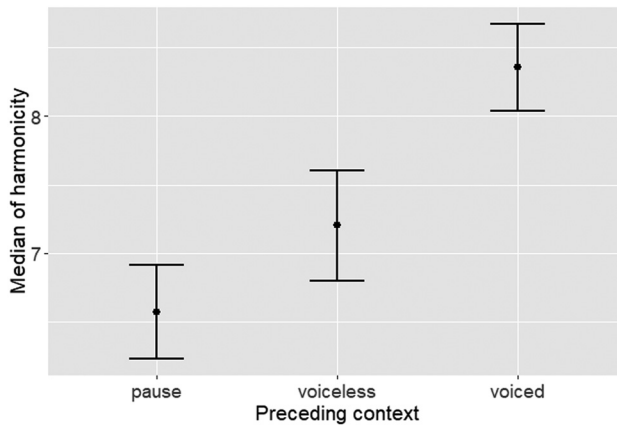


Figure 4: Effects of PRECEDING CONTEXT on the acoustic properties of <wh-> tokens.

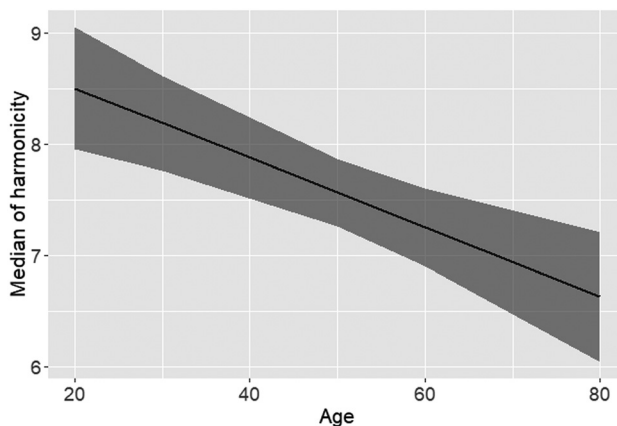


Figure 5: Effects of AGE on the acoustic properties of <wh-> tokens.

preceding context of a voiceless sound, while a preceding voiced context does not promote the production of [ɹ] ($p < 0.001$).

The factors AGE, FOLLOWING CONTEXT and PRESENCE OF SCRIPT did not reach statistical significance in the model. However, Figure 5 illustrates a tendency of the median of harmonicity to decrease with increasing speaker age. Older speakers tend to produce more [ɹ] compared with younger speakers.

Equally, Figure 6 suggests potential effects of the following context on the realisation of <wh-> words. The production of [ɹ] is favoured when followed by a KIT vowel (e.g. in *which*). Following PRICE (e.g. in *why*, *while*), DRESS (e.g. in *when*,

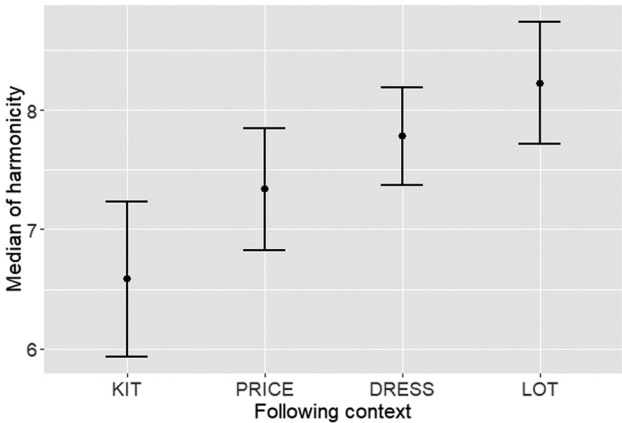


Figure 6: Effects of FOLLOWING CONTEXT on the acoustic properties of <wh-> tokens.

where) or LOT (e.g. in *what*) constitute less promoting contexts for the realisation of <wh-> as [m].

Figure 7 shows the realisation of <wh-> and <w-> words by individual speakers observed in the data. For this, we selected those 41 speakers who produced at least ten <wh-> tokens in the corpus.

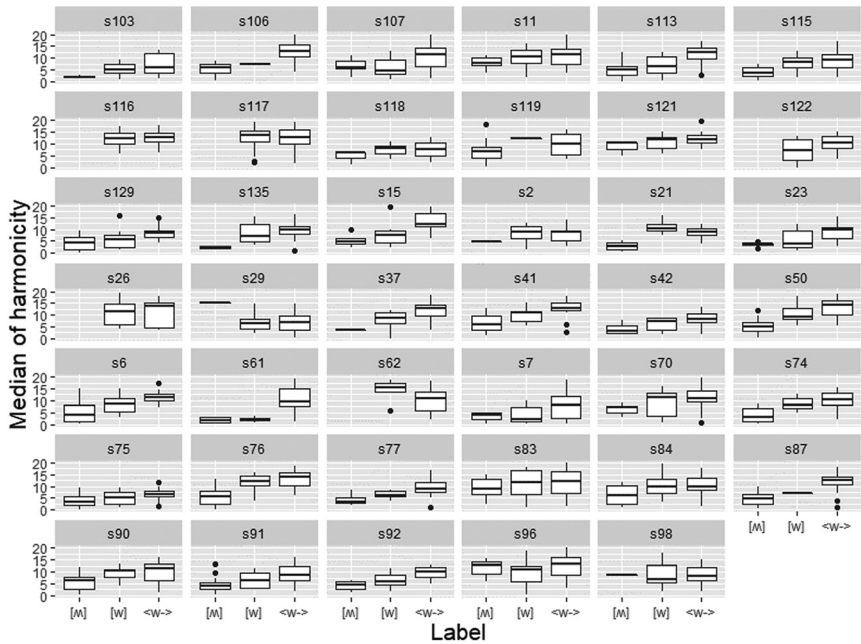


Figure 7: Individual realisations of <wh-> words transcribed as [w], <wh-> words transcribed as [m] and all <w-> words in the corpus.

As can be seen in Figure 7, considerable inter-speaker variation exists in the realisation of <wh-> and <w-> words. Among the 41 speakers, 36 produced both [w] and [ɱ] for <wh-> words, with [ɱ] typically having a lower median of harmonicity than [w], except for s107, s29, s7, s96 and s98, where no acoustic difference between their [w] and [ɱ] realisations for <wh-> words was found. Furthermore, for most speakers [w] produced for <wh-> words is acoustically different from [w] for <w-> words. To be more specific [w] occurring in <wh-> words showed a lower median of harmonicity than [w] occurring in <w-> words (except for s21 and s29). The speakers seemed to consistently maintain measurable voiceless friction in producing [w] for <wh-> words, but not for <w-> words.

The remaining five speakers did not produce [ɱ] at all for <wh-> words (s116, s117, s122, s26 and s62) and thus merged /ɱ/ and /w/ entirely in their speech. Interestingly, unlike those who maintained the contrast [w] produced by these five speakers showed very similar acoustic properties for <wh-> and <w-> words. Specifically, both were characterised by a relatively high degree of median of harmonicity that is having very little friction in their consonants. S62 even showed a higher median of harmonicity for <wh-> than for <w->. As the five speakers are of different genders, from different age groups and regions, there is nothing obvious to explain the observed patterning.

6 Discussion

This study investigated the presence of the /ɱ/-/w/ contrast and their acoustic properties in the formal text categories of ICE-Scotland. The results show that our first hypothesis, i.e. that /ɱ/ is still present in the Standard Scottish variety of English as claimed in Giegerich (1992: 36), Jones (2002: 27) and Stuart-Smith (2008: 63) can be confirmed: 37% of the auditorily analysed <wh-> tokens were realised with the traditional [ɱ] variant. However, significant inter-speaker variability was found in the corpus. About 12% of the SSE speakers seem to have a complete merger of /ɱ/ and /w/, producing [w] categorically for all <wh-> words. All of the other speakers produced both [w] and [ɱ] in words beginning with <wh->, thus showing no difference between these SSE speakers and the middle-class speakers from Edinburgh and Aberdeen studied by Schützler (2010) and Brato (2007, 2014). It is important to note that none of the speakers in the corpus exclusively used the traditional variant [ɱ] for <wh->, which would have reflected the maintenance of the traditional Scottish /w/-/ɱ/ contrast.⁴ The speech of the speakers of the standard variety of English

⁴ Although it might occasionally be difficult to find unambiguous [ɱ] variants in spontaneous speech because of potential coarticulatory effects of the following vowel, intervocalic context or unstressed position.

spoken in Scotland thus evidences an ongoing merger of /w/ and /ɹ/. In other words, it appears that [ɹ] is only one variant of realising <wh-> for the speaker group studied here.

The results of the acoustic analyses of the corpus confirm our second hypothesis, which predicted that [ɹ] and [w] have distinct acoustic properties, with more friction being produced for [ɹ] than for [w] in <wh-> words. However, the acoustic measurements also showed that [ɹ] and [w] are not categorically distinct and that a similar amount of friction can be found in some tokens that had been transcribed as [ɹ] and [w] in the corpus. We interpret this as corroboration of the impressionistic observations made by Stuart-Smith et al. (2007) and Timmins et al. (2004), who described variants of <wh-> that sounded “as if they were neither properly [w] nor [hw]” (Timmins et al. 2004: 16).

The most striking of the results of the acoustic analysis are however that, overall, there are significant, albeit small differences in the median of harmonicity between a [w] occurring in <wh-> and in <w-> words. Thus [w] when produced by an SSE speaker at the beginning of *witch* differs acoustically from [w] being produced at the beginning of *which*. This finding suggests that SSE speakers do make a difference between <wh-> and <w-> words, even if they do not produce a distinctly voiceless labiovelar fricative in all instances of the <wh-> words. It is worth mentioning that the observed difference between [w] realised in <wh-> and <w-> words can only be found for SSE speakers who alternate between [ɹ] and [w] for <wh-> words. Those that completely merged /ɹ/ and /w/, on the other hand, showed similar acoustic properties of [w] occurring in <wh-> and <w-> words. Future research could investigate whether this acoustic difference between ‘different kinds of [w]’ also exists in other varieties of Scottish English. The assumption that this is some way triggered by the orthographic differences between <wh-> and <w-> words, however, cannot be confirmed by the results of our corpus analysis. When comparing the realisation of <wh-> words in free speech to scripted speech, i.e. texts being read out by speakers as in the category broadcast news, no acoustic differences were found.

The third research aim of the present study was to explore possible social and language-internal factors influencing the realisation and distribution of [ɹ] and [w] in SSE. The output of our statistical model shows that both speakers’ gender and preceding phonetic contexts exert main effects on predicting the acoustic patterns of [ɹ]. In terms of gender, male speakers produced a significantly lower degree of median of harmonicity, thus more instances of [ɹ], than their female counterparts. This shows that gender plays a similar role in SSE to middle-class Glasgow English, where men strongly favour the production of the /ɹ/ variant (Lawson and Stuart-Smith 1999; Stuart-Smith et al. 2007; Timmins et al. 2004). Gender, however, seems to influence SSE differently than the Edinburgh speakers

studied by Schützler (2010), who reported that male speakers in Edinburgh tend to merge /m/ and /w/.

Our results further show that the realisation of [m] in SSE is constrained by a linguistic factor, namely preceding phonetic context. A preceding pause or a preceding voiceless sound was found to favour the production of /m/, with a preceding voiced context being the least promoting of the variant. This corroborates the findings of Schützler (2010) and Brato (2014) for Edinburgh and Aberdeen speakers that the /m/ variant favours postpausal position since “a pause preceding /m/ will give the speaker time to articulate the slightly more effortful [m]” (Schützler 2010: 15). It is also unsurprising that the voiceless preceding context favours the realisation of /m/ since phonetically there is a measurable voiceless friction for [m].

Although the effects of age did not reach statistical significance in our model, there is a tendency for older SSE speakers to produce more /m/ variants than younger speakers. This finding suggests similar effects of age on SSE to was found for the English spoken by Scots in Livingston (Robinson 2005), Glasgow (Timmins et al. 2004), Aberdeen (Brato 2014) and Edinburgh (Schützler 2010), where /m/ was consistently shown to be favoured by older speakers and younger speakers generally appeared insensitive to the /m/-/w/ contrast.

7 Conclusion

This study constitutes the first corpus-based study of the phonological and phonetic properties of /m/ and /w/ in SSE. It has shown that phonemically annotated corpora can prove a precious source for the analysis of the properties of accents of English. We were able to confirm with acoustic measurements not only previous impressionistic observations of ‘mixed’ or ‘in-between’ realisations of [m] and [w], but also found the first evidence of acoustic differences between [w] realised for <wh-> and for <w-> in SSE. Our finding suggests that most speakers do maintain a representational, i.e. phonemic difference between /m/ and /w/ although this might be almost merged for some on the phonetic level.

Acknowledgements: We are very grateful for the insightful comments by our two reviewers on an earlier version of this article.

Research funding: This research was funded by the German Research Council (DFG; GU 548/13-1).

Appendix: Word frequency list

Word	Frequency	Percent
What	330	26.5
Which	325	26.1
When	199	16
Where	126	10.2
Why	66	5.4
Whether	42	3.4
What's	30	2.4
While	28	2.3
Whatever	13	1
Whilst	9	0.7
Anywhere	6	0.5
Meanwhile	6	0.5
Elsewhere	5	0.4
Whisky	5	0.4
Somewhere	4	0.3
White	4	0.3
White	3	0.2
White's	2	0.2
Whitehouse	2	0.2
Whyte	2	0.2
Everywhere	2	0.2
Nowhere	2	0.2
Overwhelming	3	0.2
Somewhat	3	0.2
Whatsoever	3	0.2
Whereabouts	2	0.2
Whereby	2	0.2
Wherever	3	0.2
Whittled	2	0.2
Worthwhile	2	0.2
Whitehall	1	0.1
Whiting	1	0.1
Overwhelmingly	1	0.1
Whatsoever	1	0.1
Whenever	1	0.1
Where'd	1	0.1
Where's	1	0.1
Whereas	1	0.1
Whined	1	0.1
Whispered	1	0.1
Total	1,241	100

References

- Bates, Douglas, Martin Mächler, Ben Bolker & Steve Walker. 2015. Fitting linear mixed-effects models using lme4. *Journal of Statistical Software* 67(1). 1–48.
- Boersma, Paul. 1993. Accurate short-term analysis of the fundamental frequency and the harmonics-to-noise ratio of a sampled sound. *Proceedings of the Institute of Phonetic Sciences*, 17, 97–110. Amsterdam: University of Amsterdam.
- Boersma, Paul & David Weenink. 2017. *Praat: Doing phonetics by computer* [Computer program]. Version 6.0.31. Available at: <http://www.praat.org/>.
- Brato, Thorsten. 2007. Accent variation in adolescents in Aberdeen: First results for (hw) and (th). In Jürgen Trouvain & William Barry (eds.), *Proceedings of the 16th International Congress of Phonetic Sciences*, 1489–1492. Saarbrücken: Universität des Saarlandes.
- Brato, Thorsten. 2014. Accent variation and change in North-East Scotland: The case of (hw) in Aberdeen. In Robert Lawson (ed.), *Sociolinguistics in Scotland*, 32–51. Houndmill & New York: Palgrave Macmillan.
- Chirrey, Deborah. 1999. Edinburgh: Descriptive material. In Paul Foulkes & Gerard Docherty (eds.), *Urban voices: Accent studies in the British Isles*, 223–229. London: Arnold.
- Douglas, Fiona. 2020. English in Scotland. In Cecil L. Nelson, Zoya G. Proshina & Daniel R. Davis (eds.), *Blackwell handbooks in linguistics. The handbook of world Englishes*, 17–33. New York: John Wiley & Sons.
- Giegerich, Heinz. 1992. *English phonology: An introduction*. Cambridge: Cambridge University Press.
- Grant, William. 1914. *The pronunciation of English in Scotland*. Cambridge: Cambridge University Press.
- Greenbaum, Sidney. 1991. ICE: The international corpus of English. *English Today* 7(4). 3–7.
- Hamann, Silke & Anke Sennema. 2005. Acoustic differences between German and Dutch labiodentals. *ZAS Papers in Linguistics* 42. 33–41.
- Johnston, Paul. 1997. Regional variation. In Charles Jones (ed.), *The Edinburgh history of the Scots language*, 433–513. Edinburgh: Edinburgh University Press.
- Jones, Charles. 2002. *The English language in Scotland: An introduction to Scots*. East Linton: Tuckwell Press.
- Lawson, Eleanor & Jane Stuart-Smith. 1999. A sociophonetic investigation of the ‘Scottish’ consonants (/x/ and /m/) in the speech of Glaswegian children. *Proceedings of the International Congress of Phonetic Sciences*, 2541–2544. San Francisco: University of California, Berkeley.
- Macafee, Caroline. 1983. *Glasgow*. Amsterdam & Philadelphia: John Benjamins.
- Minkova, Donka. 2012. Philology, linguistics, and the history of [hw]~[w]. *Studies in the history of the English language II*, 7–46. Berlin & New York: De Gruyter Mouton.
- Robinson, Christine. 2005. Changes in the dialect of Livingston. *Language and Literature* 14(2). 181–193.
- Schiel, Florian. 2004. MAUS goes iterative. *Proceedings of the IV International Conference on Language Resources and Evaluation*, 1015–1018. Lisbon: University of Lisbon.
- Schützler, Ole. 2010. Variable Scottish English consonants: The cases of /m/ and non-prevocalic /r/. *Research in Language* 8. 5–21.

- Schützler, Ole, Ulrike Gut & Robert Fuchs. 2017. New perspectives on Scottish Standard English. Introducing the Scottish component of the International Corpus of English. In Sylvie Hancil & Joan Beal (eds.), *Perspectives on northern Englishes*, 273–301. Berlin: Mouton de Gruyter.
- Stuart-Smith, Jane. 2008. Scottish English: Phonology. In Bernd Kortmann & Clive Upton (eds.), *Varieties of English. The British Isles*, 48–70. Berlin & New York: Mouton de Gruyter.
- Stuart-Smith, Jane, Claire Timmins & Fiona Tweedie. 2007. “Talkin’ Jockney?”: Accent change in Glaswegian. *Journal of Sociolinguistics* 11. 221–261.
- Timmins, Claire, Fiona Tweedie & Jane Stuart-Smith. 2004. *Accent change in Glaswegian (1997 corpus): Results for consonant variables*. Glasgow: University of Glasgow.
- Wells, John. 1982. *Accents of English*, vol. 1. Cambridge: Cambridge University Press.