

# THE EFFECTIVE BORROWING OF A PHONEMIC CONTRAST

Guðrún Duna Gylfadóttir

A DISSERTATION

in

Linguistics

Presented to the Faculties of the University of Pennsylvania

in

Partial Fulfillment of the Requirements for the Degree of Doctor of Philosophy

2018

Supervisor of Dissertation

---

Meredith Tamminga, Assistant Professor of Linguistics

Graduate Group Chairperson

---

Eugene Buckley, Associate Professor of Linguistics

Dissertation Committee:

David Embick, Professor of Linguistics

Donald Ringe, Professor of Linguistics

# THE EFFECTIVE BORROWING OF A PHONEMIC CONTRAST

COPYRIGHT

2018

Guðrún Duna Gylfadóttir

This work is licensed under the  
Creative Commons Attribution-  
NonCommercial-ShareAlike 4.0  
International License

To view a copy of this license, visit

<http://creativecommons.org/licenses/by-nc-sa/4.0/>

*Para la gente de Sevilla, cuyo arte espero que se me haya pegado, aunque sea un poquito.*

# Acknowledgments

The path to a dissertation is long, and many people have touched my journey along it.

Looking back, in a sense it began during my undergraduate study abroad at the University of Seville. There, I took a class with Juan Pablo Mora Gutiérrez, and we stayed in touch. Juan Pablo generously provided me with institutional support for this dissertation work, without which it would have been much more difficult to carry out. I am very grateful. During that same semester, I was fortunate to attend Pedro Carbonero Cano's class on the sociolinguistics of Seville Spanish, which first brought my attention to, among other things, the rich variation in anterior fricatives in Seville.

Several years later, I had the opportunity to attend the University of Pennsylvania for graduate school. Penn Linguistics has been an incredible learning environment, and in retrospect it was the ideal place to gain the understanding necessary to appreciate the linguistic variation I had learned about and observed in Seville. I'm grateful to Bill Labov, both for challenging and encouraging discussions as this project developed, and for providing an example of the kind of academic I want to be: one that finds the good in everyone's work and is genuinely interested in the people whose speech they study. This work and my academic perspective have both hugely benefited from his influence.

I want to thank Meredith Tamminga, who saw potential in this project from the beginning and has been an all-around great advisor. She patiently pushed me to flesh out my half-formed thoughts and connect them to the big picture. Her guidance has considerably broadened my outlook, methodologically and conceptually. I'm proud and grateful to have been her first advisee. I would also like to thank my committee members. Dave Embick

kept me on my toes, rapidly thinking things through to the next step and emphasizing the need to write to multiple audiences. My meetings with Don Ringe always left me newly impressed with the sheer interestingness of this topic, which came as an important reminder at certain junctures.

This work has benefited from comments and conversations with many people. This includes audience members at several conferences: LSA 2018, ICLAVE 9, and the *Investigando las hablas andaluzas* conference, among them Alfredo Herrero de Haro, Itxaso Rodríguez Ordóñez, Nicholas Henriksen, Lorenzo García Amaya, and Elena Jaime Jiménez. Thanks especially to Brendan Regan for his helpful comments and correspondence. I presented this work at various stages to the Penn community. Thanks to the attendees of Common Ground, particularly Joe Toscano and Dan Swingley, and to the members of the Variation and Cognition Lab for their comments. I've had the opportunity to discuss this project with a number of fellow Penn graduate students, including Yosiane White, Pleng Chanchaochai, Ava Creemers, Wei Lai, Betsy Sneller, Lacey Wade, Ruairidh Purse, Robert Wilder, and Amy Goodwin Davies. Their input was essential to this project. Betsy's and Amy's enthusiasm in my early results was very much appreciated. Finally, I'm grateful to Gillian Sankoff, Gareth Roberts, Robin Clark, Gene Buckley, Erez Levon, Molly Babel, Arthur Samuel, and Andries Coetzee for sharing thoughts on this work.

A number of people and groups aided in practical aspects of this project. Thank you to Marina Barrio Parra in the University of Seville Phonetics Laboratory for assistance with recording equipment, advice with recruitment, and for facilitating my extensive use of the laboratory space. Fernando G. Luna at the Universidad Nacional de Córdoba readily shared a semantic association corpus with me. Eric Wilbanks provided generous technical support in the use of his forced aligner. Special thanks go to Rafael Gilabert Errazquin, my *conejiillo de indias*, for testing my experiments, and to Eduardo Rodríguez Galán for sitting uncomplainingly through some truly stifling recording sessions in the September heat. I am indebted to the members of the Experimental Morphology Lab at Penn for sharing their technical knowledge with me. Akiva Bacovcin wrote and Robert Wilder and Amy Goodwin

Davies contributed to the scripts used in the running of experimental portions of this work. Robert additionally shared his scripts and expertise on processing lexical decision data. Robert and Amy also provided extensive written comments on multiple drafts. Thanks to Amy Forsyth at Penn for her administrative expertise, and to Gene Buckley, who has been both reliable and accommodating as graduate advisor. Finally, this work was supported by a Doctoral Dissertation Improvement Grant from the National Science Foundation.

Les quiero agradecer a mis participantes, quienes se abrieron a una extranjera haciéndoles preguntas raras y mostraban una gran paciencia con su portátil moribundo. Aquí tan sólo empiezo a sacarle el valor a las observaciones que me hicieron sobre su habla y su comunidad. A todos los investigadores no les tocan sujetos que les dejan al final de la sesión con palabras de ánimo y apoyo para su proyecto.

This project required a lot of non-linguistic support too. I am privileged to have enjoyed the friendship of a number of *sevillano/as*, both native and adopted. The lindy hop community of Seville welcomed me and gave me a dance outlet during my fieldwork trips. Thanks to Lola Rueda Montero for being sure that I would write this before I ever set foot in grad school. To Eduardo Rodríguez Galán and to Elena Jiménez Durán: your absurdly generous hospitality made this work financially possible. My house is forever open to you. Elena, I feel like I know your beloved city mainly because I have experienced it with you – on rollerblades, in *trajes de flamenca*, and on countless nights out. Thank you for sharing it with me.

Thanks to the members of Phuntax: Haitao Cai, Eric Doty, Einar Freyr Sigurðsson, Anton Karl Ingason, Sabriya Fisher, Robert Wilder, Betsy Sneller, and Amy Goodwin Davies. I can't imagine how different the first few years of grad school would have been without my amazing cohort. Thanks to fellow members of the Linguistics Lab and to Sue Sheehan for a wonderful experience of community. I'd also like to give a shoutout to IRCS IRCS. Outside of the academic sphere, the West Philly Swingers kept me hale and hearty through grad school. There is nothing like mandatory rehearsals to make sure you don't spend all of your time in front of a computer. Thank you to Office Nomads in Seattle, and

especially to my lunch pals Mike Morita, Kellen Piffner, Beth Howard, and Saul Pwanson, for providing the work community I needed as I was finishing up.

Many people provided needed words of encouragement: thank you to Soohyun Kwon, Sabriya Fisher, Betsy Sneller, Aaron Freeman, Einar Freyr Sigurðsson, Aaron Ecay, Andrea Ceolin, Ava Irani, Pleng Chanchaochai, Kajsa Djärv, Michael Chirico, Aaron Farance, Kinsey Miller, Desey Tziortzis, Chelsea Dole, Tanya Davis-Castro, and Claudia Tussey (a non-exhaustive list!). A very special thanks to Amy Goodwin Davies and Robert Wilder, my teammates through this entire journey, for endless advice and support, last-minute edits at many stages, and general sanity-keeping measures. Rob, #teamBaton forever.

My parents, Ingibjörg Sigrún Einisdóttir and Gylfi Ólafsson, are a constant source of love and support. They always believed I could do this, and that was a big reason why I could. *Takk*.

Finally, to my husband<sup>1</sup> Pau: though my grant has come to an end, I will always be a proud Pereira Fellow. The further along I got in this process, the more I leaned on you. So thank you: for literally getting me out of bed in the morning, for getting me through existential crises, and for patiently listening to me talk for five years about “S”.

---

<sup>1</sup>It still looks strange on paper!

## ABSTRACT

### THE EFFECTIVE BORROWING OF A PHONEMIC CONTRAST

Guðrún Duna Gylfadóttir

Meredith Tamminga

Language change often leads to the merger of phonemic categories, but the addition of new categories is rarely attested. In the Spanish city of Seville, as in other cities across the southern region of Andalusia, contact with the prestigious standard variety to the north is leading to the emergence of a contrast between /s/ and /θ/.

The 24 young adult speakers from Seville analyzed in this dissertation produce a mixture of [s] and the incoming sound [θ] in standard /θ/ contexts. However, when naturalistic production is examined individually, a systematic pattern is revealed: most of the speakers categorically limit [θ] to /θ/ contexts. They are also able to distinguish [θ] from [s] in low-level perception experiments. I argue that this is evidence for an underlying phonological contrast whose effective borrowing may have been made possible, as others have suggested, by the high level of social awareness it enjoys.

In a semantic priming experiment with stimuli produced by a local talker, the same speakers show a processing advantage for words with standard /θ/ that are produced with the local variant [s] over words produced with [θ]. I hypothesize that this seemingly incongruent result may be evidence for talker-based phonemic flexibility in perception.

This dissertation underscores the complexity of the notion of phonemic contrast and motivates the investigation of more ongoing changes involving consonants. In general, it highlights the potential benefits of combining naturalistic and experimental approaches to variable phenomena.



# Contents

|                                                                           |             |
|---------------------------------------------------------------------------|-------------|
| <b>List of Tables</b>                                                     | <b>xii</b>  |
| <b>List of Figures</b>                                                    | <b>xiii</b> |
| <b>1 Introduction</b>                                                     | <b>1</b>    |
| 1.1 Mergers and splits . . . . .                                          | 2           |
| 1.1.1 Social pressure and merger reversal . . . . .                       | 5           |
| 1.1.2 Merger as an individual property . . . . .                          | 7           |
| 1.2 Decomposing phonemic contrast . . . . .                               | 8           |
| 1.2.1 The notion of variable merger . . . . .                             | 8           |
| 1.2.2 Perception . . . . .                                                | 9           |
| 1.3 This study . . . . .                                                  | 10          |
| 1.3.1 A note on methodology . . . . .                                     | 11          |
|                                                                           | <b>12</b>   |
| <b>2 Background: <i>Seseo</i> in Seville</b>                              | <b>13</b>   |
| 2.1 Introduction . . . . .                                                | 13          |
| 2.2 Historical context . . . . .                                          | 13          |
| 2.2.1 Castilian in Seville . . . . .                                      | 13          |
| 2.2.2 The emergence of an Andalusian Spanish . . . . .                    | 16          |
| 2.2.3 Historical development of the Spanish anterior fricatives . . . . . | 16          |
| 2.2.4 Andalusian Spanish and the Americas . . . . .                       | 19          |
| 2.2.5 The modern Andalusian anterior fricatives . . . . .                 | 20          |
| 2.3 Modern context . . . . .                                              | 22          |
| 2.3.1 Divergent norms in the east and west . . . . .                      | 23          |
| 2.3.2 The rise of <i>distinción</i> . . . . .                             | 23          |
| 2.3.3 Previous data from Seville . . . . .                                | 25          |
| 2.3.4 Linguistic attitudes and awareness . . . . .                        | 26          |
| 2.3.5 The acoustics of the /s/-/θ/ contrast . . . . .                     | 27          |
| 2.4 The value of examining the individual . . . . .                       | 28          |
|                                                                           | <b>29</b>   |

|          |                                                                          |           |
|----------|--------------------------------------------------------------------------|-----------|
| <b>3</b> | <b>Demographics and methods</b>                                          | <b>30</b> |
| 3.1      | Recruitment . . . . .                                                    | 30        |
| 3.2      | Study methods . . . . .                                                  | 31        |
| 3.2.1    | General procedures . . . . .                                             | 31        |
| 3.2.2    | Sentence completion task . . . . .                                       | 32        |
| 3.2.3    | Sociolinguistic interviews . . . . .                                     | 32        |
| 3.3      | Information about the participants . . . . .                             | 34        |
| 3.3.1    | Demographic and identity questionnaires . . . . .                        | 34        |
| 3.3.2    | Evidence of awareness . . . . .                                          | 34        |
| 3.3.3    | Attitudes towards <i>seseo</i> . . . . .                                 | 35        |
|          |                                                                          | <b>36</b> |
| <b>4</b> | <b>Production and perception</b>                                         | <b>37</b> |
| 4.1      | Introduction . . . . .                                                   | 37        |
| 4.1.1    | Production and perception of a contrast at different levels . . . . .    | 38        |
| 4.1.2    | Lexical considerations . . . . .                                         | 39        |
| 4.1.3    | The role of orthography . . . . .                                        | 40        |
| 4.1.4    | Consonants versus vowels . . . . .                                       | 40        |
| 4.2      | Social predictions . . . . .                                             | 41        |
| 4.3      | Production analysis . . . . .                                            | 42        |
| 4.3.1    | Preparation of data for analysis . . . . .                               | 42        |
| 4.3.2    | Auditory coding . . . . .                                                | 43        |
| 4.3.3    | Coding results . . . . .                                                 | 44        |
| 4.3.4    | Lexical distribution . . . . .                                           | 53        |
| 4.3.5    | Acoustic analysis . . . . .                                              | 56        |
| 4.3.6    | Acoustic results . . . . .                                               | 57        |
| 4.3.7    | Speaker awareness of <i>seseo</i> . . . . .                              | 59        |
| 4.3.8    | Discussion . . . . .                                                     | 60        |
| 4.4      | Identification and discrimination . . . . .                              | 61        |
| 4.4.1    | Near merger . . . . .                                                    | 61        |
| 4.4.2    | Mergers and methodology . . . . .                                        | 64        |
| 4.4.3    | Materials and procedure . . . . .                                        | 65        |
| 4.4.4    | Results . . . . .                                                        | 67        |
| 4.4.5    | Discussion . . . . .                                                     | 69        |
| 4.5      | General discussion . . . . .                                             | 71        |
| 4.6      | Conclusion . . . . .                                                     | 73        |
|          |                                                                          | <b>74</b> |
| <b>5</b> | <b>Mental representation</b>                                             | <b>75</b> |
| 5.1      | Introduction . . . . .                                                   | 75        |
| 5.2      | Decomposing the perception of a contrast . . . . .                       | 78        |
| 5.2.1    | Investigating mental representations with experimental methods . . . . . | 79        |
| 5.2.2    | Mental representations and merger mechanisms . . . . .                   | 82        |
| 5.2.3    | Methodological considerations . . . . .                                  | 84        |

|          |                                   |            |
|----------|-----------------------------------|------------|
| 5.3      | Methods . . . . .                 | 90         |
| 5.3.1    | Design . . . . .                  | 90         |
| 5.3.2    | Stimuli . . . . .                 | 91         |
| 5.3.3    | Talker . . . . .                  | 92         |
| 5.3.4    | Procedure . . . . .               | 93         |
| 5.4      | Results . . . . .                 | 93         |
| 5.5      | Discussion . . . . .              | 97         |
| 5.6      | Conclusion . . . . .              | 100        |
|          |                                   | <b>100</b> |
| <b>6</b> | <b>Conclusions and extensions</b> | <b>101</b> |
|          |                                   | <b>104</b> |
|          | <b>Bibliography</b>               | <b>116</b> |

# List of Tables

|     |                                                                                                                                            |    |
|-----|--------------------------------------------------------------------------------------------------------------------------------------------|----|
| 2.1 | The medieval Spanish anterior fricatives (examples from Sánchez Méndez 2003). . . . .                                                      | 16 |
| 2.2 | The Castilian anterior fricatives by 1650. . . . .                                                                                         | 17 |
| 2.3 | The single Andalusian anterior fricative by the late 17th century. . . . .                                                                 | 17 |
| 3.1 | The target words in the elicitation task. . . . .                                                                                          | 32 |
| 4.1 | Number and percent of segments in the overall data given each auditory coding label, separated by phonemic (etymological) context. . . . . | 44 |
| 4.2 | The number of speakers in this study by gender and education level. . . . .                                                                | 45 |
| 4.3 | A linear mixed-effects logistic regression model predicting realization in Z-contexts (standard realization = 1). . . . .                  | 49 |
| 4.4 | A linear mixed-effects logistic regression model predicting realization in S-contexts (standard realization = 1). . . . .                  | 49 |
| 4.5 | Percentage [s] in fricative steps; following Eisner and McQueen (2005). . . . .                                                            | 66 |
| 4.6 | A linear mixed-effects logistic regression model predicting response (“S” = 1, “Z” = 0). . . . .                                           | 68 |
| 4.7 | A linear mixed-effects logistic regression model predicting response (“Different” = 1, “Same” = 0). . . . .                                | 69 |
| 5.1 | Example critical stimuli from Experiment 1. . . . .                                                                                        | 90 |
| 5.2 | Example filler stimuli from Experiment 1 . . . . .                                                                                         | 93 |
| 5.3 | A linear mixed-effects regression model predicting the log of reaction time. . . . .                                                       | 94 |
| 5.4 | A linear mixed-effects regression model predicting the log of reaction time. . . . .                                                       | 94 |

# List of Figures

|      |                                                                                                                                                                                                                                  |    |
|------|----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|----|
| 2.1  | Map of modern Spain. . . . .                                                                                                                                                                                                     | 14 |
| 2.2  | Changes in the Spanish anterior fricatives taking place in Andalusia from the 15th to the 17th century, adapted from Penny (2000):90. . . . .                                                                                    | 18 |
| 2.3  | Changes in the Spanish anterior fricatives taking place in Castile from the 15th to the 17th century, adapted from Penny (2000):88. . . . .                                                                                      | 18 |
| 2.4  | Map showing the extension of <i>seseo</i> and <i>ceceo</i> in Andalusia, from Penny 2000, adapted from data from the ALEA dialect survey (Alvar 1973). . . . .                                                                   | 20 |
| 3.1  | Subjects' ratings of the acceptability of <i>ceceo</i> and <i>seseo</i> on a 7-point scale. .                                                                                                                                    | 36 |
| 4.1  | Percent of segments in the overall data given each auditory coding label, separated by phonemic (etymological) context. . . . .                                                                                                  | 45 |
| 4.2  | Percent segments in the spontaneous speech data coded [θ] in Z-contexts by phonemic context, faceted by the gender of the speaker. . . . .                                                                                       | 46 |
| 4.3  | Percent segments in the elicited data coded [θ] in Z-contexts in the elicited portion by etymological context, faceted by the gender of the speaker. . . .                                                                       | 47 |
| 4.4  | Percent segments in the spontaneous speech data coded [θ] in Z-contexts by phonemic context, faceted by the education level of the speaker. . . . .                                                                              | 47 |
| 4.5  | Percent segments in the elicited data coded [θ] in Z-contexts in the elicited portion by etymological context, faceted by the education level of the speaker. .                                                                  | 48 |
| 4.6  | Percent tokens coded [θ] in Z-contexts overall by speaker, with [x] realizations excluded and the percent of intermediate tokens shown in a transparent gray. . . . .                                                            | 50 |
| 4.7  | Percent elicited tokens coded [θ] in Z-contexts by speaker, with [x] realizations excluded and the percent of realizations as intermediate shown in a transparent gray. . . . .                                                  | 51 |
| 4.8  | Percent tokens coded [s] in S-contexts overall by speaker, with [x] realizations excluded and the percent of intermediate tokens shown in a transparent gray. . . . .                                                            | 52 |
| 4.9  | Percent elicited tokens coded [s] in S-contexts by speaker, with [x] realizations excluded and the percent of intermediate tokens shown in a transparent gray. . . . .                                                           | 53 |
| 4.10 | Percent tokens coded [θ] by speaker of the five most frequent lexical items containing /θ/: <i>dice</i> 'say', <i>entonces</i> 'then', <i>hace</i> 'do', <i>hacer</i> 'to do', <i>veces</i> 'times' (as in 'sometimes'). . . . . | 54 |
| 4.11 | Percent [θ] in Z-contexts for each lexical item (each word is represented by one point) plotted by frequency in the Subtlex-ESP corpus, shown with a fitted loess line. . . . .                                                  | 55 |

|      |                                                                                                                                                                                 |    |
|------|---------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|----|
| 4.12 | Percent [s] in S-contexts for each lexical item (each word is represented by one point) plotted by frequency in the Subtlex-ESP corpus, shown with a fitted loess line. . . . . | 55 |
| 4.13 | Tokens plotted by center of gravity value and mean intensity and colored by auditorily coded value.. . . .                                                                      | 58 |
| 4.14 | Tokens plotted by the first two principal components and colored by auditorily coded value. . . . .                                                                             | 58 |
| 4.15 | Proportion segments coded [θ] in /θ/-contexts by self-reported <i>seseo</i> rate. .                                                                                             | 60 |
| 4.16 | Percent “S” responses overall by continuum step. . . . .                                                                                                                        | 67 |
| 4.17 | Percent “S” responses by continuum step by subject. . . . .                                                                                                                     | 68 |
| 4.18 | Percent “Different” responses to Different pairs. . . . .                                                                                                                       | 70 |
| 4.19 | Percent “Different” responses to Same pairs. . . . .                                                                                                                            | 70 |
| 4.20 | Percent “Different” responses to Different pairs. . . . .                                                                                                                       | 71 |
| 5.1  | Reaction times by condition. . . . .                                                                                                                                            | 95 |
| 5.2  | Mean difference of reaction times by condition. . . . .                                                                                                                         | 95 |
| 5.3  | Mean difference of reaction time (unrelated - related) by experimental condition and by subject. Subjects are ordered from left to right by rate of [θ] in Z-contexts. . . . .  | 96 |

# Chapter 1

## Introduction

This dissertation seeks to contribute to our understanding of two issues in phonology. One involves what Weinreich et al. (1968) call the “Transition Problem”: it is easy to look at a language variety over time and observe that Structure A became Structure B, but it is difficult to observe the intervening stage. There has been some success in observing this intermediate stage of phonemic mergers, a common type of structural change in which two phonemes collapse into one. However, changes that involve the division of one phonological category into two are much less attested, particularly those that involve the addition of a phonemic category across all contexts. Their rarity means that they are not well understood, and there is much that is not yet known about how they progress. An ongoing change in the south of Spain, the emergence of a contrast between /s/ and /θ/, presents the opportunity to observe the intervening stage of a change that results in an extra phoneme. What kind of phonemic system does a speaker in the midst of such a change have?

The second issue is a theoretical one that relates to the phonemic system of an individual: namely, what production and perception behaviors accompany phonemic contrast? This is a straightforward question to answer for stable, established contrasts, but in cases of dialect contact and ongoing change it becomes less tractable and more interesting, challenging our understanding of what a phoneme is.

With these goals in mind, I present an analysis of the production and perception of

/s/ and /θ/ in a sample of 24 young speakers native to the city of Seville, where the local variety of Spanish has gone from showing a majority pattern of a single anterior fricative phoneme [s] (attested in the 1980s: Carbonero Cano 1982) to a majority pattern of contrast between /s/ and /θ/ (Santana 2016b). The following sections of this chapter will give some background, motivate the choice of this change as an object of study, and present some methodological considerations.

## 1.1 Mergers and splits

Mergers are a special kind of change in that they result in both a change to the size of the phonemic inventory and to the addition of previously non-existing homophones. Despite the potential confusions they create, however, the general observation has been that “mergers spread at the expense of distinctions” (HERZOG’S PRINCIPLE, Labov 1994:313). This is based on the observation of a great number of phonemic mergers by historical linguists, but also on contemporary mergers that have been documented as they progress.

Another type of change that affects the phonemic inventory involves the addition of a category, rather than the subtraction. This has been commonly known to occur with phonemes that contrast only in some contexts in a given language. Loanwords from another language containing the phoneme in the non-contrastive context, when enough speakers are able and motivated to imitate it, can lead to the de-neutralization of the phoneme in the context. This is what led to the emergence of /ʒ/ in Modern English from French borrowings, and to the loss of a voicing rule after nasals in Zoque due to Spanish loanwords (Ringe and Eska 2013:62). New allophonic contrasts can arise spontaneously over time as phonetic processes like assimilation are mistaken for deliberate phonological gestures (Ohala 1981).

The change that is the subject of this dissertation is not contextually limited in the way these examples are. Rather, it appears to be a wholesale addition of a phoneme to the inventory, and with it, a new lexical class. This type of change is very rarely attested in either the historical or the sociolinguistic literatures, leading scholars to venture that



they are not possible under normal circumstances: Labov (1994) cites Paul Garde as saying that “innovations can create mergers, but cannot reverse them” (p. 311). He sums this up as GARDE’S PRINCIPLE, which states that “mergers are irreversible by linguistic means”. Labov (1994) points out that the difficulty of reversing merger is derivable from the inherent arbitrariness of language: specifically, that in order for a merger reversal to take place, speakers have to learn to assign words to new lexical classes (and unlearn old associations along the way). He emphasizes that the *impossibility* of merger reversal posited by Garde’s principle rests on empirical observation alone; that is, counterexamples would call its validity into question.

The historical record contains a handful of reports that would seem to do that: for example, the reversals of the mergers of /i/ and /ei/ and of /ai/ and /oi/, respectively, in English. However, Labov (1994) makes a case that these categories were never truly merged. He bases this on a set of contemporary empirical observations of the behaviors of speakers in areas with an ongoing merger, which have yielded examples of what he calls “near-merger”. This term refers to a linguistic situation in which speakers produce a small but reliable acoustic contrast that they themselves are not aware of. This remarkable phenomenon was first documented in a few individual cases in areas of the United States where a vowel merger was underway. Labov, Yaeger, et al. (1972) report the case of a man named Bill Peters in Pennsylvania, where the merger of /ɑ/ and /ɔ/ was ongoing. This speaker declared relevant minimal pairs to be the same when asked in a procedure known as a “minimal pair task”, but nevertheless produced them differently. This finding was a surprise to linguists, including to the researchers themselves. The discovery of the possibility of near-merger continues to inspire some amount of skepticism (see e.g., Hickey 2004). However, near-merger has since been documented in association with a variety of mergers, in large groups of individuals, and using rigorous experimental methods (Janson and Schulman 1983; Di Paolo and Faber 1990; Di Paolo and Faber 1990; Herold 1990; Yu 2007; Johnson 2010; Bullock and Nichols 2017). Labov (1994)’s argument is that the linguists documenting a merger preceding an apparent reversal often simply mistake near-merger for true merger, whether because they

lack the relevant contrast themselves or because the acoustic contrast is so slight as to be barely perceptible. He puts forth this explanation for the two apparent historical English merger reversals mentioned above.

Near-merger is not a convincing account for all challenges to Garde’s and Herzog’s principles. Trudgill (2003) discusses the reversal of the merger of /w/ and /v/ in English, which seem to have merged and demerged in the southeast of England. They report that it is still remembered as a stereotypical feature in Norwich by older speakers. Based on its synchronic status as a true merger in a large number of colonial Englishes, they argue for its having been a true merger.<sup>2</sup> However, the explanation that Trudgill (2003) puts forth is one of dialect contact. He observes that several of the colonial examples of the merger were restricted to subparts of the population, and hypothesizes that ongoing contact between individuals with and without the merger allowed for merger reversal in both the colonial cases and in Norwich.

A reported merger reversal taking place closer to the present-day is the disappearance of the merger of the vowels in NEAR and SQUARE<sup>3</sup>, a traditional feature of the dialect of Charleston, South Carolina that receded in the early 20th century. Baranowski (2007) notes that the reversal seems to predate the post-WWII migratory wave that would otherwise provide a compelling language contact explanation, giving some weight to the possibility an internal mechanism. However, Baranowski (2007):122 ultimately concludes that “there is not enough evidence to exclude the possibility of a near-merger in the traditional dialect”.

Even more recently, a move toward contrast between /s/ and /θ/ seems to have been taking place across the south of Spain (Villena-Ponsoda 2001; Carbonero Cano 2003; García-Amaya 2008; Lasarte Cervantes 2010; Moya Corral and Wiedemann 1995; Hernández-Campoy and Villena-Ponsoda 2009; Santana 2016b; Regan 2017b; Ruiz Sanchez 2017, *inter alia*) This case has been put forth as an exception to Garde’s Principle (Villena-Ponsoda

---

<sup>2</sup>They further note its ongoing reversal in the British island of Anguilla and its apparent historic demerger in an area of coastal North Carolina (Wolfram and E. Thomas 2008).

<sup>3</sup>Although generally I will refer to phonemic contrasts using IPA symbols, in cases of conditioned merger where there is a historic tradition of referring to the change with words that represent the relevant lexical sets (Wells 1982), I follow this convention.

2003; Villena-Ponsoda 2008; Regan 2017b). The contrast is characteristic of the prestigious Castilian Spanish variety to the north, which has minimal pairs like CASA [kasa] ‘house’ and CAZA [kaθa] ‘hunt’. It steadily gained ground in the region, particularly in urban areas of east, relegating the traditional system of a single phoneme (whose phonetic character varies) to the non-urban eastern coast and to the west. Newer evidence shows that these last places are moving toward the contrast as well (Regan 2017b; García-Amaya 2008; Santana 2016b). This is the change this dissertation is concerned with: specifically, with its progression in the southwestern city of Seville (*Sevilla* in Spanish). It is not precisely a reversal of two phonemes that merged in the past; as we will see, its history is more complex. Although the variation in /s/ and /θ/ has a rich history of study in Andalusian dialectology and sociolinguistics (see e.g., Navarro Tomás et al. 1933; Alvar 1973; Salvador 1980; Villena-Ponsoda and Santos Requena 1996; Navarro Tomás 1996), it has not received attention from the broader field (particularly English-speaking parts of it) until recently (García-Amaya 2008; Regan 2017a; Regan 2017b).

### 1.1.1 Social pressure and merger reversal

It is worth emphasizing that Garde’s principle does not assert that merger reversals are impossible in general, but only that they are not purely *linguistic* innovations. Maguire et al. (2013):234 point out that this is somewhat “blurry concept”, and it seems to leave another path for merger reversals: that of contact with another language variety that possesses a contrast. Labov (1994) contends that a very particular set of circumstances are required for this to take place: “for social pressures to be brought to bear on the reversal of mergers, there must be an overt campaign to bring the problem to social attention and bestow prestige on the distinction” (p. 342). As an example, he discusses northern English schoolboys that would reverse their merger of /ʌ/ and /ʊ/ during their stays at boarding schools that emphasized the Received Pronunciation variety, in which the vowels contrast (Wyld 1936). Labov sees this kind of social pressure placed on merger as requiring a rare set of circumstances. This is related to another observation: namely, that mergers and splits

themselves do not tend to attract social attention. In the communities involved in the many mergers Labov and colleagues have observed over the years, speakers tend to limit their overt commentary to specific phonetic instances of words, and do not comment directly on the homophony or contrast of words pairs. Thus, the raised allophone in the conditioned split of /æ/ in Philadelphia (see Labov 2001) is referred to as “harsh”, outsiders remark on the presence of [ɪ] in *pen* in Southerners’ speech, and *coffee* pronounced with [ɔ] is a New York stereotype. At the same time, they observe that laypeople fail to remark on the homophony of *cot* and *caught* and of *pin* and *pen*, and on the differentiation of the vowels in *mad* and *sad* in Philadelphia. According to Labov (1994) “the evidence for the absence of social affect of splits and mergers is massive and overwhelming”, but at the same time, “almost entirely negative” (p. 343). That is to say, it is an empirical question whether it is possible for mergers to attract explicit social attention. Why shouldn’t they do so? Eckert and Labov (2017) put forward two explanations: one is that detecting phonemic differences involves “paratactic examination of the phonemic inventory... [which does] not form a part of ordinary discourse” (Eckert and Labov 2017). The other is that many of the mergers that have been studied are conditioned mergers<sup>4</sup> and for these reason do not occur very frequently. The merger of /ɑ/ and /ɔ/ is the exception, but even it does not attract the kind of the direct social attention that Labov views as necessary for demerger to occur.

There is evidence that social awareness of merger may not be as exceptional as Labov (1994) would lead us to believe. Baranowski (2013) describes several speakers that demonstrate acute awareness of their own PIN-PEN merger. For example, David B., an 82-year-old speaker from Charleston asserts: “In Charleston you’ll get a difference between pin and pen. In Columbia [South Carolina] you might not. In Charleston it would be different. Pin and pen. And the same thing for him and hem.” (Baranowski 2013:287). In Utah, Di Paolo 1992 interprets the results of her matched-guise task involving the /ɑ/-/ɔ/ contrast as evidence of social evaluation of the merger there.

---

<sup>4</sup>North America by itself boasts 15 vowel documented mergers, and all but one of them (/ɑ/ and /ɔ/) are phonologically conditioned (many before /r/): see Eckert and Labov 2017 for a list.

### 1.1.2 Merger as an individual property

Merger, in contrast to split, takes place without any sort of social stigmatization or even awareness. Yang (2009) argues that a sufficient level of merged input (whether through migration/contact or merger progression) is enough to lead to community-level merger, the idea being that the merged grammar (which relies on contextual cues rather than phonetic value to interpret lexical items) is, at a specific threshold of merger in the community, better able to deal with the child's input.

Here I have brought the notion of a single child learner into what has so far been framed as change in community patterns. The discussion of Garde's principle in Labov (1994) refers to merger, and of merger reversal, as a community-wide phenomenon. But, as Maguire et al. (2013) point out, there is something odd about this. Surely merger describes the status of an individual speaker's phonology, albeit with reference to a previous state of generalized contrast. When mergers are thought of as properties of individuals' mental grammars, a few examples of their reversal emerge. Many of these, unsurprisingly, involve children.

There is evidence that school-aged children can acquire difficult phonological features: younger children in families who migrated to a suburb of Philadelphia had generally acquired the complex conditioned allophonic split of /æ/ from their peers in Payne (1980)'s study. Those who arrived as adolescents had more mixed results, failing to acquire some of the more complex conditioning factors. Johnson (2010), in exploring the movement of the boundary of the /ɑ/-/ɔ/ merger in New England, describes a number of children who are exposed to the contrast at home and not at school, or vice versa, and sums up his findings this way: "younger children generally learn new patterns better, but there are no absolute rules for acquisition under various conditions of exposure." Of a sample of 6 Canadian children living in England, Chambers (1992) found that older adolescents did not acquire the local /ɑ/-/ɔ/ contrast at all: of the six, two – the youngest (age 9) and one of the two 13-year-olds – produced mostly unmerged speech, and the others mostly merged speech, despite showing acquisition of other local features.

Reports of merger reversal in adults tend to involve relocation from an area of merger

to an area of contrast. Nycz (2011) found Canadian adults in New York City to be partially successful in acquiring the low-back distinction – in general, those that displayed any change had only applied it to a restricted set of lexical items, highlighting the difficulty of this relearning. Other studies of adult transplants have found even less success in unmerging. Evans and P. Iverson (2007) found that when English speakers from Sheffield, where /ʌ/ and /ʊ/ are merged to /ʊ/, moved to attend university in an area with contrast, they centralized their read productions over the first two years – of both vowels. That is, although they accommodated by changing their production of these vowels, this accommodation did not include merger reversal.

## 1.2 Decomposing phonemic contrast

### 1.2.1 The notion of variable merger

Can a merger be “variable”? When conceived of as a property of a speech community it can: a merger is variable if some but not all speakers exhibit it (Maguire et al. 2013:231). But what about variable merger within the individual?

The semi-successful cases of contrast acquisition described by Johnson (2010) and Nycz (2011), as well as many of the speakers in Baranowski (2007)’s study, all involve within-speaker variation of a specific kind: some lexical items have the phonetic value expected in the merged system, and others have the phonetic value of the unmerged system. However, there is another way in which an individual can be thought of as “variably merged” in production: consider the case of Warren Maguire himself, who produces [u] for /u/ but both [u] and [ʊ] for /ʊ/. Maguire et al. (2013) contends that his is a merger in the phonetic sense, but that since he can “easily tell that members of FOOT form one set and members of GOOSE another”, the phonological situation is “more complicated” (p. 231): “when children grow up in a community where there are both merging and non-merging speakers, the potential exists for them to internalise both systems, even if only one of the systems surfaces in their production” (Maguire et al. 2013:235).

### 1.2.2 Perception

Near-merger is not only interesting because it provides explanations for otherwise inexplicable historical events. It challenges our understanding of what it means for an individual to have a contrast. Above I have described one kind of pattern that has been called near-merger, in which speakers produce a contrast that they are not aware of. However, the “Bill Peters” effect is not the only unexpected combination of perception and production that has been found in the context of phonological variation or change involving contrasts. Johnson (2010) describes essentially the reverse pattern in New England: children and teenagers with merged speech who are able to perform the /ɑ/-/ɔ/ contrast in the context of a minimal pair task.

So far, all of these examples of near-merger have been described as involving a mismatch of perception and production. But a speaker’s metalinguistic awareness of a contrast is only one aspect of their perception of it. In order to fully explore the perception of a contrast, minimal pair tasks are not sufficient. Online tasks are an improvement. Using an auditory forced-choice identification task, Hay, Warren, et al. (2006) finds that listeners who showed the /iə/-/eə/ merger in recorded word pairs nevertheless performed well above chance identifying tokens produced by talkers who have the distinction, though their performances was somewhat worse than that of unmerged listeners. This implicit kind of knowledge challenges assumptions about phonemic contrast as much or more than lack of explicit awareness.

Just because we can demonstrate that a speaker has knowledge of a contrast, whether implicit or explicit, does not mean that the knowledge is being used during actual speech processing the way we expect it to in the case of a robust contrast. As Labov (1994) notes, “as long as our data come only from methods that involve labeling, we cannot know whether or not the normal phonemic function - the use of a phonemic distinction to distinguish words in unreflecting interpretation - is also impaired.” In fact, experimental methods that require lexical access have been found to differentiate groups of speakers that perform very similarly on other types of perceptual tasks. Catalan has a pair of vowel contrasts (/e/-/ɛ/ and /o/-/ɔ/) that Spanish lacks. Spanish-Catalan bilinguals in Barcelona occupy a

continuum of language dominance. Amengual (2016) shows that although speakers who are Spanish-dominant are able to perceive the acoustic difference just as well as Catalan-dominant speakers, they make more errors when asked to decide whether stimuli containing the vowels are real words or not. The groups' perception of these vowels is clearly different at one level of processing, even if it may be similar at another.

### 1.3 This study

The case of /s/ and /θ/ is interesting for three reasons: first, it is a rare clear case of change toward a previously absent contrast, whether its origins are internal or not (I will agree with previous scholars of the change that they are not). Second, it enjoys a remarkably high level of social awareness in Spain, as has been previously reported. Merger and contrast of /s/ and /θ/ themselves do attract explicit social commentary in Andalusia, and my speakers' comments will give further support for this. Spanish words that explicitly refer to the contrast and merger of these consonants are in common use, and the fact that /s/ and /θ/ are consonants that generally have a low degree of confusability (Jongman et al. 2000) is certainly relevant. Trudgill (2003) notes that the case of /w/-/v/ is exceptional in reported cases of merger in that it involves consonants. This fact, rather than being a testament to the rarity of consonant mergers, may be simply a result of the fact that many of the languages that have, for arbitrary reasons, dominated studies of variation (especially English) have also had unstable vowel systems and have undergone a number of vowel mergers. What we know about mergers and merger reversal is disproportionately based on changes involving vowels, as Regan (2017b) also points out. It may be the case that some of the generally observed properties of mergers will not hold for consonant mergers. This remains to be seen; but for now there are very few documented mergers or merger reversals that attract social evaluation, and possibly none to the degree that the /s/-/θ/ contrast does.

Finally, previous studies of the /s/-/θ/ contrast in Seville report a significant amount of intraspeaker variation. What is the nature of this variability? Seville provides the opportunity to examine speakers that exhibit a different combination of production and



perception patterns than has been observed before.

This study presents data collected from 24 speakers (12 men and 12 women) of Seville Spanish, mostly college-educated, ranging in age from 20–35 years old. The data consists of sociolinguistic interviews along with the results of a set of perceptual experiments; details on the overall data collection procedures and on the background of the participants will be given in Chapter 3. Chapter 4 reports the results of analyses of the participants’ production, both elicited and spontaneous, and their performance in a pair of tasks designed to measure their sensitivity to the relevant acoustic contrast. Chapter 5 presents the results of a semantic priming experiment examining their processing of merged vs. unmerged productions of words. Prior to all this, Chapter 2 will give an introduction to the specific linguistic situation in question, situating it both historically and sociolinguistically.

There is an inherent expositional complication in this case when differentiating between the speaker’s phonemes, the phonemes of the target variety, and the sounds that are being produced. To mitigate it, this dissertation will use the following conventions: “Z-context” or “Z-word” will refer to contexts or words in which the relevant standard variety has /θ/, and “S-context” or “S-word” where the standard has /s/. /θ/ and /s/ are thus reserved for use when referencing a phonemic category in a speaker’s system, and [θ] and [s] to characterize realizations of the sounds.

### 1.3.1 A note on methodology

Sociolinguistics, when it comes to the examination of speaker production, has a tradition of careful consideration of the circumstances of speech and the methods used to elicit it, as well as of quantitative rigor in its analysis. While sociolinguists have long been occupied with questions about perception, the methods they have used to examine it have until recently been somewhat limited and informal.<sup>5</sup> Meanwhile, psycholinguistic examination of perception have a long history of carefully designed experiments and varied methodology. However, psycholinguistic experiments investigating perception, even those specifically

---

<sup>5</sup>See E. R. Thomas 2002 and Campbell-Kibler 2010 for reviews of work that attempts to fill this gap.

focused on variation, rarely do more than a cursory examination of their participants' production of socially or regionally-stratified variable phenomena, if they consider it at all. One reason for this is practical: the labor involved in collecting production data, especially the kind that approaches natural speech, is considerable. However, I argue that, in order to understand the connection between perception and production in cases of ongoing change, it is essential to combine the best of both of these two worlds and collect accurate production from the participants of well-designed perceptual experiments, and to make sure that the perceptual methodology explores processing at different levels.

## Chapter 2

# Background: *Seseo* in Seville

### 2.1 Introduction

This chapter introduces the linguistic phenomenon under discussion, the move toward contrast of /s/-/θ/ in Seville, in its historical and sociolinguistic context. Seville is located in the south of Spain (see Figure 2.1), in the region of Andalusia, which is politically one of 17 so-called autonomous communities of the Spanish state.

### 2.2 Historical context

#### 2.2.1 Castilian in Seville

“Spanish” is a label that eventually came to be applied to a variety of Vulgar Latin and then Ibero-Romance spoken by the Castilians. Castilian spread from its origins in the mountains of Cantabria in the north of the Iberian Peninsula all the way to the southern coast, and then across the seas to the Canary Islands and to the Americas. It was and is by no means the only language spoken on that peninsula (among them are Catalan, Galician, Basque, Portuguese, and Asturian), but it is the one that would come to dominate politically in most of the peninsula, and it would become the only official language of dictatorship-era Spain during the 20th century. The topic of this dissertation is a phenomenon taking place in a variety of the Castilian language that is spoken today in Seville.



Figure 2.1: Map of modern Spain.

The history of Castilian in Seville begins in the 13th century, when the southwestern part of Al-Andalus (Muslim Spain) is conquered by the Kingdom of Castile. The northern Christian kingdoms (Castile, Portugal, León, and Aragón) had been retaking control of the peninsula bit by bit from the Muslim people that had controlled the territory since the 8th century. This so-called *reconquista* ('reconquest') spanned over seven centuries. The city of Seville was taken by Ferdinand II of Castile in 1248. The *reconquista* brought with it a wave of migration out of Castile and León, and the bulk of the Muslim people who were not immediately displaced were quickly expelled. The remainder converted, often forcibly. By 1500, there were around 2,000 Muslim inhabitants out of about 750,000 total people in Andalusia (González Jiménez 2003:40).

The Kingdom of Castile along with its allies now controlled the bulk of the peninsula, and political and military power at the time was centered in its capital, Toledo. However, by the late 13th century, the southern city of Seville is a leader of Hispanic commerce (conducted chiefly with Holland and the Italian states), and one of the most cosmopolitan cities on the peninsula (De Cortázar and Vesga 1994:180). By the 15th century, Seville boasts 121,000 inhabitants, the most of any city in the peninsula, and is the nucleus of a rich urban center (De Cortázar and Vesga 1994:226), but it isn't until the 16th century

that Seville reaches the height of its importance. The European “discovery” of the West Indies in 1492 led to a surge in commerce and wealth taken from the Americas, and the city of Seville was once of its greatest beneficiaries. In 1503 the crown established *Casa de la Contratación* ‘House of Commerce’ in Seville and gave it broad powers over overseas financial matters. Seville became the principal port of entry and exit to and from the Americas, and for this reason, Seville constitutes the single most common place of origin for the earliest European immigrants to the Americas.<sup>6</sup> By some figures, Andalusia is the source of 78% of early emigration to the Indies, and Sevillians constituted 58% of those (Comellas 1992:138).<sup>7</sup> Also in 1492, the *reconquista* comes to an end with the conquest of the Kingdom of Granada in eastern. Granada is repopulated partly by Castilians, especially in the north, but there was also significant migration from western Andalusia (Narbona et al. 2003:55). Andalusia thus doubles in size as a territory, with Seville as its *de facto* center (Sánchez Méndez 2003:71). The court of the now-united Kingdom of Spain moves to Madrid in the 16th century, and Seville and Madrid compete as centers of powers of prestige during this time, known as the Golden Age of Spanish literature (Penny 1991:21).

The 17th century marks the point at which Seville’s size and influence begins to wane. Between the 16th and 17th centuries, the population drops from 112,000 to 60,000, both from American emigration and as a result of financial recession partly caused by a drop in the once-prodigious flow of precious metals from the Americas. The House of Commerce is moved to Cádiz in 1680 (De Cortázar and Vesga 1994:271), and power is centralized to Madrid gradually in the 18th and 19th centuries.

---

<sup>6</sup>Setting aside the failed Icelandic colony of Vínland in Newfoundland in the 11th century.

<sup>7</sup>The dominance of Andalusians among early emigrants to the Americas (and the related conclusion that many Latin American Spanish linguistic features have an Andalusian source) was the hotly debated in the early 20th century, but is generally accepted today in light of newer historical data. For an overview, see Sánchez Méndez 2003:85-95.

## 2.2.2 The emergence of an Andalusian Spanish

As stated above, the *reconquista* brings Castilian, then a regional language of the peninsula<sup>8</sup>, to Andalusia, and the speech of the descendants of these migrants will show little trace of the Arabic variety that once was spoken there.<sup>9</sup> A sweeping series of linguistic changes takes place during the centuries following the *reconquista*, but there is no evidence that they are the result of language contact. Some of these changes originate in the north and are carried south. Early on, Castile serves as a source of prestige norms for the newly conquered lands to the south, and migration still flows from it. However, the enormous cultural, political, and economic importance of Seville in the 16th and 17th centuries results in its emergence as a separate source of linguistic prestige and forms (Sánchez Méndez 2003:91), and a series of changes take place to the south that begin to differentiate Andalusian Spanish from its northern sibling varieties. A major restructuring of the Spanish sibilants takes divergent paths in Castile and Andalusia, giving rise to the linguistic variation that this dissertation is concerned with.

## 2.2.3 Historical development of the Spanish anterior fricatives

|              |                   |                     |                     |                  |
|--------------|-------------------|---------------------|---------------------|------------------|
| Phoneme      | /d͡z/             | /t͡s/               | /z̺/                | /s̺/             |
| Orthography  | <z>               | <ç>                 | <s>                 | <ss>             |
| Examples     | <i>hazer</i> ‘do’ | <i>caçar</i> ‘hunt’ | <i>casa</i> ‘house’ | <i>assi</i> ‘so’ |
| Latin origin | /kʲ/ & /tʲ/       | /kʲ/ & /tʲ/         | /s/                 | /ss/             |

Table 2.1: The medieval Spanish anterior fricatives (examples from Sánchez Méndez 2003).

The sibilant system of Medieval Spanish (see Table 2.1) underwent a shift in the 15th century such that the affricates /d͡z/ and /t͡s/ weakened to dental fricatives (/z̺/ and /s̺/).<sup>10</sup>

<sup>8</sup>Castilian eventually becomes the primary language of the Kingdom of Spain, and begins to be called “Spanish”. This term is sometimes avoided in regions of Spain where other languages are spoken in favor of the term “Castilian”, due to the implications of the label “Spanish” that it is the only language in the country, as well as in some parts of the Americas because of its colonial associations. Note that, in addition to referring to Spanish in general, however, the term Castilian is sometimes used to refer specifically to the variety of Spanish that is spoken today in the region of Castile.

<sup>9</sup>Despite some early accounts to the contrary, the linguistic influence of Arabic on the Spanish of the region is essentially limited to borrowings and toponyms.

<sup>10</sup>At the same time, /j/ and /ɟ/ became velar and subsequently merge to [x], which is in turn debuccalized in Andalusia.

|                    |                                        |                                      |
|--------------------|----------------------------------------|--------------------------------------|
| Phoneme            | /θ/                                    | /s/                                  |
| Modern orthography | <z>, <c>                               | <s>                                  |
| Examples           | <i>hacer</i> ‘do’, <i>cazar</i> ‘hunt’ | <i>casa</i> ‘house’, <i>así</i> ‘so’ |

Table 2.2: The Castilian anterior fricatives by 1650.

|                    |                                                                              |
|--------------------|------------------------------------------------------------------------------|
| Phoneme            | /s <sup>θ</sup> /                                                            |
| Modern orthography | <z>, <c> , <s>                                                               |
| Examples           | <i>hacer</i> ‘do’, <i>cazar</i> ‘hunt’, <i>casa</i> ‘house’, <i>así</i> ‘so’ |

Table 2.3: The single Andalusian anterior fricative by the late 17th century.

The result was an unstable contrast between dental and alveolar fricatives that played out differently in different regions of the peninsula (Lapesa and Pidal 1981; Penny 1991; Alvar-López 1996). In Andalusia, the dental and alveolar sibilants merged in the 16th century, and the voicing contrast between the resulting pair of dental sibilants was subsequently lost in the 17th century (see Figure 2.2). This resulted in a phonological system with a single voiceless anterior fricative (see Table 2.3).<sup>11</sup> To the north in Castile, a different sequence took place. The voicing contrast was lost first, in the 16th century.<sup>12</sup> The resulting fragile /s/-/s̺/ contrast was preserved by way of /s̺/ fronting to /θ/ in the 17th century, maintaining the historic affricate-alveolar fricative contrast (see Figure 2.3). This resulted in a system of two distinct voiceless anterior fricatives [s] and [θ] that persists today (see Table 2.2).

In Andalusia, this merger gives rise to 16th century misspellings like *disen*, *calsas*, *Blaz*, *inglez* (Pidal 1962), and the labels *çeçeo*, or less often, *zezeo*, start being applied in the 16th and 17th centuries to describe Andalusians. The original articulation of the single merged fricative in Andalusia is likely to have been nonsibilant and somewhere between dental and interdental. However, a fricative with a sibilant quality quickly becomes the urban norm of Seville, though this new [s] is not identical to the apicoalveolar /s/ of Castile, inherited from Latin (Lloyd 1987:331). The non-apical sibilant articulation expands north to nearby Córdoba. In the rural surroundings, meanwhile, a non-sibilant dental or interdental articulation takes root and spreads east toward Granada (Sánchez Méndez 2003:249). The

<sup>11</sup>A competing account is that the Andalusian anterior fricative had always been /s̺/ rather than /s/, so that the deaffrication resulted in immediate merger (Penny 1991:102).

<sup>12</sup>This loss of the sibilant voicing distinction is unique in western Romance (Lloyd 1987:268).

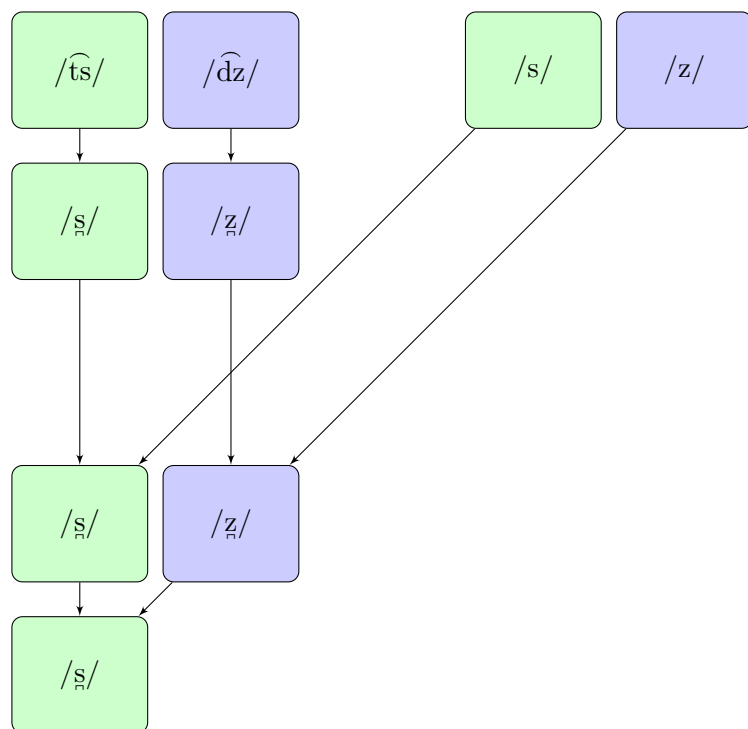


Figure 2.2: Changes in the Spanish anterior fricatives taking place in Andalusia from the 15th to the 17th century, adapted from Penny (2000):90.

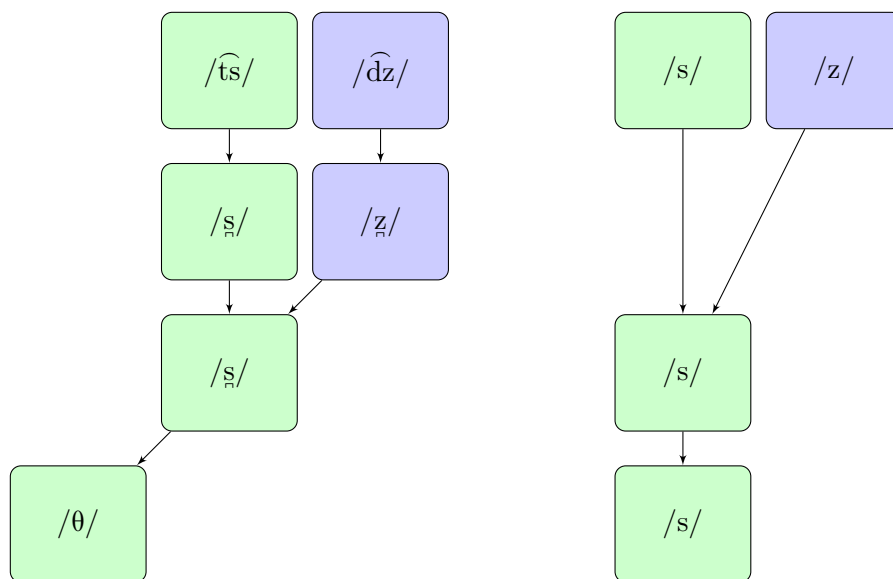


Figure 2.3: Changes in the Spanish anterior fricatives taking place in Castile from the 15th to the 17th century, adapted from Penny (2000):88.



non-sibilant variant seems to have had a rural association quite early on.

#### 2.2.4 Andalusian Spanish and the Americas

By the time portions of the Americas were being conquered and populated by the Spanish, the dental/alveolar contrast had been lost in Andalusia. The speech of most early Latin American colonizers would have reflected this, as well as the variety in articulation in this single phoneme. A single interdental fricative is attested to have been present in Puerto Rico, Santo Domingo, Colombia, Venezuela, Argentina, y Chile (Sánchez Méndez 2003:253), but the continued connection of the American ports with the prestigious norms of Seville (Pidal 1962) led to the eventual replacement of the non-sibilant variant in the Americas, resulting in the current American uniformity<sup>13</sup> of a single sibilant anterior fricative phoneme (with an Andalusian predorsal rather than apical articulation). The posterior /s/-/θ/ contrast never gains a firm foothold in the Americas, though it would have been present in the speech of many of the later colonizers and would begin to be associated with Spain.<sup>14</sup> The same is true of the apical [s] quality, which is restricted to northern Spain.

A series of lenitive processes that differentiate Andalusian Castilian from its sibling in Castile take place around this time, and are carried to some parts of the Americas. These include (Narbona et al. 2003:23):

- Debuccalization of the velar fricative (/x/ → /h/)
- Lenition of implosive (coda) /s/ (/s/ → /h/)
- Lenition of the interdental voiced fricative (allophone of /d/ in intervocalic position)  
[ð] → [∅]

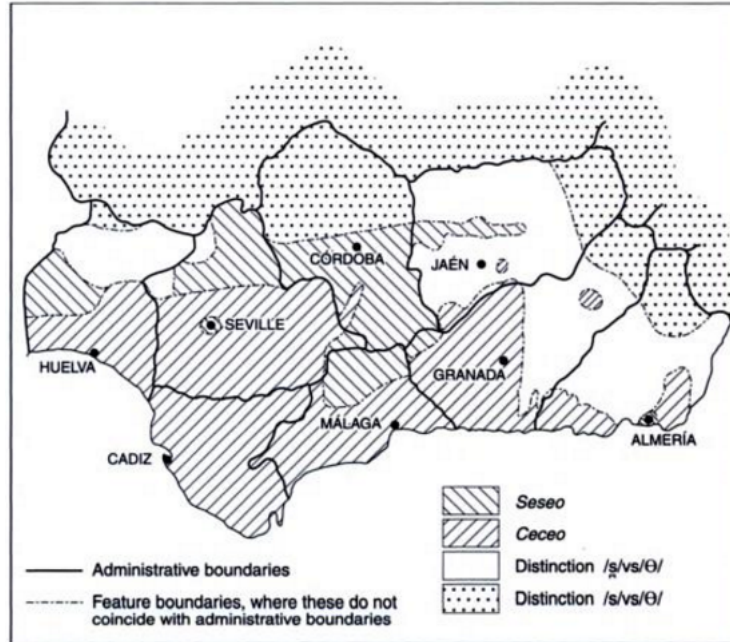


Figure 2.4: Map showing the extension of *seseo* and *ceceo* in Andalusia, from Penny 2000, adapted from data from the ALEA dialect survey (Alvar 1973).

### 2.2.5 The modern Andalusian anterior fricatives

Back in Andalusia, the single phoneme anterior fricative system with varying articulations settles in, coexisting with the contrast that dominates the north and spreads into some of the northern border lands of Andalusia, especially in the east. The geographic picture of the Andalusian anterior fricatives in the era of modern Spanish is therefore a complex one. The single phoneme has a range of articulations, forming a continuum that can (and often is, as we will see) be divided into two: a sibilant [s] and a non-sibilant articulation [θ].

I have already stated that a single sibilant phoneme /s/ is characteristic of the city of Seville. This sibilant articulation stretches north into some of the countryside and into Córdoba province, as well as into western Huelva away from the coast. There is also an area of [s] in the north of Málaga, and several other smaller pockets in the east. This non-apical [s] is not itself articulatorily uniform: Seville [s], the more common variety, is

<sup>13</sup>A marginal exception in Cuzco Peru involving a few words is reported by Caravedo (1992).

<sup>14</sup>For example, it serves as a shibboleth during the independence wars in South America (Sánchez Méndez 2003:254)

characterized by a convex predorsal tongue position, while the [s] of Córdoba has a flat coronal articulation (Narbona et al. 2003:157). Both of these sound quite different from the apical [s] found in the north.

A non-sibilant articulation [θ] occupies the surrounding countryside to the south of Seville and stretches down to the coast to Huelva, covering all of Cádiz province, and the coastal portions of Málaga, Granada and Almería, covering a large portion of the region. This [θ] ranges from a dental to a true interdental.

Finally, where the contrast does exist in Andalusia, in the east and north, /s/ is produced with a predorsal rather than an apical [s]. Divided in this way, we have four possible patterns: a single sibilant phoneme, a single non-sibilant phoneme, contrast with an apical [s], and contrast with a predorsal [s]. See Figure 2.4 for a map of Andalusia showing the mid-20th century geographic distribution of these patterns. It uses two labels for the first two patterns: “*seseo*” [seseo] and “*ceceo*” [θeθeo], and another, “*distinción*” for the /s/-/θ/ contrast. These are terms that are in common use in Spanish. The standard Spanish orthography that is eventually adopted reflects northern Castilian norms, and thus (even in America) represents the contrast: (/θ/ is represented by <z> and by <c> (when followed by /e/ or /i/), and /s/ by <s>). The rise in prestige of the Castilian variety spoken in Castile established contrast between /s/ and /θ/ as the national standard and led the Andalusian situation to be eventually labeled this way. These terms are standardly used in Spanish philology and will be used in this text. However, it is worth emphasizing that these are two names for the same phenomenon, with the exact same historical origin. The existence of two labels and their use by linguists is itself a manifestation of the influence of the northern standard, in this case as a point of reference that the Andalusian varieties are constantly being compared to (Morillo-Velarde 1997:207; Rodríguez Domínguez 2017:122–124). An additional term, *jejeo*, refers to lenition of *non-coda* /s/ or /θ/ to [x] or [h], which appears to be a pattern that occurs with some geographic extension (Prieto 2014), but is not the normative pattern anywhere.

### 2.3 Modern context

The centralization of Spain between the medieval period and the modern one and other sociocultural factors lead the northern varieties of Castilian to occupy a position of higher prestige, with the southern varieties being subordinate to it in the linguistic hierarchy (Villena-Ponsoda 2001:20). This centralization reached its peak, perhaps, under the dictatorship of Francisco Franco (1936–1975), during which a spirit of national unity was brutally enforced, along with strongly centralized linguistics norms. Not all of the dialect levelling can be attributed to the dictatorship, of course; Spain conforms to a greater European trend of dialect levelling and change toward prestigious natural varieties (Auer and Hinskens 1996). Italian and German regional varieties retreated in the face of a growing standardization since national unification of those countries in the 19th century, and this process continues. At the same time, growing mobility within the region of Andalusia leads to regional convergence (Villena-Ponsoda 2001:25).

In the years following the transition to democracy in Spain, language revitalization blossomed in areas speaking languages other than Castilian, with legislative measures taken to protect their use. Regional varieties of Castilian did not receive this kind of institutional attention, but a growing attention to regional identity had linguistic effects in Andalusia all the same. For example, it led to CanalSur, Andalusia’s regional television channel adopting a style guide in 1991 in which it “seeks to allay feelings of linguistic insecurity which might be felt by its broadcasters, [advising them] not to imitate the pronunciation of Valladolid... and to guard against hypercorrect pronunciation. The guide is a model of tolerance of language variety” (Stewart 2012:26). This marked a shift from the media policies of the previous decades.

The vast majority of professional Spanish voices in media employ the /s/-/θ/ contrast, and the style guide notwithstanding, exceptions to this even on CanalSur are difficult to find. Although *seseo* and *ceceo* can be heard in interviews with locals and from presenters on locally-oriented entertainment programming like flamenco talent shows on Canalsur, the

reporters and anchors exhibit nearly exceptionless distinction (Kvafil 2010).

### 2.3.1 Divergent norms in the east and west

Seville today is the largest city and capital of Andalusia, meaning that it is an economic and administrative center, particularly for the western part of the region. Scholars have long noted that the eastern and western halves of Andalusia have been divergent as to which linguistic norm they have oriented toward (see Villena-Ponsoda 2001:25, Villena-Ponsoda 2008:10, *inter alia*). In western Andalusia, the strength of Seville as a cultural center drives both migration into the city and the status of the linguistic norms of Seville as a regional standard. Regional convergence in the west then, means convergence or *koinization* (Villena-Ponsoda 2001:24) toward the norms of Seville, divergent from the national standard.

In contrast, eastern Andalusia lacks such a strong center of gravity and is more distant from Seville, and this has led to a greater orientation toward the north and toward the national linguistic norms. Subsequent levelling (Villena-Ponsoda 2001) has led to the formation of a regional variety that Hernández-Campoy and Villena-Ponsoda (2009) call “*español común*” (‘common Spanish’), which “acts as a buffer between the national standard – based on northern Castilian dialects (*castellano*) – and southern innovative varieties (*andaluz*, ‘Andalusian’ and, in particular, *sevillano*, ‘Sevillian’)” (p. 186).

### 2.3.2 The rise of *distinción*

The map of variation in the anterior fricatives shown in Figure 2.4, based on data collected by Manuel Alvar and colleagues as a part of the seminal Linguistic and Ethnographic Atlas of Andalusia (ALEA, Alvar 1973), would look quite different if it were to be redone today. For the past half century or so, Andalusian scholars have been documenting a retreat of *seseo* and *ceceo* in the face of the Castilian /s/-/θ/ contrast (see Chapter 2 Regan 2017b for an excellent review). The change is quite advanced in the eastern cities of Málaga (Ávila Muñoz 1994; Villena-Ponsoda, Sánchez Sáez, et al. 1995; Villena-Ponsoda

and Santos Requena 1996; Villena-Ponsoda 2001; Villena-Ponsoda 2007; Lasarte Cervantes 2010; Lasarte Cervantes 2012) and Granada (Salvador 1980; Moya Corral and Wiedemann 1995; Moreno 2005; Melguizo Moreno 2009; Moya Corral and Sosinski 2015). For example, Moya Corral and Wiedemann (1995)’s large study of 103 speakers in Granada finds an overall distinction rate there of 55%. Their results show a robust change in apparent time, especially among those with a college education: the rate of contrast rises from 63% in those 55 and older to 100% in the speakers ages 15-25. The youngest of these are born around 1978. Ávila Muñoz (1994)’s social network-based study of 30 speakers Málaga finds strong effects of age, gender, and social class with younger, more educated, and higher SES speakers tending toward realizations consistent with the standard contrast. Both [s] in Z-contexts and [θ] in S-contexts are present in this data, and there is evidence of an association traditionally held for some areas of men with *ceceo* and women with *seseo*. Social stratification of the use of the contrast has generally been found: “Data from eastern urban areas show that [the contrast] is a prestigious pattern frequently used by young, middle-class speakers. It is also clear that women favor this pattern in accordance with their leading role in changes from above” (Villena-Ponsoda 2008:14, citing Villena-Ponsoda 2001). Villena-Ponsoda (2005) finds that individuals in Málaga that are more integrated in their local community are more likely to use *ceceo* (along with [ʃ] and [h]): “the stronger his or her network ties are... the more frequently he or she uses the variables the local community members feel to be their own” (Villena-Ponsoda 2001:320). Villena-Ponsoda (2001) emphasizes that the findings of move toward contrast in Andalusia appear to challenge Garde’s principle of the irreversibility of mergers, and that the rise of *distinción* appears to be a “rapid and conscious” change (p. 129). He and Moya Corral (1997) point to dialect contact as the source of this change. As stated above, the consensus has been that western Andalusia, in contrast, tends to adhere to regional norms (e.g., Hernández-Campoy and Villena-Ponsoda 2009; Villena-Ponsoda 2008). Studies from the past decade challenge this consensus, at least with regard to the /s/-/θ/ contrast. Retreat of the merged forms has been documented in the western cities of Córdoba (Uruburu 1990), Jerez de la Frontera (García-Amaya 2008 as

compared to Carbonero Cano et al. 1992), and Huelva (Regan 2015; Regan 2017a; Regan 2017b, showing change since Heras et al. 1996), as well as in Seville (Carbonero Cano 1982 vs. Santana 2016b; Santana 2016a) and Seville province (Ruiz Sanchez 2017). Both García-Amaya (2008) and Regan (2017b) point to media and education as potential driving forces for this change.

### 2.3.3 Previous data from Seville

In the city of Seville, data from the late 20th century showed *seseo* remaining dominant: Narbona et al. (2003):162 states that the rate of distinction in Seville is around 30%, and that is is used disproportionately by “cultured, middle aged speakers”. 74% of college-educated speakers in Carbonero Cano (1982)’s Seville sample displayed *seseo*, along with 87% of the general population.<sup>15</sup> Lamíquiz and Cano (1985) indicate a high degree of acceptance of the use of *seseo* in survey data. All studies up to here point to the persistence of the Seville *seseo* pattern. However, as in the other western cities, this picture has begun to shift with newer data.

Santana (2016b) analyzes 24 sociolinguistic interviews collected in Seville between 2009 and 2015 as a part of the larger PRESEEA project (Moreno Fernández 1993). Her speakers are a sample of college educated speakers stratified into three age groups. She finds an overall rate of 71% [θ] in /θ/-contexts in the young speakers and 74% in the older speakers. She concludes that the contrast is not a recent phenomenon in Seville, at least among educated speakers. In the individual speakers, the rate of [θ] in Z-contexts ranges from 0.5% to 99%. Santana (2016b) does not examine tokens from S-contexts, reporting that they are categorically realized as [s] in the sample (p. 262). Santana (2016a) reports the data from the 24 less educated speakers, where she finds a rate of [s] in Z-contexts of 77%, which is more in line with Carbonero Cano (1982)’s educated speakers in the 1980s. The contrast is therefore not universal in Seville, but the city does seem to be undergoing a robust change toward it despite the high degree of acceptance of *seseo*. There is some

---

<sup>15</sup>An additional 6% of the latter used *ceceo*.

evidence for this outside the city as well. Ruiz Sanchez (2017), in a sociolinguistic sample from a traditionally *ceceo* Seville satellite town, also reports a majority pattern of contrast, with a low rate of [s] in Z-contexts (15%) and a lower rate [θ] in S-contexts (9.5%).

### 2.3.4 Linguistic attitudes and awareness

Seville is traditionally described as an island of *seseo* in a sea of *ceceo* (Alvar 1973), with both patterns competing with the contrast. The resulting complexity led to Seville being described a “confusing mixture of *seseo*, *ceceo*, and... distinction. It is, like *seseo* and *ceceo*, being the two allophones in free variation...” (Sawoff 1980:241). However, as Gonzalez-Bueno (1993) notes, this description was made by an outside observer and does not capture the complex sociolinguistic situation. In areas where *ceceo* predominates, speakers spontaneously move into both *seseo* and *distinción*, mediated by sociolinguistic factors. In contrast, in *seseo* areas, the variation tends to be limited to a mixture of *seseo* and distinction (Narbona et al. 2003).

Due to the strength of the Seville regional standard (and perhaps because of its international presence as well), *seseo* enjoys a certain acceptance in Spain, even outside Andalusia. *Ceceo*, in contrast, is limited to Andalusia and has a lower status (Lapesa 1957:6) “*seseo* is said to characterize more refined, educated speech, and *ceceo*, rural, relatively uneducated speech” (Dalbor 1980). In plays and novels, *ceceo* is used as a comic device (Navarro Tomás 1996:109). This differentiation between *seseo* and *ceceo* seems to have already been in place in the early 20th century “... although *distinción* and *seseo* are found to exist in a uniform and general fashion in their respective areas, *ceceo* shares its dominion with *seseo*, the former as a common form in popular speech and the latter as a finer, more careful usage. It is not easy to determine from the impressions formed during a quick visit the proportions by which the people of each place distribute themselves between the two forms of pronunciation.” (Navarro Tomás et al. 1933:239, translation mine).

In general, studies report that speakers are well aware of the differing patterns in anterior fricatives in the region (Villena-Ponsoda 2001; Regan 2017b:257). García-Amaya (2008)



states that in Jerez, some of his participants, especially younger women and those in contact with other social networks, made the contrast and were aware of it. Regan (2017b; 2017a; 2015) reports a high level of awareness of *ceceo* in his speakers, and points to this as a driving force behind the change toward *distinción*. Regan (2015) gives an extensive discussion of the social meaning of *ceceo* in Huelva and reports a number of sociolinguistic comments from speakers articulating explicit associations between *ceceo* and a local identity.

### 2.3.5 The acoustics of the /s/-/θ/ contrast

What differentiates /s/ from /θ/? In articulatory terms, both sounds are fricatives, characterized by turbulent airflow that is the result of a narrowing in the vocal tract while air is being pushed out by the lungs. In most classifications (Ladefoged and Maddieson 1996; Chomsky and Halle 1968), [s] and [θ] differ in the feature *sibilance*, or large amounts of acoustic energy at high frequencies: [s] is a sibilant, or strident fricative, while [θ] is not.<sup>16</sup> Sibilants, like other fricatives, involve turbulence generated from a constriction in the vocal tract, but in sibilants “a high velocity jet of air formed at a narrow constriction [goes] on to strike the edge of some obstruction such as the teeth” (Ladefoged and Maddieson 1996:138).

Usually, /s/ and /θ/ differ in place of articulation as well; /s/ being typically alveolar and /θ/ typically interdental. However, the exact place of articulation of both varies from language to language. [θ] in Andalusia has been described as a dental or an interdental nonsibilant fricative,<sup>17</sup> and [s] as an alveolar (Navarro Tomás 1996), predorso-alveolar (Hualde 2005:153) or a dentoalveolar (Navarro Tomás et al. 1933) sibilant. Sources agree that it is acoustically quite different from alveolar Castilian [s], and the source of that differences is usually attributed to tongue shape: Castilian [s] is consistently described as being apical, with the constriction formed with the tongue tip, unlike Andalusian (and American) [s]. The Andalusian [s] can be further broken down into several types: the typical

---

<sup>16</sup>Other terms that have been used to classify sibilants include “strident” (Chomsky and Halle 1968) and “obstacle” (Shadle 1985).

<sup>17</sup>Ladefoged and Maddieson 1996:143 state that they are not aware of any languages that contrast dental and interdental non-sibilant fricatives, and that some languages may consistently use one and some the other.

[s] of the city of Seville, according to Navarro Tomás et al. (1933):240, is dentoalveolar in place with a predorsal, convex tongue shape that gives it a “softer” frication and a low, almost emphatic resonance, while the characteristic [s] of nearby Córdoba has a coronal, convex tongue shape.

Acoustically, the contrast of [s] and [θ] is generally robust. In a study of English fricative perception, classification errors rarely crossed the sibilant/nonsibilant distinction (Jongman et al. 2000). A number of properties have been observed to differentiate [s] and [θ]. One is peakiness: “The non-sibilant fricatives... are characterized by relatively flat spectra lacking the pronounced peaks of other fricatives” (Gordon et al. 2002:30). Other are center of gravity (Jongman et al. 2000), intensity (Shadle 1985), and skewness and kurtosis (Jongman et al. 2000).

In Spanish specifically, Celdrán and Planas (2007) report that [s] about 20dB louder than [θ] and has a prominent peak at a high frequency (above 3000 Hz). Since they do not differ greatly in place of articulation, transitions are not helpful in distinguishing them, but according to Celdrán and Planas (2007):108, “in the case of [s] and [θ]... the difference in intensity is so great that there is no possibility of confusion” (translation mine).<sup>18</sup> The two acoustic examinations of the Andalusian /s/-/θ/ contrast that I am aware of have put forward intensity (Lasarte Cervantes 2010) and spectral peak (Regan 2017b) as cues that distinguish [s] and [θ] in Andalusian specifically.

## 2.4 The value of examining the individual

The bulk of the work on patterns in the anterior fricatives of Andalusia has taken a community-focused approach, from which there is of course much to be learned. Of more recent studies, many do report individual production (e.g., García-Amaya 2008; Santana 2016b), and Brendan Regan (Regan 2017b; Regan 2015) in particular takes an approach that focuses on individuals’ degree of merger (along with their possible social motivations). As discussed in Chapter 1, I view the patterns of the individual as giving important insight

---

<sup>18</sup>Note that these numbers are based on measurements taken of apical [s] in Spanish, which is both more intense than the Seville predorsal [s] and has its prominent at a higher frequency (6,230 vs. 4,522 Hz).

into the community-wide patterns of change as well as having inherent implications for theory, and I will therefore also take such an approach as well, with an added goal of giving a holistic picture of both perception and production for the individual speaker.

## Chapter 3

# Demographics and methods

The study was conducted over the course of two fieldwork trips to Seville, the first in September of 2016 and the second in May of 2017.

### 3.1 Recruitment

The *seseo* of Seville is in contact with the *ceceo* of many of the neighboring towns, as discussed in Chapter 1. This situation of contact is interesting in its own right. However, the focus of the current study is the loss of *seseo* in the city of Seville, and it is more difficult to evaluate the rate of unmerged speech in speakers that mix *seseo* with *ceceo* realizations (such as those in Ruiz Sanchez 2017's study). For this reason, I tried to avoid having participants with parents from neighboring towns. However, Seville as the capital attracts large amounts of migrants within the province who come to work and study, and requiring both parents to be city locals proved to be impractically restrictive. The final criteria were that participants (1) spent the majority of their childhoods until the age of 14 in the city, and (2) had at least one parent raised in the city.

Participants were recruited through flyers, word of mouth, and Facebook. The advertisement stated that participants were needed for a study of the Seville dialect. Participation in the study was compensated at 9 euros per hour.

## 3.2 Study methods

### 3.2.1 General procedures

The following tasks were conducted: an elicitation task, a sociolinguistic interview, a lexical decision task, an AX task, and a phoneme identification task.

The study was conducted over two, or in a few cases three, sessions. Each participant was met at a central location in the University of Seville Rectory building and accompanied to a quiet classroom or to the Phonetics Lab. They gave informed consent, and then they were then asked to complete a demographic questionnaire (see Appendix B). The production portion of the interview was conducted first in order to best avoid the participants' becoming aware of the focus of the study. This consisted of the elicitation task and then the sociolinguistic interview, recorded on a Zoom H1 Ultra-Portable Digital Audio Recorder using a Audio-Technica PRO70 Cardioid XLR Lavalier microphone, sampling at 44,000 Hz.

The interview was followed by the lexical decision experiment. The second session then began with the phoneme identification task and then the AX task. These two tasks were conducted last as, in asking speakers to listen to many instances of [s] and [θ], they made it fairly apparent what the object of the study was. Finally, the participant completed a local identity questionnaire (Appendix B), and I conducted a debriefing conversation immediately afterward, which was also recorded.

All portions of the study sessions were conducted by me. I am non-native speaker of Spanish who has spent a total of 3 years in Seville. I have a near-native proficiency in Spanish and am often taken for a native speaker. I achieved fluency in Spanish while studying in Seville, and my Spanish shows a strong influence of the local variety, perhaps most saliently in lenition of coda /s/. Spaniards (linguists and non-linguists) from both Andalusia and elsewhere in Spain comment that my Spanish sounds markedly Andalusian. I has a self-reported rate of 100% *seseo*, having been first taught Latin American Spanish and then exposed to mixed *seseo/distinción* input in Seville. When conducting sociolinguistic interviews, a general concern is the potential influence of the presence of the interviewer

| Z-words         |               | S-words       |         |
|-----------------|---------------|---------------|---------|
| <i>buceo</i>    | ‘scuba’       | <i>suelo</i>  | ‘floor’ |
| <i>zurdo</i>    | ‘left-handed’ | <i>basura</i> | ‘trash’ |
| <i>cielo</i>    | ‘sky’         | <i>sonó</i>   | ‘rang’  |
| <i>precio</i>   | ‘price’       | <i>seco</i>   | ‘dry’   |
| <i>cinturón</i> | ‘belt’        | <i>sorda</i>  | ‘deaf’  |

Table 3.1: The target words in the elicitation task.

on the speech of the interviewee, a phenomenon known as the Observer’s Paradox (Labov, Yaeger, et al. 1972). A specific worry here is that a foreign interviewer may induce a heightened pressure to employ a more standard style. However, my impression is that my use of local dialect features generally served to reassure participants that it was OK to speak normally.

### 3.2.2 Sentence completion task

The elicitation task was conducted before the sociolinguistic interview. This was done with the intent to maximize attention paid to speech (Labov 1972:112) and thus pressure to make the contrast. Participants generally relax over the course of a sociolinguistic interview. An elicitation task was chosen over a reading task in order to minimize the affect of orthography, which informs and/or reminds the speaker which sound is standard in the context.

The elicitation task was a cloze task (Taylor 1953) in which the participant had to provide the final word of a sentence. The sentence was read aloud by the interviewer, and repeated once if necessary. There were 29 sentences (see Appendix A for a complete list). 6 of the target items were S-words and 6 were  $\theta$ -words (see Table 3.2.2. Since the subjects did not always produce the target word, the amount of S- and  $\theta$ -words actually produced during the task varied.

### 3.2.3 Sociolinguistic interviews

The current study is primarily concerned with the amount of *seseo* speakers use in everyday speech, best determined by eliciting naturalistic speech in a setting that minimizes attention

paid to speech (Labov 2001). The sociolinguistic interview is the “gold standard” method used by sociolinguists to obtain speech data (Tagliamonte 2011).

For the present study, the investigator prepared list of topics to select from during the interview. The full list is below. Two are related to yearly events that are culturally important in Seville. *Semana Santa* ‘Holy Week’ refers to the week of Easter. Seville is home to what is probably Spain’s most well-known Easter celebration. Most of the city has the week off from work, and processions are put on by 60 *hermandades* ‘brotherhoods’. Each *hermandad* carries three multiple-thousand-pound structures called *pasos* bearing an image of the Virgin Mary or of Jesus on the cross, through the streets of the city on the backs on men of the brotherhood. They are preceded by a string band and a long procession. They form an impressive sight that quite literally takes over the city for the week. Holy Week is a point of strong local pride, even for non-religious Sevillans (Mitchell 1990). Not everyone enjoys the event, but being an aficionado of Holy Week is part of a certain Sevillian stereotype. The same stereotypical Sevillian will also be a loyal attendee of the “April Fair”, an weeklong annual event in Seville, taking place on the outskirts of the city. It originated as a simple trade fair, but developed over time into a celebratory gathering of huge proportions. The city virtually empties out into the fairgrounds for the week. It is much more of a party than solemn Holy Week; attendees eat, drink sherry, and dance to traditional music. Women don a type of colorful ruffled dress that is only worn to the fair, and those who are well-off enough arrive to the event in a horse-drawn carriage.

Other topics were drawn from the standard set of sociolinguistic modules from (Labov 1972). The list included:

- Childhood games
- Childhood friendships
- Childhood accidents
- Local festivals
- Premonitions and dreams

Of these, the topics that were covered in every interview were local festivals and childhood games. The interviews lasted between 30 and 60 minutes.

### 3.3 Information about the participants

#### 3.3.1 Demographic and identity questionnaires

Demographic information was collected for each speaker, including: age, gender, neighborhood, profession, education (highest achieved level), whether they were language teachers or studied language, whether they had ever lived outside of Seville, and where. I also collected their parents' place of origin and professions. I asked if speakers were bilingual (none were) or had any known hearing problems (none were reported).

The identity questionnaire was developed with reference to the one used in Reed (2016)'s study of local identity in Appalachia. I asked speakers to rate on a scale of 1-7 how much they identified with a set of local characteristics: as a Sevillian, as a fan of Holy Week, as a Feria-goer, and as a flamenco fan. They were also asked whether they would consider moving away from Seville. Based on their responses, participants were binned into two categories of local identity: High and Low.

Selected responses to both questionnaires for all the participants can be found in Appendix B.

#### 3.3.2 Evidence of awareness

Subjects generally exhibited a high degree of awareness of *seseo*, its links to Seville, and of the phonemic contrast. This manifested itself in overt commentary, both during the debriefing and during the session itself. For example, at the end of the elicitation task, participant M01, completely unprompted, pointed out that Sevillans say the word *cielo* (a Z-word) 'sky, heaven' with [s], and that in the context of a particular phrase connected to *Semana Santa*, a [θ] in that word would be inappropriate:

Entonces aquí hay una cosa que tú dices [θ]ielo y la gente lo dice normal, pero



por ejemplo en Semana Santa, ¿tú no sabes la Semana Santa cómo es? Cuando van a levantar los palios dicen, ‘¡Al [s]ielo con ella!’ Entonces todo el mundo – si le pones esa frase todo el mundo diría [s]ielo con S... ‘¡Al [θ]ielo!’ y todo el mundo se queda ‘¿Qué ha dicho? ¿Qué ha dicho?’

‘So here, there’s this thing that if you say [θ]ielo, people say it normally, but for example during Holy Week, you know how Holy Week is? When they go to lift up the image they say, ‘[up] to the [s]ielo with her!’<sup>19</sup> And everybody – if you give them that phrase everybody would say [s]ielo with an S... [If you were to say] ‘Up to the [θ]ielo!’ and everybody<sup>20</sup> is like, ‘What’d he say? What’d he say?’”

In this statement, the speaker demonstrates a keen awareness that speakers adapt their use of this variable to the context, and explicitly connects *seseo* productions to Seville and to Holy Week.

Subjects were asked about their own use of *seseo*:

|                         |                                         |
|-------------------------|-----------------------------------------|
| Te consideras:          | ‘Do you consider yourself’              |
| Algo seseante           | ‘Somewhat <i>seseante</i> ’             |
| Muy seseante            | ‘Very <i>seseante</i> ’                 |
| Ni ceceante ni seseante | ‘Neither ceceante nor seseante’         |
| Muy ceceante            | ‘Very <i>ceceante</i> ’                 |
| Seseante y ceceante     | ‘ <i>Seseante</i> and <i>ceceante</i> ’ |

17 reported that they were “somewhat” *seseantes*, 3 “very”, and 4 ‘Not at all’. They were asked to rate their awareness of their use of *seseo* on a scale of 1-7. The mean rating was 5.29.

### 3.3.3 Attitudes towards *seseo*

The identity questionnaire included several questions aimed at gauging speakers’ attitudes toward the linguistic variation at issue. In one, the participant was asked to choose between

<sup>19</sup>*Her* refers to the image of the Virgin Mary, carried on a heavy wooden structure known as a *paso*.

<sup>20</sup>That is, the bearers of the *paso*, known as *costaleros*.

three options:

1. *Seseo* is a local feature and should never be corrected.
2. *Seseo* should only be corrected in some environments.
3. *Seseo* is incorrect and should be corrected as much as possible.

The majority (17) of the subjects responded that it should never be corrected, and 7 gave the intermediate option that it should be corrected in certain contexts. The questionnaire also asked whether they had ever had anyone correct their use of *seseo*.<sup>21</sup> 21 out of 24 subjects responded that no one had.

Finally, subjects were asked to rate the acceptability of *seseo* and *ceceo*, respectively, along with their own awareness of their use of *seseo*, on a scale of 1-7. Figure 3.1 shows the responses. The mean acceptability rating of *seseo* was 6, while the mean rating of *ceceo* was 4.75.

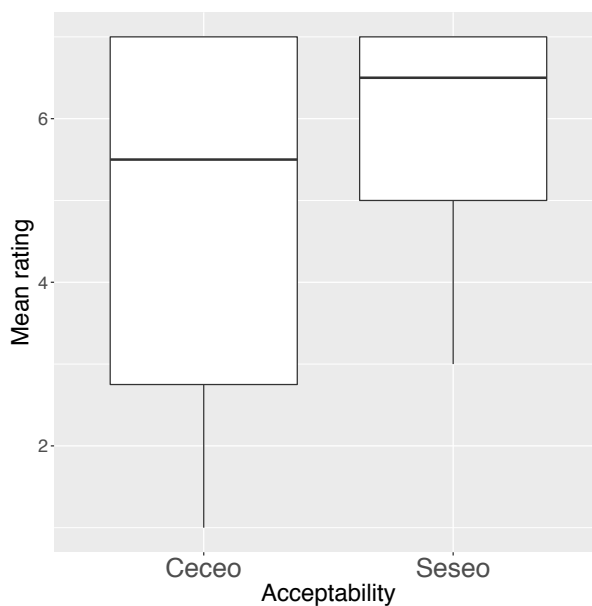


Figure 3.1: Subjects' ratings of the acceptability of *ceceo* and *seseo* on a 7-point scale.

---

<sup>21</sup>The questionnaire actually asks whether anyone had ever corrected their use of "Z". While administering it, I realized this was ambiguous and always clarified this question verbally, but it is possible some of the subjects thought I was asking about *ceceo*.

## Chapter 4

# Production and perception

### 4.1 Introduction

What does the emergence of a new phoneme look like at the individual level? As seen in Chapter 2, 20th century studies of the emerging /s/-/θ/ contrast tended to examine patterns at the group level, and undoubtedly there is much to be learned from such an approach. However, the examination of the individual gives rise to interesting questions. The high degree of intraspeaker variability in the production of Z-words in Seville reported by Santana (2016b) is intriguing: when a speaker produces a contrast only some of the time, what is that status of that contrast in their linguistic system?

The answer to this will come partly from observing the distribution of the sounds within each etymological context and across a variety of lexical items. This chapter presents an analysis of a relatively large amount (~30 minutes) of spontaneous conversational speech for each of the 24 Sevillian Spanish speakers introduced in the last chapter. However, to approach a complete picture, we need to combine production data with perceptual experiments performed with the same speakers. It may seem obvious that any speaker that is capable of making a contrast in production will be able to tell apart the sounds in perception. However, there is ample reason to question this assumption, as demonstrated by the many documented examples of so-called “near-merger”, introduced in Chapter 1 and discussed further below in Section 4.4.1. With this in mind, two experimental tasks were

administered to ascertain whether the speakers are able to identify and discriminate between the two acoustic categories. A third property relevant to phonemic contrast, the use of a contrast during lexical retrieval, is arguably separate from the properties already mentioned, and is addressed in Chapter 5 using a semantic priming paradigm with the same participants.

#### 4.1.1 Production and perception of a contrast at different levels

In order to properly describe the status of a phonological contrast in a given language variety, several facets of its speakers' production and perception should be taken into consideration. The necessity of this becomes apparent when considering evidence from a number of studies over the last fifty years showing that contrast and merger are not as clearly delineated as might be expected. The majority of those studies document examples of phenomena known as "near-merger": cases of ongoing phonemic change in which speakers exhibit unexpected combinations of behaviors. At this point, it will be useful to discuss what behaviors *are* expected given a stable phonological inventory. In the normal state of affairs, a speaker with two phonological categories is expected:

1. to have phonetically different sound distributions for the two categories;
2. to produce sounds from the two distributions with systematic reference to lexical context;
3. to perceive the contrast in their own and others' speech; and
4. to use the difference to access lexical meanings.

Generally, these behaviors cluster reliably together when it comes to stable phonemic contrast, and tend to be absent in the case of long-completed mergers (or contrasts that have never been present in a language). Situations of variation and change tell a different story, as they often do. Many studies of mergers in progress have found speakers who show some, but not all, of these characteristics. We may expect this during change toward a new

contrast as well. This chapter will address the first three points above as they apply to the case of the /s/-/θ/ contrast in Seville, and Chapter 5 is dedicated to the fourth.

### 4.1.2 Lexical considerations

Above I refer to the mapping of two sound categories to different items in the lexicon as a property of contrast. It is worth thinking about what the absence of this property looks like. For example, just because two speakers produce a lexical item with the opposite sounds does not mean that one of them lacks one of the two categories. There are many examples of contrasts whose lexical distribution varies geographically. The /ɑ/-/ɔ/ contrast in the United States shows this kind of variation. The lexical class membership of words containing Middle English /o/ followed by /g/ (e.g., *dog*) “varies in Pennsylvania not only from county to county, but even among speakers from a single county” (Herold 1990:6).<sup>22</sup> It would be a mistake to compare two speakers from two different counties and call either of them merged because they produce /ɑ/ in a few contexts where the other produces /ɔ/. A careful examination of speaker behavior will distinguish these scenarios from true merger. Specifically, examining multiple tokens of the same lexical item will be illuminating: if the variability is simply a matter of differences in lexical class membership, a single speaker should show consistency in production within lexical items. There is some evidence that speakers in communities undergoing merger reversals can exhibit patterns that resemble lexical diffusion. For example, some of the speakers in Baranowski (2007)’s study of NEAR and SQUARE in Charleston, South Carolina seem to have two vowel categories but not all words are in the expected category; that is, they pronounce *here* with a higher vowel, but pronounce *hair*, *beer*, *bear*, *fear*, and *fair* with the same lower vowel. Similarly, Jespersen (1954), describing the situation of /oi/ and /ai/ in 18th century England, cites Kenrick (1773) as saying that it would “now appear affectation” to produce *boil* and *join* with [oi], but [oi] in other words like *oil* and *toil* seems to be *de rigueur*: Kenrick calls the tendency for [ai] in these contexts a “vicious custom” (Jespersen 1954:330). The case of /q/, which

---

<sup>22</sup>She adds that “The behavior of words in which ME /a/ follows a /w/ (e.g. *swamp*, *water*, *watch*, *wash*) is also especially subject to regional variation, as is the pronunciation of *on* and *off*.” (Herold 1990:6).

was a phoneme in classical Arabic but subsequently merged with /g/, /k/, and /ʔ/, appears in a lexically limited fashion in the Arabic variety spoken in Cairo (Haeri 1991). Finally, in examining the acquisition of the /ɑ/-/ɔ/ contrast by Canadian adults, Nycz (2011) describes adults that acquire the contrast consistently for some lexical items and not others. It seems that merger reversal can be partial in the the sense that it doesn't create as large a lexical class it is expected to.

If there is variability within lexical item, does that then constitute merger? Maguire et al. (2013) argues that it does not, citing his own system of producing only [u] for /u/, but both [u] and [ʊ] for /ʊ/ (p. 231). Such regularity satisfies the above criterion of a systematicity that makes reference to lexical context. This is a compelling argument for contrast; how would such a system be possible without information about the phonemes being present in the mental lexicon? In order to observe the systematicity within this kind of variability, however, a careful examination of production is necessary. In this chapter, I do that with the Seville speakers in my sample.

### 4.1.3 The role of orthography

As mentioned in Chapter 2, the /s/-/θ/ contrast is faithfully represented in standard Spanish orthography (with S and C or Z, respectively), even in Latin America where the contrast is completely absent. This may very well be a factor that is enabling speakers to borrow the contrast. García-Amaya (2008) cites growing literacy as a potential explanation behind the change that he finds in apparent time in Jerez away from *ceceo* and toward /s/-/θ/ contrast, as does Regan (2017b) for the same change in Huelva.

### 4.1.4 Consonants versus vowels

Descriptions of phonemic merger in the sociolinguistic literature have been vowel-centric for what are probably accidental reasons, as Regan (2017b) notes. This may be important for our understanding of merger because of the broad acoustic differences between consonants and vowels and the perceptual differences that result from them. Listeners perform

differently on categorical perception tasks involving vowels than those with consonants. The categorization curve for consonants is steeper (a very small acoustic difference along a category boundary results in category differences for consonants), and within-category discrimination is poorer (listeners are better at telling apart tokens of vowels in the same category than consonants in the same category: Pisoni 1973). This suggests that “the articulatory reality of the vowel-consonant distinction... translates to a perceptual reality that causes strong category binding for consonants but certain degree of category insecurity, or flexibility for vowels” (Cutler 2012:8). This leads to listeners being more prone to misperceiving consonants in casual conversation (Bond 1999), while being better at identifying vowels than consonants in noisy experimental conditions (Cutler et al. 2004). The potential implications for vowel vs. consonant merger are large, especially if misunderstandings are a driving force of merger as e.g. Yang (2009) suggests. Broad phonetic differences between vowels and consonants may also correspond to different degrees of tendency toward social awareness of variation.

## 4.2 Social predictions

Social stratification was expected in the use of the /s/-/θ/ by this sample of speakers. Linguistic changes “from above” involve “the importation of a prestige feature from outside the speech community...” (Labov 2001:274). Changes from above have certain characteristics: they are favored in careful speech, and they are used more by women than by men (Labov 1966; Milroy and Milroy 1978, *inter alia*). The change toward /s/-/θ/ contrast in Andalusia has been called a change from above (Villena-Ponsoda 2001), and gender, education, and social class have all been found to be related to its use (Moya Corral and Wiedemann 1995; Villena-Ponsoda 2001; García-Amaya 2008; Regan 2017b). Accordingly, I expected gender, education (No College vs Some College), and social class as estimated by the median income in the speaker’s neighborhood to be associated with a higher rate of /s/-/θ/ contrast in production. I also hypothesized that a strong local identity and having a parent who comes from an area with the contrast might influence production of Z-words,

and that having contacts that use *ceceo* might influence production of S-words. Finally, those whose college majors or profession deal directly with language were predicted to show more contrast, due to heightened awareness of standard norms.

This chapter proceeds as follows: Section 4.3 presents the production analysis, and Section 4.4 presents two perception experiments.

## 4.3 Production analysis

The procedures with which this data was collected are described in Chapter 3.

### 4.3.1 Preparation of data for analysis

All of the speech produced by the participant in each of the elicitation tasks and interviews was transcribed by the investigator in standard orthography using ELAN and subsequently force-aligned using the *faseAlign* software (Wilbanks 2018). Sections where the interviewer alone was talking or where the interview’s speech overlapped significantly with the subject’s were not transcribed. Material that clearly consisted of a direct quote was marked in the transcription process and subsequently excluded from analysis. Since *faseAlign* uses Latin American Spanish-based acoustic models, Z-context fricatives were labeled “S” by the model. I used the Python programming language to write a script replacing the relevant labels with “Z” using the orthographic representation of the word. The alignment of relevant tokens was not hand corrected due to time constraints. I extracted all instances of S and Z in onset position. S and Z in coda position were excluded from analysis because both /s/ and /θ/ are very often lenited in coda position in Andalusian Spanish (Penny 2000).

A combination of impressionistic auditory coding and subsequent acoustic measurement was used in order to assess how much the new sound [θ] is being used and in which contexts. Auditory coding has both advantages and disadvantages when compared to acoustic analysis (Milroy and Gordon 2008). It is more time-consuming than modern acoustical techniques, especially for large datasets. Auditory judgments are inherently subjective, and researcher bias is a risk. On the other hand, while acoustic measurements are precise and unbiased,



perception is a complex process that takes advantage of a variety of cues in the signal, both during and adjacent to the segment in question. Without a solid basis of perception work on a particular contrast, it is hard to know which pieces of acoustic information are most relevant and which are superfluous. A researcher’s ear may be better able to perceive a contrast than an uninformed acoustic measurement. The challenge is whether the researcher is tuning into the same cues as members of the speech community, as the latter is the real subject of interest. Some form of auditory coding is used in most of the studies of the /s/-/θ/ contrast discussed in Chapter 2. This study follows Lasarte Cervantes (2010) and Regan (2017b) in combining auditory coding with an acoustic analysis.

### 4.3.2 Auditory coding

The auditory coding analysis was performed by me, a trained phonetician. I am not a native speaker of Spanish. My high level of proficiency in the local variety is discussed in Section 3.2.1, and I am a native speaker of two languages that have a /s/-/θ/ contrast: English and Icelandic. This coding analysis would ideally be carried out by several phonetically-trained native speakers of the Seville variety and their inter-reliability checked, but this was not possible due to practical constraints.

The data was acoustically coded using a Praat script that plays each token inside a 1-second window and simultaneously displays a spectrogram of that window. Each token was coded as being an instance of [s], [θ], or [x]/[h] (Santana 2016a reports <1% of [h] realizations across both contexts in her data). Tokens were additionally tagged as being intermediate when I had difficulty making a decision. Grossly misaligned tokens were marked and later excluded from acoustic analysis, as were tokens with noticeable overlapping background noise or speech. The coding for S-contexts and Z-contexts was conducted consecutively.<sup>23</sup>

---

<sup>23</sup>I found it easier to focus on the acoustic quality rather than the orthographic value when the task was grouped this way.

|     | [θ]           | [s]            | intermediate | [x]         |
|-----|---------------|----------------|--------------|-------------|
| /s/ | 1.42% (158)   | 95.62% (10667) | 1.6% (178)   | 1.37% (153) |
| /θ/ | 59.91% (2286) | 35.98% (1373)  | 2.96% (113)  | 1.15% (44)  |

Table 4.1: Number and percent of segments in the overall data given each auditory coding label, separated by phonemic (etymological) context.

### 4.3.3 Coding results

The procedure above resulted in a total of 11,974 coded tokens of standard /θ/ and 4,005 tokens of standard /s/ (/s/ is about seven times more frequent than /θ/ in standard Castilian (Narbona et al. 2003:164)).<sup>24</sup> The overall rates of each auditory coding value in the two contexts are shown in Figure 4.1 and in Table 4.3.3 with Intermediate tokens separated out into a separate category. Ns are shown above each bar in this and the following figures in this section. 0.5% of tokens were coded as intermediate; this is well below the rate in Lasarte Cervantes (2010)’s study in Málaga, in which auditory coding was conducted by native speakers and 12.3% of tokens were coded as intermediate (p. 492).<sup>25</sup>

It is immediately clear from Figure 4.1 that [θ] is not being produced indiscriminately across S- and Z-contexts, but rather is mostly limited to Z-contexts. The rate of [θ] in Z-contexts is similar to the global rate of 74.22% found in the same age group by Santana (2016b). This study lends support to her finding that [θ] is the dominant pronunciation of /θ/ in this demographic group. The rate of *jejeo* ([x]/[h]) is quite similar to the rate in Santana (2016a). She reports that this realization is limited to lexical items that are serving as discourse-markers or other conversational functions, and an impressionistic examination suggests that this is true of my data as well. I exclude these realizations from statistical analysis.

Fricatives in Z-contexts were expected to be realized as [θ] more often by women than by men and more often as neighborhood income increased. They were also predicted to occur more often in the speech of participants with some college education, those with a

<sup>24</sup>My data has a less extreme ratio because the frequencies are especially imbalanced in coda position, which I excluded.

<sup>25</sup>The exact term used by Lasarte Cervantes (2010) is *[realización] dudosa* (‘dubious/questionable’).

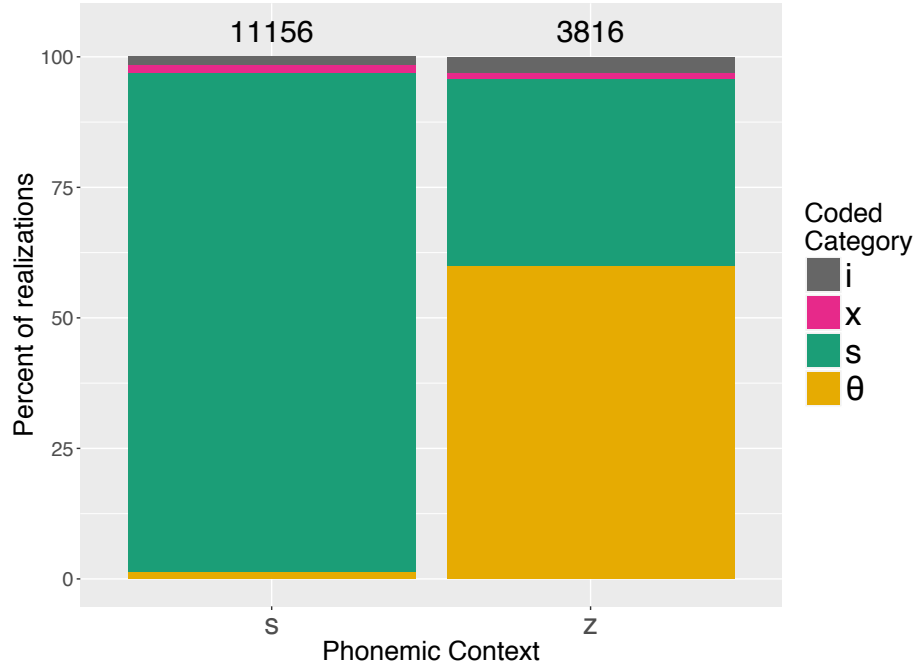


Figure 4.1: Percent of segments in the overall data given each auditory coding label, separated by phonemic (etymological) context.

| Education    | Gender | N |
|--------------|--------|---|
| No College   | F      | 3 |
| Some College | F      | 9 |
| No College   | M      | 4 |
| Some College | M      | 8 |

Table 4.2: The number of speakers in this study by gender and education level.

parent from a contrast area, those who studied or worked with language, and those with a less local identity.

Figures 4.2 and 4.3 show the distribution separated out by Gender and by the type of data: spontaneous (conversational interview) data shown in Figure 4.2 and the data elicited in the sentence completion task (described in Chapter 3) in Figure 4.3. These figures are not suggestive of any patterning of this variable according to either factor: both the left and right halves of each figure and the two figures look quite similar. Elicited tokens were less likely to be coded as Intermediate and as [x]. This is expected: the elicited tokens are closer to careful speech, which is characterized by more deliberate articulation and fewer

lenition processes (Labov 2001:158, Hanique et al. 2013).

Figures 4.4 and 4.5 show the same data by education level, divided by whether or not the participant had any college education (recall that many participants were current college students). Note that the distribution of participants across education categories was not even (see Table 4.3.3).

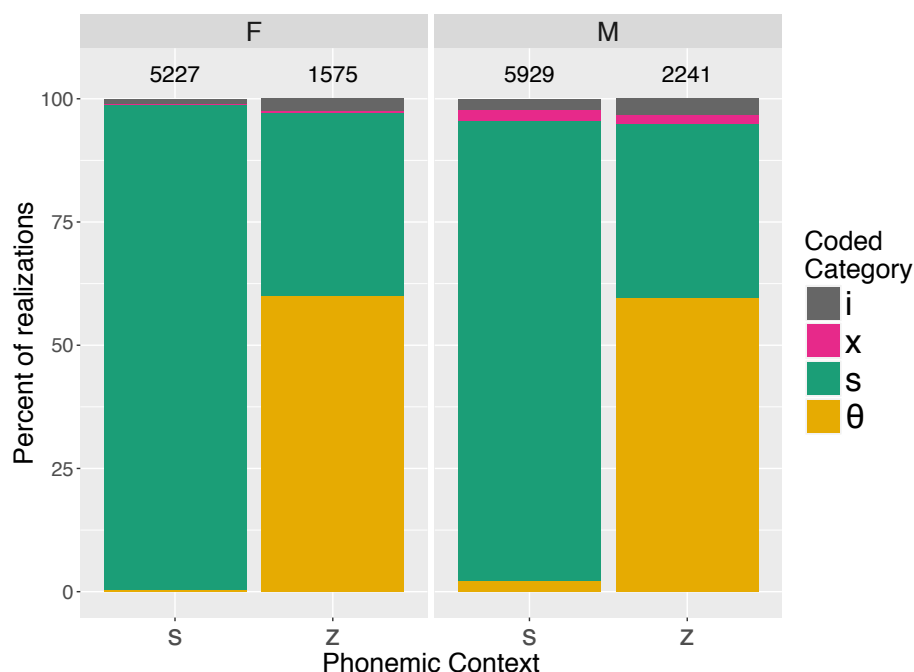


Figure 4.2: Percent segments in the spontaneous speech data coded [θ] in Z-contexts by phonemic context, faceted by the gender of the speaker.

Table 4.3 shows a binomial linear mixed-effects model predicting production of fricatives in Z-contexts as determined by auditory coding. The levels of the predicted value are [s] (0) and [θ] (1); tokens coded intermediate and [x] are excluded. The model contains random effect terms for Speaker and Word. The external fixed effects that were considered for the model were: Age, Gender, Word Frequency, Median Neighborhood Income, Education (Some vs. No College), whether the speaker had significant contact with speakers from surrounding towns where *ceceo* is used (Ceceo Contacts), whether their work involves language (e.g., studying philology or teaching English as a second language), whether they have a parent who is from an area of contrast, and finally, strength of local identity (High

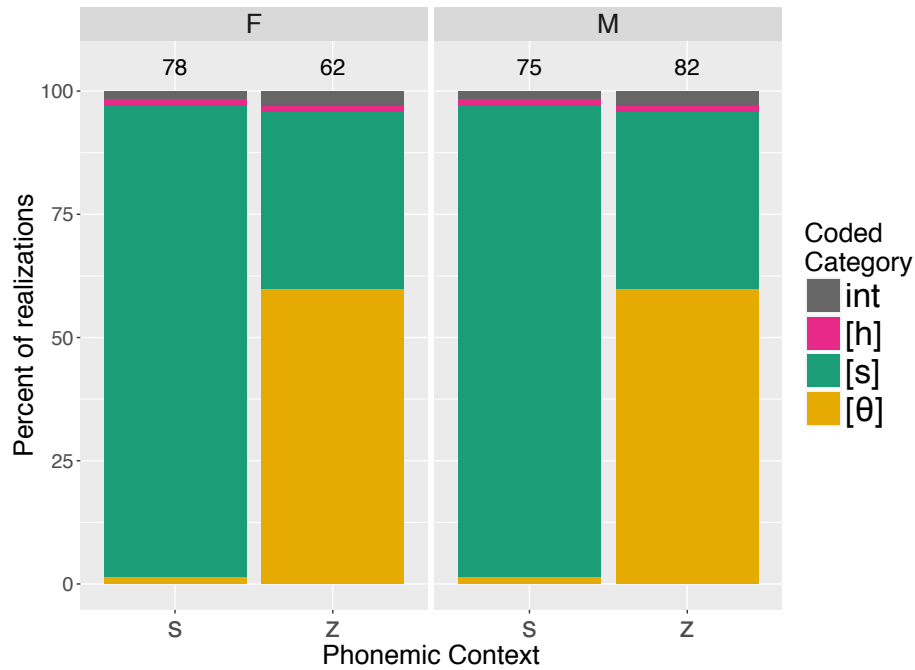


Figure 4.3: Percent segments in the elicited data coded [θ] in Z-contexts in the elicited portion by etymological context, faceted by the gender of the speaker.

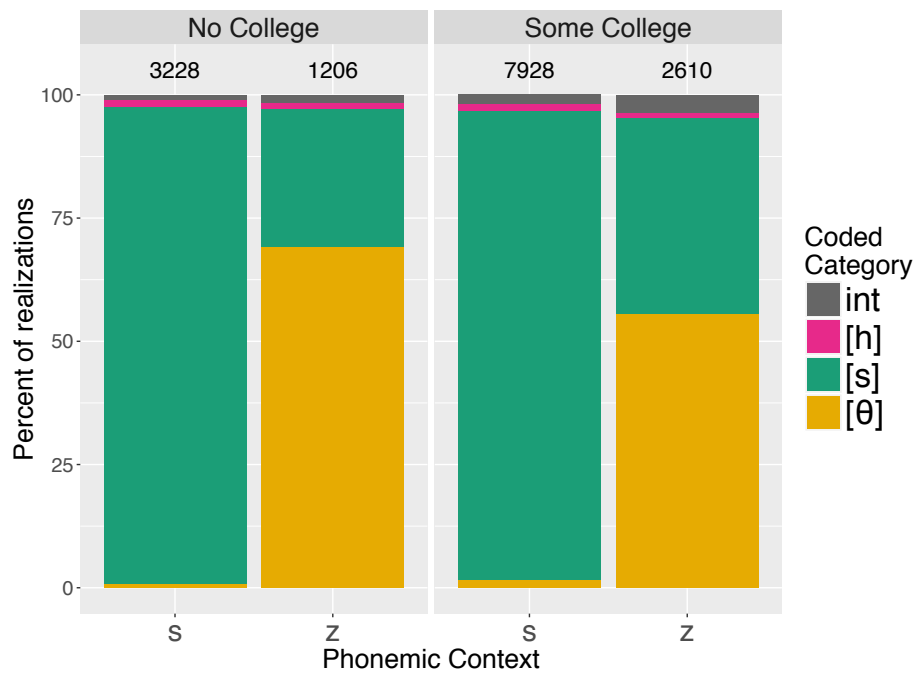


Figure 4.4: Percent segments in the spontaneous speech data coded [θ] in Z-contexts by phonemic context, faceted by the education level of the speaker.

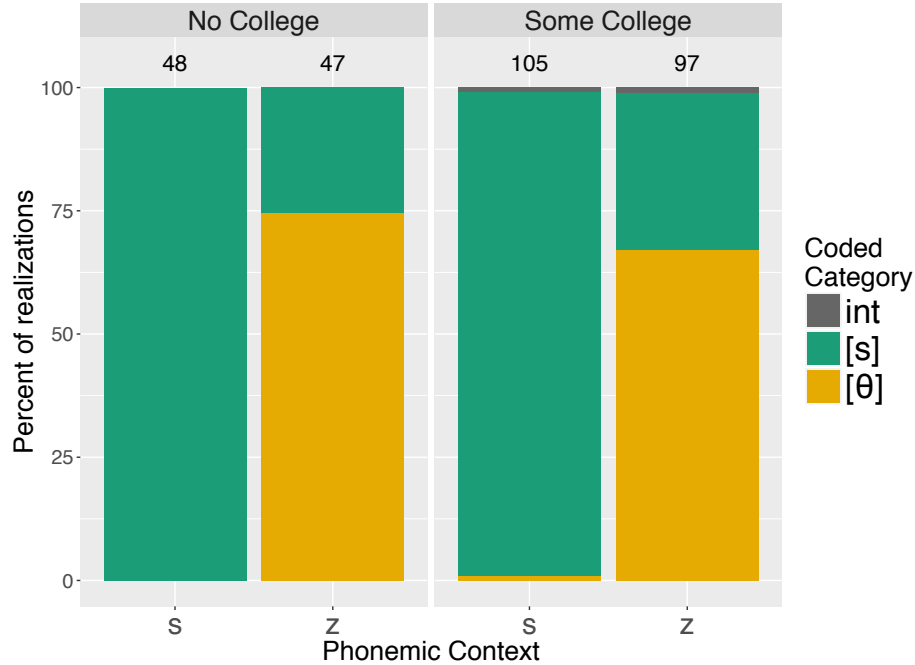


Figure 4.5: Percent segments in the elicited data coded [θ] in Z-contexts in the elicited portion by etymological context, faceted by the education level of the speaker.

or Low; see Chapter 3). Internal effects included Position in the word and Frequency (from the Vega et al. 2011 corpus).<sup>26</sup> The full model failed to converge; fixed effects were added in one-by-one in a step-up fashion. Any significant factors were retained, as well as any factors that did not raise the AIC criterion.

The final model shows significant effects of Language Profession ( $\beta = 2.35$ ,  $p < 0.05$ ) and of Ceceo Contacts ( $\beta = 3.87$ ,  $p < 0.00$ ). That is, both having a language-related professional focus and having contact with speakers from neighboring *ceceo* towns were associated with *less* [θ] in Z-contexts (more *seseo*). Age, neighborhood income, gender, word frequency, and word position were not significant factors.

The same model was run on the S-contexts, and it shows a different picture. Gender was a significant factor ( $\beta = -1.54$ ,  $p < 0.004$ ); men were less likely to produce [s] in S-contexts (more *ceceo*). Those who dealt professionally with language were significantly more likely

<sup>26</sup>Previous studies (Moya Corral and Wiedemann 1995; Villena-Ponsoda 2007; Regan 2017a; Regan 2017b) have found an effect of the presence of another anterior fricative in the word. This factor was not included here, but it is a planned addition for future analysis.

to produce [θ] in S-contexts ( $\beta = -1.329$ ,  $p < 0.039$ ).

|                            | Estimate | Std. Error | z-value | p-value |
|----------------------------|----------|------------|---------|---------|
| (Intercept)                | 0.07     | 1.11548    | 0.058   | 0.953   |
| Log Frequency              | 0.11     | 0.06126    | 1.727   | 0.084   |
| Median Neighborhood Income | 0.28     | 0.52718    | 0.537   | 0.591   |
| Position: Start            | 0.18     | 0.17452    | 1.044   | 0.269   |
| Language Profession: Yes   | 2.36     | 1.18       | 1.998   | 0.045*  |
| Education: Some College    | -2.14738 | 1.15       | -1.864  | 0.062   |
| Ceceo Contacts: Yes        | 3.86720  | 1.25       | 3.088   | 0.000** |
| Gender: Male               | 1.12097  | 0.98       | 1.141   | 0.254   |

Table 4.3: A linear mixed-effects logistic regression model predicting realization in Z-contexts (standard realization = 1).

|                            | Estimate | Std. Error | z-value | p-value    |
|----------------------------|----------|------------|---------|------------|
| (Intercept)                | 5.653    | 0.638      | 8.860   | < 0.00 *** |
| Log Frequency              | -0.146   | 0.059      | -2.472  | 0.013 *    |
| Median Neighborhood Income | -0.269   | 0.273      | -0.985  | 0.324      |
| Position: Start            | 0.322    | 0.136      | 2.361   | 0.018 *    |
| Language Profession: Yes   | -1.329   | 0.644      | -2.063  | 0.039 *    |
| Education: Some College    | 0.278    | 0.622      | 0.447   | 0.654      |
| Ceceo Contacts: Yes        | -0.789   | 0.663      | -1.190  | 0.233      |
| Gender: Male               | -1.541   | 0.535      | -2.877  | 0.004 **   |

Table 4.4: A linear mixed-effects logistic regression model predicting realization in S-contexts (standard realization = 1).

Since the question of interest is whether these speakers possess a full-fledged phonological contrast, it will be instructive to examine the data at the level of the speaker, in addition to the aggregate. We have observed the sharp difference in the overall rate at which this group of speakers produce a standard phonetic realization in the two etymological contexts: 65.8% standardly realized Z-words vs. 98.7% standardly realized S-words. However, the left bar in Figure 4.1 could represent a group of speakers who each produce Z-words with [θ] at a rate around 65%, or it could a group made up of two halves, one made of up speakers with a rate of of 100% and the other a rate of 30% [θ]. This is why it is important to look at individual production.

Figures 4.6 and 4.8 show by-speaker rates of the standard fricative in S- and Z-contexts, respectively. Intermediate tokens are included in both of these figures (that is, they are

essentially double-counted for the purposes of these two visuals) and are represented by a semi-transparent gray portion on top of the bar. The arrangement of the subjects on the x-axis is by their rate of [θ] in Z-contexts, from least on the left to most on the right. This ordering is used in all subsequent figures that display data by subject. The letter (M or F) preceding the subject ID number denotes the participant’s gender. Note that these genders are distributed throughout the ordered x-axis, reflected the lack of a clear gender-based trend in the realization of Z-words. The large divergence in use of the standard fricative across two contexts is again readily apparent in comparing the two figures. Figure 4.6 shows that there is a large range in the rate of standard [θ] by speaker: from categorical *seseo* on the left (F24) to categorical standard production on the right (M25). Most speakers are neither; this again is in line with Santana 2016b’s findings that most speakers in Seville today appear to be what she calls “two-solution” speakers, mixing standard [θ] and local [s] in Z-contexts.

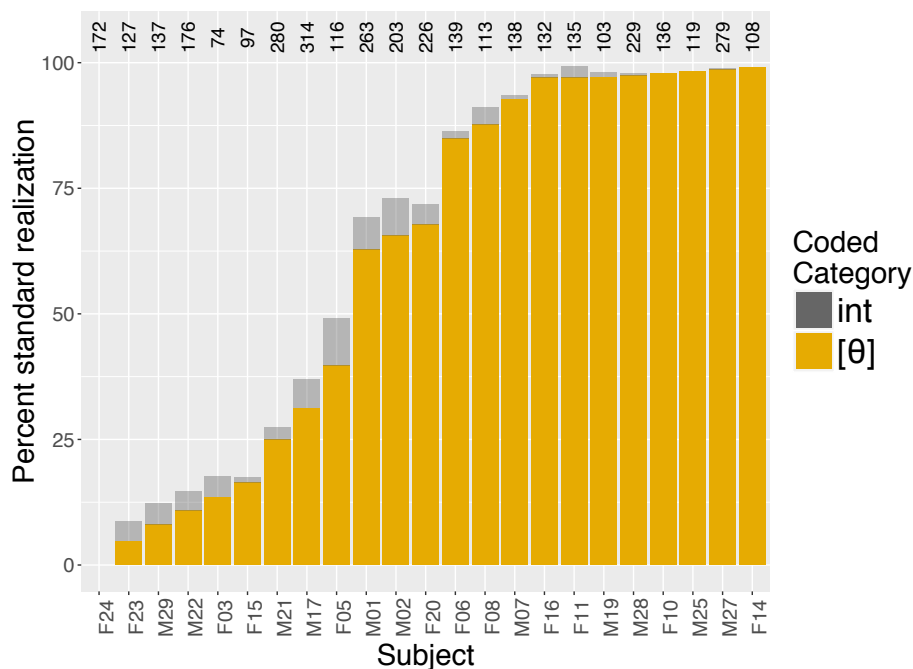


Figure 4.6: Percent tokens coded [θ] in Z-contexts overall by speaker, with [x] realizations excluded and the percent of intermediate tokens shown in a transparent gray.

A subset of this data is shown in Figure 4.7: these are tokens from the sentence elicitation



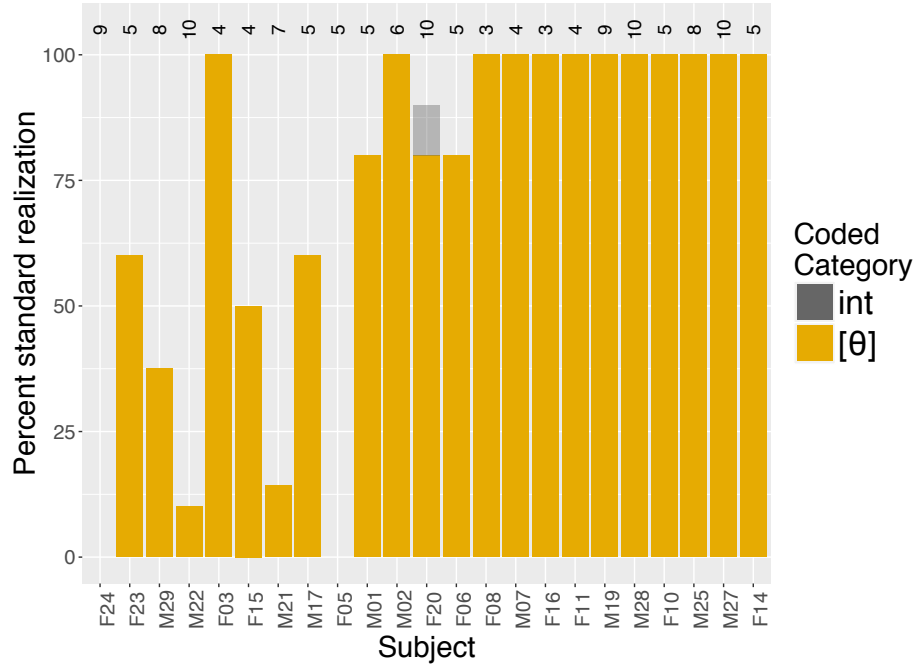


Figure 4.7: Percent elicited tokens coded [θ] in Z-contexts by speaker, with [x] realizations excluded and the percent of realizations as intermediate shown in a transparent gray.

task. When it is compared to Figure 4.6, Figure 4.7 highlights how a handful of elicited data (such as might be acquired with a recorded same/different task) can be misrepresentative of a speaker’s actual production. For the speakers who exhibit very few merged ([s]) tokens of /θ/ in speech, elicited tokens tended to be standard as well. However, for the speakers exhibiting variability, the elicited tokens are simply too few to give an accurate picture. Even if we assume that speakers will produce about the same rate of *seseo* during an elicitation task as during spontaneous speech, which not a good assumption to begin with, there is too much noise when a rate of e.g. 15% *seseo* is sampled only five times. Examining only the elicited data would mean overestimating the number of speakers with categorically merged or unmerged production.

There is nevertheless a generalization to be made, which is that speakers whose /θ/ production was closest to categorical maintained that pattern in the elicitation task: F24 produces categorical [s] across the board, regardless of context, and speakers F16, F11, M19, M12, M28, F10, M25, M27, and F14 produce nearly categorical [θ] in the interview and

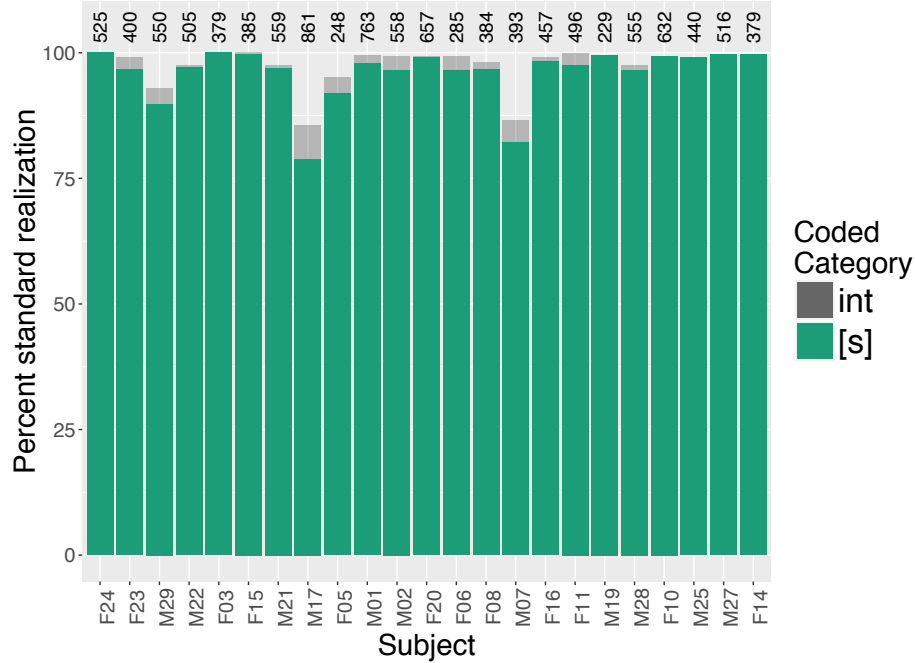


Figure 4.8: Percent tokens coded [s] in S-contexts overall by speaker, with [x] realizations excluded and the percent of intermediate tokens shown in a transparent gray.

categorical [θ] in the elicitation task.

Turning to Figure 4.8, we see that rate of use of the standard [s] in S-contexts is very high, ranging from 88 to 100% (excluding tokens coded as [x] or Intermediate). Only 4 out of 24 speakers have a rate of [θ] for in S-contexts of over 1%. Here we can see that the bulk of the nonstandard [θ] tokens come from three speakers (M29, M17, and F05), while more than half of the speaker overall have 1 or fewer tokens of [θ] in S-contexts. In the elicitation task, only F05 produces any S-word tokens coded as [θ] (Figure 4.9).

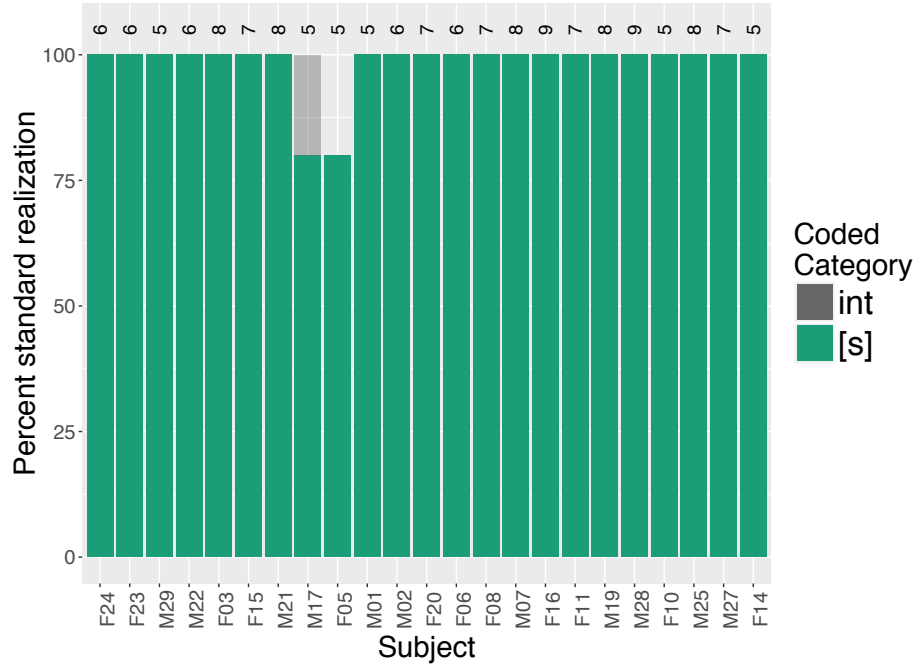


Figure 4.9: Percent elicited tokens coded [s] in S-contexts by speaker, with [x] realizations excluded and the percent of intermediate tokens shown in a transparent gray.

#### 4.3.4 Lexical distribution

As discussed above, the lexical distribution of realizations within each speaker is an important piece of evidence to the status of the sounds in their inventory. Figure 4.10 plots the rate of standard realization by each speaker of the five most frequent lexical items containing standard /θ/ in the corpus (each lexical item appears multiple times in the plot). What this figure shows is that the production data does not support the possible scenario referenced in Section 4.1.2, that of apparent variability resulting from regional (or idiolectal) differences in the lexical distribution of /s/ and /θ/. Rather, there is significant intraspeaker variability within lexical item. The speakers are again arranged from most *seseante* on the left to most standard on the right, and the speakers on those edges tend to be more consistent within lexical item. Here again we must consider the size of the sample: if the rate at which a speaker produces a variant is 90%, many tokens of the same word will be needed to see any variability in that word. For the speakers that have an intermediate rate of *seseo*, the

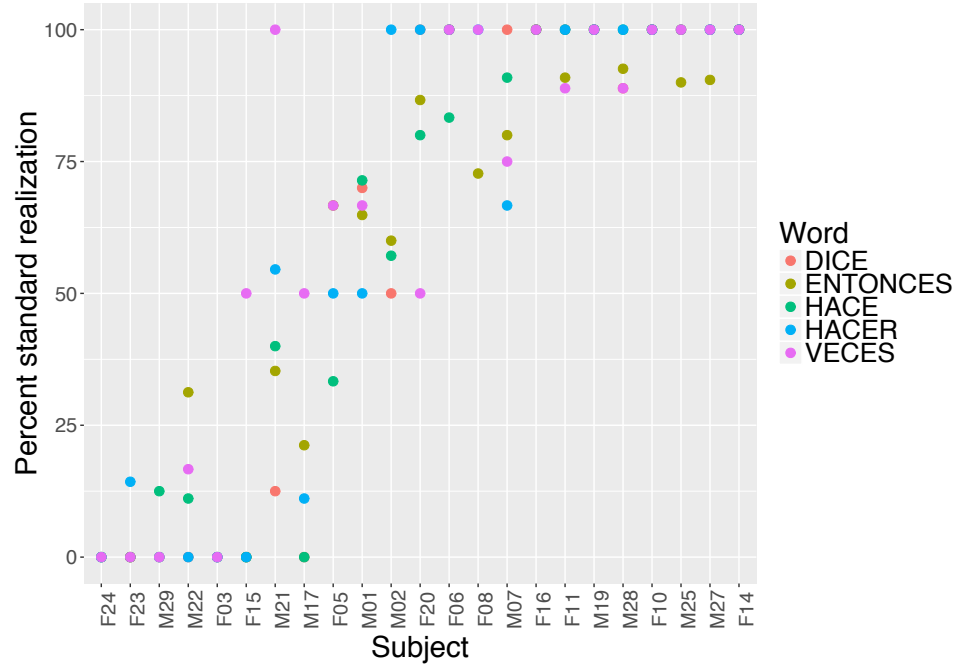


Figure 4.10: Percent tokens coded [θ] by speaker of the five most frequent lexical items containing /θ/: *dice* ‘say’, *entonces* ‘then’, *hace* ‘do’, *hacer* ‘to do’, *veces* ‘times’ (as in ‘sometimes’).

lexical variability comes out nicely in this figure: many of the points toward the center of the x-axis are also located in the middle of the y-axis, signaling that the speaker produced that word with multiple variants of the fricative.

The lack of a frequency effect in the production of Z-words in Table 4.3 is visualized in Figure 4.11: there is no clear upward trend as word frequency increases. Frequent words are often the first words to undergo change (Hooper 1976), and we might have expected an effect. There is, however, a significant main effect of frequency in the model predicting production in Z-contexts ( $\beta = -0.146$ ,  $p < 0.013$ ), and this is reflected in a perceptible upward trend in Figure 4.12 when the words that are produced with 100% [s] by a given speaker (i.e., the bulk of the data) are ignored.

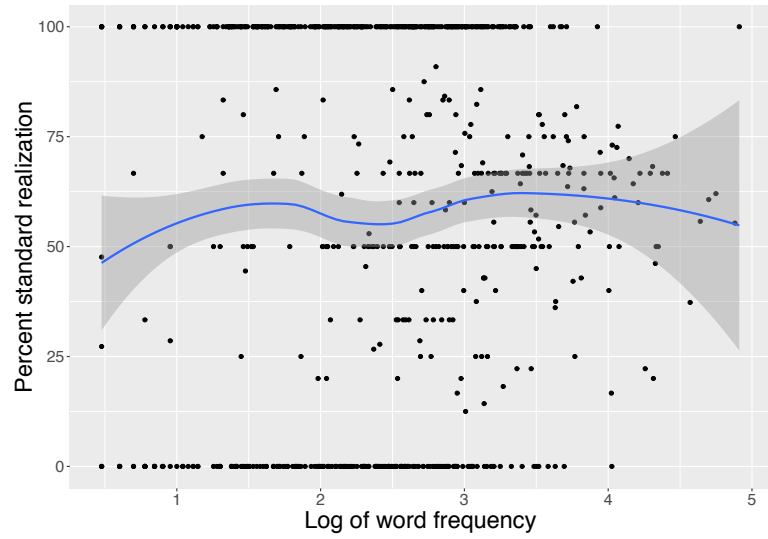


Figure 4.11: Percent [θ] in Z-contexts for each lexical item (each word is represented by one point) plotted by frequency in the Subtlex-ESP corpus, shown with a fitted loess line.

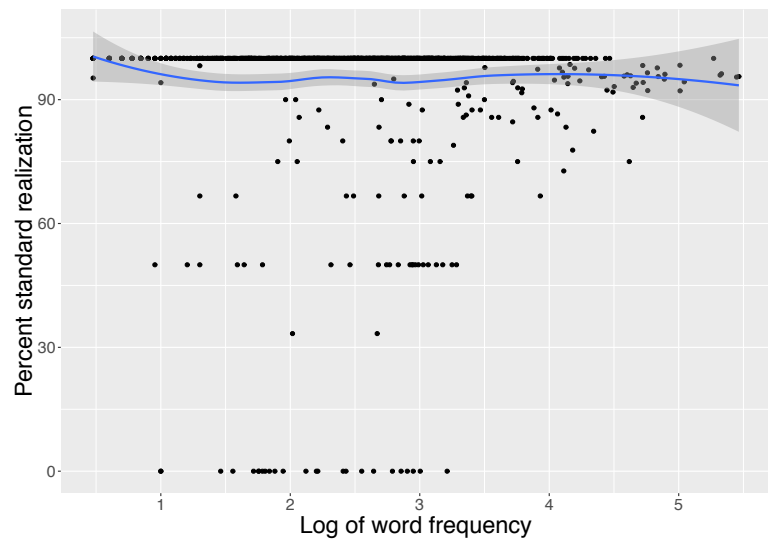


Figure 4.12: Percent [s] in S-contexts for each lexical item (each word is represented by one point) plotted by frequency in the Subtlex-ESP corpus, shown with a fitted loess line.

### 4.3.5 Acoustic analysis

I follow Lasarte Cervantes (2010)’s study of *ceceo* in Málaga in comparing the acoustic characteristics of the fricatives as grouped by impressionistic coding, rather than by etymological context. The reason is that the intraspeaker variability in employment of the contrast is such that averaging acoustic characteristic across /θ/-context means averaging over both [s] and [θ] tokens, thereby obscuring the acoustic cues that constitute the contrast (when present). To give an analogy, analyzing a group of naturally occurring tokens of /θ/ spoken by a Londoner who exhibits TH-fronting would not be the best way to determine what the acoustic properties distinguish [f] from [θ] (or indeed from any other sound) in English.

The goal of the acoustic analysis is twofold: first, I aim to contribute to understanding the acoustic correlates of /s/ and /θ/ as produced by this population. My second goal is to make an objective assessment of the reality of the contrast that I perceived during the auditory coding. If the two categories are at least somewhat separated along some acoustic dimensions, I can be more confident that the impressionistic coding was picking up on a systematic acoustic contrast.

I follow Regan (2017b) in the general procedure that I used for measuring the fricative tokens, utilizing a Praat script from Elvira (2014). It takes a series of measurements of the fricative from a filtered Hamming window, discarding tokens shorter than 12 msec in duration. It calculates the four spectral moments (spectral mean, standard deviation, skewness and kurtosis) representing the spectrum’s central tendency, dispersion, tilt, and peakedness, respectively. It also measures both mean and maximum intensity, maximum frequency, and center of gravity. As stated above, tokens were not hand aligned. Those that were marked as severely misaligned during the auditory coding process were excluded, but some misalignment persisted in the data. I used an acoustic measurement of voicing (the fraction of unvoiced frames, from the “Voice report” function) to attempt to filter out the more misaligned tokens, with the assumption that the fricative proper will generally be voiceless. This eliminates some of the noise introduced by other segments being included

in the measurement window. Tokens with a proportion of voiceless frames of less than 0.80 were discarded (36% of the data).

#### 4.3.6 Acoustic results

The two measurements that best separated the means of the two coded categories [s] and [θ] were center of gravity and mean intensity, consistent with Regan 2017b's findings. Figure 4.13 shows each token along axes of these two measurements, plotted by subject. For the subject with few tokens of [θ], separation is difficult to judge. However, most of the most standard speakers show two roughly separable clouds; speakers M27, M19, M01, F16, and M19 in particular have two discernible distributions.

Since the human ear is likely to be picking up on a number of acoustic cues in perceiving the [s]-[θ] contrast, combining different acoustic measurements has the potential to yield better results. The difficulty is in deciding which measurements to include. Using Principal Components Analysis (PCA), that decision can be handed off to a statistical procedure that breaks down a set of variables into uncorrelated components. Figure 4.14 again plots tokens for each speaker, this time by the first two Principal Components. The result is an increased, but still generally incomplete, separation of [s] and [θ], suggesting that there is something to be gained by considering more than two acoustic dimensions.

What I present here is a preliminary acoustic analysis with clear limitations. However, it lends some objective support to the idea that the Seville speakers have more than one acoustic target. What should we make of the lack of a clean separation between categories? There is some previous evidence that speakers exhibiting merged production of /s/ and /θ/ do not separate their articulations as robustly. Lasarte Cervantes (2010), examining read speech, found that the acoustic differences between tokens coded auditorily as /s/ and /θ/ were not as robust in the rural speakers that exhibited *ceceo* as they were in the urban/educated speakers that exhibited invariable phonemic contrast. The existence of intermediate targets in addition to two separate [s] and [θ] targets is not something that I can rule out, and I leave this possibility for future research. The is, to my best of knowledge,

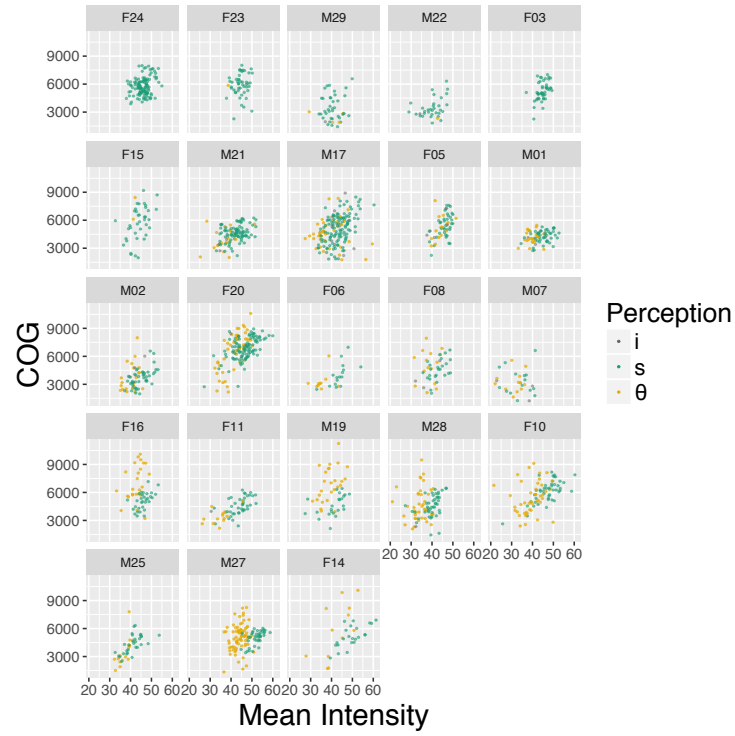


Figure 4.13: Tokens plotted by center of gravity value and mean intensity and colored by auditorily coded value..

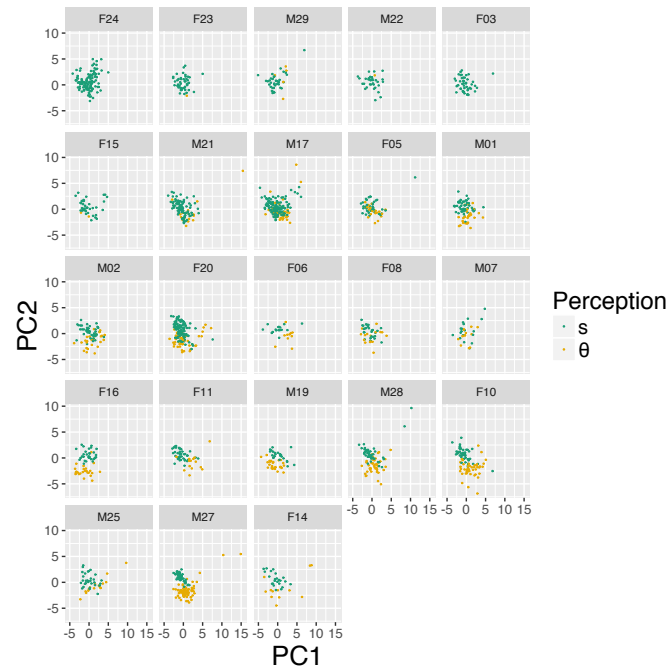


Figure 4.14: Tokens plotted by the first two principal components and colored by auditorily coded value.



the first acoustic analysis of /s/-/θ/ variation taken from a corpus of spontaneous, rather than read, speech. Interview speech is inherently more messy than read speech. This factor as well as the lack of hand-correction may be influencing the results.

#### 4.3.7 Speaker awareness of *seseo*

To what degree are Seville speakers aware of their own use of *seseo*? Figure 4.15 plots the rate of [θ] in Z-contexts by the response that the subject gave to a question from the identity questionnaire:

|                         |                                                |
|-------------------------|------------------------------------------------|
| Te consideras:          | ‘Do you consider yourself’                     |
| Algo seseante           | ‘Somewhat <i>seseante</i> ’                    |
| Muy seseante            | ‘Very <i>seseante</i> ’                        |
| Ni ceceante ni seseante | ‘Neither <i>ceceante</i> nor <i>seseante</i> ’ |
| Muy ceceante            | ‘Very <i>ceceante</i> ’                        |
| Seseante y ceceante     | ‘ <i>Seseante</i> and <i>ceceante</i> ’        |

The result is not a random picture: those that described themselves as “very” *seseante* indeed had low rates of standard realizations of Z-words (and therefore high rates of *seseo*), and the reverse is true of those that described themselves as being neither *seseante* nor *ceceante*. The production of those who described themselves as “somewhat” *seseantes* spreads across the entire range. It might be tempting to conclude from this plot that speakers that are closer to being categorically merged or unmerged are likelier to be aware of their own use, but an alternative interpretation is possible: that a certain proportion of speakers are fairly aware of their own *seseo* rate, and that those that are not aware *know* that they are not aware and chose the in-between response (lacking an option for ‘I don’t know’).

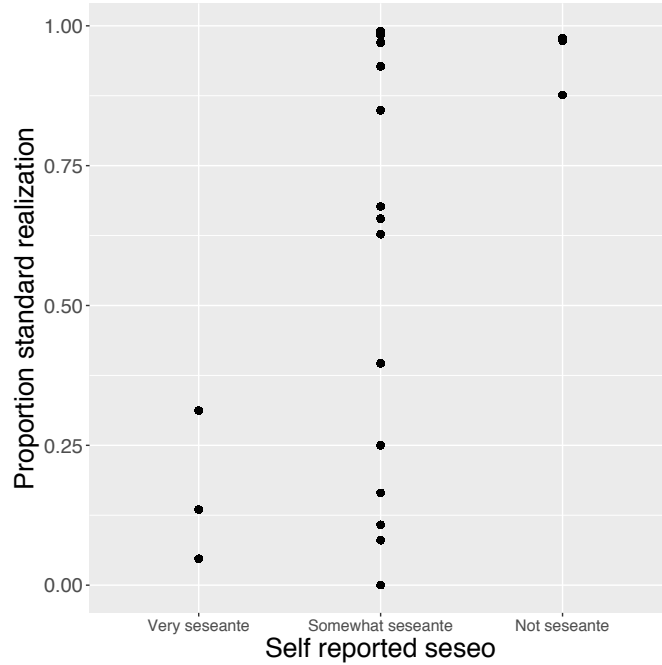


Figure 4.15: Proportion segments coded [θ] in /θ/-contexts by self-reported *seseo* rate.

#### 4.3.8 Discussion

The lack of social stratification of the use of *seseo* along any of the expected lines is puzzling. The absence of an age effect may simply be due to the small amount of age differentiation in the sample. However, the lack of gender education effects is surprising, and is a departure from Santana (2016b)’s findings. The influence of contact with speakers from *ceceo* areas could have been predicted to have an effect in either direction: either toward or away from the standard contrast. Contact with the merger could be predicted to be associated with more merged production. However, contact with *ceceo* means more exposure to [θ] productions in Z-contexts as well. This may be leading speakers to increased awareness of their own *seseo*, and to a subsequent bias away from it. The bias of language professionals toward rather than away from *ceceo* is also difficult to explain.

## 4.4 Identification and discrimination

This section turns to these speaker’s perception, presenting the results of two perceptual experiments. I will preface it with a discussion of a number of examples of near merger which demonstrate why it is desirable to conduct this kind of experiment in this case.

### 4.4.1 Near merger

Near-merger is a term applied to cases in which researchers observe mismatched behaviors (in the sense introduced above) in the context of a merger. These demonstrate how perception and production can come apart in interesting ways.

The merger of /ɑ/ and /ɔ/ in North American English is a widespread and ongoing process (Labov, Ash, et al. 2006), referred to variously as the LOT-THOUGHT merger, the COT-CAUGHT merger, or “low-back” merger. One of the earliest documented cases of this kind of “near-merger” is Labov, Yaeger, et al. (1972)’s description of Dan Jones in Albuquerque, New Mexico, where they report an ongoing merger of /ul/ and /ʊl/. In a verbal “minimal pair test”, in which the speaker is asked to pronounce relevant pairs and then say aloud whether they are the same or different, Dan Jones labeled several pairs of words containing /ul/ and /ʊl/ as being the same, and pronounced them the same as well. However, when he was recorded pronouncing several tokens each of *fool* and *full* for another task, as well as in an interview context, an acoustic analysis found a small but consistent difference in the second formant (Labov, Yaeger, et al. 1972:236–242). The same text describes another individual from near Harrisburg, Pennsylvania named Bill Peters who showed a similar pattern in his production and perception of /ɑ/ and /ɔ/. (Labov, Yaeger, et al. 1972:235–236).

These individual cases were received with some skepticism by the linguistic community, but the phenomenon of speakers producing contrasts that they profess to be unaware of has subsequently been documented in many individuals and in some large scale studies. Many of these involve the low-back merger. Decades after Labov’s study in Pennsylvania, Herold

(1990) visited a nearby area in her study of the same merger in coal mining towns of eastern Pennsylvania and found a small number of speakers “impossible to classify” as distinct or merged” (Herold 1990:27). Four elderly speakers from Tamaqua who were distinct in their production showed evidence of a “Bill Peters”-type effect in their responses on a minimal pair task (Herold 1990:185). Di Paolo 1992 examined the low-back merger in Utah. She conducted a vowel categorization task and recorded a word list with 122 participants, and found widespread lack of awareness of the contrast along with significant differences in F1 and F2 in production. Around the same time, Di Paolo and Faber (1990) examined the apparent merger of /i-ɪ, e-ɛ, u-ʊ/ before tautosyllabic /l/ in the same area and again found evidence for an acoustic contrast that listeners were not aware of, although this time the acoustic cue seemed to be phonation, rather than in the expected dimensions of F1-F2.

These studies so far primarily rely on same/different tasks, but near-merger has also been observed using more formal experimental methods. An early example of the latter involves the contrast between the short vowels /e/ and /ɛ/ in Swedish.<sup>27</sup> These sounds are generally distinct in the north and merged in Stockholm, but in a town called Lycksele in northern Sweden they seem to be neither: speakers from this town in the 1980s were reported to produce a distinction, but were not able to reliably identify synthesized tokens of these vowels (Janson and Schulman 1983) any better than merged Stockholmers. The authors draw a parallel between this “non-distinctive” feature in the Lycksele dialect and allophonic contrasts: both involve reliable acoustic differences that the speaker is unaware of, with the difference of course being that allophones are predictably distributed. Similarly, Yu (2007)’s study of Hong Kong Cantonese found a significant acoustic difference in F0 between two tones that are usually described as being the same: words with lexical mid-rising tone and words whose mid-rising tone is the result of a morphological process called *pinjam*, Native subjects nevertheless performed at chance at identifying the (near-)minimal pairs.

A phenomenon related to near-merger is that of incomplete neutralization. This term refers to a phonemic distinction that seems to be neutralized (i.e., not phonetically realized)

---

<sup>27</sup>Note that in the long vowels, this contrast was robustly maintained in all of Sweden at the time.

in a given phonological context, but for which a phonetic cue is found to reliability exist. A classic example is the contrast between American English /t/ and /d/ in contexts where they are subject to flapping; a number of studies have found that preceding vowel length reflects the underlying voicing status of [r] (Patterson and Connine 2001; Braver 2014). Braver (2014) found that listeners were nevertheless unable to use this information to help them in acoustic identification tasks. Other contrasts that are said to be incompletely neutralized, such as underlying final voicing in languages with final obstruent devoicing like German and Catalan, have been shown to be perceptible to a certain extent (Warner et al. 2004; Kleber et al. 2010). However, performance in these experiments is far from perfect and highly variable, subject to acoustic influence, talker-specific, and utterance-specific effects. The small effect sizes found in production studies of incompletely neutralized contrasts have led to skepticism about their reality (e.g., Manaster Ramer 1996), and the debate about incomplete neutralization is ongoing (see e.g., Roettger et al. 2014).

So far I have discussed cases a contrast in production is not accompanied by the expected degree of awareness and perceptual identification ability. The reverse effect has also been documented: Johnson (2010)’s large-scale study of the low-back merger in two towns in New England describes children of non-merged parents who are able to distinguish and even accurately imitate the contrast using appropriate lexical items, but who produce no reliable contrast in normal speech (Johnson 2010:153, 164). Studies of vowel mergers in New Zealand have yielded similar effects. One is the ongoing merger of the diphthongs /iə/ and /eə/ (the NEAR-SQUARE merger). Hay, Warren, et al. (2006) found that listeners who showed the /iə/-/eə/ merger in recorded word pairs nevertheless performed well above chance at an auditory forced-choice identification task of tokens produced by distinct talkers, although distinct listeners performed better. Babel et al. (2013), investigating the same merger, found that some speakers produced more separated tokens of /iə/ and /eə/ after shadowing an unmerged Australian talker as compared to a baseline reading task. Finally, in a study examining another New Zealand merger, the ELLEN-ALLEN merger, or the conditional merger of /ɛ/ and /æ/ before /l/, Hay, Drager, and B. Thomas (2013) found

a puzzling lack of correlation between participants’ production of tokens containing these vowels in a reading task and their performance a forced-choice identification task.

Some have suggested that these phenomena are an effect of orthography, and indeed some studies have found that durational differences in incompletely neutralized contrasts increase when attention is drawn to orthography (e.g., Fourakis and G. K. Iverson 1984; Jassem and Richter 1989). However, orthography cannot be the source of all of these observations. A few studies of incomplete neutralization explicitly control for possible orthographic influence and still find durational differences (e.g., Warner et al. 2004). Some of the vowel mergers above, like the /ɑ/-/ɔ/ merger, are inconsistently reflected in orthography. Finally, Yu (2007)’s study found near-merger phenomena for a contrast that is completely absent from the written language. Standard Chinese orthography does not represent tonal information, and so it certainly cannot explain the effect.

In sum, speakers that can be classified as merged based on one behavior often simultaneously exhibit behaviors that are characteristic of contrast. This finding is robust and repeated across many contrasts undergoing change. Furthermore, there is not one pattern of mismatch, but rather several.

#### **4.4.2 Mergers and methodology**

As seen in the previous section, there have been many production studies of a few mergers. Some of these (Johnson 2010) have included questions or tasks that probe metalinguistic awareness of the contrast in question. In the sociolinguistic literature, this has traditionally been referred to as a ‘perception’ task, but note that the meaning of the term here is not in line with the use of this term in phonetics and psycholinguistics, where it is generally applied to tasks that involve auditory presentation of external stimuli rather than meta-judgments of a subject’s own production. Gordon (2013) refers to minimal pair tasks and commutation tests as “fairly blunt instruments that explore perception in artificial contexts” (p. 2017). Orthography can influence performance in metalinguistic tasks involving mergers. When researchers have included non-homograph homophone pairs (WHOLE-HOLE, PAWS-

PAUSE) as fillers in written same/different questionnaires, a significant fraction (10-20%) of participants rate them as being “Different” (Gordon 2013; Johnson 2010:218).

Examples of near-merger like those of Dan Jones and Bill Peters show that naturalistic speech is crucial in establishing a realistic picture of speaker production, including the production of phoneme contrasts. If researchers only observe maintenance of the contrast in question in careful speech or minimal pair-type tasks, they cannot be sure that the contrast is present in the speaker’s normal speech - and vice versa.

It is worth mentioning here that 24 is not as many subjects as would be ideal for the experimental portions of this study, both the ones in this chapter and for the experiment in Chapter 5. However, this shortcoming must be balanced against the practical constraints of recruiting from a specific demographic, as well as the labor it takes to record, transcribe, and analyze spontaneous speech from each of the subjects, and the value inherent in having both kinds of data from the same speakers.

In order to investigate their ability to perceive the relevant contrast and its boundary, the participants were asked to complete a forced-choice identification task and an AX discrimination task, following the methodology in Amengual (2016).

### 4.4.3 Materials and procedure

A binary forced-choice identification task was designed to test the participants’ ability to label the sounds [s] and [θ], and an AX discrimination task tested their ability to tell the sounds apart. These two experiments were similar in design, and their results are reported together.

The stimuli for these experiments were created by taking a recording of the word *atún* ‘tuna’ and splicing in a fricative token, as follows: Two fricatives were extracted from recordings of the word *azúcar* ‘sugar’, one produced with [s] and one with [θ]. The intensity of these two word tokens was matched, and then the fricative was spliced from each at a zero crossing. The duration of these two fricatives was manipulated to match (116 ms), and they were then mixed (their intensities scaled and the resulting sounds layered on top of

|      |    |    |    |    |    |    |
|------|----|----|----|----|----|----|
| S1   | S2 | S3 | S4 | S5 | S6 | S7 |
| 100% | 66 | 50 | 33 | 25 | 10 | 0  |

Table 4.5: Percentage [s] in fricative steps; following Eisner and McQueen (2005).

one another) using Praat. The proportions of each fricative making up each stimulus step are in Table 4.4.3. These 7 steps were embedded into the *atún* frame to create 7 versions of a pseudoword *asún*.

For the AX task, 12 pairs were created, with an ISI of 1000 ms: 7 consisted of each step paired with itself (the Same condition), and the remaining 5 (the Different condition) consisted of all of the possible combinations such that the pair were 2 steps apart (S1-S3, S2-S4, S3-S5, S4-S6, and S5-S7). Once the subject made a response, the next pair began playing 1000 ms later.

The stimuli for all of the experiments conducted as a part of this dissertation were recorded by the same talker, a 30-year-old male native of Seville, over two sessions. The recordings were conducted in a soundproof booth at the University of Seville with an Ritmix RR-600 digital recorder (bitrate 128 kb/sec) and an external microphone.

The experiments were conducted in a quiet room on a 2012 MacBook Pro using PsychoPy experimental software and KRK systems KNS-8400 over-the-ear headphones. For the identification task, participants were instructed verbally (in Spanish) that they would be listening to sounds and deciding what they were. Written instructions told the participants that they would be hearing a consonant between vowels. They were asked to press the ‘Z’ key if the consonant was a  $\theta$ , and the ‘S’ key if the sound was an S. Participants responded to 77 stimuli: 7 practice stimuli (one full block of stimuli) + 70 randomized test stimuli (7 stimuli \* 10 repetitions). In total there were 7 stimuli \* 10 repetitions x 24 participants which resulted in 1680 responses.

The AX task was conducted immediately after the identification experiment. Participants were told they would be listening to a pseudoword again, but that now they had to decide if a pair of sounds were the same or not. Written instructions told the participants that they would be hearing sounds in pairs consisting of a consonant between two vowels,



and instructed them to press one of two keys according to whether the sounds in the pair were the “Same”, and the other if they were “Different”. They were presented in 8 blocks consisting of the “same” pairs in order from 1–7, followed by the ‘different pairs’ in order from most to least [s]-like.<sup>28</sup> There were 12 trials x 8 repetitions x 25 participants, which resulted in a total of 2375 responses.

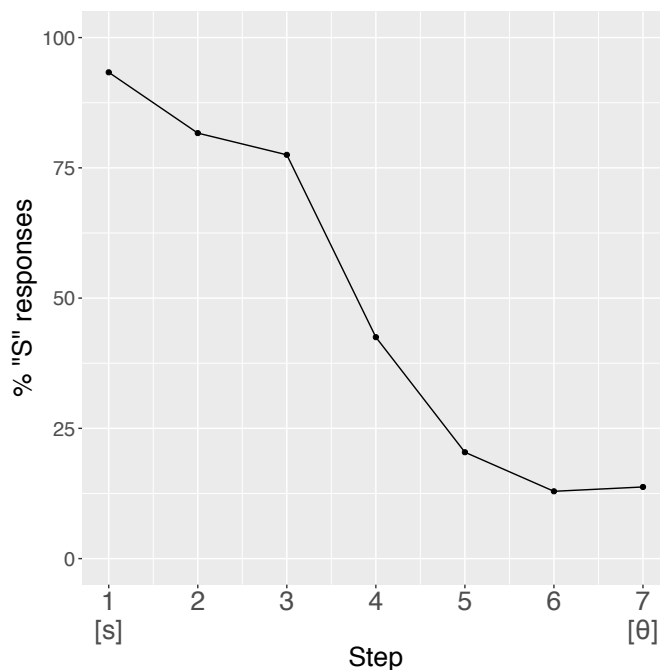


Figure 4.16: Percent “S” responses overall by continuum step.

#### 4.4.4 Results

In the identification task, overall mean accuracy at the continuum endpoints was 90%. A mixed-effects logistic regression model (Table 4.6) predicting response with a random intercept by subject had a significant main effect of stimulus ( $\beta = -0.155$ ,  $p = 0.00$ ). The results indicate that subjects were generally able to identify tokens /s/ and /θ/, and suggest that there is an acoustic boundary between [s] and [θ] for this population (see Figure 4.16).

Results by subject are shown in Figure 4.17. This figure suggests that all but two subjects, Subjects M1 and M21, exhibit a categorical response pattern. These subjects had

<sup>28</sup>The intent was to randomize within each 12-trial block, but this was not done due to an error in coding.

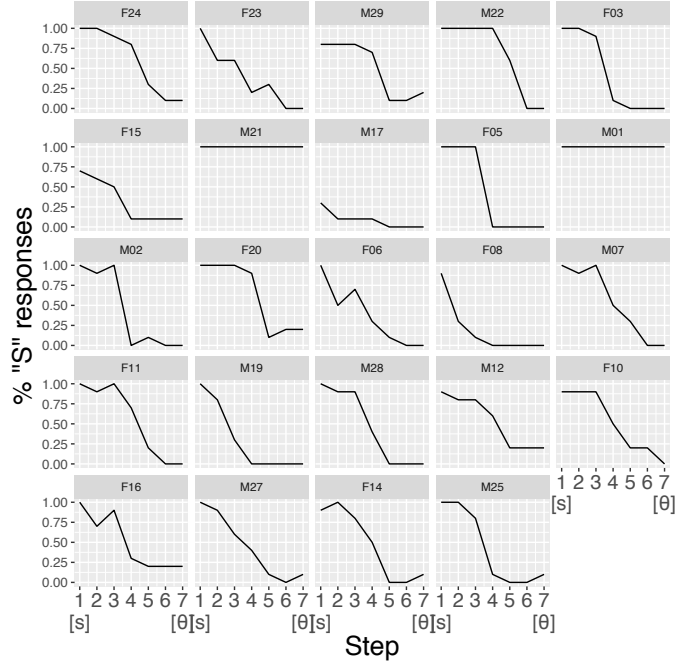


Figure 4.17: Percent “S” responses by continuum step by subject.

|              | Estimate | Std. Error | df  | t-value | p-value |
|--------------|----------|------------|-----|---------|---------|
| (Intercept)  | 1.099    | 0.0857     | 362 | 12.82   | 0.00*** |
| Step         | -0.155   | 0.0041     | 165 | -37.49  | 0.00*** |
| ISI          | 0.000    | 0.0001     | 165 | 0.98    | 0.327   |
| Trial Number | -0.001   | 0.0004     | 165 | -4.23   | 0.00*** |

Table 4.6: A linear mixed-effects logistic regression model predicting response (“S” = 1, “Z” = 0).

a *seseo* rate of 14% and 26%, respectively. Accurate identifications at the endpoints rises to 93.4% when these two subjects are excluded.

For the AX task, a mixed-effects regression model (Table 4.7 predicting response with a random intercept by subject had a significant main effect for Pairs S1-S3 ( $\beta = 0.12$ ,  $p = 0.00$ ), S2-S4 ( $\beta = 0.37$ ,  $p = 0.00$ ), S3-S5 ( $\beta = 0.30$ ,  $p = 0.00$ ), and S4-S6 ( $\beta = .0905$ ,  $p = 0.008$ ). These steps were informative to the contrast for these subjects, but correct responses to Different pairs close to the boundary were around 50% (see Figure 4.18).

By-subject discrimination of the Different pairs is shown in Figure 4.20. Responses are perceptibly less uniform than in the identification task. A few subjects (F05, F08,

|              | Estimate | Std. Error | df   | t-value | p-value  |
|--------------|----------|------------|------|---------|----------|
| (Intercept)  | .1726    | .0673      | 2247 | 2.562   | 0.010*   |
| Pair 1-3     | .1206    | .0339      | 2675 | 3.549   | 0.000*** |
| Pair 2-2     | -.0222   | .0348      | 2675 | -0.636  | 0.525    |
| Pair 2-4     | .3780    | .0339      | 2675 | 11.127  | 0.000*** |
| Pair 3-3     | -3.752   | .0348      | 2675 | -1.077  | 0.282    |
| Pair 3-5     | .3034    | .0339      | 2675 | 8.930   | 0.000*** |
| Pair 4-4     | .0138    | .0348      | 2675 | 0.396   | 0.692    |
| Pair 4-6     | .0905    | .0339      | 2675 | 2.665   | 0.008**  |
| Pair 5-5     | -.0377   | .0348      | 2675 | -1.084  | 0.278    |
| Pair 5-7     | -.0096   | .0339      | 2675 | -0.283  | 0.777    |
| Pair 6-6     | -.0559   | .0348      | 2675 | -1.605  | 0.109    |
| Pair 7-7     | -.0746   | .0348      | 2675 | -2.139  | 0.033*   |
| ISI          | -.0001   | .0001      | 2683 | -0.833  | 0.405    |
| Trial Number | .00015   | .0002      | 2675 | 0.685   | 0.493    |

Table 4.7: A linear mixed-effects logistic regression model predicting response (“Different” = 1, “Same” = 0).

F11, M28) respond “Same” overwhelmingly to different pairs taken from the edges of the continuum, but show a sharp rise in “Different” responses for pairs 2-4 and 3-5. Others show a less pronounced rise, or an erratic pattern (F15, F10, M19). Of the two subjects with a flat identification curve, M01 shows a discrimination pattern consistent with a failure to perceive the contrast, while M21 shows increased “Different” responses to S2-S4 and S3-S5.

#### 4.4.5 Discussion

The results from these two tasks are consistent with a general ability to distinguish [s] and [θ] in perception. The participants identified tokens of [s] and [θ] at a rate of 90%. This accuracy rate is not at ceiling, but it is in the range of rates found for the endpoints of continua in experiments showing categorical perception. It is tempting to attribute the inconsistent by-subject results to phonemic merger of the subjects. However, most studies of categorical perception show only aggregate results. It is hard to know how many published studies of categorical phonemic perception included speakers who failed to exhibit the effect, because this data is simply not reported.

A potential explanation for the AX task results is that the set of Different pairs did not

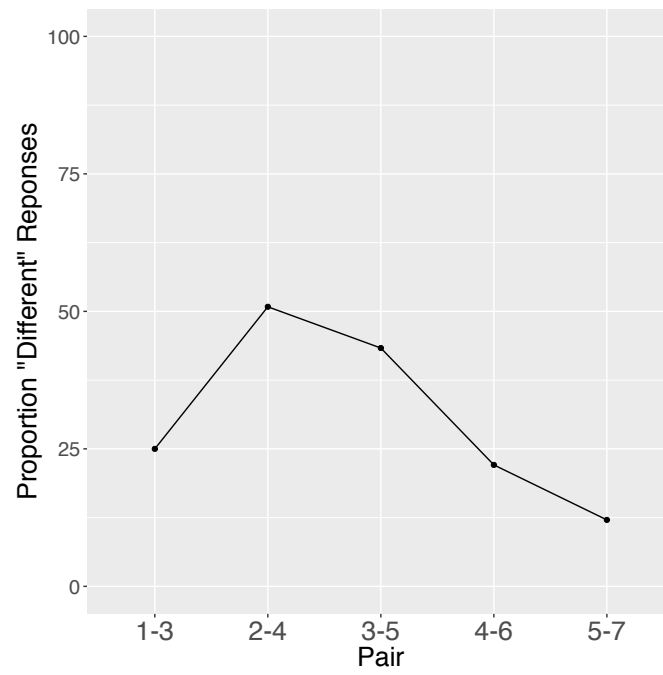


Figure 4.18: Percent “Different” responses to Different pairs.

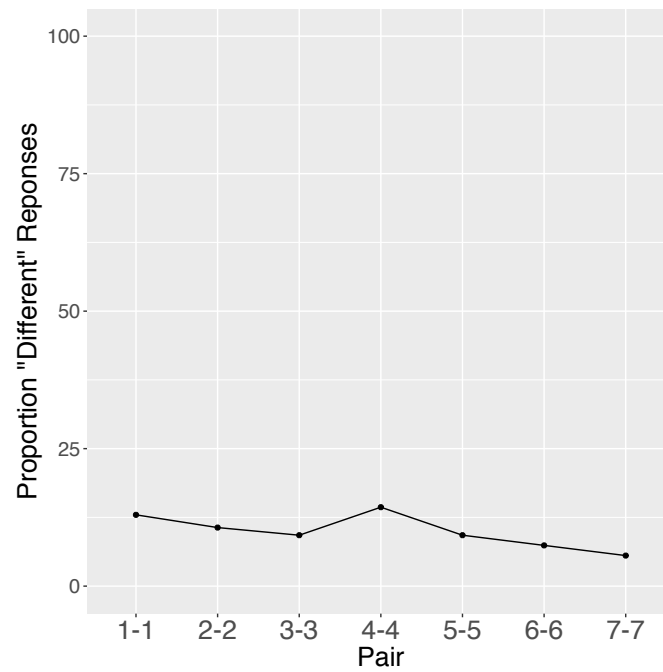


Figure 4.19: Percent “Different” responses to Same pairs.

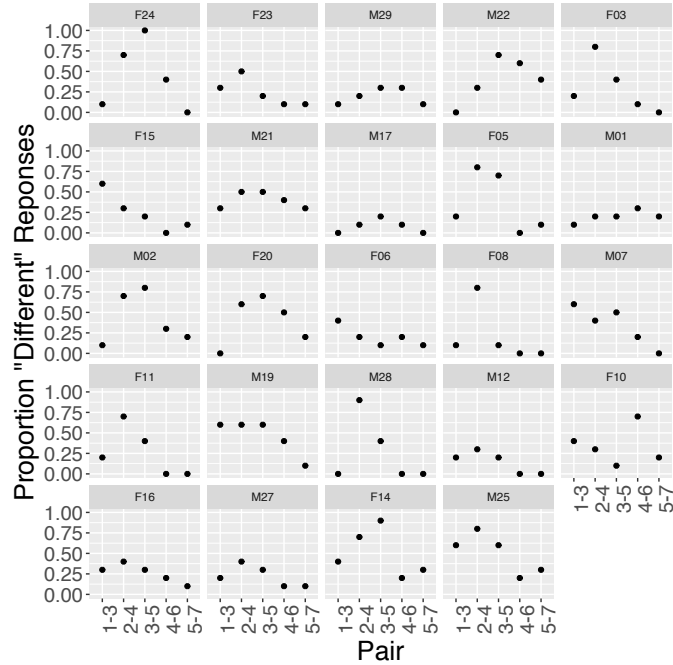


Figure 4.20: Percent “Different” responses to Different pairs.

include a pair whose two members were located firmly on either side of the boundary. These two tasks share the limitation of being performed on a single fricative continuum embedded in a single frame, leaving open the possibility that the patterns found are a product of some property of the specific stimuli.

## 4.5 General discussion

This contrast departs from the vowel contrasts that dominate the sociolinguistic literature in a number of ways. Some are discussed at length in Chapters 2 and 3, and I will give a summary here. Firstly, it is generally not reported that speakers struggle to hear the difference between [s] and [θ] or are in any way unaware of the variability. This might be rooted in the sheer acoustic distinctiveness of [s] and [θ] compared to most vowel contrasts involved in merger, and it is perhaps this property, in combination with a long and rich history of geographic variation, that has allowed this merger to achieve the remarkable social prominence that it seems to enjoy. Canonical examples like the /ɑ/-/ɔ/ merger

operate mostly below the level of consciousness. A handful of speakers in contact with a pattern different from their own have been documented commenting on an unexpected phonetic value in another’s speech or on the occasional misunderstanding, but this is a very limited, surface-level awareness at best (Labov 1994:344). The situation of /s/-/θ/ in Spain is very different. There, lay people have access to a specific lexical label for not only the standard lack of merger (*distinción*), but also two labels for merger with different phonetic values (*seseo* and *ceceo*). It has been previously suggested that this possibly exceptional social awareness, along with orthographic support, that has facilitated the borrowing of the standard pattern over and over in different areas of Andalusia. The sociolinguistic situation therefore qualifies as the type of “intense social evaluation” that allows a phonemic category to be borrowed through contact (Labov 1994:346).

If we accept this premise, what can we make of the absence of social conditioning on the production of the /s/-/θ/ contrast in this data? Consider that despite the intense awareness of this contrast as a sociolinguistic variable, it is not necessarily clear that it is overtly prestigious for these speakers. In Chapter 3, I report that they rate *seseo* as being highly acceptable and “correct”, and that most of them respond that they do not change their speech at all when talking to someone from another area of Spain. 21 out of the 24 participants said that no one had ever explicitly corrected their use of *seseo*. In the debriefing conversation, participants give impassioned defenses of *seseo* along with other local features. So although standard Castilian Spanish certainly still exerts plenty of influence, the strengthening of the Andalusian regional identity discussed in Chapter 2 has perhaps led to a restrengthening of the “Seville norm”. At the same time, increased mobility and education seem to have had their effects as evidenced by the change apparently in progress. Some participants recognized that certain contexts exert a pressure to make the contrast.

Although the /s/-/θ/ contrast seems exceptional in that the phonological facts rise to the level of awareness, at the end of the day it is the phonetic value of [θ] that attracts stigma: *ceceo* is much more stigmatized than *seseo*. Participants gave *seseo* an average

acceptability rating of 6 out of 7 (compared to 4.5 for *ceceo*).

The relative prestige of *distinción* and *seseo* is more difficult to get a handle on, and this is partly because prestige is inherently flexible and can be assigned differentially in different places (Milroy and Gordon 2008), and in different situations: “linguistic variables that are prestigious or stigmatized in the abstract are full of positive meaning in the concrete everyday” (Eckert 2000:226). It may be that the opposing forces of local and national prestige balance each other out, in a sense, and that women are not particularly motivated to avoid *seseo* (at least, any more). It appears that these forces exert an influence on individual use of *seseo* in ways that are more complex than macrosocial categories or questions about Seville stereotypes are able to capture.

We can describe many of the speakers in this study, then, as individuals who have access to the contrast, but who do not to always make use of it. This might seem less surprising when we draw a parallel to other socially-meaningful features that show probabilistic variation. For example, English /ɹ/-vocalization, the absence of syllable-final /ɹ/, occurs variably in many speakers (Wells 1982). It, is in a certain sense, a merger to a null realization. The contrast is not maintained in this contexts some of the time; but we do not attribute conditional merger to these speakers. Many speakers in the south of England exhibit variability in what is called “TH-fronting” (the production of /θ/ as [f], Kerswill and Williams 2002), but they are not thought to have a merger of these phonemes. A parallel can also be drawn to grammar competition (Kroch 1989), usually applied to syntactic changes. Competing grammars describes the situation in which usually exclusive parameter settings coexist within the grammar of a single speaker. In this case, we can think of the phonological system of contrast between the phonemes /s/ and /θ/ coexisting with a system of lack of contrast in the Seville speakers.

## 4.6 Conclusion

This chapter presents behavioral evidence that elucidates the phonological systems of young speakers of Seville Spanish. Recall the four characteristics associated with having two

phonological categories from the introduction:

1. having phonetically different sound distributions for the two categories;
2. reliably producing the sound from the appropriate distribution according to lexical context
3. perceiving the contrast in their own and others' speech; and
4. using the difference to access lexical meanings.

I sum up the results in this chapter as giving evidence on three of these four criteria. Firstly, there is evidence that the Seville speakers do have (at least) two targets that separate in acoustic space, if imperfectly. Their productions can be sorted by ear as one type or the other with relative ease. There is also evidence that the third criterion is satisfied; the Seville speakers can identify the two sounds in question, and they are generally aware of this phoneme's absence or presence in their own speech.<sup>29</sup> For the second point, the speakers show evidence of lexical systematicity. I argue that this evidence, taken together, suggests a phonological contrast in the systems of these speakers. The fourth and final criterion is left for the next chapter.

---

<sup>29</sup>This was not explicitly asked of every speaker, but no speaker reacted to questions about *seseo*, etc. with confusion; rather, they also seemed very acquainted with the terms used to discuss the phenomenon.



## Chapter 5

# Mental representation

### 5.1 Introduction

Chapter 4 established that the young Seville speakers that are the subjects of this study are able to identify instances of [s] and the recently borrowed sound [θ], and that they produce both sounds in their speech in a way that reflects knowledge of their standard lexical distribution. However, there is an aspect of perception that the methodology in Chapter 4 does not investigate, and that is the involvement of the /s/-/θ/ contrast in real-time lexical access. This chapter examines how the speakers' auditory recognition of recorded Z-words (words with etymological /θ/) is affected by the identity of the fricative in the recorded word ([s] or [θ]). This is a separate question from whether or not speakers are able to simply tell the fricatives apart, and I will argue that it is an equally, if not more, informative one.

To begin with, it provides an additional window onto the question of whether or not the speakers' phonologies contain two separate categories, /s/ and /θ/. Theoretical linguistics is mainly concerned with implicit knowledge of language, not with the cognitive processes of perception and production *per se*. The field uses speaker intuitions to gather facts from a wide variety of languages with which to construct theories (Boland 2017). However, as discussed in Chapter 1, there are special challenges involved with describing the phonologies of speakers participating in an ongoing phonological change. When describing stable

phonemic contrasts, the linguistic researcher can simply ask for speakers' intuitions and briefly observe their production. They can expect these to line up and can assume that perception of the contrast at all levels will proceed accordingly. In cases of ongoing change, such methods will fall short of capturing the empirical facts. This can be seen from the observations made in Chapter 4: after brief observation of the speakers in this study, a researcher might conclude from the variation in their production of Z-words that they do not have a standard /s/-/θ/ contrast and are therefore merged. However, they could conduct an identification task and conclude that the speakers are not merged. These conclusions would be equally premature. Identifying sounds on command is not a natural linguistic activity, and possessing this ability is not the same thing as having phonemes in one's grammar. The next section will discuss evidence from previous studies that a speaker's auditory discrimination of two sounds and their use of a phonemic contrast during the process of speech perception are in fact two different abilities, motivating the use of methods that more closely resemble naturalistic speech perception in order to determine whether there are two abstract categories in this case.

The dissociation of phonemic contrasts into different aspects of production and perception is relevant to those who study the psychological process of speech perception and its relationship to production. Psycholinguists have created tasks that specifically target the process of word recognition, and linguists may benefit by applying these to cases of linguistic variation. Studies of systematic variation in speech were noticeably sparse in the psycholinguistic literature until relatively recently (Bürki 2018). Even the recent set of studies that have begun to address this tend to focus on variation that is a result of internal linguistic factors, such as assimilation (Marslen-Wilson and Warren 1994), quick speech (Mitterer and McQueen 2009; Ranbom and Connine 2007), and conditional insertion phenomena (Tuinman et al. 2011), and not on purely socially and/or regionally-conditioned variation. In a recent review of psycholinguistic work on variation, unconditioned phoneme replacement (e.g., the variation between [s] and [θ] in Z-words shown in this study) does not even make it into the otherwise comprehensive typology of phonetic/phonological variation

phenomena (Bürki 2018). The move away from *seseo* in Seville is a case of social variation that can be informative for psycholinguistic research because the rich phonological variation involved creates the potential for unusual patterns of production and perception behavior. Languages are constantly changing, and the behavior of speakers in the middle of changes is something that we should seek to describe and understand as a part of the human capacity for language. Language change aside, the perception of nonstandard variants and how it compares to and interacts with the perception of canonical variants has rightly become an area of interest within psycholinguistics in the last few years (Boland et al. 2016; Sumner, Kim, et al. 2013). In this study, the case of the anterior fricatives of Seville, which I argue is an especially intriguing case, is added to the variable phenomena that have been studied in this way.

Finally, very little is known about the mechanism through which phonemic splits proceed, given their rarity and the scarcity of studies examining them. Work on the mechanisms of merger ascribes an important role to speech perception – not only the acoustic recognition of the relevant sounds, but to whether speakers use that information to interpret words in everyday life and the point at which they cease to do so (Herold 1990; Labov 1994; Yang 2009). By asking how and when speakers of a variety with a split that is in progress *begin* to use the acoustic information that is available to them for speech comprehension, we can learn something about how phonemic split proceeds, a question that is of keen interest to sociolinguists and historical linguists alike.

The experiment that will be described in this chapter explores how Z-words are represented in the mental lexicons of the individual speakers in this study. It consists of an auditory lexical decision task, a standard psycholinguistic task that will be introduced in the next section. I use this task to investigate whether Z-words are represented with [s], with [θ], or both. To answer this question, I compare how well the speakers access Z-word meanings when they hear a Z-word pronounced with either [s] or with [θ] by comparing the semantic priming that is induced by the two forms. The results will show a processing advantage for [s] pronunciations.

## 5.2 Decomposing the perception of a contrast

Chapter 4 showed that Seville speakers differentiate two acoustic categories, [s] and [θ], in perception. This finding might lead to the assumption that their mental lexicon reflects this differentiation. However, examples from second language acquisition and the literatures dealing with mergers and bilingualism indicate that differentiation of acoustic categories and subsequent differentiation within the mental lexicon are related but separate properties.

and the literatures dealing with mergers and bilingualism

Speakers can learn to discriminate new phonemic contrasts in learning a second language (L2). Depending on the acoustic confusability of the sounds and how well they map onto contrasts in the L1, they can do this quite quickly. For example, Finnish has phonemic vowel length but lacks lax /ɪ/. Ylinen et al. (2010) trained Finnish English learners to improve their discrimination of /i/ and /ɪ/. The training improved accuracy overall and was successful in encouraging listeners to attend to vowel quality over duration.<sup>30</sup> ERP responses to stimuli after training were indistinguishable from those of native English speakers. However, it would be a mistake for the researchers to conclude that the Finnish listeners came away from the study with native-like English /i/ and /ɪ/ categories. They note that “more authentic exposure and practice with English is certainly needed” (Ylinen et al. 2010:1329). The training did not teach the Finnish English learners which English words contain which vowel, a crucial aspect of a native-like contrast.

In the same vein, Ota et al. (2009) found evidence of weak separation of lexical entries for pseudo-homophones for second language learners: Japanese-speaking learners (who have difficulties with the /r-/l/ contrast) had more false positives and slower reaction times when deciding whether unrelated pairs of words like ROCK-KEY were related as compared to control pairs. The words were orthographically presented, so there was no question of the initial sound being misperceived. Rather, the authors interpret the results to mean that ROCK and LOCK (the latter being the competitor that is actually related to KEY) are not

---

<sup>30</sup>Vowel length is not a dependable cue for what is traditionally called the tense/lax distinction in American English (Hillenbrand et al. 2000).

well separated in the lexicons of the Japanese English learners because the forms to which the meanings are mapped are similar, if not identical. There is evidence that even words that do not contain problematic phonemes (i.e., whose sound categories are also present in the learner's L1) appear to have “fuzzy” form-to-meaning mappings in an L2. Cook et al. (2016) found that Russian learners of English were slower at deciding whether an English word is a good translation of a Russian word when the target word was more similar to its implied competitor, unlike a control group of native English-speaking learners of Russian. The differences between the target word and the competitor included contrasts present in Russian, and so perception difficulties are unlikely to explain the result.

These examples show that mere discrimination of two sounds, or even the ability to map them onto phonemic categories, is not sufficient for speakers to possess a phoneme contrast in the way that we normally think of it. Although the examples come from late second language learners and may not be directly comparable to speakers acquiring a sound within the context of their native language, it is worth considering that prestige-driven changes within an L1 are sometimes acquired late, at school-age or even later (Sankoff 2004), and the acquisition of new phonemes in an L1 and an L2 may in some respects be similar. In any case, this phenomenon has also been documented in a few studies of speakers who have a contrast that is marginal in some way, not because they are non-native speakers but because of language contact or change (Amengual 2016; Labov 1994). With such speakers, creative methods must be used in order to evaluate whether or not the contrast is fully present in those speakers. This includes determining whether the phonemic categories are fully linked to the appropriate items in the lexicon. At this point, it will be useful to briefly introduce some psycholinguistic methodology that has been used to investigate the lexical aspect of phonemic contrasts.

### **5.2.1 Investigating mental representations with experimental methods**

Speech perception is the psychological process that converts auditory signals to meaning. Behavioral results from experimental tasks that require subjects to perceive speech can be

used to give evidence for how the mental grammar is structured. Lexical decision is one of the most widely used tasks in psycholinguistics (Goldinger 1996; Cutler 2012). The term refers to a task in which subjects are asked to make a “lexical decision” on a given stimulus; that is, to decide whether the stimulus is a real word or not. The stimuli consist of both words and pseudowords, presented in isolation or preceded by priming or other contextual stimuli. They can be presented visually (written on a screen), auditorily (played from a recording), or both (“cross-modal” presentation). Lexical decision implicitly requires full lexical processing, and this has made it a useful paradigm for exploring what factors affect the success and speed of lexical access. Lexical decision experiments were used to demonstrate the influence of word frequency and neighborhood density on lexical processing, and other priming phenomena have been demonstrated with the paradigm. Priming can be defined as speeded, or facilitated, recognition of a stimulus after a previous stimulus has caused it to be activated. This previous stimulus can be identical to the primed stimulus (“repetition priming”, Slowiaczek and Pisoni 1986), phonologically or morphologically similar (Radeau et al. 1995), or semantically related or associated (Radeau 1983).

Several examples of research that use psycholinguistic methods to investigate a possible phonemic contrast come from studies of Catalan-Spanish bilinguals. This population provides an attractive case study for marginal phonemic contrast because of a phonological difference between Catalan and Spanish: front and back open-mid vowels contrast with their close-mid counterparts in Catalan but not in Spanish, which has only the close-mid vowels. These two phonemes have been the subject of extensive psycholinguistic investigation, taking advantage of the large pool of Spanish-Catalan bilinguals. Virtually all speakers of Catalan are bilingual to some degree, and there exists a spectrum from Catalan-dominant, early-acquisition bilinguals to Spanish-dominant, late-acquisition bilinguals (formal schooling is conducted primarily in Catalan). By comparing speakers at different points along this spectrum, researchers have been able to observe different patterns of the /e/-/ɛ/ and /o/-/ɔ/ contrasts through production and perception behaviors. Pallier et al. (1999), using a lexical decision task with subjects from Barcelona, found repetition priming of minimal

pairs differing in the stressed vowel (e.g., [ˈse.βa] ‘onion’ and [ˈse.βa] ‘his/hers/its’), but only for Spanish-dominant bilinguals. Catalan-dominant listeners appeared to process the minimal pairs as different words and did not exhibit repetition priming. This is evidence that the Spanish-dominant listeners do not have fully separate representations of these phonemes. The paper does not give direct evidence on whether the speakers can detect the acoustic difference, but states that the contrasts are emphasized in school. Amengual (2016), looking at the same phenomenon in Majorca, compared Catalan-dominant and Spanish-dominant bilinguals’ performance on both low-level acoustic tasks and lexical tasks. In discrimination and categorization tasks, the performance of both groups was at ceiling, suggesting that neither group has trouble hearing the the front and back mid vowel contrasts. However, in the lexical decision task, the Spanish-dominant listeners had higher error rates responding to words and pseudowords containing those vowels. That is, although they had no difficulty discriminating the close-mid and open-mid contrasts, and plenty of exposure to them, they apparently had difficulty discarding both real-word competitors of pseudowords and pseudoword minimal pair competitors of real words. This set of experiments speaks to a degree of independence between the ability to discriminate two sounds and the ability to use them for comprehension.

The same observation had been made by the 1970s through the study of phonemic merger. Labov distinguishes what he calls the “normal phonemic function – the use of a phonemic distinction to distinguish words in unreflecting interpretation” from the abilities that speakers showed in successfully completing same/different tasks and even commutation tasks (Labov 2001:403). He designed a task in 1976 to examine this ability in Philadelphian speakers with the MERRY-MURRAY (ɛɪ-ʌɪ) merger (Labov 1994). It consisted of an auditorily presented story in which a baseball coach puts a player in a game. According to the listeners’ interpretation of the string [mɛɪɪjənðɛɪ], the coach either played “Marion there” or “Murray in there”. The story was followed by a question to determine the listener’s interpretation of the vowel. 33% (7/21) of the Philadelphia speakers indicated that they had understood “Murray”, as compared with 7% (1/14) of a non-Philadelphian control

group. Four of the seven Philadelphians who responded “Murray” showed non-overlapping productions of the two vowels and responded “Different” to minimal pairs (Labov 1994: Table 14.8). This task is cleverly designed, but its main drawback is that a single observation (with chance accuracy 50%) from a small number of participants is very difficult to interpret. However, the key ideas, namely, that speech comprehension is both interesting to linguists studying merger and distinct from more metalinguistic tasks, are fully present in this early study.

### 5.2.2 Mental representations and merger mechanisms

Herold (1990) proposes that merger of phonemic categories arises when there is a decline in the usefulness of a phonemic contrast in speech comprehension. In her words, “a distinction ceases to be useful for making semantic distinctions when one is in contact with people who do not reliably produce it” (Herold 1990:92). She describes a scenario in which natives from town A who distinguish two phonemes X and Y come into extensive contact with natives from town B, who do not. Misunderstandings as a result of incorrect interpretations of words with either X or Y spoken by town B residents lead town A residents to “recategorize the distinction... as one which no longer allows them to distinguish lexical items in the speech of others, or to convey semantic information themselves.” (Herold 1990:94). Real-life misunderstandings provide some support for this idea. Examples collected as a part of the Cross-Dialectal Comprehension project (Labov 1989; Labov 2010) include 27 naturally occurring misunderstandings involving various American English vowel mergers. Labov (1994) points out that the large majority (23/27) involved a two-phoneme listener misunderstanding words produced by a one-phoneme talker. This is anecdotal evidence for the idea that misunderstandings by two-phoneme speakers drive merger by pushing them toward a single-phoneme style comprehension strategy: one that uses semantic, syntactic, and pragmatic information from the context to disambiguate minimal pairs, instead of acoustic information. There is experimental evidence that some kinds of acoustic information can be worse than no information at all when there are dialectal differences between



listener and talker: a series of gating experiments conducted as a part of the Cross-Dialectal Comprehension project (Labov 1989; Labov 2010) took recordings of advanced tokens of the Northern Cities vowel Shift, or NCS, (Labov, Ash, et al. 2006) from sociolinguistic interviews. These tokens were played for listeners to transcribe. The NCS is a vowel rotation that results in the phonetic quality of /a/ being closer to that of standard /æ/, /æ/ being closer to standard /e/, and so on. In this task, the correct response to the stimulus “You had to wear [sæks], no sandals” is SOCKS. Control subjects were given only the sentence frame written on paper with a blank for the key word. The acoustic quality of the vowel diverged so far from what the test subjects expected that they were less accurate in writing down the key word than the control subjects, although they had all the information the control subject had plus the acoustic information. Although the NCS is not a merger, the example serves to demonstrate why it might be beneficial for unmerged speakers to ignore what might seem like an informative acoustic contrast.

Yang (2009) takes the mechanism of misunderstandings and operationalizes it using the idea of relative fitness of grammars, analogous to the fitness that drives the natural selection of species. Under his model, there is a frequency threshold of failures given the input (misunderstandings) that a two-phoneme language learner (that is, a child) can tolerate before she abandons it for the one-phoneme grammar. This frequency is modulated both by the functional load of the phoneme contrast in question (the number of minimal pairs that the contrast disambiguates and their frequency in the lexicon, Martinet 1952) and by the density of one-phoneme speakers in the speech community. With this model, a tipping point can be calculated for a given phonemic contrast: the density of merged speakers in a community that, when reached, will be followed by merger of the whole community. Labov (1994) raises the question of whether two-system speakers are able to change their comprehension system according to their interlocutor, using a one-phoneme grammar to interpret speech from one-phoneme talkers and a two-phoneme grammar with two-phoneme talkers. Labov (1994) doubts the viability of such a solution, both because of how one-phoneme and two-phoneme speakers tend to be “intimately mixed” in the speech community

and as such indistinguishable, and because of the general lack of social consciousness of merger discussed in Chapter 1. We have already seen that the case of *seseo* is exceptional with regards to the latter point; whether there are also systematic cues that distinguish *seseantes* from *distinguidores* is an empirical question. In the discussion, I will return to the possibility of speakers having multiple solutions in perception mirroring the multiple solution pattern found in their production in Chapter 4.

The mechanisms of phonemic split are poorly understood, chiefly because of the scarcity of studies of phonemic splits in progress discussed in Chapter 1. It is relatively clear that the social prestige of the contrast plays a large role, and that exposure to a variety with the contrast is essential. Under what conditions is a phonemic split likely to persist? In the case of canonical phonemic contrast, speakers use distinguishing acoustic information to disambiguate between minimal pairs and dynamically narrow down lexical candidates while a word is being spoken (Cutler 2012). At what point do speakers of a variety with a split in progress use the acoustic information to guide lexical access in this way? Arguably, the presence or absence of this kind of processing will be an indicator of whether a contrast is robust and potentially stable. Determining whether this, the normal phonemic function, is present will necessarily involve methods that depend on lexical access.

### 5.2.3 Methodological considerations

Seville speakers must have plenty of experience with the /s/-/θ/ contrast or they could not possess the lexical knowledge and discrimination ability that we saw evidence for in Chapter 4. However, given the evidence above, it is not safe to assume that they use these for speech processing. Given the variability of production in Z-words that most of them exhibit, we must assume that their representations of these words are connected to both [s]-like and [θ]-like articulations. Are they likewise equally able to process an instance of a word with [s] or [θ] as an instance of a Z-word? Or do Seville speakers strictly map instances of [θ] to Z-words and instances of [s] to S-words as a speaker from Madrid would, meaning that a Z-word (with no corresponding minimal pair S-word) with [s] will be processed as a

pseudoword?

In this chapter I report a semantic priming experiment aimed at evaluating the mental representation of Z-words. Semantic priming can be defined as the speeded recognition of a lexical item after a semantically related lexical item has already been processed. This speedup in recognition is presumably a result of some degree of activation of the item through its conceptual representation, in turn linked with the conceptual representation of the first item. Whatever the exact nature of conceptual activation, semantic priming is a well-known finding in psycholinguistics. The current study is not about semantic relatedness, but rather, the established phenomenon of semantic priming here is leveraged as a tool to investigate the processing of phonological variation (drawing from Sumner and Samuel 2009, discussed below). The research question from above, whether Z-words are represented with [s] or with [θ], can be tested by measuring the amount of semantic priming of related words that is induced by Z-words that are produced with [s], as compared to Z-words produced with [θ]. This is done here through a lexical decision experiment with a paired priming design: during a critical trial, the participant hears a Z-word prime that is *either* produced with [s] *or* with [θ], followed 500 ms later by a (non-Z-word) target that is either related or unrelated to the Z-word. In this way, any speedup in recognition of the second word when preceded by a related word can be compared across the two test conditions: Z-word produced with [s] and Z-word produced with [θ].

The auditory modality was a clear choice for this study. Lexical decision experiments with visually presented stimuli are easy to design and execute and are ubiquitous in the psycholinguistic literature (Cutler 2012). While auditory stimuli are more labor intensive to create, evidence suggests that visual and auditory lexical decision do not necessarily tap into the same processes. Holcomb and Neville (1990) compared semantic priming in visual and auditory modalities in an ERP study. The distribution of the N400 effect was similar but different in its lateral distribution, indicating “non-identical” sources of semantic priming. Additionally, the visual modality is a poorer window into the phonological representations that are of interest in the case of *seseo*, and do not allow for direct examination

of phonological variation that is not reflected in the orthography.

In designing this experiment, an important consideration was that the two realizations [poso] and [poθo] of a word like *pozo*, differ in prestige: the merged variant [s] is regionally marked in Spain and less prestigious, and the unmerged [θ] represents the national standard, as discussed in Chapter 2. It is known that nonstandard variants in psycholinguistic experiments can create certain complications. Early psycholinguistic literature was preoccupied with the deviation from canonical form that is always present in the speech signal (due to acoustic, physiological, and environmental factors) and the difficulty it creates for the listener. From this perspective, it was assumed that non-canonical variants carry a processing cost. Researchers began to critically examine this idea a few decades ago by examining the processing of words containing phonological variants that are non-canonical in some way (e.g., Maye et al. 2008; McLennan et al. 2003; Floccia et al. 2006). There is some evidence that forms with predictable, regular variation (often lenition) are accessed as easily as their canonical counterparts (Sumner and Samuel 2005; Sumner, Kim, et al. 2013; Connine 2004). These and similar results (McLennan et al. 2003) has led Sumner, Kim, et al. (2013) to posit a dual-route model with stronger encoding of standard forms, due to their social salience, in order to explain this equivalence in access between frequent, nonstandard forms and less frequent standard forms.

What about variation that is not widespread or predictable, but rather socially and/or regionally conditioned? Only a handful of experiments have investigated the processing of this kind of variation using psycholinguistic methodology. Prominent among these is work by Meghan Sumner and colleagues on variation on rhoticity in English. The historic English rhotic phoneme /r/ varies in its reflex in postvocalic position: in General American, the rhoticity was maintained, but in southern British English varieties, postvocalic /r/ was lost<sup>31</sup>. The resulting non-rhotic pronunciation was carried to American East Coast port cities like New York in the 19th century. Today, the rhotic variant is gaining ground in New York, but the non-rhotic variant persists and has associations with an authentic, local

---

<sup>31</sup>Although there is evidence for its phonological presence: it reemerges in many non-rhotic varieties when immediately preceding a vowel-initial word, a phenomenon known as “linking /r/”.

identity (Becker 2009). Sumner’s work on the processing of non-rhotic tokens suggests that the subject’s degree of experience with a nonstandard variant is a key factor. Sumner and Samuel (2009) found that [slendə] (*slender*) primed *thin* equally well as [slendə̃] did for speakers from New York City (regardless of whether they used non-rhotic forms themselves). General American control subjects that did not have much exposure to non-rhotic speakers did not exhibit priming from non-rhotic forms. The New York City subjects’ own use of the variant did seem to have an effect on storage of the form in memory: for NYC speakers whose own production was rhotic, long-distance repetition priming only occurred with the standard variant, while the New Yorkers who used the non-rhotic form, the standard and nonstandard forms primed equally well at all distances.

Another study involving regional difference involving English /r/ gives similar evidence for the importance of exposure. “Intrusive /r/” is a phenomenon in non-rhotic varieties in which /r/ is inserted after a non-high vowel and before a vowel-initial word (Giegerich 1992; Cruttenden 1994). Tuinman et al. (2011) found, using a forced-choice comprehension task, that British English listeners use acoustic cues to distinguish intrusive /r/ from canonical /r/ but American and Dutch L2 speakers of English speakers do not.

Studies of cross-dialectal speech processing are still relatively few, and this may be partly because there are methodological barriers and uncertainties involved in the study of regional and social variation. Nonstandard speech in lexical decision tasks, for example presents a difficulty in that subjects can vary in their interpretations of “real word” and may feel inhibited from expressing overt wordhood judgments to stigmatized forms. In this case, the question of whether a stimulus, e.g. [poso], is a word is further complicated in that it is dependent on the talker; coming from a speaker from Madrid it is a pseudoword, but in Seville it is a word. For these reasons, I avoided asking subjects to make lexical decision to Z-words or to items containing [θ] by asking the participants to only make a lexical decision to the second member of the pair (the target), and not to the Z-word prime, drawing from one of the experiments in Sumner and Samuel (2009).

There is ample evidence that context is an essential consideration when it comes to

the processing of nonstandard variants. “Context” here can be defined at different levels. Sumner (2013) found that when a nonstandard form is spliced into a frame that previously contained the standard form, a processing cost is incurred as a result of the incongruity. Sumner attributes previous findings of cost for frequent nonstandard forms (e.g., Andruski et al. 1994) to this practice of cross-splicing reduced forms into unreduced frames. That is, listeners “expect” other low-level acoustic properties to co-occur with the nonstandard phone, and their absence in an experiment can interfere with processing, creating the illusion of a generalizable processing cost of nonstandard forms themselves. To account for this, Sumner, Kim, et al. (2013) include the interaction between the frequency of the frame and the frequency of the variant in their dual-route model.

Higher levels of context have also been demonstrated to have an effect on speech processing. A series of studies have shown effects of social information about the talker on the parsing of speech. This information can be external, such as indications of a talker’s place of origin in the form of stuffed toys associated with nationalities (Hay and Drager 2010), explicitly provided information (Niedzielski 1999), or faces shown along with the auditory stimulus (Staum Casasanto 2009; McGowan 2011; Strand 2000; Koops et al. 2008). The information can also come from the speech signal itself. Hay, Drager, and Warren (2010) found that distinct New Zealand listeners performed better at distinguishing between NEAR and SQUARE words in a written task when the experimenter spoke an unmerged dialect of English (Received Pronunciation or American English), but that merged listeners performed worse. In an ERP study, Hanulíková et al. (2012) found that Dutch sentences with grammatical gender errors produced a P600 (late posterior positivity) effect only when the talker was a native speaker of Dutch. The P600 is associated with processes of syntactic reanalysis (Osterhout and Holcomb 1992). This suggests that the listeners had different expectations with regards to grammatical well-formedness when listening to the L2 speakers, presumably triggered by other features in the signal.

The speech signal contains information about all kinds of social characteristics beyond just nativeness. There is evidence that listeners are able to identify regional identity from

speech alone (Clopper and Pisoni 2004). It is therefore not implausible that listeners are making assessments about regional and other identities of the talkers and generating expectations based on them. I conducted a pilot study (Gylfadottir 2016) in which two native Spanish speakers recorded single-word stimuli: one talker from Barcelona (where /s/ and /θ/ are distinguished in Spanish) and one talker from Mexico (where *seseo* is universal). The critical items were Z-words, which were always produced by both talkers with [θ]. Subjects performed lexical decisions on a set of stimuli that were repeated in two separate-talker blocks. The (Spanish, non-Andalusian) listeners were easily able to identify one talker as being Spanish and the other as Latin American without having been given explicit information about their origins. Their reaction times suggest that they were using this information in processing the stimuli: the listeners exhibited a significant slowdown when responding to Z-words produced by the Mexican talker (whose use of [θ] is at odds with his variety of Spanish), but not to those produced by the Spanish talker.

With these considerations in mind, the experiment was designed to create a context in which *seseo* would not be unexpected or incongruous. Coming from a standard-sounding talker from Castilla, *seseo* is as unexpected as standard [θ] is coming from a Mexican talker. Given the evidence above of listeners' sensitivity to context, if this study were to use a standard talker to evaluate whether *seseo* induces a processing cost, I argue that its conclusions would be questionable at best. After all, the question of interest is how the speakers in this study perceive [s] and [θ] in speech in their everyday lives, and not how they might perceive them from a standard-sounding speaker imitating a feature of their dialect. The talker for the experiment was therefore carefully chosen to be a speaker of the local dialect whose speech was variable with regards to the relevant contrast, so that *seseo* and *distinción* would both be congruous with the other characteristics of his speech, and so that he be able to produce both forms naturally. An effort was made to evoke a register or style in which *seseo* would be expected by asking the talker to use other regional phonological characteristics. Informal questioning of the participants' judgment of the talker's place of origin during the debriefing session suggests strongly that the participants identified him

Table 5.1: Example critical stimuli from Experiment 1.

|             | Related     |                | Unrelated      |                        |
|-------------|-------------|----------------|----------------|------------------------|
|             | Prime       | Target         | Prime          | Target                 |
| S-condition | <b>pozo</b> | <b>agua</b>    | <b>chorizo</b> | <b>escuela</b>         |
|             | [poso]      | [ayua]         | [tʃoriso]      | [ek <sup>h</sup> uela] |
|             | ‘well’      | ‘water’        | ‘sausage’      | ‘school’               |
| θ-condition | <b>zumo</b> | <b>naranja</b> | <b>belleza</b> | <b>arena</b>           |
|             | [θumo]      | [naranʝa]      | [βejeθa]       | [arena]                |
|             | ‘juice’     | ‘orange’       | ‘beauty’       | ‘sand’                 |

as being from Andalusia and from Seville province specifically, though responses varied in whether they pinpointed the city, referenced the province generally, or firmly believed he was from the countryside (possibly due to his use, at my request, of the stigmatized /l/ → [r] before obstruents feature).

### 5.3 Methods

#### 5.3.1 Design

Participants completed a paired lexical decision task with a response to the second member of the pair. In the critical items, the first member of the pair was always a Z-class item, and the second member was never a Z-class item. The experiment had a 2 x 2 design: the prime and target were either semantically Related or Unrelated, and the Z-class prime was assigned to either the S-condition or the θ-condition, that is, it was produced either with [s] or with [θ]. See Table 5.1 for examples of stimuli from each condition. The list was the same for all subjects and distributed in such a way that critical trials were non-adjacent. A single list was used rather than a between-subjects design, in order to standardize order and item effects and thereby minimize the sources of variation in the results across individual subjects. This was done with the goal of exploring potential patterns of individual differences in perception and production.



### 5.3.2 Stimuli

There were a total of 106 critical pairs of primes and targets, related and unrelated. Each of these was assigned to the S-condition or the  $\theta$ -condition as stated above. The four conditions were balanced by prime and target frequency to the largest degree possible. The critical pairs were all taken from two Spanish semantic association corpora (Luna et al. 2016; Callejas et al. 2003). Every word that contained (1) one or two instances of standard / $\theta$ / in an onset and (2) no instances of /s/ in an onset<sup>32</sup> was extracted from the corpora along with its commonly associated word. Pairs that involve a regionally limited word or in which the related target was also a Z-class word were discarded. Z-class primes that are minimal pairs with S words (e.g., CAZA) were avoided as much as possible<sup>33</sup>. Critical stimuli and real word fillers were all nouns. Half of the pairs of related words became critical primes and targets. To create the unrelated pairs, Z-words were paired with the related word from the next line, and the resulting pair was examined to make sure they were not very phonologically similar, in which case the next word was taken until an adequate pair was created. This procedure was continued until the two corpora were exhausted. During data collection, concerns were raised about whether the Related and Unrelated pairs were sufficiently different in degree of semantic relatedness. To address this concern, criteria were established before analysis of the full results began, and these were used to exclude a subset of the critical pairs after data collection was complete. These criteria were based on a norming study conducted on Amazon Mechanical Turk. Subjects were asked to rate the relatedness of the critical pairs on a scale of 1 to 7. Only responses from Turkers located in Spain who correctly answered a simple comprehension question were included, leaving 40 responses. Criteria were established such that all Unrelated pairs had a mean rating lower than the 25th percentile of the Related ratings, and all Related pairs had a mean rating higher than the 75th percentile of the Unrelated pairs.<sup>34</sup> The analysis reflects these

---

<sup>32</sup>Recall that both /s/ and / $\theta$ / are lenited in codas in this variety of Spanish.

<sup>33</sup>Despite best efforts, a few primes have a highly infrequent S counterpart.

<sup>34</sup>These criteria led to the exclusion of 32 of 138 original critical pairs. The excluded stimuli were well balanced across the four experimental conditions.

exclusions.

There were 411 filler pairs. Real word filler targets were all taken from the semantic association corpora. Critical targets and real word filler targets were balanced for mean and standard deviation of frequency, so that critical targets would not be more or less frequent overall than filler targets. Half of the real word primes came from the same corpora and half were taken from SUBTLEX-ESP (Vega et al. 2011). Finally, half of the pseudoword fillers were taken from a list of bisyllabic pseudowords conforming to Spanish phonotactics that was generated for a previous study using the Wuggy application. The other half was generated by replacing phonemes in the real word stimuli used in the present study. 62 of the pseudoword pairs consisted of the same stimulus repeated. Examples and proportions of real and pseudoword fillers can be found in Table 5.2.

When stimuli were S-class words, they were always produced with [s] (as they are expected to be in this variety). Around 10% of the stimuli overall contained [s] in an onset, in order to prevent the presence of [s] in the prime or the target from giving information as to the correct target response.

Both /s/ and /θ/ appeared freely in codas, since these are in any case elided to [h] or deleted in this variety and in the stimulus recordings. Half of the critical primes (by design) and 8% of real and pseudoword filler primes contained [s] in an onset. Additionally, 10% of the pseudoword filler primes contained [θ], to discourage an association of [θ] in the prime to real word targets. No filler primes were Z-class words, so that no items except the critical items included contexts that gave information about the talker’s use of *seseo* or distinction.

### 5.3.3 Talker

The speaker chosen to be the model talker for the experimental stimuli was a 28-year-old male Sevillano with native Sevillano parents. When recording the stimuli, the talker was encouraged to use local phonological features, including coda /s/ “aspiration” or lenition, [h] for /x/, and /l/ → [r] in word-internal codas. The stimuli were recorded in the same session as the stimuli from Chapter 2 over two sessions. Stimuli were read from a list on

Table 5.2: Example filler stimuli from Experiment 1

| Number | Fillter pair type | Prime           | Gloss    | Target         | Gloss  |
|--------|-------------------|-----------------|----------|----------------|--------|
| 138    | Real - Real       | <b>violín</b>   | ‘violin’ | <b>tarta</b>   | ‘cake’ |
| 69     | Real - Pseudo     | <b>bufanda</b>  | ‘scarf’  | <b>pranujo</b> |        |
| 142    | Pseudo - Pseudo   | <b>apuerejo</b> |          | <b>fumbre</b>  |        |
| 62     | Repeated Pseudo   | <b>tago</b>     |          | <b>tago</b>    |        |

paper and produced three times each. The talker was instructed to pronounce anterior fricatives according to how they were written on the page.

### 5.3.4 Procedure

The subjects were the same 24 speakers described in Chapter 2. The task was conducted in a quiet room on a MacBook Pro 2012 using PsychoPy experimental software and KRK systems KNS-8400 over-the ear headphones. Stimuli were presented in pairs with a 500 ms interval in between. Participants were instructed to make a lexical decision only on the second stimulus they heard. Once the subject made a response, the next pair began playing 1000 ms later. The entire experiment consisted of 549 pairs, and the subjects were given two breaks over the course of the experiment.

## 5.4 Results

The data was prepared for analysis according to the procedure outlined in Baayen and Milin (2010): it was transformed to maximize the normality of the distribution of reaction times, in this case with a log transformation, and outliers were trimmed. This resulted in the exclusion of 9% of critical trials.

A linear mixed-effects regression model predicting the logarithm of reaction time to critical trials (Table 5.3) was fitted with a random intercept by Target and a random intercept of Subject with a random slope by Trial Number. It included all of the following fixed terms: Relatedness, Z-Word Condition, Frequency of Target, Frequency of Prime, Target Duration, and two interaction terms: Relatedness x Condition and Frequency of

|                              | Estimate | Std. Error | df    | p-value |
|------------------------------|----------|------------|-------|---------|
| (Intercept)                  | 6.91     | 0.04       | 31.35 | 0.00**  |
| Target Duration              | 0.04     | 0.01       | 91.59 | 0.00**  |
| $\theta$ -condition          | -0.02    | 0.02       | 90.85 | 0.34    |
| Related                      | -0.05    | 0.02       | 92.58 | .02*    |
| $\theta$ -condition: Related | 0.05     | 0.03       | 91.34 | 0.09    |

Table 5.3: A linear mixed-effects regression model predicting the log of reaction time.

|                      | Estimate | Std. Error | df    | p-value |
|----------------------|----------|------------|-------|---------|
| (Intercept)          | 6.89     | 0.041      | 31.35 | 0.00**  |
| Target Duration      | 0.05     | 0.00       | 91.59 | 0.00**  |
| S-condition          | 0.02     | 0.02       | 90.85 | 0.34    |
| Related              | 0.00     | 0.022426   | 92.58 | 0.95    |
| S-condition: Related | -0.05    | 0.031      | 91.34 | 0.09    |

Table 5.4: A linear mixed-effects regression model predicting the log of reaction time.

Target x Relatedness. Stepwise model selection was conducted using Likelihood Ratio Tests. The following paragraph summarizes the final model.

Reaction times in the S-condition were faster for Related pairs than for Unrelated pairs ( $\beta = -0.05$ ,  $p < 0.001$ ). When the model is run with the  $\theta$ -condition as the reference level (see Table 5.4), Relatedness is not a significant predictor. These relationships are reflected when the raw data is visualized: reaction times were similar for Related and Unrelated pairs in the  $\theta$ -condition (see Figure 5.1). Stimuli in the S-condition induced a mean priming effect of 81.48 ms, and stimuli in the  $\theta$ -condition did not induce a priming effect (see Figure 5.2). That is, subjects were faster to respond to related targets than to unrelated ones after an S-condition prime, but there is no evidence that stimuli in the  $\theta$ -condition induced priming.

The standard interpretation of a speedup effect like this one is that subjects successfully and rapidly accessed the meaning of the prime, in this case a Z-word produced with *seseo*. It seems that words with *seseo* are not only processed like words rather than pseudowords, but they appear to be processed more easily than their standardly-produced counterparts.

When examined in combination with the proportion of [ $\theta$ ] in / $\theta$ /-contexts that each speakers showed in their sociolinguistic interviews in Chapter 2, no clear pattern emerges. Figure 5.3 shows the priming effects by condition and by subject. The speakers, as in the

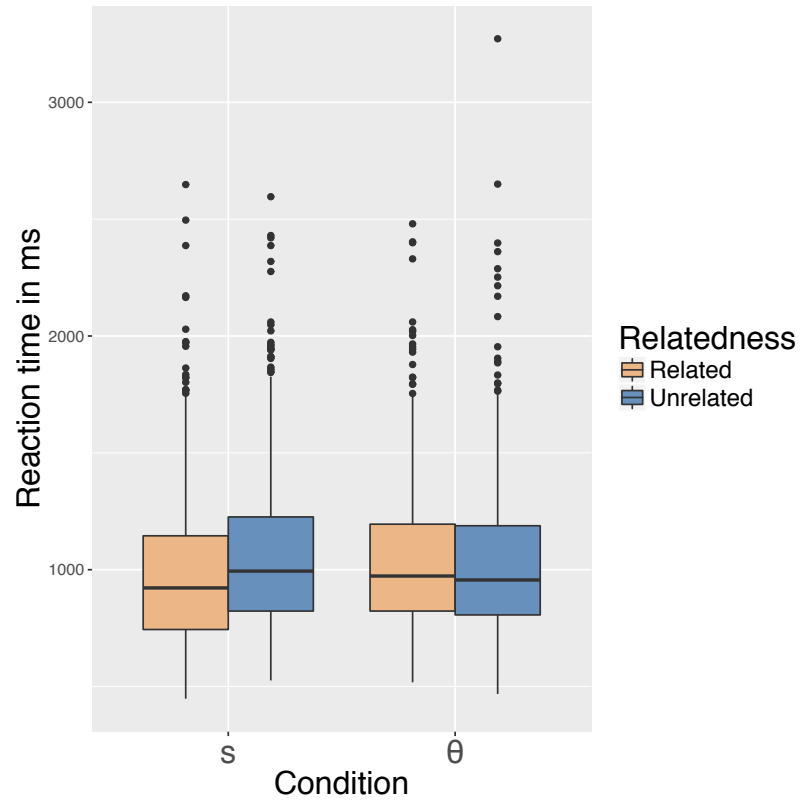


Figure 5.1: Reaction times by condition.

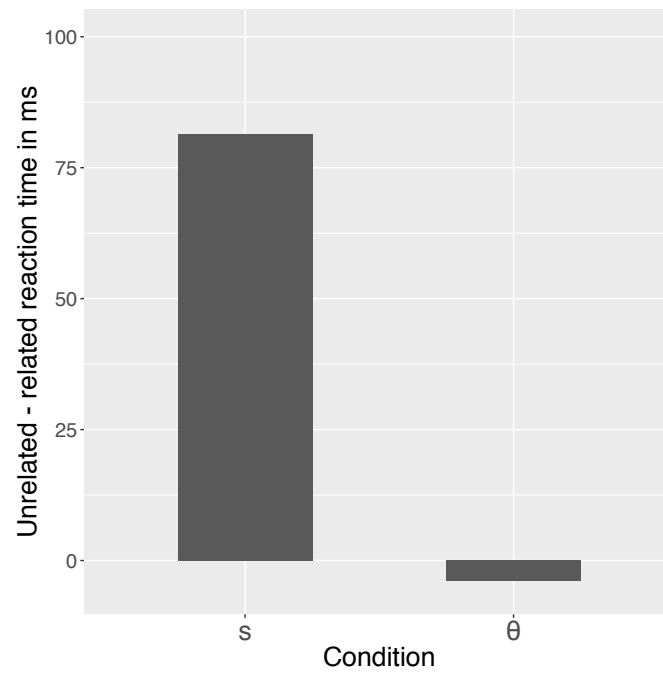


Figure 5.2: Mean difference of reaction times by condition.

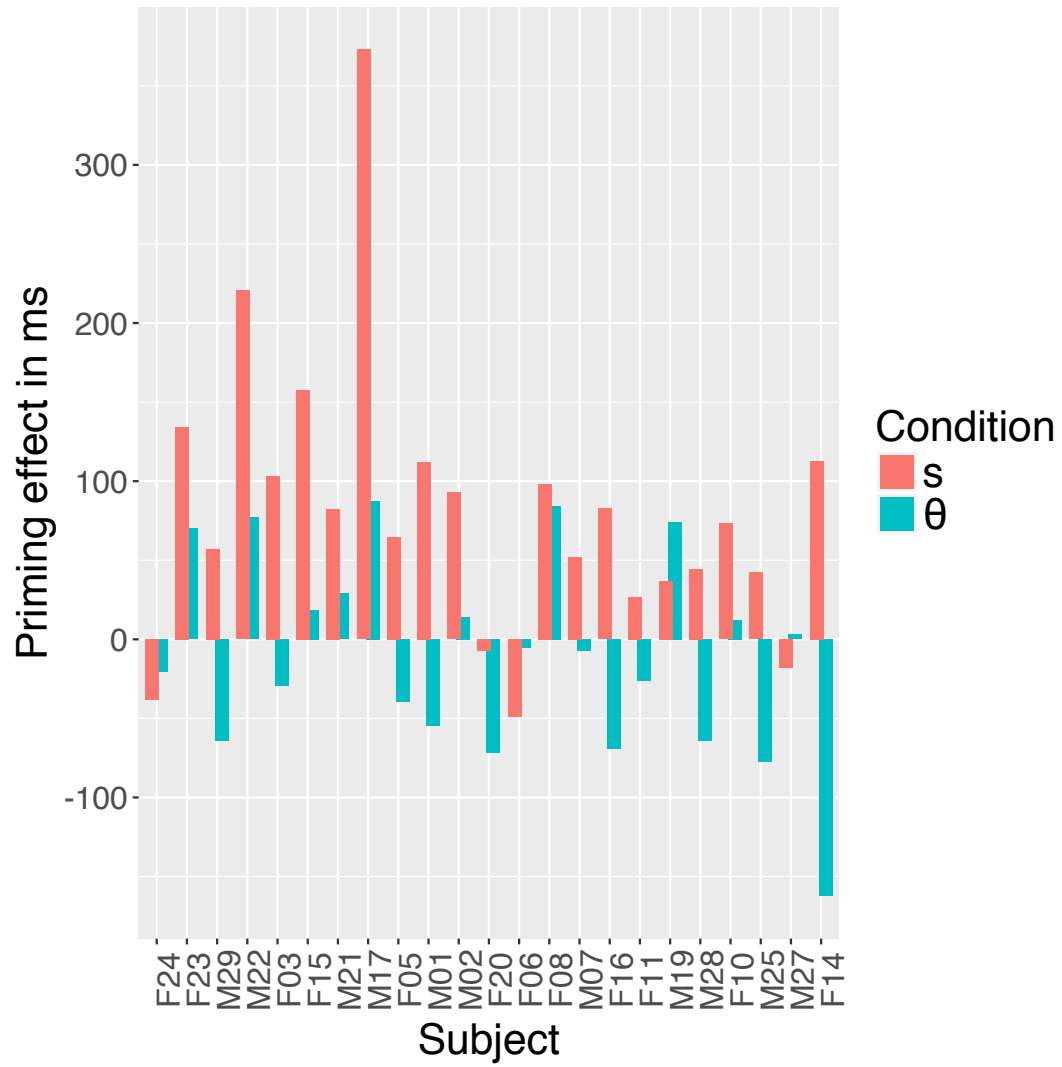


Figure 5.3: Mean difference of reaction time (unrelated - related) by experimental condition and by subject. Subjects are ordered from left to right by rate of  $[\theta]$  in Z-contexts.

previous chapter, are ordered such that the most *seseante* speakers are on the left and the most standard speakers are on the right. The priming effects are not distributed such that the advantage for S-condition stimuli is limited only to the most *seseante* speakers. Rather, this generalization holds for most, but not all of the speakers generally. The only real tendency that can be observed is that the priming effect in the S-condition tends to be larger in the *seseo* dominant speakers.

## 5.5 Discussion

The priming manipulation was successful; subjects showed faster recognition of target words when they were preceded by a related prime than by an unrelated prime. This effect appears to be driven by the primes in the S-condition, which showed a greater degree of facilitation than primes in the  $\theta$ -condition. This result is unusual in that it shows an advantage for what can clearly be characterized as a local, if not a nonstandard, form over a standard form (cf. Andruski et al. 1994; Pitt 2009; Sumner and Samuel 2009). This is striking given that the local form appears to no longer be dominant in the community.

This study represents an addition to a set of studies on marginal phoneme contrasts in which speakers show discrimination abilities that are not reflected in their lexical access behavior. Like Labov (1994) and Amengual (2016)’s speakers, they show that *producing* a contrast does not mean that they use it in speech perception. Like the Catalan-Spanish bilinguals, they show that hearing the contrast does not mean that they do either. These are surprising facts, and should be taken seriously by researchers interested in theoretical and psychological aspects of language. As Labov points out, “stubborn facts that resist explanation by any available theory... provide the most powerful stimulus to the development of new methods and insights into the operation of the world around us” (Labov 1994:368).

The speakers in this study, then, present an apparent paradox: they appear to have difficulty processing lexical items containing a sound, but that sound is present in their own productive inventories, is linked to appropriate lexical items in their production, and they are able to distinguish it acoustically from its close neighbor. Knowledge of the so-

ciolinguistic background explored in Chapter 2 tells us that the speakers in this study are mostly likely able to handle instances of Z-words produced standardly with [θ] perfectly well. They consume national media dominated by *distinción* every day through television and are exposed to some amount of such tokens in their local context as well, sometimes from their own mouths. One possibility that would explain this unexpected result is that all or some of the speakers in this study are perfectly capable of processing [θ], but not tokens that occur in the context that this study created. Again, context can be defined at different levels. One context is the experiment itself, including my presence (and use of *seseo*) and expectations of what the experiment was about (the Seville dialect). Another is the talker: this includes his gender and individual voice characteristics as well as characteristics like social class and place of origin that are signaled in the stimuli over the course of the experiment. As several critical items have been heard, they begin to occur in a context in which the listener now has heard evidence that the talker is a user of *seseo*. Finally, the token occurs in the immediate context of a word frame that has characteristics like duration and pitch but also potentially contains regional segmental features, all of which may occur more or less often when *seseo* is used. All of this combined to create an expectation of *seseo*, which the θ-condition stimuli did not match.

The proposed explanation is essentially that the speakers have the ability to switch between one-phoneme and two-phoneme systems of perception according to who is talking. If there exists both awareness of *seseo* and features that can be relied on to co-occur with it, this kind of multiple system solution in perception is a possibility. A speech perception model that accounts for this must include not just information about the immediate phonetic frame that the variant appears in, as in Sumner, Kim, et al. (2013)’s model, but also about the larger context of other variants produced previously by the talker. If listeners do in fact modulate their expectations of how the fricative in Z-words will sound in response to talker features, we would expect different results with a Castilian talker, or even with one that used local features minimally. This kind of talker manipulation is a direction for future work.



Recall that care was taken to evoke the local dialect as much as possible through the choice of talker and through his use of local features. A reasonable concern is whether this resulted in the  $\theta$ -condition primes sounding unnatural, or incongruous. Two points are relevant here: the first is that the talker’s peers all use at least some standard  $[\theta]$  in Z-words, according to the production data from this study and the data in Santana (2016b), in which a mixture of *seseo* and contrast was found to be the most common pattern among young, educated people in Seville. That is, standard  $[\theta]$  is not incongruent with the local dialect as it is spoken today. The second point is that even if the talker was using a style that lowers expectations of the standard  $/s/-/\theta/$  contrast, it is important to remember that a  $[\theta]$  in a Z-class word is also compatible with *ceceo*, which is found in some surrounding towns as discussed in Chapter 1.

The differential rates of *seseo* in production do not seem to reflect deep differences in acoustic representation, to the extent that the processes accessed by this study reflect the general phonemic function in perception. Although use of  $[\theta]$  appears to be the norm among college-educated people in Seville, rates of *seseo* remain relatively high in other social groups (Santana 2016a). The input of all of these speakers may be dominated by  $[s]$ -forms.

Methodologically, the choice to not counterbalance the item conditions between subjects, done in order to maximize meaningful comparisons between individuals, does place a limitation on the interpretability of the results. Care was taken to balance word characteristics between the conditions as much as possible, but the possibility of a difference in processing difficulty between the conditions that is unrelated to the experimental manipulation cannot be fully ruled out. Another consideration is that the lexical decision task creates an expectation for pseudowords that are not present in normal communication/speech perception which might change how Z-words are processed. Using a different task which does not involve pseudowords to further probe representations, such as a semantic classification task, is a promising route for future research.

## 5.6 Conclusion

The cases of perception/production mismatch from the beginning of this chapter show that changes in progress can challenge assumptions about how production and perception pattern together. Speakers of Sevillian Spanish, a dialect undergoing a phonemic split in the front anterior fricative space, provide an unusual case study that is a departure from even these cases. These speakers appear to have created a new phonemic category with an incoming sound and pattern their production accordingly, deploying the old sound freely alongside the new sound, but carefully limiting the new sound to its standard lexical class. However, in perception the established sound appears to result in “better” tokens of Z-class words, at least given the conditions in this experiment.

Taken together with the results from Chapter 4, the results highlight the complexity of the notion “perception” and of its link to production. Individual phonological knowledge cannot be properly captured by a binary inclusion or exclusion of phonemes in an inventory, but must rather be captured by nuanced exploration of perception and production behavior. The processing advantage for the local [s] variant can be taken as evidence against the universal application of a speech model that privileges standard forms. Finally, this study underscores the importance of collecting data on both the production and perception of variation in individuals, as well as the need for studies on changes involving contrasts between a variety of sounds.

## Chapter 6

# Conclusions and extensions

In this dissertation, I examine a reported change in progress (Santana 2016b) in a speech community in the south of Spain toward a phonemic contrast between /s/ and /θ/ and away from a previous situation of merger to /s/ (*seseo*). I take an individual-focused approach, examining a group of 24 young Seville speakers by analyzing both their production and perception of the contrast.

The change in question has been put forth as an exception to Garde’s principle (Labov 1994) of the impossibility of merger reversal. In my view, contact-based reversal is not *contra* the principle, and this change seems to have contact-based rather than internal origins, as others have concluded. Nevertheless, it is still one the few attested cases of such a borrowing taking place. The /s/-/θ/ contrast enjoys a remarkably high level of general awareness and is reflected in the orthography. These factors, along with contact with the contrast in standard Castilian Spanish through contact through migration, travel, media, and formal education, seem to have enabled the contrast to be borrowed. The impression that this change is exceptional may reflect an existing bias in the literature toward vowel changes in English, as others examining the progression of the same change in nearby cities (e.g., Villena-Ponsoda 2001; García-Amaya 2008; Moya Corral and Wiedemann 1995) have pointed out (Hernández-Campoy and Villena-Ponsoda 2009; Regan 2017b).

My results add to a growing set of studies giving evidence that, *contra* earlier claims

that emphasize their orientation toward regional norms, speakers in western Andalusian cities are moving toward the standard /s/-/θ/ contrast along with, if a little bit behind, their eastern counterparts. The rate of use of the contrast in my sample is very much in line with the that of the PRESEEA sample of educated Seville speakers reported in Santana (2016b), lending independent support to her findings that the move toward /s/-/θ/ Seville is robust and probably has been proceeding for some time.

Though most of them produce [s] in standard /θ/ contexts at least some of the time, I contend that the speakers in my sample have two phonemes /s/ and /θ/. In a sense, then, this population will not help solve Weinreich et al. 1968’s Transition Problem, because the change from one phoneme to two has already taken place in these speakers. I base this claim on the results described in Chapter 4, which give acoustic evidence of separate targets, demonstrates the speakers’ ability to identify and discriminate between the sounds in question, and shows the systematic patterning of their use of the borrowed sound [θ] in conversational speech: [θ] is essentially limited to contexts where it is standard. I argue that this necessarily requires the linking of items in the mental lexicon to two separate phonological categories, and that this is most appropriately described as phonological contrast. I suggest that individual speakers are able to switch between systems of contrast and merger, much in the way that Maguire et al. (2013)’s discussion of “variable” merger implies (p. 235). This switching can be likened to the syntactic idea of competing grammars (Kroch 1989), which has in recent work been applied to phonological change (Sneller 2018; Fruehwald et al. 2009).

The more puzzling question is perhaps not how speakers in Seville might have come to have the contrast, but rather why *seseo* would persist in their speech in spite of it. Usual predictors of change from above like gender do not seem to be related to the use of *seseo* in this sample of speakers. I believe the answer to this question may lie in the interaction of the regional and national prestige of *seseo* and contrast in Spain. Unlike the move from *ceceo* to *distinción*, this may be a change “from above” only in the sense of its position above the level of awareness. The social meanings of *seseo* and *distinción* are complex, and

macrosocial categories may not be able to capture the resulting variation. Data gathered as a part of this study could shed further light on this: the results of demographic and local identity questionnaires are only partially reported here, and commentary given by the speakers during recorded debriefing sessions has yet to be transcribed and analyzed. An examination of individual speakers' attitudes and commentary alongside their demographic information and production behavior may be a promising avenue for understanding the social meanings and motivations that are at play in their production of these sounds. The speakers in this study do show social stratification in their (overall negligible) use of *ceceo* ([θ] for /s/). The lower prestige of *ceceo*, the regional prestige of *seseo*, and the overt prestige of *distinción* combine to form a fascinating sociolinguistic situation in parts of Seville province where all three patterns are robustly present (Ruiz Sanchez 2017). More study of anterior fricative variation in these places is an exciting potential area of future research.

Finally, in Chapter 5, I investigate the speakers' use of the phonemic contrast for lexical access during speech processing using a semantic priming paradigm. I find that Z-words produced with [s] show a processing advantage over ones produced with [θ], and I suggest that the result reflects the speakers' perception of the experimental talker's local background: the speakers had difficulty processing words with standard [θ] not because they are merged in perception, but rather because the talker generated an expectation of *seseo* with his use of other local phonological features consistent with *seseo*. A slightly different potential explanation is that the talker's use of stigmatized features activated associations to some sort of stereotypicalized version of a Sevillian that listeners are referencing to generate expectations. Future work can explore talker-based phonemic flexibility in perception by comparing Seville listeners' perception of different realizations of Z-words produced by both local and standard Castilian-sounding talkers (who would be predicted to generate expectations of standard contrast), and by manipulating the specific local features exhibited by local talkers. Additionally, if we hypothesize that speakers have multiple grammars or modes of perception (merged and distinct), analogous to those in their production, that are

turned on and off by talker cues, this makes interesting predictions about perception, e.g., effects of shifting lexical neighborhood size for a given word.

To conclude, in this dissertation, I distinguish four different facets of perception and production that are associated with phonemic contrast, and I emphasize that researchers should not assume that a given speaker will exhibit any one of them. Rather, we should examine them separately, and this requires a combination of methodologies. Different aspects of an individual speaker's perception and production sometimes fail to line up with one another, and this should lead us to question both what we mean when we talk about phonemic contrast, whether as a grammatical or a psychological property, and what tools we should use to diagnose it in cases of phonological change. In general, this dissertation demonstrates the benefits of combining experimental methodology with an evaluation of the experimental participants' own naturalistic productions of variable linguistic phenomena.

# Appendix A

## Elicitation sentences

1. María estaba leyendo un libro.
2. Le pidieron a Jaime que fuera el nuevo entrenador del equipo.
3. Al boli no le queda mucha tinta.
4. Al final se acaba, al principio se acaba.
5. Miraba el fondo del mar con unas gafas de buceo.
6. Fue al dentista a que le miraran los ojos.
7. No es diestro, es zurdo.
8. Ella tenía que leerle los labios porque era sorda.
9. Para que los perros no entraran, construyó una valla.
10. Por aparcar mal pagó una multa.
11. Un paraguas te ayuda a mantenerte seco.
12. La ducha llenó el cuarto de baño de agua.
13. Le pidió que repitiera lo que había dicho.
14. Vertió el vaso y tiró agua por todo el suelo.

15. No podía resistir comprarlo por un tan buen precio.
16. No podía tomar el café caliente porque quemaba.
17. Mandaron el paquete por correo.
18. Julio esperaba que no le reñieran por haber llegado tan tarde.
19. Tuvo un día largo y estaba de mal humor.
20. Tiró la comida mala a la basura.
21. Le caían los pantalones por falta de un cinturón.
22. No esperaba ninguna llamada cuando el teléfono sonó.
23. Llevaba un abrigo gordo por el frío.
24. No tiene el pelo corto, lo tiene largo.
25. En la primera línea, escribe tu nombre.
26. Los surfistas temen a los tiburones.
27. No se veía ni una nube en el cielo.
28. Daniel fue a coger leña para la hoguera.
29. Llevaba un pañuelo blanco en el cuello.



# Appendix B

## Demographic questionnaire

Encuesta demográfica

Número de sujeto:

Fecha:

1. Edad

\_\_\_\_\_

3. ¿Tienes conocimiento de padecer algún problema de audición?

2. Sexo

\_\_\_\_\_

Sí

No

4. Aparte de español, ¿eres hablante nativo de algún otro idioma?

Sí

No

En caso afirmativo, ¿qué idioma(s)?

\_\_\_\_\_

5. ¿A qué te dedicas?

\_\_\_\_\_

6. Marca los que te aplican.

Bachillerato   FP   Carrera empezada   Carrera   Máster

Filólogo/a   Profesor/a de lengua

7. ¿En qué lugar(es) pasaste la infancia (hasta los catorce años)?

| Localidad | Provincia | País | Edades |
|-----------|-----------|------|--------|
|-----------|-----------|------|--------|

| Localidad | Provincia | País | Edades |
|-----------|-----------|------|--------|
|-----------|-----------|------|--------|

| Localidad | Provincia | País | Edades |
|-----------|-----------|------|--------|
|-----------|-----------|------|--------|

8. ¿En que otro(s) lugar(s) has vivido o permanecido más de un mes?

|           |           |      |                      |
|-----------|-----------|------|----------------------|
| Localidad | Provincia | País | Duración de estancia |
|-----------|-----------|------|----------------------|

|           |           |      |                      |
|-----------|-----------|------|----------------------|
| Localidad | Provincia | País | Duración de estancia |
|-----------|-----------|------|----------------------|

|           |           |      |                      |
|-----------|-----------|------|----------------------|
| Localidad | Provincia | País | Duración de estancia |
|-----------|-----------|------|----------------------|

9. ¿Tienes amigos o familia (personas con las has tenido siempre mucha relación) de sitios que no sean Sevilla capital?

Sí                      No

10. En caso afirmativo, ¿de dónde son y quiénes son? Utiliza el espacio abajo para explicar tus respuestas si es necesario.

|           |           |                           |          |
|-----------|-----------|---------------------------|----------|
| Localidad | Provincia | Frecuencia de interacción | Relación |
|-----------|-----------|---------------------------|----------|

|           |           |                           |          |
|-----------|-----------|---------------------------|----------|
| Localidad | Provincia | Frecuencia de interacción | Relación |
|-----------|-----------|---------------------------|----------|

|           |           |                           |          |
|-----------|-----------|---------------------------|----------|
| Localidad | Provincia | Frecuencia de interacción | Relación |
|-----------|-----------|---------------------------|----------|

11. Piensa en las personas que te han criado (con las que pasaste más tiempo de niño/a). ¿Dónde han crecido? ¿A qué se dedican/dedicaban?

|           |           |           |          |
|-----------|-----------|-----------|----------|
| Localidad | Provincia | Profesión | Relación |
|-----------|-----------|-----------|----------|

|           |           |           |          |
|-----------|-----------|-----------|----------|
| Localidad | Provincia | Profesión | Relación |
|-----------|-----------|-----------|----------|

12. ¿Hay alguna influencia en tu forma de hablar que aquí no se ha preguntado? Por favor menciónala utilizando el espacio en blanco.

## Local identity questionnaire

## Encuesta demográfica

Número de sujeto:

Fecha:

1. ¿Hay algunas circunstancias en las que te ves viviendo fuera de Sevilla?

Sí No

En caso afirmativo, ¿qué circunstancias serían?

¿Dentro o fuera de España?

2. Ordena los siguientes lugares de 1-7 según tu identificación con los mismos. Barrio se refiere a tu zona de Sevilla, y pueblo al pueblo con el que sientes más vinculación; escríbe por favor los dos en los blancos.

España Andalucía Sevilla                       Andalucía occidental  
Barrio Pueblo

3. ¿Te identificas como sevillano/a?

|      |   |   |      |   |   |            |
|------|---|---|------|---|---|------------|
| 1    | 2 | 3 | 4    | 5 | 6 | 7          |
| Nada |   |   | Algo |   |   | Totalmente |

4. ¿Aficionado/a a Semana Santa?

| 1    | 2 | 3 | 4    | 5 | 6 | 7          |
|------|---|---|------|---|---|------------|
| Nada |   |   | Algo |   |   | Totalmente |

5. ¿Feriante?

1      2      3      4      5      6      7  
Nada                      Algo                      Totalmente

6. ¿Flamenco/a?

|      |   |   |      |   |   |            |
|------|---|---|------|---|---|------------|
| 1    | 2 | 3 | 4    | 5 | 6 | 7          |
| Nada |   |   | Algo |   |   | Totalmente |

7. ¿Hasta qué punto tu forma de hablar habitual se puede caracterizar como sevillano?

1      2      3      4      5      6      7  
Nada                      Algo                      Totalmente

Encuesta demográfica

Número de sujeto:

Fecha:

8. ¿Cambias tu forma de hablar cuando hablas con españoles no andaluces?

|      |   |   |      |   |   |            |
|------|---|---|------|---|---|------------|
| 1    | 2 | 3 | 4    | 5 | 6 | 7          |
| Nada |   |   | Algo |   |   | Totalmente |

9. Te consideras:

|                         |              |                         |
|-------------------------|--------------|-------------------------|
| Algo seseante           | Muy seseante | Ni ceceante ni seseante |
| Algo ceceante           | Muy ceceante | Seseante y ceceante     |
| No conozco los términos |              |                         |

10. Escribe 5 palabras que asocias con una persona que sesea.

\_\_\_\_\_

\_\_\_\_\_

11. ¿Alguien te ha corregido la pronunciación de la Z?

|                        |                            |                      |
|------------------------|----------------------------|----------------------|
| Maestros en el colegio | Profesores en el instituto | Profesores en la uni |
| Familiares de aquí     | Familiares de fuera        | Otros _____          |
| Nadie                  |                            |                      |

12. ¿Cómo ves el seseo de aceptable?

|      |   |   |      |   |   |            |
|------|---|---|------|---|---|------------|
| 1    | 2 | 3 | 4    | 5 | 6 | 7          |
| Nada |   |   | Algo |   |   | Totalmente |

13. ¿Cómo ves el ceceo de aceptable?

|      |   |   |      |   |   |            |
|------|---|---|------|---|---|------------|
| 1    | 2 | 3 | 4    | 5 | 6 | 7          |
| Nada |   |   | Algo |   |   | Totalmente |

14. ¿Eres consciente de tu uso o no uso del seseo de día en día?

|      |   |   |      |   |   |            |
|------|---|---|------|---|---|------------|
| 1    | 2 | 3 | 4    | 5 | 6 | 7          |
| Nada |   |   | Algo |   |   | Totalmente |

15. Elige la opinión con la que estás más de acuerdo.

El seseo es incorrecto y se debe corregir todo lo posible

El seseo sólo se debe corregir en ciertos ámbitos

El seseo es una cosa local y no se debe corregir nunca

### Selected demographic and identity data

| Subject | Age | Neighborhood       | Highest Ed.  |
|---------|-----|--------------------|--------------|
| M01     | 31  | Calle Feria        | Carrera      |
| M02     | 21  | Amate              | Máster       |
| F03     | 21  | Triana             | Bachillerato |
| F05     | 23  | Pino Montano       | Máster       |
| F06     | 27  | Triana             | Bachillerato |
| M07     | 24  | Santa Aurelia      | Bachillerato |
| F08     | 32  | Centro             | Máster       |
| F10     | 20  | La Oliva           | Bachillerato |
| F11     | 24  | Centro             | Bachillerato |
| M12     | 22  | San Bernardo       | Bachillerato |
| F14     | 25  | Centro             | Carrera      |
| F15     | 29  | Sevilla Este       | Carrera      |
| F16     | 25  | Polígono San Pablo | Bachillerato |
| M17     | 30  | Centro             | Carrera      |
| M19     | 30  | Macarena           | Bachillerato |
| F20     | 29  | Macarena           | Carrera      |
| M21     | 34  | Pino Montano       | FP           |
| M22     | 29  | Triana             | Máster       |
| F23     | 24  | Macarena           | Bachillerato |
| F24     | 33  | Triana             | Máster       |
| M25     | 24  | San Pablo          | Carrera      |
| M27     | 35  | Macarena           | Bachillerato |
| M28     | 30  | Huerta de la Salud | Máster       |
| M29     | 21  | Parque Flores      | Bachillerato |

| Subject | Sevillano | Flamenco | Lived Abroad | Would Relocate | Feria | Holy Week |
|---------|-----------|----------|--------------|----------------|-------|-----------|
| M01     | 7         | 1        | No           | Yes            | 6     | 7         |
| M02     | 7         | 2        | No           | Yes            | 5     | 7         |
| F03     | 4         | 5        | No           | Yes            | 5     | 2         |
| F05     | 7         | 3        | Yes          | Yes            | 4     | 4         |
| F06     | 6         | 2        | No           | Yes            | 4     | 2         |
| M07     | 4         | 1        | No           | Yes            | 2     | 1         |
| F08     | 4         | 1        | Yes          | Yes            | 3     | 3         |
| F10     | 6         | 5        | No           | No             | 7     | 1         |
| F11     | 7         | 4        | No           | Yes            | 6     | 3         |
| M12     | 7         | 4        | Yes          | Yes            | 7     | 5         |
| F14     | 4         | 3        | Yes          | Yes            | 3     | 2         |
| F15     | 7         | 7        | No           | No             | 7     | 6         |
| F16     | 6         | 7        | No           | Yes            | 2     | 1         |
| M17     | 7         | 1        | No           | Yes            | 1     | 5         |
| M19     | 4         | 2        | No           | Yes            | 7     | 1         |
| F20     | 7         | 6        | Yes          | Yes            | 2     | 1         |
| M21     | 4         | 3        | No           | Yes            | 3     | 1         |
| M22     | 5         | 3        | No           | Yes            | 2     | 2         |
| F23     | 7         | 1        | No           | Yes            | 4     | 6         |
| F24     | 6         | 3        | Yes          | Yes            | 6     | 2         |
| M25     | 5         | 2        | No           | Yes            | 3     | 1         |
| M27     | 6         | 1        | No           | Yes            | 1     | 1         |
| M28     | 3         | 3        | Yes          | Yes            | 4     | 2         |
| M29     | 6         | 1        | No           | Yes            | 3     | 1         |

| Subject | Adjusts | Seseo Rating | Ceceo Rating | Correction | Should Correct |
|---------|---------|--------------|--------------|------------|----------------|
| M01     | 1       | 6            | 2            | Nobody     | Sometimes      |
| M02     | 4       | 6            | 4            | Everybody  | Never          |
| F03     | 6       | 7            | 7            | Nobody     | Never          |
| F05     | 3       | 7            | 7            | Nobody     | Sometimes      |
| F06     | 4       | 5            | 4            | Nobody     | Sometimes      |
| M07     | 5       | 4            | 1            | Nobody     | Never          |
| F08     | 4       | 7            | 7            | Nobody     | Never          |
| F10     | 2       | 7            | 7            | Nobody     | Never          |
| F11     | 2       | 4            | 1            | Family     | Never          |
| M12     | 5       | 6            | 5            | Nobody     | Never          |
| F14     | 2       | 7            | 7            | Nobody     | Never          |
| F15     | 1       | 6            | 6            | Nobody     | Sometimes      |
| F16     | 1       | 7            | 7            | Teachers   | Never          |
| M17     | 1       | 7            | 1            | Nobody     | Never          |
| M19     | 6       | 5            | 3            | Nobody     | Sometimes      |
| F20     | 1       | 7            | 7            | Nobody     | Never          |
| M21     | 3       | 5            | 1            | Nobody     | Sometimes      |
| M22     | 2       | 7            | 7            | Nobody     | Sometimes      |
| F23     | 1       | 7            | 7            | Nobody     | Never          |
| F24     | 4       | 6            | 5            | Nobody     | Never          |
| M25     | 1       | 7            | 7            | Nobody     | Never          |
| M27     | 1       | 3            | 2            | Nobody     | Never          |
| M28     | 6       | 7            | 6            | Nobody     | Never          |
| M29     | 5       | 5            | 3            | Nobody     | Never          |



| Subject | Northern Parent | Northern Contacts | Ceceo Contacts | Ceceo Parent |
|---------|-----------------|-------------------|----------------|--------------|
| M01     | Yes             | Yes               | No             | No           |
| M02     | Yes             | Yes               | No             | No           |
| F03     | Yes             | Yes               | No             | No           |
| F05     | No              | Yes               | No             | No           |
| F06     | No              | No                | Yes            | No           |
| M07     | No              | No                | Yes            | No           |
| F08     | No              | Yes               | No             | No           |
| F10     | No              | No                | No             | No           |
| F11     | No              | No                | No             | Yes          |
| M12     | No              | No                | Yes            | No           |
| F14     | No              | Yes               | Yes            | No           |
| F15     | No              | No                | No             | No           |
| F16     | No              | No                | Yes            | No           |
| M17     | No              | No                | No             | No           |
| M19     | No              | No                | No             | Yes          |
| F20     | No              | No                | No             | No           |
| M21     | No              | No                | No             | Yes          |
| M22     | No              | No                | No             | No           |
| F23     | No              | No                | No             | No           |
| F24     | Yes             | No                | No             | No           |
| M25     | No              | No                | Yes            | No           |
| M27     | No              | Yes               | No             | Yes          |
| M28     | Yes             | Yes               | No             | No           |
| M29     | No              | No                | No             | No           |

# Bibliography

- Alvar-López, Manuel (1996). *Manual de dialectología hispánica. El español de España*. Barcelona: Ariel.
- Alvar, Manuel (1973). *Atlas Lingüístico y Etnográfico de Andalucía*. Universidad de Granada. Consejo superior de investigaciones científicas.
- Amengual, Mark (2016). “The perception of language-specific phonetic categories does not guarantee accurate phonological representations in the lexicon of early bilinguals”. In: *Applied Psycholinguistics* 37.5, pp. 1221–1251.
- Andruski, Jean E, Sheila E Blumstein, and Martha Burton (1994). “The effect of subphonetic differences on lexical access”. In: *Cognition* 52.3, pp. 163–187.
- Auer, Peter and Frans Hinskens (1996). “The convergence and divergence of dialects in Europe. New and not so new developments in an old area”. In: *Sociolinguistica* 10, pp. 1–30.
- Ávila Muñoz, Antonio (1994). “Variación reticular e individual de s/z en el Vernáculo Urbano Malagueño: Datos del barrio de Capuchinos”. In: *Analecta Malacitana* 17.1994, pp. 343–367.
- Baayen, R Harald and Petar Milin (2010). “Analyzing reaction times”. In: *International Journal of Psychological Research* 3.2, pp. 12–28.
- Babel, Molly, Michael McAuliffe, and Graham Haber (2013). “Can mergers-in-progress be unmerged in speech accommodation?” In: *Frontiers in Psychology* 4, p. 653.
- Baranowski, Maciej Andrzej (2007). *Phonological variation and change in the dialect of Charleston, South Carolina*. Duke University Press.
- (2013). “On the role of social factors in the loss of phonemic distinctions”. In: *English Language & Linguistics* 17.2, pp. 271–295.
- Becker, Kara (2009). “/r/and the construction of place identity on New York City’s Lower East Side 1”. In: *Journal of Sociolinguistics* 13.5, pp. 634–658.
- Boland, Julie E (2017). “Cognitive Mechanisms and Syntactic Theory”. In: *Twenty-first century psycholinguistics: Four cornerstones*. Ed. by Anne Cutler. Routledge.

- Boland, Julie E, Edith Kaan, Jorge Valdés Kroff, and Stefanie Wulff (2016). “Psycholinguistics and variation in language processing”. In: *Linguistics Vanguard* 2.s1.
- Bond, Zinny S (1999). *Slips of the ear: Errors in the perception of casual conversation*. Academic Press.
- Braver, Aaron (2014). “Imperceptible incomplete neutralization: Production, non-identifiability, and non-discriminability in American English flapping”. In: *Lingua* 152, pp. 24–44.
- Bullock, Barbara E and Jenna Nichols (2017). “Return to Frenchville”. In: *Romance Languages and Linguistic Theory 11: Selected papers from the 44th Linguistic Symposium on Romance Languages (LSRL), London, Ontario*. Vol. 11. John Benjamins Publishing Company, p. 229.
- Bürki, Audrey (2018). “Variation in the speech signal as a window into the cognitive architecture of language production”. In: *Psychonomic Bulletin & Review*, pp. 1–32.
- Callejas, Alicia, Ángel Correa, Juan Lupiáñez, and Pío Tudela (2003). “Normas asociativas intracategoriales para 612 palabras de seis categorías semánticas en español”. In: *Psicológica* 24.2.
- Campbell-Kibler, Kathryn (2010). “The sociolinguistic variant as a carrier of social meaning”. In: *Language Variation and Change* 22.03, pp. 423–441.
- Carbonero Cano, Pedro (1982). “El habla de Sevilla”. In: *Sevilla, Biblioteca de Temas Sevillanos*.
- (2003). *Estudios de sociolingüística andaluza*. Sevilla: Universidad.
- Carbonero Cano, Pedro, José Luis Álvarez, Joaquín Casas, and Isabel M Gutiérrez (1992). “El habla de Jerez. Estudio sociolingüístico”. In: *Cuadernos de Divulgación* 5.
- Celdrán, Eugenio Martínez and Ana María Fernández Planas (2007). *Manual de fonética española: articulaciones y sonidos de español*. Editorial Ariel.
- Chambers, Jack K (1992). “Dialect acquisition”. In: *Language* 68.4, pp. 673–705.
- Chomsky, Noam and Morris Halle (1968). *The sound pattern of English*. Harper & Row.
- Clopper, Cynthia G and David B Pisoni (2004). “Some acoustic cues for the perceptual categorization of American English regional dialects”. In: *Journal of Phonetics* 32.1, pp. 111–140.
- Comellas, José Luis (1992). *Sevilla, Cádiz y América: el trasiego y el tráfico*. MAPFRE.
- Connine, Cynthia M (2004). “It’s not what you hear but how often you hear it: On the neglected role of phonological variant frequency in auditory word recognition”. In: *Psychonomic Bulletin & Review* 11.6, pp. 1084–1089.

- Cook, Svetlana V, Nick B Pandža, Alia K Lancaster, and Kira Gor (2016). “Fuzzy nonnative phonolexical representations lead to fuzzy form-to-meaning mappings”. In: *Frontiers in Psychology* 7, p. 1345.
- Cruttenden, Alan (1994). *Gimson’s pronunciation of English*. London: Edward Arnold.
- Cutler, Anne (2012). *Native listening: Language experience and the recognition of spoken words*. MIT Press.
- Cutler, Anne, Andrea Weber, Roel Smits, and Nicole Cooper (2004). “Patterns of English phoneme confusions by native and non-native listeners”. In: *The Journal of the Acoustical Society of America* 116.6, pp. 3668–3678.
- Dalbor, John B (1980). “Observations on Present-Day Seseo and Ceceo in Southern Spain”. In: *Hispania*, pp. 5–19.
- De Cortázar, Fernando García and José Manuel González Vesga (1994). *Breve historia de España*. Alianza.
- Di Paolo, Marianna (1992). “Hypercorrection in response to the apparent merger of (ɔ) and (ɑ) in Utah English”. In: *Language & Communication* 12.3-4, pp. 267–292.
- Di Paolo, Marianna and Alice Faber (1990). “Phonation differences and the phonetic content of the tense-lax contrast in Utah English”. In: *Language Variation and Change* 2.2, pp. 155–204.
- Eckert, Penelope (2000). *Language variation as social practice: The linguistic construction of identity in Belten High*. Wiley-Blackwell.
- Eckert, Penelope and William Labov (2017). “Phonetics, phonology and social meaning”. In: *Journal of Sociolinguistics* 21.4, pp. 467–496.
- Eisner, Frank and James M McQueen (2005). “The specificity of perceptual learning in speech processing”. In: *Perception & Psychophysics* 67.2, pp. 224–238.
- Evans, Bronwen G and Paul Iverson (2007). “Plasticity in vowel perception and production: A study of accent change in young adults”. In: *The Journal of the Acoustical Society of America* 121.6, pp. 3814–3826.
- Floccia, Caroline, Jeremy Goslin, Frédérique Girard, and Gabrielle Konopczynski (2006). “Does a regional accent perturb speech processing?” In: *Journal of Experimental Psychology: Human Perception and Performance* 32.5, p. 1276.
- Fourakis, Marios and Gregory K Iverson (1984). “On the ‘incomplete neutralization’ of German final obstruents”. In: *Phonetica* 41.3, pp. 140–149.
- Fruehwald, Josef, Jonathan Gress-Wright, and Joel Wallenberg (2009). “Phonological rule change: The constant rate effect”. In: *Proceedings of NELS*. Vol. 40.

- García-Amaya, Lorenzo J (2008). “Variable norms in the production of  $\theta$  in Jerez de la Frontera, Spain”. In: *IULC Working Papers* 8.3.
- Giegerich, Heinz J (1992). *English phonology: An introduction*. Cambridge University Press.
- Goldinger, Stephen D (1996). “Auditory lexical decision”. In: *Language and Cognitive Processes* 11.6, pp. 559–568.
- González Jiménez, Manuel (2003). “Andalucía, una realidad histórica”. In: *Actas de las jornadas sobre el habla andaluza: El español hablado en Andalucía*.
- Gonzalez-Bueno, Manuela (1993). “Variaciones en el tratamiento de las sibilantes Inconsistencia en el seseo sevillano: Un enfoque sociolingüístico”. In: *Hispania*, pp. 392–398.
- Gordon, Matthew (2013). “Investigating chain shifts and mergers”. In: *The Handbook of Language Variation and Change*. Vol. 11. Blackwell Malden, MA/Oxford, pp. 203–219.
- Gordon, Matthew, Paul Barthmaier, and Kathy Sands (2002). “A cross-linguistic acoustic study of voiceless fricatives”. In: *Journal of the International Phonetic Association* 32.02, pp. 141–174.
- Gylfadottir, Duna (2016). *Effects of talker dialect on lexical decision involving a merged phoneme*. Poster presented at the 2016 Annual Meeting of the Linguistics Society of America.
- Haeri, Niloofar (1991). “Sociolinguistic variation in Cairene Arabic: Palatalization and the *qaf* in the speech of men and women”. PhD thesis. University of Pennsylvania.
- Hanique, Iris, Mirjam Ernestus, and Barbara Schuppler (2013). “Informal speech processes can be categorical in nature, even if they affect many different words”. In: *The Journal of the Acoustical Society of America* 133.3, pp. 1644–1655.
- Hanulíková, Adriana, Petra M Van Alphen, Merel M Van Goch, and Andrea Weber (2012). “When one person’s mistake is another’s standard usage: The effect of foreign accent on syntactic processing”. In: *Journal of Cognitive Neuroscience* 24.4, pp. 878–887.
- Hay, Jennifer and Katie Drager (2010). “Stuffed toys and speech perception”. In: *Linguistics* 48.4, pp. 865–892.
- Hay, Jennifer, Katie Drager, and Brynmor Thomas (2013). “Using nonsense words to investigate vowel merger”. In: *English Language & Linguistics* 17.2, pp. 241–269.
- Hay, Jennifer, Katie Drager, and Paul Warren (2010). “Short-term exposure to one dialect affects processing of another”. In: *Language and speech* 53.4, pp. 447–471.
- Hay, Jennifer, Paul Warren, and Katie Drager (2006). “Factors influencing speech perception in the context of a merger-in-progress”. In: *Journal of Phonetics* 34.4, pp. 458–484.

- Heras, Jerónimo Borrero de las, José Romero Delgado, María Dolores Bardallo Bardallo, Valentín Torrejón Moreno, María Carmen Castrillo Díaz, Juan Gallego Blanca, José María Padilla Valencia, and Carmen Vacas Muñoz (1996). "Perfil sociolingüístico del habla culta de la zona periurbana de Huelva". In: *Aestuarium: revista de investigación* 4, pp. 109–124.
- Hernández-Campoy, Juan Manuel and Juan Andrés Villena-Ponsoda (2009). "Standardness and nonstandardness in Spain: dialect attrition and revitalization of regional dialects of Spanish". In: *International Journal of the Sociology of Language*.
- Herold, Ruth (1990). "Mechanisms of merger: The implementation and distribution of the low back merger in Eastern Pennsylvania". PhD thesis. University of Pennsylvania.
- Hickey, Raymond (2004). "Mergers, near-mergers, and phonological interpretation". In: *New Perspectives on English Historical Linguistics: Selected papers from 12 ICEHL, Glasgow, 21–26 August 2002. Volume II: Lexis and Transmission* 252, p. 125.
- Hillenbrand, James M, Michael J Clark, and Robert A Houde (2000). "Some effects of duration on vowel recognition". In: *The Journal of the Acoustical Society of America* 108.6, pp. 3013–3022.
- Holcomb, Phillip J and Helen J Neville (1990). "Auditory and visual semantic priming in lexical decision: A comparison using event-related brain potentials". In: *Language and cognitive processes* 5.4, pp. 281–312.
- Hooper, Joan B (1976). *An introduction to natural generative phonology*. Academic Press.
- Hualde, José Ignacio (2005). *The Sounds of Spanish*. Cambridge University Press.
- Janson, Tore and Richard Schulman (1983). "Non-distinctive features and their use". In: *Journal of Linguistics* 19.02, pp. 321–336.
- Jassem, Wiktor and Lutoslaw Richter (1989). "Neutralization of voicing in Polish obstruents". In: *Journal of Phonetics* 17.4, pp. 317–325.
- Jespersen, Otto (1954). *A Modern English Grammar: On Historical Principles: Part I - Sounds and Spelling*. Routledge.
- Johnson, Daniel Ezra (2010). *Stability and change along a dialect boundary: The low vowels of Southeastern New England (Publication of the American Dialect Society 95)*. Duke University Press.
- Jongman, Allard, Ratree Wayland, and Serena Wong (2000). "Acoustic characteristics of English fricatives". In: *The Journal of the Acoustical Society of America* 108.3, pp. 1252–1263.
- Kerswill, Paul and Ann Williams (2002). "'Salience' as an explanatory factor in language change: evidence from dialect levelling in urban England". In: *Language change: The interplay of internal, external and extra-linguistic factors*, pp. 81–110.

- Kleber, Felicitas, Tina John, and Jonathan Harrington (2010). "The implications for speech perception of incomplete neutralization of final devoicing in German". In: *Journal of Phonetics* 38.2, pp. 185–196.
- Koops, Christian, Elizabeth Gentry, and Andrew Pantos (2008). "The effect of perceived speaker age on the perception of PIN and PEN vowels in Houston, Texas". In: *University of Pennsylvania Working Papers in Linguistics* 14.2, p. 12.
- Kroch, Anthony S (1989). "Reflexes of grammar in patterns of language change". In: *Language variation and change* 1.3, pp. 199–244.
- Kvapil, Jan (2010). "La situación actual del seseo y su reflejo en los medios de comunicación". MA thesis. Masarykova univerzita.
- Labov, William (1972). "Some principles of linguistic methodology". In: *Language in society* 1.1, pp. 97–120.
- (1989). "The limitations of context: Evidence from misunderstandings in Chicago". In: *Chicago Linguistic Society*. Vol. 25, pp. 171–200.
- (1994). *Principles of Linguistic Change. Vol. 1: Internal factors*. Oxford: Blackwell.
- (2001). *Principles of Linguistic Change. Vol. 2: Social Factors*. Oxford: Blackwell.
- (2010). *Principles of Linguistic Change. Vol. 3: Cognitive and cultural factors*. Oxford: Blackwell.
- Labov, William, Sharon Ash, and Charles Boberg (2006). "Atlas of North American English: Phonology and Phonetics". In: *Berlin: Mouton de Gruyter*.
- Labov, William, Malcah Yaeger, and Richard Steiner (1972). *A Quantitative Study of Sound Change in Progress*. US Regional Survey.
- Ladefoged, Peter and Ian Maddieson (1996). *The sounds of the world's languages*. Vol. 1012. Blackwell Oxford.
- Lamíquiz, Vidal and Pedro Carbonero Cano Cano (1985). *Sociolingüística andaluza: Metodología y estudios*. 59. Universidad de Sevilla.
- Lapesa, Rafael (1957). *Sobre el ceceo y seseo andaluces*. Universidad de La Laguna, Canarias.
- Lapesa, Rafael and Ramón Menéndez Pidal (1981). *Historia de la lengua española*. 9th. Escelicer.
- Lasarte Cervantes, María de la Cruz (2010). "Datos para la fundamentación empírica de la escisión fonemática prestigiosa de /θ<sup>s</sup>/ en Andalucía". In: *Nueva Revista de Filología Hispánica*, pp. 483–516.

- Lasarte Cervantes, María de la Cruz (2012). “Variación social en la percepción del contraste meridional entre /s/ y /θ/ en Málaga”. In: *Estudios sobre el español de Málaga. Pronunciación, Vocabulario y Sintaxis*, pp. 167–188.
- Lloyd, Paul M (1987). *From Latin to Spanish: Historical phonology and morphology of the Spanish language*. Vol. 1. American Philosophical Society.
- Luna, Fernando Gabriel, Julián Marino, Joaquín Darío Silva, and Alberto Acosta Mesas (2016). “Normas de asociación léxica e índices psicolingüísticos de 407 palabras en español en una muestra latinoamericana”. In: *Psicológica* 37.1.
- Maguire, Warren, Lynn Clark, and Kevin Watson (2013). “Introduction: what are mergers and can they be reversed?” In: *English Language & Linguistics* 17.2, pp. 229–239.
- Manaster Ramer, Alexis (Oct. 1996). *A Letter from an Incompletely Neutral Phonologist*.
- Marslen-Wilson, William and Paul Warren (1994). “Levels of perceptual representation and process in lexical access: words, phonemes, and features”. In: *Psychological review* 101.4, p. 653.
- Martinet, André (1952). “Function, structure, and sound change”. In: *Word* 8.1, pp. 1–32.
- Maye, Jessica, Richard N Aslin, and Michael K Tanenhaus (2008). “The weckud wetch of the wast: Lexical adaptation to a novel accent”. In: *Cognitive Science* 32.3, pp. 543–562.
- McGowan, Kevin B (2011). “The role of socioindexical expectation in speech perception”. PhD thesis. University of Minnesota.
- McLennan, Conor T, Paul A Luce, and Jan Charles-Luce (2003). “Representation of lexical form”. In: *Journal of Experimental Psychology: Learning, Memory, and Cognition* 29.4, p. 539.
- Melguizo Moreno, Elisabeth (2009). “Una aproximación sociolingüística al estudio del ceceo en un corpus de hablantes granadinos”. In: *Estudios de Lingüística Aplicada* 49, pp. 57–78.
- Milroy, Lesley and Matthew Gordon (2008). *Sociolinguistics: Method and interpretation*. Vol. 13. John Wiley & Sons.
- Mitchell, Timothy (1990). *Passional culture: emotion, religion, and society in Southern Spain*. University of Pennsylvania Press.
- Mitterer, Holger and James M McQueen (2009). “Processing reduced word-forms in speech perception using probabilistic knowledge about speech production”. In: *Journal of Experimental Psychology: Human Perception and Performance* 35.1, p. 244.
- Moreno Fernández, Francisco (1993). *Proyecto para el estudio del Español de España y América (PRESEEA). Presentación*.



- Moreno, Elisabet Melguizo (2005). “Estudio del patrón no sibilante (“ceceo”) en un grupo de inmigrantes pineros instalados en Granada”. In: *Interlingüística* 16, pp. 777–789.
- Morillo-Velarde, Ramón (1997). “Seseo, ceceo y seceo: problemas metodológicos”. In: A. Narbona y M. Ropero, *Actas del Congreso del habla andaluza*, pp. 201–219.
- Moya Corral, Juan Antonio (1997). “Desarraigo social y cambio lingüístico. El ejemplo de Granada”. In: *El habla andaluza. Actas del Congreso de habla andaluza*, pp. 623–634.
- Moya Corral, Juan Antonio and Marcin Sosinski (2015). “La inserción social del cambio. La distinción s/θ en Granada. Análisis en tiempo aparente y en tiempo real”. In: *LEA: Lingüística española actual* 37.1, pp. 33–72.
- Moya Corral, Juan Antonio and Emilio J García Wiedemann (1995). *El habla de Granada y sus barrios*. Universidad de Granada.
- Narbona, Antonio, Rafael Cano, and Ramón Morillo (2003). *El español hablado en Andalucía*. Editorial Ariel.
- Navarro Tomás, Tomás (1996). *Manual de pronunciación española*. Consejo Superior de Investigaciones Científicas.
- Navarro Tomás, Tomás, Aurelio Macedonio Espinosa, and Lorenzo Rodríguez-Castellano (1933). “La frontera del andaluz”. In: *Revista de filología española* 20, pp. 225–277.
- Niedzielski, Nancy (1999). “The effect of social information on the perception of sociolinguistic variables”. In: *Journal of language and social psychology* 18.1, pp. 62–85.
- Nycz, Jennifer (2011). “Second dialect acquisition: Implications for theories of phonological representation”. PhD thesis. New York University.
- Ohala, John J (1981). “Articulatory constraints on the cognitive representation of speech”. In: *Advances in Psychology*. Vol. 7. Elsevier, pp. 111–122.
- Osterhout, Lee and Phillip J Holcomb (1992). “Event-related brain potentials elicited by syntactic anomaly”. In: *Journal of Memory and Language* 31.6, pp. 785–806.
- Ota, Mitsuhiro, Robert J Hartsuiker, and Sarah L Haywood (2009). “The KEY to the ROCK: Near-homophony in nonnative visual word recognition”. In: *Cognition* 111.2, pp. 263–269.
- Pallier, Christophe, Nuria Sebastián-Gallés, and Angels Colomé (1999). “Phonological representations and repetition priming”. In: *Proceedings of Eurospeech '99*. Budapest, Hungary, pp. 1907–1910.
- Patterson, David and Cynthia M Connine (2001). “Variant frequency in flap production”. In: *Phonetica* 58.4, pp. 254–275.

- Payne, Arvilla (1980). “Factors controlling the acquisition of the Philadelphia dialect by out-of-state children”. In: *Locating language in time and space* 1, pp. 143–78.
- Penny, Ralph (1991). *A history of the Spanish language*. 2nd. Cambridge University Press.
- (2000). *Variation and change in Spanish*. Cambridge University Press.
- Pidal, Ramón Menéndez (1962). *Sevilla frente a Madrid: Algunas precisiones sobre el español de América*. Universidad de La Laguna.
- Pisoni, David B (1973). “Auditory and phonetic memory codes in the discrimination of consonants and vowels”. In: *Perception & Psychophysics* 13.2, pp. 253–260.
- Pitt, Mark A (2009). “The strength and time course of lexical activation of pronunciation variants”. In: *Journal of Experimental Psychology: Human Perception and Performance* 35.3, p. 896.
- Prieto, Juan Pablo Rodríguez (2014). “Distribución geográfica del “jejeo” en Español y propuesto de reformulación y extensión del término”. In: *Revista española de lingüística* 38.2, pp. 129–144.
- Radeau, Monique (1983). “Semantic priming between spoken words in adults and children.” In: *Canadian Journal of Psychology/Revue canadienne de psychologie* 37.4, p. 547.
- Radeau, Monique, José Morais, and Juan Segui (1995). “Phonological priming between monosyllabic spoken words”. In: *Journal of Experimental Psychology: Human Perception and Performance* 21.6, p. 1297.
- Ranbom, Larissa J and Cynthia M Connine (2007). “Lexical representation of phonological variation in spoken word recognition”. In: *Journal of Memory and Language* 57.2, pp. 273–298.
- Reed, Paul (2016). “Sounding Appalachian: /aI/ Monophthongization, Rising Pitch Accents, and Rootedness”. PhD thesis. University of South Carolina.
- Regan, Brendan (2015). “Stylistic variation of non-standard merger in Andalucía: A Speaker Design Approach”. Unpublished manuscript (under 2nd review).
- (2017a). “A study of ceceo variation in Western Andalusia (Huelva)”. In: *Studies in Hispanic and Lusophone Linguistics* 10.1, pp. 119–160.
- (2017b). “The effect of dialect contact and social identity on fricative demerger”. PhD thesis. The University of Texas at Austin.
- Ringe, Don and Joseph F Eska (2013). *Historical Linguistics: Toward a Twenty-First Century Reintegration*. Cambridge University Press.
- Rodríguez Domínguez, Manuel (2017). *El andaluz, vanguardia del español*. Alfar.

- Roettger, Timo B, Bodo Winter, Sven Grawunder, James Kirby, and Martine Grice (2014). “Assessing incomplete neutralization of final devoicing in German”. In: *Journal of Phonetics* 43, pp. 11–25.
- Ruiz Sanchez, Carmen (2017). “Seseo, ceceo, and distinción in Andalusian Spanish: Free variation or sociolinguistic variation?” In: *Linguistics Vanguard* 3.1.
- Salvador, Francisco (1980). “Niveles sociolingüísticos de seseo, ceceo y distinción en la ciudad de Granada”. In: *Español actual: Revista de español vivo* 37, pp. 25–32.
- Sánchez Méndez, Juan P (2003). *Historia de la lengua española en América*. Tirant lo Blanch.
- Sankoff, Gillian (2004). “Adolescents, young adults and the critical period: Two case studies from ‘Seven Up’”. In: *Sociolinguistic variation: Critical reflections*, pp. 121–39.
- Santana, Juana (2016a). “La realización de /s/ y /θ/ iniciando sílaba en los materiales de PRESEEA Sevilla: estudio en el sociolecto bajo”. XII Congreso Internacional de Lingüística General. Alcalá de Henares.
- (2016b). “Seseo, ceceo y distinción en el sociolecto alto de la ciudad de Sevilla: nuevos datos a partir de los materiales de PRESEEA”. In: *Boletín de filología* 51, pp. 255–280.
- Sawoff, Adolf (1980). “A sociolinguistic appraisal of the sibilant pronunciation in the city of Seville”. In: *Festgabe für Norman Denison, Grazer Linguistische Studien*, pp. 11–12.
- Shadle, Christine Helen (1985). “The acoustics of fricative consonants”. PhD thesis. Massachusetts Institute of Technology, released as MIT-RLE Technical Report No. 506.
- Slowiacek, Louisa M and David B Pisoni (1986). “Effects of phonological similarity on priming in auditory lexical decision”. In: *Memory & Cognition* 14.3, pp. 230–237.
- Sneller, Betsy (2018). “Mechanisms Of Phonological Change”. PhD thesis. University of Pennsylvania.
- Staum Casasanto, Lauren (2009). “Experimental investigations of sociolinguistic knowledge”. PhD thesis. Stanford University.
- Stewart, Miranda (2012). *The Spanish Language Today*. Routledge.
- Strand, Elizabeth A (2000). “Gender stereotype effects in speech processing”. PhD thesis. The Ohio State University.
- Sumner, Meghan (2013). “A phonetic explanation of pronunciation variant effects”. In: *The Journal of the Acoustical Society of America* 134.1, EL26–EL32.
- Sumner, Meghan, Seung Kyung Kim, and Kevin B McGowan King (2013). “The socially weighted encoding of spoken words: a dual-route approach to speech perception”. In: *Frontiers in Psychology* 4.

- Sumner, Meghan and Arthur G Samuel (2005). “Perception and representation of regular variation: The case of final /t/”. In: *Journal of Memory and Language* 52.3, pp. 322–338.
- (2009). “The effect of experience on the perception and representation of dialect variants”. In: *Journal of Memory and Language* 60, pp. 487–501.
- Tagliamonte, Sali (2011). *Variationist sociolinguistics: Change, observation, interpretation*. Vol. 40. John Wiley & Sons.
- Taylor, Wilson L (1953). “‘Cloze procedure’: A new tool for measuring readability”. In: *Journalism Bulletin* 30.4, pp. 415–433.
- Thomas, Erik R (2002). “Sociophonetic applications of speech perception experiments”. In: *American Speech* 77.2, pp. 115–147.
- Trudgill, Peter (2003). “On the reversibility of mergers: /w/, /v/ and evidence from lesser-known Englishes”. In: *Folia Linguistica Historica* 24.1-2, pp. 23–45.
- Tuinman, Annelie, Holger Mitterer, and Anne Cutler (2011). “Perception of intrusive /r/ in English by native, cross-language and cross-dialect listeners”. In: *The Journal of the Acoustical Society of America* 130.3, pp. 1643–1652.
- Uruburu, Agustín (1990). “Seseo en el habla juvenil de la ciudad de Córdoba (España)”. In: *Estudios sobre la lengua española en Córdoba*, pp. 125–134.
- Vega, Fernando Cuetos, María González Nosti, Analía Barbón Gutiérrez, and Marc Brysbaert (2011). “SUBTLEX-ESP: Spanish word frequencies based on film subtitles”. In: *Psicológica: Revista de metodología y psicología experimental* 32.2, pp. 133–143.
- Villena-Ponsoda, Juan Andrés (2001). *La continuidad del cambio lingüístico. Tendencias conservadoras e innovadoras en la fonología del español a la luz de la investigación sociolingüística urbana*. Granada: Editorial Universidad de Granada.
- (2005). “How similar are people who speak alike? An interpretive way of using social networks in social dialectology research”. In: *Dialect change: Convergence and divergence in European languages*, pp. 303–334.
- (2007). “Interacción de factores internos y externos en la explicación de la variación fonológica. Análisis multivariable del patrón de pronunciación no sibilante [θ] de la consonante fricativa coronal /s/ en el español hablado en Málaga”. In: *Las hablas andaluzas y la enseñanza de la lengua*, p. 69.
- (2008). “La formación del español común en Andalucía. Un caso de escisión prestigiosa”. In: *Fonología instrumental: patrones fónicos y variación*. El Colegio de México, pp. 211–256.
- Villena-Ponsoda, Juan Andrés, José María Sánchez Sáez, and Antonio Manuel Ávila Muñoz (1995). “Modelos probabilísticos multinomiales para el estudio del ceceo, seseo y dis-

- tinción de /s/ y /θ/: datos de la ciudad de Málaga”. In: *ELUA. Estudios de Lingüística*, N. 10, pp. 391–435.
- Villena-Ponsoda, Juan Andrés and Félix Santos Requena (1996). “Género, educación y uso lingüístico: la variación social y reticular de S y Z en la ciudad de Málaga”. In: *Lingüística* 8, pp. 5–52.
- Warner, Natasha, Allard Jongman, Joan Sereno, and Rachèl Kemps (2004). “Incomplete neutralization and other sub-phonemic durational differences in production and perception: Evidence from Dutch”. In: *Journal of Phonetics* 32.2, pp. 251–276.
- Weinreich, Uriel, William Labov, and Marvin Herzog (1968). “Empirical foundations for a theory of language change”. In:
- Wells, John C (1982). *Accents of English*. Vol. 1. Cambridge University Press.
- Wilbanks, Eric (2018). *faseAlign*. <https://github.com/EricWilbanks/faseAlign>. Version 1.1.9.
- Wolfram, Walt and Erik Thomas (2008). *The Development of African American English*. John Wiley & Sons.
- Wyld, Henry (1936). *A History of Colloquial English*. Oxford Press.
- Yang, Charles (2009). “Population structure and language change”. In: *Ms., University of Pennsylvania*.
- Ylinen, Sari, Maria Uther, Antti Latvala, Sara Vepsäläinen, Paul Iverson, Reiko Akahane-Yamada, and Risto Näätänen (2010). “Training the brain to weight speech cues differently: A study of Finnish second-language users of English”. In: *Journal of Cognitive Neuroscience* 22.6, pp. 1319–1332.
- Yu, Alan (2007). “Understanding near mergers: The case of morphological tone in Cantonese”. In: *Phonology* 24.1, pp. 187–214.