

The vowel system of Western Eivissan youngsters

rMA thesis

Student: Alessandro Pecoraro 12523615

Thesis supervisors: dr Silke Hamann, dr Francesc Torres-Tamarit

Department of Linguistics

Academic year 2024/2025



UNIVERSITY OF AMSTERDAM
Amsterdam Center for Language and
Communication

Table of Contents

Acknowledgements	2
1. Introduction	3
1.1 Vowel quality	4
1.1.1 Pillai scores	6
1.2 Vowel duration	6
1.3 Sociolinguistic influences on Catalan	7
1.4 Research questions	9
2. Methods	10
2.1 Participants	10
2.2 Stimuli	11
2.3 Materials	12
2.4 Procedure	12
2.5 Data analysis	13
2.5.1 Setting up the analyses	13
2.5.2 Vowel quality analysis	13
2.5.3 Vowel duration analysis	14
3. Results	15
3.1 Bilingual Language Index	15
3.2 Vowel quality results	16
3.3 Vowel duration results	20
4. Discussion	26
5. Conclusions	27
Appendices	29
References	74

Acknowledgements

I would like to thank my supervisors dr Silke Hamann and dr Francesc Torres-Tamarit for their support during these many months of research. Additional thanks to dr Paul Boersma for providing feedback on the analytical parts of the study.

It feels like a given, but it is not: many many thanks to Jaume Fitó and Roser Marí at the *Sant Agustí Institute for Secondary Education* for selecting the participants and of course to each and every one of the students who provided instrumental data for a phonetic analysis of Western Eivissan.

Finally, I am greatly indebted to my family and to my girlfriend Marta for loving me the way I need.

1. Introduction

Catalan is a Western Romance language spoken primarily along the eastern coast of the Iberian peninsula in Catalunya and Valencia, as well as on the Balearic Islands, the Roussillon region of France and Alghero in Sardinia (Eberhard et al., 2025).

On the Balearic Islands, as is the case for the other Catalan-speaking regions, further subvarieties are present. Eivissan Catalan, spoken on the homonymous island, is one of said subvarieties and the topic of this paper. Despite some research being published on the vowel system differences within Eivissa (Alcover, 1908; Torres Torres, 1983), only the eastern half of the island has been studied instrumentally (Hamann & Torres-Tamarit, 2023). The present paper thus sets out to document and describe the inventory of stressed vowels of the western part of Eivissa as present in the speech of youngsters, to provide a clearer picture of the reported plurality of speech systems within the Eivissan Catalan subvariety. Following is a list of examples for each vowel of Western Eivissan, taken from the stimuli set used in this paper. The 7-vowel inventory is the same as proposed in Torres Torres (1983) for this variety of Catalan.

(1) a.	/i/	<i>fibra</i>	[ˈfiβrə]	‘fiber’
b.	/e/	<i>febre</i>	[ˈfeβrə]	‘fever’
c.	/ɛ/	<i>peu</i>	[pɛw]	‘foot’
d.	/ɛ/ (/ə/ ¹)	<i>beu</i>	[bɛw]	‘(s)he drinks’
e.	/a/	<i>pau</i>	[paw]	‘peace’
f.	/ɔ/	<i>bou</i>	[bɔw]	‘ox’
g.	/o/	<i>pou</i>	[pow]	‘water well’
h.	/u/	<i>fum</i>	[fum]	‘smoke’

A key difference between this and Hamann & Torres-Tamarit’s (2023) studies on the one hand and most exploratory phonetic research on the other is the use of young participants. Dialectological work in general has traditionally focussed on older speakers as being the holders of ‘purer’ forms of local dialects (Lippi-Green, 1989; Beaulieu & Cichocki, 2002). The choice of young speakers in this project is motivated by the desire to compare the phonetic situation Alcover’s (1908) and Torres Torres’s (1983) papers report with what is present at the time of

¹ As will be explained in section 2.2, a distinction was made between /ɛ/ words that historically had /ə/ in the stimuli so as to facilitate analysis.

20 writing. Further, if any changes are to be found over a century after the earliest reports of the vowels of Eivissan, they would be even likelier found in young speakers than in older ones².

1.1 Vowel quality

As of now, the vowel phonemes of Catalan have been described as differing from each other in terms of frequency ranges, as no variety of Catalan has been found to make a phonological distinction between short and long vowels.

25 With that in mind, the main phonological difference between the numerous varieties of Catalan is the exact number of vowels allowed in stressed position, spanning from as few as five to as many as eight. Table 1 below groups different varieties of Catalan based on their vowel systems.

Balearic Catalan	/i, e, ε, ə, a, o , ɔ, u/
Western and Eastern Catalan, Western Eivissan and Eastern Minorcan	/i, e, ε, a, o , ɔ, u/
Eastern Eivissan	/i, e, E, a, O, u/
Felanitx Majorcan Catalan	/i, E, ə, a, o , ɔ, u/
Girona Catalan	/i, e, ε, a, O, u/
Northern and Sitges Catalan	/i, E, a, O, u/

Table 1: the stressed vowel inventories of different varieties of Catalan (Recasens & Espinosa, 2009 and Hamann & Torres-Tamarit, 2023). In Felanitx Majorcan Catalan, Northen Catalan and Sitges Catalan, /E/ is a front mid vowel resulting from the merger of /e/ and /ε/, while in Eastern Eivissan it is the merger of /ə/ and /ε/. /O/ is the result of the merger between /o/ and /ɔ/ in all the above varieties.

It is of interest to note that the front mid vowels behave differently in different systems.

30 In some systems, e.g. Balearic Catalan, both /e/ and /ε/ as well as the central vowel /ə/ are distinct phonemes. In all other varieties in Table 1, at least one of these three vowels is not present. In some, two of the three vowels have merged into one vowel. These front mergers can also have a back mid vowel merger counterpart within the same system, though this is not always the case.

² Conversely to what is said above regarding exploratory phonetic research, in variationist sociolinguistics, younger speakers have been largely preferred as the supposed main vehicles of variation and change in speech (Pichler et al., 2018).

35 The instability of stressed /ə/ described by Hamann & Torres-Tamarit (2023) might very well be the reason for the frequency of mergers with or between the front mid vowels. Indeed, Recasens (2019, 2023) notes that Latin /e/ (Ī, Ē) shifted to /ə/ in Balearic Catalan through continental Eastern Catalan. Later, while the shift from /ə/ to /ɛ/ happened sooner in Eastern Catalan, the Balearic Catalan varieties kept /ə/. Today, there appears to be a split between
40 varieties of Balearic that keep a central /ə/ like Majorcan (Recasens & Espinosa, 2006) and ones whose /ə/ went on to merge with /ɛ/ like Eastern Eivissan (Hamann & Torres-Tamarit, 2023). This shift from /e/ to /ə/ might be the reason why stressed as well as unstressed /ə/ in Eastern Eivissan are rather front (so much so that stressed /ə/ merged with /ɛ/). A complementary, sociolinguistic explanation for the mid vowel mergers in Catalan is given further in section 1.3
45 below.

 Focussing on the stressed vowels of Eivissan Catalan, the picture is once again not unitary. The standard for the whole of Balearic Catalan points to eight stressed vowels, namely /i, e, ɛ, ə, a, ɔ, o, u/. Hamann & Torres-Tamarit (2023), however, show that Eastern Eivissan has reduced its inventory to a six vowel system – /i, e, E, a, O, u/³ – whose back vowel merger was
50 hinted at by Torres Torres (1983:67) and whose /E/ merger is novel in the eastern variety. As for Western Eivissan, Torres Torres (1983) reports a similar merger of /ə/ and /ɛ/ as having already taken place – since at least the beginning of the century, according to Alcover (1908). Crucially, the two /ə/-/ɛ/ mergers are yet to be confirmed as identical or even comparable. Torres Torres (1983) reports the merger to result in /ɛ/, although no instrumental data is given either in
55 his paper or Alcover's (1908). The phonetic status of the merger in Hamann & Torres-Tamarit (2023), whose report follows the acoustic analysis of data from a production experiment, is not as clearly fully aligning with /ɛ/. In their study, some speakers' merged /E/ aligns with either the pretonic or final schwa (both unstressed), with most speakers' /E/ instead overlapping with both unstressed schwas⁴. Of note is that the Eastern Eivissan unstressed schwas are themselves
60 more fronted than a prototypical central vowel as schwa normally is. The overall picture is thus of a fairly fronted merged /E/ in Eastern Eivissan, likely due to the fronted starting point of schwa, both stressed and unstressed.

 Because of a lack of instrumental data to clarify the vowel system of Western Eivissan, this project analysed all Balearic stressed vowels across a selection of words.

³ /E/ and /O/ are used by the authors for the vowels resulting from the mergers of /ɛ/ and /ə/ and that of /ɔ/ and /o/ respectively.

⁴ Two of their speakers' /E/ was clearly separate from both unstressed schwas.

1.1.1 Pillai scores

Pillai scores, also known as the Pillai-Barlett trace, for Multivariate Analyses of Variance (MANOVAs) were used to measure the degree of overlap between the members of the two vowel pairs. Pillai scores range from 0 to 1, where 0 means that the two items being compared are not different at all and 1 means they are completely different. Pillai scores have extensively been used to analyse vowel mergers (for English, see Schmidt et al. 2021 *inter alia*; for Galician, see Amengual & Chamorro, 2016; for Italian, see Nadeu & Renwick, 2016; for Bangla, see Islam & Ahmed, 2020; for Hawai‘ian, see Ketting, 2021; for Swiss German, see Joo et al., 2018; for Austrian German, see Sloos, 2014;), though some alternative uses are seen in Fung & Lee (2019) and Tse (2018) for analysing tones in varieties of Cantonese and in Regan (2020) for analysing a fricative split in Andalusian Spanish.

Before Pillai scores, Euclidean distances were employed as a metric for mergers between the mean F1-F2 points of one vowel and those of another. Because means were used, this method does not consider the degree of overlap between vowel categories and the distribution of all tokens for said categories. Pillai scores, however, have been found to be the overall best option for measuring overlap between vowels (Kelley & Tucker, 2020). Despite that, existing research tends not to agree on a cut-off point for the distinction between merged and non-merged sounds, or not even explicitly establish one for the study at hand. In Stanley & Sneller (2023), the authors simulate re-iterated analyses of vowel overlap with different sample sizes and conclude that sample size is inversely proportional to the analysis’ resulting Pillai score. Based on this finding, they propose a formula for calculating a Pillai score threshold based on any given sample size: $p_{95} = e/m$, ‘where p_{95} is the 95th percentile of Pillai scores given the average sample size per group m ’ (Stanley & Sneller, 2023:61). This formula will be used to establish Pillai score thresholds per participant.

1.2 Vowel duration

Even in languages without a phonological long-short vowel contrast, a feature found to play a complementary role in distinguishing consonantal minimal pairs is the duration of the stressed syllable’s vowel (for English, see Ladefoged & Johnson, 2010; for Dutch, see Ernestus & Baayen, 2006; for Brazilian Portuguese, see Ribeiro, 2017 and Alves & Brisolara, 2020; for Italian, see Renwick, 2024).

The notion that vowel duration is linked with vowel height is well known. Said phenomenon is explained as mechanical in nature, with lower vowels being longer than higher

95 vowels as a consequence of the longer time needed to lower the jaw to articulate the former as opposed to the latter vowels.

That said, more than articulation has been found to be at play. Toivonen et al. (2015) report two experiments on English and Swedish where the mechanically informed duration difference between high and low vowel categories did not manifest within vowel categories (speaker and flanking consonants were controlled for). If lower F1 values (taken as proxy for jaw height) for realisations of the same vowel category do not correlate with a longer vowel duration, more than jaw movement must be at play. Indeed, Bai & Scarborough (2025) present evidence for lexical influences such as minimal pairhood and neighbourhood density of words on vowel duration.

105 To the author's knowledge, no analysis of vowel duration differences in Catalan has been carried out. It would then be of interest for the question of vowel mergers in Eivissan Catalan to study whether vowel duration differences are present in the consonantal minimal pairs with the previously reported merged vowels (further explanation of the choice of *consonantal* minimal pairs is given in section 2.2 below). If that were the case, the implication would be that the merged vowels of Eivissan are in fact not yet merged but rather nearly merged. In this paper, the differentiating role of vowel duration was searched for in Western Eivissan as a secondary area of study, utilising the consonantal minimal pairs present in our stimuli. Due to the preliminary nature of this study, the role of neighbourhood density on vowel duration in Eivissan was not investigated.

1.3 Sociolinguistic influences on Catalan

115 Catalan, of any variety, does not exist in a vacuum. All varieties of Catalan, with the possible exception of Alguerés (spoken in the city of Alghero in Sardinia) share Spanish as the main contact language. Spanish does not have a mid vowel contrast, instead including only high mid vowels in its vowel inventory.

In Simonet (2019), the author gives a comprehensive review on the role of Catalan-Spanish bilingualism on phonetic perception and articulation of the Catalan mid vowel contrast. In it, many perception as well as production studies show that Spanish-dominant bilinguals are either unable to perceive the difference between high and low mid vowels of Catalan (for perception research, see Pallier et al., 1997 and Navarra et al., 2005; for production research, see Simonet, 2011) or not as capable as Catalan-dominant bilinguals (Sebastián-Gallés & Soto-Faraco, 1999).

These studies point to the importance of the first year(s) of life for speaker-listeners' ability to tell apart and produce high mid and low mid Catalan vowels. Still, the author acknowledges production studies by Amengual (2016) and Mora & Nadeu (2012) for challenging the assumption of the first year(s) of life as the *sole* predictor of speakers' ability to differentiate the two sounds in production as well as perception. Amengual's (2016) Spanish-dominant bilinguals were indeed able to produce two different vowel categories for the high and the low mid vowels of Catalan, though not as phonetically distinct as the Catalan-dominant bilinguals'. Mora & Nadeu's (2012) exclusively Catalan-dominant bilinguals were split into two groups: those who used Catalan 90% of the time and those whose used Catalan 60% of the time. They found that, despite both groups keeping high and low mid vowels separate in their production, the group who used Catalan 60% of the time produced mid vowels that were closer to each other than they were in the speech of the group who used Catalan 90% of the time.

The two studies described above differ from previous perception research in at least two ways. Firstly, the bilinguals sampled by Amengual (2016) and Mora & Nadeu (2012) were specifically from monolingual Catalan families *and* predominantly spoke Catalan in their daily lives (as was the case in Simonet, 2011). Secondly, Amengual's (2016) study differs from previous production work such as Simonet (2011) in that the former uses a picture-naming task, whereas a sentence-reading task was conducted by the participants in the latter.

These different and contradicting results do not take away from the hypothesis that early exposure to a language is a strong predictor of differentiation of high and low mid vowels. It is clear, though, that the degree of language use throughout speakers lives also play a crucial role in the degree of differentiation of said vowels. Finally, it is important to acknowledge how different tasks, in conjunction with different participant samples, impact the results of studies looking into the same phenomenon.

Focussing on Eivissa, qualitative and quantitative data by Castell et al. (2023) seem to categorise young speakers on the island as being more balanced bilinguals than other young Balearic speakers. This label is not so much to do with the degree of proficiency in either Catalan or Spanish; rather it reflects young Eivissan speakers' openness to use either language or even switch from one language to the other depending on the communicative context, such as whether a different language or a different variety of the same language was used by the interlocutor, the relationship they have with said interlocutor – e.g. stranger v. friend, teacher v. peer – or even the social context – e.g. school, home, work etc.

Keeping all this in mind, the present research sought to minimise the influence of Spanish on the Catalan mid vowels through a number of preventive measures. As will be

160 explained in detail in section 2.1, participants were sampled from a school in Western Eivissa by the school staff. On top of that, the researcher who collected the data, dr Francesc Torres-Tamarit, is a native speaker of the Eastern, urban variety of Eivissa. Admittedly, the researcher's own speech could have potentially affected the participants' speech, due to the former speaking a different variety of Eivissan from the participants' (a variety which was notably found to
165 merge back mid vowels, as opposed to what was known about Western Eivissan before the writing of this paper). Still, neither this difference in speech nor the status of adult-teacher-stranger on the part of the researcher prevented the participants from speaking Catalan throughout the whole experiment.

1.4 Research questions

This research focussed on the island of Eivissa to record and document the vowel system spoken
170 in its western section and to establish whether mergers comparable to those attested in Eastern Eivissan Catalan are present in its western counterpart. The project's primary research question is: *what are the stressed vowels in the vowel inventory of young Western Eivissan Catalan speakers?* This research question inherently answers the question of whether the /ə/-/ɛ/ merger reported by Alcover (1908) and located by Torres Torres (1983) in a number of sites in the
175 western portion of the island is present in said portion to this day, as well as evaluating the status of the /o/-/ɔ/ vowel pair, which as of now is only confirmed to have merged in the east of Eivissa. Additionally, a preliminary comparison of the data with those from Hamann & Torres-Tamarit (2023) as regards Eivissan Catalan is made possible. So is a comparison with data from previous research on the other varieties of Catalan by Recasens & Espinosa (2006).

180 The null hypothesis for this research question is that the vowel inventory of Western Eivissan does not include a stressed schwa, instead having a (historically already) merged /E/ in the system of seven stressed vowels /i, e, E, a, ɔ, o, u/, thus remaining the same as that proposed by Torres Torres (1983). The present study's hypothesis is that a back mid vowel merger is actually present, as growing contact between the two varieties of Eivissan likely led
185 the western variety to absorb the eastern back merger. No predictions for the specific phonetic quality of the merged vowels were made, i.e. to which of the original vowels in each pair the merged vowels were closer.

If the phonetic quality of the aforementioned vowels does indeed point to the established merger, the secondary research question is then: *in Western Eivissan, do vowel duration
190 differences occur between consonantal minimal pairs whose stressed vowels are merged?*

The starting hypothesis was that a small difference in vowel duration would be found between the members of both vowel pairs, either as a secondary cue if the vowels differ in quality or as a compensatory cue in case the vowels have merged in quality.

This research project complements previous research in supplying academic knowledge of inner varieties of a minoritised language such as Catalan.

2. Methods

2.1 Participants

Ten participants were sampled at the *Sant Agustí Institute for Secondary Education* located in the village of *Sant Agustí des Vedrà*, which is in the municipality of *Sant Josep de sa Talaia*, in the southwest of the island of Eivissa. In order to participate, students were required to speak predominantly Catalan, have normal or corrected-to-normal vision and hearing, and not have language impairments. Additionally, the students' parents should also be Catalan-dominant speakers who are native to the island of Eivissa.

Five female and five male participants took part in the experiment (16-18, mean age 16.8 years, median 17). All of the participants were residents of the locales in the western part of Eivissa explicitly mentioned by Torres Torres (1983:1) as having already undertaken a /ə/-/ɛ/ merger (*Els Cubells, Sant Josep de sa Talaia, Sant Agustí des Vedrà, Sant Antoni de Portmany*).

In addition to require Catalan dominance over Spanish (i.e. the major contact language for Catalan) when sampling participants, the Spanish-Catalan version of the Bilingual Language Profile (BLP) questionnaire (Birdsong & Amengual, 2012) was employed for better assessing participants' bilingualism and language dominance levels (see Appendix B for the full list of questions). The BLP questionnaire contains questions across four topics (language history, language use, language proficiency and language attitudes), which are asked for both the languages under study. For all translations of the questionnaire – e.g. Catalan-Spanish in the case of this study – a version in each language of interest is available to participants. All the participants in the present research filled out the questionnaire in Catalan.

2.2 Stimuli

The target vowels in this study's stimuli set are those of the standard Balearic Catalan inventory: /i, e, ε, ə, a, o, ɔ, u/. Since /ə/ and /ε/ in Western Eivissan have been anecdotally and academically considered to have merged to /ε/ since at least Alcover's (1908) report, the inclusion of the stressed schwa in our stimuli set serves different purposes. Firstly, having a separate label for words that used to have a stressed schwa in Western Eivissan but nowadays do not have one aids in documenting clearly the situation of the hypothesised western merger at the time of writing, i.e. whether the two vowels are still phonetically close or if they have drifted apart again. Secondly, it makes a comparison between Western Eivissan epsilon and (historical) schwa and Eastern Eivissan's merged epsilon and schwa. From here on as regards Western Eivissan, the original epsilon will be represented by /ε/ and the reported /ə/-/ε/ merger will be represented by /ə/ to emphasise the historical origin of the vowel change and to aid in discussing the data analysis.

For each vowel in the standard Balearic Catalan inventory, three words were used for each of the following phonological contexts following Recasens & Espinosa's (2006, 2009) methodology: labial, dentoalveolar, palatal and liquid (/r/, /l/, /r/). Additionally, 12 words were used for the schwa in the pretonic context. The choice of separating the liquids /l/ and /r/ is explained by the authors as being due to both phonemes being specified for 'some predorsum lowering and postdorsum retraction' (Recasens & Espinosa, 2006:649; 2009:245). This secondary velar articulation thus justifies the separation of /l/ and /r/ from the alveolar category. Having no hard data on the velarity of liquids in Eivissan Catalan, the present paper follows the phonotactic context as established by Recasens & Espinosa and adds /r/ to the stimuli set.

When possible, both consonants flanking the vowel aligned with the context (labial, dentoalveolar, palatal, or liquid); otherwise, only the consonant following the vowel did (see Appendix A).

The total number of stimuli thus amounted to 8 vowels x 3 words x 4 contexts + 12 words with pretonic schwa = 108. Each participant realised each target word 3 times. This resulted in 108 items x 3 repetitions x 10 speakers = 3240 items in total. From the 3240 maximum, 1 item's repetition was missing in a participant's recording (palatal *corregeix* 's/he corrects/marks') and 9 other stimuli were excluded due to the lengthening of the stressed or pretonic syllable's onset caused by either hesitation or the production of e.g. *un nus* 'a nose' instead of the word in isolation (dentoalveolar *nus* x 3, /l/-/r/ *cintura* 'waist' x 1, dentoalveolar *sud* 'south' x 1, pretonic *nasset* 'little nose' x 1, /l/-/r/ *sol* 'sun' x 2, dentoalveolar *net* 'grandson' x 1 = 9). Finally, outliers were filtered in R upon data collection (see section 2.5.1 for an

explanation of the filtering technique), excluding 242 outliers. 92% of the original stimuli, 2988,
250 were used in the statistical analyses of the data.

As mentioned in section 1.4 above, the minimal pairs for the vowel duration analysis are not actual minimal pairs, depending on whether the vowels in each member of the pair have merged or not. In this research, the assumption was that the vowels in both vowel pairs had in fact merged, hence the categorisation of these pairs as minimal pairs and specifically
255 consonantal minimal pairs. Because the second research question was developed after data collection took place, the stimuli used in this study were not set up accordingly and only included the single vocalic minimal pair /net-nət/ ‘grandson-clean’ (whose members, in regard to the historical merger, would have been possible homophones and would arguably have left as little room as possible for other phonological influences on the vowels). For this reason, the
260 consonantal minimal pairs shown in (2) below were employed, which could have influenced vowel duration differences by means of the preceding consonant’s voicing.

- (2)
- | | | |
|--------------|-----------|------------------------------|
| a. peu-beu | /pɛw-bəw/ | ‘foot’ - ‘(s)he drinks’ |
| b. mel-pèl | /mɛl-pəl/ | ‘honey’ - ‘hair’ |
| c. bou-pou | /bɔw-pow/ | ‘ox’ - ‘water well’ |
| d. moll-poll | /mɔʎ-poʎ/ | ‘wet’ ⁵ - ‘louse’ |

2.3 Materials

A Zoom H4n Pro device was used to record speech via a RØDE NTG-1 microphone (Phantom +48V option) microphone at 44.1 kHz. The randomised test stimuli were presented to the participants in a slide show on a MacBook Pro (M4).

2.4 Procedure

265 Participants sat in a quiet room together with the interviewer dr. Francesc Torres-Tamarit, facing the laptop monitor and speaking into the microphone. Each participant was shown the slide show and asked to name the object/activity in each slide three times⁶ in Eivissan Catalan. Following the task, each participant filled out the BLP questionnaire.

⁵ As was learnt when collecting data from participant STE-000, *moll* has shifted its meaning to ‘soft’ at least in Eivissan. To elicit the target word *moll*, participants were asked to give another word for the picture targeting *tendre* ‘tender’.

⁶ Due to a lack of pilot experiments, the procedure changed slightly starting from participant STE-002. Participants STE-000 and STE-001 went through a modified version of the original procedure, according to which each participant would be shown the slide show three times in a row, so as to not produce the three repetitions for each word back to back. They were shown the slide show twice, repeating the word twice on the second viewing.

The whole procedure, including the picture-naming task and the questionnaire as well as all the instructions given to the participants, was entirely in Eivissan Catalan and took around 30 minutes.

2.5 Data analysis

2.5.1 Setting up the analyses

Data analysis closely followed the procedure employed by Hamann & Torres-Tamarit (2023). The recordings were manually annotated in the Praat software (Boersma & Weenink, 2025) using TextGrids. Vowel duration as well as the F1 and F2 from the middle 50% of the target vowels were extracted with a Praat script. Vowel duration data was extracted in milliseconds and log-transformed, while the formant data were extracted in Hz and transformed into ERB.

For the filtering of outliers, the same probabilistic mixture modelling as that in Sandoval et al. (2013) was used, where an outlier threshold of 1.5 was adopted, so that all datapoints with fit values more than 1.5 times the interquartile range (IQR) below the first quartile would be flagged as outliers (Stehr, 2018).

Turning to the BLP questionnaire, the scoring system by Birdsong et al. (2012) was followed. Points were assigned to the answers to each question. The answers for each topic were summed and the resulting points were summed across topic for each language. The total score for Spanish dominance was coded as a negative number and the total score for Catalan dominance was coded as a positive number. The difference between the two dominance scores would then show the level of dominance of one language over the other. As explained in section 1.3, Spanish only having high mid vowels could play a role in the production of low mid vowels by Catalan speakers. For this reason, Catalan dominance was sought out in the experiment's participants, over Spanish dominance or even balanced bilingualism.

2.5.2 Vowel quality analysis

To allow for a balanced comparison with the results presented by Hamann & Torres-Tamarit (2023), two separate MANOVAs per speaker were fitted in the R statistical analysis software (R Core Team, 2025) for the /ə/-/ɛ/ and /o/-/ɔ/ vowel pairs, using F1 and F2 values as dependent

Showing the slide show three times was intended to prevent participant from changing their pitch as it happens when listing things in order. While this was the better setup, noticing the participants' fatigue upon viewing the same slide show more than once, we opted for a less ideal procedure which did not hinder participants' attention as much.

variables and the vowel pair ('Vowel' in our data frame) as the independent variable. Pillai scores were used as the output of the MANOVAs, with values ranging from 0 to 1, where 0 signifies close similarity between the two vowels and 1 no similarity. Because no universal cut-off point exists for Pillai scores, the formula for calculating Pillai score thresholds per participant, based on sample size – i.e. number of tokens for each vowel pair per participant – was applied to the filtered data (see Appendix F for a full list of Pillai score thresholds per participant, set against the participants' actual Pillai scores for each vowel pair).

Further, the same criteria used by Hamann & Torres-Tamarit were used: if the analyses of the vowel pairs had *p*-values lower than 0.0025, said vowel pairs would not be considered merged. Each member of the potentially merged pairs, /ə/-/ɛ/ and /o/-/ɔ/, was also paired with one other neighbouring vowel: MANOVAs were thus also run on /e/-/ɛ/, /e/-/ə/, /u/-/o/ and /u/-/ɔ/. Lowering the *p*-value threshold to 0.0025 instead of the generally used 0.05 was the consequence of applying a Bonferroni correction to the standard *p*-value based on the number of participants (10) and the number of possible mergers that were expected to be found (2). This was done to avoid false positives.

As a further sanity check, linear mixed effects models were run on each formant separately with vowel pair as fixed effect, word and speaker as random intercepts. Vowel by speaker as random slope. The R packages used to this end were *lme4* (Bates et al., 2015) and *lmerTest* (Kuznetsova et al., 2017). While MANOVAs are useful when two dependent variables are investigated together, they do not compute random effects like linear mixed effects models do. Running linear mixed effects models on each formant separately also grants a more fine-grained view of any differences between the two vowels of each pair, where the weight of frontness and height on the difference – or similarity, for that matter – can be singled out. In the case of the present research, this is however only partially useful: the lack of difference between the two vowels of each pair along either the F1 or F2 axes would in any case point to a merger, albeit potentially of two different types.

2.5.3 Vowel duration analysis

The duration analysis of both vowel pairs was carried out on the respective consonantal minimal pairs (2 for /ə/-/ɛ/, 2 for /o/-/ɔ/, see (2) above for the full list). First a linear mixed effects model was run on both consonantal minimal pairs per vowel pair. For /ə/-/ɛ/, log-transformed vowel duration was the dependent variable, while vowel pair and the *preceding* consonant's voicing and their interactions were the independent variables; speaker was set as random intercept. For

/o/-/ɔ/, voicing could not be used as an independent variable, as will be explained in the following paragraph.

As mentioned in section 2.2 above, the consonantal minimal pairs differed by the voicing of the consonant *before* the vowels of interest to this research. This is a crucial difference from the vowel quality analysis' variable 'consonantal context', where the context was based on the consonant *following* the vowels of interest. Such a shift in consonantal context is admittedly problematic for generalisations of results across the vowel quality and vowel duration analyses. Still, this recategorization resulted in the analysis of only one consonantal context across all 4 minimal pairs, namely the labial context. By way of this categorisation of all minimal pairs as being of the labial consonantal context, some balancing was possible for the /ə/-/ɛ/ minimal pairs: from (2a) /pɛw-bəw/ to (2b) /mɛl-pəl/, the vowel-consonant voicing pairing is swapped. Unfortunately, this was not possible for the two /o/-/ɔ/ minimal pairs (2c) /bɔw-pow/ and (2d) /mɔʎ-poʎ/, where the consonant preceding /ɔ/ was always voiced and the consonant preceding /o/ was always voiceless.

After this cumulative linear model, individual linear models were run for each minimal pair individually. In these models, only the vowel pair was taken as independent variable, the rest of the model being the same.

3. Results

3.1 Bilingual Language Index

Before moving on to the analyses of the data from the production experiment, the results of the BLP questionnaire will be presented briefly.

All but one participant indicated that Catalan was the only language spoken with their family and that it was spoken within the family for 20+ years (the maximum given in the Likert scale). The language dominance range in our sample spanned from a maximum of 197.98 points for Catalan to -130.756 for Spanish. All participants proved to be Catalan dominant (average language dominance score = 64.202). Figure 1 below visualises the data for each participant. Of note is that despite results by most participants gravitated around the mean, at least two participants showed differing levels of bilingualism and therefore of Catalan dominance – namely participant STE-001 on the higher end of the score range and STE-005 on the lower end of the score range.

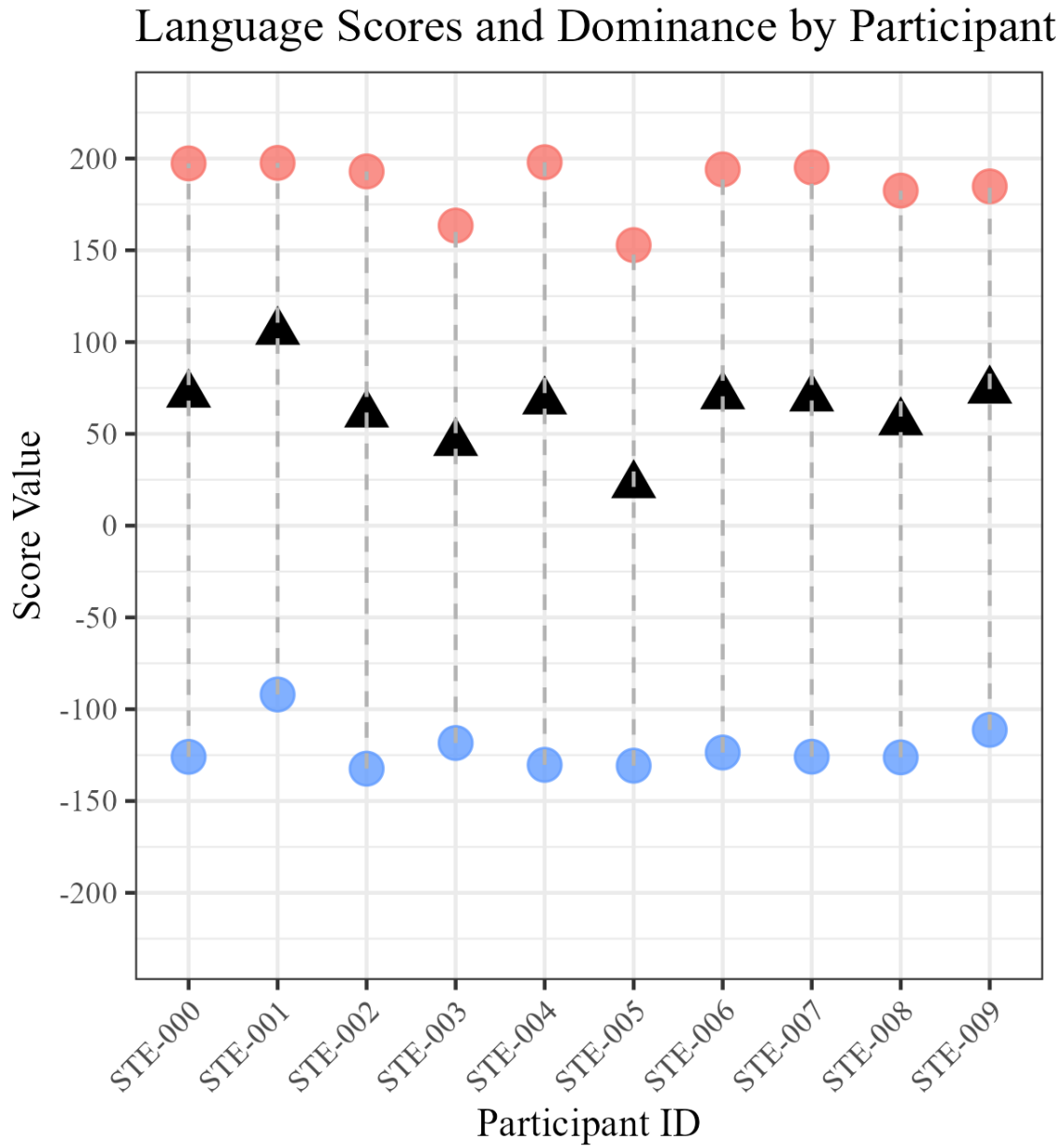


Figure 17. Catalan score (red circles), Spanish score (blue circle) and language dominance (black triangle) per participant.

3.2 Vowel quality results

As visible in Figure 2, the vowel space as averaged across words, repetitions and participants paints an interesting picture.

⁷ For this and all subsequent plots, the R package *plyr* (Wickham, 2011), the R package collection *tidyverse* (Wickham et al., 2019) and the R package *mclust* (Scrucca et al., 2023) were used, with the exception of Figure 7 which was created in Praat.

As for the ‘front’ mid vowels, /ə/ and /ɛ/ are very close to each other, while still not
 355 overlapping completely. /e/, on the other hand, is in the general area of a prototypical high mid
 vowel and not in proximity of any vowel. The pretonic schwa /pə/ is not one of the stressed
 vowels of Western Eivissan, but one type of unstressed schwa (the other being schwa in word-
 and syllable-final position). Still, it too is fronted and seemingly approaching /ə/ and /ɛ/, though
 not to the degree of /pə/ in Hamann & Torres-Tamarit (2023:59).

360 Interestingly, the back mid vowels do not show any degree of merger, with /o/ standing
 not much closer to /ɔ/ than the front mid vowels /e/ and /ɛ/ are to each other.

Another difference between this graph and that in Hamann & Torres-Tamarit (2023) is
 that these vowels occupy an overall slightly lower F1-F2 space than the vowels in their study.

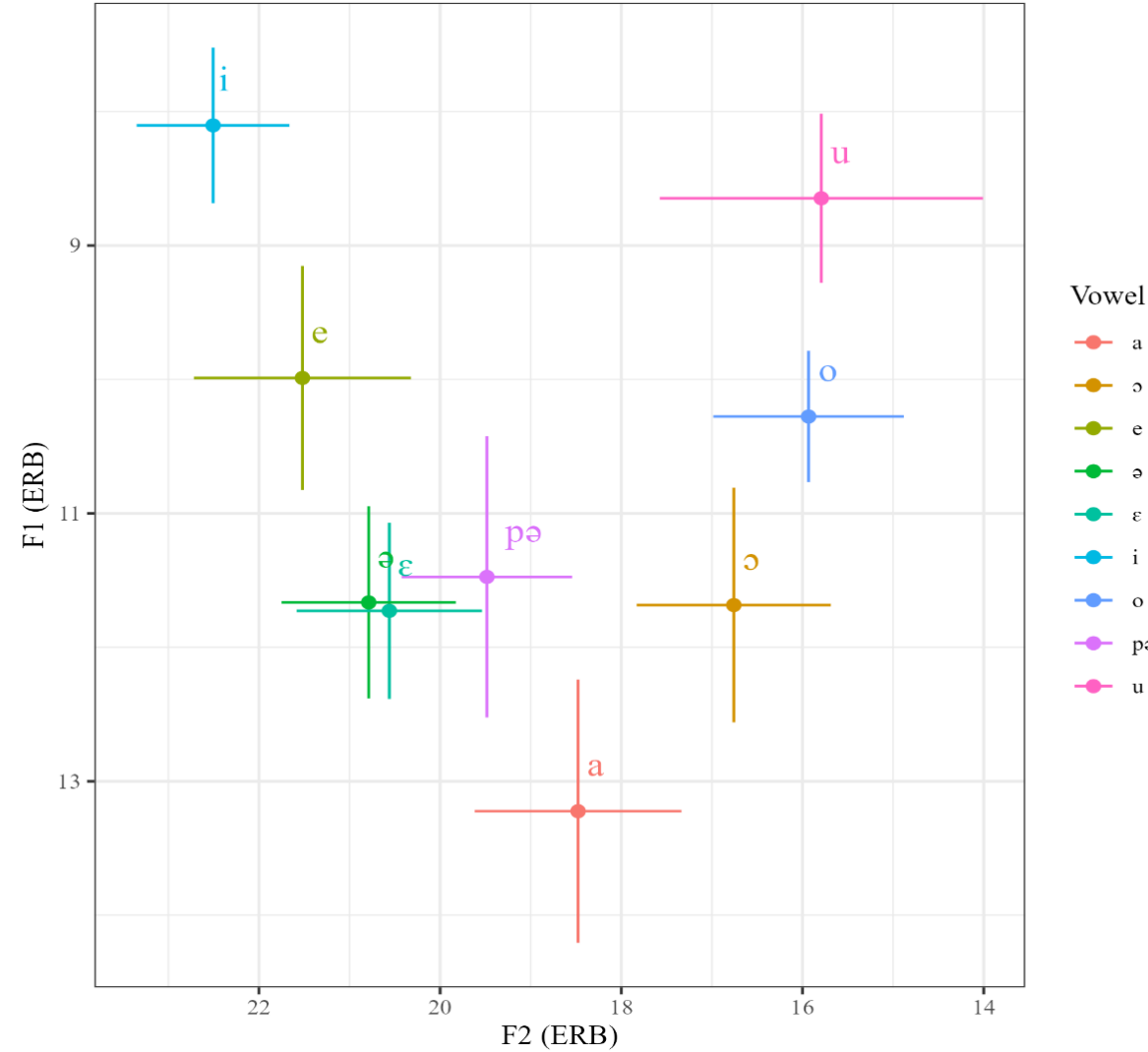


Figure 2. The stressed vowels of Western Eivissan as produced by our participants.

For all the vowel pairs presented in section 2.5, the Pillai following scores agree with
365 the plotted data above.

Figure 3 below shows a stark difference between /ə/ and /ɛ/ as opposed to both /e/ and /ɛ/ and /e/ and /ə/. The mean Pillai score for the vowel pair /ə/-/ɛ/ was 0.138 ($p = 0.04$). On the other hand, the mean Pillai scores of /e/-/ɛ/ and /e/-/ə/ are noticeably higher than those of the first vowel pair, 0.760 ($p < 0.0025$) and 0.725 ($p < 0.0025$) respectively. Still, as per Figure 2
370 above, /ə/ and /ɛ/ are not as close to each other in our data as they are in Hamann & Torres-Tamarit (2023:58-60). In their case, the mean Pillai score for /ə/-/ɛ/ is 0.084 ($p = 0.443$), the mean Pillai score for /e/-/ɛ/ is 0.853 ($p < 0.001^8$) and that for /e/-/ə/ is 0.848 ($p < 0.001$). Curiously, the mean Pillai scores for the /e/-/ɛ/ and /e/-/ə/ in this study are slightly lower than the mean Pillai scores for the same pairs in Hamann & Torres-Tamarit (2023).

Figure 3 speaks to /ə/ and /ɛ/ being quite close to each other. Still, some individual variation occurs in our sample. According to Stanley & Sneller's (2023) Pillai score threshold formula, only the /ə/-/ɛ/ Pillai scores for participants STE-001 and STE-004 are below their respective thresholds. Notably, while STE-001's Pillai score for /ə/-/ɛ/ is only marginally lower than the threshold (Pillai score = 0.081; Pillai threshold = 0.083), STE-004's Pillai score for the
380 same vowel pair is well below the threshold (Pillai score = 0.043; Pillai threshold = 0.081). The Pillai scores by participant for /e/-/ɛ/ and /e/-/ə/ are inevitably always above their respective thresholds.

The results of the linear mixed effects models run on F1 and F2 for /ə/-/ɛ/ both confirm what was found from the MANOVA analysis. In this model, /ɛ/ had a 0.137 times higher F1 (in ERB) than /ə/ in our sample, although the difference was not statistically significant (95% confidence interval = -0.186 .. 0.357 ERB; $p = 0.538$). Conversely, /ɛ/ had a 0.231 times lower F2 (in ERB) in our sample, though again the effect was found to be not statistically significant (95% confidence interval = -0.451 .. -0.011 ERB; $p = 0.047$).

⁸ In Hamann & Torres-Tamarit (2023), a Bonferroni correction of factor 50 is applied, as their study included 25 participants and two mergers. The resulting p-value threshold for their analyses is thus 0.001, while the p-value threshold for this analysis is 0.0025, as mentioned .

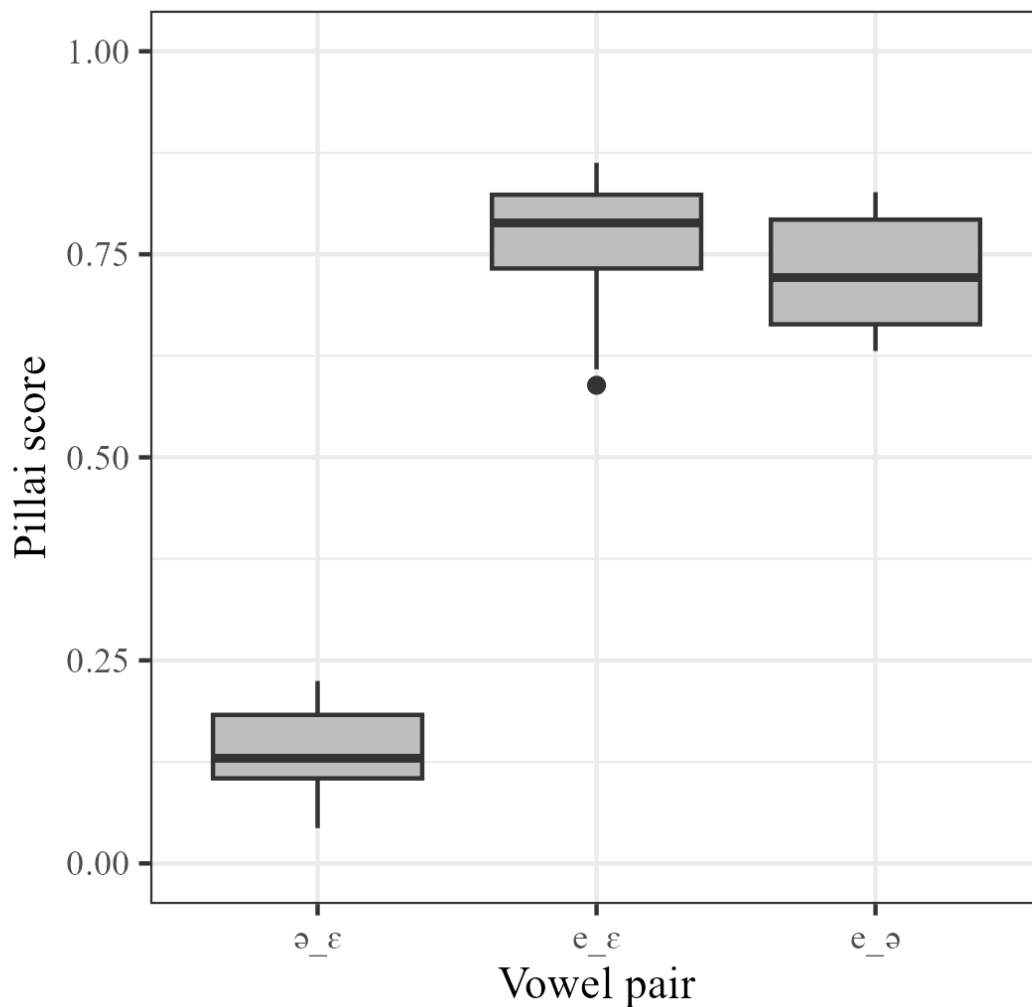


Figure 3. Pillai score boxplots for /ə/-/ε/, as compared to /e/-/ε/ and /e/-/ə/.

Moving to the back mid vowels, Figure 4 below agrees with the general vowel graph above in that /o/ and /ɔ/ at the time of writing are phonetically quite distinct in the speech of our young speakers of Western Eivissan. The mean Pillai score for /o/-/ɔ/ is 0.713 ($p < 0.0025$), which is indicative of a difference in quality between the two vowels. Similarly, the mean Pillai scores for /u/-/ɔ/ and /u/-/o/ are not much higher, resting at 0.899 ($p < 0.0025$) and 0.789 ($p < 0.0025$) respectively. Again, these results differ considerably from those of Hamann & Torres-Tamarit (2023:58-60). Most notably, /o/-/ɔ/ have merged in their Eastern Eivissan population, with a mean Pillai score of 0.095 ($p = 0.442$). /u/-/ɔ/ and /u/-/o/ have much higher mean Pillai score relative to the back mid vowel pair, respectively 0.727 ($p < 0.001$) and 0.771 ($p < 0.001$). As regards /u/-/ɔ/ and /u/-/o/, the present research shows a small difference in mean Pillai score size from the Eastern Eivissan for the former pair but hardly any difference from said Eivissan variety for the latter pair.

395 As said above, /o/ and /ɔ/ do not show overlap in our sample. Inevitably, all of the Pillai scores for this pair are above their respective thresholds. See Appendix F for a full overview.

Again, the results of the linear mixed effects models on F1 and F2 for /o/-/ɔ/ are in accordance with the MANOVA results overall: the two vowels are not merged. /ɔ/ has a 1.405 times higher F1 (in ERB) than /o/ in our sample (95% confidence interval = 1.067 .. 1.743 ERB; $p < 0.0025$). /ɔ/ has a 0.816 times higher F2 (in ERB) than /o/ in our sample, though this effect was found to be not significant (95% confidence interval = 0.121 .. 1.510 ERB; $p = 0.0295$).

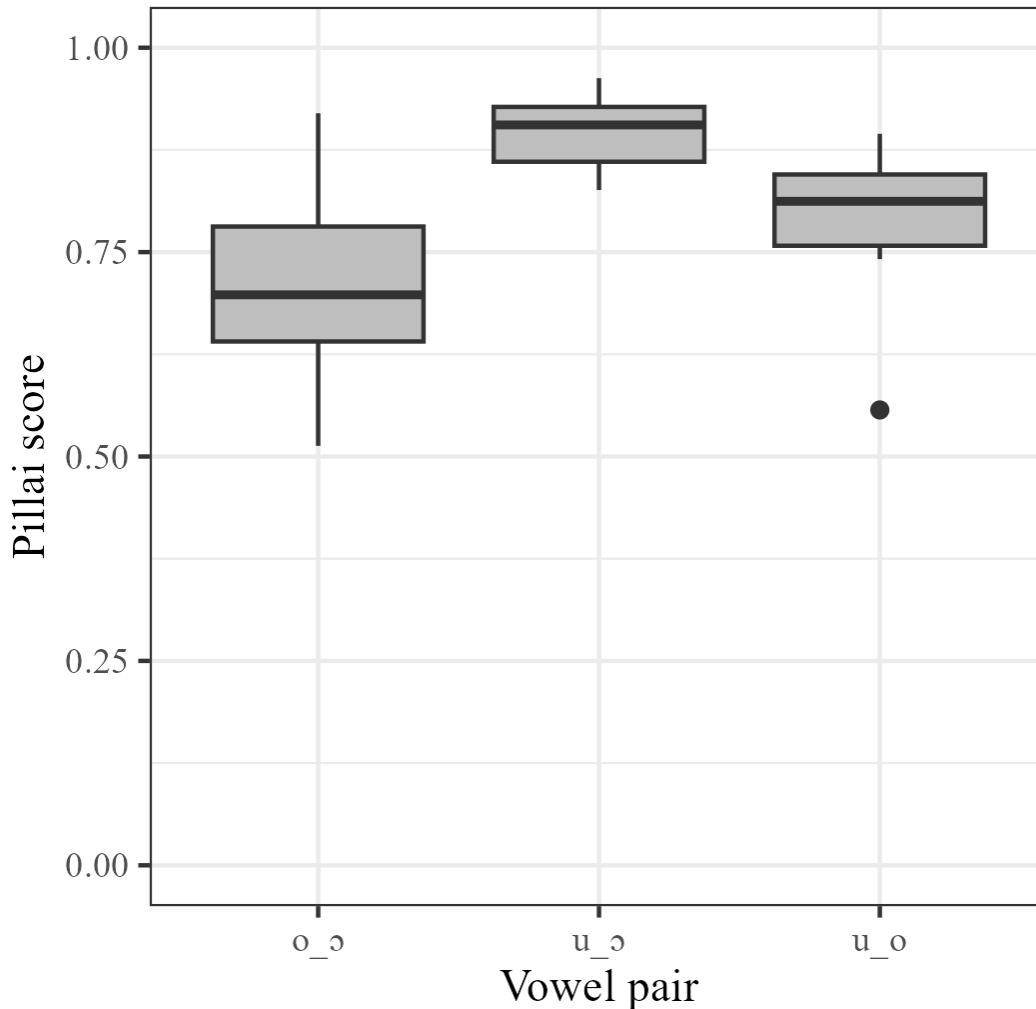


Figure 4. Pillai score boxplots for /o/-/ɔ/, as compared to /u/-/ɔ/ and /u/-/o/.

3.3 Vowel duration results

As mentioned in section 2.2 above, the stimuli were not designed to account for an analysis of vowel duration. For this reason, only 5 consonantal minimal pairs were available for analysis (3 for /ə/-/ɛ/, 2 for /o/-/ɔ/). Balancing of the consonantal context and the consonant's voicing was therefore also not possible. While the following results can and should be questioned on

the basis of generalisability, some significant results were found and could justify further, more structured research on the topic.

Figure 5 below offers a visualisation of the data for the /ə/-/ɛ/ pair. The linear mixed effects model run across both /ə/-/ɛ/ minimal pairs yielded the following results: /ɛ/ was found
410 to be 0.009 log-transformed milliseconds longer than /ə/ in our sample, though this was found to be statistically not significant (95% confidence interval = -0.018 .. 0.037 log-transformed ms; $p = 0.512$). As for the effect of voice, the voiced words were found to be 0.029 log-transformed milliseconds shorter than voiceless words in our sample, though this effect was not significant either according to the corrected p-value threshold (95% confidence interval = -0.057 .. -0.0008
415 log-transformed ms; $p = 0.047$). Unsurprisingly, the interaction between vowel and the preceding consonant's voice was also not significant (estimate: 0.043; 95% confidence interval = -0.013 .. 0.099 log-transformed ms; $p = 0.138$).

The results of the linear mixed effects models for the individual minimal pairs are as follows: in the /bəw-pəw/ minimal pair, /ɛ/ was found to be 0.038 log-transformed milliseconds
420 longer than /ə/ in our sample, although the effect was not statistically significant (95% confidence interval = 0.001 .. 0.075 log-transformed ms; $p = 0.044$). In the /mɛl-pəl/ minimal pair, /ɛ/ was found to be 0.019 log-transformed milliseconds shorter than /ə/ in our sample, although the effect was not statistically significant (95% confidence interval = -0.062 .. 0.024 log-transformed ms; $p = 0.389$).

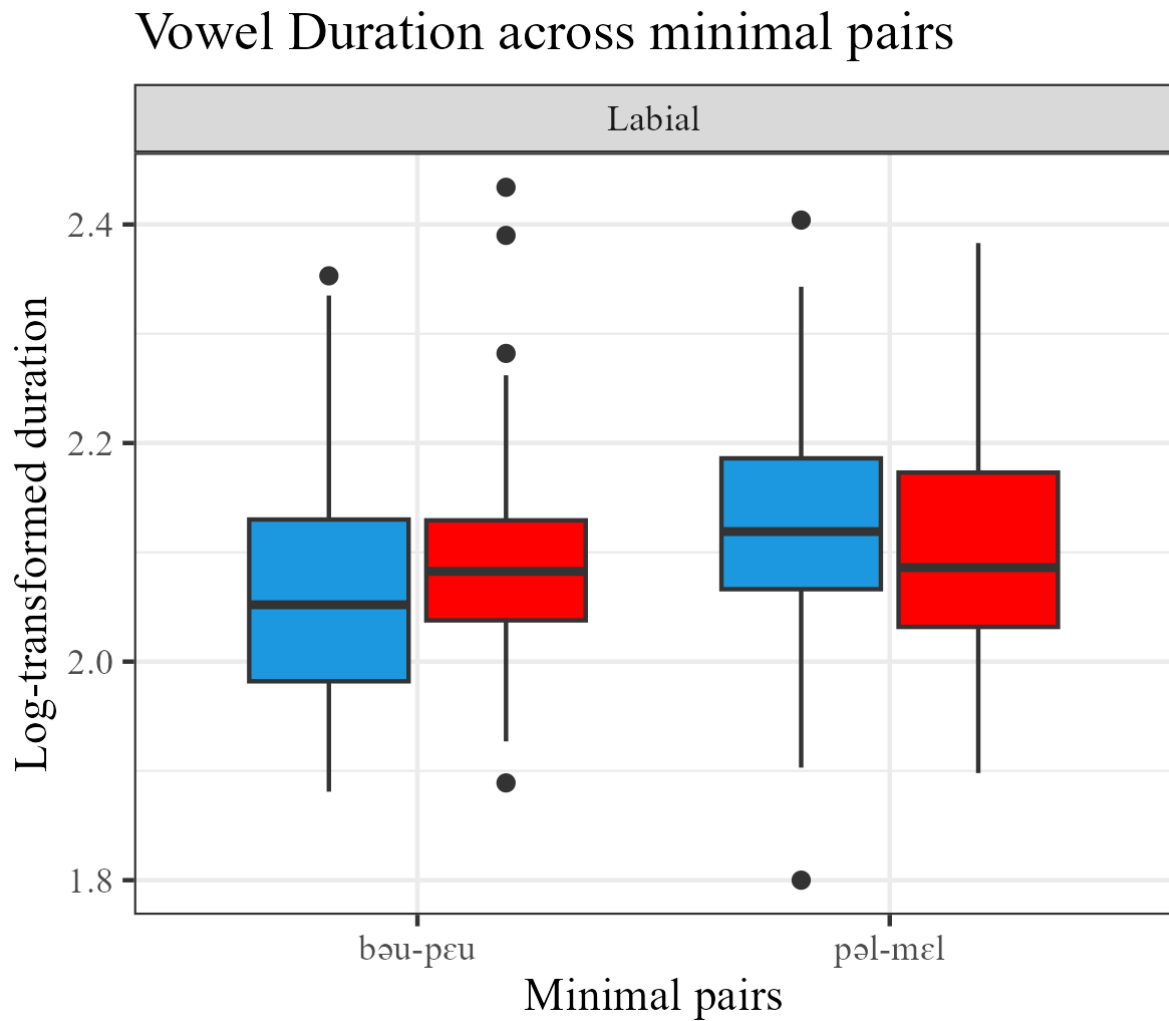


Figure 5. Vowel duration lmer boxplots. /ɔ/ is coded as blue and /ɛ/ as red.

Figure 6 below offers a visualisation of the data for the /o/-/ɔ/ pair. In the model run across both minimal pairs, /ɔ/ was found to be 0.009 log-transformed milliseconds longer than /o/ in our sample, though the effect was not statistically significant (95% confidence interval = -0.028 .. 0.047 log-transformed ms; $p = 0.627$).

Zooming in on the individual minimal pairs for /o/-/ɔ/, the situation seems to change somewhat. In the /bɔw-pow/ pair, /ɔ/ was found to be 0.072 log-transformed milliseconds shorter than /o/ in our sample, though the effect just exceeded the 0.0025 p-value threshold and cannot be considered statistically significant (95% confidence interval = -0.117 .. -0.027 log-transformed ms; $p = 0.0026$). Among all minimal pairs, the only one showing a significant effect of vowel type on duration was /mɔʌ-poʌ/. In this pair, /ɔ/ was found to be 0.056 log-transformed milliseconds longer than /o/ in our sample (95% confidence interval = 0.023 .. 0.089 log-transformed ms; $p < 0.0025$).

Vowel Duration across minimal pairs

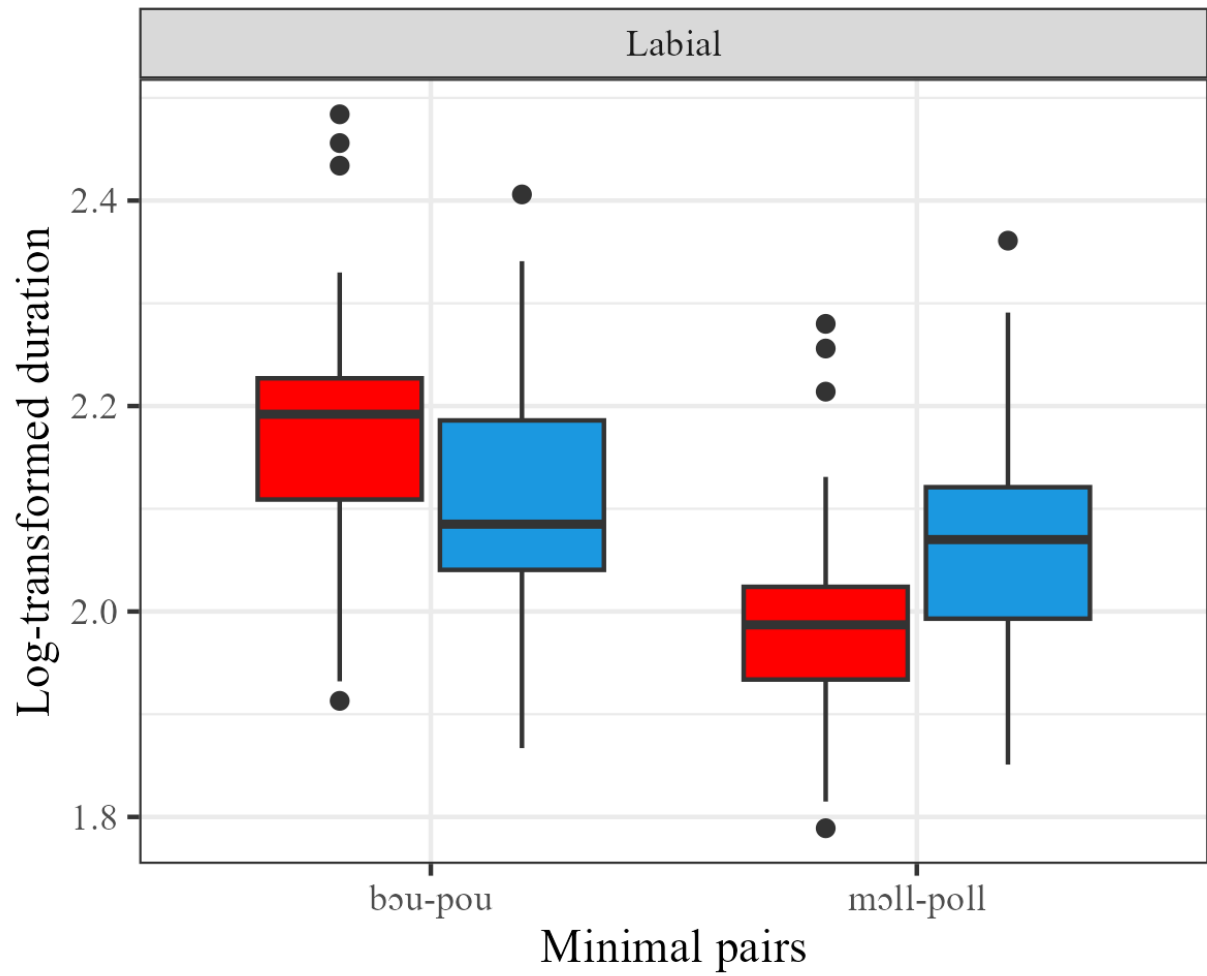


Figure 6. Vowel duration lmer boxplots. /o/ is coded as blue and /ɔ/ as red.

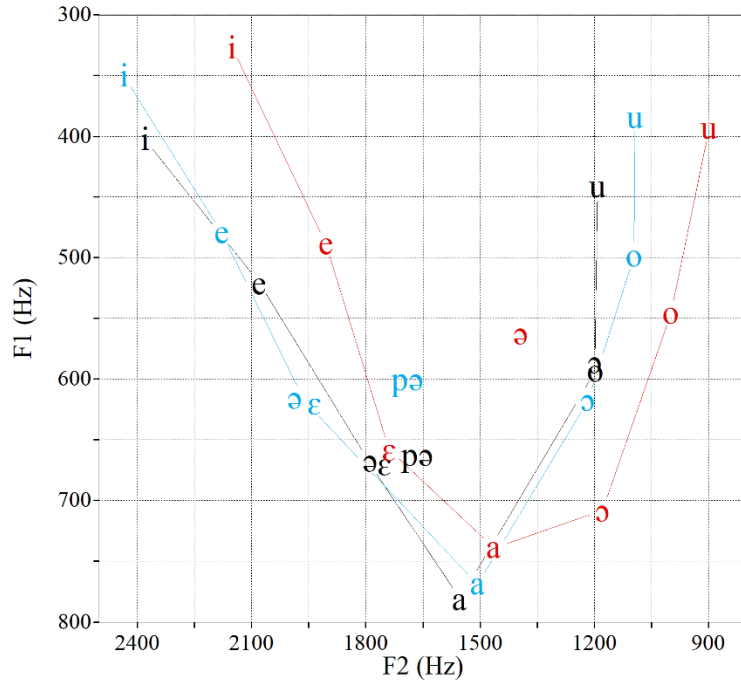


Figure 7. Mean values of stressed vowels of Western Eivissan (present study, cyan), Eastern Eivissan (black) and Majorcan (red). The data for Eastern Eivissan was offered by Hamann & Torres-Tamarit (p.c.) and the data for Majorcan was obtained from Recasens & Espinosa (2006:655). Unlike the rest of the graphs, Figure 7 was made in Praat. The plotting script can be found in Appendix E.

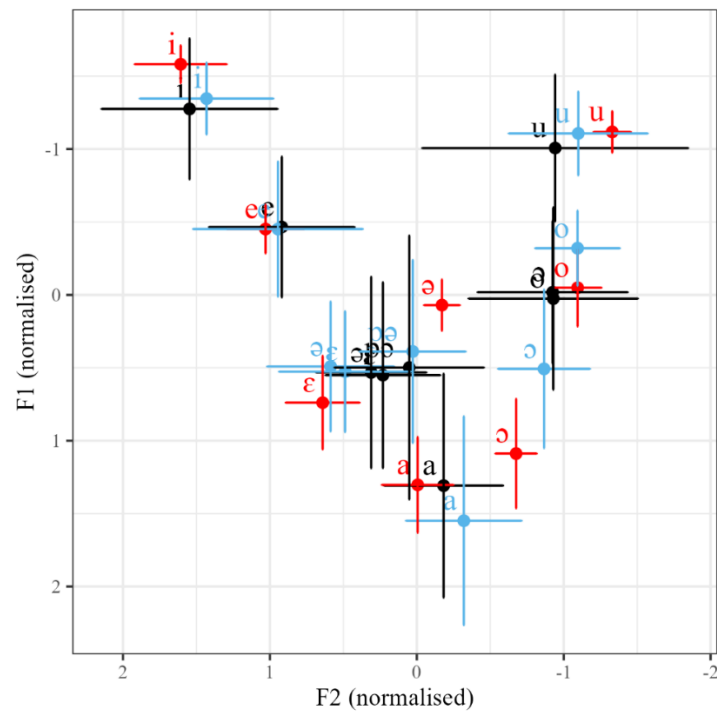


Figure 8. Stressed vowels of Western Eivissan (cyan), Eastern Eivissan (black) and Majorcan (red). The data for Eastern Eivissan was offered by Hamann & Torres-Tamarit (p.c.) and the data for Majorcan was obtained from Recasens & Espinosa (2006:655). Normalised values via z-scores.

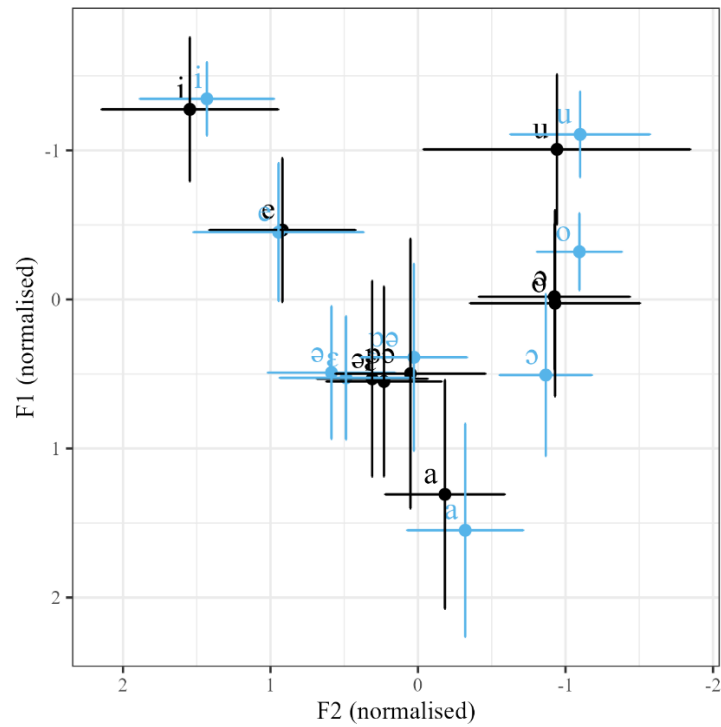


Figure 9. Stressed vowels of Western Eivissan (cyan) and Eastern Eivissan (black). The data for Eastern Eivissan was offered by Hamann & Torres-Tamarit (p.c.). Normalised values via z-scores.

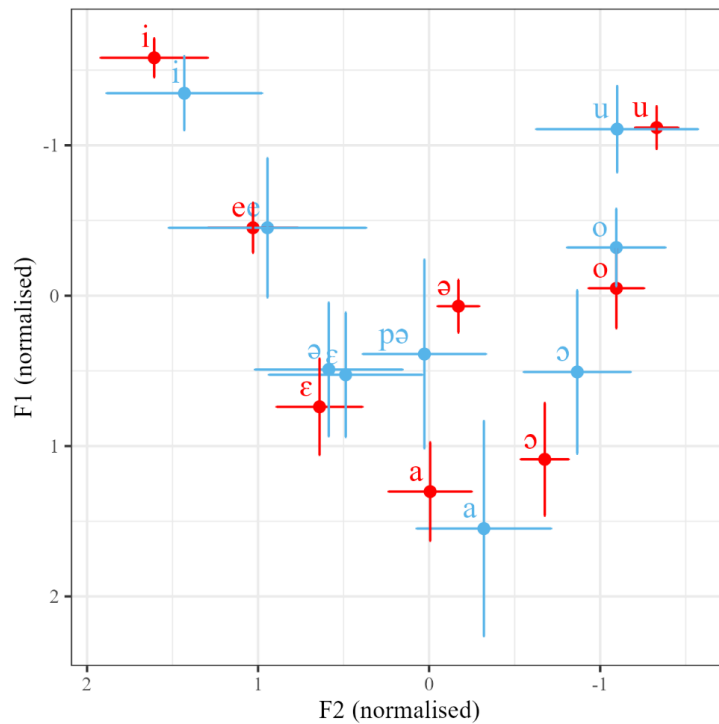


Figure 10. Stressed vowels of Western Eivissan (cyan) and Majorcan (red). The data for Majorcan was obtained from Recasens & Espinosa (2006:655). Normalised values via z-scores.

4. Discussion

As the graphs above show, the western and eastern /ə/-/ɛ/ pairs are very close to each other, though not fully overlapping. After comparing participants' Pillai scores against their respective Pillai score thresholds, the Western Eivissan pair cannot be said to have fully merged. Still, it
440 can be confidently claimed that the two vowels have very nearly merged in the speech of our participants. Moreover, in the western variety, the vowels are slightly more apart from each other than in the eastern one.

Interestingly, Figure 9 most clearly shows a greater degree of dispersion in Eastern Eivissan than in the western variety. Such a difference in consistency across words and speakers
445 might be attributed to the western merger having occurred much earlier than the eastern one and thus having had more time to solidify. While the group in the present paper is surely comparable to that in Hamann & Torres-Tamarit's (2023), the stimuli sets are neither identical nor of the same size – 2988 tokens here v. 4122 in their study. The differences in number and exact type of lexical items might thus play a role in the differences in formant values between
450 the western and the eastern pairs.

At any rate, it is clear that Eivissan does not have a prototypical central schwa, neither in stressed nor in unstressed positions. This sets it apart from the Majorcan data of Recasens & Espinosa (2006) (Figures 8, 10), where the stressed schwa very much remains a central vowel.

As for the back mid vowels, a clear split between Western and Eastern Eivissan is at
455 play. The speech of young Western Eivissan speakers does not point to a merger, as opposed to the eastern variety's youngsters. The linear mixed effects model for F2 run on the back vowel pair found that they were not significantly different along the F2 axis. This, however, does not by itself indicate a merger or overlap between the two vowels. Additionally, the back mid vowels are slightly closer to each other along the F1 axis than their front mid counterparts /e/
460 and merged /E/ are, though it is essential to state that this closeness was not reflected in the statistical models run. If the two back vowels were merging at all, the figures above show this phenomenon would only be in the very early stages.

The lack of vowel duration differences in all but one analysis, while unsurprising, is still interesting as it shows different strategies in vocalic minimal pair discrimination are employed
465 cross-linguistically. It is curious that at least one minimal pair, /məʎ-poʎ/, did show a statistically significant difference in duration between the two vowels. Upon examining the two words in the *Corpus Textual Informatitzat de la Llengua Catalana* (CTILC) via the NIM search engine (Guasch et al., 2013), a noticeable difference in absolute word frequency between *moll*

‘wet’ (1240 entries) and *poll* ‘louse’ (240 entries) emerges. Unfortunately, a number of factors lead the present author to believe this is not the ultimate cause of the distance between /o/ and /ɔ/ found in the /mɔʎ-poʎ/ minimal pair. Firstly and most impactfully, *moll* in Eivissan Catalan means ‘soft’. Secondly, the Parts Of Speech (POS) labels for *moll* in the CTILC corpus include the categories *noun*, *adjective* and *other*. Finally, comparing the members of the other minimal pairs studied here seems to indicate that word frequency alone cannot explain the difference in duration for the vowels of /mɔʎ-poʎ/: in the /bəw-pəw/ pair, *beu* ‘(s)he drinks’ has 827 entries in the CTILC corpus and *peu* ‘foot’ has 8975.

5. Conclusions

This research project sought to document the system of stressed vowels of Western Eivissan proposed by Torres Torres (1983) and first explored by Alcover (1908). To do so, young speakers of Western Eivissan were recorded and their speech was analysed acoustically and statistically.

The results of the present study show that no clear difference between /ə/ and /ɛ/ in voice quality nor duration are present. The vowel quality analysis reflects the findings of Hamann & Torres-Tamarit (2023) on the east of Eivissa. Despite this vowel pair not falling within the ‘merger range’ established per participant by using Stanley & Sneller’s formula. The available data seems to indicate that a full merger is not far. Also, it should be stressed that while a full merger in production is not yet present, this does not exclude the possibility that listeners have two distinct categories for these vowel. It is then left to longitudinal research down the line to document any further changes between /ə/ and /ɛ/ and to perception research to investigate whether young Western Eivissan listeners can distinguish the two sounds and, if so, to what degree.

Contrary to what was found in Eastern Eivissan, our Western Eivissan participants did indeed make a distinction between /o/ and /ɔ/ in their speech. This sets them apart from their eastern peers in interesting ways, as this back vowel distinction reflects a more traditional speech. In turn, a speech perceived as traditional could also bring about a rural v. urban speech framing of not only the varieties of Eivissan but of their speech communities as a whole. Such a crucial difference in the speech of two neighbouring areas of the island of Eivissa will hopefully inspire further sociolinguistic work to document how these different systems interact with each other, if at all.

As is often the case, this project is not without its faults. A possible issue could have
500 been the speech of the researcher collecting data not matching the participants'. Depending on
what language attitudes exist within Eivissa towards the two varieties, this might have affected
the participants' production of Western Eivissan vowels. Nevertheless, at least one of the
mergers previously found in Eastern Eivissan was not present in the speech of the Western
Eivissan participants. Until further research produces more production data from a similar
505 population sample, it remains an open question to what extent the role of the interviewer's
speech influenced our participants.

Secondly, the BLP questionnaire contained a question about the participants' gender.
However, biological sex was the relevant information needed when extracting the data from the
recordings. Inevitably, sex and gender were conflated, as this mistake was noticed only well
510 after data collection had been over. Such a conflation must be avoided in further research, as
acoustic data will be directly impacted by it.

Finally, the vowel duration analysis did not rest on solid grounds. The unbalanced nature
of the stimuli set produced too many variables that could have impacted the results. Further
research should remedy this by including more minimal pairs in the stimuli.

515 Hopefully this study offered enough evidence for the existence of at least two distinct
vowel systems in Eivissa, such that more phonetic research is carried out on Eivissan. The
present data, in conjunction with that from Hamann & Torres-Tamarit (2023) also provides
more data for future sound change research in Romance, especially in regard to the changes of
schwa.

Appendices

Appendix A: Stimuli set

	Labial	Dentoalveolar
/i/	<i>fibra</i> fiber <i>pipa</i> pipe <i>vibra</i> vibrates (the telephone)	<i>disc</i> disc (album) <i>dit</i> finger <i>nit</i> night
/e/	<i>febre</i> fever <i>jueva</i> Jewish (feminine) <i>novembre</i> November	<i>cent</i> one hundred <i>dent</i> tooth <i>nét</i> grandson
/ɛ/	<i>europ<u>e</u></i> European <i>trofeu</i> trophy <i>peu</i> foot	<i>cendra</i> ash <i>set</i> seven <i>tendre</i> tender
/ɛ/ (/ə/)	<i>beu</i> (s)he drinks <i>pebre</i> pepper <i>veu</i> voice	<i>ant<u>e</u>na</i> antenna <i>sed</i> thirst <i>net</i> clean
/a/	<i>bava</i> drool <i>Papa</i> Pope <i>pau</i> peace	<i>nas</i> nose <i>sant</i> saint <i>tassa</i> cup
/ɔ/	<i>bou</i> ox <i>moble</i> furniture <i>poble</i> village	<i>dona</i> woman <i>soci</i> business partner <i>son (tenir)</i> sleep (noun) or to be sleepy
/o/	<i>bomba</i> bomb <i>poma</i> appel <i>pou</i> well (water source)	<i>dos</i> two <i>sostre</i> ceiling <i>tos</i> cough
/u/	<i>fum</i> smoke <i>puma</i> puma <i>república</i> republic	<i>duna</i> dune <i>nus</i> knot <i>sud</i> south

Appendix A continued (1)

520 * = only one of the flanking consonant has the intended place of articulation, usually the post-vocalic one

** = taps are also included (Recasens & Espinosa had only one trill and alveolar laterals)

	Palatal	/l, r/**
/i/	<i>gin</i> gin <i>desig*</i> desire <i>fitxa*</i> tile (a domino tile)	<i>carril</i> lane/track <i>goril-la**</i> gorilla <i>lila</i> purple (color)
/e/	<i>lleig</i> ugly <i>llenya</i> firewood <i>canyella</i> cinnamon	<i>carrera**</i> race/career <i>cirera**</i> cherry <i>cremallera* **</i> zipper
/ɛ/	<i>escabetx</i> pickle (pickled mussel) <i>txec</i> Czech <i>xerra</i> (s)he chats	<i>arrel</i> root <i>ferro*</i> iron <i>mel*</i> honey

<i>/ɛ/</i> <i>(/ə/)</i>	<i>afegeix</i> (s)he adds <i>corregeix</i> (s)he corrects <i>llegeix</i> (s)he reads	<i>pèl*</i> hair <i>tela*</i> fabric <i>vel*</i> veil
<i>/a/</i>	<i>fitxatge</i> signing (in sports, transfer) <i>maquillatge</i> makeup <i>llavi</i> lip	<i>bala*</i> bullet <i>pala*</i> shovel/paddle <i>sala*</i> room/hall
<i>/ɔ/</i>	<i>coll*</i> neck <i>moll*</i> wet <i>rellotge</i> clock/watch	<i>dol* (anar de)</i> mourning (mourning dressing) <i>fillol</i> godson <i>sol*</i> sun
<i>/o/</i>	<i>coix*</i> lame <i>cotxe*</i> car <i>poll*</i> louse	<i>dolç*</i> sweet <i>morro*</i> snout (a pig's snout) <i>torre*</i> tower
<i>/u/</i>	<i>juny</i> June <i>jutge</i> judge <i>lluny</i> far	<i>cintura* **</i> waist <i>natura* **</i> nature <i>pintura* **</i> painting

Appendix A continued (2)

pretonic /ə/	
<i>nasset</i>	small nose
<i>santet</i>	little saint
<i>tasseta</i>	small cup
<i>anteneta</i>	small antenna
<i>pelet</i>	little hair
<i>carril</i>	lane/track
<i>saleta</i>	small room/hall
<i>denteta</i>	small tooth
<i>netet</i>	young grandson
<i>cirereta</i>	small cherry
<i>canyella</i>	cinnamon
<i>cremallera</i>	zipper

Appendix B: Bilingual Language Questionnaire and scores

Table B1: Catalan-Spanish BLP template (in Catalan)

Participant # _____

Bilingual Language Profile: Català- Espanyol

Ens agradaria demanar la seva ajuda per contestar les preguntes següents sobre el seu historial lingüístic: ús, actituds i competència lingüística. Aquesta enquesta conté dinou preguntes i la durada és d'uns deu minuts. Això no és una prova, per tant no hi ha respostes correctes ni incorrectes. Per favor contesti cada pregunta i respongui amb sinceritat, ja que només així es garanteix l'èxit d'aquesta investigació. Moltes gràcies per la seva atenció.

I. Informació biogràfica

Nom _____ Data d'avui ____/____/____

Edat ____ ☐ Home/ ☐ Dona/ ☐ Altre Lloc de residència actual: ciutat/estat _____ País _____

Nivell més alt de formació acadèmica:	☐ Menys que l'escola secundària	☐ Escola Secundària
	☐ Un poc d' universitat	☐ Universitat (diplomatura, llicenciatura.)
	☐ Un poc d'escola graduada	☐ Màster
	☐ Doctorat	☐ Altres: _____

Please cite as :

Birdsong, D., Gertken, L.M., & Amengual, M. *Bilingual Language Profile: An Easy-to-Use Instrument to Assess Bilingualism*. COERLL, University of Texas at Austin. Web. 20 Jan. 2012. <<https://sites.la.utexas.edu/bilingual/>>.

II. Historial lingüístic

En aquesta secció, ens agradaria que contestés algunes preguntes sobre el seu historial lingüístic marcant la casella adequada.

1. A quina edat va començar a **aprendre les llengües** següents?

Català

☐ Des del Naixement ☐ 1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐ 6 ☐ 7 ☐ 8 ☐ 9 ☐ 10 ☐ 11 ☐ 12 ☐ 13 ☐ 14 ☐ 15 ☐ 16 ☐ 17 ☐ 18 ☐ 19 ☐ 20+

Espanyol

☐ Des del Naixement ☐ 1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐ 6 ☐ 7 ☐ 8 ☐ 9 ☐ 10 ☐ 11 ☐ 12 ☐ 13 ☐ 14 ☐ 15 ☐ 16 ☐ 17 ☐ 18 ☐ 19 ☐ 20+

2. A quina edat va començar a **sentir-se còmode** emprant les llengües següents?

Català

☐ Tan aviat com record ☐ 1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐ 6 ☐ 7 ☐ 8 ☐ 9 ☐ 10 ☐ 11 ☐ 12 ☐ 13 ☐ 14 ☐ 15 ☐ 16 ☐ 17 ☐ 18 ☐ 19 ☐ 20+ ☐ encara no

Espanyol

☐ Tan aviat com record ☐ 1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐ 6 ☐ 7 ☐ 8 ☐ 9 ☐ 10 ☐ 11 ☐ 12 ☐ 13 ☐ 14 ☐ 15 ☐ 16 ☐ 17 ☐ 18 ☐ 19 ☐ 20+ ☐ encara no

3. Quants **anys de classes (gramàtica, història, matemàtiques, etc.)** ha tingut en les llengües següents (des de l'escola de primària fins a la universitat).

Català

☐ 0 ☐ 1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐ 6 ☐ 7 ☐ 8 ☐ 9 ☐ 10 ☐ 11 ☐ 12 ☐ 13 ☐ 14 ☐ 15 ☐ 16 ☐ 17 ☐ 18 ☐ 19 ☐ 20+

Espanyol

☐ 0 ☐ 1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐ 6 ☐ 7 ☐ 8 ☐ 9 ☐ 10 ☐ 11 ☐ 12 ☐ 13 ☐ 14 ☐ 15 ☐ 16 ☐ 17 ☐ 18 ☐ 19 ☐ 20+

4. Quants anys **ha passat en un país/ regió** on es parlen les llengües següents?

Català

☐ 0 ☐ 1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐ 6 ☐ 7 ☐ 8 ☐ 9 ☐ 10 ☐ 11 ☐ 12 ☐ 13 ☐ 14 ☐ 15 ☐ 16 ☐ 17 ☐ 18 ☐ 19 ☐ 20+

Espanyol

☐ 0 ☐ 1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐ 6 ☐ 7 ☐ 8 ☐ 9 ☐ 10 ☐ 11 ☐ 12 ☐ 13 ☐ 14 ☐ 15 ☐ 16 ☐ 17 ☐ 18 ☐ 19 ☐ 20+

5. Quants anys s'han parlat les llengües següents dins la **família**?

Català

☐ 0 ☐ 1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐ 6 ☐ 7 ☐ 8 ☐ 9 ☐ 10 ☐ 11 ☐ 12 ☐ 13 ☐ 14 ☐ 15 ☐ 16 ☐ 17 ☐ 18 ☐ 19 ☐ 20+

Espanyol

☐ 0 ☐ 1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐ 6 ☐ 7 ☐ 8 ☐ 9 ☐ 10 ☐ 11 ☐ 12 ☐ 13 ☐ 14 ☐ 15 ☐ 16 ☐ 17 ☐ 18 ☐ 19 ☐ 20+

6. Quants anys ha passat en un **ambient de treball** on es parlen les llengües següents?

Català

☐ 0 ☐ 1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐ 6 ☐ 7 ☐ 8 ☐ 9 ☐ 10 ☐ 11 ☐ 12 ☐ 13 ☐ 14 ☐ 15 ☐ 16 ☐ 17 ☐ 18 ☐ 19 ☐ 20+

Espanyol

☐ 0 ☐ 1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐ 6 ☐ 7 ☐ 8 ☐ 9 ☐ 10 ☐ 11 ☐ 12 ☐ 13 ☐ 14 ☐ 15 ☐ 16 ☐ 17 ☐ 18 ☐ 19 ☐ 20+

III. Ús de les llengües

En aquesta secció, ens agradaria que contestés algunes preguntes sobre el seu ús de les llengües que s'anomenen a continuació. Marqui la casella adequada. L'ús total de totes les llengües ha d'arribar al 100% a cada pregunta.

7. En una setmana normal, quin percentatge de temps empra les llengües següents amb els **seus amics**?

Català	☐	☐	☐	☐	☐	☐	☐	☐	☐	☐
	0%	10%	20%	30%	40%	50%	60%	70%	80%	90%
Espanyol	☐	☐	☐	☐	☐	☐	☐	☐	☐	☐
	0%	10%	20%	30%	40%	50%	60%	70%	80%	90%
Altres llengües	☐	☐	☐	☐	☐	☐	☐	☐	☐	☐
	0%	10%	20%	30%	40%	50%	60%	70%	80%	90%

8. En una setmana normal, quin percentatge de temps usa les llengües següents amb **la seva família**?

Català	☐	☐	☐	☐	☐	☐	☐	☐	☐	☐
	0%	10%	20%	30%	40%	50%	60%	70%	80%	90%
Espanyol	☐	☐	☐	☐	☐	☐	☐	☐	☐	☐
	0%	10%	20%	30%	40%	50%	60%	70%	80%	90%
Altres llengües	☐	☐	☐	☐	☐	☐	☐	☐	☐	☐
	0%	10%	20%	30%	40%	50%	60%	70%	80%	90%

9. En una setmana normal, quin percentatge del temps usa les llengües següents a **l'escola/la feina**?

Català	☐	☐	☐	☐	☐	☐	☐	☐	☐	☐
	0%	10%	20%	30%	40%	50%	60%	70%	80%	90%
Espanyol	☐	☐	☐	☐	☐	☐	☐	☐	☐	☐
	0%	10%	20%	30%	40%	50%	60%	70%	80%	90%
Altres llengües	☐	☐	☐	☐	☐	☐	☐	☐	☐	☐
	0%	10%	20%	30%	40%	50%	60%	70%	80%	90%

10. Quan es parla **a si mateix**, amb quina freqüència **es parla** amb les llengües següents?

Català	☐	☐	☐	☐	☐	☐	☐	☐	☐	☐
	0%	10%	20%	30%	40%	50%	60%	70%	80%	90%
Espanyol	☐	☐	☐	☐	☐	☐	☐	☐	☐	☐
	0%	10%	20%	30%	40%	50%	60%	70%	80%	90%
Altres llengües	☐	☐	☐	☐	☐	☐	☐	☐	☐	☐
	0%	10%	20%	30%	40%	50%	60%	70%	80%	90%

11. Quan està fent **càlculs**, amb quina freqüència **compta** amb les llengües següents?

Català	☐	☐	☐	☐	☐	☐	☐	☐	☐	☐
	0%	10%	20%	30%	40%	50%	60%	70%	80%	90%
Espanyol	☐	☐	☐	☐	☐	☐	☐	☐	☐	☐
	0%	10%	20%	30%	40%	50%	60%	70%	80%	90%
Altres llengües	☐	☐	☐	☐	☐	☐	☐	☐	☐	☐
	0%	10%	20%	30%	40%	50%	60%	70%	80%	90%

En aquesta secció, ens agradaria que considerés la seva competència lingüística marcant la casella de 0 a 6..

6=molt bé

- En aquesta secció, ens agradaria que contestés les següents afirmacions sobre les actituds lingüístiques. Marqui les caselles de 0 a 6.

6= estic d'acord

- 4

Table B2: BLP scores

				Catalan				Spanish						
Speaker	Age	Gender	Place	Language History	Language Use	Language Proficiency	Language Attitudes	Language History	Language Use	Language Proficiency	Language Attitudes	Catalan Global Score	Spanish Global Score	Dominance Score
STE-000	17	Home	Els Cubells	44.946	47.96	54.48	49.94	33.142	6.54	54.48	31.78	197.326	125.942	71.384
STE-001	18	Home	Els Cubells	44.038	46.87	52.21	54.48	33.142	4.36	49.94	54.48	197.598	91.982	105.616
STE-002	17	Home	Sant Josep	44.038	49.05	49.94	49.94	34.958	4.36	49.94	43.13	192.968	132.388	60.58
STE-003	17	Home	Sant Josep	44.946	43.6	45.4	29.51	33.596	7.63	47.67	29.51	163.456	118.406	45.05
STE-004	17	Dona	Sant Agustí	44.946	49.05	47.67	52.21	38.59	5.45	49.94	36.32	197.98	130.3	67.68
STE-005	16	Dona	Sant Antoni	42.222	42.51	43.13	24.97	40.406	10.9	49.94	29.51	152.832	130.756	22.076
STE-006	17	Dona	Els Cubells	44.946	44.69	49.94	54.48	26.332	8.72	49.94	38.59	194.056	123.582	70.474
STE-007	17	Dona	Els Cubells	44.946	45.78	54.48	49.94	33.142	8.72	54.48	29.51	195.146	125.852	69.294
STE-008	16	Home	Els Cubells	43.584	43.6	47.67	47.67	33.596	10.9	47.67	34.05	182.524	126.216	56.308
STE-009	16	Dona	Sant Josep	43.584	43.6	45.4	52.21	29.964	10.9	36.32	34.05	184.794	111.234	73.56

Appendix C: Praat script for extracting data

```
# this script is an improved version by dr Silke Hamann of a previous script
# I'd coded
if numberOfSelected ("Sound") <> 1 or numberOfSelected ("TextGrid") <> 1
  exit "Please select a Sound and a TextGrid first."
endif
textGrid = selected ("TextGrid")
sound = selected ("Sound")

form speakerinfo
  choice Gender 1
    button female
    button male
endform

if gender$ = "female"
  maximum = 5500
else
  maximum = 5000
endif

table = Create Table with column names: "table", 0, { "Speaker", "Gender", "Vowel", "Word",
"Context",
```

```
... "beginVowel", "Duration_ms", "Duration_log", "AverageF1_Hz", "AverageF2_Hz",
"F1_ERB", "F2_ERB" }
```

```
selectObject: sound
fileName$ = selected$ ("Sound")
```

```
selectObject: textGrid
numberOfIntervals = Get number of intervals: 2
```

```
for interval from 1 to numberOfIntervals
  label$ = Get label of interval: 2, interval
  if label$ <> ""
    starttime = Get start time of interval: 2, interval
    endtime = Get end time of interval: 2, interval
    duration_ms = (endtime - starttime)*1000 ; ms
    condition$ = Get label of interval: 3, interval
```

```
#take starttime of vowel to access interval on wordtier and get label of interval there
```

```
interval2 = Get interval at time... 1 starttime
word$ = Get label of interval... 1 interval2
```

```
#create formant object for 25%-75% duration of vowel (this way we know for sure that the
formant values outside that time range are not influencing our measure)
```

```
selectObject: sound
t25 = starttime + duration_ms/4000
t75 = endtime - duration_ms/4000
vowel = Extract part: t25, t75, "Rectangular", 1.0, "no"
my.Formant = To Formant (burg): 0.001, 5, maximum, 0.025, 50 ;if too many
undefines then change window length from 0.025 to shorter
f1 = Get quantile: 1, 0, 0, "Hertz", 0.50
f2 = Get quantile: 2, 0, 0, "Hertz", 0.50
f1ERB = hertzToErb (f1)
f2ERB = hertzToErb (f2)
removeObject: my.Formant, vowel
```

```
#write results into the table that was created at the beginning
```

```
selectObject: table
Append row
rowNumber = Get number of rows
Set string value: rowNumber, "Speaker", fileName$
Set string value: rowNumber, "Gender", gender$
Set string value: rowNumber, "Vowel", label$
Set string value: rowNumber, "Word", word$
Set string value: rowNumber, "Context", condition$
Set string value: rowNumber, "beginVowel", fixed$ (starttime, 6)
Set string value: rowNumber, "Duration_ms", fixed$ (duration_ms, 3)
Set string value: rowNumber, "Duration_log", fixed$ (log10(duration_ms), 3)
Set string value: rowNumber, "AverageF1_Hz", fixed$ (f1, 3)
```

```

        Set string value: rowNumber, "AverageF2_Hz", fixed$ (f2, 3)
        Set string value: rowNumber, "F1_ERB", fixed$ (f1ERB, 3)
        Set string value: rowNumber, "F2_ERB", fixed$ (f2ERB, 3)

        selectObject: textGrid
    endif
endfor

selectObject: table
Save as tab-separated file: "results.Table"
View & Edit
selectObject: sound, textGrid

```

Appendix D: R script for plotting and data analysis

```

# this long script is an amalgamation made by me of previous scripts used for managing and
# analysing data made by dr Silke Hamann and dr Francesc Torres-Tamarit.
# Some plotting scripts were devised or modified with the aid of generative artificial
# intelligence. The vowel duration analysis scripts were a novel addition by me.

```

```

---
title: "alesRmaThesisBigStatsFile"
author: "Alessandro Pecoraro 12523615"
date: "2025-05-30"
output: html_document
---

```

```

# General setup

```

```

Loading libraries:
``{r}
library("plyr")
library("tidyverse")
library("mclust")
library("lme4")
library("lmerTest")
``

```

```

# Bilingual Language Profile analysis

```

```

Plotting the BLP questionnaire results:
``{r}
# Create the data frame
language_data <- data.frame(
  Participant = c("STE-000", "STE-001", "STE-002", "STE-003", "STE-004",
    "STE-005", "STE-006", "STE-007", "STE-008", "STE-009"),
  Language_Dominance = c(71.384, 105.616, 60.58, 45.05, 67.68,

```

```

        22.076, 70.474, 69.294, 56.308, 73.56),
Spanish_Score = c(-125.942, -91.982, -132.388, -118.406, -130.3,
                  -130.756, -123.582, -125.852, -126.216, -111.234),
Catalan_Score = c(197.326, 197.598, 192.968, 163.456, 197.98,
                  152.832, 194.056, 195.146, 182.524, 184.794)
)

# Load libraries
library(ggplot2)
library(tidyr)
library(dplyr)

# Create the plot with points - NO LONG DATA TRANSFORMATION
ggplot(language_data, aes(x = Participant)) +
  # Spanish scores as points
  geom_point(aes(y = Spanish_Score, color = "Spanish Score"),
             size = 4, alpha = 0.8) +
  # Catalan scores as points
  geom_point(aes(y = Catalan_Score, color = "Catalan Score"),
             size = 4, alpha = 0.8) +
  # Language dominance as points
  geom_point(aes(y = Language_Dominance, color = "Language Dominance"),
             size = 4, shape = 17) +
  # Connecting lines (using original data)
  geom_segment(aes(x = Participant, xend = Participant,
                  y = Spanish_Score, yend = Catalan_Score),
              color = "gray70", linetype = "dashed") +
  # Custom colors
  scale_color_manual(values = c(
    "Catalan Score" = "#F8766D", # Orange-red
    "Spanish Score" = "#619CFF", # Teal
    "Language Dominance" = "black" # Blue
  )) +
  scale_y_continuous(limits = c(-225, 225),
                    breaks = seq(-250, 250, by = 50)
  ) +
  # Labels and theme
  labs(title = "Language Scores and Dominance by Participant",
       y = "Score Value",
       x = "Participant ID",
       color = "Measure") +
  theme_minimal() +
  theme_bw(base_family = "serif") +
  theme(axis.text.x = element_text(angle = 45, hjust = 1),
        legend.position = "none", aspect.ratio = 1)

ggsave("Figures/BLP_results.png", width = NA, height = NA, dpi = 300, limitsize = TRUE)
```

```

Reading the experiment data:

```

```{r}
data = read.delim ("000-009.Table", stringsAsFactors = TRUE)
```

```

Dataset snippet:

```

```{r}
head(data)
```

```

Adding a column "outlier" with value 0:

```

```{r}
dim(data)
data['outlier'] <- 0
summary(data) # CANYELLA, CARRIL AND CREMALLERA COUNT AS 60 IN THE
WORD CATEGORY BECAUSE THEY ARE BOTH USED FOR PRETONIC AND ONE
OTHER CONTEXT
```

```

Detecting outliers:

```

```{r}

# I tried different ways to hide the output density models but I have not been able to find any
combination of code to get rid of them from the summary. Apologies :(
outlierThreshold=1.5
for(indiv in levels(data$Speaker)){
  for(Vowel in levels(data$Vowel)){
    select=data$Vowel==Vowel&data$Speaker==indiv
    dataF=data[select,c("F1_ERB","F2_ERB")]
    gmmfit <- densityMclust(dataF, 1)
    log_density <- log(gmmfit$density)
    lowerq <- quantile(log_density)[2]
    iqr <- IQR(log_density)
    threshold.lower <- lowerq - (outlierThreshold * iqr)
    data[select,"outlier"]=ifelse(log_density < threshold.lower, 1, 0)
  }
}
data$outlier <- factor(data$outlier, levels = c(0, 1), labels = c("valid", "outlier"))
write.table(data, file = paste("filtered_", "000-009.Table", sep=""), sep = "\t", row.names =
FALSE, fileEncoding="UTF-8")
```

```

Plotting the outlier flagging results for visual check:

```

```{r}
outlier_plot <- ggplot(data, aes(x = F2_ERB, y = F1_ERB, shape = outlier, color = Vowel,
size = outlier)) +
geom_point(alpha = 0.9) +
scale_x_reverse() +

```



```

scale_y_reverse() +
labs(x = "F2 (ERB)", y = "F1 (ERB)") +
theme_minimal() +
theme_bw()
print(outlier_plot)
ggsave(
  "Figures/WE_outliers_plot.png",
  plot = outlier_plot,
  width = 10,
  height = 8,
  dpi = 300
)
...

```

Creating a table with only valid items:

```

```{r}
datafiltered=data[data$outlier=="valid",]
datafiltered$outlier <- as.factor (as.character(datafiltered$outlier))
summary(datafiltered)
write.table(datafiltered, file = paste("valid_only_", "000-009.Table", sep=""), sep =
"\t", row.names = FALSE, fileEncoding="UTF-8")
...

```

Plotting mean values for valid vowels:

```

```{r}

datafiltered_summary <- datafiltered %>%
  group_by(Vowel) %>%
  summarise(
    F1_mean = mean(F1_ERB),
    F1_sd = sd(F1_ERB),
    F2_mean = mean(F2_ERB),
    F2_sd = sd(F2_ERB)
  )

ggplot(datafiltered_summary, aes(x = F2_mean, y = F1_mean, color = Vowel)) +
  geom_point(size = 2) +
  geom_errorbar(aes(ymin = F1_mean - F1_sd, ymax = F1_mean + F1_sd, width = 0) +
  geom_errorbarh(aes(xmin = F2_mean - F2_sd, xmax = F2_mean + F2_sd, height = 0) +
  geom_text(aes(family = "serif", label = Vowel), vjust = -1.5, hjust = -0.5, size = 5,
show.legend = FALSE) +
  scale_x_reverse() +
  scale_y_reverse() +
  labs(
    x = "F2 (ERB)",
    y = "F1 (ERB)"
  ) +

```

```

theme_minimal() +
theme(
  legend.position = "none",
  panel.grid.major = element_line(linetype = "solid", color = "gray80"),
  panel.grid.minor = element_blank(),
  panel.border = element_rect(
    colour = "black",
    fill = NA,
    linewidth = 0.5),
  aspect.ratio = 1) +
theme_bw(base_family = "serif")
```

```

Creating a plot for word distribution along the vowel space:

```

```{r}
word_summary <- datafiltered %>%
  group_by(Speaker, Word, Vowel) %>%
  summarise(
    F1_mean = mean(F1_ERB),
    F2_mean = mean(F2_ERB),
    .groups = 'drop'
  ) %>%
  group_by(Word, Vowel) %>%
  summarise(

F1 = mean(F1_mean),
  F2 = mean(F2_mean),
  .groups = 'drop'
)

ggplot(word_summary, aes(x = F2, y = F1, color = Vowel)) +
  geom_text(aes(label = Word),
    family = "serif",
    size = 4,
    check_overlap = FALSE) +
  scale_x_reverse() +
  scale_y_reverse() +
  coord_cartesian(
    xlim = c(24, 13),
    ylim = c(15, 6)
  ) +
  labs(x = "F2 (ERB)", y = "F1 (ERB)") +
  theme_bw(base_family = "serif") +
  theme(
    legend.position = "none",
    panel.grid.major = element_line(color = "gray80"),
    panel.grid.minor = element_blank(),
    panel.border = element_rect(colour = "black", fill = NA, linewidth = 0.5),

```

```

    aspect.ratio = 1
  )

ggsave("Figures/WE_word_chart.png", width = 8, height = 8, units = "in")
```

```

Now for schwa\_eps and o-O and e and u (vowels across speakers, not words and not averaged)

```

```{r}
selected_vowels <- c("o", "ɔ", "u", "e", "ɛ", "ə")

vowel_plot <- datafiltered %>%
  filter(Vowel %in% selected_vowels) %>%
  ggplot(aes(x = F2_ERB, y = F1_ERB, color = Vowel)) +
  geom_point(size = 2, alpha = 0.6) +
  geom_text(aes(label = Vowel),
            vjust = -0.8, hjust = -0.3,
            size = 4, family = "serif",
            check_overlap = TRUE) +
  scale_x_reverse() +
  scale_y_reverse() +
  coord_cartesian(
    xlim = c(24, 11),
    ylim = c(15, 6)
  ) +
  labs(x = "F2 (ERB)", y = "F1 (ERB)") +
  theme_bw(base_family = "serif") +
  theme(
    legend.position = "none",
    panel.grid.major = element_line(color = "gray80"),
    panel.grid.minor = element_blank(),
    panel.border = element_rect(colour = "black", fill = NA, linewidth = 0.5),
    aspect.ratio = 1
  )

# Save the plot
ggsave("Figures/all_vowel_instances.png",
      plot = vowel_plot,
      width = 8,
      height = 8,
      units = "in",
      dpi = 300)
print(vowel_plot)
```

```

Now for the "front" vowels:

```

```{r}
selected_vowels <- c("e", "ɛ", "ə")

```

```

vowel_plot <- datafiltered %>%
  filter(Vowel %in% selected_vowels) %>%
  ggplot(aes(x = F2_ERB, y = F1_ERB, color = Vowel)) +
  geom_point(size = 2, alpha = 0.6) +
  geom_text(aes(label = Vowel),
            vjust = -0.8, hjust = -0.3,
            size = 4, family = "serif",
            check_overlap = TRUE) +
  scale_x_reverse() +
  scale_y_reverse() +
  coord_cartesian(
    xlim = c(24, 11),
    ylim = c(15, 6)
  ) +
  labs(x = "F2 (ERB)", y = "F1 (ERB)") +
  theme_bw(base_family = "serif") +
  theme(
    legend.position = "none",
    panel.grid.major = element_line(color = "gray80"),
    panel.grid.minor = element_blank(),
    panel.border = element_rect(colour = "black", fill = NA, linewidth = 0.5),
    aspect.ratio = 1
  ) +
  scale_color_manual(values = c("e" = "black", "ɛ" = "red", "ə" = "#56B4E9"))

# Save the plot
ggsave("Figures/front_vowel_instances.png",
       plot = vowel_plot,
       width = 8,
       height = 8,
       units = "in",
       dpi = 300)
print(vowel_plot)
``

```

Now for the back vowels:

```

`` {r}
selected_vowels <- c("o", "ɔ", "u")

vowel_plot <- datafiltered %>%
  filter(Vowel %in% selected_vowels) %>%
  ggplot(aes(x = F2_ERB, y = F1_ERB, color = Vowel)) +
  geom_point(size = 2, alpha = 0.6) +
  geom_text(aes(label = Vowel),
            vjust = -0.8, hjust = -0.3,
            size = 4, family = "serif",
            check_overlap = TRUE) +
  scale_x_reverse() +
  scale_y_reverse() +
  coord_cartesian(

```

```

    xlim = c(24, 11),
    ylim = c(15, 6)
  ) +
  labs(x = "F2 (ERB)", y = "F1 (ERB)") +
  theme_bw(base_family = "serif") +
  theme(
    legend.position = "none",
    panel.grid.major = element_line(color = "gray80"),
    panel.grid.minor = element_blank(),
    panel.border = element_rect(colour = "black", fill = NA, linewidth = 0.5),
    aspect.ratio = 1
  ) +
  scale_color_manual(values = c("u" = "black", "ɔ" = "red", "o" = "#56B4E9"))

# Save the plot
ggsave("Figures/back_vowel_instances.png",
  plot = vowel_plot,
  width = 8,
  height = 8,
  units = "in",
  dpi = 300)
print(vowel_plot)
``

```

Creating a vowel chart with data grouped by participant sex:

```

``{r}

datafiltered_summary <- datafiltered %>%
  group_by(Gender, Vowel) %>%
  summarise(
    F1_mean = mean(F1_ERB),
    F1_sd = sd(F1_ERB),
    F2_mean = mean(F2_ERB),
    F2_sd = sd(F2_ERB)
  )

ggplot(datafiltered_summary, aes(x = F2_mean, y = F1_mean, color = Gender)) +
  geom_point(size = 2, show.legend = FALSE) +
  geom_errorbar(aes(ymin = F1_mean - F1_sd, ymax = F1_mean + F1_sd), width = 0) +
  geom_errorbarh(aes(xmin = F2_mean - F2_sd, xmax = F2_mean + F2_sd), height = 0) +
  geom_text(aes(family = "serif", label = Vowel), vjust = -1.5, hjust = -0.5, size = 5,
show.legend = FALSE) +
  scale_x_reverse() +
  scale_y_reverse() +
  coord_cartesian(
    xlim = c(24, 13),
    ylim = c(15, 6)
  ) +
  labs(
    x = "F2 (ERB)",

```

```

    y = "F1 (ERB)"
  ) +
  theme_minimal() +
  theme_bw(base_family = "serif") +
  theme(
    legend.position = "none",
    panel.grid.major = element_line(linetype = "solid", color = "gray80"),
    panel.grid.minor = element_blank(),
    panel.border = element_rect(
      colour = "black",
      fill = NA,
      linewidth = 0.5), aspect.ratio = 1)

ggsave("Figures/WE_vowel_chart_by_sex.png", width = 6, height = 6, units = "in")
```


same but with normalised data:


```

```{r}
datalob0 = datafiltered %>% group_by(Gender) %>% mutate(F1scale = scale(F1_ERB),
F2scale = scale(F2_ERB))

f1.f2.summary_by_sex <- ddply(datalob0, c("Vowel", "Gender"), summarise,
 F1mean = mean(F1scale, na.rm = TRUE), F1sd = sd(F1scale, na.rm =
TRUE), # z-scores
 F2mean = mean(F2scale, na.rm = TRUE), F2sd = sd(F2scale, na.rm =
TRUE)) # z-scores

ggplot(f1.f2.summary_by_sex, aes(x = F2mean, y = F1mean, label = Vowel, color = Gender))
+
 geom_point(size = 2, show.legend = FALSE) +
 geom_text(aes(family = "serif"), hjust = 1.5, vjust = -.5, size = 5) +
 scale_y_reverse() +
 scale_x_reverse() +
 xlab("F2 (normalised)") + ylab("F1 (normalised)") +
 geom_errorbarh(data = f1.f2.summary_by_sex, aes(xmin = F2mean - F2sd, xmax = F2mean
+ F2sd, y = F1mean),
 height = .01, show.legend = FALSE) +
 geom_errorbar(data = f1.f2.summary_by_sex, aes(ymin = F1mean - F1sd, ymax = F1mean +
F1sd, x = F2mean),
 width = .01, show.legend = FALSE) +
 theme_bw(base_family = "serif") +
 theme(legend.position = "none", aspect.ratio = 1)

ggsave("Figures/WE_vowel_chart_by_sex_normalised.png", width = 6, height = NA, dpi =
300, limitsize = TRUE)
```

```


```

Creating vowel charts for each speaker:

```
```{r}
datafiltered_summary <- datafiltered %>%
  group_by(Speaker, Vowel) %>%
  summarise(
    F1_mean = mean(F1_ERB),
    F1_sd = sd(F1_ERB),
    F2_mean = mean(F2_ERB),
    F2_sd = sd(F2_ERB),
    .groups = 'keep'
  )

speakers <- unique(datafiltered$Speaker)
for(spr in speakers) {
  p <- datafiltered_summary %>%
    filter(Speaker == spr) %>%
    ggplot(aes(x = F2_mean, y = F1_mean, color = Vowel)) +
    geom_point(size = 2) +
    geom_errorbar(aes(ymin = F1_mean - F1_sd, ymax = F1_mean + F1_sd, width = 0) +
    geom_errorbarh(aes(xmin = F2_mean - F2_sd, xmax = F2_mean + F2_sd, height = 0) +
    geom_text(aes(label = Vowel), vjust = -1.5, hjust = -0.5, size = 5,
      family = "serif", show.legend = FALSE) +
    scale_x_reverse() +
    scale_y_reverse() +
    coord_cartesian(
      xlim = c(24, 13),
      ylim = c(15, 6)
    ) +
    labs(x = "F2 (ERB)", y = "F1 (ERB)", title = paste("Speaker:", spr)) +
    theme_bw(base_family = "serif") +
    theme(legend.position = "none",
      aspect.ratio = 1)

  print(p)
  ggsave(paste0("Figures/perSpeaker/WE_vowel_chart_", spr, ".png"), plot = p, width = 6,
  height = 6)
}
```
```

The same thing but creating a single picture with all plots:

```
```{r}

datafiltered_summary <- datafiltered %>%
  group_by(Speaker, Vowel) %>%
  summarise(
    F1_mean = mean(F1_ERB),
    F1_sd = sd(F1_ERB),
    F2_mean = mean(F2_ERB),
    F2_sd = sd(F2_ERB),
    .groups = 'keep'
```

```

)

faceted_plot <- ggplot(datafiltered_summary,
  aes(x = F2_mean, y = F1_mean, color = Vowel)) +
  geom_point(size = 2) +
  geom_errorbar(aes(ymin = F1_mean - F1_sd, ymax = F1_mean + F1_sd),
    width = 0, linewidth = 0.3) +
  geom_errorbarh(aes(xmin = F2_mean - F2_sd, xmax = F2_mean + F2_sd),
    height = 0, linewidth = 0.3) +
  geom_text(
    aes(label = Vowel),
    family = "serif",
    vjust = -0.6,
    hjust = -0.3,
    size = 3.5,
    nudge_x = 0.1,
    nudge_y = 0.1
  ) +
  scale_x_reverse(expand = expansion(mult = 0.1)) +
  scale_y_reverse(expand = expansion(mult = 0.1)) +
  coord_cartesian(xlim = c(24, 13), ylim = c(15, 6.5)) +
  facet_wrap(~ Speaker, ncol = 3) +
  labs(x = "F2 (ERB)", y = "F1 (ERB)") +
  theme_bw(base_family = "serif") +
  theme(
    legend.position = "none",
    aspect.ratio = 1
  )
)

# Save plot
ggsave("Figures/perSpeaker/WE_vowel_chart_by_speaker.png",
  plot = faceted_plot,
  width = 10,
  height = ceiling(length(unique(datafiltered_summary$Speaker))/3) * 3.5,
  units = "in",
  dpi = 300)
print(faceted_plot)
```

```

Because I was ironically taking forever to make a script that appended selected data ("vowel", "Dialect", "speaker", "F1" (in Hz) and "F2" (in Hz)) from datafiltered to a file with the data on Majorcan and Eastern Eivissan from Recasens & Espinosa (2006:665) Hamann & Torres-Tamarit (2023: p.c.), I just did it manually and called the new file "Ales\_3dialects.txt".

Creating a table with normalised F1-F2 values per vowel across all speakers.

```

```{r}
dataDialects = read.delim ("Ales_3dialects_with_pretonic_schwas.txt")
datalob = dataDialects %>% group_by(Dialect) %>% mutate(F1scale = scale(F1), F2scale =
scale(F2))
```

```



Summarising the data:

```
```{r}
f1.f2.summary <- ddpoly(datalob, c("vowel", "Dialect"), summarise,
  F1mean = mean(F1scale, na.rm = TRUE), F1sd = sd(F1scale, na.rm =
TRUE), # z-scores
  F2mean = mean(F2scale, na.rm = TRUE), F2sd = sd(F2scale, na.rm =
TRUE) # z-scores
)

f1.f2.summary
```
```

Plotting the three dialects in the same graph for comparison (normalised values):

```
```{r}
ggplot(f1.f2.summary, aes(x = F2mean, y = F1mean, label = vowel, color = Dialect)) +
  geom_point(size = 2, show.legend = FALSE) +
  geom_text(aes(family = "serif"), hjust = 1.5, vjust = -.5, size = 5) +
  scale_y_reverse() +
  scale_x_reverse() +
  xlab("F2 (normalised)") + ylab("F1 (normalised)") +
  geom_errorbarh(data = f1.f2.summary, aes(xmin = F2mean - F2sd, xmax = F2mean + F2sd,
y = F1mean),
  height = .01, show.legend = FALSE) +
  geom_errorbar(data = f1.f2.summary, aes(ymin = F1mean - F1sd, ymax = F1mean + F1sd, x
= F2mean),
  width = .01, show.legend = FALSE) +
  theme_bw(base_family = "serif") +
  theme(legend.position = "none", aspect.ratio = 1) +
  scale_color_manual(values = c("Eastern Eivissan" = "black", "Majorcan" = "red", "Western
Eivissan" = "#56B4E9"))
ggsave("Figures/Ales_3dialects_R_normalised.png", width = 6, height = NA, dpi = 300,
limitsize = TRUE)
```
```

Making the same normalised plot but with only Eivissan varieties:

```
```{r}
data2 <- dataDialects[dataDialects$Dialect == "Western Eivissan" | dataDialects$Dialect ==
"Eastern Eivissan" , ]
data2$Dialect <- as.factor (data2$Dialect)
levels (data2$Dialect)

```{r}
datalob2 = data2 %>% group_by(Dialect) %>% mutate(F1scale = scale(F1), F2scale =
scale(F2))
```
```

```

```{r}
f1.f2.summary2 <- ddpoly(datalob2, c("vowel", "Dialect"), summarise,
 F1mean = mean(F1scale, na.rm = TRUE), F1sd = sd(F1scale, na.rm =
TRUE),
 F2mean = mean(F2scale, na.rm = TRUE), F2sd = sd(F2scale, na.rm =
TRUE)
)

f1.f2.summary2
```

```{r}
ggplot(f1.f2.summary2, aes(x = F2mean, y = F1mean, label = vowel, color = Dialect)) +
 geom_point(size = 2, show.legend = FALSE) +
 geom_text(aes(family = "serif"), hjust = 1.5, vjust = -.5, size = 5) +
 scale_y_reverse() +
 scale_x_reverse() +
 xlab("F2 (normalised)") + ylab("F1 (normalised)") +
 geom_errorbarh(data = f1.f2.summary2, aes(xmin = F2mean - F2sd, xmax = F2mean +
F2sd, y = F1mean),
 height = .01, show.legend = FALSE) +
 geom_errorbar(data = f1.f2.summary2, aes(ymin = F1mean - F1sd, ymax = F1mean + F1sd,
x = F2mean),
 width = .01, show.legend = FALSE) +
 theme_bw(base_family = "serif") +
 theme(legend.position = "none", aspect.ratio = 1) +
 scale_color_manual(values = c("Eastern Eivissan" = "black", "Western Eivissan" =
"#56B4E9"))
ggsave("Figures/Ales_2dialects_EEWE_R_normalised.png", width = 6, height = NA, dpi =
300, limitsize = TRUE)
```

```

Making the same normalised plot but only with Western Eivissan and Majorcan:

```

```{r}
data3 <- dataDialects[dataDialects$Dialect == "Western Eivissan" | dataDialects$Dialect ==
"Majorcan",]
data3$Dialect <- as.factor (data3$Dialect)
levels (data3$Dialect)
```

```{r}
datalob3 = data3 %>% group_by(Dialect) %>% mutate(F1scale = scale(F1), F2scale =
scale(F2))
```

```{r}
f1.f2.summary3 <- ddpoly(datalob3, c("vowel", "Dialect"), summarise,
 F1mean = mean(F1scale, na.rm = TRUE), F1sd = sd(F1scale, na.rm =
TRUE),
 F2mean = mean(F2scale, na.rm = TRUE), F2sd = sd(F2scale, na.rm =
TRUE)
)

```

```

F2mean = mean(F2scale, na.rm = TRUE), F2sd = sd(F2scale, na.rm =
TRUE)
)

f1.f2.summary3
```
```{r}
ggplot(f1.f2.summary3, aes(x = F2mean, y = F1mean, label = vowel, color = Dialect)) +
 geom_point(size = 2, show.legend = FALSE) +
 geom_text(aes(family = "serif"), hjust = 1.5, vjust = -.5, size = 5) +
 scale_y_reverse() +
 scale_x_reverse() +
 xlab("F2 (normalised)") + ylab("F1 (normalised)") +
 geom_errorbarh(data = f1.f2.summary3, aes(xmin = F2mean - F2sd, xmax = F2mean +
F2sd, y = F1mean),
 height = .01, show.legend = FALSE) +
 geom_errorbar(data = f1.f2.summary3, aes(ymin = F1mean - F1sd, ymax = F1mean + F1sd,
x = F2mean),
 width = .01, show.legend = FALSE) +
 theme_bw(base_family = "serif") +
 theme(legend.position = "none", aspect.ratio = 1) +
 scale_color_manual(values = c("Majorcan" = "red", "Western Eivissan" = "#56B4E9"))
ggsave("Figures/Ales_2dialects_MCWE_R_normalised.png", width = 6, height = NA, dpi =
300, limitsize = TRUE)
```

```

Vowel quality analysis

schwa_eps

Creating a sub-dataset for schwa_eps:

```

```{r}
schwa_eps <- droplevels(datafiltered [datafiltered$Vowel %in% c("ə", "ɛ"),]) %>%
 select(Speaker, Vowel, F1_ERB, F2_ERB, Word, Gender, Duration_log, Context) %>%
 rename (F1 = F1_ERB, F2 = F2_ERB)
summary(schwa_eps)
```

```

MANOVA with F1 and F2 as dependent variables and Vowel as independent variable:

```

```{r}
schwa_eps_manova <- manova (cbind(F1, F2) ~ Vowel, data = schwa_eps)
summary(schwa_eps_manova)
help(manova)
```

```

MANOVA with F1 and F2 as dependent variables and Vowel, Gender, Speaker and Word as independent variables:

```

```{r}

```

```
schwa_eps_manova <- manova (cbind(F1, F2) ~ Vowel + Gender + Speaker + Word, data =
schwa_eps)
summary(schwa_eps_manova, correlation = TRUE)
```

```

Setting contrasts for F1 and F2 lmer:

```
```{r}
contrast <- cbind (c(-0.5, +0.5))
colnames (contrast) <- c("-ə+ε")
contrasts (schwa_eps$Vowel) <- contrast
contrasts (schwa_eps$Vowel)
```

```

lmer with F1 as dependent variable, Vowel as an independent variable and word and speaker as random intercepts and vowel pair by speaker as random slope:

```
```{r}
schwa_eps_f1_lmer <- lmer(F1 ~ Vowel + (1|Word) + (1|Speaker) + (0 + Vowel | Speaker),
data = schwa_eps)
summary(schwa_eps_f1_lmer)
confint(schwa_eps_f1_lmer)
```

```

lmer with F2 as dependent variable, Vowel as an independent variable and word and speaker as random intercepts and vowel pair by speaker as random slope:

```
```{r}
schwa_eps_f2_lmer <- lmer(F2 ~ Vowel + (1|Word) + (1|Speaker) + (0 + Vowel | Speaker),
data = schwa_eps)
summary(schwa_eps_f2_lmer)
confint(schwa_eps_f2_lmer)
```

```

e_eps

Creating a subset for e_eps:

```
```{r}
e_eps <- droplevels(datafiltered [datafiltered$Vowel %in% c("e", "ε"),]) %>%
 select(Speaker, Vowel, F1_ERB, F2_ERB, Word, Gender, Duration_log) %>%
 rename (F1 = F1_ERB, F2 = F2_ERB)
summary(e_eps)
```

```

MANOVA with F1 and F2 as dependent variable and Vowel as independent variable:

```
```{r}
e_eps_manova <- manova (cbind(F1, F2) ~ Vowel, data = e_eps)
summary(e_eps_manova)
```

```

```
```
```

MANOVA with F1 and F2 as dependent variable and Vowel, Speaker, Gender and Word as independent variables:

```
```{r}
e_eps_manova <- manova (cbind(F1, F2) ~ Vowel + Gender + Speaker + Word, data = e_eps)
summary(e_eps_manova)
```
```

```
e_schwa
```

Creating a subset for e\_schwa:

```
```{r}
e_schwa <- droplevels(datafiltered [datafiltered$Vowel %in% c("e", "ə"), ]) %>%
  select(Speaker, Vowel, F1_ERB, F2_ERB, Word, Gender, Duration_log) %>%
  rename (F1 = F1_ERB, F2 = F2_ERB)
summary(e_schwa)
```
```

MANOVA with F1 and F2 as dependent variable and Vowel as independent variable:

```
```{r}
e_schwa_manova <- manova (cbind(F1, F2) ~ Vowel, data = e_schwa)
summary(e_schwa_manova)
```
```

MANOVA with F1 and F2 as dependent variable and Vowel, Speaker, Gender and Word as independent variables:

```
```{r}
e_schwa_manova <- manova (cbind(F1, F2) ~ Vowel + Gender + Speaker + Word, data =
e_schwa)
summary(e_schwa_manova)
```
```

```
Pillai score threshold calculations
```

Using Stanley & Sneller's (2023:6) formula in R:

```
```{r}
schwa_eps_pillai_threshold = schwa_eps %>%
  group_by(Speaker) %>%
  summarize(schwa_eps_per_speaker = n(),
            m = n()/2, # Each row in the data frame corresponds to a vowel, so n() works without
my Vowel variable
            pillai_threshold = exp(1)/m)
print(schwa_eps_pillai_threshold, n = Inf)
```
```

```
e_eps_pillai_threshold = e_eps %>%
 group_by(Speaker) %>%
 summarize(e_eps_per_speaker = n(),
 m = n()/2, # Each row in the data frame corresponds to a vowel, so n() works without
my Vowel variable
 pillai_threshold = exp(1)/m)
print(e_eps_pillai_threshold, n = Inf)
```

```
e_schwa_pillai_threshold = e_schwa %>%
 group_by(Speaker) %>%
 summarize(e_schwa_per_speaker = n(),
 m = n()/2, # Each row in the data frame corresponds to a vowel, so n() works without
my Vowel variable
 pillai_threshold = exp(1)/m)
print(e_schwa_pillai_threshold, n = Inf)
```

```
listOfTables0 <- list(entries0 = cbind(schwa_eps_pillai_threshold, e_eps_pillai_threshold,
e_schwa_pillai_threshold))
```

```
big_pillai_threshold_table = list(entries0 = 0)
for (pillai_threshold_table in listOfTables0) {

 big_pillai_threshold_table$entries0 <- pillai_threshold_table
}
big_pillai_threshold_table
````
```

Defining two functions for Pillai scores and p-values (non-corrected values):

```
````{r}
schwa_eps_pillai <- function(...) {
 summary(manova(...))$stats["Vowel", "Pillai"]
}
schwa_eps_pvalue <- function(...) {
 summary(manova(...))$stats["Vowel", "Pr(>F)"]
}
e_eps_pillai <- function(...) {
 summary(manova(...))$stats["Vowel", "Pillai"]
}
e_eps_pvalue <- function(...) {
 summary(manova(...))$stats["Vowel", "Pr(>F)"]
}
e_schwa_pillai <- function(...) {
 summary(manova(...))$stats["Vowel", "Pillai"]
}
e_schwa_pvalue <- function(...) {
```

```
summary(manova(...))$stats["Vowel", "Pr(>F)"]
}
...
```

Calling the values just to check (non-corrected values):

```
```{r}
schwa_eps_pillai(cbind(F1, F2) ~ Vowel, data = schwa_eps)
schwa_eps_pvalue(cbind(F1, F2) ~ Vowel, data = schwa_eps)
e_eps_pillai(cbind(F1, F2) ~ Vowel, data = e_eps)
e_eps_pvalue(cbind(F1, F2) ~ Vowel, data = e_eps)
e_schwa_pillai(cbind(F1, F2) ~ Vowel, data = e_schwa)
e_schwa_pvalue(cbind(F1, F2) ~ Vowel, data = e_schwa)
```
```

Table split by speaker (non-corrected values):

```
```{r}
table1 = schwa_eps %>%
  group_by(Speaker) %>%
  summarize(ə_ε_pillai = schwa_eps_pillai(cbind(F1, F2) ~ Vowel),
    ə_ε_pvalue = schwa_eps_pvalue(cbind(F1, F2) ~ Vowel))
```

```
table2 = e_eps %>%
  group_by(Speaker) %>%
  summarize(e_ε_pillai = e_eps_pillai(cbind(F1, F2) ~ Vowel),
    e_ε_pvalue = e_eps_pvalue(cbind(F1, F2) ~ Vowel))
```

```
table3 = e_schwa %>%
  group_by(Speaker) %>%
  summarize(e_ə_pillai = e_schwa_pillai(cbind(F1, F2) ~ Vowel),
    e_ə_pvalue = e_schwa_pvalue(cbind(F1, F2) ~ Vowel))
```

Extremely imperfect way to print a dataframe. At this stage, making a tsv/txt file with proper order
 # (three columns: Speaker, Pillai, score) by hand will be quicker for me than looking up how to let R do it

```
listOfTables <- list(entries = cbind(table1, table2, table3))
```

```
bigtable = list(entries = 0)
for (pillaiTable in listOfTables) {
```

```
  bigtable$entries <- pillaiTable
}
bigtable
```

```
...
```

Reading the cumulative mid vowels data frame from a hand-made file from previous table:

```
```{r}
pillaiData = read.delim ("WE_Pillai_scores.txt", stringsAsFactors = TRUE)
head(pillaiData)
```
```

Checking the values in the Pillai column:

```
```{r}
levels (pillaiData$Pillai)
```
```

Creating a subset with only front vowels:

```
```{r}
data.front <- pillaiData[pillaiData$Pillai=="ə_ε"| pillaiData$Pillai=="e_ε"|
pillaiData$Pillai=="e_ə",]
data.front$Pillai <- as.factor (as.character(data.front$Pillai))
levels (data.front$Pillai)
```
```

Assigning different order to vowel pairs, schwa first, then reference pairs:

```
```{r}
data.front$Pillai <- factor(data.front$Pillai, levels = c("ə_ε", "e_ε", "e_ə"))
```
```

Plotting the Pillai score differences between the three vowel pairs:

```
```{r}
ggplot(data.front, aes(x = Pillai, y = Score)) +
 geom_boxplot(fill = "grey") +
 xlab ("Vowel pair") + ylab ("Pillai score") + scale_y_continuous(limits = c(0, 1)) +
 theme_bw(base_family = "serif") +
 theme(aspect.ratio = 1)
```
```

Saving the plot to file:

```
```{r}
ggsave("Figures/WE_Pillai_scores_front.png", width = 4, height = 3.5, dpi = 300, limitsize =
TRUE)
```
```

closeO_openO

Creating a sub-dataset for closeO_openO:

```
```{r}
closeO_openO <- droplevels(datafiltered [datafiltered$Vowel %in% c("o", "ɔ"),] %>%
```



```

select(Speaker, Vowel, F1_ERB, F2_ERB, Word, Context, Duration_log, Gender) %>%
 rename (F1 = F1_ERB, F2 = F2_ERB)
summary(closeO_openO)
```

```

MANOVA with F1 and F2 as dependent variables and Vowel as independent variable:

```

```{r}
closeO_openO_manova <- manova (cbind(F1, F2) ~ Vowel, data = closeO_openO)
summary(closeO_openO_manova)
```

```

MANOVA with F1 and F2 as dependent variables and Vowel, Speaker and Word as independent variables:

```

```{r}
closeO_openO_manova <- manova (cbind(F1, F2) ~ Vowel + Gender + Speaker + Word, data
= closeO_openO)
summary(closeO_openO_manova)
```

```

Setting contrasts for F1 and F2 lmer:

```

```{r}
contrast <- cbind (c(+0.5, -0.5))
colnames (contrast) <- c("-o+o")
contrasts (closeO_openO$Vowel) <- contrast
contrasts (closeO_openO$Vowel)
```

```

lmer with F1 as dependent variable, Vowel as an independent variable and word and speaker as random intercepts and vowel pair by speaker as random slope:

```

```{r}
closeO_openO_f1_lmer <- lmer(F1 ~ Vowel + (1|Word) + (1|Speaker) + (0 + Vowel |
Speaker), data = closeO_openO)
summary(closeO_openO_f1_lmer)
confint(closeO_openO_f1_lmer)
```

```

lmer with F2 as dependent variable, Vowel as an independent variable and word and speaker as random intercepts and vowel pair by speaker as random slope:

```

```{r}
closeO_openO_f2_lmer <- lmer(F2 ~ Vowel + (1|Word) + (1|Speaker) + (0 + Vowel |
Speaker), data = closeO_openO)
summary(closeO_openO_f2_lmer)
confint(closeO_openO_f2_lmer)
```

```

u_openO

Creating a subset for u_openO:

```
```{r}
u_openO <- droplevels(datafiltered [datafiltered$Vowel %in% c("u", "ɔ"),]) %>%
 select(Speaker, Vowel, F1_ERB, F2_ERB, Word, Gender, Duration_log) %>%
 rename (F1 = F1_ERB, F2 = F2_ERB)
summary(u_openO)
```
```

MANOVA with F1 and F2 as dependent variable and Vowel as independent variable:

```
```{r}
u_openO_manova <- manova (cbind(F1, F2) ~ Vowel, data = u_openO)
summary(u_openO_manova)
```
```

MANOVA with F1 and F2 as dependent variable and Vowel, Speaker, Gender and Word as independent variables:

```
```{r}
u_openO_manova <- manova (cbind(F1, F2) ~ Vowel + Gender + Speaker + Word, data =
u_openO)
summary(u_openO_manova)
```
```

u_closeO

Creating a subset for u_closeO:

```
```{r}
u_closeO <- droplevels(datafiltered [datafiltered$Vowel %in% c("u", "o"),]) %>%
 select(Speaker, Vowel, F1_ERB, F2_ERB, Word, Gender, Duration_log) %>%
 rename (F1 = F1_ERB, F2 = F2_ERB)
summary(u_closeO)
```
```

MANOVA with F1 and F2 as dependent variable and Vowel as independent variable:

```
```{r}
u_closeO_manova <- manova (cbind(F1, F2) ~ Vowel, data = u_closeO)
summary(u_closeO_manova)
```
```

MANOVA with F1 and F2 as dependent variable and Vowel, Speaker, Gender and Word as independent variables:

```
```{r}
u_closeO_manova <- manova (cbind(F1, F2) ~ Vowel + Gender + Speaker + Word, data =
u_closeO)
summary(u_closeO_manova)
```
```

```
## Pillai score threshold calculations
```

Using Stanley & Sneller's (2023:6) formula in R:

```
```{r}
closeO_openO_pillai_threshold = closeO_openO %>%
 group_by(Speaker) %>%
 summarize(closeO_openO_per_speaker = n(),
 m = n()/2, # Each row in the data frame corresponds to a vowel, so n() works without
my Vowel variable
 pillai_threshold = exp(1)/m)
print(closeO_openO_pillai_threshold, n = Inf)
```

```
u_openO_pillai_threshold = u_openO %>%
 group_by(Speaker) %>%
 summarize(u_openO_per_speaker = n(),
 m = n()/2, # Each row in the data frame corresponds to a vowel, so n() works without
my Vowel variable
 pillai_threshold = exp(1)/m)
print(u_openO_pillai_threshold, n = Inf)
```

```
u_closeO_pillai_threshold = u_closeO %>%
 group_by(Speaker) %>%
 summarize(u_closeO_per_speaker = n(),
 m = n()/2, # Each row in the data frame corresponds to a vowel, so n() works without
my Vowel variable
 pillai_threshold = exp(1)/m)
print(u_closeO_pillai_threshold, n = Inf)
```

```
listOfTables2 <- list(entries2 = cbind(closeO_openO_pillai_threshold,
u_openO_pillai_threshold, u_closeO_pillai_threshold))
```

```
big_pillai_threshold_table2 = list(entries2 = 0)
for (pillai_threshold_table2 in listOfTables2) {

 big_pillai_threshold_table2$entries2 <- pillai_threshold_table2
}
big_pillai_threshold_table2
```
```

Defining two functions for Pillai scores and p-values (non-corrected values):

```
```{r}
closeO_openO_pillai <- function(...) {
 summary(manova(...))$stats["Vowel", "Pillai"]
}
closeO_openO_pvalue <- function(...) {
 summary(manova(...))$stats["Vowel", "Pr(>F)"]
}
```

```

u_openO_pillai <- function(...) {
 summary(manova(...))$stats["Vowel", "Pillai"]
}
u_openO_pvalue <- function(...) {
 summary(manova(...))$stats["Vowel", "Pr(>F)"]
}
u_closeO_pillai <- function(...) {
 summary(manova(...))$stats["Vowel", "Pillai"]
}
u_closeO_pvalue <- function(...) {
 summary(manova(...))$stats["Vowel", "Pr(>F)"]
}
...

```

Calling the values just to check (non-corrected values):

```

```{r}
closeO_openO_pillai(cbind(F1, F2) ~ Vowel, data = closeO_openO)
closeO_openO_pvalue(cbind(F1, F2) ~ Vowel, data = closeO_openO)
u_openO_pillai(cbind(F1, F2) ~ Vowel, data = u_openO)
u_openO_pvalue(cbind(F1, F2) ~ Vowel, data = u_openO)
u_closeO_pillai(cbind(F1, F2) ~ Vowel, data = u_closeO)
u_closeO_pvalue(cbind(F1, F2) ~ Vowel, data = u_closeO)
```

```

Table split by speakers (non-corrected values):

```

```{r}
table4 = closeO_openO %>%
  group_by(Speaker) %>%
  summarize(o_o_pillai = closeO_openO_pillai(cbind(F1, F2) ~ Vowel),
    o_o_pvalue = closeO_openO_pvalue(cbind(F1, F2) ~ Vowel))

```

```

table5 = u_openO %>%
  group_by(Speaker) %>%
  summarize(u_o_pillai = u_openO_pillai(cbind(F1, F2) ~ Vowel),
    u_o_pvalue = u_openO_pvalue(cbind(F1, F2) ~ Vowel))

```

```

table6 = u_closeO %>%
  group_by(Speaker) %>%
  summarize(u_c_pillai = u_closeO_pillai(cbind(F1, F2) ~ Vowel),
    u_c_pvalue = u_closeO_pvalue(cbind(F1, F2) ~ Vowel))

```

Extremely imperfect way to print a dataframe. At this stage, making a tsv with proper order
 # (three columns: Speaker, Pillai, score) by hand will be quicker than looking up how to let R
 do it

```
secondListOfTables <- list(entries = cbind(table4, table5, table6))
```

```
secondBigTable = list(entries = 0)
```

```

for ( secondPillaiTable in secondListOfTables) {

  secondBigTable$entries <- secondPillaiTable
}
secondBigTable
'''

```

Creating a subset with only front vowels:

```

'''{r}
data.back <- pillaiData[pillaiData$Pillai=="o_ɔ"| pillaiData$Pillai=="u_o"|
pillaiData$Pillai=="u_ɔ", ]
data.back$Pillai <- as.factor (as.character(data.back$Pillai))
levels (data.back$Pillai)
'''

```

Assigning different order to vowel pairs, schwa first, then reference pairs:

```

'''{r}
data.back$Pillai <- factor(data.back$Pillai, levels = c("o_ɔ", "u_ɔ", "u_o"))
'''

```

Plotting the Pillai score differences between the three vowel pairs:

```

'''{r}
ggplot(data.back, aes(x = Pillai, y = Score)) +
  geom_boxplot(fill = "grey") +
  xlab ("Vowel pair") + ylab ("Pillai score") + scale_y_continuous(limits = c(0, 1)) +
  theme_bw(base_family = "serif") +
  theme(aspect.ratio = 1)
'''

```

Saving the plots to file:

```

'''{r}
ggsave("Figures/WE_Pillai_scores_back.png", width = 4, height = 3.5, dpi = 300, limitsize =
TRUE)
'''

```

Duration analysis

schwa_eps

Creating a subset of schwa_eps with the relevant minimal pairs:

```

'''{r}
schwa_eps_minimalish_pairs0 <- droplevels(schwa_eps[schwa_eps$Word %in% c("beu",
"peu", "mel", "pèl"), ]) %>%
  mutate(

```

```

Voicing = case_when(
  Word %in% c("beu", "mel") ~ "voiced",
  Word %in% c("peu", "pèl") ~ "voiceless",
  TRUE ~ NA_character_
),
pre_cons_context_front = case_when(
  Word %in% c("beu", "peu", "mel", "pèl") ~ "Labial",
  TRUE ~ NA_character_),
minimal_pairs_front = case_when(
  Word %in% c("beu", "peu") ~ "bəu-peu",
  Word %in% c("mel", "pèl") ~ "pəl-məl",
  TRUE ~ NA_character_
),
Voicing = factor(Voicing),
pre_cons_context_front = factor(pre_cons_context_front),
minimal_pairs_front = factor(minimal_pairs_front)
) %>%
select(Speaker, Vowel, Word, Duration_log, Voicing, pre_cons_context_front,
minimal_pairs_front)

```

```

summary(schwa_eps_minimalish_pairs0)
write.table(schwa_eps_minimalish_pairs0,
  file = paste("valid_only_", "schwa_eps_minimalish_pairs0.Table", sep = ""),
  sep = "\t",
  row.names = FALSE,
  fileEncoding = "UTF-8")
...

```

Plotting the schwa_eps minimal pair data:

```

```{r}
library(ggplot2)
ggplot(schwa_eps_minimalish_pairs0, aes(x=minimal_pairs_front, y=Duration_log, fill =
Vowel)) +
 geom_boxplot() +
 facet_wrap(~pre_cons_context_front) +
 labs(title = "Vowel Duration across minimal pairs") +
 xlab("Minimal pairs") + ylab("Log-transformed duration") +
 scale_fill_manual(breaks = schwa_eps_minimalish_pairs0$Vowel,
 values = c("ə" = "#1b98e0", "ɛ" = "red")) +
 theme_bw(base_family = "serif") +
 theme(legend.position = "none")
ggsave("Figures/mirrored_voicing_front_vowels_duration_differences.png", width = 4,
height = 3.5, dpi = 300, limitsize = TRUE)
...

```

Setting contrasts:

```

```{r}

```

```
contrast <- cbind (c(-0.5, +0.5))
colnames (contrast) <- c("-ə+ε")
contrasts (schwa_eps_minimalish_pairs0$Vowel) <- contrast
contrasts(schwa_eps_minimalish_pairs0$Vowel)
```

```
contrast <- cbind (c(+0.5, -0.5))
colnames (contrast) <- c("-voiceless+voiced")
contrasts (schwa_eps_minimalish_pairs0$Voicing) <- contrast
contrasts(schwa_eps_minimalish_pairs0$Voicing)
```
```

Running a lmer on all schwa\_eps minimal pairs:

```
```{r}
schwa_eps_duration_model0 <- lmer(Duration_log ~ Vowel * Voicing + (1 | Speaker), data =
schwa_eps_minimalish_pairs0)
summary(schwa_eps_duration_model0)
```

```{r}
confint(schwa_eps_duration_model0)
```
```

Doing the same for the first minimal pair:

```
```{r}
schwa_eps_minimalish_pairs1 <- droplevels(schwa_eps_minimalish_pairs0
[schwa_eps_minimalish_pairs0$Word %in% c("peu", "beu"), ]) %>%
  select(Speaker, Vowel, Word, Duration_log)
summary(schwa_eps_minimalish_pairs1)
```

```{r}
contrast <- cbind (c(-0.5, +0.5))
colnames (contrast) <- c("-ə+ε")
contrasts (schwa_eps_minimalish_pairs1$Vowel) <- contrast
contrasts (schwa_eps_minimalish_pairs1$Vowel)
```
```

```
```{r}
schwa_eps_duration_model1 <- lmer(Duration_log ~ Vowel + (1 | Speaker), data =
schwa_eps_minimalish_pairs1)
summary(schwa_eps_duration_model1)
```

```{r}
confint(schwa_eps_duration_model1)
```
```

Doing the same for the second minimal pair:

```
```{r}
schwa_eps_minimalish_pairs2 <- droplevels(schwa_eps_minimalish_pairs0
[schwa_eps_minimalish_pairs0$Word %in% c("pèl", "mel"), ]) %>%
  select(Speaker, Vowel, Word, Duration_log, Voicing)
summary(schwa_eps_minimalish_pairs2)
```

```

```
contrast <- cbind (c(-0.5, +0.5))
colnames (contrast) <- c("-ə+ε")
contrasts (schwa_eps_minimalish_pairs2$Vowel) <- contrast
contrasts (schwa_eps_minimalish_pairs2$Vowel)
```

```{r}
schwa_eps_duration_model2 <- lmer(Duration_log ~ Vowel + (1 | Speaker), data =
schwa_eps_minimalish_pairs2)
summary(schwa_eps_duration_model2)
```

```{r}
confint(schwa_eps_duration_model2)
```

#### closeO_openO

Creating a subset of closeO_openO with the relevant minimal pairs:
```{r}
closeO_openO_minimalish_pairs0 <- droplevels(closeO_openO[closeO_openO$Word %in%
c("bou", "pou", "moll", "poll"),]) %>%
mutate(
 Voicing = case_when(
 Word %in% c("bou", "moll") ~ "voiced",
 Word %in% c("pou", "poll") ~ "voiceless",
 TRUE ~ NA_character_
),
 pre_cons_context_back = case_when(
 Word %in% c("bou", "pou", "moll", "poll") ~ "Labial",
 TRUE ~ NA_character_
),
 minimal_pairs = case_when(
 Word %in% c("bou", "pou") ~ "bou-pou",
 Word %in% c("moll", "poll") ~ "moll-poll",
 TRUE ~ NA_character_
),
 Voicing = factor(Voicing),
 pre_cons_context_back = factor(pre_cons_context_back),
 minimal_pairs = factor(minimal_pairs)
) %>%
select(Speaker, Vowel, Word, Duration_log, Voicing, pre_cons_context_back,
minimal_pairs)

summary(closeO_openO_minimalish_pairs0)
write.table(closeO_openO_minimalish_pairs0,
 file = paste("valid_only_", "closeO_openO_minimalish_pairs0.Table", sep = ""),
 sep = "\t",
 row.names = FALSE,

```



```

 fileEncoding = "UTF-8")
 }

```

Plotting the closeO\_openO minimal pair data:

```

 {r}
 library(ggplot2)
 ggplot(closeO_openO_minimalish_pairs0, aes(x=minimal_pairs, y=Duration_log,
 fill=Vowel)) +
 geom_boxplot() +
 facet_wrap(~pre_cons_context_back) +
 labs(title = "Vowel Duration across minimal pairs") +
 xlab("Minimal pairs") + ylab("Log-transformed duration") +
 scale_fill_manual(breaks = closeO_openO_minimalish_pairs0$Vowel,
 values = c("o" = "#1b98e0", "ɔ" = "red")) +
 theme_bw(base_family = "serif") +
 theme(legend.position = "none")
 ggsave("Figures/back_vowels_duration_differences.png", width = 4, height = 3.5, dpi = 300,
 limitsize = TRUE)
 }

```

Setting the contrasts:

```

 {r}
 contrast <- cbind(c(+0.5, -0.5))
 colnames(contrast) <- c("-o+ɔ")
 contrasts(closeO_openO_minimalish_pairs0$Vowel) <- contrast
 contrasts(closeO_openO_minimalish_pairs0$Vowel)
 }

```

Running the lmer across closeO\_openO minimal pair data:

```

 {r}
 closeO_openO_duration_model0 <- lmer(Duration_log ~ Vowel + (1 | Speaker), data =
 closeO_openO_minimalish_pairs0)
 summary(closeO_openO_duration_model0)
 }

 {r}
 confint(closeO_openO_duration_model0)
 }

```

Doing the same for the first minimal pair:

```

 {r}
 closeO_openO_minimalish_pairs1 <- droplevels(closeO_openO_minimalish_pairs0
 [closeO_openO_minimalish_pairs0$Word %in% c("bou", "pou"),]) %>%
 select(Speaker, Vowel, Word, Duration_log, Voicing)
 summary(closeO_openO_minimalish_pairs1)
 }

 {r}
 contrast <- cbind(c(-0.5, +0.5))
 colnames(contrast) <- c("-o+ɔ")

```

```

contrasts (closeO_openO_minimalish_pairs1$Vowel) <- contrast
contrasts(closeO_openO_minimalish_pairs1$Vowel)
```

```{r}
closeO_openO_duration_model1 <- lmer(Duration_log ~ Vowel + (1 | Speaker), data =
closeO_openO_minimalish_pairs1)
summary(closeO_openO_duration_model1)
```

```{r}
confint(closeO_openO_duration_model1)
```

```

Doing the same for the second minimal pair:

```

```{r}
closeO_openO_minimalish_pairs2 <- droplevels(closeO_openO_minimalish_pairs0
[closeO_openO_minimalish_pairs0$Word %in% c("poll", "moll"),]) %>%
 select(Speaker, Vowel, Word, Duration_log, Voicing)
summary(closeO_openO_minimalish_pairs2)
```

```{r}
contrast <- cbind (c(-0.5, +0.5))
colnames (contrast) <- c("-o+ɔ")
contrasts (closeO_openO_minimalish_pairs2$Vowel) <- contrast
contrasts (closeO_openO_minimalish_pairs2$Vowel)
```

```{r}
closeO_openO_duration_model2 <- lmer(Duration_log ~ Vowel + (1 | Speaker), data =
closeO_openO_minimalish_pairs2)
summary(closeO_openO_duration_model2)
```

```{r}
confint(closeO_openO_duration_model2)
```

```

Appendix E: Praat script for plotting

Adapted plotting script by Hamann & Torres-Tamarit (2023).
The contents of Ales_3dialects_averaged.txt were created separately in R.

```
;Read from file... Ales_3dialects_with_pretonic_schwas_averaged.txt
```

```

Erase all
Select outer viewport... 0.24 6 0.4 5.6
Select inner viewport... 1 5.8 0.5 5

```

Line width... 0.8

Black

12

Axes: 2500, 800, 800, 300

Draw line: 1191, 442, 1199, 587 ; u-c Eastern Eivissan

Draw line: 1197, 593, 1553, 781 ; o-a

Draw line: 1553, 781, 1770, 679 ; a-E

Draw line: 1770, 679, 2079, 521; E-e

Draw line: 2079, 521, 2379, 403; e-i

Red

Draw line: 899, 394, 999, 546 ; u-o Majorcan

Draw line: 999, 546, 1178, 708 ; o-c

Draw line: 1178, 708, 1463, 739 ; c-a

Draw line: 1463, 739, 1739, 658 ; a-eps

Draw line: 1739, 658, 1905, 489; eps-e

Draw line: 1905, 489, 2151, 327; e-i

Cyan

Draw line: 1093, 385, 1096, 499 ; u-o Western Eivissan

Draw line: 1096, 499, 1217, 618 ; o-c

Draw line: 1217, 618, 1507, 769 ; c-a

Draw line: 1507, 769, 1961, 618 ; a-E

Draw line: 1961, 618, 2178, 480; E-e

Draw line: 2178, 480, 2435, 351; e-i

Black

Paint circle: "white", 1191, 442, 35; u Eastern Eivissan

Paint circle: "white", 1199, 587, 35; o

Paint circle: "white", 1197, 593, 35; c

Paint circle: "white", 1553, 781, 35; a

Paint circle: "white", 1750, 670, 35; eps

Paint circle: "white", 1788, 667, 35; schwa

Paint circle: "white", 1665, 663, 35; pretonicSchwa

Paint circle: "white", 2079, 521, 35; e

Paint circle: "white", 2379, 403, 35; i

Paint circle: "white", 899, 394, 35; u Majorcan

Paint circle: "white", 999, 546, 35; o

Paint circle: "white", 1178, 708, 35; c

Paint circle: "white", 1463, 739, 35; a

Paint circle: "white", 1739, 658, 35; eps

Paint circle: "white", 1393, 563, 35; schwa

Paint circle: "white", 1905, 489, 35; e

Paint circle: "white", 2151, 327, 35; i

Paint circle: "white", 1093, 385, 35; u Western Eivissan

Paint circle: "white", 1096, 499, 35; o

Paint circle: "white", 1217, 618, 35; c

Paint circle: "white", 1507, 769, 35; a

Paint circle: "white", 1935, 621, 35; eps

Paint circle: "white", 1987, 616, 35; schwa

Paint circle: "white", 1691, 601, 35; pretonicSchwa
Paint circle: "white", 2178, 480, 35; e
Paint circle: "white", 2435, 351, 35; i

select Table Ales_3dialects_with_pretonic_schwas_averaged
table = Extract rows where column (text)... Dialect "is equal to" Eastern Eivissan
Scatter plot: "F2mean", 2500, 800, "F1mean", 800, 300, "vowel", 18, "no"
Marks bottom every... 1 150 no yes yes
Marks bottom every... 1 300 yes no no
Marks left every... 1 100 yes no no
Marks left every... 1 50 no yes yes
select table
Remove

Cyan
select Table Ales_3dialects_with_pretonic_schwas_averaged
table2 = Extract rows where column (text)... Dialect "is equal to" Western Eivissan
Scatter plot: "F2mean", 2500, 800, "F1mean", 800, 300, "vowel", 18, "no"

select table2
Remove

Red
select Table Ales_3dialects_with_pretonic_schwas_averaged
table3 = Extract rows where column (text)... Dialect "is equal to" Majorcan
Scatter plot: "F2mean", 2500, 800, "F1mean", 800, 300, "vowel", 18, "no"
Black
Line width... 0.4
Draw inner box
14
Text left: "yes", "F1 (Hz)"

Text bottom: "yes", "F2 (Hz)"
12

select table3
Remove

select Table Ales_3dialects_with_pretonic_schwas_averaged

Appendix F: Results of MANOVAs per speaker

Table F1: Pillai scores, Pillai thresholds and p -values for /ə/-/ɛ/, /e/-/ɛ/ and /e/-/ə/.

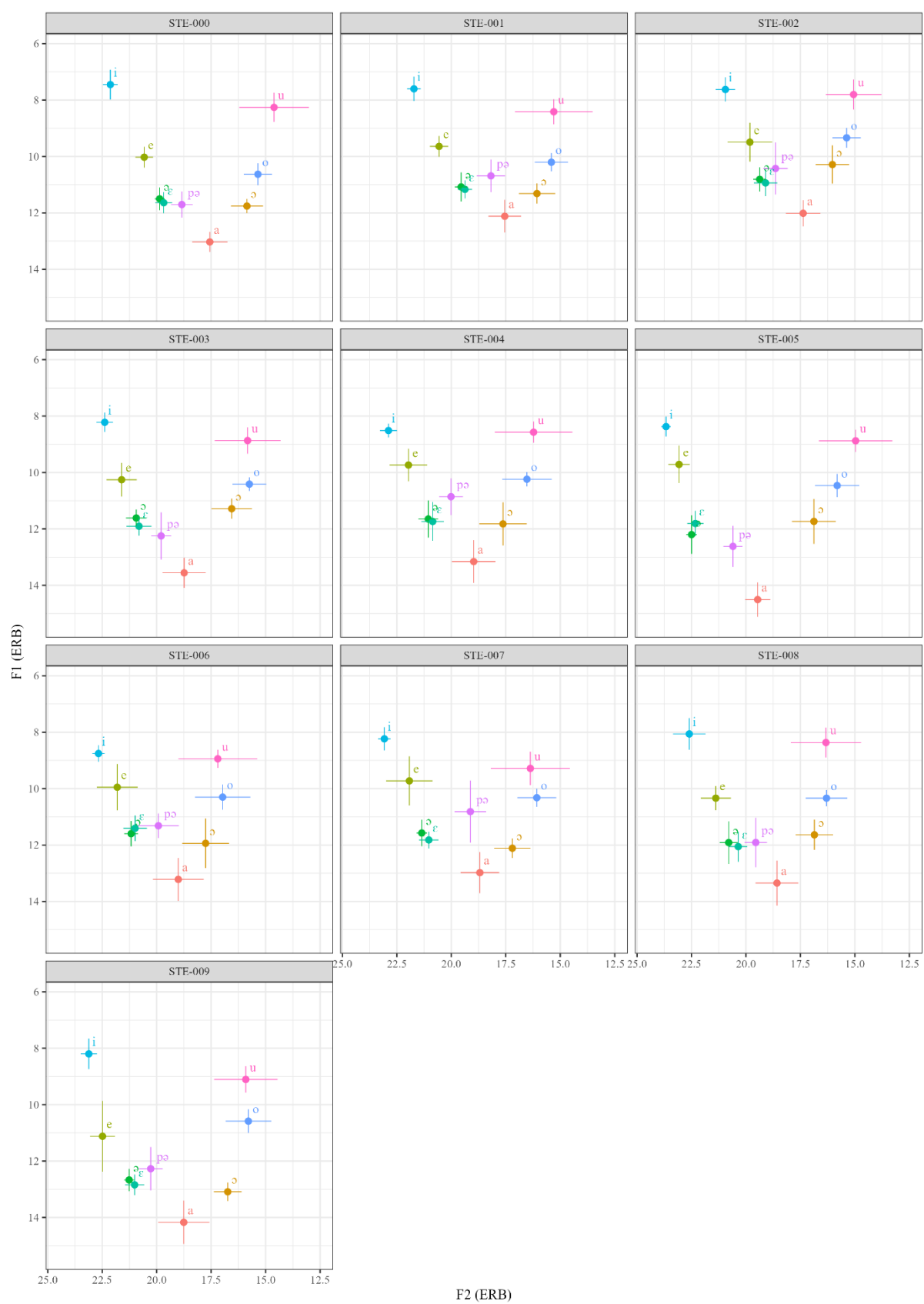
Pillai scores below the thresholds, as well as results below the 0.0025 p -value threshold are in bold.

| Speaker | ə ɛ | | | e ɛ | | | e ə | | |
|--------------|--------------------|--------------------|---------------------|-------------------|--------------------|------------|-------------------|--------------------|------------|
| | Pillai score | Pillai threshold | p -value | Pillai score | Pillai threshold | p -value | Pillai score | Pillai threshold | p -value |
| STE-000 | 0.10424018 | 0.0776652 | 0.025027523 | 0.8433655 | 0.08494631 | 2.78E-25 | 0.797783 | 0.08237218 | 1.36E-22 |
| STE-001 | 0.0812886 | 0.08363944 | 0.072202509 | 0.8626436 | 0.08363944 | 1.88E-27 | 0.7979081 | 0.08237218 | 1.33E-22 |
| STE-002 | 0.11754348 | 0.08237218 | 0.0194688158 | 0.6084585 | 0.07994947 | 5.82E-14 | 0.6513286 | 0.07994947 | 1.34E-15 |
| STE-003 | 0.18689865 | 0.08237218 | 0.0014776022 | 0.7850263 | 0.08363944 | 2.01E-21 | 0.7241924 | 0.08114274 | 1.26E-18 |
| STE-004 | 0.04347133 | 0.08114274 | 0.2411769134 | 0.7231613 | 0.07994947 | 7.45E-19 | 0.7185985 | 0.08114274 | 2.39E-18 |
| STE-005 | 0.14149 | 0.07879078 | 0.0065102374 | 0.8245381 | 0.08114274 | 6.51E-25 | 0.8265147 | 0.07994947 | 1.89E-25 |
| STE-006 | 0.10625813 | 0.08494631 | 0.0325064207 | 0.5887711 | 0.08629466 | 2.65E-12 | 0.6336518 | 0.08629466 | 8.26E-14 |
| STE-007 | 0.21244346 | 0.08114274 | 0.0004797509 | 0.7607166 | 0.07879078 | 3.19E-21 | 0.7013875 | 0.07994947 | 8.73E-18 |
| STE-008 | 0.22462989 | 0.08237218 | 0.0003307769 | 0.8197334 | 0.08912399 | 2.64E-22 | 0.6309554 | 0.08114274 | 1.40E-14 |
| STE-009 | 0.17091562 | 0.08363944 | 0.0029961309 | 0.7922941 | 0.08494631 | 1.52E-21 | 0.7775752 | 0.08363944 | 5.79E-21 |
| Means | 0.138917934 | 0.081808319 | | 0.76087085 | 0.083242261 | | 0.72598952 | 0.081795509 | |

Table F2: Pillai scores, Pillai thresholds and p -values for /o/-/ɔ/, /u/-/ɔ/ and /u/-/o/.

| Speaker | o ɔ | | | u ɔ | | | u o | | |
|--------------|-------------------|--------------------|------------|-------------------|-------------------|------------|-------------------|-------------------|------------|
| | Pillai score | Pillai threshold | p -value | Pillai score | Pillai threshold | p -value | Pillai score | Pillai threshold | p -value |
| STE-000 | 0.7547004 | 0.08629466 | 4.91E-19 | 0.9575656 | 0.08494631 | 1.40E-42 | 0.8946652 | 0.08363944 | 5.01E-31 |
| STE-001 | 0.7901416 | 0.07994947 | 9.17E-23 | 0.9333552 | 0.08237218 | 8.88E-38 | 0.8463284 | 0.08237218 | 2.39E-26 |
| STE-002 | 0.5130424 | 0.07879078 | 4.87E-11 | 0.8259179 | 0.07994947 | 2.11E-25 | 0.7574209 | 0.07879078 | 5.01E-21 |
| STE-003 | 0.6960707 | 0.07879078 | 8.54E-18 | 0.9034502 | 0.07879078 | 3.14E-34 | 0.8238692 | 0.07994947 | 3.09E-25 |
| STE-004 | 0.6962984 | 0.08114274 | 2.74E-17 | 0.888659 | 0.08114274 | 3.11E-31 | 0.8742564 | 0.08237218 | 4.30E-29 |
| STE-005 | 0.5470647 | 0.07994947 | 6.62E-12 | 0.8511507 | 0.08237218 | 8.74E-27 | 0.8004215 | 0.08237218 | 8.99E-23 |
| STE-006 | 0.6222425 | 0.08363944 | 7.83E-14 | 0.8506757 | 0.08237218 | 9.66E-27 | 0.7588885 | 0.08114274 | 1.70E-20 |
| STE-007 | 0.8917104 | 0.07994947 | 4.21E-32 | 0.9103003 | 0.08114274 | 3.09E-34 | 0.5569308 | 0.07879078 | 2.16E-12 |
| STE-008 | 0.6992352 | 0.08237218 | 3.67E-17 | 0.907671 | 0.0776652 | 2.18E-35 | 0.8407903 | 0.07994947 | 1.16E-26 |
| STE-009 | 0.9198599 | 0.08363944 | 1.05E-34 | 0.9627576 | 0.08363944 | 5.04E-45 | 0.7414477 | 0.08237218 | 3.13E-19 |
| Means | 0.71303662 | 0.081451843 | | 0.89915032 | 0.08143922 | | 0.78950189 | 0.08117514 | |

Appendix G: Vowel spaces by speaker



Appendix H: Vowel spaces by sex (females in light red, males in light blue)

Figure H1: Averaged values for each vowel by sex.

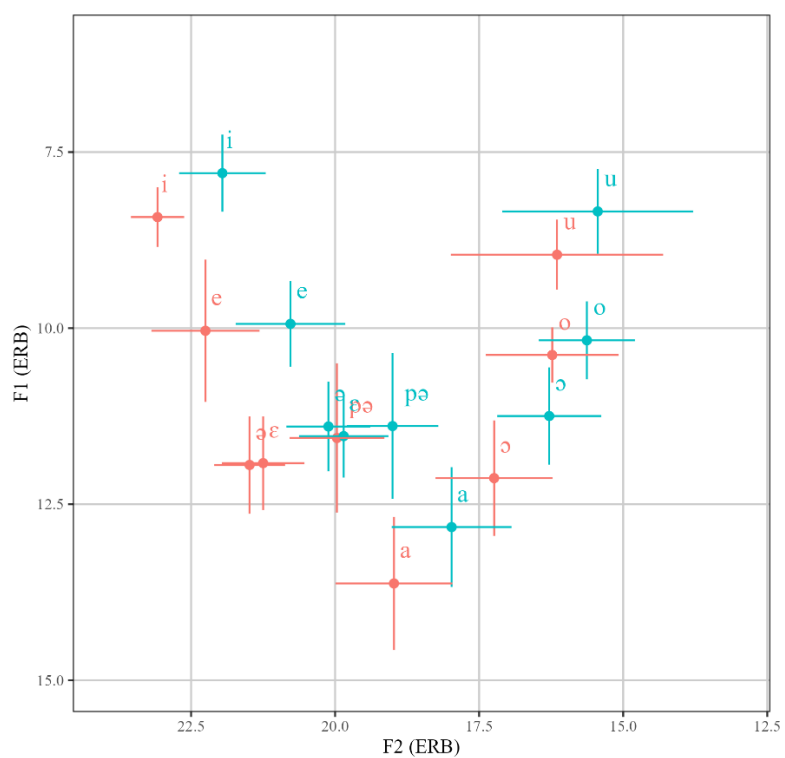
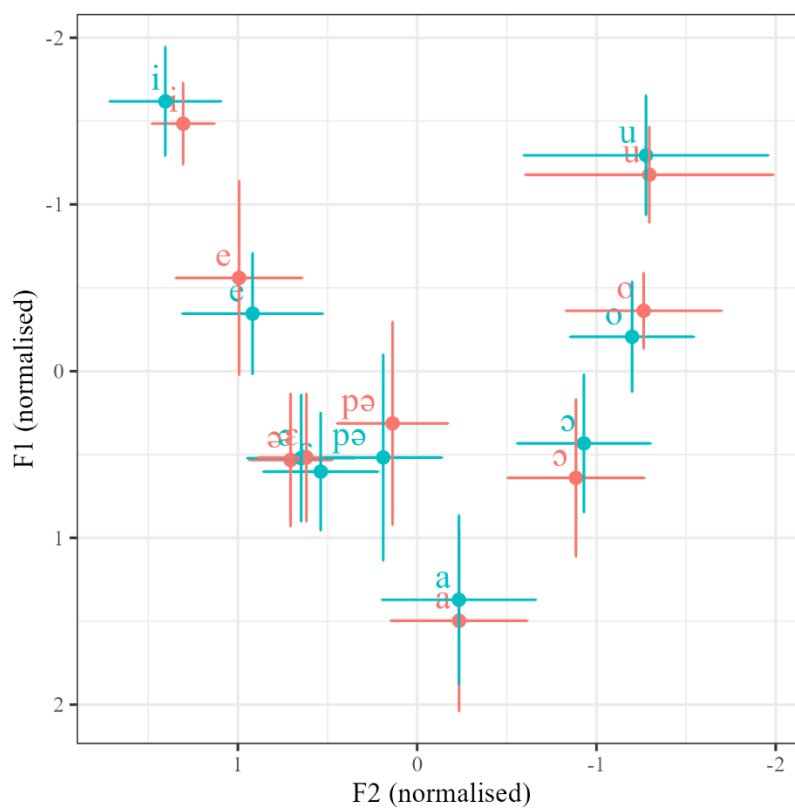
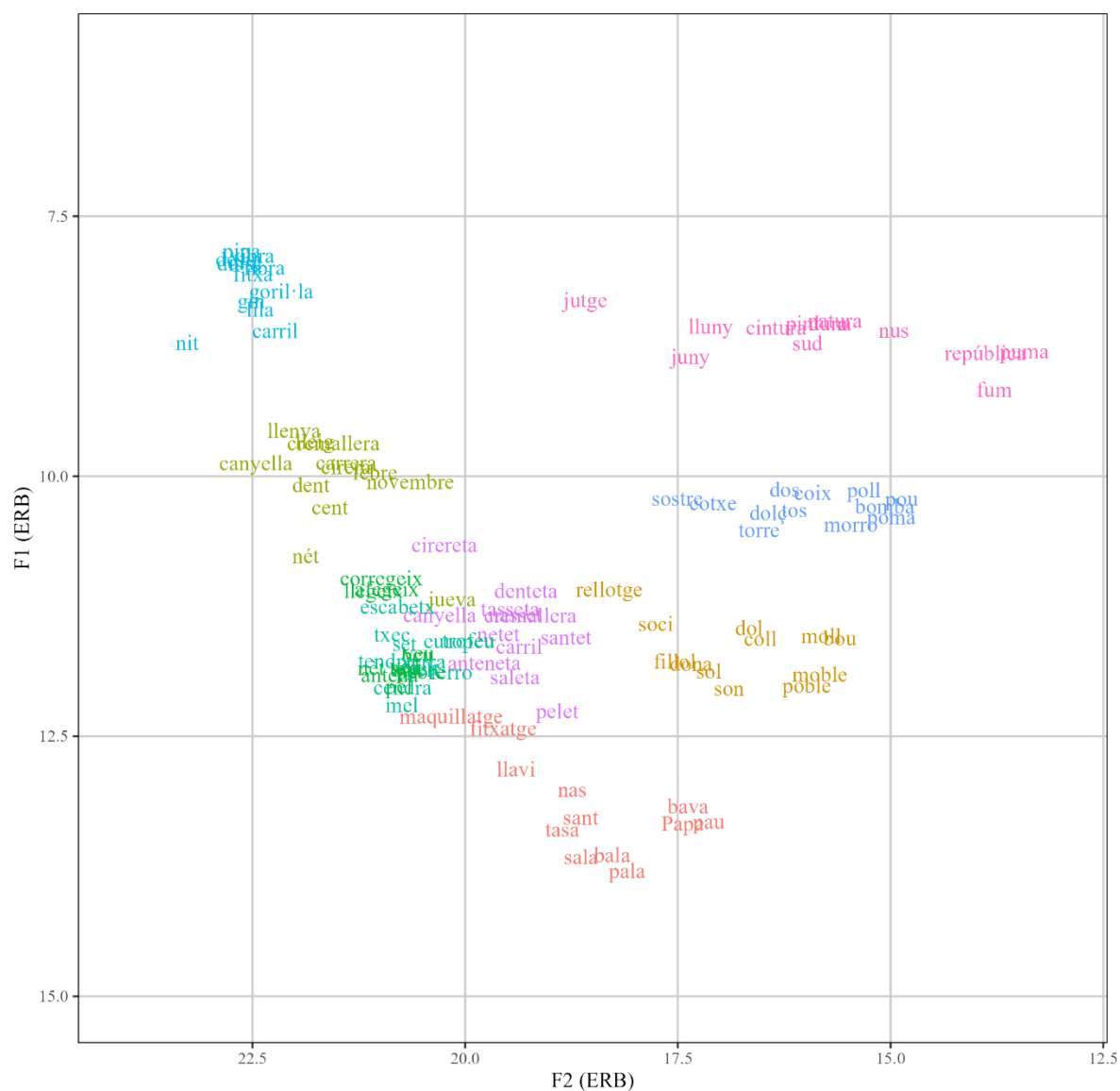


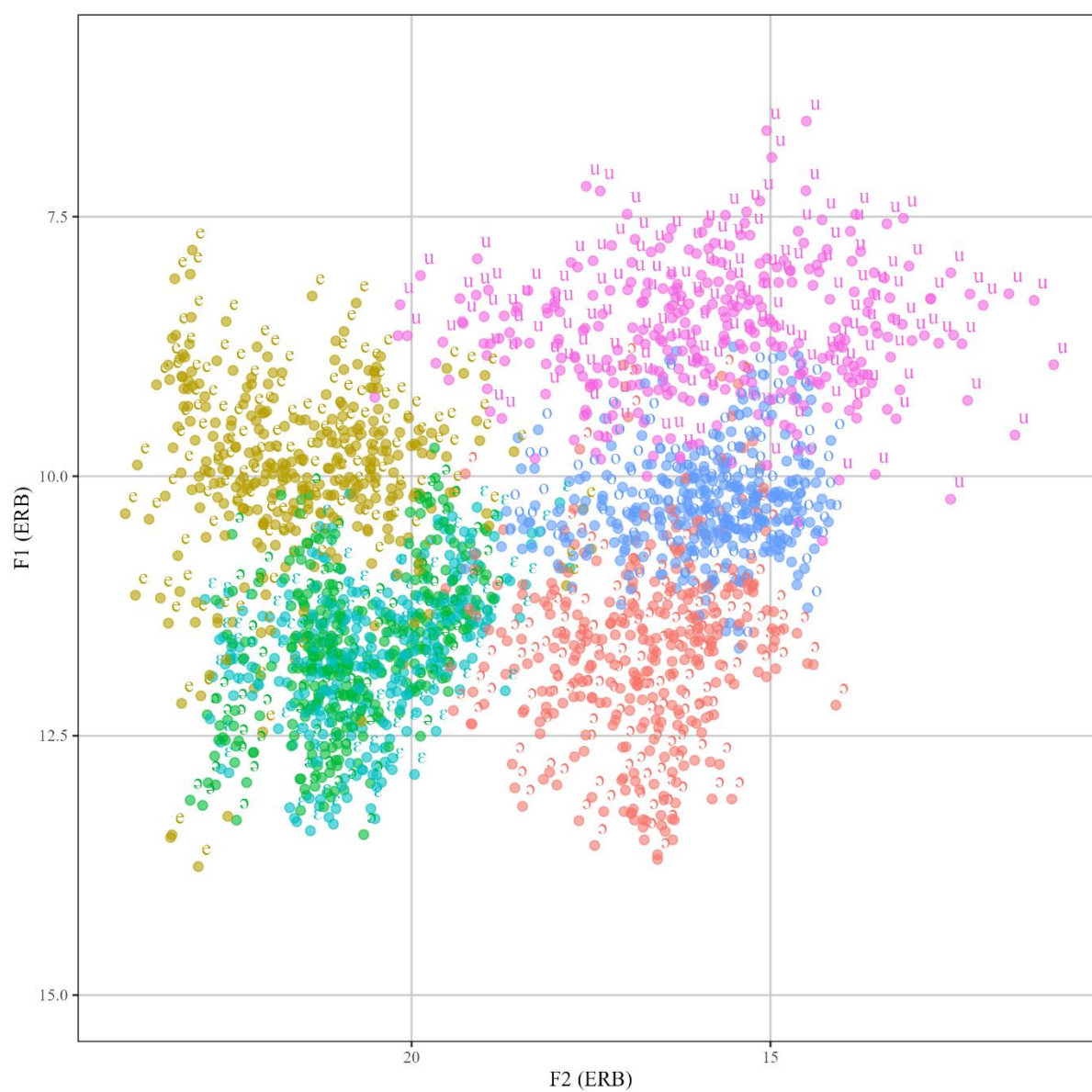
Figure H2: z-scores normalised values for each vowel by sex.



Appendix I: averaged word-specific realisations



Appendix J: All instances of /e/, /ɛ/, /ə/, /ɔ/, /o/, /u/



Appendix K: Vowel pair instances

Figure K1: front vowel instances (/e/ in black, /ɛ/ in red, /ə/ in blue)

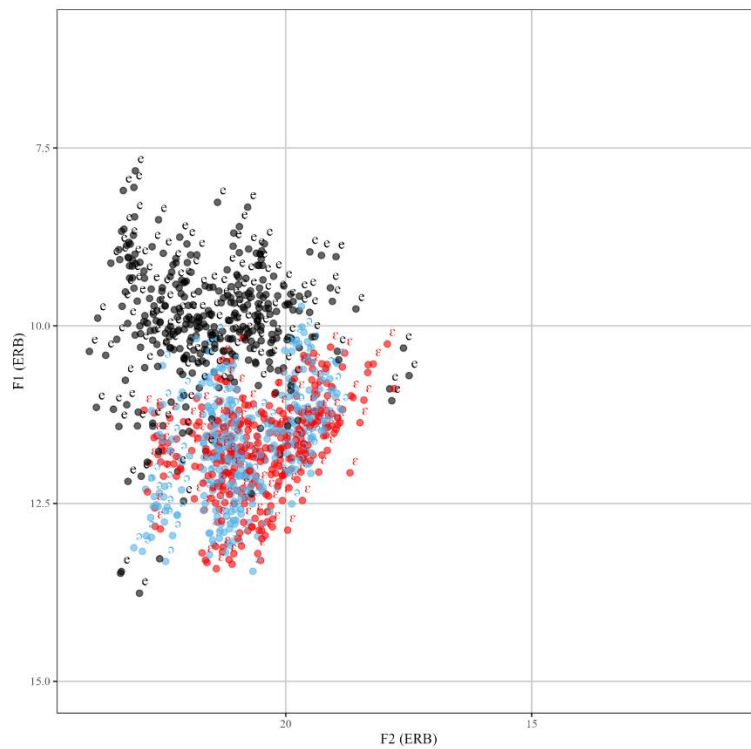
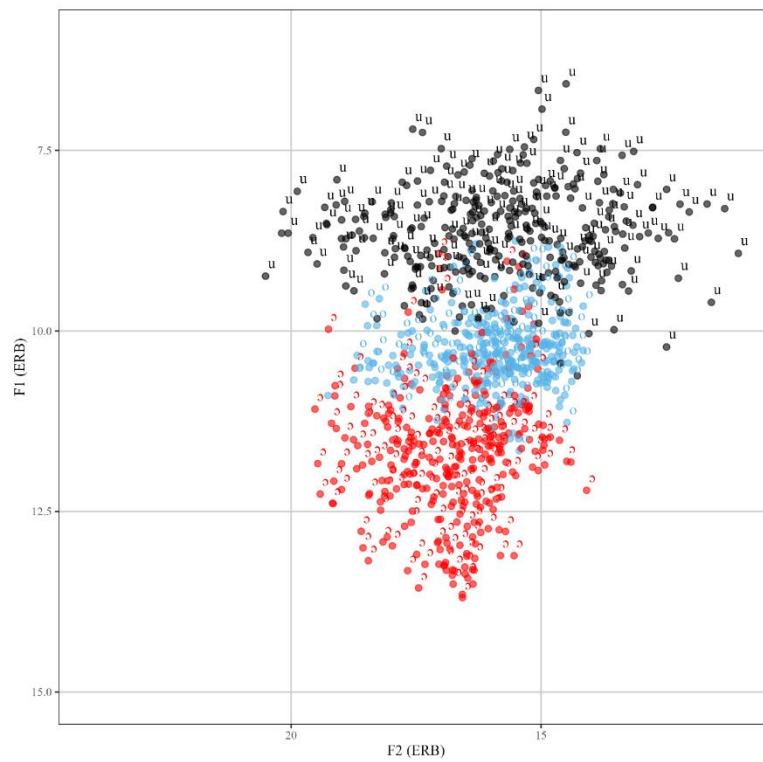


Figure K2: Back vowel instances (/u/ in black, /ʊ/ in red, /o/ in blue)



References

- Alcover, Antoni Maria. 1908. Una mica de dialectologia catalana. *Botlletí del Diccionari de la Llengua Catalana* IV. Palma. 126–211.
- Alves, Ubiratã & Brisolara, Lucienne. 2020. Listening to accented speech in Brazilian Portuguese: On the role of fricative voicing and vowel duration in the identification of /s/ – /z/ minimal pairs produced by speakers of L1 Spanish. *Journal of Portuguese Linguistics* 19(1). 6.
- Amengual, Mark. 2016. The perception and production of language-specific mid-vowel contrasts: shifting the focus to the bilingual individual in early language input conditions. *International Journal of Bilingualism*, 20(2). 133–152. [first published in 2014]
- Amengual, Mark, and Chamorro, Pilar. 2016. The effects of language dominance in the perception and production of the Galician mid vowel contrasts. *Phonetica* 72(4). 207–236.
- Bai, Justin & Scarborough, Rebecca. 2025. Neighborhood density and minimal pair effects on English pre-voicing vowel duration. Manuscript.
- Bates, Douglas & Maechler, Martin & Bolker, Ben & Walker, Steve 2015. Fitting Linear Mixed-Effects Models Using lme4. *Journal of Statistical Software* 67(1). 1–48.
- Beaulieu, Louise & Cichocki, Wladyslaw. 2002. Le concept de réseau social dans une communauté acadienne rurale. *Canadian Journal of Linguistics/Revue canadienne de linguistique* 47. 123–150.
- Birdsong, David, Libby M. Gertken & Mark Amengual. 2012. *Bilingual Language Profile: An Easy-to-Use Instrument to Assess Bilingualism*. COERLL, University of Texas at Austin. <https://sites.la.utexas.edu/bilingual/> (4 December 2024).
- Boersma, Paul & Weenink, David (2025). Praat: doing phonetics by computer [Computer program]. Version 6.4.31. <http://www.praat.org/> (3 May 2025).
- Eberhard, David M. & Simons, Gary F. & Fennig, Charles D. (eds.). 2025. *Ethnologue: Languages of the World*. Twenty-eighth edition. Dallas, Texas: SIL International. Online version: <https://www.ethnologue-com.proxy.uba.uva.nl/language/cat/> (21 May 2025)
- Ernestus, Mirjam, & Baayen, Harald. 2006. The functionality of incomplete neutralization in Dutch: the case of past-tense formation. *Laboratory Phonology* 8. 27–49.
- Fung, Roxana. S. Y. & Lee, Chris K. C. 2019. Tone mergers in Hong Kong Cantonese: An asymmetry of production and perception. *The Journal of the Acoustical Society of America* 146(5). EL424–EL430.

- Guasch, Marc & Boada, Roger & Ferré, Pilar, & Sánchez-Casas, Rosa. 2013. NIM: A Web-based Swiss Army knife to select stimuli for psycholinguistic studies. *Behavior Research Methods* 45. 765–771.
- Hamann, Silke & Torres-Tamarit, Francesc. 2023. Merger in Eivissan Catalan - an acoustic analysis of the vowel systems of young native speakers. *Phonetica* 80(1-2). 43–78.
- Islam, Md Jahurul., and Ahmed, Iftakhar. 2020. Mid-front and back vowel mergers in Mymensingh Bangla: An acoustic investigation. *Linguistics Journal* 14(1). 206–232.
- Joo, Hyoun-A & Schwarz, Lara, & Page, B. Richard. 2018. Nonconvergence and divergence in bilingual phonological and phonetic systems: Low back vowels in Moundridge Schweitzer German and English. *Journal of language contact* 11(2). 304–323.
- Kelley, Matthew C. & Tucker, Benjamin. V. 2020. A comparison of four vowel overlap measures. *The Journal of the Acoustical Society of America* 147(1). 137–145.
- Ketting, Thomas T. 2021. *Ha 'ina 'ia Mai Ana Ka Puana: The Vowels of 'Ōlelo Hawai'i*. Manoa: University of Hawai'i. (Doctoral dissertation.)
- Kuznetsova, Alexandra & Brockhoff, Per B. & Christensen, Rune H. B. 2017. lmerTest Package: Tests in Linear Mixed Effects Models. *Journal of Statistical Software* 82(13). 1-26.
- Ladefoged, Peter, & Johnson, Keith. 2010. *A Course in Phonetics*. Boston: Cengage Learning.
- Lippi-Green, Rosina L. 1989. Social network integration and language change in progress in a rural alpine village. *Language in Society*. 18. 213–234.
- Mora, Joan Carles & Nadeu, Marianna. 2012. L2 effects on the perception and production of a native vowel contrast in early bilinguals. *International Journal of Bilingualism* 16(4). 484–499.
- Nadeu, Marianna & Renwick, Margaret. E. L. 2016. Variation in the lexical distribution and implementation of phonetically similar phonemes in Catalan. *Journal of Phonetics* 58. 22–47.
- Navarra, Jordi & Sebastián-Gallés, Núria & Soto-Faraco, Salvador. 2005. The perception of second language sounds in early bilinguals: new evidence from an implicit measure. *Journal of Experimental Psychology: Human Perception and Performance* 31(5). 912–918.
- Pallier, Christophe & Bosch, Laura & Sebastián-Gallés, Núria. 1997. A limit on behavioral plasticity in speech perception. *Cognition* 64(3). B9–B17.
- Pichler, Heike & Wagner, Suzanne Evans, & Hesson, Ashley. (2018). Old-age language variation and change: Confronting variationist ageism. *Language and Linguistics Compass*, 12(6).

- R Core Team. 2025. R: A Language and Environment for Statistical Computing. R Foundation for Statistical Computing, Vienna, Austria. <https://www.R-project.org/>. (21 May 2025)
- Recasens, Daniel. 2019. Stressed /e/ centralization into schwa and related mid vowel developments in Catalan and elsewhere in the Romania. *Transactions of the Philological Society* 117. 294–316.
- Recasens, Daniel. 2023. *Consonant-induced sound changes in stressed vowels in Romance: Assimilatory, dissimilatory and diphthongization processes*. Berlin, Boston: De Gruyter.
- Recasens, Daniel & Aina Espinosa. 2006. Dispersion and variability of Catalan vowels. *Speech Communication* 48. 645–666.
- Recasens, Daniel & Aina Espinosa. 2009. Dispersion and variability in Catalan five and six peripheral vowel systems. *Speech Communication* 51. 240–258.
- Regan, Brendan. 2020. Extending Pillai scores to fricative mergers: Advancing a gradient analysis of a split-in-progress in Andalusian Spanish. *University of Pennsylvania Working Papers on Linguistics* 26(2). 109–118.
- Renwick, Margaret E. L. 2024. Robustness and Complexity in Italian Mid Vowel Contrasts. *Languages*, 9(4). 150.
- Ribeiro, Rodrigo S. 2017. *Duração de vogais tônicas antecedentes a consoantes plosivas no Português Brasileiro* (BA thesis at Universidade Federal do Rio Grande do Sul, Brazil).
- Sandoval, Steven & Berisha, Visar & Utianski, Rene L. & Liss, Julie M. & Spanias, Andreas. 2013. Automatic assessment of vowel space area. *The Journal of the Acoustical Society of America* 134.5_Supplement. EL477–483.
- Schmidt, Penelope & Diskin-Holdaway, Chloé & Loakes, Debbie. 2021. New insights into /el/-/æɪ/ merging in Australian English. *Australian Journal of Linguistics* 41(1). 66–95.
- Scrucca, Luca & Fraley, Chris & Murphy, T. Brendan & Raftery, Adrian E. 2023. *Model-Based Clustering, Classification, and Density Estimation Using mclust in R*. New York: Chapman and Hall/CRC.
- Sebastián-Gallés, Núria & Soto-Faraco, Salvador. 1999. Online processing of native and non-native phonemic contrasts in early bilinguals. *Cognition* 72. 111–123.
- Simonet, Miquel. 2011. Production of a Catalan-specific vowel contrast by early Catalan Spanish bilinguals. *Phonetica* 68(1-2). 88–110.
- Simonet, Miquel. 2019. Phonetic behavior in proficient bilinguals: Insights from the Catalan-Spanish contact situation. In Mark Gibson & Juana Gil (eds.), *Romance phonetics and phonology*, 395–406.

- Sloos, Marjoeleine. 2014. The reversal of the BÄREN-BEEREN merger in Austrian Standard German. *Mental Lexicon* 8(3). 353–371.
- Stehr, Daniel. 2018. A tutorial on recent methods for estimating working vowel space from connected speech. <https://rpubs.com/dstehr/autoVSA>. (11 June 2025)
- Toivonen, Ida & Blumenfeld, Lev & Gormley, Andrea & Hoiting, Leah & Logan, John & Ramlakhan, Nalini & Stone, Adam. 2015. Vowel height and duration. In *Proceedings of the 32nd west coast conference on formal linguistics* 32. 64-71. Somerville, MA: Cascadilla Proceedings Project.
- Torres Torres, Marià. 1983. Aspectes del vocalisme tònic eivissenc. *Eivissa* 14. 22–23.
- Tse, Holman. 2018. *Beyond the monolingual core and out into the wild: A variationist study of early bilingualism and sound change in Toronto heritage Cantonese*. Pittsburgh: University of Pittsburgh. (Doctoral dissertation.)
- Wickham, Hadley. 2011. The Split-Apply-Combine Strategy for Data Analysis. *Journal of Statistical Software* 40(1). 1-29.
- Wickham, Hadley & Averick, Mara & Bryan, Jennifer & Chang, Winston & McGowan Lucy D’Agostino & François, Romain & Grolemund, Garrett & Hayes, Alex & Henry, Lionel & Hester, Jim & Kuhn, Max & Pedersen, Thomas Lin & Miller, Evan & Bache, Stephan Milton & Müller, Kirill & Ooms, Jeroen & Robinson, David & Seidel, Dana Paige & Spinu Vitalie & Takahashi, Kohske & Vaughan, Davis & Wilke, Claus & Woo, Kara & Yutani, Hiroaki. 2019. Welcome to the tidyverse. *Journal of Open Source Software* 4(43). 1686.