

The use of songs in the ESL / EFL classroom as a means of teaching pronunciation: A case study of Chilean university students

Karen Elizabeth Brawnwyn Borland

A dissertation submitted to the
Faculty of Graduate and Postdoctoral Studies
in partial fulfillment of the requirements
for the PhD degree in Spanish (Linguistics)

Modern Languages and Literatures
Faculty of Arts
University of Ottawa

Abstract

In this thesis, I set out to investigate whether the use of songs can help L2 speakers learn to better perceive and produce suprasegmental phenomena. Effective pronunciation skills are necessary for successful communication and as such can greatly impact one's personal, social, and professional life. Studying the use of songs for teaching pronunciation is interesting because as a linguistically rich material, songs can enhance learning due to their positive affective, social, and cognitive influence in the L2 classroom. Using songs to teach pronunciation within a Communicative Language Teaching (CLT) framework constitutes a novel approach to an underexplored area of classroom research. In order to learn how using songs might help native Spanish speakers learn English suprasegmentals, I conducted a mixed methods exploratory short-term case study of Chilean university students studying English Language and Literature at the Universidad Católica de Chile. Using three groups: a control, songs, and no-songs group, the pre- to post-course progress was measured first with the two treatment groups combined and then with them separated. In this way we were able to measure the effectiveness of songs compared to other materials as well as to no intervention whatsoever. After two weeks of instruction, we found that using songs can significantly help in the production of the schwa when reading a text and of thought groups when speaking freely. Results obtained in listening tests were not statistically significant. However, closer examination of the performance of individual songs-group participants showed not only a greater than average progress in different suprasegmental areas in both listening and speaking, but also an appreciation of songs as an effective and enjoyable means of learning pronunciation. It would be advantageous for future research to explore the effects of teaching the pronunciation areas using the same methodology but for longer periods of time with delayed post-course testing to determine whether the effects are long-

term. In addition, further exploration into the relationship between pronunciation perception and production could provide insight for the development of more effective teaching techniques.

Table of Contents

Abstract	ii
Table of Contents	iv
List of Tables	xi
List of Figures	xv
Acknowledgements	xvi
1. Introduction	1
Overview	1
The Origins of This Thesis	2
Pronunciation and Its Neglect in L2 Teaching and Research	8
Reasons for the Lack of Pronunciation Instruction in CLT	10
The Challenge for Learners to Find Help	15
Some Consequences of Poor Pronunciation	17
Discrimination	18
The Musicality of Language	20
Dimensions of the Language Classroom	22
Songs and the Linguistic Aspect	23
Natural Speech	24
Levels of Formality	27
Language Awareness	27
The Affective Aspect	28
Motherese	29
Motivation	30

Meaningful Learning Material	31
The Social and Cultural Aspect	33
A Unifying Influence	33
Facets of Culture	35
The Cognitive Aspect	36
The Din and the SSIMH Phenomena	37
Automatic and Fluent Output	37
Intelligences and Learning	38
The Noticing Hypothesis	40
Memory	41
Problem Statement	44
Purpose of the Study	45
Research Questions	46
Limitations of the Study	49
Assumptions	50
Definitions	50
Summary	53
2. Review of the Literature	55
Introduction	55
Pronunciation's Subordinate Status	55
Accent	59
Non-Native Speakers' Opinions of Their Accents in L2 and L1	60
Environments	

Accent Addition Versus Accent Reduction	63
The Nativeness Principle and the Intelligibility Principle	65
Consequences of Accented Speech	69
Accent Discrimination	70
Reverse Linguistic Stereotyping	73
Educating the public	74
Pronunciation Teaching	74
Segmental Versus Suprasegmental	74
The Prosody Pyramid	76
Articles	78
Studies on Pronunciation Teaching	80
Traditional Materials	91
Methodologies and Approaches to L2 Pronunciation Teaching	95
Teaching With Songs	98
Methodologies and Approaches That Use Songs in L2 Teaching	98
Songs and Their General Usefulness in L2 Teaching	100
Literature Reviews	100
Textbooks	103
Articles	103
Studies on Songs in Language Teaching	105
Songs and Their Usefulness for Teaching and Learning Pronunciation	107
Literature reviews	108
Textbooks	108

	Articles	109
	Studies on Songs in Pronunciation Teaching	112
	Summary	116
3.	Methods and Procedures	118
	Introduction	118
	Research Paradigm	119
	Research Design	120
	Criteria	120
	Sampling	120
	Groups	121
	Listening	122
	Pronunciation areas	122
	Overview of the Course	124
	Pre-Intervention Steps	125
	The Participants	126
	Assignment to Groups	128
	The Schedule	129
	Data Collection Sources	131
	Questionnaires	131
	Tests	133
	Syllabus	138
	Methodology	142
	Designing the Course	144

	Lesson Plans and Classes	145
	Data Analysis Techniques	167
	Questionnaires	167
	Listening Tests	168
	Speaking Tests	173
	Summary	181
4.	Research Findings	182
	Introduction	182
	Analyses of Research Questions	183
	Research Question # 1a	183
	Listening Tests	183
	Summary of Research Question # 1a	190
	Research Question # 1b	191
	Listening Tests	191
	Comments in Speaking Test # 2	200
	Questionnaire Results for Usefulness of Songs	201
	Consolidation of Listening Data for the Songs Group	202
	Summary of Research Question # 1b	207
	Research Question # 2a	207
	Speaking Tests	207
	Summary of Research Question # 2a	214
	Research Question # 2b	215
	Speaking Tests	215

	Comments in Speaking Test # 2	228
	Questionnaire Results for Usefulness of Songs	229
	Consolidation of Speaking Data for the Songs Group	231
	Summary of Research Question # 2b	241
	Additional Findings	241
	No-Songs Group Results and Opinions	241
	Questionnaire Results	245
	Participants' Opinions on Accent	246
	Participants and Musical Aptitude	250
	Summary of the Findings	253
5.	Summary, Discussion, and Implications	256
	Introduction	256
	Limitations of the Study	257
	Assessment of Research Questions	261
	Research Question # 1a	261
	Research Question # 1b	263
	Research Question # 2a	265
	Research Question # 2b	267
	General Discussion of the Study	270
	Implications for Future Study	272
	Summary	277
	References	280
	Appendices	306

Appendix A	Ethics Approval and PUC Authorization	306
Appendix B	Questionnaires and Consent Forms	309
Appendix C	Lesson Plans and Materials	334
Appendix D	Listening and Speaking Tests	430

List of Tables

Table 1.1	The Usefulness of Songs According to the Dimensions of the L2 Classroom	43
Table 2.1	Guiding Principles in Pronunciation Teaching	91
Table 3.1	Content and Materials Used in English Pronunciation Course	142
Table 3.2	Listening Test # 2 Marking Example	170
Table 4.1	Listening Test # 1 Descriptive Statistics for the Control and Treatment Groups	184
Table 4.2	Listening Test # 2 Descriptive Statistics for the Control and Treatment Groups	184
Table 4.3	Listening Test # 1 Mean Scores and Mean Percentage Increase for the Control and Treatment Groups	185
Table 4.4	Listening Test # 2 Mean Scores and Mean Percentage Increase for the Control and Treatment Groups	186
Table 4.5	Listening Test # 1 and 2 Mean Scores and Mean Percentage Increase in Each Pronunciation Area for the Control and Treatment Groups	188
Table 4.6	Listening Test # 1 Descriptive Statistics for Each Group	192
Table 4.7	Listening Test # 2 Descriptive Statistics for Each Group	192
Table 4.8	Listening Test # 1 Mean Scores and Mean Percentage Increase for Each Group	193
Table 4.9	Listening Test # 1 One-Way Analysis of Variance	195
Table 4.10	Listening Test # 2 Mean Scores and Mean Percentage Increase for Each Group	196

Table 4.11	Listening Test # 2 One-Way Analysis of Variance	198
Table 4.12	Listening Test # 1 and 2 Mean Scores and Mean Percentage Increase in Each Pronunciation Area for Each Group	199
Table 4.13	Listening Tests Mean Percentage Increase for Each Group in Each Pronunciation Area	203
Table 4.14	Listening Tests Percentage Increase in Each Pronunciation Area for the Songs Group	203
Table 4.15	Speaking Test # 1 Descriptive Statistics for the Control and Treatment Groups	208
Table 4.16	Speaking Test # 2 Descriptive Statistics for the Control and Treatment Groups	209
Table 4.17	Speaking Test # 1 Mean Errors and Mean Percentage Difference for the Control and Treatment Groups	209
Table 4.18	Speaking Test # 2 Mean Errors and Mean Percentage Difference for the Control and Treatment Groups	210
Table 4.19	Speaking Test # 1 Mean Errors and Mean Percentage Difference in Each Pronunciation Area for the Control and Treatment Groups	212
Table 4.20	Speaking Test # 2 Mean Errors and Mean Percentage Difference in Each Pronunciation Area for the Control and Treatment Groups	213
Table 4.21	Speaking Test # 1 Descriptive Statistics for Each Group	216
Table 4.22	Speaking Test # 2 Descriptive Statistics for Each Group	216
Table 4.23	Speaking Test # 1 Mean Errors and Mean Percentage Difference for Each Group	216

Table 4.24	Speaking Test # 1 One-Way Analysis of Variance	218
Table 4.25	Speaking Test # 1 Post Hoc Tests	219
Table 4.26	Speaking Test # 2 Mean Errors and Mean Percentage Difference for Each Group	220
Table 4.27	Speaking Test # 1 Mean Errors and Mean Percentage Difference in Each Pronunciation Area for All Groups	221
Table 4.28	Speaking Test # 1 One-way Analysis of Variance for the Schwa	223
Table 4.29	Speaking Test # 1 Post Hoc Tests for the Schwa	223
Table 4.30	Speaking Test # 2 Mean Errors and Mean Percentage Difference in Each Pronunciation Area for All Groups	225
Table 4.31	Speaking Test # 2 One-way Analysis of Variance for Thought Groups	226
Table 4.32	Speaking Test # 2 Post Hoc Tests for Thought Groups	227
Table 4.33	Speaking Tests # 1 and 2 Mean Percentage Decrease in Errors for the Songs Group	231
Table 4.34	Speaking Test # 1 Mean Percentage Difference in Errors in Each Pronunciation Area for All Groups	233
Table 4.35	Speaking Test # 1 Percentage Decrease in Errors in Each Pronunciation Area for the Songs Group	233
Table 4.36	Speaking Test # 2 Mean Percentage Difference in Errors for Each Pronunciation Area for All Groups	235
Table 4.37	Speaking Test # 2 Percentage Decrease in Errors in Each Pronunciation Area for the Songs Group	236
Table 4.38	Listening Tests # 1 and 2 Mean Percentage Increase for the No-Songs	242

	Group	
Table 4.39	Speaking Tests # 1 and 2 Mean Percentage Decrease in Errors No-Songs	242
	Group	
Table 4.40	Musical Aptitude	251
Table 4.41	Participant Ranking and Musical Aptitude	252

List of Figures

Figure 2.1	The Prosody Pyramid	76, 123
Figure 3.1	The Groups and Data Sources	131
Figure 3.2	Pronunciation Survey	150
Figure 4.1	Listening Test # 1 Mean Percentage Increase for the Control and Treatment Groups	186
Figure 4.2	Listening Test # 2 Mean Percentage Increase for the Control and Treatment Groups	187
Figure 4.3	Listening Test # 1 Mean Percentage Increase for Each Group	193
Figure 4.4	Listening Test # 2 Mean Percentage Increase for Each Group	197
Figure 4.5	Speaking Test # 1 Mean Percentage Difference in Errors for the Control and Treatment Groups	210
Figure 4.6	Speaking Test # 2 Mean Percentage Difference in Errors for the Control and Treatment Groups	211
Figure 4.7	Speaking Test # 1 Mean Percentage Difference in Errors for Each Group	217
Figure 4.8	Speaking Test # 2 Mean Percentage Difference in Errors for Each Group	220
Figure 4.9	Views Toward English Accent	247
Figure 4.10	Views Toward Spanish Accent	168, 248
Figure 4.11	Desire to Keep Spanish Accent	249
Figure 4.12	Views on the Accents of Others	250

Acknowledgements

The successful completion of this thesis would not have been possible without the kind and generous assistance of many.

First of all, I would like to express my sincere gratitude to my advisor, Dr. Rodney Williamson, for his unwavering academic and emotional support throughout this whole ordeal. His numerous astute and pertinent comments were fundamental for the thesis, but it was his kindness, patience, understanding, and sense of humour that were essential for the soul.

In addition to my advisor, I would like to thank Dra. Miriam Cid Uribe and Dr. José Luis Samaniego for making it possible for me to conduct my research at the Pontificia Universidad Católica de Chile. I could not have asked for a better or friendlier place to be. The generous support and assistance from Sra. Paula Ross, Dra. Ana María Burdach, the English language professors, the secretaries, Don Guillermo, and Pablo, the lab assistant were much appreciated. Although she is no longer with us, I thank Dra. Ana María Harvey for her sage advice. I am grateful to my participants for being such wonderful students. My sincere thanks also go to Marcela and Pepe for generously opening their home to me and for their friendship and ongoing emotional support.

I would like to express my heartfelt appreciation to Dr. Jérémie Séror for compassionately showing me a way forward when I did not know how to proceed. His detailed feedback was simply invaluable.

I thank Dr. Luise von Flotow and Dr. Jaffer Sheyholislami for graciously offering their collaboration and insightful comments.

I am grateful to Marcus Vinícius Mota Amorim, my hard-working research assistant who was always eager to help whenever I needed him. I honestly do not know what I would have done without him.

To my dear friends and colleagues: Thank you for listening to me, encouraging me, and believing in me. You all helped me in many different and important ways.

To my family, and especially my parents, I owe a special thank you. I am truly grateful for your support throughout these years. Thank you for allowing me the freedom to study whatever I wished and for encouraging me to pursue a Ph.D. I would not and could not have done it without you.

Last but not least, I thank my sweet Haboobi for being my faithful little lap snuggler and for reminding me of the need to take breaks to eat and play.

Chapter One

Introduction

Overview

This thesis consists of five chapters plus References and Appendices. The Introduction provides background information leading to the study, which is on the use of songs for teaching English pronunciation to non-native speakers. The history and status of pronunciation teaching is explored as are some of the problems faced by non-native speakers with less than ideal pronunciation skills. The link between language and music is outlined followed by a discussion of how music and song may be beneficial in the language classroom. After discussing the need for pronunciation instruction as well as our lack of knowledge as to which techniques and materials might be most effective, this study proposes a way to teach pronunciation and seeks to determine whether songs can be an effective material in the teaching of English suprasegmentals. In the Introduction we also present the research questions, the limitations of our study, and definitions of key terms used throughout the thesis. “Chapter Two: A Review of the Literature” provides more detail and context regarding pronunciation’s status in the field. The notion of accent is explored including opinions regarding accents. Two opposing principles, The Nateness Principle and the Intelligibility Principle, which have to do with the teaching of second/foreign language (L2) pronunciation, are discussed as are the various consequences of accented speech. Then, L2 pronunciation teaching and L2 teaching involving songs are explored before reviewing literature specifically on the use of songs for teaching pronunciation. “Chapter

Three: Methods and Procedures” details the research design and how it was applied. This includes a detailed description of the materials used before, during, and after our brief study of teaching pronunciation to Chilean university students. “Chapter Four: Research Findings” provides an analysis of the research findings and answers each of the research questions. “Chapter Five: Summary, Discussion, Implications” provides a summary of the findings, discusses the data further, and mentions some of the practical implications for pronunciation teaching as well as possible future studies.

The Origins of This Thesis

This thesis essentially began many years ago when I first moved to Chile. I had already completed an undergraduate degree in Philosophy and Spanish Literature, during which I had spent my second year in Spain, so I felt confident in my ability to communicate in this new and exciting land. My self-assurance waned, however, shortly after arriving as I found myself unable to understand the Chilean accent. It was like nothing I had heard in school. And although I was able to make myself understood while speaking Spanish, at times I encountered problems and ridicule because of my pronunciation. I spoke like a “gringa¹” and even though I wanted to blend in orally, my accent always indicated that I was a foreigner.

After a few months, I began a master’s degree in Linguistics at the Pontificia Universidad Católica de Chile and became especially fascinated by one of my courses, Phonetics and Phonology. To learn that the International Phonetic Alphabet (IPA) was something that could capture and represent how people spoke meant that I had a means of learning the pronunciation

¹ Gringo/a: A term used in Chile and much of Latin America to refer to an individual from an English-speaking culture, especially an American.

that I so wanted to emulate. Through patient tutoring and the use of antiquated and very inadequate textbooks, I was able to change my accent so that I did not sound so much like an English speaker. Yet still I struggled. In my desire to avoid having to answer the same old tiring questions, such as, “Where are you from?” and “Why did you come to Chile?” I tried to fool even taxi drivers into believing that I was Chilean – but to no avail. Something was still wrong and I did not know what.

Then, the week before I was to leave Chile to begin a life in Brazil, I found myself listening to the songs of Victor Jara playing on the radio. That was when I began to realize the power of music. Suddenly I could hear the sounds and melody of the language in a way that I had not been able to before. As a result, I began taking more of an interest in Spanish-speaking music because I had found a way to practice the new-found sounds and sound patterns that had previously eluded my ears. Over and over I listened and sang to Victor Jara, Shakira, Buena Vista Social Club, and Los Jaivas. Finally, through singing, I was feeling and connecting with the language in a brand new way.

This experience with songs as a language learner was when I began to take music seriously as an educational tool. During the first year that I was in Chile, as a teacher, I had used music in my EFL classes but I was completely unaware of its merits. I had an open curriculum where I taught, and while I endeavoured to plan effective grammar and vocabulary lessons, my students requested classes filled with songs. Despite their claims that they learned through singing and discussing the lyrics, I assumed that they just wanted to spend their lunch hours having fun. Even many months later, when I was about to leave Chile and had discovered that songs could help *me* as a language learner, I did not make the connection as to why my students had insisted upon having songs in every class.

In Brazil, where I continued my journey as both language learner and teacher, songs did not form a part of my teaching and only after a while became part of my learning. As an instructor, I still did not take them seriously nor was their use ever advocated in the textbooks and curricula that I followed. While I was learning Portuguese, most of the classes involved grammar and vocabulary. The turning point, however, was when one of my colleagues exposed me to a very famous Brazilian song. At the time I was overwhelmed by the social and political importance of it, but then after listening to it for the language, I was finally able to hear the nasalisation in Brazilian Portuguese vowels. This was very important to me, because I had been told that I was not pronouncing certain vowels correctly, and as a result, some things that I said completely interfered with meaning. That is, I inadvertently said things that were different from what I had intended. Once again I realized the usefulness of songs for learning, and once again it was related to pronunciation.

Later, after leaving Brazil and returning to Chile for the next four years to teach at a university, I began to include songs from time to time in my teaching but usually as a grammatical, lexical or thematic tool as opposed to one for pronunciation. (My students had separate pronunciation classes, and since I taught grammar and “lengua” (language) classes, it was expected that I would focus on these areas.) By then I was more adept at using music and once again my students continually asked me to include more songs in my classes. I happily complied because it was while using songs that I noticed the spark in my students’ eyes and their interest in and engagement with the language. They stopped chattering to each other, they paid more attention, and they asked questions about the lyrics. Finally, as a result of these experiences, I became convinced of the power of songs for use in language teaching.

Upon my return to Canada, songs continued to be a material that I used in both the Spanish and ESL classes that I taught. It was in the Spanish classes especially at Trent University where my Canadian undergraduate students regularly asked me for more songs and requested to bring their own – complete with activities – for the class. That they appreciated classes in which they were able to take some ownership in what they were learning was clear because they eagerly volunteered to bring in their favourite ones to discuss and present in class.

After teaching at Trent for two years, I moved to Ottawa to begin my Ph.D. By this time I knew that I wanted to discover what it was about songs that seem to make them so well-liked and useful as a language-learning tool. While in Ottawa, I have been fortunate to teach both Spanish and English and my students have been, in addition to undergraduates and those gaining their English credentials to study at the university, business people, diplomats, and EFL teachers. No matter what the language or the kind of student, songs have continued to be popular and in demand.

Just across the river from Ottawa in Gatineau, at the Universidad Nacional Autónoma de México – Escuela de Extensión en Canadá (UNAM-ESECA), I found myself not only using songs in my regular language classes, but I led workshops for Mexican teachers on how to use songs in the foreign language classroom. I also gave these same teachers separate workshops on how to teach pronunciation. The final assignment for the training on songs required each teacher to choose a song and then develop three activities with which to use it in their classes back in Mexico. Some of the teachers decided to develop a pronunciation activity for their chosen song. Before they left to go back home, I compiled all of the corresponding materials and activities for their songs into a book for them to use with their students. All in all, there were 48 songs and over 254 activities that they could use with those songs. The student-teachers appreciated the

instruction and take-home materials and this became even more apparent in their answers to a questionnaire that I sent to them after two years. The teachers commented on the different ways that they found music to be helpful. What follows is a sampling of some of their general comments:

“Using the songs in English classes, it has helped to students to increase their capacity in many ways...”

“Using songs as part of learning English is fun and kids relate to the songs, plus they are catchy and easy to remember, improves their pronunciation and vocabulary. An depending of the song, they have a good message that sticks.”

“The use of songs in the classes has helped me to motivate their learning into the foreign language, even with unwilling students.”

“To me, the activities with songs (that they like, or that they are prone to like) is awesome for my students. It is as they get in to a door that takes them into a friendlier world in which they feel an emotional conection.”

“When I worked with a song I have noted that most of the students like to learn English in this way.”

“I’ve gladly found that my students interest in English grow by the day, and I occasionally drop a tear when they offer to sing a song in a school event, without being self conscious or ashamed. Thank you, from the bottom of my heart.”

Meanwhile at the University of Ottawa, I was also using songs from time to time in my Spanish and ESL classes and even had the freedom to design and teach not only two separate course components in which songs formed a part of each day's lesson, but also part of a pronunciation course in which I included songs as a material. Again I received positive comments from students regarding how helpful they were for them as a language-learning tool:

"I like the songs because it helped us to understand native speakers in Canada."

"It was a new style and unique way of teaching it help us a lot to learn more new words and improve our accent."

"In general, I find this course is excellent. I like the way to learn english by listen to some song and find the vocab. Also I like the way we learn grammar by song. I think the course was useful."

"It is a very helpful class and you have a good way to teach us. Lesition [sic] to the songs is the best part of the class, because it help me how to listien [sic] clearly and how to understand each word."

"Spanish music videos helped me enjoy learning a new language and learn more effectively."

"I like that you put videos in the beginning and during the course. They are interesting and make us learn more on the culture."

"Thanks for providing those video clips in your lectures – they made class much more interesting and engaging."

Finally, after many years of teaching many students in a number of different countries and contexts, I became more and more convinced that songs could be a powerful teaching tool and one that might help students reach their language-learning goals. As a language learner, pronunciation is, for me, a fundamental skill that can greatly enhance or hinder one's life in a foreign language environment. This has led me to want to learn more about how songs might be an effective tool for teaching pronunciation. Therefore, in this thesis, I decided to study the use of songs in the ESL / EFL classroom as a means of teaching pronunciation within the framework of Communicative Language Teaching (CLT).

Pronunciation and Its Neglect in L2 Teaching and Research

Pronunciation is a language skill which involves the oral production of sounds and groups of connected sounds that pertain to language. When these sounds are articulated in such a way that a listener can understand what is being said (assuming an accompanying adequate use of grammar and vocabulary), then communication may be achieved. Indeed, according to Morley (1991), "Intelligible pronunciation is an essential component of communicative competence" (p. 488). For Setter and Jenkins (2005), "Pronunciation is the major contributor to successful spoken communication" (p. 13). Another crucial element of communication, however, is the ability of the listener to correctly interpret the sounds and sound patterns. As Wong (1993) points out, "listeners need to know how speech is organized and what patterns of intonation mean in order to interpret speech accurately" (p. 46). Otherwise, a breakdown in communication can occur.

Despite being an important language skill, in the area of language teaching, pronunciation appears to have been somewhat disregarded over the years. Kelly (1969), in her book, *25 centuries of language teaching: an inquiry into the science, art and development of language teaching methodology, 500 B.C.-1969*, referred to pronunciation as the “Cinderella” area, indicating that it had been neglected in language teaching. Thirty years later, Fraser (2000) suggested that this might still be true, noting that second language pronunciation “has been out of fashion for some decades” (p. 35). In their discussion of the various language-teaching approaches that have been used from the 1940s to today, Celce-Murcia, Brinton, & Goodwin (2010) illustrate how the teaching of pronunciation has come in and out of favour depending on the teaching method in vogue at the time. Even since the emergence of Communicative Language Teaching (CLT), though, pronunciation has not generally received the attention it deserves (Gilbert, 2010; Morley, 1991; Wei, 2006). Despite Morley’s statement that “it is clear that pronunciation can no longer be ignored” (1991, p. 513), other authors have recently indicated that English pronunciation continues to be neglected (Cheng, 1998; Wei, 2006).

For someone who wants to learn pronunciation, its apparent disregard in language teaching is difficult to understand, given its importance for achieving communicative competence. Wei (2006) indicates that “anyone who wants to gain communicative competence has to study pronunciation” (p. 2) and Cheng (1998) reminds us that “in order to make oneself intelligible and to understand the spoken language, one must have a good working knowledge of the pronunciation of that language” (p. 41). Although there are some students who seem to have a natural aptitude for pronunciation, they still, presumably, develop their pronunciation skills in their own way. Nevertheless, as mentioned above, pronunciation teaching has been somewhat

overlooked, and according to Wei (2006), “has fallen far behind that of the four basic skills in English” (p. 17).

With regard to research, Fraser (2000) mentions “a great need for increased scholarly research on ESL pronunciation and ESL pronunciation teaching” (p.2). In addition, Derwing and Munro (2005) indicate that it has been marginalized in applied linguistics and the results of some research that has been done are often not cited or interpreted for use in teaching. They say that “much less research has been carried out on L2 pronunciation than on other skills” (p. 380) and add that although there is more and more literature on L2 speech in speech production and perception journals, “this work is rarely cited or interpreted in teacher-oriented publications” (p. 382). They go on to indicate that reasons for this include the inaccessibility of the material for those without specialized knowledge of phonetics and the questionable practical nature of studies conducted in laboratory settings as opposed to the classroom.

Reasons for the Lack of Pronunciation Instruction in CLT

Most people learn a language in order to speak it and yet they are rarely taught how to speak it from a phonetic point of view. There are many reasons for this.

On the one hand, it may be assumed that learners will naturally acquire acceptable pronunciation through the course of their studies or contact with the language – even though this only tends to occur with learners that have begun their language training at a young age (Long, 1990). There are, however, researchers who contend that native-like pronunciation is possible for learners who do not begin their training as children (Bongaerts, T., Susan M., & Slik, F.,

2000).² Therefore, in the case of adolescent and adult learners, the process of achieving an intelligible or comprehensible pronunciation is often unknown and left to chance.

On the other hand, many instructors still tend to devote little time and space to pronunciation training in their classes (Breitkreutz, Derwing & Rossiter, 2001; Foote, Holtby & Derwing, 2011; Wei, 2006). Reasons for this vary and reflect the fact that teachers have quite different points of view regarding pronunciation. Derwing and Munro (2005) indicate that “as a result of pronunciation’s marginalized status, many ESL teachers have no formal preparation to teach pronunciation” (p. 389). Moreover, other researchers suggest that some instructors perceive pronunciation to be less important than other skills (Foote et al., 2011; Lin, Fan & Chen, 1995), or may not consider pronunciation teaching to be effective (Foote et al., 2011). Of those instructors who do believe in pronunciation training, some feel that they do not have available to them appropriate techniques, strategies and classroom activities (Henderson et al. (2012); Kirkova-Naskova et al. (2013); Wei, 2006). Investigating reasons for this in a survey of ESL teachers, Darcy, Ewert, and Lidster (2012) indicate a lack of available time and training as well as a need for more institutional support and guidance. Similarly, Fraser (2000) found that many instructors lacked confidence regarding teaching pronunciation, some because of a lack of training and others because of a use of ineffective methods. Breitkreutz et al. (2001) and Foote (2011) also noted that there were teachers who expressed a desire for more training specifically in teaching pronunciation. Derwing and Munro (2005) refer to “the lack of attention to pronunciation teaching in otherwise authoritative texts” as resulting in “limited knowledge about how to integrate appropriate pronunciation instruction into second language classrooms” (p. 383). Furthermore, Gilbert (2010) indicates that it has not been integrated into course books or

² The concepts of intelligibility and nativelikeness will be discussed in Chapter Two.

most language programs. Foote et al. (2011) state that “it is clear that very few universities devote a whole course to L2 pronunciation” (p. 22). Finally, Kirkova-Naskova et al. (2013) indicate that this may be “a global problem” since EFL and ESL teachers in Europe, USA, Canada and Australia “feel they lack training in how to teach pronunciation” and as such, might result in a reflection on ways in which teacher training programmes can be improved (p. 37).

In sum, we see that there are many reasons why there is a lack of pronunciation instruction in the language classroom. On the one hand there is the assumption that learners will just pick it up, while on the other, many teachers are inclined to spend little time on pronunciation teaching for the following reasons: a lack of formal training; a belief that pronunciation is less important than other skills; the idea that pronunciation teaching is ineffective; a lack of suitable techniques, strategies, and activities; a lack of time; a need for more institutional support and guidance; and a lack of knowledge of how to integrate pronunciation into the lessons given the lack of information in course books and language programs.

This lack of attention to pronunciation can be traced to the current major trend in second language teaching. Since the 1980s when Communicative Language Teaching (CLT) eclipsed such methods as Audiolingualism, the Direct Method and the Silent Method, among others, pronunciation had been virtually ignored until very recently, in the last decade or so, during which there has been a renewed interest in the teaching of pronunciation. The ability to communicate through the integration of the four major language skills (listening, speaking, reading and writing) became the focus of language teaching in CLT and the insistence on phonetic and grammatical accuracy, which the previous methods required, ceased to be as important so long as communication was achieved.

Furthermore, as CLT has evolved, it has been difficult to situate pronunciation teaching within the approach, since, according to Richards (2006):

Current communicative language teaching theory and practice draws upon a number of different educational paradigms and traditions. And since it draws on a number of diverse sources, there is no single or agreed upon set of practices that characterize current communicative language teaching (p. 22).

With regard to pronunciation specifically, Fraser (2000) says that “there has been little development of a communicative method of pronunciation teaching” (p. 25) and Celce-Murcia et al. (2010) echo this by saying that “most proponents of this approach have not dealt adequately with the role of pronunciation in language teaching, nor have they developed an agreed-upon set of strategies for teaching pronunciation communicatively” (p. 9). With this lack of consensus in CLT in general, and pronunciation in particular, knowing its position and teaching it have become elusive goals.

Typically in CLT, learners are taught a foreign language without actually learning the sound system of that language. That is, the sounds and sound patterns of the language are generally not explicitly taught. Since the operative word in Communicative Language Teaching is “communicative”, students learning a language with this approach are predominately given activities that are communicative in nature, that is, tasks that require that they interact with fellow classmates. Any pronunciation areas that may be taught will be to facilitate the communication of ideas using the relevant grammar and vocabulary related to the topic or theme of the lesson, but pronunciation areas *themselves* will not form the topic or theme. By learning a language this way, it is difficult for students to progressively acquire a general phonetic and

phonological understanding of the language the way they might eventually be able to do with grammar. For example, in (communicative) four-skills textbooks, grammar (verb tenses and verb aspects, as in the *perfect* or *progressive* forms) tend to be integrated into the chapters, whereas pronunciation is “addressed in a scattered way, [and] as an ‘add on’” (Gilbert, 2010, para. 26). Derwing and Munro (2005), in reference to pronunciation textbooks and software, for their part, contend that “most materials have been designed without a basis in pronunciation research findings” (p. 388).

In second-language environments, where students may have different L1s, the pronunciation areas covered in CLT may or may not pose a problem, depending on the sound system of the students’ native languages. Besides, with there being more of an emphasis on communication and less on the (phonetic) accuracy of that communication, many language learners end up having little understanding or awareness of their weaknesses in pronunciation. Derwing and Rossiter (2002) found that the participants who acknowledged that their pronunciation posed a problem for communication were worried most about individual consonants and vowels that actually had little bearing on their intelligibility. Furthermore, the participants did not mention prosodic elements as being difficult for them. In another Canadian study of adult ESL learners’ perceptions of their accents, Derwing (2003) found that “many of these ESL students have no idea what their own pronunciation problems might be” (p. 559), while Fraser (2000) indicated a similar level of unawareness in Australia on the part of ESL learners.

To sum up, because of the very nature of CLT, pronunciation teaching has been somewhat left aside. First, in CLT there is more of a focus on communication and less on (phonetic) accuracy. Secondly, since CLT itself comes from a variety of educational traditions

and theories, there is no agreed upon way of teaching pronunciation for this popular method. As a result, pronunciation may be addressed in a rather haphazard way and by teachers who have little or no training in the area. Finally, the lack of attention given to teaching the sound system of a language results in learners not acquiring a general idea of the phonetic features of the target language. In turn, this can result in students being unaware of their specific pronunciation shortcomings.

The Challenge for Learners to Find Help

For those learners who are aware of their pronunciation shortcomings, help might be difficult to obtain. This is because, as mentioned above, many language teachers stem from teacher training programs where articulatory phonetics and pronunciation-teaching training were not addressed in detail. Fraser (2000), when discussing ESL teachers, says that “many have received little training in how to teach pronunciation, and even where they have, many standard methods of pronunciation teaching are less than ideally effective” (p. 8). Breitkreutz et al. (2001) reported that a mere 30% of the teachers interviewed had received any formal instruction on teaching pronunciation, while Foote et al. (2011) reported that only 20% of their teachers interviewed had taken courses on teaching pronunciation. Furthermore, it was found that many teachers did not seem to have sufficient background knowledge or confidence to critically evaluate pronunciation beliefs and practices that are dubious, and 67% of those surveyed mentioned that they wanted more training in pronunciation instruction (Thomson, 2013). Finally, Fraser (2000) said that “Phonetics and phonology courses were gradually dropped from many teacher training programs, and pronunciation was, in general, covered briefly if at all” (p. 33). Therefore, when

pronunciation fell from grace in the 1980s, not only were ESL/EFL programs affected, as we saw earlier, but also TESL/TEFL programs. The effect of this downturn is still apparent: Derwing has indicated that “there are very few TESL programs that offer a full course in teaching pronunciation” (2010, p. 27), and according to Wei, pronunciation is “still neglected or ignored at many universities and colleges around the world” (2011, p. 1). In sum, for students with pronunciation difficulties, help could be hard to find. There may be language teachers who, due to a lack of training, are not in a position to be able to effectively help their students or there might be little room for pronunciation lessons in many language programs as they exist now.

In addition to a lack of opportunity to study pronunciation at universities, it is important to mention that there are institutions that do not require correct or adequate pronunciation in official proficiency tests. This is not to say that this occurs everywhere, but it does, for example, happen in the capital of Canada. The University of Ottawa, for example, does not oblige foreign students to take or pass the speaking component of the CanTEST (a Canadian English proficiency exam) when applying for admission to undergraduate programs. Why this is so is unknown (Dr. A. Hope, personal communication, August 27, 2015). For potential students who opt to take other language proficiency tests such as the TOEFL, IELTS, EPT (MELAB), CAEL, or PTE, the university requires a certain overall score as well as a minimum writing score; however, there is no mention of a minimum speaking score (University of Ottawa, n.d.). Students may regard not having to prove their English speaking skills as something positive; however, by not requiring an oral component, an institution could be sending out the wrong signal to students. That is, that it is not important or necessary to have good second-language speaking skills. The problem is that in real life communicative situations, foreign students do not

just read, write and listen – nor can they be successful if that is all they do. Foreign students have to perform the same academic and everyday social speaking tasks that native speakers do.

Some Consequences of Poor Pronunciation

The marginalization of pronunciation teaching unfortunately does not allow language learners to function successfully in our modern global and multicultural context/world. If their pronunciation in the target language is such that it is difficult to understand or if it is heavily accented with their L1³, they may find themselves disadvantaged. Many people with foreign accents, especially those in second language environments, suffer economic, social and emotional difficulties because of their accents. Baluja (2011), Fraser (2000), Immen (2011), and Lev-Ari and Keysar (2010) indicate the ramifications that a strong accent can have on a non-native speaker's professional life while Derwing (2003) exposes the social and emotional difficulties that they can face in a second-language context. Munro, referring to language learners who are difficult to understand, says that “apart from the personal frustration they may feel, communication difficulties can damage educational and career opportunities. They can also lead to negative social evaluation, such that others avoid interactions...and that can lead to isolation from the host community and lost opportunities to use the L2” (2011, p. 9). If poor pronunciation skills can lead to such negative experiences, then it only makes sense that language learners should be provided with opportunities for better pronunciation instruction that would serve to help them avoid the kinds of consequences mentioned above.

³ A more detailed account of *accent* can be found in Chapter 2: A Review of the Literature, pp. 59-64.

Discrimination

Discrimination on the basis of linguistic accent is part of the reality of language use. Unfortunately, it can degenerate into racist attitudes and the assumption that people with different accents are somehow less language competent. Fraser (2000) says that (compared to other linguistic skills) “Pronunciation is the aspect that most affects how the speaker is judged by others and how they are formally assessed in other skills” (p. 7). Fraser also goes on to say that, in Australia, “many native speakers are consciously or unconsciously prejudiced against those who speak with an unfamiliar accent - whether that is a non-standard native accent, an established non-native accent, or a learner’s accent” (pp. 9-10). Foote (2011) found that 70% of ESL teachers surveyed in Canada agree that “a heavy accent is a cause of discrimination against ESL speakers” (p. 21) and Derwing (2003) found that 30% of the Canadian immigrants surveyed felt that they were discriminated against because of their accents. Munro (2003) reports on instances of accent discrimination that have arisen in Canadian human rights cases and makes a number of recommendations for ESL teachers including the need to be more informed about important issues in phonetics and how they are relevant to learners.

It would seem ironic then that despite the very real problem of accent discrimination, there appears to be a point of view that it is somehow racist to expect people to work on modifying their accents so that they are closer an acceptable local variety. An example is a comment by a former professor of Anti-Racism Studies at the University of Toronto, who says that the idea of reducing accents “speaks to a certain kind of linguistic racism” (Baluja, 2011). On the other hand, not everyone shares such a point of view. A Canada Research Chair in immigration and governance indicates that it is “a kind of vacuous political correctness to suggest that acquiring a local accent won’t provide, in general terms, an advantage in the

business world” (Baluja, 2011). We will discuss in more detail the debate surrounding accent and other types of discrimination in addition to the concepts of accent reduction and addition in “Chapter Two: A Review of the Literature”.

Discrimination aside, whether speakers find themselves in personal or professional situations, if they are from different language backgrounds, they can have tremendous difficulty understanding one another because of their accents. When native speakers (NSs) are the listeners, inadequate suprasegmental pronunciation is what causes most of their problems in understanding, whereas for non-native speakers (NNSs), it is the pronunciation of segmentals (Setter & Jenkins, 2005). Therefore, depending on the communicative context, an L2 speaker’s pronunciation skills will have a bearing on whether or not they are intelligible to their interlocutors.

In a business context, consequences of poor L2 pronunciation skills can be experienced by professionals who are otherwise qualified for employment positions. According to a survey conducted by Compas, hiring executives considered inadequate language skills⁴ to be “the biggest barrier to employment” (2009, p. 9) whereas only 8 out of 91 of the IEPs (Internationally Educated Professionals) who were candidates for hiring reported having difficulties communicating. What this may mean is that the IEPs did not recognise that their language skills were negatively affecting their ability to communicate. Once again, like language learners who are unaware of their pronunciation shortcomings, professionals who are second language speakers might not have developed self-awareness with regard to their ability to communicate. While for some at an ideological level it may seem racist to expect people to approach an acceptable speech norm, such a viewpoint may not help at a practical level, nor is it beneficial

⁴ In the Compas study, while *language skills* is a general expression, *pronunciation* is specifically mentioned a number of times in the comments section (pp. 29-30) when referring to the survey question regarding barriers to employment.

for language learners to be unaware of their linguistic shortcomings, for such ignorance can carry over to their professional lives.

In sum, as we have seen, recent studies in Canada and literature from around the world are bringing these issues to light and pronunciation instruction is beginning to be seen as a language skill that can no longer be ignored as it has been for the last 30 or so years.

How to most effectively teach and integrate pronunciation into curricula needs careful consideration and investigation. In “Chapter Two: A Review of the Literature”, we will see that there are a number of principles for pronunciation teaching which have emerged in the literature. These principles also consider the materials that are used in the L2 classroom. One material, though, that has received very little attention are songs. As we will discover, there are a number of reasons why they could be useful in the pronunciation classroom. Before examining these reasons, though, it is worthwhile to consider how music is a natural part of language.

The Musicality of Language

In the adult L2 classroom, language and music can come together in songs, which when introduced, help to both relax the students and provide a change of pace in the class (Murphey, 1995; Wilcox, 1995). While relaxation and change of pace are positive elements, the usefulness of songs as a tool for language learning can go much further. For this reason, we should consider the inherent musicality of language and the possible important role that it might play in language learning and communication.

A (spoken) language does not exist without a musical element. This idea can be traced back to Saussure (1966) when he refers to the *sound-image* or *signifier*. He stresses that it is

more “intimately united” with the *concept* and is not a matter of vocalizing but rather is the “psychological imprint of the sound, the impression that it makes on our senses” (p. 66). Therefore, this signifier is part of the concept and is revealed as part of the concept. Similarly, the music of a language is intimately tied to the language itself. For us, the only time there is no music in a language is if we see the writings of a language that uses an alphabet or system of symbols unknown to us. In that case, we cannot put the graphic form to sound, whether it be internal or external. Normally when we think of sound, we think of it as something that is heard with our ears, but when we mouth words or have a song in our heads, we hear with our minds. There may be no vibrations reaching our ears, but we are still able to perceive the pitch, tempo and rhythm of whatever would be spoken or sung. In that way, language necessarily has music as an inherent quality.

More recently, Mora (2000) reminds us of the shared features of music and language. She indicates that

both stem from the processing of sounds...both are used by their authors/speakers to convey a message...[both] music and language have intrinsic features in common, such as pitch, volume, prominence, stress, tone, rhythm, and pauses... [and] we learn both of them through exposure (p. 147).

When we understand that musical elements are actually a part of language, the idea of studying a language with the help of music and song may begin to be seen as a logical and reasonable way to approach L2 teaching in general, and pronunciation teaching, specifically. However, in order to more fully appreciate the potential that songs may have, an examination of the different dimensions that are at play in the L2 classroom is necessary.

Dimensions of the Language Classroom

Schoepp (2001) provides a number of reasons – all of which he indicates are grounded in learning theory – for using songs in the ESL/EFL classroom and he groups these reasons into the *affective, cognitive* and *linguistic*. Since we know that an L2 classroom involves teachers, students, and learning, we should recognise that there are necessarily underlying classroom dynamics at work in such an environment. By identifying the affective, cognitive, and linguistic as *reasons* for using songs, Schoepp is effectively referring to these classroom dynamics. If we consider that the dynamics pertain to certain *dimensions* in the classroom, we can see that not only are there affective, cognitive, and linguistic dimensions, but also a social/cultural one, all of which require careful thought while planning and executing classes and courses. Indeed, this is in accordance with Duff (2008) who asserts that L2 learning “involves linguistic, cognitive, affective, and social processes. That is, it is an ongoing interplay of individual mental processes, meanings and actions as well as social interactions that occur within a particular time and place, and learning history.” (p. 37). Therefore if we consider this complexity of an L2 classroom and examine these dimensions, we can see how the use of songs relates to the ongoing processes. That is, by reviewing research that pertains to music and how it relates to each dimension, it will be possible to arrive at a better understanding of the role that songs can play in language classrooms. Through such close inspection, it should become clear that songs could be ideal texts for pronunciation teaching and thus should not be overlooked.

According to Ebong and Sabbadini, “songs provide examples of authentic, memorable and rhythmic language” (2006, para. 1) and Stansell (2005) succinctly indicates some of the ways that music connects the various classroom dimensions when he says,

Music codes words with heavy emotional and contextual flags, evoking a realistic, meaningful, and cogent environment, and enabling students to have positive attitudes, self-perceptions, and cultural appreciation so they can actively process new stimuli and infer the rules of language. (p. 35)

In the discussion that follows, the dimensions are examined separately. However, it is important to keep in mind that in a language classroom, all of the dimensions are simultaneously present and interacting on various levels. What this means is that although we have attempted to divide the dimensions and discuss each one independently, there will be overlap into other dimensions due to the interconnectedness of language and culture and the attitudes toward linguistic registers that arise in different social and cultural contexts.

Songs and the Linguistic Aspect

The linguistic aspect considers the conventionally-named “major” and “minor” linguistic skills of *listening*, *speaking*, *reading*, *writing*, and *grammar* and *vocabulary*. *Pronunciation* is generally relegated to sub-skill status under *speaking*. Songs are authentic and live oral representations of the language, but by also being written, language learning activities involving songs inevitably integrate and exploit the different language skills. Additionally, because they are simultaneously a spoken and written text, they indicate the sound-spelling correspondence of the phonological and graphic elements of the language. Of course other spoken-and-written-texts do the same, but songs especially emphasize the melodic or suprasegmental aspects of the language and, according to Wilcox (1995), help to “establish the prosody of the language” (p.

118) – those pronunciation features that may contribute most to the comprehensibility and intelligibility of its speakers⁵.

Songs may be considered an ideal kind of oral text to listen to over and over, not only because they provide the linguistic repetition that language learners require in order to learn the grammatical, lexical and phonetic features of the language, but also because (if we like a song), it is enjoyable and it is natural to want to listen to it many times. Moreover, they are small, manageable texts that can be used in a class, and which can show different levels of formality in the use of the language as well as different aspects of the culture, depending on the musical genre that is chosen. In the case of pop music, whose genres are dear to many young language-learning adults, songs may be chosen to reflect some of the grammar and vocabulary that ESL students will face while in the target culture. Murphey (1992) found that the language of pop songs is “repetitive” and “conversationlike” (p. 771) and Schoepp (2001) reminds us that “using songs can prepare students for the genuine language they will be faced with” (p. 3) as does Chunxuan (2009).

Many ESL teachers will recognise that a common complaint among newcomer ESL students is that they cannot understand native speakers. While there may be various reasons for this, such as a lack of grammatical or lexical knowledge, some of them can be traced back in part to inadequate or inappropriate listening and pronunciation training.

Natural speech. Too often, levels of formality in pronunciation are not considered and foreign (and even second) language students may be exposed to an artificial kind of English

⁵ It is important to acknowledge that other texts that are spoken and written, such as poems and limericks, can also emphasize the melodic or suprasegmental aspects of the language.

through textbook audios and teacher talk which might not contain all features of connected speech which are the norm in natural speech. Cauldwell (2013) compares controlled or artificial speech to natural speech by using plants as a metaphor. He says that “the speech that learners encounter outside the classroom is more like jungle vegetation than garden or greenhouse plants, much wilder than the forms they encounter in the classroom” (p. 2) and that we need to prepare students for jungle listening. Furthermore, Vandergrift says that if students are not taught how to listen, then they will have trouble learning while listening (2004, 2007). Vandergrift (2007) also stresses the importance of using authentic listening materials not only because they are “relevant to the learners’ lives” but they “reflect real-life listening” and “allow for exposure to different varieties of language” (p. 200). As a teacher who regularly uses textbooks with accompanying CDs and DVDs, I often encounter artificial-sounding and slowly-spoken audio and video files, even in advanced-level materials.

With regard to teacher talk, while it may be conducive to learning at certain times in the classroom situation, too much of it does not serve to prepare students for authentic spoken language. A case in point can be found in Cho and Reich (2008) who suggested that since some students have trouble understanding native speakers of English, teachers should pronounce words with equal stress. This, however, is not reflective of natural spoken English. Exposing students to this kind of artificial stress pattern will not help them understand native speakers who obviously do not speak this way. Mora (2000) mentions how the exaggerated melodic contours of teacher talk tend to occur when new structures are being presented for repetition purposes or when the teacher is modelling the language for correction purposes. When used in this way and for specific purposes, teacher talk can indeed be helpful, especially when it serves to direct the learners’ attention to specific pronunciation patterns. However, if teachers otherwise do not

speak at a natural pace which includes reductions and other connected speech phenomena, they are not adequately preparing their students for authentic spoken interaction.

In natural spoken discourse, sounds and words do not occur by themselves but “run together” in *connected speech*. Furthermore, in connected English speech, some syllables are squeezed between stressed ones so that rhythm can be maintained (Celce-Murcia et al., 2010). When syllables are squeezed, changes in sound occur which are referred to as *reductions*. According to the authors, reductions in connected speech can take on many different forms: from the “disappearance of a sound” (deletion), to the “distortion of word boundaries” (contractions, blends and reductions), to “the smooth connection of sounds” (linking), to “the change in adjacent sounds” (assimilation, dissimilation) to, finally, the “addition of a sound” (epenthesis) (p. 164). Songs, as authentic texts, contain these features of connected speech, and when used appropriately, can help language learners recognise that the segmental changes that occur are normal, natural, and acceptable.

What is interesting is that, as an ESL / EFL teacher with 19 years of experience, I consistently meet resistance from students when trying to encourage them to use reductions in English. When asked why they do not use contractions, for example, learners from different language backgrounds tell me that they were taught that it was “bad English”. This corresponds with Miyake (2004) who mentions that students consider such speech to be “‘lazy’, ‘casual’ or ‘sloppy’” (p. 74). For students to have such notions of spoken English, one can only assume that some of their teachers must have used a lot of teacher talk with them. Nevertheless, contrary to what some language learners or teachers may think, the occurrence of multiple types of reductions in everyday spoken English is the norm rather than the exception.

Levels of formality. Weinstein (2001) mentions three levels of reduced speech, of which the third and most informal level is by far the most prevalent. The first level involved no pronunciation change from what was written, for example, “What do you, what do (we / they), what are you”. The second involved one change: “Whadda you; wha do (we / they”; what’re ya”. The third involved more than one of the above types of (reduction) changes: “whaddaya; whadda’ whaddaya” (p. 119). The difference in frequency of occurrence of reductions among the three levels was 8, 47 and 258 times. Weinstein, indicates that the third level occurs 82% of the time, whereas the first and second occur 3% and 15%, respectively.

Songs from many popular genres reflect this natural, authentic level of pronunciation, and thus their use enables students to decode and use features of connected speech. Schön et al. (2008) found that songs can help beginning learners perceive word boundaries, while Van den Berg (2011) indicates that songs are useful for helping learners perceive and produce suprasegmentals. According to Chunxuan (2009), “songs serve as a medium through which these [phonetic] rules can be made concrete and accessible” (p. 92) by internalizing them through repetition and imitation. Murphey, when considering the doubts that instructors might have regarding the use of songs beyond just singing says, “whatever you can do with a text, recording, or film, you can probably do with songs” (1995, p. 69). That is, Murphey suggests that songs can be considered serious, academic texts that can be used for much more. When employed properly, songs can help learners appreciate the different levels of formality of the language.

Language awareness. Songs, if used appropriately as a text, have a way of making learners pay attention to the language, which can help to further promote the learning process.

This is what is referred to as *language awareness*. “Language awareness is a mental attribute which develops through paying motivated attention to language in use, and which enables language learners to gradually gain insights into how languages work” (Bolitho et al., 2003, p. 251). The authors indicate that the main goal of gaining language awareness is that learners will notice the (target) language and notice how their use of it differs from that of native speakers, which will then lead them to achieve a “learning readiness” (p. 252). Schmidt (2010) refers to this same notion as “‘noticing the gap’, the idea that in order to overcome errors, learners must make conscious comparisons between their own output and target language input” (p. 724). Murphey (1995) mentions that, if exploited creatively, songs can help with learning the sound system of the language as well as be a way to “bridge the gap between the pleasurable experience of listening/singing and the communicative use of language” (p. 6).

The Affective Aspect

With respect to the affective aspect, which deals with emotion and how it influences behaviour, many authors indicate that anxiety, self-esteem, interest and motivation⁶ can be improved by using songs in the language classroom (Adkins, 1997; Anton, 1990; Baoan, 2008; Cheng, 1998; Chunxuan, 2009; Francis, 2010; Lê, 1999; Miyake, 2004; Murphey, 1995; Murray, 2005; Pardede, 2010; Richards, 1993; Schoepp, 2001; Van den Berg, 2011). As for anxiety, language learners, as students, are often under tremendous pressure to succeed when learning a new language. Whether they are foreign or domestic students, their parents might be under

⁶ With regard to motivation, we recognise that there are different kinds of motivation and that this is a psycholinguistic topic that we cannot fully address in this thesis (see Dornyei & Ushioda, 2011 and Hadfield & Dornyei, 2013). We should point out that in this study, reference is made largely to emotional or internal motivation (see Williams & Burden, 1997). As will be seen from our results, it turned out not to be a substantive distinguishing factor among our participants, given the high level of motivation that they started out with.

considerable financial stress if they are supporting them or paying their tuition. Furthermore, students might also be under time constraints to achieve a certain level of English in order to pursue their studies or advance in their employment. In addition, communicating in a new language can be a frightening experience, and one in which the ability to do it well may decrease as one's anxiety level increases, especially if the anxiety is what Scovel (1978) refers to as *debilitating* anxiety. According to the author, debilitating anxiety can have a negative impact on student learning whereas facilitating anxiety can be motivating. Adding music to the classroom may help to reduce the level of debilitating anxiety that students feel in the classroom while at the same time help to take their minds off their troubles. After all, according to Stansell (2005), it is a well-known fact that music has the ability to affect our emotions and consequently our experience of events involving music.

Motherese. Murphey (1995) refers to songs as “adolescent motherese.” Motherese, more commonly known as *baby talk* is the “highly affective and musical language that adults use with infants”, which provides teenagers (and adults) with the “affective attention” they may be lacking (p. 7). While affective, the main goal of motherese is communication (Murphey, 1985). This communication can be felt in songs; the high prevalence of first and second personal pronouns means that the song “creates a situation...and involves the listener in a type of pseudo-dialogue or conversation’ (p. 794). Furthermore, Murphey states that since “most pop songs (and probably many other musical genres) do not have precise people, place or time references”, the songs can be “appropriated by listeners for their own purposes” (p. 8). In other words, learners can find meaning in songs that is personal to them.

Motivation. Music can be used to enliven, subdue or simply engage learners, and in language teaching, it has been pointed out by both (adult) students and teachers that songs are a motivating force for learning which helps enable them to participate more and with less fear (Lê, 1999). Chunxuan (2009) says that “listening to songs can knock down the learner’s psychological barriers, such as anxiety, lack of self-confidence and apprehension as well as fire the learner’s desire to grasp the target language” (p. 94), while Celce-Murcia et al. (2010) indicate that songs “can be an excellent tool for motivating learners and practicing pronunciation” (p. 353). Furthermore, Chunxuan (2009) indicates, “Songs are abundant in themes and expressions which will echo in the learner’s heart” (p. 88). Thus, with respect to the affective aspect in the language classroom, when students are learning with songs, they are interested and engaged, and the learning is more enjoyable and less stressful. Many authors assert that it is the motivation to learn which will play a key role in their eventual success or failure in learning the target language (Cheng, 1998; Chunxuan, 2009; Schön et al., 2008; Van den Berg, 2011). The fact that songs are motivating for students is indicated by a number of authors (Adkins, 1997; Baoan, 2008; Borland, 2012; Chunxuan, 2009; Ebong & Sabbadini, 2006; Lê, 1999; Murphey, 1995; Stansell, 2005; Van den Berg, 2011).

The effect that songs can have on a learners’ level of anxiety and motivation are important. MacIntyre (2007) indicates that anxiety and motivation are factors that play a role in a willingness to communicate (WTC), but that in addition, at any given time learners are also simultaneously subject to “driving” and “restraining” forces, i.e. “the interplay of the features of the situation with the psychology of the individual speaker” that will affect whether or not the learners decide to communicate (p. 573). Therefore, while anxiety and motivation are important, there

are, understandably, other factors at play, such the social dynamics of the situation, which we will examine as we continue our discussion of the dimensions of the language classroom.

Meaningful learning material. In many cultures around the world, music permeates daily life and takes centre stage during holidays, traditions and celebrations. Murphey (1995) reminds us of the many places where music and song appear in our lives (e.g. hospitals, restaurants, cafés, malls, sports events, and cars). Being rich in culture and history, musical pieces and songs are used as a learning material for children. Ironically, however, as learners grow up and become more and more self conscious, songs tend to disappear from, at least, the North American classroom and are replaced by other materials, primarily written ones. While there is obviously nothing wrong with using written materials, in our view it is a mistake to abandon the use of combination oral-written genres like songs. For one, songs can bring to the classroom meaningful and widely-known examples of the culture such as Christmas carols, jingles, and children's and folk songs which can help students feel a sense of understanding and belonging in the target culture.

The themes in songs can be meaningful to learners and reflect everyday spoken language, yet there are also many songs that can be found to reflect written forms of language and which address a variety of topics. It is important that language teaching keep up with the rapid evolution of language and how language is now being used in technology for quickly and efficiently spreading information and socializing virtually. Tweeting, texting and Facebooking are part of a written genre that is a fusion or crossover of written and spoken language. This is a kind of language and register that young adults especially want to learn – but most importantly,

need to learn to communicate appropriately in informal contexts. Schoepp (2001) reminds us that “the majority of language most ESL students will encounter is in fact informal [and that] using songs can prepare students for the genuine language they will be faced with” (para. 11). Moreover, Couper (2002) found that ESL students themselves indicated the need to learn informal spoken language and pronunciation. Songs can also be examples of this efficient and popular use of language and are therefore meaningful to learners.

As already indicated, the informal, colloquial, and non-standard language found in some songs can be a way of introducing learners to linguistic forms that they would not normally encounter in their language classes but would in their daily lives or online when they are socializing. Words that undergo reductions and are informally spelled the way they are pronounced, such as “wanna”, “gonna”, “hafta”, for example, regularly show up in songs. This graphic representation of “want to”, “going to”, and “have to” are often written this way in the above-mentioned forms of social media. Part of communicative competence involves being able to shift one’s register according to the sociolinguistic variables of the communicative context. If students are exposed to and taught about colloquial language forms, it will help them be more competent and confident in more informal contexts. While this is certainly important, teachers need to keep in mind what the course objectives are. For example, if students are taking an English for Academic Purposes (EAP) course, time spent on teaching aspects of informal registers should probably be minimal and clearly related to future academic situations the students would encounter.

The Social and Cultural Aspect

The social and cultural aspect of the language classroom includes how the students and teacher relate to one another. The role of culture itself is an important part of this dimension because the cultural (and educational) background and interests of the individuals involved influence the social dynamics in the classroom as well as the kind of teaching that might be perceived to be most effective.

A unifying influence. In my experience as a language teacher, I have used songs with a variety of groups of students, such as homogeneous groups of Muslims or Catholics, as well as heterogeneous multinational groups consisting of both religious and nonreligious learners. In every course that I have used music, I have found that songs are the one material that brings students together, that is, songs help create a feeling of solidarity in the classroom. This is in accordance with Stansell (2005) who says that “songs help to relax and unify a class” (p. 34) as well as Lê (1999) who indicates that they enhance social harmony in the classroom. Murphey (1995) indicates that they “encourage harmony within oneself and within a group” (p. 8). Lê (1999) says that “music bridges the gap between teachers and students” (p. 5). Although every class and student may be different, music is common to all cultures and is something that can bring people together. As Francis (2010) indicates when referring to singing, it is “good for classroom ‘cohesion’” when they are “all working together on the same thing” (p. 159). Singing aside, when students engage with music, they engage with each other and in this way music is a true communicative tool. To adolescent and young adult language learners, music may very well

be the most interesting aspect of the target culture, no matter where they are from, and this shared interest can bring together students who may have little else in common.

For most language learners, popular songs have been seen to be an interesting, relevant and accessible genre and there are surveys that indicate that students prefer to have teachers that use such music in the classroom (e.g. Lê, 1999). Lê (1999) found that it “bridges the gap between teachers and students” and had student testimonials saying that they liked to be in class with a teacher who is not so serious and who likes to sing (p. 5). This is important because it is possible for generation gaps and culture gaps between teachers and students to negatively affect the learning and teaching environment in the classroom. In my experience in both ESL and EFL settings, when faced with groups of students from cultures that I had not taught before, there was always a “getting-to-know-you-stage” for a period of time as my students and I learned and understood better how to relate to each other. In addition, teachers age, and in some way may find each new generation of learners a little more different or distant from the last, unless they are able to relate to the interests of the students. Yet, when students and/or teachers are from very different cultures, music can nevertheless be a common, unifying thread which helps to remind us that we are not so different after all and that we have shared interests, no matter how old we are or where we come from.

There are also specific reasons why this power of music to promote sharing and social bonding can be especially important in an L2 environment, for instance the phenomenon of *anomie*. Lake (2002-3) mentions that a number of ESL students experience *anomie* which he describes as a “feeling of homelessness” in which students “feel cut off from their native cultures and find a struggle in adapting to a new culture” (p. 2). Brown (2000) refers to this as a symptom

of culture shock in which there is “a feeling of homelessness where one feels neither bound firmly to one’s native culture nor fully adapted to the second culture” (p. 184).

We have seen from the authors cited above that songs can not only bring people together and contribute to a feeling of belonging, but also that songs are rich sources of cultural information which can help students better understand their new world and the people in it. By using music in the classroom, teachers can simultaneously create a positive social experience while introducing students not only to the culture, but also to an aspect of the culture that they may be most interested in.

Facets of culture. When students sing songs in another language, they are participating in the target culture (Francis 2010). According to Stansell (2005), “music also makes cultural ideas accessible to students” (p. 34). This is extremely important because culture shock, cultural ignorance and culture discrimination⁷ can be considered barriers to learning. Songs, however, can help with this because, depending on the songs chosen and how they are used, they can be a rich source of linguistic and cultural material and are live and authentic representations of the language and culture to which they belong. Murray (2005) says that “music can provide L2 students with a unique and exciting opportunity to both explore the language and culture of a foreign country” (p. 5). Indeed, according to Chunxuan (2009), “language and music are interwoven in songs to communicate cultural reality in a very unique way” (p. 88). Lê (1999) and Baoan (2008) discuss how, in Vietnam and China, along with the rise of interest in English, there is an interest, on behalf of English learners, in the cultural forms associated with the

⁷ I use *culture discrimination* to refer to a bias against other cultures, especially the practices of other cultures.

language, namely, music. Since language and culture are inseparable, it seems strange that songs of the culture would not be regularly included in the curriculum. Furthermore, for language teaching to be truly student centred, it makes sense for more music/songs to be included. According to the author of *The ESL Songbook*, Adamowski, (1997) “popular songs show cultural behaviour and cultural attitudes” and in this way the culture enters the classroom through the language (lyrics) contained in songs (pp. x-xi). She adds that discussing these cultural aspects can help the many newcomers who may feel resistance to the target culture be more at ease. For many, the exploration of songs could very well be a key to the target culture, one that paves the way for (linguistic) inquiry, experimentation, participation – and, ultimately, understanding.

The Cognitive Aspect

The cognitive aspect concerns variables related to the functioning of the brain in relation to the process of language learning, and automaticity, general intelligence, and memory fall under this aspect (Richards & Schmidt, 2002). Traditionally it was thought that language and music processing occur in different parts of the brain. However, today, neuroscientific research indicates that music and language processing are not entirely separate after all, but rather that the different parts of the brain collaborate or task share. According to Levitin (2006), “musical activity involves nearly every region of the brain that we know about, and nearly every neural subsystem” (pp. 85-86). He adds that when we listen to or remember the words to a song, different language centres of the brain are involved. This makes perfect sense if we keep in mind the fact that spoken language has a musical, i.e., suprasegmental component to it that co-exists with and is inseparable from the segmental component. Therefore, by bringing song into the

classroom, students at the very least are provided with an opportunity to improve their suprasegmental listening skills.

The Din and the SSIMH. The Din and the Song-Stuck-in-my-Head (SSIMH) phenomenon are two mental processes that occur without our conscious control. Murphey (1990) refers to the Din as “involuntary rehearsal of a foreign language in one’s mind” and the SSIMH phenomenon (which he coined) is the involuntary mental rehearsal of a song (p. 53). (The layman term for the SSIMH phenomenon is *earworm*.) Murphey explains that both the Din and the SSIMH are linguistic input and can both lead to linguistic intake on the part of the learner. The differences, he states, however, are that the SSIMH phenomenon does not require comprehensible input and can be triggered within a few minutes whereas the Din does require comprehensible input and takes one to two hours to start.

Automatic and fluent output. Regardless of whether or not a song causes the SSIMH phenomenon, there are authors who indicate that the repetitive nature of songs, in terms of both the language as well as their tendency to be listened to over and over can lead to automatic and fluent output on the part of the language learner (Celce-Murcia et al., 2010; Murphey, 1990; Schoepp, 2001). Schoepp (2001) indicates that “using songs can help automatize the language development process” (p. 2) by providing learners with the opportunity to rewrite the lyrics so that they are more meaningful to them. Murphey (1990), when discussing lyrics that cause the SSIMH phenomenon is referring to a process that can lead to automaticity. What is interesting in such a case is that, as Salcedo (2010) points out, a learner is involuntarily rehearsing the

language to the sound of another person's voice. As a linguistic model for students, songs are a material that they have at their disposal to hear correct pronunciation of specific sounds and patterns over and over again without becoming bored – especially if they like the song. Furthermore, Salcedo (2010) indicates that rather than students hearing their own (incorrect) pronunciation while reading texts, repeatedly listening to the (correct) pronunciation via songs could improve the potential for better pronunciation. Therefore songs not only have the potential to provide the language learner with entertaining and meaningful input, but repeatedly hearing the correct pronunciation of the language structures and vocabulary may allow them to internalize the language, that is, commit it to memory – and then be creative with it. At the same time, because songs incite us to sing along, learners are constantly encouraged with the use of songs to practice producing the language which can help foster automaticity. In fact Schoepp (2001) says that automaticity “is the main cognitive reason for using songs in the classroom” (p. 2). Celce-Murcia et al. (2010), echo this when saying that “songs can give learners chunks of language they can produce with automaticity – rapidly without pausing – rather than word by word” (p. 353).

Intelligences and learning. Using songs in the language classroom can also be important for general intelligence and the ability to learn more efficiently. According to Gardner and his theory of Multiple Intelligences, people have a number of different types of intelligences that are developed to one degree or another (Gardner, 1993). One of these intelligences, *Musical intelligence*, which is more clearly explained by Smith (2002, 2008), “involves skill in the performance, composition, and appreciation of musical patterns. It encompasses the capacity to recognise and compose musical pitches, tones, and rhythms... [and it] runs in an almost structural

parallel to linguistic intelligence” (para. 15). Gardner (1993), though, points out the communicative aspect of music when quoting Stravinsky who said,

When I compose something, I cannot conceive that it should fail to be recognised for what it is and understood. I use the language of music and my statement in grammar will be clear to the musician who has followed music up to where my contemporaries and I have brought it” (p. 104).

Gardner continues with this line of thought when he refers to a composer who “attempts to create a new idiom” and a listener who “tries to make sense of nursery rhymes” as a way of indicating the different degrees of skill involved in musical intelligence (p. 104). That is, he states that composing is more difficult than performing which is more difficult than listening. What is important here is that he indicates that communication can occur with music *itself*⁸, and not only with the music (prosody) and words of language.

Since people have different cognitive learning styles based on different intelligences (Silver, Strong & Perini, 2000), using a combination of materials that address these different intelligences simultaneously makes sense because then the learning will be more integrated and holistic. Christison (1996), indicates that “the traditional second or foreign language classroom has favored visual and verbal delivery systems,” which is not ideal for students that “exhibit other intelligences” (p.10), while Saglam et al. (2010) contend that “these intelligences can be nurtured and strengthened” (p. 12) if lessons are designed to address the different intelligences. Therefore, using songs in the language classroom should facilitate the learning process for more

⁸ Exploring the communicative nature of music itself could help us better understand how this aspect of music can add to the language learning experience. This, however, is a topic which falls outside the scope of this thesis.

students because the combination of both language and music will appeal to more learners (both cognitively and culturally) while strengthening both of these intelligences (musical and linguistic) at the same time.

What this means specifically for pronunciation is that through the use of song, learners can become more adept at recognising and hearing the various segmental and suprasegmental features of connected speech. Richards (1993), an elementary school teacher, asserts that “general classroom music activities that include singing and rhythm help enhance the development of auditory discrimination skills, including integration of letter sounds, syllabification, and pronunciation of words” (p. 109). In the case of adult learners, having an awareness of correct speech patterns is important for developing appropriate pronunciation, especially since adult learners will typically have much more difficulty acquiring acceptable pronunciation skills than will children (Flege, Munro & Mackay, 1995). Listening to (and even singing along to) songs does not require the cognitive load that is involved when using the language to communicate because we are not struggling to formulate ideas, find words, or use the correct grammatical structure. If, while listening to a language, the pressure to comprehend and the need to communicate specific messages are removed, it would make sense that the brain would be freer to direct its attention to the actual sounds.

The Noticing Hypothesis. Paying attention to the target language is especially important for adult learners. Schmidt (2010), discusses his Noticing Hypothesis (Schmidt 1990, 2001) which is that “input does not become intake for language learning unless it is noticed, that is, consciously registered” (p. 722). By choosing and using songs carefully, a teacher can direct

students' attention to specific pronunciation phenomena that pose a challenge to the learners. Whether it is with a new song or one that students are already familiar with, teachers can help learners become aware of and take in previously unnoticed areas of pronunciation, especially because songs lend themselves to being listened to repeatedly. According to Schmidt (2010), paying attention is very important for language learning, and he refers to it as "the pivotal point at which learner external factors (including the complexity and distributional characteristics of input, the discoursal and interactional context, instructional treatment, and task characteristics) and learner internal factors (including motivation, aptitude, learning styles and strategies, current L2 knowledge and processing ability) come together" (p. 731). If paying attention is indeed as important as he says, then it would make sense for instructors to employ materials such as songs that would facilitate this.

Memory. Regarding memory, Stansell (2005) indicates that if songs are carefully chosen so that the lyrics and rhythm are properly paired, then this helps the mind to remember the song. This is important for language learning because the linguistic structures have been taken in. Lake (2011) says that "adding rhythm and melody to chunks of language invites rehearsal and transfers words into long-term memory" (para. 30). Adkins (1997) mentions the same, adding that "a song is 'chunked' with rhythm and rhyme [and that] 'chunking' material means that the ideas are broken down into memorable segments and when these 'chunks' are rhythmical, so much the better" (p. 10). Since songs can carry linguistic information which is taught in the language classroom, choosing a song that has the appropriate combination of rhythm and language can have much pedagogical value. To lend further support to the idea, Levitin (2010) reminds us that this is something that our ancestors have known for ages. He indicates that if we

draw upon the knowledge of anthropologists, we can see that historically, before written language existed and became widespread, music was an effective tool for remembering things and passing down important information from one generation to the next. Furthermore, he reminds us that although humans have been on Earth for fifty to a hundred thousand years, written language has only existed for about five thousand years. He remarks that, since it was necessary for our ancestors to pass down knowledge from one generation to the next, music was a good way to transfer that knowledge, adding that remembering things is easier when they are set to music.

Finally, Saglam et al. (2010) make an important point that is relevant to the above discussion of the various dimensions of the language classroom: “If music plays such a vital role and has such a deep reaching effect on people’s mood, motivation, emotions, socialization, behavioral outcome and involvement, all of which are crucial components in education and in particular, in language learning, it is ironic that music still [sic] waiting at the threshold of the classroom to be invited in” (p. 2). Engh (2013a, 2013b) found as well that although there is ample theoretical support for the use of songs in language teaching, in practice it is quite limited. Furthermore, he states that “from an educational standpoint, music and language not only can, but should be studied together” (2013a, p. 121).

To sum up, through looking at the linguistic, affective, social/cultural, and cognitive dimensions of the language classroom, we have traced a number of ways in which songs might facilitate learning in an adult language classroom. In the following table, it is possible to see at a glance the different dimensions along with the ways that using songs might be beneficial. While it may appear on the surface that some of the entries in the table are synonymous with others, there are in fact differences, albeit subtle but relevant.

Table 1.1 – The Usefulness of Songs According to the Dimensions of the L2 Classroom

Dimension	Usefulness of songs		
Linguistic	integrate and exploit the different language skills	provide aural and oral linguistic repetition	emphasize the prosodic aspects of the language
	a means by which phonetic rules can be made tangible	show different aspects of the culture	show different levels of formality
	reflect everyday common features of connected speech	reflect grammar and vocabulary that ESL students face	bridge the gap between listening and communicating
	show sound-spelling correspondence of phonological and graphic elements	help with noticing the language and noticing gaps between L1 and L2	assist with achieving language awareness
Affective	raise self esteem	increase motivation	reduce anxiety and fear
	increase enjoyment of the language learning process	improve classroom experience	raise interest
	provide personally relevant meaning	engage learners	boost participation
	illustrate interesting / authentic registers	present interesting vocabulary	have themes that are interesting to students
	create a situation/pseudo-dialogue with the listener	give a sense belonging	provide affective attention (motherese)
Social / Cultural	bring students together	bridge the gap between students and teacher	enhance individual and social harmony
	create a feeling of belonging	assist in building solidarity	help students engage with each other
	are a way to explore the culture	contribute to a student-centred classroom	show cultural behaviour and attitudes
	make cultural ideas accessible	help with culture shock	provide direct participation in the culture
Cognitive	provide sub-audio input via SSIMH phenomenon	lead to automatic language output	lead to fluent output
	improve hearing of segmentals, connected speech, suprasegmentals	are conducive to internalizing correct pronunciation	help in noticing the language in order for intake to occur
	enhance auditory discrimination of syllables and words	provide memorable chunks of language	are chunked language with rhythm and rhyme to help learners remember
	benefit different learning styles	strengthen linguistic intelligence	strengthen musical intelligence
	engage more areas of the brain	invite rehearsal and transfer into long-term memory	incite singing i.e. practicing the structures

Problem Statement

In order to achieve the best pronunciation possible when learning an L2, one of the elements that adult students logically require is exposure to effective teaching techniques which employ materials that will facilitate the perception and production of the sounds and sound patterns of the L2.

However, given that the teaching of pronunciation has been marginalized over the years both in terms of research and practice, there is a gap in the knowledge as to which combination of techniques and materials might be most effective. Our preliminary review of the literature has determined that, although good pronunciation is important to language learners, not only are there many L2 teachers who are not aware of how to effectively teach pronunciation, but there is no agreed upon way of teaching pronunciation in CLT. That is, although ways of teaching pronunciation do exist, there are not many accounts of their application in real language classrooms. This is a situation which contributes to the problems that pronunciation is having in terms of being integrated more fully within the curriculum. Should this situation continue, it is unlikely that some of the difficulties that L2 learners face as a result of their accents will be resolved.

In response to this problem, our study proposes to investigate a way of teaching pronunciation that not only might be easy for teachers to employ but also effective for learners. We plan to use Gilbert's Prosody Pyramid (Gilbert, 2008) as the basic model for teaching English suprasegmentals along with the use of songs as a material to illustrate and provide practice of the elements contained in the pyramid in addition to a few common features of connected speech. Should this method be helpful to students, then teachers will have a viable

means of integrating pronunciation teaching into their classrooms. A detailed description of the Prosody Pyramid will be provided in “Chapter Two: A Review of the Literature”.

Purpose of the Study

The purpose of the study is to help determine whether songs can be a helpful material in the teaching of English suprasegmentals and to provide a descriptive account of the use of the techniques employed.

Utilizing Gilbert’s Prosody Pyramid as the basic model of the English prosodic system allows for a simple, step-by-step way to present individual suprasegmental features to students. Coupling the presentation of these features with songs and providing separate perception and production tasks for each feature constitute a novel way for the prosodic elements of English to be made clear and tangible for learners. Careful scaffolding and recycling of the different features in student-oriented activities allow for assimilation and practice of the pronunciation elements.

I decided to put this method to the test in a brief study of Spanish-speaking university students learning English in Chile. My research involved a two-week pronunciation course on suprasegmentals administered to two groups of students at the Pontificia Universidad Católica de Chile in Santiago in October 2012. The course was taught to one group of ten students using songs as material, and to a second group of five using texts other than songs. A control group of eight students who did not take the pronunciation course was included for comparison of results. In this way we were able to explore: (1) whether teaching suprasegmental pronunciation in this way aided the students in question and (2) the efficacy of the use of songs compared to traditional materials for this purpose.

The decision to have participants who were homogeneous in terms of L1, age, and exposure to the L2 environment allowed for a better understanding of which areas of the Prosody Pyramid and features of connected speech pose problems for the students and which areas did not. Having a clear idea of this information with regard to Spanish speakers will provide for future targeting of the problematic areas in terms of both teaching and research. We also hope that the lesson plans and materials provided in the appendices constitute a source of information for researchers and teachers in choosing texts and sequencing activities for teaching and practicing the different pronunciation features.

Teaching suprasegmentals with the use of songs and Gilbert's Prosody Pyramid is easy and because of this, it may be an attractive way for language instructors to address the pronunciation needs of their students. If teachers are aware of how they can use songs to teach pronunciation, doing so may benefit both students and teachers, namely because songs have the potential to permeate the various dimensions of the language classroom and may help in ways that go beyond the linguistic, i.e. affective, social / cultural, cognitive. In this way, this study also aims to add legitimacy to the use of songs in adult language teaching in general.

Research Questions

As outlined in the previous section, the central research question of our investigation is whether the use of songs in the ESL / EFL classroom can help improve the pronunciation of Chilean university students learning English. The research objectives that emerge from this problem are the following:

1. a) Can teaching suprasegmental phenomena in a two-week pronunciation course enable L2 learners to better perceive suprasegmental phenomena of the target language?

Classroom research on pronunciation perception appears to be scarce. Abe (2009) and (2011)⁹ conducted longer-term studies. With regard to short-term studies, Muller Levis and Levis (2012) in a study of contrastive focus (focus words) involving a group of 18 participants who had four hours of instruction which included listening, production and prediction exercises, did not note a significant improvement in focus word discrimination. However, in a study on teaching Mandarin, Wang, Spence, Jongman, and Sereno (1999) who gave eight classes over a period of two weeks with the goal of seeing whether the participants could distinguish four tones in Mandarin words noted a 21% increase from the pre-test in accurately identifying the different tones.

b) Are songs as a material beneficial for this process?

As we will see in Chapter Two: A Review of the Literature, few actual studies have been done on this area, although there have been numerous articles written on the subject. For example, Anton (1990) indicates that students learn the “rhythm, intonation, and pronunciation in a natural way” (p. 1169), while Van den Berg (2011) indicates that songs are useful for helping learners perceive suprasegmentals. With regard to studies, Schön et al. (2008) found that songs can help beginning learners perceive word boundaries. It is important to note that this literature treats suprasegmental phenomena as a whole; therefore, with respect to the components of suprasegmentals, it is unclear whether certain ones are more easily acquired than others.

2. a) Can teaching suprasegmental phenomena in a two-week pronunciation course enable L2 learners to better produce suprasegmental phenomena of the target language?

⁹ More details regarding these studies are provided in Chapter Two: A Review of the Literature.

While there have been longer-term pronunciation studies in which improvement was achieved i.e Abe (2009, 2011); Couper (2003); Derwing, Munro, and Wiebe (1997, 1998); Derwing and Rossiter (2003); Lord (2008); Sardegna (2011), and Sardegna and McGregor (2013)¹⁰, there have been fewer short-term ones. Couper (2006) conducted a two-week study on (vowel) epenthesis and (consonant) absence and found that long-term gains had been made. Sturm (2013), who devoted two days in a semester to liaison instruction in French, found that at the end of the term, participants made less forbidden liaison errors. Muller Levis and Levis (2012) (the same study mentioned above) found that after four hours of instruction which included listening, production and prediction exercises, a group of 18 participants significantly improved in their use of focus in reading sentences.

2. b) Are songs as a material beneficial for this process?

Likewise, there have been few studies that tried to determine whether songs can help with production. With regard to articles, Anton (1990), Borland (2012), Jolly (1975) and Lake (2002-2003) assert that songs are good for reinforcing pronunciation in general. As for studies, Chunxuan (2009) found that using songs resulted in learners who spoke more clearly and fluently. Rengifo (2009) also indicated that the use of songs improved students' pronunciation. As was the case with perception, with respect to production, suprasegmentals have generally been treated as a whole, so it is unknown which aspects of prosody songs can help with. One exception to this is Fischler (2009) who after a month-long course of 32 hours found that students who had studied English word and sentence stress with the help of rap music were more intelligible afterwards.

¹⁰ More details regarding these studies are provided in Chapter Two: A Review of the Literature.

Limitations of the Study

As with any investigation, this one had some limitations. First of all, two weeks is a very short time to improve a language skill. In order to learn to perceive and consistently and accurately produce both the segmental and suprasegmental aspects of a target language in addition to certain additional features of connected speech, a pronunciation course would need to have contained many more hours and been much longer in length. Because it was not feasible to offer a longer course with more teaching (and practice) hours, only select suprasegmental areas and connected speech aspects were taught and practiced. Therefore, the participants could not have been expected to improve in all of the pronunciation areas where improvement was necessary. As I mentioned earlier, modifying one's speech habits is not a task that can be completely accomplished over a period of two to three weeks; it is a long-term undertaking and one in which learners progress at varying speeds. Therefore, it was not possible to discover the extent to which there might be long-term beneficial effects for the students who took this course.

In addition, limitations of time and financial resources, and the fact that this was an individual research project inevitably imposed restrictions. Travelling to another country in order to have a relatively homogeneous group of participants required not only cooperation of the receiving institution but also a considerable expenditure of time and expense on the part of the researcher, working with no institutional funding. Despite these limitations, we hope to have established grounds for future longer term individual and collective research. By choosing to study (individual) receptive and productive prosodic aspects of a relatively homogeneous group

of participants, we hope to have indicated a certain direction and clarity for the future study of the benefits of different ways of teaching pronunciation.

Assumptions

In the study, it can be assumed that the participants answered all of the questionnaires honestly and did their best when taking the tests.

*Definitions*¹¹

L1. A speaker's native language.

L2. The target language. The language that is being taught or the one that non-native speakers are learning.

Language classroom. Refers to a classroom in which the target language (L2) is taught.

SSIMH Phenomenon. The song-stuck-in-my-head (SSIMH) phenomenon is the involuntary mental rehearsal of a song (Murphey, 1990).

Earworm. An earworm is the layman term for a song that is stuck in your head.

Segmentals. The individual phonemes and allophones that pertain to a language.

Suprasegmentals. Refer to phonetic elements beyond the level of the segmentals, i.e. rhythm, stress, and intonation.

Rhythm. "The pattern of long and short syllables" (Gilbert, 2005b, p. 89).

¹¹ Note that some of these terms will be discussed in further detail in Chapter Two.

Stress. Refers to “the pronunciation of a syllable or word with more respiratory energy or muscular force than other syllables or words in the same utterance” (Richards & Schmidt, 2002, p. 516). In Gilbert’s Prosody Pyramid, the term *stress* refers to *syllabic* stress, namely the syllable of the *focus word* which receives the most stress. The words that receive *sentence* stress, according to Richards and Schmidt (2002), are commonly those in an utterance (or sentence) that contain new information (p. 516). *These* words are referred to as *focus words*.

Focus word. It is the word in an utterance that receives the most stress. Gilbert (2005b) defines the focus word as “the most important word in a thought group. Focus words are emphasized with a pitch change and a long, clear vowel in the stressed syllable” (p. 88).

Word stress. “The pattern of stressed and unstressed syllables in a word” (Richards & Schmidt, 2002, p. 516).

Sentence stress. “The pattern of stressed and unstressed words in a sentence or utterance” (Richards & Schmidt, 2002, p. 516).

Thought group. A short sentence, clause, or phrase in which the words are chunked together and separated by brief pauses and often a drop in pitch at the end. Thought groups “are the organization of the speaker’s thoughts into groups” (Gilbert, 2008, p. 11).

Intonation. Refers to pitch patterns, that is, “the musical change in the voice, up or down, that helps the listener notice the important words” (Gilbert, 2005b, p. 88)

Melody. Pitch patterns, intonation.

Prosody. The collection of suprasegmental features, commonly referred to as the “musical” aspects of a language, that is, the combination of both rhythm and melody (Gilbert, 2008).

Connected speech. Refers to some of the segmental modifications that occur as a result of suprasegmental patterns, such as linking, assimilation, and elision.

Intelligibility. Refers to how much a person’s speech is understood by listeners (Munro, 2011).

Comprehensibility. Refers to how difficult it is for the listener to understand a person’s speech (Munro, 2011).

Accentedness. Munro (2011) explains this as “how different someone’s speech seems (often from the listener’s variety)” (p. 9).

Native-like. Indicates that a non-native speaker’s speech is close to that of a native speaker.

Fossilisation. “A process which sometimes occurs in which incorrect linguistic features become a permanent part of the way a person speaks or writes a language.” (Richards & Schmidt, 2002, p. 211)

Linguistic stereotyping. “A mechanism of social judgement whereby listeners ascribe a myriad of traits to speakers based often on only very thin samples of pronunciation” (Rubin, 2012, p. 11) Also referred to as *accent stereotyping* (Munro, 2003).

Accent discrimination. Discrimination against a person based on his or her accent.

Critical period. The critical period for language acquisition refers to “a period of time when learning a language is relatively easy and typically meets with a high degree of success. Once this period is over, at or before the onset of puberty, the average learner is less likely to achieve nativelike ability in the target language.” (Marinova-Todd, Marshall & Snow, 2012, p.9).

Critical Period Hypothesis for second language acquisition (CPH/L2A). The idea that a person’s susceptibility to language input is dependent on age, meaning that late-start L2 learners are unlikely to attain native-like speech.

Summary

This chapter began with a discussion of my background as a language learner and teacher and how my discovery of songs as a tool for learning and teaching a second or foreign language led me to want to study the use of songs for teaching pronunciation. Then pronunciation was defined and was followed by a brief discussion of how L2 pronunciation teaching and research has been neglected over the years. Following that, some reasons were given for the lack of pronunciation instruction in the CLT classroom and how this has resulted not only in students who do not have a clear understanding of their pronunciation shortcomings, but also in a challenge for them to find help. There is a brief discussion on some of the consequences of poor pronunciation along with mention of certain professional obstacles that can result from this. Then, the inherent musicality of language was discussed before introducing the dimensions of the language classroom in relation to songs. During the discussion of the first aspect, the linguistic aspect, an overview was provided as to how songs integrate and exploit the different language skills as well as how they assist with the process of noticing and language awareness. Next, the affective

aspect was examined with respect to how songs can have a positive influence on learners' emotions, confidence, and desire to learn while at the same time incorporating meaningful language, themes, and facets of the target culture that prevail on a daily basis. The social and cultural aspect relates to how songs are a reflection of the target culture and how they can help create positive social dynamics in the classroom. The last dimension, the cognitive, discusses how songs relate to different mental processes such as the Din, SSIMH phenomenon, automaticity, fluency, multiple intelligences, and memory. Following the table that sums up the usefulness of songs in relation to each dimension, there is the research problem statement which explains how this thesis aims to fill a gap in the area of L2 pronunciation teaching through using songs to teach English suprasegmentals. The contributions of this thesis to advancing knowledge are explained followed by the research questions that this study hopes to answer. After that, assumptions and limitations of the study are mentioned and the definitions of key terms are provided.

Chapter Two

Review of the Literature

Introduction

In this chapter, we revisit some of the areas mentioned in *Chapter One: Introduction*, namely pronunciation's subordinate status and some of the consequences of accented speech. We also review what different scholars have said about accent, pronunciation training principles, and educating the public as a way of helping non-native speakers in professional contexts. Then pronunciation teaching itself is explored through a look at what authors say about which pronunciation areas should be taught and how to go about teaching them. Articles and studies in the area of pronunciation teaching are reviewed, as are traditional pronunciation teaching materials. After that, the different methodological approaches to pronunciation teaching as well as teaching with songs are examined, followed by a review of the literature regarding songs in language teaching in general and in teaching pronunciation specifically.

Pronunciation's Subordinate Status

Several authors have pointed out the subordinate status of pronunciation as a language skill since the emergence of Communicative Language Teaching (CLT), which has been the major language teaching approach since the 1980s. Gilbert, in her 2010 article *Pronunciation as orphan: What can be done?*, sheds light on this subordinate status through examining literature

on the methodological history of language teaching, why teachers generally do not teach pronunciation, the kind of pronunciation teaching that is done, the actual amount of knowledge that is produced on the subject, as well as providing recommendations for improving the teaching of pronunciation. She suggests that the combination of these factors has resulted not only in a decrease in the ability of language teachers to help learners in this area but also a lack of awareness on the part of the learners as to their personal pronunciation weaknesses.

The use of CLT, whose goal is to enable students to communicate in the target language, tends to ignore the importance of phonetics and the essential role that pronunciation plays in effective communication. As Celce-Murcia et al. (2010) point out: “Communicatively adequate pronunciation is generally assumed to be a by-product of appropriate practice over a sufficient period of time” (p. 449). In addition to the assumption, what is considered to be appropriate practice and what is a sufficient period of time are problematic. Presumably appropriate practice would be that which has been shown to improve pronunciation and which has been accomplished on a regular basis over a minimum period of time. Furthermore, as mentioned in the Introduction of this thesis, Celce Murcia et al. indicate that in CLT “most proponents of this approach have not dealt adequately with the role of pronunciation in language teaching, nor have they developed an agreed-upon set of strategies for teaching pronunciation communicatively” (p. 9). In light of this, we might speculate that if adult learners are never taught how to say the notions, grammatical structures and vocabulary with the correct articulation of sounds, rhythm and intonation according to the communicative context and allowed to regularly practice for a certain period of time, it is unlikely that they will acquire this knowledge and skill on their own.

To many a teacher it will be obvious that adult learners do not simply acquire by themselves acceptable pronunciation in the target language. Morley (1991), referring to adult language learners, indicates that their “educational, occupational, and personal/social language needs, including reasonable intelligible pronunciation, [should] be served with instruction that will give them communicative empowerment – effective language use that will help them not just to survive, but to succeed” (p. 489). Despite Morley’s recommendations, this in general does not seem to be happening although there has been increased attention to pronunciation in the past decade or so.

Additionally, as mentioned earlier in this thesis, more research in teaching pronunciation is clearly needed. Even though Morley (1991) cited numerous pioneering studies in the 1980s and indicated the need for more, even now, Deng et al. (2009), while acknowledging that more attention has recently been paid to the area, says that “it is clear that pronunciation is still not a priority for most L2 researchers or teachers” (p. 11).

Regarding the state of the Canadian situation of pronunciation teaching as of 2010, Foote et al. (2011) present the most comprehensive explanation and coverage and compare these results with those of Breitzkreutz et al. (2001) who had conducted a survey of 67 teachers and/or program coordinators a decade earlier. The purpose of this earlier study was to determine the attitudes, beliefs and state of affairs regarding pronunciation teaching in Canada. What Foote et al. found was that although the earlier study had articulated the desires of ESL teachers, namely, more instruction in pronunciation training, more materials and more curriculum development related to pronunciation, by 2010, not a lot had changed. That is, while there has been an increase in available materials, still more teacher training is needed as is instruction in pronunciation for adult ESL learners. Regarding teacher training, 75% of the teachers surveyed

indicated a wish to have more training in pronunciation instruction, with only 58% stating that they were confident teaching segmentals and 56% suprasegmentals. As for the latter, of the teachers surveyed, an average of 6% (mode of 2%) of class time was spent on pronunciation instruction while 52% of the instructors made use of the pronunciation activities in their textbooks. With regard to reporting how many teachers in their programs teach pronunciation, that figure dropped from 73% in the Breitzkreutz et al. study to 46% in Foote et al. (2011). Finally, 43% percent of the teachers indicated that their institutions offered stand-alone pronunciation courses. These statistics appear to indicate that pronunciation is not considered to be as important as other language skills. Nevertheless, 85% of the teachers felt that pronunciation instruction could be effective and 92% stated that a stand-alone course would benefit some of their learners.

Thomson (2013) conducted a survey of 58 teachers (43 from Canada and 15 from the United States) regarding specific beliefs and techniques as expressed predominantly on websites devoted to accent reduction/modification as well as YouTube videos devoted to this. The study reported on 24 of a total of 131 statements regarding pronunciation beliefs and practices. What the author found was that many English Language teachers “seem not to have the background knowledge, and lack the confidence necessary, to critically assess questionable pronunciation beliefs and practices – beliefs and practices that they may encounter in the materials that they choose to use” (p. 225). Moreover, the author even found this to be the case with English language teachers who had taken a university credit course in teaching pronunciation and he indicated that “there was no evidence that they were better equipped to recognise the more questionable statements” (p. 229).

The above surveys indicate that teachers are in need of more training and support with regard to pronunciation teaching, not to mention a possible general need for an increase of awareness surrounding pronunciation and the teaching of it. Before reviewing the literature on some of the salient topics in pronunciation teaching, it is important to have a look at accent and distinguish between native and foreign accents. This distinction is necessary because it is directly related to some of the conflicting opinions related to pronunciation and people's judgements of accents.

Accent

Whenever we speak a language, we necessarily speak with an accent. The lexical items and grammatical structures along with the phonetic segments (allophones¹²) and the suprasegmental patterns that we use while speaking a language generally identify us as belonging to a particular group. We may either sound like we belong to a native group of speakers or a non-native or foreign group. If we sound like a native speaker of English, our accent will indicate where we might be from, in addition to our age, sex and socioeconomic level.¹³ In other words, our accent will systematically reflect the dialect, or linguistic variety, of English we speak. If we sound like a non-native speaker, that means that we are using some linguistic structures and phonetic patterns that are either not native to English or not used in the way that native English speakers would. Such patterns do not normally belong to any dialect and bear no relation to sociolinguistic indicators. They are mostly quite simply classified as mistakes. Both native and non-native accents may undergo value judgments by those who hear the accents spoken.

¹² An allophone refers to the actual manifestation of a phoneme, which is a purely theoretical construct.

¹³ See Labov (1972), for example, for more information on accent and sociolinguistic variables.

The status of an accent – and its acceptability – will depend on the social context in which the accent is situated, with political correctness and tolerance being merely variables of such a context. While one would expect that it is acceptable to have an accent in a country such as Canada, in the eyes of a person with an accent, this is not always the case, especially if it is one that is difficult for most people to understand. Take for example the words of Mayank Bhatt (2013), a Canadian immigrant, proficient in English, but possessing a strong Indian accent:

Accent is a strange thing. It gives variety to a language, identity to a person, adds spice to a conversation. Accent is diversity and multiculturalism personified. But make no mistake: it's not desirable. Those who have it would prefer to lose it. Those who don't can't understand the misery that those who have it experience. It's a barrier to communication. Even the most educated and articulate person will find it hard to communicate if he or she has a pronounced accent. (Canadian Immigrant, 2013, para. 4)

Non-Native Speakers' Opinions of Their Accents in L2 and L1 Environments

The above point of view about having an accent appears to be a common one, at least in Canada and Australia. There are studies that show that ESL learners want and need more pronunciation training and that pronunciation “is of great concern to many L2 learners in Canada” (Foote et al., 2011, p. 5). Derwing and Rossiter (2002) found that 55% of learners recognise that breakdowns in communication occur due to pronunciation problems and 90% of them say they would take pronunciation courses if they were offered. Derwing (2003) also discovered that 95% of the L2 immigrant speakers surveyed would speak with a native English accent if they could.

Zielinski (2012) interviewed 26 migrants in Australia about their personal perceptions of their English accents. The participants, who had taken ESL classes¹⁴ in Australia for a year, were “overwhelmingly negative about their pronunciation skills and most indicated that they felt their pronunciation affected their ability to be understood when they spoke English” (p. 18). There were 14 beginner and 12 intermediate level participants and all but one from each group made negative comments about their own pronunciation and both groups felt that they were “hard to understand” while the beginners mentioned that their pronunciation problems “affected their confidence to speak” (p. 21).

What we can conclude from the above sentiments expressed by non-native speakers in an L2 environment is that their perception of the acceptability of their accents is a result of their experience with native speakers. In a foreign language environment, where non-native speakers may have little or no contact with native speakers of the target language, their opinion of their own accent may not be linked to any experience with a native speaker.

In Chile, Véliz Campos (2011) performed a survey of 15 English language teachers-in-training in a five-year program in Santiago. The qualitative study sought to collect data which included information on accent, native speakers, and pronunciation materials. The pre-service teachers indicated that they had all been taught British English and three of them held the notion that it was more prestigious and formal than American English. With regard to how English teachers in Chile should speak, half of them said that British English or Received Pronunciation (RP) (used interchangeably) is the accent that English language teachers should use. The other half of the participants was divided; half of them said that one accent only should be used and the

¹⁴ The study does not provide details as to the kind of ESL classes the participants had received, that is, whether or not or to what extent they had training specifically in pronunciation.

other half indicated that many accents should be employed. Of the 15, four believed that teachers should speak with a native English accent; three said that teachers should aspire to native-likeness, and six felt it was more important to speak the language “in the best way possible” and master the teaching of it (p. 228). As far as speaking English with a Chilean accent, nearly all of the respondents judged people harshly for this, especially if they were teachers and thus models of the language. As well, it was indicated that “in most cases they see a Chilean accent as a factor that might undermine the person’s potential in the job market or simply when travelling” (p. 229). The respondents were unanimous in saying that it was important to be exposed to other English accents. With regard to the materials used in their pronunciation and phonetics classes, all of the materials feature British English and the author attributes part of their attachment to RP as a result of this¹⁵.

What is interesting about comparing non-native speakers’ opinions of their accents in L2 and L1 environments is that it serves to remind us of the different realities and pressures in different language environments. It is clear in both situations that learners are concerned about their accents and wish to speak in a way that is both acceptable and accepted. That Chileans themselves consider English pronunciation to be important within Chile may be due in part to the Chilean Ministry of Education’s educational reform initiatives¹⁶ to increase the level of English of its citizens with the hopes of eventually becoming a bilingual (Spanish-English) country.

¹⁵ That British English is considered the ideal model in Chile, is partially in line with Setter and Jenkins (2005) who indicate that British or American English are generally the models adopted around the world, with local varieties being looked down upon. However, although they say that American English tends to be the adopted model in South America, that is not the case in Chile.

¹⁶ The Programa Inglés Abre Puertas (English Opens Doors Program), launched in 2003 by the Chilean Ministry of Education, and which is supported by the Estrategia Nacional de Inglés 2014-2030 (National English Strategy 2014-2030) of the Government of Chile (Piñera, 2014).

Accent Addition Versus Accent Reduction

With regard to pronunciation training there are two terms, accent addition and accent reduction, which are used by different scholars to refer to the process of changing a person's accent. The terms have different connotations depending on who uses them.

Kjellin (1999) uses the term *accent addition* to refer to making a non-native speaker's speech sound closer to that of a native speaker. The author insists that the term accent addition constitutes a change of attitude, since he says, in reference to accent reduction, "what somebody has cannot be taken away" (p. 2). He contends that accent addition "will serve to play down the process, which may make some people apprehensive" and he adds that "fossilization prevention and remediation is rather achieved by the addition of an ability to play-act, as it were, in a native-like manner whenever the speaker wishes to speak in that way" (p. 2). Gilbert (2005b) explains the difference in the two terms in another way. She says that "instead of trying to remove mispronunciations...it is more helpful for teachers to concentrate on adding new elements required by the target language" (p. xiii) and therefore prefers the term accent addition. Munro (2011) exposes the negative connotation associated with the term accent reduction that Kjellin hints at when he states, "Regrettably, the 'accent reduction' industry often exploits immigrants' insecurities about their accents as a way of marketing their dubious and expensive services" (p. 12). Jenkins (2004) mentions that the "prevailing concept of 'accent reduction' [has a] tendency to regard learners as subjects for speech pathology and to exhort them to lose all traces of their L1 accent in their L2" (p. 115). As we can see, these four authors prefer to use the term accent addition and agree in their negative assessment of the term accent reduction¹⁷. However,

¹⁷ Note that Kjellin is the only one who adheres to the nativeness principle, whereas Gilbert, Munro, and Jenkins are proponents of the intelligibility principle. These principles are discussed below.

Schmidt and Sullivan's (2003) survey of American graduate programs in speech language pathology does not indicate that the goals of foreign accent modification (FAM)¹⁸ are to eradicate all traces of a foreign accent. In fact, these authors discuss barriers to communication and intelligibility, the future of FAM, what the goal should be for the clients, as well as the need for FAM clinicians to consult TESL journals in order to further their knowledge of pronunciation instruction. In light of this, one might question whether the prevailing connotation of accent reduction is an accurate one. Furthermore, with regard to accent addition, if we consider the contrast between changing one's L1 pronunciation and one's L2 pronunciation, the literal notion of the term seems to be called into question.

Even though the above authors are referring to non-native speakers and the process they undergo in L2 pronunciation training, it is important to remember that native speakers themselves may also have the need to undergo accent training for professional reasons, especially if their native accent is one that is not accepted as appropriate where they happen to be living. Native English accents that could be considered unacceptable and barriers to professional employment in Canada might be a Jamaican or African American Vernacular, for example. In the performing arts, actors, who are native speakers of the accent where they live, might need to learn to speak another accent for a part in a movie, or a play. Cases such as these could also be seen to be ones of accent addition, because the people would be adding a new accent to their repertoire. In this way they would speak two different accents of one language and would be able to switch between the two according to the sociolinguistic and professional context. Therefore, both accent addition and accent reduction can lead to the non-use of a particular accent. What distinguishes the two situations, i.e. L2 and L1 pronunciation learning, however,

¹⁸ Accent modification and accent reduction are often used interchangeably. The term *elocution* seems to be used less often than *accent addition*, *accent modification*, *accent reduction* when referring to changing the way one speaks.

is that while native speakers might easily alternate between their two accents, one must question whether non-native speakers actually have two L2 accents between which they can (or ever would) alternate between. This now leads us to consider pronunciation training for non-native speakers in which there are two different ideologies: the nativeness principle and the intelligibility principle.

The Nativeness Principle and the Intelligibility Principle

The *nativeness principle* “holds that it is both possible and desirable to achieve native-like pronunciation in a foreign language” (Levis, 2005, p. 370). While this may be attainable for L2 speakers who begin their learning before the age of six, it is generally accepted to be unlikely for older learners, especially those who begin learning during what is referred to as the *sensitive period* which is between the ages of six and 12 (Long, 1990). There are also studies that contend that native-like speech is not possible for late learners (i.e. after puberty) (e.g. Abrahamsson & Hyltenstam, 2009).

The *intelligibility principle* “holds that learners simply need to be understandable” (Levis, 2005, p. 370). Moreover, this principle asserts that the degree of an accent does not always affect intelligibility, that is, that intelligibility may be sustained despite a heavy accent (Munro & Derwing, 1999). In addition, it maintains that pronunciation instruction “should focus on those features that are most helpful for understanding”, such as suprasegmentals (Levis, 2005, pp. 370-371). On the other hand, Jenkins (2000) proposes the teaching of a Lingua Franca Core (LFC) for communication between non-native speakers of English who have different L1s. The LFC is considered to contain those phonological features that are most important for intelligibility between non-native speakers with different L1s. The term *intelligibility* is one that

is currently very salient in the literature on second language teaching and acquisition and, as Munro (2011) indicates, has been around at least as far back as Sweet (1900). Munro defines intelligibility as “how much of the speech is understood by interlocutors” (p. 9). In pronunciation teaching today, it is advocated that intelligibility and comprehensibility¹⁹ be the goals that second language learners should have as opposed to attaining a native-like status (Derwing, 2010). While in most cases, it might be clear that the intelligibility principle is the most reasonable one for language teachers and learners to adhere to, some of the viewpoints raised by different authors challenge this notion.

Rajagopalan (2010) refers to intelligibility as “an evaluator adjective...[which] automatically invoke[s] the figure of an evaluator” (p. 468). To make his point he goes on to say that “no variety is intelligible or otherwise in and of itself” (p. 469). While Rajagopalan’s article is controversial (see Munro 2011), he is not alone in pointing out the inherently problematic subjective nature of the term. According to Rubin (2012) “any assessment of a speaker’s speech performance could very well reflect nearly as much about the listener as about the speaker” (p. 11). That is, if a listener is prejudiced against the group to which a speaker belongs, then there could very likely be linguistic discrimination when that person talks, resulting in a loss of intelligibility. As well, even if a listener is not biased, if he or she lacks experience or exposure to a certain accent, the speaker might be considered less intelligible than would be the case with an interlocutor who is familiar with the accent. That is to say that “listeners with certain experience, background, and perhaps aptitude may be more successful than others at comprehending L2 speech (Munro, 2011, p. 11). In this way, if an L2 speaker’s intelligibility can

¹⁹ *Comprehensibility* refers to how easy or difficult a person’s pronunciation is to understand (Derwing, 2010).

be so variable, one might question whether a goal of intelligibility is sufficient for all L2 learners.

The concept of nativelikeness, however, is also problematic due to the subjectivity of the interlocutor's assessment of the speech based on his or her knowledge of linguistics and native varieties. Ajioka (2010) found that grammar and pronunciation, upon which academic judgements of nativelikeness are traditionally based, are less important for laypeople who considered factors related to interaction to be more important. Therefore, speech that may seem native-like to a layperson could quite possibly not stand up to the rigorous criteria imposed by linguists.

Ajioka's study raises the question as to what nativelikeness entails and what is required in order to achieve native-like speech. Birdsong (2005) warns against having criteria for nativelikeness that are too strict so as to protect the Critical Period Hypothesis for second language acquisition (CPH/L2A).²⁰ To further compound the issue, Birdsong questions whether all instances of non-nativelikeness are because of a late start and could perhaps instead be due to the influence of bilingualism. That is to say, he wonders whether errors that are considered to be a result of faulty (i.e. late) learning may actually be due to linguistic interference. Marinova-Todd, Marshall and Snow (2012) contend that "age does influence language learning, but primarily because it is associated with social, psychological, educational, and other factors that can affect L2 proficiency, not because of any critical period that limits the possibility of language learning by adults." (p. 28). Furthermore, while acknowledging that most adult L2 learners do not achieve native-like proficiency, they assert that most of them "fail to engage in the task with

²⁰ This is important to consider, because if the CPH/L2A is highly questionable, then opponents to the nativeness principle might need to consider what factors other than age account for less than native-like L2 performance.

sufficient motivation, commitment of time or energy, and support from the environments in which they find themselves to expect high levels of success” (p. 27). Despite the age factor, there are studies that have found that native-like pronunciation is possible for late learners (Bongaerts, van Summeren, Planken, Schils, 1997; Bongaerts, Mennen and van der Slik, 2000). The authors suggest that a combination of plenty of exposure to “authentic” target language, high motivation, and intensive training in the perception and production of the sounds of the L2 can compensate for a late start and result in a native-like pronunciation (Bongaerts et al., 2000, p. 306). In addition, they ask whether typological proximity of the native and L2 languages might also work in the favour of late learners, meaning L2 learners whose L1 is similar, as in the case of English and Dutch, for example. In another study, Ioup, Boustagui, El Tigi, and Moselle (1994), found that formal instruction was not a factor in the attainment of native-like speech. (Nor could typological similarity have been a factor because the person in their study was a native speaker of English who learned Arabic.) In addition to motivation and plenty of L2 use and input, the authors attributed the success of the person in their case study to another factor: talent in learning languages, which is hypothesized to be associated with “unusual brain organization” which permits the individual to be more “cognitively flexible in processing L2 input and ultimately organizing it into a system” (p. 92).

Finally, fossilisation is another concept presented in the literature on pronunciation and related to the principles of intelligibility and nativelikeness. There are some studies that indicate that fossilisation can be overcome (e.g. Couper, 2006; Derwing, Munro, & Wiebe, 1997) although there still seems to be the idea that pronunciation fossilisation is a permanent affliction which can never be cured. Kjellin (1999) addresses this issue and believes that “fossilization in second learners...is more due to insufficient instruction and training at the beginner’s level than

to any biological constraints, and thus is preventable” (p. 2). He even goes on to say that “traditional teachers and learners of a second language usually aim too low...probably because many people believe in the myth [that native-like L2 pronunciation is impossible]” (p. 4). Whether the achievement of native-like pronunciation by most late-start learners is a myth, we still do not know. Furthermore, while there are studies that indicate that fossilisation can be overcome, the extent to which it may be overcome is still unknown. However, as we can see from the above discussion of nativelikeness, achieving native-like speech may not necessarily be the impossible dream that adherents to the intelligibility principle suggest. This is not to say that learners should be given false hope but limiting L2 speakers to a goal of intelligible speech might not be regarded as sufficient by some. Véliz Campos (2011) contends that “a distinction must be made based upon the language use(s) that the learner pursues, i.e. the degrees of linguistic depth greatly vary between a businessman language learner and a prospective teacher of English as a language learner” (p. 231).

Consequences of Accented Speech

As we have seen, many L2 learners are concerned about their accents. Munro, Derwing and Sato (2006) indicate that one of the consequences for those with accented speech is negative social evaluation. They point out that people may be evaluated negatively on the one hand because of their accent and the misunderstandings and miscommunications that occur, and on the other because of “the stereotypes or prejudices that accent can evoke in a listener” (p. 71). By considering personal accounts of the manifestations of negative judgements of non-native speakers, it is possible to understand that speaking with an accent brings with it consequences that extend beyond simply being difficult to understand.

Accent discrimination. Derwing (2003) quotes typical negative comments from a group of ESL learners, 30% of whom felt they were discriminated against because of the way Canadians reacted to them in different situations. The comments the ESL students reported were grouped into the following themes: “a lack of attention, rudeness, anger, and deliberate misunderstanding” (p. 557). Furthermore, the situations in which this type of treatment occurred often involved “educated”, “professional” Canadians in the workplace. In addition, Derwing indicated that almost a third of the ESL learners said that they had been discriminated against because of their accent while 53% felt that “Canadians would respect them more” if they spoke better (p. 555). What we can suppose from this information provided by Derwing, is that since these people are immigrants, the discrimination they are subjected to is quite possibly ongoing and is permeating most areas of their lives. Rubin (2012) refers to this kind of discrimination as *linguistic stereotyping*. According to the author, the Linguistic Stereotype Hypothesis stresses that “speech style is a powerful emblem of social identity” (Rubin, 2012, p. 12). Furthermore, he indicates that it is natural for listeners to assign or attach a social identity to speakers and then judge them according to the listener’s stereotypes of the speaker’s social group.

Munro (2003), in his discussion of accent discrimination in Canada, refers to this phenomenon as “accent stereotyping”, which is one of three types of accent discrimination that appear in human rights cases, the other two being “discrimination in employment due to inappropriate concern with accent” and “harassment based on accent” (p. 38). Foote et al. (2011) found that 70% of the teachers surveyed considered a heavy accent as a cause of discrimination against ESL learners. Furthermore, according to Derwing (2003), this negative reaction is more likely to occur the more different the accent is, that is, the “further the accent is from their own”

(p. 548). What is more, Derwing cites documented examples by authors who show that native speakers of a language who do not speak the standard accent will also suffer negative judgements. In fact, it is important to remember that sociolinguistics teaches us how discrimination and stigma are part of most language situations in real contexts, and no amount of political correctness is going to exempt foreign speakers from that.

Hahn (2004), in a study involving three groups of first-semester freshman students, found that they evaluated more positively an international teaching assistant who had delivered a lecture with correct sentence stress, than with no or incorrect sentence stress. That is, although the only difference in the three lectures was the placement of sentence stress, the listeners evaluated more positively the lecture as well as the speaker's ability to communicate²¹ when the sentence stress was correct. Even more disconcerting is a study by Lev-Ari and Keysar (2010), who found that people with foreign accents are considered to be less credible, and that the stronger their accent, the less credible they are seen to be. When we read this, one might think that racism is the root of the issue. What compounds the issue, however, is that Lev-Ari and Keysar said that non-native speakers with accents were seen as less credible "even when prejudice against foreigners could not play a role" (p. 1095). They say this because the non-native speakers were just relaying information from native speakers, and the listening (judging) participants were made aware of that. The researchers believe that the listeners "misattribute the difficulty of understanding the speech to the truthfulness of the statement" (p. 1094). What this means is that whether or not there is actual racism involved, speakers with a strong non-native accent²² will be discriminated against simply because they are hard to understand²³.

²¹ See Hahn (2004), page 213 for the list of questions that the listeners were asked.

²² It should be noted that the researchers did not refer to any relative intelligibility ratings of the heavy accents.

Consequences of this could be that speakers with a strong accent will be denied jobs, overlooked for promotions, and forced to accept lower-paying service positions – even though they may be highly skilled and educated²⁴.

This phenomenon can be seen with Canadian immigrants who have less than acceptable accents but are otherwise professionally and technically qualified for positions. In *The Globe and Mail*, Immen's article, "Maybe your English isn't as good as you think", reported that the Chinese Professionals Association of Canada "estimates that just 10 per cent of immigrating Chinese professionals land a job in their field of expertise in their first year in Canada" and that "a large proportion of those...have worked on improving their language skills to make themselves better understood...[and that] those whose accent holds them back find they often have to accept a less-well-paying job just to survive" (2006, para. 11-13).

Lippi-Green (2012) in *English with an accent: Language ideology and discrimination in the United States* provides an account of the situation for foreigners residing in that country. With regard to the workplace, she admits that there is a lack of proper studies but does mention a survey by the General Accounting Office of the United States Government in which it was found that of the 461,000 companies surveyed, 10% openly admitted to discriminating against foreigners based on their appearance or accent. Most certainly, however, it is argued that this percentage is higher, as it is likely that many employers did not openly admit to discriminating. Moreover, in hiring audits designed to detect this kind of discrimination, the situation was found to be "prevalent" (p. 153). According to the author, possible reasons for there not being more

²³ It should be noted that the authors do not distinguish between *comprehensibility* (how difficult it is for the listener to understand a person's speech) and *accentedness* (how different someone's speech seems from the local variety).

²⁴ Obviously there are certain kinds of jobs in which the ability to speak very intelligibly would be fundamental to the job and included in the job description. Professions in which this would presumably be the case are air traffic controllers, pilots, and police officers, to name a few, and in situations such as these, partiality based on the person's pronunciation would be both legal and acceptable.

documented cases or for the discrimination to go undetected include: sophisticated and subtle ways of discriminating on part of the employers, a lack of reaction on the part of the foreigners because they are used to being poorly treated, a lack of knowledge of legal recourse and how to pursue it, and complaints that are resolved through mediation or an outside party.

Reverse Linguistic Stereotyping. If accent discrimination were the sole type of bias that foreigners experience, then the situation would be less complicated. Unfortunately, however, there are studies that suggest that even if someone has “good” pronunciation, if his or her appearance is that of a visible minority, then discrimination (supposedly) based on accent can still occur. Reverse Linguistic Stereotyping (RLS) is a phenomenon in which “listeners attribute a speech style to a speaker based not on what they hear, but on what they believe is the speaker’s social identity” (Rubin, 2012, p. 12). Rubin shows that consequences of RLS for the listeners are that not only do they anticipate that they will have difficulty understanding a non-native speaker, but that they actually do understand less. Studies on RLS have taken place in university, business, and health care settings (Rubin, 2012) and although RLS is not always shown to be existent (see Lima, 2012), it is definitely a concern because, according to Rubin, “no amount of speech training or therapy will erase the effects of RLS” (pp. 14-15). Perhaps the author is correct in asserting this; however, it does not mean to say that speech training or pronunciation instruction will not lessen the severity of RLS. Indeed, he goes on to say that pronunciation research and teaching has “the ultimate goal of mitigating (if not erasing) negative prejudices that arise simply because certain speakers’ talk mark them as the ‘other’” (p. 15). As such, pronunciation and speech training should be a viable option for those non-native speakers with an accent who wish to modify it.

Educating the public. Effecting widespread public education with respect to foreign accents would certainly be ideal but it is most likely a long-term goal. Nevertheless, in the short term, it could be possible to help some non-native speakers who work in professional settings. For example, Derwing (2010) suggests following the Danish example of putting up posters in the workplace which contain ten suggestions on how to interact with non-native speakers. Essentially the suggestions reflect common courtesy and commonsense, such as, “Look at the person you’re speaking to” and “Don’t mumble” (p. 31). However, such simple etiquette can easily be forgotten during the course of a busy workday. Furthermore, there are experiments that have been done in select professional settings designed to train people to listen more carefully to non-native speakers (e.g. Derwing, Rossiter, & Munro, 2002). As we can see, targeted portions of society could be trained on how to deal with accented speech. While such education cannot be expected to penetrate the society widely or deeply enough to benefit all of the foreigners, it does provide a possible solution for those accented speakers who are in certain types of professions who otherwise do not make use of or do not have access to pronunciation training.

Pronunciation Teaching

Segmental Versus Suprasegmental

The focus on which aspects of pronunciation should be taught has changed over the years. Derwing and Rossiter (2002) mentioned that “for two decades” the focus on contrastive segmental teaching involving minimal pair drilling shifted to “a more global approach” which

emphasized those areas that most affect comprehensibility, namely suprasegmental (p. 156). This shift can be noticed in the pronunciation teaching materials that are on the market. Experts, such as Gilbert (2005a, 2005b, 2008), Grant (2007, 2010), Celce-Murcia et al. (2010) lean toward this more suprasegmental approach while also including certain segmentals. Nevertheless, there are still pronunciation textbooks that focus primarily on consonants and vowels, such as, Bareither (2007), Dale and Poms (2005) and Lane (2013a, 2013b, 2013c), although Lane does present a somewhat more balanced approach. Worth mentioning, is an early text by Adamowski, *The ESL Songbook*²⁵ (1997) in which the pronunciation areas are purely suprasegmental.

Derwing and Rossiter (2002) indicate that the reason for this shift is due to studies that suggest that native speakers' ability to understand non-native speech is very much tied to their prosody. However, both Derwing and Munro (2005) and Levis (2005) indicate that despite the belief that a speaker's prosody affects their intelligibility, there are (still) few studies that actually support this. One that does is Hahn (2004), which (as mentioned above) indicates that stress made a difference in the intelligibility of non-native speech. In the study, three groups of undergraduate students listened to one of three versions of an international teaching assistant's speech: one with correct sentence stress, one with incorrect sentence stress, and one with no sentence stress. The author found that "with correct primary stress, the participants recalled significantly more content...than when primary stress was aberrant or missing" (p. 201). A study which supports the importance of word stress was carried out by Field (2005), who compiled a list of two-syllable words, in some of which the syllabic stress had been shifted to either the left or right. After presenting these words to a group of native speakers (students at a British high

²⁵ This book will be discussed in more detail later in the chapter.

school) and a group of non-native speakers (students at private British language schools), Field found that intelligibility was impaired (most) when stress was shifted to the right and less when it was shifted leftward, especially if this left shift was accompanied by a change in vowel quality.

The Prosody Pyramid

A way of teaching English suprasegmentals is presented by Gilbert (2005a, 2005b, 2008). Gilbert (2005a, 2005b) are the student text and teacher resource guide for classroom instruction while Gilbert (2008), *Teaching Pronunciation Using the Prosody Pyramid*, explains the approach. Figure 2.1 provides a visual image of the areas in the pyramid.

Figure 2.1. The Prosody Pyramid²⁶



Thought groups are language segments which reflect how speakers organize their thoughts, and they consist of phrases, clauses or short sentences which are separated at the beginning and end by pauses and/or changes in intonation. For example, the English utterance “My CAT / is sleeping on top of the BOOKSHELF”, has two thought groups which are shown to be separated by a forward slash. Each thought group contains what is called a *focus word*. A

²⁶ Gilbert (2008)

focus word, indicated here through the use of capital letters, is the one word in the thought group that is most important for the listener to hear²⁷, and so it is pronounced in such a way that it will stand out from the other words in the thought group. This is achieved mostly through a change in pitch, a lengthening of the stressed syllable, and/or an increase in volume. Syllabic stress²⁸ refers to the different types of stress in English words. That is, words with two or more syllables will have one syllable that is more stressed than the others and therefore is said to be *stressed*. To show that a syllable is stressed, the syllable may be pronounced louder, with a pitch change, and with a longer clearer vowel sound. The stressed syllable in the focus word “represents the peak of information in the thought group” and so “must be heard clearly” (Gilbert, 2008, p. 14). In the example above, the stressed syllable of the second focus word is the one that has been underlined and is the *peak* in the second thought group. The other syllables in a multisyllabic word will either be unstressed or reduced. When *reduction* occurs, it is often through use of the schwa, although contractions or the elision of sounds and syllables are some of the other means of achieving reduction. The *schwa* is, interestingly enough, the most common vowel sound and yet is the most difficult one to actually hear, because it is necessarily very short in duration and obscure in quality. What it does is provide the contrast that is needed for the stressed vowels to be highlighted (Gilbert, 2008, p. 17). An example of the schwa sound is the second vowel sound in the word “nation” or the first one in “alone”. Together, thought groups, focus words, stress, and peak are the elements that form Gilbert’s Prosody Pyramid. Illustrating and presenting the English prosody system in this way, should help make it comprehensible and accessible for language learners because it breaks down suprasegmental phenomena into easily understandable units.

²⁷ Sometimes referred to as *sentence stress*.

²⁸ Sometimes referred to as *word stress*.

Before reviewing more of the literature related to pronunciation, namely articles on classroom teaching, studies, and materials used for pronunciation teaching, we should bear in mind the observation of Derwing and Munro (2005) that more research is needed along with better communication of the complexity involved in pronunciation studies, research and their application. Among a number of recommendations, the authors discuss the importance of carrying out research that is relevant for the classroom and which will further knowledge in the field. With this in mind, some articles presented in the following discussion of the literature are included even though they are not based on studies but rather solely on classroom experience, since they are still helpful in illustrating some of the existing guiding principles of pronunciation teaching. The guiding principles that are indicated will be presented together in the form of a table on page 91.

Articles

Cheng (1998) outlines the steps involved in teaching pronunciation in an English phonetics course for Chinese speaking student teachers studying to be English language teachers at a teachers' college in China. The author states that the faculty members of the college value the course and that their experience indicates that "after a year of systematic study of English phonetics, the student teachers have made great progress in their pronunciation and intonation" (p. 37). He also mentions that English language teachers who studied at this college are more successful than those with little training in phonetics and produce students who perform better in the annual English Speech Contest in Binzhou (p. 39). This course, although referred to as an English phonetics course, is essentially an English pronunciation course, as we will see from the following description. That is, while the two sounds systems are initially compared and the areas

that are most likely to be problematic are focussed on when the teachers provide instruction in articulatory phonetics, much of the course involves listening and speaking activities using dialogues, songs²⁹, tongue twisters, games as well as articles from their textbooks. The author follows the premise that “production would be impossible without perception” and so considers perception to be more important because it is necessary for the acquisition of the various segmental and suprasegmental features of speech (p. 37). Therefore, in the actual course, perception training precedes production training. In addition, Cheng mentions the need for teaching in a meaningful and motivating way by ensuring that students have a realistic and useful context within which to practice words, dialogues and conversations. Finally, the author asserts that periodic assessment is important so that students are aware of their progress. In order to do this he uses recordings of the students’ speech, role plays, discussions, communication games, story-telling activities, and speech contests. (Interestingly enough, Pardede (2010) outlines the exact same steps to pronunciation teaching as does Cheng – often word-for-word!)

Isaacs (2009) discusses the difficulty in integrating targeted and sufficient practice of specific pronunciation features with meaning in L2 pronunciation teaching. She talks about how narrowing the gap between a focus on pronunciation and Communicative Language Teaching’s focus on communication has been elusive. The author provides a semi-communicative CLT teaching example of how the gap can be narrowed through employing Firth’s (1992) zoom principle which advocates a constant shift in focus to and from the target pronunciation forms and the communicative activities in which these forms appear. In this way repetitive practice of the forms is achieved through the use of realia (in this case an advertisement) in which specific

²⁹ A discussion of how Cheng uses songs in his pronunciation course is included later in this chapter under the subheading *Songs and Their General Usefulness in L2 Teaching*.

pronunciation areas are focussed on along with various opportunities for the students to practice reciting the ad.

Tan and Woodworth (2011), indicate that “the development of awareness of sounds and rhythms in English, coupled with the strategies to identify and auto-correct errors will create more productive learners” (p. 3). This “productivity” presumably means that these learners will speak more intelligibly and with more confidence, and will not hold back or be reticent to speak due to inhibitions regarding their L2 speaking ability or a weakness in listening comprehension.

Darcy et al. (2012) propose a set of six principles for designing a curriculum for pronunciation for communication classes for an intensive English program at Indiana University that spans the levels from true beginner to low advanced. These principles include the need to lean on research, include both listening and speaking, embed pronunciation within the curriculum as well as in each lesson, begin at lower levels, adapt the teaching throughout the levels, and consider the teachers’ development as the curriculum is incorporated.

Studies on Pronunciation Teaching

As indicated in the previous chapter (Introduction), studies on pronunciation teaching are not as prevalent as those in other skill areas. What follows are examples of the types of pronunciation classroom studies relevant to ours that have been done in different pronunciation areas. While our research is concerned with the teaching of English pronunciation, a few studies on other languages are also included.

In an early classroom study of pronunciation teaching by Macdonald, Yule, and Powers (1994), the pronunciation of words (taken from mini lectures that the participants gave) were judged by 23 native English speakers. The participants, 23 L1 Mandarin speakers, were

identified in TOEFL scores as being at a high-intermediate to low-advanced proficiency level, but with pronunciation problems. They were randomly assigned to four different groups, with six in each of the three treatment groups and five in the control group. The treatment groups were divided according to types of activities that were typical teaching practices at the time: drilling (10 minutes), self study with tape-recordings (30 minutes), and interactive activities (10 minutes). There was a pre-, post-, and delayed post-test. Results of the immediate post-test indicated that only the self study group had significantly improved but this was not maintained in the delayed post-test. Had the intervention periods been longer, it is possible that the results might have been different, given that 10-30 minutes is a very short period of time to learn anything as complex as pronunciation.

In another study of the pronunciation of words, Miller (2012) studied the sound (segmental) discrimination skills of two groups of beginner L2 French learners. The 11 participants in section one and 12 in section two received four 15-minute pronunciation lessons in minimal pair sound discrimination. Section one had two lessons with the phonetic approach, which involved explicit (phonetics) teaching, for the first two lessons and then the reference approach, which involves repetition of words, for the last two. Section two, on the other hand, received the reference approach first for two lessons, and then the phonetic approach for the last two. The quantitative results indicated that sound discrimination significantly improved from the initial pre-test to the final post-test when the phonetic approach preceded the reference approach. Qualitative results showed that 57% of the participants, irrespective of the technique used, found the pronunciation lessons to be helpful. Since no delayed post-tests were administered, there is no way to know whether the gains in pronunciation were maintained.

Another investigation involving explicit pronunciation teaching is that of Couper (2003). His teaching was based on seven principles, which came from feedback given by previous students. These principles include: recognising that pronunciation is important; developing an awareness of their own pronunciation in contrast with native-speaker pronunciation; becoming skillful at monitoring their pronunciation, and taking steps and applying techniques that not only strengthen their auditory memory and articulatory skills (motor skills) but also their self confidence through support and positive reinforcement. In his action research study in which 15 post-intermediate adult learners of varying ages from six different language backgrounds were provided with two hours a week of pronunciation instruction for 16 weeks, Couper noticed an improvement in the pronunciation of the participants. There were pre- and post-course speaking tests which involved reading sentences and speaking freely. In these tests, the errors that participants made in the articulation of phonemes, word stress, weak and strong forms, epenthesis and absence, joining sounds, and sentence stress were noted and added up and averaged, and then the pre- and post-course averages were compared. There were no listening tests. Knowing how the participants' perception skills might have changed over the course of the study could have provided some insight into the effectiveness of the teaching, as well as what kind of connection might exist between the participants' pronunciation strengths and weaknesses in terms of perception and production. Nevertheless, surveys were conducted which asked the students about their views regarding the course, and teaching and learning pronunciation. The majority of the students liked the teaching approach.

In 2006, Couper took his research further to determine the effectiveness of specific pronunciation techniques on (vowel) epenthesis and (consonant) absence.³⁰ The author taught a

³⁰ Epenthesis refers to the addition of a consonant or vowel sound when there should not be one. Absence refers to missing sounds.

treatment group of 21 participants (mostly Chinese speakers and with an average age of 32.5) for six hours over a period of two weeks. The classes were 30 minutes long and were interspersed between general language learning classes. Couper was able to determine that the pronunciation instruction had a positive effect on the error rate of the students by giving pre- and (two) post-course listening and speaking tests which indicated that long-term gains in pronunciation had been made. The error rate dropped from 19.9% in the pre-test to 5.5% in the immediate post test and then rose to 7.5% in the delayed post-test which was given 12 weeks after the end of the pronunciation sessions. What is notable about the results is that although epenthesis and absence constitute a pronunciation area that, according to Couper (2006), “often becomes fossilized” (p. 60), long-term gains were made. Regarding perception, the researcher was unable to draw any conclusions; however, he did mention that it may be “more difficult to change perception than it is to change production” (p. 57).

With regards to studies that concentrate only on segmentals, there have been a few. Saito (2011) conducted a study of 10 L1 Japanese speakers who were taught eight English segmentals with explicit phonetic instruction over the course of four weeks (one hour per week). Saito found that in the post-tests, in which native English speakers evaluated accentedness and comprehensibility, compared to the control group (n=10), the participants in the treatment group made a significant improvement in comprehensibility when reading sentences; however there was no significant improvement in accentedness. When describing picture stories, there was no improvement in comprehensibility nor accentedness..

Later Saito and Lyster (2012) carried out a short-term segmental study of the effect of form-focussed instruction (FFI) on the pronunciation of the English /ɹ/ by Japanese speakers, The researchers gave a four-hour course over a period of two weeks that included instruction not

only on the /ɹ/ but also on argumentative skills. There were 65 participants in total: 29 in the FFI-corrective feedback group, 25 in the FFI-only group, and 11 in the control group (who had instruction but not on the /ɹ/). After the post-test, acoustic analysis of the third formant values of the /ɹ/ decreased significantly (meaning that the pronunciation of /ɹ/ significantly improved, i.e. sounded less like /l/) for the FFI group who had received corrective feedback in both controlled (word and sentence reading) and spontaneous speech (picture descriptions).

In a study of Spanish pronunciation, Lord (2005) studied the progress of 17 L1 English speakers' pronunciation of Spanish phonemes, /p/, /t/, /k/, /r/, /β/, /ð/, /ʎ/ and diphthongs beginning with /i/ or /u/. At the end of a semester in which the participants had received instruction in phonetics, practice with voice analysis software, and self analysis, the author noted significant results in paragraph readings. Later, Lord (2008), in a study of 16 L1 English speakers learning Spanish used podcasts in which the participants recorded themselves in controlled speech (reading a text and tongue twisters) and free speech over the course of a semester. Without going into detail regarding the specific pronunciation areas, it was noted that, at the end of the term, overall, progress had been made in segmentals and suprasegmentals in controlled speaking, although there were some participants who were rated the same or lower than at the beginning of the term. The judges rating the participants were three graduate students: an L1 Spanish speaker and two L1 English speakers who were at an advanced level of proficiency in Spanish.

An earlier study which also included segmental and suprasegmental pronunciation areas is that of Derwing, Munro, and Wiebe (1997) who observed a 12-week pronunciation course totalling 72 hours. There were 13 adult participants of varying L1s and proficiency levels who

came from various language backgrounds and had been in Canada for an average of 10 years. In this study, NS listeners evaluated the intelligibility, comprehensibility and accentedness of the NNS based on the readings of true and false sentences at the beginning and then at the end of the course. The researchers found that although the participants were more intelligible after the course, only the true sentences were considered to be less accented and more comprehensible. Later Derwing, Munro and Wiebe (1998) were involved in a more substantial study with 48 participants, and a control group plus two instruction groups (a segmental one and a suprasegmental one). At the end of the 12 week course, both instruction groups showed improvement in reading sentences aloud, but only the suprasegmental group showed improvement in comprehensibility and fluency when 45-second excerpts of speaking using a picture story were evaluated. The observed lack of improvement in fluency of the segmental group was seen as understandable because if instruction is focussed on segmentals, the participants' fluency could not be expected to improve.

In a similar investigation which produced similar results, Derwing and Rossiter (2003) studied 48 full-time ESL college students from different language backgrounds whom they had divided into three groups; a segmental, a prosodic, and one that did not receive any specific pronunciation training. The students had an average age of 31.7 years, were all studying at the intermediate level and had begun to learn English as adults. The segmental and suprasegmental group received 20 minutes a day of pronunciation instruction over a period of 12 weeks, resulting in a total of 20 hours. They were tested on their speech near the beginning and end of the course using the same tool: a cartoon story of eight pictures in which they had to describe what happened. Both times they were evaluated on their comprehensibility, accentedness, and fluency. The only group that was found to have improved with regard to these criteria was the

group that had received training in suprasegmentals. Although it was noted that the segmental group made fewer phonological errors at the end of the treatment period, they did not show improvement in comprehensibility, accentedness, and fluency.

With regard to research that focusses purely on suprasegmental areas, Wang, Spence, Jongman, and Sereno (1999) conducted a study of native English speakers' acquisition of non-native suprasegmental contrasts, which in this case are Mandarin tones. There were 16 participants at the American university: eight in the control group and eight in the treatment group. All had previously taken one or two semesters of Mandarin. For the study, the treatment group received training in eight 40-minute sessions over a period of two weeks in the identification of Mandarin tones in spoken words in which the tones were presented pairwise. There was a pre-, post-, and six-month delayed post-test in addition to two generalisation tests that were designed to see if improvement could be extended to new words. Statistical tests (ANOVA) showed a significant improvement overall and in the identification of each tone in the post-test and in the two generalisation tests, and this improvement was maintained in the delayed post-test.

As part of a larger investigation into the use of English intonation, Pennington and Ellis (2000) studied the ability of 30 L1 Cantonese speakers between the ages of 20-35 who were at an advanced proficiency level in English to distinguish four types of prosodic elements. They found that with training that simply involved drawing the participants' attention to the pronunciation features, they were able to significantly improve their perception of contrastive focus (focus words). The results for the other three prosodic elements, that is, the different intonations of tag questions, as well as differences in juncture and phrases were not significant.

In another study contrastive focus (focus words), Muller Levis and Levis (2012) gave four hours of instruction to 18 participants over three days which involved listening, production and prediction practice. There were pre- and post-course listening tests in which the same 15 sentences were read but the focus words were changed in the post-test. For the speaking tests the participants read 15 sentences. The results showed that although in listening the gains were not significant, they were in the reading tests.

In a study of French, Sturm (2013) examined liaison in three groups of 11 participants: a phonetics group enrolled in a French phonetics and pronunciation course, a control group enrolled in other advanced French courses, and group of native French speakers whose speech was recorded for purposes of comparison. In the semester-long course, two days were devoted to liaison instruction in French. The participants had listening, reading, and pronunciation exercises and during the classes their attention was drawn to the instances of liaison, thus following Schmidt's Noticing Hypothesis. The researcher found that at the end of the term, participants made fewer forbidden liaison errors when reading a text and attributed this improvement to the instruction they had received and especially the fact that the students' attention had been drawn to liaison.

In a study of teaching linking within and across words in English, Sardegna (2011) used the Covert Rehearsal Model (CRM) to teach English pronunciation to 38 international graduate students ranging in age from 22-47 and with 8 different L1s. The classes were for 50 minutes three times a week for four months. In addition, the students had five thirty-minute private meetings with the instructor. According to Sardegna, CRM advocates that teachers provide students with "predictive skills, pronunciation rules, and strategies that they need to work on the accuracy of their speech in private" (p. 107). The analysis was a mixed method one in which the

participants' pre-, post-, and two delayed post reading tests (over a period of 4-25 months), two questionnaires and a survey were analysed along with notes that the researcher had made about the participants' comments regarding their practice. Results showed that the participants significantly improved their linking across the time period, even though there was some back sliding from the second to third post-test. Additionally, the researcher found that practice and motivation contributed to improved performance.

Later, Sardegna and McGregor (2013) wanted to know how effective student-centered³¹ instruction combined with teacher scaffolding would be on the English pronunciation of 15 international teaching assistants (ITA) with four L1s. The areas that were targeted in the instruction were reduction, linking, primary stress, and intonation. The ITAs were taught for 90 minutes twice a week for 15 weeks. The researchers administered pre- and post reading tests and found that in the post-test, the participants had significantly improved in all areas.

In another study of different suprasegmental areas, Abe (2009) investigated whether pronunciation instruction involving input enhancement can have an effect on rhythm, linking, assimilation, and elision. The two treatment groups of 30 Japanese participants each³² were third-year high school students who were enrolled in a technical college studying English at the low to intermediate level. One group had input enhancement plus explanation (IEE) and the other had input enhancement plus interaction (IEI). In IEE, the participants were asked to listen and pay attention to an instance of connected speech and then were given an explanation of it. The IEI group first listened to two speech samples, one containing connected speech processes and one without, and then discussed them in pairs before sharing what they discovered with the

³¹ Student-centred instruction involves helping students set learning goals based on their needs, empowering them through teacher scaffolding, and providing opportunities for monitoring, practice, and reflection (Sardegna & McGregor, 2012).

³² There was a control group as well, but no information on this group is provided.

class. The participants had a pre-test, a post-test and a one-month delayed post-test involving 20 multiple choice questions³³. The researcher found that after eight sessions of 15 minutes each over a period of five weeks, both of the post-tests indicated a significant improvement in both perception and production from the pre-tests, with the IEI group significantly outperforming the IEE group.

Later, Abe (2011) investigated the use of Form-Focussed Instruction (FFI) on weak forms, i.e. function words such as pronouns, auxiliary verbs, and prepositions that have undergone reduction in speech. The participants were 60 Japanese high school students in their second year who were enrolled in English classes at a technical college and their level was low to intermediate. They were divided into a control group (NFI) and treatment group (FFI), with each having 30 students each. The control group was given an explanation of the weak forms and then listening and repeat exercises while the treatment group listened to two speech samples, one containing the weak forms and one without, and were asked, in pairs, to compare the differences between the two samples³⁴. The participants were given classes for a period of a month³⁵. There were pre-, post-, and delayed post-tests in both perception and production of the weak forms which were a series of 20 sentences. The results indicated that the treatment group had a significant improvement in perception and production in the post-test that was maintained four weeks later in the delayed post-test, and in the case of perception, improved.

In a different kind of study, Gordon (2012) investigated certain extra-linguistic factors that are at play in a pronunciation class and which can affect both the students' learning as well as their perception of the efficacy of the class and the teacher. Within an English intensive

³³ How the multiple choice questions were used is not indicated in the article.

³⁴ Note that this is the same procedure that the IEI group followed in Abe (2009).

³⁵ The total number of class hours is not provided.

program at a large university in the United States, he studied a separate ESL pronunciation class that met five days a week for seven weeks. The class, which was offered to students in the upper two levels of the seven-level program, consisted of nine learners from seven different language and cultural backgrounds who were in their early to mid twenties. Through classroom observations, detailed notes, recordings of the classes and personal interviews, the researcher discovered three categories in which there was disagreement between the students and the teacher. Firstly, there was a conflict between the goals of the students, which was to obtain a native accent, and that of the teacher, which was for them to achieve intelligibility. The second area of conflict was in the choice of classroom activities, some of which negatively affected the students' motivation because they did not consider the activities in question to be pedagogically valuable. The third conflict had to do with the students' expectations of the teacher, which were not always met. They expected the teacher to provide both positive and negative feedback; however, she only provided positive feedback. These additional factors and issues studied here provide insight for teachers as they may find that the realities of the classroom are quite different from the theory that they learn in teacher training.

The above studies reflect an encouraging trend in studies related to pronunciation teaching. Not only have there recently been more studies conducted, but researchers are finding that various pronunciation areas can be successfully taught using different techniques over varying time periods spanning short- and long-term. Furthermore, through examining the studies and some of the preceding literature, we can see that there are a number of guiding principles in the teaching of pronunciation. A summary of these principles can be found in Table 2.1.

Table 2.1 Guiding Principles in Pronunciation Teaching

Recognise that perception training should precede production training	Help learners recognise that pronunciation is important	Help learners notice how their pronunciation is in contrast to that of native speakers
Include both listening and speaking activities	Assist students in setting realistic goals based on social and professional needs	Teach learners to become skillful at monitoring their pronunciation
Concentrate on suprasegmentals but also cover important segmentals	Help students perceive the pedagogical value of activities	Assist students in developing an awareness of sounds and rhythms
Provide intensive training in perception and production	Provide both positive and negative feedback	Help students develop strategies to identify and self correct errors
Ensure that learners have plenty of practice in realistic and useful contexts	Help learners strengthen self confidence through support and positive reinforcement	Teach steps and techniques to strengthen auditory memory, i.e. ability to retain and mimic sounds
Teach in a meaningful and motivating way	Provide regular assessment	Teach steps and techniques to strengthen articulatory skills, i.e. motor skills
Begin pronunciation teaching at lower levels	Adapt pronunciation teaching throughout the levels	Integrate pronunciation within the curriculum and in each lesson
Expose students to authentic target language	Employ the zoom principle (shifting between the form and communicative activity)	Undergo ongoing teacher development and lean on research
Equip students with prediction strategies	Teach pronunciation rules	Encourage students to also practice alone
Explain the pronunciation phenomena to students	Allow students time to discuss the pronunciation areas with classmates	Draw students' attention to the pronunciation areas through visual and auditory means

Traditional Materials

There are a number of different kinds of materials that can be used for teaching pronunciation, although it is important to keep in mind that, according to Isaacs (2009), “currently existing instructional materials on pronunciation do not fit the bill in terms of providing authentic, context-rich activities that provide focussed practice for the specific area of

pronunciation to be targeted, nor do they always draw on research evidence” (p. 4). In this section we will review a number of materials that are used in the teaching of pronunciation as well as some of the tasks that authors have suggested be used to teach and practice the different pronunciation areas.

Minimal pairs³⁶ are mentioned by only a few authors (Tan & Woodworth, 2011; Avery & Ehrlich, 1992). Avery and Ehrlich (1992) indicate that minimal pairs can be useful for making students aware of sounds (segments) that are problematic for them and they suggest using words that students are familiar with so as to help them understand the importance of saying them well. The different tasks that they suggest having the learners do include de-contextualized listening and speaking tasks but the authors do suggest to not overuse minimal pairs and recommend that the sounds be contextualized by incorporating them into communicative activities.

Tongue twisters are another material that can be used not only for practicing segmentals but also suprasegmental aspects such as sentence stress (focus words) and thought groups. Cheng (1998) and Wei (2006) indicate that tongue twisters motivate students. Bareither (2007), *Tongue Twisters, Rhymes & Songs to Improve your English Pronunciation*, advocates tongue twisters not only because they are fun and because her clients ask for them, but also because they may be more effective than non-melodic practice materials. When using tongue twisters, however, the author stresses the need for the learners to be able to correctly articulate the individual contrasting sounds first before practicing them in the context of a tongue twister.

Advertisements are mentioned by Isaacs (2009) as a material that can be used to practice both perception and production of segmental and suprasegmental aspects. For example, they can be used to direct students’ attention to question intonation as well as the communicative use of contrastive stress (changing focus words). Furthermore, when practicing reciting advertisements,

³⁶ Minimal pairs are two words that differ with respect to one sound, for example, the vowel sound in *bit* and *bite*.

students can employ drama techniques which require them to play with volume, voice tenor, and emotion. Avery and Ehrlich (1992) indicate that advertisements are a material that can be used to practice persuasion and tag questions.

Games can be used not only to practice a wide range of pronunciation areas ranging from segmentals to suprasegmentals, but also to raise students' awareness of pronunciation and help them become less dependent on the teacher as a model (Hancock, 1995). Games can also be a means of motivating students (Cheng, 1998; Wei, 2006). They can be used to practice intonation in questions and answers in a game like Twenty Questions (Isaacs, 2009). Avery and Ehrlich (1992), Celce-Murcia et al. (2010) and Cheng (1998) indicate that they can be a way to practice contrasts between individual vowel and consonant sounds, especially in a game like Bingo, where minimal pairs contrasting two sounds are used in place of numbers. Question and answer type games allow learners to practice connected speech phenomena such as assimilation, for example, $d + y = /dʒ/$ as in "Did you...?" (Avery & Ehrlich, 1992).

Poems and limericks are considered good materials for teaching rhythm, linking, as well as vowel sounds, not to mention sound-spelling correspondence in words that rhyme (Celce-Murcia et al., 2010).

Jazz Chants are short chants in the form of a dialogue or conversation in which there is a strong rhythmic beat. According to the author, Carolyn Graham (1978), American spoken English reflects the same rhythm as traditional American jazz. Celce-Murcia et al. (2010) advocate the use of jazz chants because they allow students to use body movements (i.e. tapping, snapping fingers) to show the rhythmic beat or stress timing of English.

Dialogues and conversations are advocated for pronunciation practice by a number of authors (Avery & Ehrlich, 1992; Celce-Murcia et al., 2010; Cheng, 1998; Couper, 2003; Wei,

2006). Dialogues and conversations can be useful for working on stress and intonation as well as segmentals, provided that there is sufficient concentration of the target sounds (Cheng, 1998). Avery and Ehrlich (1992) mention that dialogues are opportunities to practice reduced expressions such as *gonna*, *wanna*, *hafta*, etc. In addition, they stress the importance of modeling the dialogues for the students, while also commenting that, when possible, the students themselves should try to generate dialogues. The authors also suggest providing words that end in consonants and prepositions that begin with vowels as a way of having the students practice linking in their dialogues, for example: “I saw Bobu in the bookstore”, “Did he buy that booku about Atomic Energy?” (p. 169). In addition, they indicate that dialogues can be used to practice language functions (i.e. asking for/giving directions, making suggestions, etc.) and wh- and yes/no question intonation, which provide “contextualized communicative practice” (p. 193). For dialogues, Celce-Murcia et al. (2010) suggest having students first listen for meaning before practicing. They suggest highlighting the stress, pitch, or intonation, for example, or eliciting it from the students, thereby encouraging them to “feel the pronunciation by having them tap on the desk for stress and rhythm, use sweeping hand gestures for intonation, and stretch a rubber band apart for vowel length in stressed syllables” (p. 293).

The above materials that have been discussed are perhaps the most common materials that are used in pronunciation teaching. The list above, however, is not by any means exhaustive. Cartoons, comics, TV and movie clips, as well as monologues and even photographs or artwork can be used as materials for teaching and practicing pronunciation, although they rarely surface in the literature. While possibly any material, if carefully selected and used appropriately in pronunciation teaching, can help learners improve their L2 pronunciation, we are most interested in how songs might be used to achieve this goal. As such, we turn now to

some of the methodologies and approaches to pronunciation teaching before looking at those that use songs. After this we will review the literature on how songs are used in language teaching in general and pronunciation teaching in particular.

Methodologies and Approaches to L2 Pronunciation Teaching

Situational Language Teaching (The Oral Approach) was developed by British applied linguists starting from the 1920s and had its roots in the Direct Method, which was made famous by the Berlitz schools of language learning. Situational Language Teaching was based on a structural view of language and a behaviourist theory of language learning. Therefore, the goal is for learners to use – with automaticity - correct pronunciation and correct grammar and the appropriate vocabulary. For this, drills are used but activities do progress from the controlled to the freer. Finally, even though it is important to learn the four skills, listening and speaking are addressed before reading and writing. With regard to the types of activities that this method uses, listening, isolation of sounds, words etc., choral and individual imitation, and drilling are ones that are typically used for teaching pronunciation. However, because most of these are not communicative in nature, they are often not employed in CLT.

Meanwhile, in the United States, another traditional approach to learning pronunciation appeared. The Audio-Lingual Method, which evolved from the Aural-Oral approach in the 1950s, was originally a response to the growing numbers of international students in the country. It too advocated learning listening before speaking and employed a lot of drills, but it was lacking in a formal psychological theory of learning, which is what Audiolingualism had. In some ways, Audiolingualism is similar to SLT. They were both based on the same view of language and theory of language learning. One major difference, however, was the contrastive

aspect of the Aural-Oral Approach that made its way into Audiolingualism. The types of activities that were used were very much like those of SLT, although dialogues and drills were the basic classroom activities. Correct pronunciation was paramount and the learning of grammatical patterns followed from there. Finally, it is interesting to note that according to Richards and Rodgers (2001), “An important tenet of structural linguistics was that the primary medium of language was oral: Speech is language” and so “it was assumed that speech had a priority in language teaching” (p. 55). This was in contrast to popular opinion at the time when people believed written language to be superior to spoken. Later, though, in the 1960s, Audiolingualism began to decline as a result of Chomsky’s criticism and rejection of structuralism and behaviourist theory with respect to language learning (p. 65).

The method of Suggestopedia or Suggestology, created by Lozanov, a Bulgarian psychiatrist and educator, has music as a salient feature. However, it is used as a learning aid for its affective benefits, namely, relaxation and creation of a positive learning environment, as opposed to any linguistic or cognitive benefits. In fact, Mozart, Bach, Haydn, Handel, Corelli, Beethoven, Vivaldi, Tchaikovsky, Brahms, Couperin and Rameau are the musical selections that the method uses³⁷. And just like its strict selection of music, Suggestopedia has very specific and exact guidelines as to the role and behaviour of both teachers and students as well as the classroom procedures. Richards and Rodgers (2001), in their discussion of the method, indicate that the texts that form the basis of each unit are dialogues containing large amounts of new vocabulary which, along with the grammar, are graded, and that these dialogues are all translated into the learners’ L1. Throughout the course, the students listen to the dialogue, mimic it and

³⁷ From <http://www.scribd.com/doc/16207609/Musical-Selections-In-Suggestopedia>. “This music list for suggestopedia was originally published in "The Foreign Language Teacher's Suggestopedic Manual" (Lozanov & Gateva, 1988) and was later copied in "Suggestopedia - Desuggestive Teaching (Communicative Method on The Level of The Hidden Reserves of The Human Mind)"(Lozanov, 2005)”

practice it orally, make changes to it, ask and answer questions and engage in conversation. Emphasis on correct pronunciation, however, does not seem to be important in the method.

Finally, there is one approach that is worth mentioning, not only because of the developer's claims, but also because of the steps and pedagogically-sound tasks. Kjellin (1999), a language teacher with over thirty years of experience, developed an approach which stems from the "prosody-based, natural acquisition" of a native language (p. 2). In fact, his hypothesis is that "the entire spoken language acquisition (L1 and L2 alike) pivots on pronunciation and particularly *prosody*" that are guided and governed by "perception and constant auditory monitoring" (p. 3). Kjellin admits that the findings of his approach are empirical, that is, based on observation and experience, and not experimental, yet he asserts that "it is surprisingly simple for adults to achieve native-like pronunciation in a second language within only a few minutes of practice" (p. 3). In order to do this he uses a process that involves three steps: *finding* the sounds and sound chunks; *automatizing* sample phrases; and *transferring* the phrases to new ones. For steps one and two, Kjellin indicates that there are three interconnected strategies that need to be implemented. Strategy A, involves using "auditory perception and auditory memory" while strategy B involves laying a "profound prosodic base", and both of these require "attentive listening" on the part of the learners (p. 5). Strategy C means making multiple, (i.e. 50-100) choral repetitions of a phrase. He adds that in order for these steps to be executed quickly and effectively, instructors and students must know that "near-native pronunciation is a realistic goal and not an elusive luxury" and that "persistence and determination is the main requirement" (p. 4). Although believing that native-like pronunciation is a possibility in L2 adults directly contradicts what well-known pronunciation experts assert, Kjellin's theory and approach are based on numerous sources across many disciplines, such as, neurophysiology, audiology,

linguistics, and second language teaching.³⁸ While Kjellin's claims of the success of his method may appear to be unrealistic, his multi-disciplinary approach and implementation of attentive listening/noticing hypothesis and prosody first require, at the very least, the recognition that he uses a method that has a logical theoretical base.

Teaching With Songs

Methodologies and Approaches That Use Songs in L2 Teaching

Communicative Language Teaching has probably been the most popular language-learning approach since the 1980s. This does not mean to say, however, that it is the *only* approach that is followed. Many instructors use an eclectic approach, that is, one that includes elements of other approaches in the different stages of the lesson. With respect to pronunciation, as we saw, there were popular approaches in the past that focussed on it, but music and song, if used in the classroom, were not principally used for teaching pronunciation. Nevertheless, there are some instructors who have developed their own “approaches” or “methods” that do involve the use of music, with some of them having the goal of improving the learners' pronunciation.

The first of these individual approaches is that of Anton's. Anton (1990), who was mentioned above, developed what he calls the *Contemporary Music Approach*, or *CMA*. In this approach he wrote and produced 10 songs specifically for the teaching of Spanish grammar and as supporting material for the main textbook of the course. His rationale for using songs

³⁸ In defense of Kjellin, in my personal experience as a language teacher, I have encountered numerous students and other language teachers who, having begun learning their L2 in their teenage years, on the surface, appear to have managed to acquire a native-like pronunciation. Although no elaborate tests or experiments were performed in order to determine how native-like they may be, the average person would probably not be able to detect a foreign accent in these individuals.

included both cognitive and affective reasons. The three stages that make up the approach are: the Song Introduction Stage, the Recording Stage and the Grading Stage – all of which involve set tasks. With respect to pronunciation, the students practice listening in the first stage, singing in the second, and speaking in the third. According to Anton, “Through the songs, students learn rhythm, intonation, and pronunciation in a natural way as they listen to the music over and over and then attempt to reproduce the sounds they hear” (p. 1169).

Chan (1999) outlines how to practice and contrast minimal pairs and demonstrates and practices the rhythm of English for Japanese speakers. First, she has the students listen to the song once or twice and then she provides them with the lyrics before explaining and demonstrating the articulation of the target sounds using vocabulary from the song. She then shows how the minimal pairs contrast in meaning using words and sentences. After that, students look for, identify and define words in the song that contain the target sound and then work together and read the lyrics to each other. After working with the vocabulary and content of the song, the teacher illustrates stress and rhythm by reading the lyrics while the students indicate the stressed syllables. This may be followed by some additional work on syllables and stress and then the lesson ends with different ideas on reading or singing the song.

In her discussion of the *Melodic Approach*, Mora (2000) reminds us that music and language have some common characteristics and functions. They both have “pitch, volume, prominence, stress, tone, rhythm, and pauses”; they are used to convey meaning; and both skills are acquired “through exposure” (p. 147). The author then draws parallels between L1 and L2 learning with respect to the musical elements of speech by reminding us that intonation is the first aspect of speech that infants acquire before they are exposed to the “highly repetitive, and...simple syntactic structures” of *motherese* in which the prosodic elements may be slowed

down, exaggerated, or otherwise distorted (p. 149). She says that EFL teachers sometimes also distort their speech exposing their students to this kind of exaggerated pronunciation. Finally, in the *melodic approach*, when teaching new language elements, she advocates plenty of repetition in addition to using this type of pronunciation while putting the utterance to music, whether it be known or simply created by the teacher according to his or her musical abilities. The example she provides is putting the question, “What do you do?” to Beethoven’s Fifth Symphony when teaching learners how to ask simple wh- questions (p. 151). When discussing additional benefits of this approach, the author mentions some of the affective and cognitive benefits of using music in the language classroom.

Through discussing the different well-known methods and approaches that have involved either a focus on pronunciation or the use of music, we can see that in only a few of them are songs used for the purpose of teaching pronunciation discrimination or production.

Songs and Their General Usefulness in L2 Teaching

As we saw in the Introduction section of this thesis, the addition of songs to a language learning class has generally been seen to have a positive effect on all of the dimensions present in the classroom. We summarize the main literature on the topic in the following pages. After that, we discuss the literature that deals specifically with songs and pronunciation training.

Literature reviews. Adkins (1997), a fan of Anton’s CMA, also used original lyrics written for language teaching but instead put them to the music of popular songs. After noticing and hearing from both students and parents how enjoyable her learners found her classes, which

validated her feeling that learning with music was beneficial for language learning, she embarked on a review of the existing literature to answer a number of questions as to why this was so. In her review she mentions cognitive and affective reasons for using songs in language teaching and she concludes that songs contain the “necessary factors for learning and memory” because “setting the new language within a familiar context forms strong associations, creates motivation on the part of the learner, and aids in memory storage.” (p. 11)

Schoepp (2001) discusses the affective, cognitive and linguistic reasons for using songs in the language classroom. While not mentioning pronunciation specifically, he makes points that are directly relevant to the teaching and learning of the skill of pronunciation. For example, when discussing the cognitive reasons for using songs, he says that they “present opportunities for developing automaticity” (presumably) because songs are “fairly repetitive and consistent” (para. 10). Regarding linguistic reasons for using songs, Schoepp reminds the reader that songs contain colloquial expressions and says that using songs in the classroom can “prepare students for the genuine language they will be faced with” (para. 11).

Stansell (2005) discusses music and language and reviews literature that reveals the inseparable link between the two. When referring to Gardner’s theory of Multiple Intelligences, Stansell indicates that, although separate, the different intelligences will share tasks, and he goes on to say that “language intonation relies heavily upon perception of musicality” (p. 9). As well, he asserts that people who are musical are more skilled at learning foreign languages because they have “an advanced ability in perceiving, processing, and closely reproducing accent” (p. 11). When referring to music itself, Stansell states that it “positively affects language accent, memory, and grammar as well as mood, enjoyment, and motivation”, and he urges language teachers and music therapists to “collaborate on their joint venture” (p. 8). He also refers to a

number of studies which indicate that using songs in language teaching can help with acquisition (e.g. Wilcox, 1995; Mora, 2000). In support of his assertion, we might quote the experience of Gabrielle Giffords, the U.S. congresswoman who was shot in the head in Arizona in January 2011, as an excellent case in point. Through music therapy, and specifically melodic intonation therapy, she was able to regain the ability to speak after the shooting left her with damage to the left hemisphere of her brain (Bambury, 2011; Giffords, 2011).

Saglam et al's (2010) *Music, Language and Second Language Acquisition: The Use of Music to Promote Second Language Teaching and Learning* touches upon a variety of topics related to music and language teaching. It discusses music in relation to the dimensions of the language classroom as well as the use of music for teaching the different language skills. Finally, it mentions music methodologies, such as CMA and Suggestopedia, as well as some research that has been done on music and language, for example, Lê (1999) and Murphey (1992).

Salcedo (2010) refers to a number of studies which confirm that many learners experience involuntary mental rehearsal, a *din*, when learning a new language (Bedford, 1985; de Guerrero, 1987; Krashen, 1983; McQuillan & Rodrigo, 1995; Parr & Krashen, 1986). In addition, she discusses the musical *din* or *Song Stuck In My Head Phenomenon* (SSIMH) following the term used by Murphey (1990), and she refers to an earlier study, Salcedo and Harrison (2002), which indicated that the "extended mental interaction" caused by songs "may have a profound effect on the second language acquisition process" (p. 23).

Engh (2013a) reviews literature that discusses why music should be used in language learning. To do this he explores sociological considerations, cognitive science, first and second language acquisition, and practical pedagogical resources. He reports on the wealth of literature

and studies that have been carried out in these different disciplines and discusses the available materials. Finally he concludes his article questioning the “disparity between theoretical support and practical application in the classroom” and the “gap in teacher pedagogical resource books supporting the use of music”. Furthermore, he wonders why “the needs of adult and teen learners [are] not reflected in the current web or print pedagogical resources” (p. 120).

Textbooks. Murphey, the leading authority on the use of music and song for language teaching, wrote a teacher’s resource book on the subject. *Music and Song* (1995) is basically a how-to book on using music and songs in the language classroom. It discusses the importance of music and song in language learning, addresses teachers’ concerns, suggests how to integrate music into lessons, provides examples of numerous activities (and variations thereof) for different ages and language levels, and even gives song suggestions for themes and grammatical categories.

Adamowski wrote *The ESL Songbook* (1997) which is an English language learning text that contains many skills activities for 10 original songs. Each song exemplifies a theme or topic that is relevant to some language learners and is accompanied by a number of different kinds of activities such as listening, cultural discussion, writing, grammar and vocabulary. One of the skills that is addressed for each song is that of pronunciation, and specifically the suprasegmental aspects of pronunciation.

Articles. Jolly (1975), who used songs as supplementary material in her Japanese conversation courses, conducted surveys in which students rated songs favourably. On a scale of

1-5, with 5 being *very useful*, in two successive semesters, 80% and then 91% of students rated songs as being “very useful”. The students mentioned that the songs “created a relaxed and enjoyable atmosphere” and “livened up the pace” as well as relieving boredom and making the learners more open to learning (p. 13).

Lems (1996) discusses the benefits of using songs and music in the ESL classroom and provides an inventory of a number of techniques, such as cloze exercises, unscrambling the lyrics, dictation, rewriting/retelling the story of the song, and responding to issues in the song, that target the major and minor language skills, including pronunciation and speaking as well as speech perception. She also outlines a number of criteria for choosing appropriate songs (e.g. clear lyrics, appropriate vocabulary and content) and includes the lyrics to a variety of songs that lend themselves well to the ESL classroom.

Cheng (1998) includes the use of songs in his pronunciation course which was developed for Chinese students of English and which has been described above. With regard to the use of songs specifically, the author says that they need to be “simple enough for the students” to practice pronunciation areas (although he does not explain what he means by *simple*) (p. 39). The procedures he employs involve: reading the words that contain the target areas, indicating the stress, singing the songs to the students, repeating words chorally while tapping on the desk to show the rhythm, and then putting the words into the tune.

Borland (2012) conducted a survey related to the use of songs for teaching English to Mexican middle-school students. In 2009, 48 Mexican English-language teachers came to Canada to further their English language and teaching skills. One of the workshops was conducted on the use of songs in the ESL classroom and the result of their final assignment was

the development of an ESL music and songs activity book that contained over 144 activities and the accompanying handouts to use with 48 different English songs from various genres. According to the survey that was sent to the teachers two years later, the top three skills that the song activities helped the most were pronunciation, listening, and speaking, respectively.

Studies on Songs in Language Teaching. Anton (1990), in what he refers to as his *Contemporary Music Approach*, or *CMA*, uses 10 songs that have been written and produced specifically for the teaching of Spanish grammar and as supporting material for the main textbook. His rationale for using songs is that they lower the affective filter³⁹, require the use of both hemispheres of the brain, and act as an aid to memory. In order to assess the benefits of CMA, Anton administered a questionnaire. Over 200 high school and college students who were beginners in Spanish completed a course using the approach and then later filled out the survey asking them to comment on their experience. According to this, 93% liked learning with CMA, 98% said it helped them learn Spanish, and 50% listened to the music during leisure time.

Lê (1999) conducted a qualitative study in Vietnam to find out how EFL students and teachers view the use and importance of using music in language teaching. He asserts that music is considered to be a very important element to include in a language class, not only because it is an integral part of Vietnamese culture and L1 learning, but also because songs simply help with language learning. In addition to other areas, interviewees mentioned that songs were helpful for linking listening to writing and speaking by having students listen to the song, write the lyrics, and then sing karaoke. It was also noted that songs appeared to help students develop linguistic

³⁹ The Affective Filter Hypothesis was put forward by Stephen Krashen and has to do with the influence of affective factors such as motivation, self-confidence, and anxiety on language acquisition. According to the hypothesis, when affective variables such as fear or anxiety, for example, are high, the affective filter will block language acquisition from taking place. Conversely, when the filter is low, language acquisition will be facilitated.

awareness, specifically in relation to the differences between speaking and writing as well as the difference between linguistic varieties by studying music from different countries. What is specifically salient in this study is the extent to which music and singing is important in Vietnam and therefore should necessarily play an important role in English-language learning, albeit with a certain degree of caution with respect to differences in values and underlying messages in Western pop songs. He says that the Vietnamese cultural context “treasures serenity of the mind, love, care, and share... [whereas] the [American] themes stressing individual freedom, sexual liberation, and social hostility contradict with Vietnamese traditional songs, which stress togetherness, respect, responsibility, and social harmony” (p. 8).

Kennedy and Scott (2005) used a variety of teaching techniques involving music and song with ESL middle-school students in order to determine whether or not they had an effect on speaking and the retelling of stories. The results showed that the treatment group had higher scores and improved significantly in their speaking skills compared to the control group. The speaking skills, a total of 14 items, however, do not include pronunciation specifically other than two which are “speaks smoothly” and “speaks clearly” (p. 261). Rather, they have to do with general oral communication skills, such as turn taking, sticking to the topic, answering questions effectively, etc.

In Salcedo (2010) one of the author’s research questions inquires whether the din is increased with the addition of song in the classroom. According to the results of the questionnaires that the participants filled out, 66.67% of the musical group (those who listened to the songs sung) experienced the din as opposed to 33.3% of the text-only group (those who listened to the songs spoken). The author says that “the use of music and songs to present

material appears to be a more efficient way to trigger mental rehearsal that may in turn be a more effective method to stimulate the language acquisition process” (p. 27).

Engh (2013b) surveyed teachers and performed a needs analysis of their attitudes and current practice with regard to music in the language classroom with the hopes of determining whether or not there truly was such a difference between the theoretical support of the use of music and the application of it in the classroom. He conducted an online survey in which he had 50 respondents of various ages and from a number of different countries and of these respondents, seven participated in interviews. The themes that emerged had to do with the beliefs and attitudes of teachers, challenges of implementing music in the classroom, teachers’ practices and potential limited views, and suggestions for more use of music in language learning. While the author did not report clear results of his findings, he does say that “use of music and song in the language-learning classroom is both supported theoretically by practicing teachers and grounded in the empirical literature as a benefit to increase linguistic, sociocultural and communicative competencies” (para. 45).

Songs and Their Usefulness for Teaching and Learning Pronunciation

The use of songs specifically to teach pronunciation is a less-explored topic. In most approaches and methodologies which have used music or songs, such as Suggestopedia, pronunciation occupied a subordinate position, and when speaking skills were a central focus of a method (e.g. Direct Method, The Silent Way), songs did not appear to be incorporated into the curriculum. As well, there have been few studies that have specifically investigated the possible benefits to L2 pronunciation that could occur through the use of song. Nevertheless, it is known

and commented on by both students and teachers alike that using songs in L2 teaching can help learners' pronunciation. What follows is an account of the various kinds of literature that deal with music or song and the teaching and learning of pronunciation.

Literature reviews. Stansell (2005) in his review of the literature on songs and language teaching says that “musical people have increased aptitude in foreign language learning due to an advanced ability in perceiving, processing, and closely reproducing accent” (p. 11); however, he does not cite any studies that have shown this to be true.

Textbooks. Songs as a material for teaching pronunciation are mentioned in a few textbooks. Avery and Ehrlich (1992), in *Teaching American English Pronunciation*, advocate the use of songs for teaching students to listen for stress and rhythm. Grant (2007) in *Well Said Intro: Pronunciation for Clear Communication* mentions the use of songs only twice: once as a listening task for recognising linking and once as a homework assignment asking students to “use English rhythm patterns” while sharing their favourite line of their favourite song with the class (p. 98). Bareither (2007) in *Tongue Twisters, Rhymes, and Songs to Improve your English Pronunciation*, despite the title of her textbook, makes no mention of songs nor how they can be used for improving one's pronunciation. Adamowski (1997) claims that without understanding the need to make changes, songs will help students make necessary modifications in connected speech as they become comfortable with the rhythm of songs. The author also asserts that “songs are especially useful for teaching rhythm and intonation because the melody of a song acts as a guide and helps with fluency and flexibility in intonation patterns” (p. x). Murphey (1995) says that songs can be used to practice “pronunciation, intonation, and stress” (p. 10).

Celce-Murcia et al. (2010) in *Teaching Pronunciation: A Course Book and Reference Guide*, have a section on how to use songs for this purpose and suggest some different activities for practicing word stress, prominence, linking, final consonants, thought groups and intonation, as well as reduced speech. Specific activities include using the lyrics to indicate word and sentence stress, reading the lyrics by paying attention to thought group boundaries and linking patterns, and filling in the blanks. The authors provide criteria for choosing songs and suggest using ones that the learners will like and want to sing along to. In addition, they indicate that the singer's voice should be clear and that it is good to start out with slower songs having "simple vocabulary, concrete images, and repetition in the lyrics" (p. 353). As well, they indicate that the words should be stressed clearly and that the rhythm pattern sound conversational.

Articles. Jolly (1975) describes songs and speech as "qualitatively similar linguistically" and states, with respect to pronunciation, that "native or folk songs of a given culture naturally follow or reflect the basic meter, pitch, dynamics or other phonological elements and patterns of its language" (pp. 12 & 13). It is, presumably for this reason that she indicates that songs are "effective" in reinforcing pronunciation (p. 13).

Anton (1990) discusses the *Contemporary Music Approach* (CMA) that he uses with his students and which includes pronunciation in one of its phases. He asserts that "through the songs, students learn rhythm, intonation, and pronunciation in a natural way as they listen to the music over and over and then attempt to reproduce the sounds they hear" (p. 1169).

Chan (1999), in her article *Teaching Pronunciation Using Song Lyrics*, outlines a lesson plan for practicing and contrasting the voiced dental fricative /ð/ as in *breathe* with the voiced

alveolar fricative /z/ as in *breeze*, and demonstrating and practicing stress-timed rhythm to intermediate and advanced Japanese-speaking adult EFL learners. When executing the lesson plan, students first listen to the song once or twice and then are provided with the lyrics to read and think about the topic of the song. Then the teacher explains and demonstrates the articulation of /ð/ using vocabulary from the song. This is followed by showing how this sound contrasts in meaning with /z/ by writing minimal pair words and sentences on the board. Students then look for, identify and define words in the song that contain the target sound. After this, the learners work in pairs or groups of three and read the lyrics to each other. Students are then provided with a worksheet containing exercises to help them understand the content in the song. After taking up these exercises, stress and rhythm are illustrated by the teacher reading the lyrics and the students indicating the stressed syllables. The students are then given the second worksheet which includes principles of syllabic stress in English as well as the activity which asks students to count the number of syllables in each line, note the stressed syllables in each line, and then reflect, comment and ask questions about this. Finally, the lesson ends with different ideas on reading or singing the song. What this tells us is that songs are flexible for classroom use because they can be utilized to teach one or many features of pronunciation.

In Lake (2002-2003), when referring to an ESL class in which he used music for 90 minutes once a week throughout a school year, it was noted by his director that “there was a dramatic improvement in pronunciation that she attributed to the singing exercises” (para. 24). In addition to singing, the author only mentions having included “repetition, pronunciation and hand motions” as in-class activities with the students, but he does not explain what they entail (para. 24).

Miyake (2004) discusses why it is important to teach connected speech patterns (especially reduced speech, contractions and elisions) and how to do this using music. As an English language instructor of native Japanese speakers, she recognises the difficulty that students who speak a syllable-timed language have when learning a language like English. She indicates that reductions, contractions, and elisions are challenging for students and that they may be misconstrued as being poor English. In addition to discussing the need that arises due to language differences, Miyake cites student interest in pronunciation as well as some of the misgivings that teachers may have regarding pronunciation instruction. According to informal surveys that the author conducted at five junior colleges and universities in Kanto, she was able to determine that students are very interested in learning pronunciation. In addition, she refers to Willing (1988) and Makarova & Ryan (2000) who also found that interest in pronunciation is very high and even outranks that of learning foreign culture or literature. When suggesting the use of different musical genres from different countries, she addresses concerns that some teachers have about teaching pronunciation, namely that it is boring and imposes a certain type of accent upon the student. Language learning activities that she suggests using with songs for teaching pronunciation are cloze activities in which words with the reduced-speech forms are replaced with underlines with space to write the full written version of the words.

Ebong and Sabbadini (2006), in their article *Developing pronunciation through songs*, discuss the merits of using songs to help students with sounds, stress patterns, connected speech, reductions, and contractions. The authors also indicate that using songs in the language classroom provides students with the repetition that they need in a form of “authentic, memorable and rhythmic language” (para. 1). They go on to say that songs help convince learners that phonetic phenomena such as weak syllables and contractions, for example, are

natural and correct English and that words are not pronounced individually but rather flow together in groups. These ideas and the activities the authors suggest are useful in English language classes because students are not used to studying how English sounds when words are strung together to form meaning.

Van den Berg (2011) discusses stress, the use of music in the language classroom, as well as two approaches that use music. He goes into detail on English rhythm as well as the complex nature of its stress patterns and he indicates that errors in stress “can cause severe problems in comprehensibility” (p. 11). Regarding music, he mentions its affective benefits and its usefulness for helping learners first perceive a suprasegmental linguistic phenomenon before attempting to produce it, as well as how perception is fundamental for self-monitoring and self-correction. In addition, he talks about how language that is put to music (i.e. songs) is more easily remembered than spoken texts. Finally, he provides an overview of two teaching methods that involve the use of music: Suggestopedia and Jazz Chants. Based on his discussion, he concludes by asserting that music can “definitely have a positive effect on the acquisition of the English stress system” (p. 26).

Studies on Songs in Pronunciation Teaching. Sposet (2008), *The role of music in second language acquisition. A bibliographical review of seventy years of research, 1937-2007*, reveals that there have been few studies done that have focussed on pronunciation. In fact, the author found only one study that specifically dealt with this, which was Karimer (1984). The study, which used rhythm to try to improve the accent of 25 Southeast Asian adults learning English found that students who had listened to three songs and three tongue twisters as opposed to those that were given minimal pairs improved more as a collective group in their post test scores than

the minimal pair group. The tests measured their ability to distinguish between nine word initial and word final phonemes. Both groups of students received instruction for 20 minutes, twice a week, for two weeks. The students in the treatment group listened to the teacher singing the verses and were instructed to clap their hands when they recognised one of the target consonant sounds. The students in the control group only heard the sounds spoken in the context of minimal pairs. The author felt that part of the reason the treatment group improved more than the control group was due to the “sensory factors which definitely encourage memory retention” (p. 6).

Schön et al. (2008) conducted three experiments to find out whether song can help beginning language learners perceive word boundaries. In each experiment, the participants were given a continuous stream of randomly repeated nonsense words to listen to over a period of seven minutes. In the first experiment, the participants listened to seven minutes of a continuous stream of six, three-syllable nonsense words randomly repeated. In the second experiment, the only difference was that the synthesizer sang the words instead of saying them by assigning a different tone to each syllable. In the third experiment, the researchers used variable syllable pitch mapping so that the linguistic and musical boundaries did not line up. In the first experiment, the participants were not able to tell the difference between words and syllables. In the second one, the participants learned 64% of the words and in the third, they learned 56%. Finally, the authors concluded that beginning foreign-language learners “may largely benefit from the motivational and structuring properties of music in song” (p. 982). Although these results are interesting, they only pertain to perception at the word and syllable level as no tests of the participants’ production of the words were carried out.

Chunxuan (2009), in his article on the theory and practice of using songs in English language teaching, mentions a very relevant way that songs can facilitate language learning,

especially with respect to pronunciation. He says that listening to songs and trying to understand them helps to promote language awareness in the learner. Furthermore, he says that when students have difficulty singing the songs in the same way that the artists do, they become aware of how native-speaker pronunciation differs from their own. In this way songs encourage learners to explore how sounds are pronounced and combined which will provide them with insights into how native speakers use them to convey meaning). He also indicates that songs are a “concrete” and “accessible” way in which students can acquire the (theoretical) phonological rules of the target language by repeating and imitating the songs and singers (p. 92). In his teaching experiment that he carried out in college English classes, he compared two groups of students; one, the control group, which had 68 hours of classes during the semester with no songs, and the other, the experimental group, which had a total of four of the 68 hours devoted to songs. (The author does not indicate exactly which kinds of activities were used with the songs.) With respect to the experimental group, the researcher indicated that “students who always listen to English songs pay more deliberate attention to pronunciation, phonological rules, stress and intonation than the others and thus pronounce more correctly and speak English more fluently” (p. 92). When discussing the results of the experiment, in particular when the author compared the students’ English scores at the end of the semester, the only information he reported was that the experimental group did significantly better than the control group, with a mean score that was 6.58% higher and with a *p* value of .02. It is interesting to learn that the treatment group’s overall English scores improved with only four hours devoted to songs. However, with respect to pronunciation specifically, the author did not specify how the group improved.

Fischler (2009) completed a month-long, 32 hour, action research project with six ESL learners aged 13-17 years whose proficiency level in English was intermediate to advanced. In

this study, the researcher wanted to determine whether intelligibility could be improved through instruction of word and sentence stress that included the use of rap songs. The researcher-teacher used phonetic instruction, contrast of correct and incorrect speech, rhythmic practice with songs, and then communicative speaking exercises. Pre- and post-course speech from readings and picture story descriptions were compared and results determined that five out of the six participants had higher intelligibility ratings post-course.

Rengifo (2009) designed and executed an action research project involving the use of karaoke in order to help improve the EFL students' pronunciation. The 12-15 (presumably, Spanish-speaking) Colombian participants were adult learners ranging in age from 18-60 studying at an English education institute. The procedures involved in the research included talking about the song, listening to the teacher singing it, the students singing it (alone or in a group), and a discussion of the lyrics. Specific procedures and tasks related to pronunciation (in which both American and British English were the target models) involved explanations using the International Phonetic Alphabet (IPA), minimal pair and intonation activities, as well as matching sounds in a sentence, and looking for sound patterns. At the end of the cycle of research, Rengifo found that using karaoke not only improved the students' pronunciation, but was also beneficial in that they reflected on their learning, showed increased motivation and more confidence and less fear when speaking. Although the author reported that the students' pronunciation, motivation, and confidence improved, no tangible results are reported nor does the author mention having administered pre- and post-intervention tests.

As we have seen in the literature review, the studies that have been done on pronunciation training whether with or without the use of song, indicate that pronunciation can indeed be improved through intervention. Although there were a number of studies that address

pronunciation training in English language acquisition, we only found a few that combined pronunciation training with the use of songs, namely Karimer (1984), Schön et al. (2008), Chunxuan (2009), Rengifo (2009) and Fischler (2009). Our thesis is intended to remedy that deficiency. In the review of the literature and research related to both pronunciation teaching and the use of songs in pronunciation teaching, it became clear to us that there were studies that were either lacking in one methodological aspect or another or they simply sought information that was different or less specific than that which we seek.

Our study intends to fill this gap in a number of ways. First of all, by teaching a number of suprasegmental areas, we hope to gain a greater understanding as to which areas may be easier or more difficult to learn in terms of perception as well as production. The inclusion of a control group and no-songs group will help us determine the efficacy of the method using both songs and non-song texts. Finally, a mixed method approach will serve to add legitimacy (or not) to the untested assertions that have been made regarding the benefits of using songs for pronunciation teaching.

Summary

This chapter began with a review of the literature surrounding pronunciation's subordinate status. Following this, we looked at the notion of accent, including L2 speakers' opinions of their accents in some L1 and L2 environments. What different scholars have to say about accent addition and accent reduction was explored, as were the nativeness principle and the intelligibility principle. After that, we reviewed further literature related to consequences of accented speech before reviewing the different kinds of literature related to pronunciation

teaching. We looked at what scholars suggest are the best areas of pronunciation to focus on, as well as a specific approach to teaching suprasegmentals. After looking at articles and studies related to pronunciation teaching and an exposition of what authors say about traditional materials that are used in pronunciation teaching, methodologies related to pronunciation teaching and teaching with songs were reviewed. Finally, we looked at articles and studies that have been written on songs and their usefulness in language teaching in general as well as for teaching pronunciation.

Chapter 3

Methods and Procedures

Introduction

In this chapter, we describe the methods and procedures that were followed in a two-week study in Chile designed to explore the effectiveness of songs in pronunciation teaching. The research involved the administration of a short pronunciation course and evaluation of questionnaires and listening and speaking tests obtained with a group of 23 university students of English. The course was taught to one group of ten students using songs as material, and to a second group of five using texts other than songs. A control group of eight students who did not take the pronunciation course or have any instruction whatsoever from the researcher was included for comparison of results. In this way we were able to explore: (1) whether teaching suprasegmental pronunciation aided the students in question and (2) whether songs in particular were effective in teaching suprasegmental pronunciation. The chapter begins with a discussion of the research and how it was designed. The criteria regarding the participants, the number of groups, and the course length, content, and approach are covered before discussing the steps involved in designing and implementing the intervention. The steps that each group of participants took are outlined as well as the execution of the pronunciation classes, along with and the reasons for techniques used. Information is provided regarding the data sources, that is, the questionnaires and tests that were given to the students before and after the treatment period. Finally, there is a discussion of how these data sources were evaluated.

Research Paradigm

The current study is an exploratory case study of first-year Chilean university students studying a Bachelor of Arts in English Language and Literature at the Pontificia Universidad Católica de Chile. Applying Gilbert's Prosody Pyramid to a group of Spanish-speaking students and evaluating their pre- and post-course performance in the pronunciation areas, I sought to provide guidance for ESL/EFL teachers as to which areas of the pyramid Spanish-speakers in particular might have more or less difficulty learning, along with whether songs can facilitate the perception or production of the pronunciation areas.⁴⁰ I chose to carry out the study in a classroom setting in which both qualitative and quantitative data were collected through questionnaires and tests.⁴¹ Additional qualitative data in the form of participant observation is included. Action research was considered to be the best means of carrying out the study because, for teachers, the purpose of action research is for "gaining a better understanding of their educational environment and improving the effectiveness of their teaching" (Dörnyei, 2007, p. 191). Since this thesis is on *teaching* pronunciation, with the use of songs to do so, the most viable way is through action research, and as such, it is research that may be accessible and meaningful for teachers wishing to improve their understanding and competence in teaching pronunciation.

⁴⁰ Duff (2008) discusses how exploratory case studies can "open up new areas for future research, by isolating variables and interactions among factors that have not previously been identified for their possible influence on the behaviour under investigation...and reveal new perspectives of processes or experiences from participants themselves" (p. 44).

⁴¹ Although case studies are generally considered to be qualitative approaches, Dörnyei (2007) mentions that quantitative data collection instruments are often included in case studies as does Duff (2008).

Research Design

Criteria. The studies mentioned in the previous chapter, though few in number, provided some insight and ideas as to how to design an investigation related to the use of songs in pronunciation teaching. Criteria related to the choice of participants, number and type of groups, the inclusion of listening tests, along with which model to use and which pronunciation areas were to be taught were considered when designing the study.

Sampling. One aspect that was important to consider was the sampling of the participants. The decision was made to use homogeneous sampling, which is a purposeful sampling technique in which participants share similar characteristics, for example, age or background (Dörnyei, 2007).

First of all, it was important that the participants have the same L1 in order to truly measure their progress, especially since a learner's L1 may affect the ease and rate with which L2 pronunciation will be acquired depending on how similar the two languages are typologically. Determining specifically which English pronunciation areas pose problems for Spanish speakers will contribute to a much-needed common pronunciation curriculum for L1 speakers of Spanish (Thomson & Derwing, 2015).

Secondly, we wanted participants with a similar proficiency level. The reasons for this are three-fold: Firstly, it was expected that their proficiency with respect to pronunciation would be similar, allowing us to more easily compare their progress. Secondly, it would be easier to design the materials, focussing on a particular proficiency level, and allowing the students to concentrate on the pronunciation phenomena being taught as opposed to unfamiliar grammatical

and lexical items.⁴² Thirdly, it was important that they all be EFL learners, not only so that they would have the same L1, as mentioned above, but also so that they would have had a similar degree and type of exposure to the L2. Finally, we wanted the participants to be of similar age. By limiting the age range of the participants, we limit any possible influence that the age factor may have within and across the groups.

Groups. Another aspect that we considered was the number of groups. For this study, it was decided that there should be three groups: a group that received pronunciation instruction with songs (songs group), a group that received pronunciation instruction without songs (no-songs group), and a group that did not receive any pronunciation instruction or intervention whatsoever (control group). This way, it would be possible to assess whether the addition of songs to a pronunciation curriculum could help, and if so, whether they could help more or less than the use of other materials. The existence of a control group would allow us to determine if any progress made by the songs and no-songs group could indeed be attributed to the pronunciation classes or if the gains were simply a result of the fact that the participants were enrolled in and simultaneously taking their regular English language classes in their undergraduate program. To be clear, all three groups were enrolled in their regular program of studies, which included the following courses: English Language, Applied Grammar, Literary and Linguistic Research Methods, and Anthropologic-Ethical Development, plus an elective from another discipline. As part of English Language, for 80 minutes once a week, the students had a module called Phonics, which involved understanding spelling conventions and homophones as well as sound discrimination and transcription of English segmentals with an

⁴² Using materials which are grammatically or lexically too easy or too difficult could serve to undermine the participants' level of motivation.

initial/partial introduction to suprasegmentals. The songs and no-songs groups were exempt from going to their English Language classes and instead attended my pronunciation classes. The control group simply attended their regular university classes but they took my pre- and post-course tests (for comparative purposes). How it was possible to manage three groups will be covered below.

Listening. An aspect of the study that I felt should be taken into consideration is that of auditory phonetic discrimination as opposed to just production because there are researchers who believe that production is not possible without perception (e.g. Cheng, 1998; Setter & Jenkins, 2005). For this reason, pre- and post listening tests were administered with the hope that they may provide insight into the relation between pronunciation listening and speaking skills. To be clear, though, a full examination of the correlation between the two skills falls outside of the scope of this research.

Pronunciation areas. Though intelligibility, comprehensibility and accentedness⁴³ are notions that apply to a person's speech in general, I set out to investigate specifically which suprasegmental and connected speech areas were the ones that were contributing to a Spanish speaker's intelligibility, comprehensibility and accentedness in English⁴⁴ – and which of these were easier or more difficult to improve through targeted instruction. That is, I wanted to know which suprasegmental features and connected speech phenomena were the most or least

⁴³ Note that these terms were defined in the Introduction.

⁴⁴ Trofimovich and Baker (2006) indicate the need for studies that investigate how suprasegmental areas influence intelligibility, comprehensibility and accentedness.

challenging for Spanish speakers, opening the way for further studies to target their specific needs.

Therefore, the decision was made to teach the pronunciation areas that are contained in Gilbert's (2008) Prosody Pyramid, namely thought groups, focus words, syllabic stress and reduction, including the schwa as explained in the Review of Literature chapter on pages 76-77. A copy of Figure 2.1 The Prosody Pyramid (Gilbert, 2008), which was presented in Chapter Two of this thesis, is provided below.

Figure 2.1 The Prosody Pyramid (Gilbert, 2008)



In addition to the elements in the Prosody Pyramid, other forms of reduction and certain features of connected speech formed part of the content taught. As previously indicated, reduction can occur through the elision of sounds and syllables. This may arise in the pronunciation of some content words⁴⁵ such as, for example “chocolate” or “Catholic” in which the second vowel letter is not pronounced. This will also happen in words or word combinations that have consonant clusters, such as “softness” in which the second consonant, /t/, is not pronounced, and “left field” in which, again, the /t/ is elided. It is much more common,

⁴⁵ Content words are: nouns, main verbs, adjectives, possessive pronouns, demonstrative pronouns, interrogatives, not/negative contractions, adverbs, adverbial particles (Celce-Murcia et al., 2010, p. 212).

however, that certain function words⁴⁶ be reduced in natural speech. For example, the word, “have” in “He should have stayed in bed” is normally pronounced as a schwa plus /v/ or simply as a schwa as in “shoulda”. The connected speech phenomena that was covered included the linking of consonants at the end of a word to vowels at the beginning of the following word, such as in the case of “stop it”, which sounds like “sto pit”. In addition, when there are words that end in a /d/ or /t/ and are then followed by a word beginning with /j/, such as “could you” or “don’t you”, these two sounds merge and become a different sound: /dʒ/ and /tʃ/, respectively. Contractions were also covered in the course because, although the participants knew about them, they are not commonly or consistently used by most language learners. Another common form of connected speech that was included is the pronunciation of the auxiliary verbs “got”, “going”, “have” and “want” followed by “to”. In natural speech, they are usually pronounced like, “gotta”, “gonna”, “hafta” and “wanna”. Again, these forms are not unknown to language learners; however some do not realize that it is natural, normal and *acceptable* to use these forms in everyday speech.

Overview of the Course. In order to meaningfully determine the impact of pronunciation classes on the pronunciation of second or foreign language learners, and specifically the usefulness of songs in that context, I developed and taught a nine-hour pronunciation course to two groups of EFL students over a period of two weeks with each class being 80 minutes long. One group received pronunciation instruction with the use of songs included while the other group received pronunciation instruction without songs. The control group did not receive any pronunciation

⁴⁶ Function words are: articles, auxiliary verbs, personal pronouns, possessive and demonstrative adjectives, prepositions and conjunctions (Celce-Murcia et al., 2010, p. 212).

training or other form of intervention from me during the two week period that the pronunciation classes were taking place.

With regard to the length of the course, nine hours of instruction over a period of two weeks, while short has nevertheless been shown to be beneficial (see Couper, 2006; Karimer 1984; Muller Levis and Levis, 2012; Sturm, 2013; Wang et al., 1999). In my case, although my pronunciation course was short in length and with a limited number of hours, it was carefully designed with regard to the content taught as well as the presentation, sequencing, materials, activities, and review and recycling of pronunciation points. Having seven 80-minute classes over a period of two weeks allowed for a certain continuity and concentration on the material taught, which hopefully provided the students with the foundation needed to continue to learn independently.

Pre-Intervention Steps. The week before the classes began, the students received an information session in a classroom on the San Joaquín campus of the Universidad Católica de Chile. Almost all of the students from the course, English Language II, were there. I was introduced to the students by Dr. Miriam Cid Uribe and then was given the opportunity to talk to them about the investigation. In this session, the nature of the study and their possible participation in it were explained and discussed. As well, they were given a recruitment text in Spanish which they were able to keep (see Appendix B). After this, those who expressed interest in participating in the investigation signed the consent forms and filled out the initial questionnaire (see Appendix B, which includes copies of University of Ottawa Ethics forms and Pontificia Universidad Católica de Chile permission form). The rest of the week I spent analysing the responses on the

questionnaires and assigning the students to the control group, the songs group, or the no-songs group.

The Participants. Creating a potential homogeneous group of participants fortunately was not difficult. After living, studying, and working in Chile for six-and-a-half years, I wanted to return to the country and work with a similar profile of students that I had had in the past. I spent four years of that time teaching core English language courses to students enrolled in a five-year English language pedagogy program in which the students graduated with degrees allowing them to teach the English language at the high school level. Because I taught courses in each year of the program, I regularly had the same students, which allowed for a strong student-teacher relationship and rapport to be built up and maintained. At that time, I was a rather inexperienced teacher who found herself in a very demanding position. Because of the generous and patient support and guidance that I received from my Chilean professors while completing my Master's degree at the Pontificia Universidad Católica de Chile, and colleagues and students at the Universidad del Bío-Bío, I wanted, in some way, to be able to contribute to English-language education in that country by studying a case which was similar to that in which I taught.⁴⁷ Luckily, I had former colleagues and professors who generously offered to help me with this. Dr. Miriam Cid Uribe helped me find the potential participants at the Universidad Católica de Chile. In addition, with the cooperation of the English language professors in the program as well as the consent of the Dean, Dr. José Luis Samaniego, it was possible to find the ideal group of participants and set the dates for when the investigation would be carried out.

⁴⁷ The reason for the above brief discussion is the following: Duff (2008) mentions how it is unfortunate that many researchers "do not mention how or why a particular case was selected or what the relationship or history was between researcher and researched and what bearing that relationship had on the research process or interpretations" (p. 118).

The 23 participants as a whole were a homogeneous group in the sense that they were Chilean, Spanish-speaking, first-year students in their second term studying a Bachelor of Arts in English Language and Literature at the Pontificia Universidad Católica de Chile in Santiago at the San Joaquín campus. This involved taking five courses that term: English Language, Applied Grammar, Literary Research Methods, Anthropologic-Ethical Development, plus an elective from another discipline. They were considered to be at a lower-intermediate to intermediate level because they had already completed a full semester in which all of their English classes were conducted entirely in English at the pre-intermediate-intermediate level. Placement in those classes was based on the Cambridge Quick Placement Test (QPT) that the students had to take before enrolling in the program. The QPT is a diagnostic test of English language proficiency in which test takers answer a series of multiple-choice questions. The scores are then linked to the Association of Language Testers in Europe (ALTE) and Council of Europe Levels to enable students to be accurately placed into language courses and programmes. Furthermore, none of the participants had any hearing or speaking disabilities. In addition, as will be seen in further detail when we look at the 23 participants' responses to the Initial English Pronunciation Questionnaire, we will see that they range in age from 18 to 20 years old; all but 2 have never lived abroad; 19 are female and 4 are male; only 2 have any in-depth knowledge of any foreign languages other than English; and they all have a high level of desire to improve their English pronunciation, and confidence that they can. With regard to the two that lived abroad, one spent three months in the United States when she was 15 years old and five months when she was 18 years old. The other person spent seven months in the United States when she was 17 years old. As for in-depth knowledge of foreign languages, one person claimed to be an advanced learner of French and another had a pre-intermediate level of German. Finally, it is

important to mention that the majority of the participants have career aspirations that would require a very high level of English. Nine of them wish to teach English and five want to use it to study and/or live abroad. Of the others, one wants to become an interpreter, another either an interpreter or a teacher of English, one a translator, another a grammarian or phonetician, one an editor of English publications, and one wants to complete a Ph.D. in an English-speaking country. Finally, one indicated the desire to use it to achieve goals, another to have English as a special skill, another for continuing studies, and one did not answer the question.

Assignment to Groups. The participants were, for the most part, randomly distributed to one of the three groups. In certain circumstances they were assigned to a particular group partially based on their answers to some of the questions in the Initial Pronunciation Questionnaire⁴⁸. This questionnaire allowed me to ascertain the students' age and native language, as well as to have an idea of the extent of their knowledge of other languages and their taste in music. Firstly, it was important to know what kinds of English accents they were interested in learning and which ones they were not. If they indicated that they did not wish to learn a General American / Canadian accent, then they were immediately placed in the control group, because General American and Canadian are the accents that are in the course and of the researcher and I did not want to expose them to a model that they were not interested in emulating. Had any students indicated that they did not like the music of English-speaking countries, then they would have been placed in the control group or the no-songs group. However, as it turned out, they all said that they liked this music. Next, it was important to consider musical aptitude and motivation. For this it was important that each group have similar distribution of participants regarding musical aptitude and

⁴⁸ Further details about the Initial Pronunciation Questionnaire are provided on pages 130-132.

motivation and so, as much as possible, members indicating high and low musicality were distributed evenly⁴⁹. It was necessary to do this to avoid the possibility of the songs group having a disproportionate number of participants with high musicality. With regard to motivation, it turned out not to be a factor in the allotment of participants to a group because all but one student indicated a high level of motivation. (Incidentally, this particular participant also did not indicate a wish to take the pronunciation course and so was placed in the control group.) Determining why the students were learning English and how they plan to use the language in the future provided an idea of what might motivate them extrinsically, although this was more a matter of curiosity on the part of the researcher and not a factor in the division of groups. Finally, the few males in the study were spread among the groups. The end result of these placement factors was as follows: the control group had 8 members, the songs, 10, and the no-songs group, 5⁵⁰.

The Schedule. On Day one, all of the groups were together in the language lab to take the pre-course listening and speaking tests. Once these were done, the participants were informed as to which group they belonged and the members of the control group left and went to their regular program classes. The songs group had their pronunciation classes with the researcher during module (period) one, which was from 8:30am – 9:50am Monday to Thursday in week one and Monday to Wednesday in week two. The no-songs group had a less consistent schedule due to timetable conflicts with program classes that they could not miss. Their schedule during the first

⁴⁹ Musical ability was only evaluated through self-assessment questions. Though this is limited by being purely subjective, any objective test would have required input from musicians and/or musical theory which would clearly go beyond the scope of this thesis.

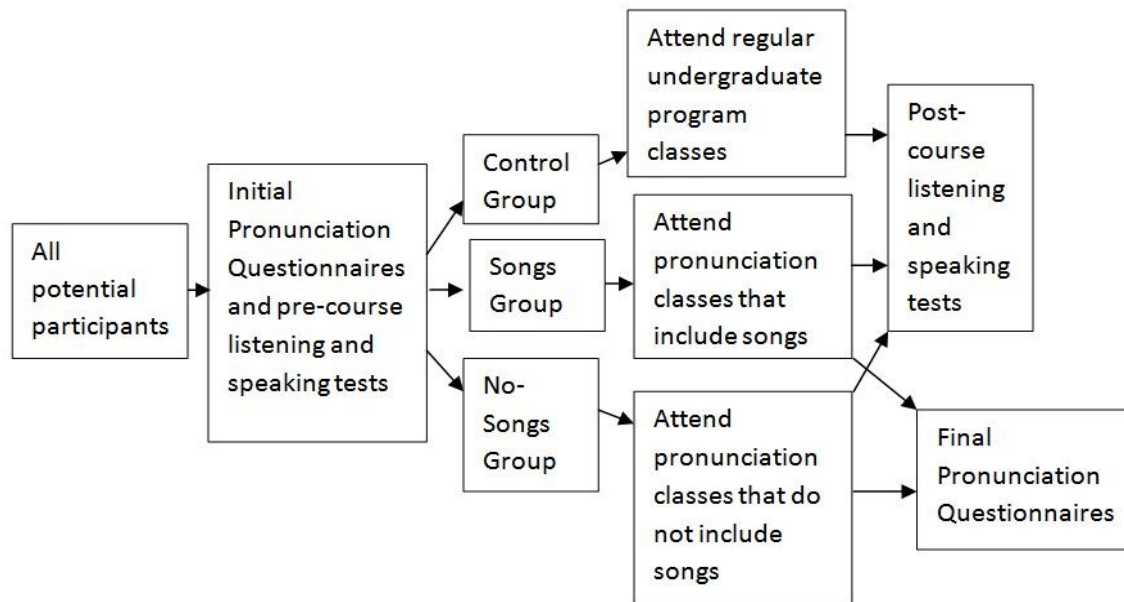
⁵⁰ Every effort was made to have an equal number of participants in each group; however, due to some timetable conflicts and illness, participants either dropped out or were asked to switch to another group.

week on Day one was in the lab with the songs group from 8:30am – 9:50am⁵¹. Then they had their pronunciation classes in a classroom from 10:00am – 11:20am Wednesday to Friday of that first week. During the second week, the no-songs group had a class from 2:00 – 3:20pm on the Monday, and then from 10:00 – 11:20am on the Wednesday and Thursday. Again on Day 8, during the Thursday module one time period for the songs group, members of the control group joined them in the lab during the second half of the period to take the post-course listening and speaking tests. The members of the no-songs group took their post-course listening and speaking tests on that same day in the lab but during the second half of module 4, which was from 2:00 – 3:20pm. After completing the post-course listening and speaking tests, members of the songs group and no-songs group completed the Final Pronunciation Questionnaire.

The diagram in Figure 3.1 shows the groups and the data sources, that is, the sequence of events that the participants underwent from the beginning to the end of the research period in Chile. A description of the questionnaires, pronunciation classes, and tests follows the diagram.

⁵¹ Note that not all of the planned material for the Day 1 class was covered during that first class. Only content that was common to both the songs and no-songs group was covered. The remainder of the Day 1 activities were covered on Day 2.

Figure 3.1 The Groups and Data Sources



Data Collection Sources. Some of the materials that were developed or used for the investigation were the Initial English Pronunciation Questionnaire, the Final English Pronunciation Questionnaire, as well as two listening and speaking assessment tools.⁵²

Questionnaires. The Initial English Pronunciation Questionnaire contains three parts. Part one is a series of questions requesting both personal and general information. As mentioned above, some of this information assisted in determining the composition of the group of participants as a whole and to plan a similar cross section across the groups. The questions in part one ask about their age, sex, place of birth, places lived, languages spoken, reasons for

⁵² The necessary approval, authorization, and consent forms required by the University of Ottawa Ethics Committee can be found in the Appendix A section of this thesis. When necessary, the documents were written in both Spanish and English. The questionnaires may be found in Appendix B, the lesson plans in Appendix C, and the listening and speaking tests in Appendix D.

learning English, plans for the future, previous pronunciation training, preferences regarding English accents, and opinion of English music.

Part two of the questionnaire asks a series of questions to be answered on a six-point scale ranging from *strongly disagree* to *strongly agree*. These questions pertain to: the students' attitudes and opinions toward accent and improving their accent; how they feel about their musical abilities, as well as language learning using music. The questions themselves do not follow a thematic order, to avoid influences between one question and another and to try and ensure more reliability in the answers given. The reason for asking these questions was to try to get a clear sense of how the participants feel about accents and pronunciation, their motivation to improve their English pronunciation, and how musically inclined they feel themselves to be.

Finally, part three of the questionnaire, in addition to asking whether or not the students would be interested in taking a pronunciation course, asks the students to list their favourite English songs and the groups or performers who sing them. Initially, when developing the questionnaire, the researcher assumed that there would be enough time between administering the questionnaire and the beginning of the course to develop pronunciation activities using songs that the participants indicated that they liked. That turned out not to be the case and so other songs were chosen.

The Final English Pronunciation Questionnaire, which only the members of the songs and no-songs groups filled out, contained 22 six-point scale questions, most of which were identical to those in part two of the Initial English Pronunciation Questionnaire, although there were also questions that asked them to reflect upon the pronunciation course and whether or not they felt it was helpful to them. It was important to ask the same questions after the course

because a comparison of the pre- and post-course responses is fundamental for determining the extent to which taking the pronunciation course affected them with respect to their level of motivation, their opinions regarding accent, as well as the perceived effect of music and pronunciation instruction on their learning.

Tests. Two listening tests were given to all of the groups before and after the course. Given that one of the research questions asked whether the use of songs in a pronunciation curriculum could enable L2 learners to better perceive suprasegmental phenomena of the target language, it was necessary to have pre- and post-course listening tests to determine if the course had an impact on their ability to better perceive the pronunciation phenomena taught. As well, since one of the guiding principles in pronunciation teaching is that perception is key to production, giving the students listening tests would give us an idea as to how the participants' listening scores compared to their speaking scores in the different pronunciation areas. While we have not set out to fully test this principle, it is interesting nonetheless to have a look at how the listening and speaking scores compare.

The reason for having two different listening tests is because the first one tested the areas in the Prosody Pyramid while the second one tested the connected speech and linking areas that were taught in the course.

The first listening test was taken from *Clear speech: Pronunciation and listening comprehension in North American English: Teacher's Resource Book* (3rd ed.) by Gilbert (2009), pp. 92-95. The seven-part test assesses consonant and vowel sounds, the number of syllables in words, word stress, focus words, reductions, and thought groups. That is, all of the areas of the Prosody Pyramid are contained in the test. Even though segmentals were not the

focus of our study, we decided not to truncate the test by eliminating the first two sections of it, which deal with consonants and vowels. Each section has its own corresponding audio of someone speaking in a General American accent to which the test-takers listen and indicate on the test what they hear. The test required the participants to listen to words, sentences and dialogues and identify the correct answer by choosing between alternatives, underlining words or syllables, or filling in the blanks. The instructions are clear and easy to follow, even for students who have no real knowledge of the different areas of pronunciation. The limitation of the test is related to the audios: the speakers do not speak in an authentic way. That is, the speech is too slow and too clear. On the one hand, this can be seen as a strength if the students have a very low proficiency level or have especially weak listening skills. However, for language learners who are used to native speakers or who have strong listening skills, the test will not provide accurate insight into which pronunciation areas they might have difficulty hearing when faced with an authentic speaking situation. Since the participants were EFL learners and identified as being at a lower-intermediate to intermediate level, I felt that the test would be appropriate, especially since the text from which it came was designed for intermediate and high intermediate students (Gilbert, 2005b). Despite the participants' identified level of proficiency, it seems that the test was a little too easy for them because they did exceptionally well. That is, they made fewer mistakes than expected considering that the test was designed for learners with a slightly higher proficiency level.

This test was given twice, once on day one before any classes and then a second time on the last day after the classes were done. The audios were played once and one right after another. The students did not receive any feedback on their performance nor were they told that they would be given the same test after the course.

The second listening test is a dialogue that contains specific instances of the types of connected speech and linking that were covered in the course. The first four lines of it come from *Well said intro: Pronunciation for clear communication* by Grant (2007), p. 133. The rest of the dialogue was created with the course content in mind; that is, I was careful to include the specific areas of connected speech and linking that were taught. It is a short dialogue with a total of ten turns. It is essentially a dictation exercise in which the dialogue was played in its entirety and the participants had to write what they heard, so the participants were evaluated on their connected speech and linking perception skills. The weakness of the dialogue has to do with the recording itself. Only my voice was used since I had no recourse at the time to another native English speaker. While recording the test, I was conscious of the need to ensure that I produced the connected speech and linking phenomena as authentically as possible. I also tried to ensure that I spoke neither too slowly nor too quickly. In the end, after a number of attempts at making the recording, I was finally able to produce one that satisfied the above criteria.

As with the first test, this one was given twice, once on day 1 before any classes and then a second time on the last day after the classes were done. The audio was played once and the students did not receive any feedback on their performance nor were they told that they would be given the same test after the course. Finally, the listening tests were done in the university's computer-equipped language laboratory, where the participants used headphones and could not easily copy from one another.

There were also two speaking assessments that were given before and after the course. These tests, as with the listening ones, were given in the university's language laboratory. The participants used headphones and spoke into the computer while recording themselves speaking. Once they were finished, they immediately e-mailed the researcher their recordings.

The first speaking assessment was a controlled one in which the participants read a five-paragraph text that was slightly less than a page long. This particular reading, called “The Rainbow Passage”, was chosen because it is a reading that has often been used to test the production of connected speech (<http://everything2.com/title/The+Rainbow+Passage>). I liked this text because it was neither too easy nor too difficult for the level. In addition, it is a meaningful passage, rather than a collection or list of unrelated sentences as is often the case in pronunciation textbooks. The only drawback to the passage is that it does not contain all of the connected speech phenomena that were taught in the course. The two types of assimilation, /d/ + /j/ = /dʒ/ and /t/ + /j/ = /tʃ/ do not surface in the reading of the passage. Nevertheless, it does require application of all of the areas in the Prosody Pyramid as well as consonant-vowel linking. This assessment, *Speaking Test # 1*, was first given on day one and then again on day eight.

Both times, while taking the test, the students were given a photocopy of the passage which included the following instructions: *Read the following text to yourself. When you are ready, record yourself by reading it out loud into the computer. Return this piece of paper to your teacher after you have finished recording your speech. You may write on this sheet.* If the students had questions regarding the particular pronunciation of a word, I told them what it was. It was important to do this because if there are unusual words in a reading, there is more likelihood that they will be read incorrectly (Levis & Barriuso, 2012). Since my focus was on the participants’ production of suprasegmentals, allowing them to stumble over the pronunciation of unfamiliar words could have resulted in a greater number of errors (in proximity to the “strange” words) at the suprasegmental level than the participants would otherwise not have made. Nevertheless, the “The Rainbow Passage” contains few words that were troublesome for the participants and only seven students in total asked for clarification on the pronunciation. The

particular words that the students asked me to model were: *arch* and *Aristotle*. Furthermore, the students were permitted to listen to their recording and re-record themselves until they were happy with what they had produced. Because reading a passage is a controlled activity, I wanted to be sure that the participants had the opportunity to submit their best attempts at the reading. Reading aloud a text, albeit a common technique used for identifying correct pronunciation, is not always a familiar task for everyone nor will it necessarily reflect the same pronunciation patterns as free speech (Levis & Barriuso, 2012). It is, however, a means of assessing the students' application of the pronunciation phenomena covered in the course. By allowing them to submit a reading that they were happy with meant that I might have a more reliable indication of how well they learned the content covered in class.

The second speaking test, *Speaking Test # 2*, required the students to record themselves speaking freely into the computer for at least two minutes⁵³. Initially I had asked them to talk about the on-going student strikes that were happening in Chile at the time, but they indicated that they were not interested in talking about that. I had also considered the idea of using a picture story, that is, a sequenced series of scenes or events drawn on a piece of paper. I decided against this task, however, because I wanted the participants to be able to speak on a topic that would allow them to forget that they were being assessed. By allowing them to speak on a topic that was important to them, the students' attention could be on the telling of a story that was meaningful to them, rather than on the fact that they were being assessed, which would have been the case if they had had to tell a story that they had no personal connection to. I wanted them to talk to me and it was important that they be comfortable doing so. Conducting a personal interview would have also had drawbacks, the main one being a lack of time.

⁵³ Ideally to collect a true sample of spontaneous speech more than two minutes of speaking is required. Since this was one of four tests that participants had to complete within an 80 minute period in the lab, it was not possible to ask them to produce a longer sample.

Interviewing 23 participants before and after the course was simply out of the question. Furthermore, had it been possible, the social dynamics between us would have changed from the before the course to after the course. For these reasons, I decided to have the students speak into the computer and talk about something that was of interest to them. After discussing possible topics of discussion with them, most of the participants said that they would like to talk about a book that they had read or a movie they had seen. There were a few students, however, who did not find these two options appealing, so they were allowed to talk about whatever they wished. For the post-course recordings of the test, many of the participants chose to discuss other topics. One of the topics some of the students talked about was the pronunciation course itself, even though this was not one of the options that we had discussed beforehand. The students were asked to speak for a minimum of two minutes during each of these free-speaking tests, but only a sample, about a minute long, from the middle of each was evaluated. This was to avoid the first part of the recording in which the participants would be most aware of being recorded and would speak in a self-monitoring and controlled manner. Again, once they finished, they e-mailed their recordings to the investigator.

Syllabus. Parallel lesson plans and in-class materials were developed for both the songs and no-songs group. The lesson plans were developed for a 9-hour course spread out over 7 days, in which each class was 80 minutes in length. Nine hours over a period of two weeks was the maximum time that was available for the researcher to have with the participants. The content that was taught to the songs and no-songs groups were, as indicated above, the suprasegmental aspects of English outlined in Gilbert's Prosody Pyramid as well as other common features of

reduction and connected speech. Some of these phenomena are different from Spanish and can be troublesome to master for Spanish learners of English.

Word stress and sentence stress, for example, are different in the two languages and there are studies that indicate that they are important with respect to intelligibility (e.g Field, 2005; Hahn, 2004). Therefore, word stress and sentence stress are examples of suprasegmental features that needed to be covered. In English, the syllables in words can have three types of stress: stressed, unstressed, and reduced; in Spanish there are only two: stressed and unstressed. In the word *radical*, in English the first syllable is stressed, the second, unstressed, and the third, reduced, as can be seen by transcribing it: [' ɹæ dɪ kəl]. In the same word in Spanish, the first two syllables are unstressed while the third one is stressed: [ra di ' kal]. Note that in Spanish, neither one of the first two syllables are compromised in any way, such as, by having their vowels reduced to a schwa, for example. This means that vowel sounds in Spanish are clearly pronounced and, at least in Chilean Spanish, never reduced to a schwa. Therefore Spanish speakers tend to over articulate English vowels, or pronounce them too clearly, which results in them speaking English with a syllable-timed rhythm. Rhythm in English, however, is generally considered to be stress-timed, which means that the stressed syllables tend to occur at regular intervals. In order for this to happen, unstressed syllables have to be reduced, thus the fundamental importance of the schwa. With regard to sentence stress, we saw that there are focus words in English which we pronounce in such a way that they stand out to the listener. That is, we say focus words with a change of pitch and ensure that the vowel in the stressed syllable is long and clear and even perhaps louder (Gilbert, 2008). Spanish, on the other hand, will often use grammatical means to draw attention to words. For example, a subject pronoun (which is normally redundant in Spanish) can be used to emphasize or clarify to the listener who

the subject is in the utterance. To give an example, “Hace todo” for example, could mean “He does everything”, “She does everything”, “You (singular, formal) do everything” or “It does everything”. If it is not clear from the context, then the addition of the noun itself or the subject pronouns *él, ella, uno* etc. would be required. In order to emphasize, the subject pronoun is also used even when the subject is clear from the inflexion of the verb. For example, “Yo la quiero conocer” emphasizes that it is “I” and *not* someone else who may or may not have been presumed to be the subject. In this situation, an English speaker would make *I* the focus word of the utterance and would therefore use a change in pitch as well as making the vowel sound long and clear and possibly louder. Although English may at times also use grammatical means for focussing, while Spanish uses phonetic, as in the Chilean expression, “Pucha, qué laaaata!”, it is important for learners of English to perceive and produce focussing through pronunciation means as opposed to syntactical.

In addition to the suprasegmental features of the pyramid, common and salient connected speech phenomena were also taught in the lessons. Covering the linking of a word-final consonant to the vowel of the following word is important because the majority of Spanish words end in vowel sounds, whereas in English, most words end in consonants. In addition, it was necessary to touch upon the assimilation of sounds, which is when one sound will take on the characteristics of a neighboring sound, as in “don’t you” where the final “t” and initial “y” sounds merge resulting in an utterance that sounds like “donchew”. Covering specific auxiliary verbs such as, “want”, “got”, “going” etc. plus “to”, which result in “wanna”, “gotta”, “gonna” were considered particularly important not only because mastering them results in a much more natural-sounding – and thus easier to understand – speech, but also because many language learners believe that such assimilation is sloppy or so informal that they consider it to be

inappropriate slang. As such, without meaning to, Spanish language learners often tend to sound too formal or stilted when they speak English. Therefore, the decision to teach the above pronunciation areas was based primarily on a desire to help Spanish speakers be more intelligible, and also to raise their awareness of how English is spoken, so that they can begin to notice other pronunciation phenomena on their own and continue to improve their speaking skills.

In order to determine whether songs can help facilitate the learning of pronunciation, the two groups were presented with and practiced the same pronunciation phenomena, that is, the sequence and the content of the course were the same for both. However, the pronunciation training for the two groups differed with respect to the use of songs as a material: the songs group had songs as a material available to them for listening discrimination and production practice whereas the no-songs group did not. Therefore, instead of having a song that exemplified the pronunciation feature being taught, the no-songs group worked instead with non musical audio and/or written text. A course outline highlighting the content and materials used can be seen below in Table 3.1.

Table 3.1.

Content and Materials Used in English Pronunciation Course

DAY	PRONUNCIATION AREAS	SONGS	OTHER MATERIALS
1	Pre-course tests; introduction & raising awareness; thought groups	Cat Stevens – <i>Father and Son</i>	PPP; Steve Martin – Pink Panther clip; Catherine Tate – <i>The Interpreter</i> ; Phrases; sentences; radio advertisement
2	The focus word	Various Artists – <i>I Can See Clearly Now</i>	PPP; conversation; dialogue
3	Stress; peak; syllables and word stress; the schwa	K'naan – <i>Wavin' Flag</i>	PPP; poem, <i>If</i> – Rudyard Kipling; schwa video explanation; sentences; imitation video
4	Connected speech and linking 1 (-C+V-; /d/+j/=/dʒ/; contractions)	Eric Clapton – <i>Tears in Heaven</i>	PPP; tongue twister; Q&A sheet
5	Connected speech and linking 2 (/t/+j/=/tʃ/; auxiliary verbs + “to”)	Neil Young – <i>Old Man</i>	PPP; video; jazz chant; Steve Jobs video
6	Connected speech and linking 3 (-ing; sound & syllable elision)	Burton Cummings – <i>Break It To Them Gently</i>	PPP; <i>A Stolen Butterfinger</i> audio and text; worksheet
7	Sentence stress and rhythm	Parody of Gotye's – <i>Somebody That I Used To Know</i>	PPP; dialogue; photos

Methodology. The learners were taught the pronunciation features via an eclectic method as opposed to a purely communicative one. That is, they were introduced to the features by both inductive and deductive means and they also received overt explanations involving the International Phonetic Alphabet (IPA), diagrams, videos as well as other materials. As much as possible, they practiced perceiving and producing the features first in controlled, then in progressively freer contexts as is the case in CLT. However, not all of the activities were

communicative in nature. Therefore, although the lessons themselves generally followed the typical sequence of CLT methodology, the method utilized can be described as eclectic because there were activity types such as drills, chants, and mimicking, which are derived from other teaching methods. Thus, while phonetic perception and singing exercises, for example, were not communicative in nature, other activities such as question and answer tasks, short conversations, and talking about images were.

The reasons for using an eclectic method are many. Above all, as we saw earlier, there is no agreed upon way of teaching pronunciation using the communicative method. Furthermore, there is the challenge involved in providing enough effective practice of forms while still maintaining a communicative context, as indicated by Isaacs (2009). However, just because communicative language teaching is the most popular approach in general language teaching, does not mean that it is always the most effective approach with respect to pronunciation. Techniques that pertain to other methods are often necessary, especially if we take into consideration some of the pronunciation teaching principles that became apparent in the review of the literature in the previous chapter (see Table 2.1). For example, in order to help learners notice the sounds they need to produce, understand how their pronunciation differs from the model or target, discover the means to make the new sounds, understand why it is important to make them correctly, and practice repeatedly to develop the automaticity required to produce the sounds in natural speech, non-communicative tasks will sometimes be necessary. If we look at our lessons in a little more detail, the eclectic nature of the approach can be appreciated and the purpose of the activities understood. In the description of the lesson plans and classes, where it may not be entirely obvious that a procedure or activity follows CLT or the guiding principles of pronunciation teaching, further explanation is provided.

Designing the course. When designing the course, I first decided on what the teaching sequence would be of the target pronunciation areas. First, I followed the sequence of how Gilbert (2008) presents the pronunciation features of the Prosody Pyramid and then I covered the easier and more common elements of connected speech before the more difficult ones. Next, I set about to find the most suitable song for each lesson. That is, I looked for a song that would exemplify and serve as a model for the pronunciation elements being taught and have a number of instances of those elements.

When choosing a song for teaching pronunciation, there are a number of other factors which needed to also be taken into consideration. First of all, it was important that the song be grammatically and lexically appropriate for the L2 proficiency level of the students. That is, I wanted them to be able to focus on the pronunciation elements as opposed to unfamiliar grammar or words. I was also careful to make sure that the songs had appropriate content, i.e. I was careful not to choose songs that involved nudity, drugs, or violence. In addition, the lyrics needed to be clearly sung and not overwhelmed by accompanying musical instruments and the accent could not be identifiably different from a Canadian or General American one. I also tried to choose songs that were catchy and possibly earworm material. After all, the more a learner can be exposed to the song, the better, even if it is sub-vocally, because this still constitutes input. I chose songs that were classic or iconic in Canadian or American music. Although it was not possible to use songs that the students themselves had indicated in the *Initial Pronunciation Questionnaire*, due to lack of time, I did choose a parody of a popular song at the time, *Somebody that I Used to Know* (by Gotye), which the students were familiar with and quite liked. The advantage of choosing a parody of a popular song at the time is that the students are likely to immediately have a positive attitude towards it and be open to trying to understand the

humour in the unfamiliar lyrics. Some authors recommend using songs that students have indicated a preference for because this provides them with an opportunity to have some ownership of the lesson. While I do not disagree with this decision, it is also prudent to choose songs that may be unfamiliar to the students. In this way they will be exposed to more of the target culture as well as additional musical genres of that culture. Besides, homework assignments can always be assigned that ask students to listen to their favourite songs, notice, and report back on the pronunciation features that the song has.

Once the songs were chosen, I looked for and developed additional materials that provided practice for the same pronunciation phenomena taught with the song, and, when possible, also related or adhered to the theme in the song. When structuring the classes I was careful to try to follow a general CLT sequence of warm up + presentation + controlled practice + free practice + freer practice with plenty of review and recycling of content taught. As mentioned above, there were a number of activities employed that are from other approaches.

Lesson plans and classes. On the first day⁵⁴, the goals were to raise the students' awareness of how languages differ with respect to their sounds and delivery, to indicate what causes foreign accents, to introduce the Prosody Pyramid, and to teach thought groups. No particular theme was followed because of the nature of the first day of class. That is, it was also necessary to create and establish a relaxed atmosphere in the classroom, and for that reason plenty of humour was used. Therefore, the class began with a humorous video clip of Steve Martin's *The Pink Panther* indicating how the class would *not* be. The clip was the one in which Steve Martin (Inspector Clouseau) was receiving pronunciation training and was learning how to

⁵⁴ All print materials for the classes may be found in Appendix C.

properly say, “I want to buy a hamburger”. This was followed by another video, in which Catherine Tate, a British comedienne, pretends to be an interpreter, yet has no idea how to speak the seven different languages of the people she is faced with. In this video, she simply pretends to speak the languages by nonsensically using stereotypical sounds and intonation. After this, the students are placed in small groups and asked to pretend to speak American English like Catherine Tate did for the other languages. After brainstorming with the students on the sounds that they were making, they were asked to mimic an American speaking Spanish. I provided them with an example to guide them: “*Joe nou keyero nada. Eso es toe doe*”. Then I asked them to take note of what they were doing with their articulatory organs while they were engaging in the activity. The reason for these last two activities was to help the students in developing an awareness of sounds and rhythms. By pretending to speak like an American and taking note of how they were speaking, the students’ attention was drawn to the “errant” English sounds and sound patterns that they were using while speaking Spanish.

Once the students were aware of some of the sounds and sound patterns of North American English, they were then shown a PowerPoint Presentation explaining the nature of foreign accents (see Appendix C). This helped them not only understand why English speakers sound the way they do when they speak Spanish, but also to encourage them to use those sounds and sound patterns themselves when they speak English. While not a pronunciation teaching principle per se, many learners appreciate explanations that help them to understand the nature of a phenomenon. The next activity was a Quality Repetition⁵⁵ one in which the students chorally repeated one of the phrases in the Catherine Tate video, namely, “I can do *that*”. This phrase has the learners unknowingly practice focus words, and the schwa in the modal auxiliary verb “can”.

⁵⁵ Quality Repetition is a pronunciation drilling technique that helps students learn chunks of language complete with all of the suprasegmental features. This technique was developed by Ollie Kjellin (1999).

This activity provided contextual listening and speaking practice of particular pronunciation forms as a way of helping them to develop automaticity. Then the students were taught the terminology that they would need for the course and given an overview of the Prosody Pyramid and the first item in it, the thought group (see Appendix C). Up until here, the execution of the lesson was identical for both the songs group and the no-songs group.⁵⁶

At this point, the songs group were provided with the lyrics and listened to a song, Cat Stevens' "Father and son", in which they had to separate the thought groups by putting slashes between them (e.g. *It's not time / to make a change, / Just relax, / take it easy. /*). The reason this particular song was chosen was because the lyrics are simple and very clearly sung. (When initially introducing songs to students, it is very important to choose songs that will not intimidate them, thereby negatively affecting their confidence.) After taking up the answers, the students took turns reading the song to each other, making sure to respect the thought group boundaries. While the students were reading the lyrics to each other, I circulated around the room listening to the students and silently assessing their use of thought groups. It was apparent that, in general, they had no trouble with thought groups, with the exception of two students who needed to be encouraged to make more of a pause between thought groups.

The no-songs group, on the other hand, listened to phrases and then an advertisement and also had to show the separation of the thought groups on their handouts. As with the songs group, we took up the answers and then the students took turns reading the phrases and advertisement to each other, making sure to respect the thought group boundaries. Again, I listened to them practice and was assessing how they were doing. Two students in this class

⁵⁶ This was the point in the Day 1 lesson plan where the class ended and the rest of the activities were actually covered in Day 2.

naturally speak quite quickly, and although their use of thought groups was detectable, they were encouraged to slightly extend the pauses between thought groups. Next there was a brainstorming activity in which the students of both groups needed to list things they could do inside and outside of class in order to improve their pronunciation. The reason for having the students do this was so that they could take some ownership in the process of their learning by identifying opportunities for practice. Finally, their homework was to watch an American TV show or movie and note down at least five different things that Americans do when they speak. That is, do they use their hands a lot? How far from each other do they stand? Do they make faces? One of the purposes of this activity was to help the students appreciate that speaking involves our whole body and is part of a communicative activity with another person. Another purpose was to provide the students with some extended, meaningful contact with authentic English.

On the second day of class⁵⁷, the underlying theme was clarity and clarifying. In this lesson, the goals were to review thought groups and teach focus words.

In this class and in all that follow, pronunciation areas that were previously taught were reviewed and recycled through additional practice. Having the students engage in review activities allowed me to regularly assess their learning. When taking up the listening review activities, it was possible to determine whether the learners were hearing the pronunciation areas taught according to the answers they provided. Similarly, when the students were engaging in speaking review activities, their use of speech allowed me to hear whether they were correctly producing the pronunciation areas. When the students demonstrated that they had perceived the pronunciation features and/or when they correctly produced the features, they were provided

⁵⁷ All print materials for the Day 2 class may be found in Appendix C.

with positive feedback through praise. When they were not, they were made aware of it through negative feedback in the form of correction. Explanation was also provided when necessary as a way to help the students understand and recognise their errors. Regular assessment, though, also occurred whenever students were engaging in tasks.

At the beginning of the second class, the students started out by briefly discussing the homework from the day before. Thought groups were then reviewed, and then the students performed a controlled speaking activity in which they had to give their phone numbers to the people around them, using thought groups. This quick little task required the learners practice using thought groups in a realistic and useful context. The next activity had the songs group listen to the song “I can see clearly now” and separate the thought groups, just like the day before. After they did this, I read out the lyrics and had the students check the thought group boundaries. While reading the lyrics, I said “it’s going to be” as “it’s gonna be” as an inductive way of introducing connected speech. Then I had the students compare their answers and then sing along to the song as it was played again. Singing, as we have seen in the review of the literature, is an enjoyable activity and one that practices pronunciation forms in a meaningful context. The no-songs group, on the other hand, were given a handout with a conversation on it. Then they listened to an audio of the conversation and were asked to separate it into thought groups. After they did this, I read out the conversation and had the students check the thought group boundaries. While reading it, as with the songs group, I read the conversation so as to inductively introduce connected speech. Then the students formed groups and practiced the dialogue with each other. After this, both groups of students had a free speaking activity in which they were to use thought groups while asking each other questions and talking about what they did the night before. With both groups, while they were practicing the previous activities, I

circulated around the room and listened to how they were using thought groups. I noticed that with both groups, the students were using thought groups more accurately, in that when they made the pauses, they appeared to make them more deliberately. Once this activity was done, we moved onto focus words and to introduce them, I asked the students to complete a short, three-question, true/false survey on pronunciation and clarity of words and syllables. The survey is in Figure 3.2.

Figure 3.2 Pronunciation Survey

<i>Read and decide whether the following statements are true or false.</i> 1. _____ It is good to pronounce each syllable in a word clearly so that people will understand you better. 2. _____ It is good to pronounce each word clearly so that people will understand you better. 3. _____ It is more correct to say things like “I will” or “He is” instead of the contracted forms “I’ll” and “He’s”.
--

The questions have seemingly obvious answers⁵⁸ for English pronunciation teachers, but students almost always get them wrong, because for those without an awareness or understanding of English pronunciation conventions, it is counterintuitive to think that it is not good to pronounce every sound, syllable, and word clearly. Pointing out the students’ misconceptions about English pronunciation helps raise their awareness and prepare them for new information.

⁵⁸ The answers are all *false*.

The survey led into explaining focus words and illustrating the idea through drawings so that the students could understand not only why their answers were wrong but also the importance of contrast in the English stress system, namely, that of focus versus reduction. I thought that presenting the idea of contrast in a visual way would help some learners comprehend the idea better than by just an oral explanation. The next activity was a quick Quality Repetition exercise in which the words were partially taken from the song, “I can see clearly now” (see Appendix C). Having the students repeat, “It’s gonna be a *cold* day”, required them to practice not only focus words and the reduction of “going to”, but also the linking of two vowels with a [j] and the lengthening and continuation of a word-final consonant before a word beginning with the same consonant sound. This way the students were practicing features taught in the course as well as others not explicitly taught due to lack of time. As mentioned in Day 1, quality repetition exercises provide contextual listening and speaking practice of particular pronunciation forms as a way of helping learners develop automaticity.

After this, the songs group identified focus words while listening to a different version of the song in which many of the words were not sung clearly (only the focus words were). The reason for doing this was to draw attention to the fact that some words have to be de-emphasized so that others can stand out. Then they were asked whether they could see a pattern to the focus words, (which tend to be the last word in the thought group). The no-songs group, were presented with a printed dialogue which they listened to and had to identify the focus words and see if they noticed a pattern to them (also the last word in a thought group). The last exercise was a controlled one in which the students in the songs group sang along to still another version of the song, which is the most famous and catchy. The reason for choosing to end the class with this version of the song was to perhaps activate the SSIMH phenomena. As I walked around

and assessed the participants' use of focus words, I noticed that they used them well, and it seemed that using focus words when singing came very naturally to them. The students in the other group took turns practicing the dialogue and taking on different roles in it. They were to pay particular attention to thought groups and focus words. While assessing their use of focus words, I also noticed that the no-songs group produced them quite well. However, the fact that the focus words were highlighted in the dialogues (as they were in the song lyrics for the songs group) most probably had a positive influence on their production. For homework, the songs group was to listen for and practice using focus words in songs and speech, while the no-songs group was to listen for them in speech.

For day three of the course⁵⁹, the underlying theme of the class was freedom and struggle and as much as possible the materials reflected this. The goals of the lesson were to teach the last two parts of the Prosody Pyramid, stress and peak, as well as to introduce the schwa. Both classes started with a warm-up and review of thought groups and focus words involving a pair-work practice activity which required the students to use both thought groups and focus words. Reviewing and recycling information is a necessary technique in language teaching, as is regular practice of the language areas that have been taught. For this task, (see Appendix C for the full exercise⁶⁰) student A had to choose between two focus word options and make a statement for which student B, their partner, had to perceive which focus word was said and then correct the information that the first person had stated, for example:

Student A

Student B

a. It's a big DOG.

No, / it's a WOLF.

⁵⁹ All print materials for the Day 3 class may be found in Appendix C.

⁶⁰ From Gilbert (2005a), p. 64.

b. It's a BIG dog.

No, / it's really more **MEDIUM**-sized.

Being able to perceive and produce a change in focus word, is very important for understanding when information has been clarified or corrected and for clarifying, correcting, and disagreeing with an interlocutor. This particular activity served to help the students realize the importance of focus words for correcting or disagreeing with people. It is worth mentioning that activities such as this should ideally be practiced in two ways: first with students looking at each other as would be the case in face-to-face contact, and second, with the students back to back. In the face-to-face contact, the students should be encouraged to also use the accompanying facial movements or gestures upon articulation of the focus words (i.e. raising eyebrows, moving the head etc.). In the back-to-back practice context, the learners can only rely on what they hear, which forces the speaker to clearly employ pitch change and vowel clarity. This back-to-back practice better prepares learners for telephone conversations.

Students were then shown slides which presented the concepts of stress and peak both via explanations and images along with quotations in which the students as a group had to identify the stressed syllables in the italicized (focus) words, for example, “When we lose the right to be *different*, we lose the *privilege* to be *free*”⁶¹. Then they were presented with words that had the stressed syllables capitalized and the vowels of the reduced ones crossed out, for example, “DIFFerent, PRIVilege, FREE”. For these they practiced listening to and repeating the words. Indicating the difference in pronunciation of stressed and reduced syllables in a visual way, can help learners first realize and then notice that there are in fact differences in the way these sounds and syllables are produced and thus heard.

⁶¹ Quote by Charles Evans Hughes

The next activity had the students practice both listening for the stressed syllable of focus words and then singing or saying them. The songs group were provided with the lyrics to “Wavin’ flag” by K’naan in which the focus words were underlined in the first half of the song (e.g. *When I get older, I will be stronger. They’ll call me freedom just like a wavin’ flag*). In the first half of the song, they had to circle the stressed syllable, while in the second half, they were to sing along to the song and raise their hands every time they sang a focus word. Involving bodily movements is another technique used to reinforce language points and it is particularly relevant here because native speakers will, as mentioned earlier, often use facial, head or hand movements when saying a focus word. The no-songs group instead listened to the poem, “If”, by Rudyard Kipling and circled the peak syllable of the underlined focus words before reading the poem to one another. In this case, the students were encouraged to use gestures (i.e. facial, head, or hand movements) while reading the poem. While observing the participants perform this exercise, it was quite obvious that the members of both groups had no problems with focus words. Although most of the students in the songs group only raised one hand (and not very high), it was still clear that the hand movements and use of focus coincided, as it did with the no-songs group and the gestures they used.

For teaching the schwa, a YouTube video by Osorio (2012) was used in which the schwa was explained in Spanish to the students. In addition, there were also a couple of PowerPoint slides which provided examples of when and where it appears. The reason for presenting the schwa in this way was to help ensure that the students fully understood what the schwa is and why it is necessary to use it. Since this allophone does not exist in Chilean Spanish and is so different from any Spanish vowels, yet is the most common vowel sound in English, providing a clear and thorough explanation of it in their own language was necessary in order to ensure that

the students grasped the concept. In line with the principle that learners sometimes need help in recognising that pronunciation is important, they also sometimes require assistance in recognising why certain aspects of pronunciation are important. After the presentation of the schwa, both groups received handouts and listened to an audio of half a telephone conversation. The handout⁶² indicated two possible words and their task was to circle the one they heard, for example, “I can / can’t take you”, “I’ll ask for / four volunteers to help clean up”. Note that one of the options in each case is a word in which there is a schwa. After this, the students practiced reading both options. Finally, the homework involved the students doing an imitation exercise in which they were to view segments of a movie clip on YouTube called “Peggy’s new office” which has short parts of the conversations repeated three times followed by a span of time in which the students had to repeat what they heard. This imitation exercise is similar to Quality Repetition, but it does not have them repeat the same phrase as many times. Instead they are forced to repeat mini conversations in context and are not given time to actually think about exactly *what* is being said; they only have time to repeat *how* it is being said. As in Quality Repetition, the students practice all suprasegmental aspects. Exercises that require the students to practice the various suprasegmental aspects together are necessary because that is what natural speech is like. Of course the various features can be presented and practiced separately to a certain degree, but only with practice that integrates the features, can students improve their accuracy and fluency.

In the lesson for day four⁶³, the theme of the class was difficult life events and decisions. The goals were to review and cover remaining information pertaining to the schwa, syllables and word stress, and to teach three additional connected speech phenomena. Unfortunately, most of

⁶² From Grant (2007), pp. 89-90.

⁶³ All print materials for the Day 4 class may be found in Appendix C.

the students had not done the homework assigned from the class before, and so we took class time at the beginning to do that activity. In retrospect, I conclude that this kind of activity *is* best done a few times in class before asking the students to do it on their own. The speed of the speaking in the activity and the requirement that the students repeat what they hear without thinking about what is said can be challenging and unnerving for learners at first. As it was the first time that the students had done this kind of exercise, it was clear that it was an intimidating one. They were reticent about repeating what was said because they did not understand it. When faced with a situation like this, it is important for pronunciation teachers to provide encouragement and refrain from error correction until the students' confidence returns. After that, the tongue twister, "A tutor who tooted the flute..." was used to review material covered from the day before. This also involved teaching the pronunciation of the English /t/ which is alveolar and thus different from the Spanish one, which is dental. It was necessary to teach this and indicate that the English /d/ has the same point of articulation, i.e. alveolar, in order for the students to be able to successfully learn the assimilation of the word-final /t/ and /d/ with word-initial /j/ or "y". The students first listened as I said the tongue twister and indicated the rhythm by tapping on the desk and then they had to try to repeat it exactly as I had said it. After that, I showed an image depicting the situation in the tongue twister (a tutor with a flute teaching two tooters to toot) and showed the written version of the tongue twister so that the students could understand what it meant. Finally, they were asked to identify the words and syllables with the schwa in the tongue twister. Employing the zoom principle is important especially when students need to understand how something as small as a sound (i.e. the schwa) can be meaningful in the context of an utterance, or in this case, a tongue twister.

A review and presentation in the form of slides was then given on the schwa, syllables and word stress. This included an explanation of function words and various examples of how they are reduced, e.g. the word *and* is often pronounced [ən] or [n]. (See Appendix C for the information presented in PowerPoint). The activity after this was to be one which practiced stressed syllables but time was running short and so I decided to skip this part of the lesson. Instead, the songs group first listened to Eric Clapton's "Tears in heaven" without the lyrics and were asked to try to understand what they could. Then, they were provided with the lyrics in which many of the pronunciation areas that we saw in class were blanked out, as well as all of the cases of "would you", which was one of the new areas being taught on this day (e.g. _____ *know my name* _____ *saw you in Heaven* _____ *be the same* _____ *saw you in Heaven*). The students' task was to fill in the blanks as much as possible. They were reminded that English does not always sound the way it is spelled and they were told that some of the blanks contained reduced function words. After playing the song two to three times to allow the students to fill in as many blanks as possible, the answers were reviewed with the help of a PowerPoint. While taking up the answers, the students were asked what they noticed about how the singer said certain words and phrases. The point of asking students what they noticed about the pronunciation was to help them realize the importance of noticing as well as help to boost their confidence in their ability to do so. This led into the PowerPoint presentation of the connected speech and linking pronunciation points that were part of the day's lesson.

The no-songs group, on the other hand, listened to me read a series of hypothetical situations in which they had two choices, for example, "Would you or your partner rather have kids or would you rather adopt or live without children?". On the first listening they were told to

just try to understand what they could. Then they were provided with the handouts in which a number of words and phrases were blanked out, that is, pronunciation areas that we saw in class as well as all of the cases of “would you”, and they were told the same as the songs group above, namely, that English does not always sound the way it is spelled and that some of the blanks contained reduced function words. Similarly, they heard the situations two to three times and were asked to fill in the blanks as much as possible. When taking up the answers, the procedure was the same as above. Finally after the areas in the PowerPoint were covered, the songs group listened and sang along to the song while the no-songs group, in pairs, practiced posing the situations to one another. This activity provided for controlled practice of the elements, a technique to strengthen articulatory skills, while the last activity was a fluency practice. In this activity, the songs group got together in groups of four to discuss the song and ask each other hypothetical questions related to it as well as answer, while the no-songs group asked and answered some of the questions in the handout as well as new ones that they had invented. Both of these activities required the students to practice the connected speech items that were covered in the slides but in a meaningful communicative context.

The goals of the fifth class⁶⁴ were to review and recycle forms covered in Connected Speech and Linking 1 (see Appendix C) and to teach additional forms of connected speech. The underlying theme of the lesson was learning and success. The class started out with another imitation exercise as a warm up. Again, it was from YouTube, and the movie clip was called “People change”. Even though this imitation exercise was only the second one for the students, they were much more relaxed and adept at it than they were for the first one. The streams of speech they were producing were very much like the original. After that, the students worked on

⁶⁴ All print materials for the Day 5 class may be found in Appendix C.

a handout that dealt with linking consonants to vowels and the phrases that were used in the handout were all commands. They had to practice saying the commands, come up with their own, and then practice giving orders to their partner. The purpose of this activity was to help the students realize that consonant to vowel linking is applied not only in speech in general but also for very particular types of speech acts. The next activity that the students did was a rap. For the songs group I played the video, and just had them listen first and try to understand. The video, available on YouTube, was a rap called “Whatcha gonna do?” that a retired teacher, named Burton Crane, wrote and performed on American Idol. Afterwards we discussed it briefly and I made sure that they understood what *whatcha gonna do?* means. Then I distributed the lyrics and briefly covered the slang before playing the song again to which the students were to sing along while mimicking the singer.

The activity for the no-songs group was the same except instead of playing the video, I read out the rap to them while tapping the desk to show the rhythm. I also mentioned to the students that the rap was written and performed by a retired teacher. Then, the students read it to each other. It was funny to see how interested the students were in the slang that appeared in the lyrics and how much fun both groups had practicing the rap. I think the fact that it was performed by a man in his 70s helped make it light hearted and accessible for the students so that they themselves could let loose and engage in the activity. This particular text and activity was very appropriate to use because it worked as a lead-in to the content to be taught in Connected Speech and Linking 2, (see Appendix C) which dealt with word-final /t/ plus /y/ and “going to”. I briefly covered this with a slide and discussed why it is that students sometimes do not end up learning authentic-sounding speech, which can sometimes be due to teacher talk and (artificial) textbook audios. It was also important to discuss how there are different levels of formality and

that for the auxiliary verbs plus “to”, the changes do not normally occur in formal speaking situations. The purpose of going into such depth regarding formality and the reduction of forms is because students can be quite stubborn in their language beliefs, especially if former teachers have taught them something different. Many Chilean students, as we saw in Véliz Campos (2011) have been influenced by the attitude that British English is formal and an ideal linguistic model. That is, they may have been presented with a simplified version of the British standard without an informal level, a level which is salient in the North American variety that they would regularly be exposed to in movies, TV shows, and songs.

The next activity demonstrated the reduction of some of the auxiliary verbs through a jazz chant. For this both groups of students were provided with copies of the chant and instructed to listen first and snap their fingers to the rhythm, paying particular attention to the pronunciation of the reductions. After this they were given the time to practice it in pairs. The activity after that involved listening for much of the connected speech phenomena that had been seen up to that point in the course. Again, reviewing and recycling the forms taught with practice activities helps learners practice and remember to integrate the different forms taught and thereby to monitor their own pronunciation and begin to self correct.

For the songs group, Neil Young’s “Old man” was played and the students were asked to just try to understand what was sung. Then they were given the handout in which more than half of the lyrics were blanked out (e.g. _____ *my life*, _____ *were*). It was important at this point to tell them not to panic, and that after listening to the song a few times, they would be able to fill in most if not all of the blanks. In my experience, I have found that students working with songs progress quickly with their ability to perceive pronunciation phenomena taught, however they themselves are not aware

of this progression. It is confidence building for them to realize that they can, after listening to a song only two or three times, distinguish the missing lyrics. When the students do miss hearing some of the lyrics and how they were sung, providing the answer and then immediately playing that part of the song will enable the learner to hear what they did not hear previously. After playing the song for the students, they were put into groups of four to compare their answers and discuss the pronunciation. Then the PowerPoint slide with the full lyrics was shown for them to check their answers.

The no-songs group listened to a YouTube video clip of the late Steve Jobs, (former) Apple CEO discussing rules for success. Again, they did not have the script for the first listening. For the second and successive listenings, they were given the text with words and phrases blanked out and instructed to fill in as much as they could. Before revealing all of the missing words, the students worked in groups and compared their answers and discussed how these words and phrases were pronounced. As they did this, I circulated and listened to what the students were saying, and provided both positive and corrective feedback accordingly. Generally the group was quite accurate with their answers and their pronunciation although one member did struggle more than the others. Using Steve Jobs as an example speaker, I felt, would add some credibility to my earlier discussion of the acceptability of using reduced forms in speaking. For homework, the students were asked to go to YouTube and find videos similar to the last one seen in class. They were to listen a number of times, try to understand what they heard and then repeat or mimic what they could. In this way, their exposure to authentic speech would be increased while at the same time providing them with the opportunity to strengthen their auditory memory by mimicking chunks of meaningful speech. For both the song and the video of Steve Jobs, my intentions were to have the students read the song or the script in class and record

themselves doing it and then check their own pronunciation. Unfortunately, however, we ran out of time. Giving students the opportunity to self assess their pronunciation can provide them with insight as to how they really sound compared to a model speaker.

In the lesson for day six,⁶⁵ the underlying theme was crime and the goals were to review and recycle the connected speech and linking forms from the day before and to teach new ones. Before we started, I asked the students if they were having any problems noticing or using the connected speech phenomena that we had seen so far in the course. In addition, I asked them if there was anything they had noticed that we had not yet seen in class. The students said that they understood everything that had been covered. As for noticing pronunciation phenomena, there were some students who mentioned hearing additional kinds of assimilation and linking, such as the lengthening of [s] when it occurs in word-final + word-initial position, e.g. “the computers s seem to be...”. The review and warm up had students giving each other advice through the use of affirmative and negative commands. They were given a handout and were instructed to fill in the blanks with a given list of verbs and nouns and then practice taking turns and giving this advice to a partner. Having them practice linking with a controlled activity like this enabled them to remember to focus on their pronunciation as opposed to strictly the content of their message. As was the case the day before, the students produced the linked forms very well. After this, new material was presented, namely, the pronunciation of “-ing” and the elision of sounds and syllables, such as those in consonant clusters in words and phrases such as *kindness*, *textbook*, and *next month*, as well as the elision of vowels, in words such as *Catholic*, *chocolate*, and *family*. (see Appendix C). To practice these, the students were given a handout containing common words with omitted syllables along with the words shown in a phrase or collocation.

⁶⁵ All print materials for the Day 6 class may be found in Appendix C.

They first practiced repeating how I pronounced the items and then were given the opportunity to make up and practice an impromptu dialogue using phrases in the handout. In this way they practiced the form and then used it in a communicative context. It was interesting to notice that the students readily accepted the elision of consonant sounds in consonant clusters and easily applied the changes, whereas this was not the case with the elision of vowel sounds. This made perfect sense because, phonologically, Spanish does not have the complex consonant clusters that English does and so the students would find it easy to drop consonant sounds in such a cluster. In contrast, in the case of vowels, since some of the words in question are cognates (i.e. Catholic/católica, chocolate/chocolate, family/familia) it would be strange for them to elide the vowel sound, especially if it is a tonic vowel in Spanish. The next activity involved hearing and practicing much of the connected speech we had seen in the course.

For the songs group, the activity with the song, Burton Cumming's "Break it to them gently", followed the same procedure as the day before, only it was slightly more challenging. The difference was that instead of each blank indicating a word, there were blanks that were joined together so that the students did not know how many words each blank contained (e.g., *Break it to them gently when you tell _____ that I _____ home again*). This forced them to rely solely on what they heard without any clues from the handout. After giving the students a chance to compare their answers in a group and discuss the pronunciation of the words in the blanks, I showed them the correct version of the lyrics. Again, time was running out, and we were unable to complete the last activity which required the students to change the lyrics in the song and then practice speaking them into the computer as they recorded themselves.

The corresponding activity for the no-songs group involved listening to an audio of a person telling a story. First they heard the story and then they were given the script with two or more blanks combined together, just like the procedure in the songs group. After listening to it two to three times, the students compared their answers and discussed the pronunciation before they were shown the correct answers. Finally, for this group as well, it was necessary to forego the last part of the exercise which would have had them change the ending to the story and then record themselves telling the story.

On day seven⁶⁶, the goals were to review the different areas taught and complete the final listening and speaking tests. The underlying theme was life and learning. As a warm up, the students were asked to call out the names of different movies. As I wrote them on the board the students were to say them paying attention to the linking between the words. For the next activity, I asked the students to list the different areas of pronunciation that were taught in the course and had different volunteers explain each one and provide an example. Then we reviewed sentence stress and rhythm. The students were given a handout (see Appendix C) which showed a chart of different rhythmic patterns according to where the stressed syllable (peak) was located in the word or phrase. The students practiced the pronunciation of the different rhythmic patterns as a controlled exercise and then were asked to create and practice their own phrases for each of the patterns. For the second chart in the handout, there were five sentences which mainly differed by the addition of progressively more auxiliary verbs. Each student or group of students was given a number from one to five. Those with number one said their sentence over and over again, which was “CATS CHASE MICE”⁶⁷. While they were saying it, the number two’s joined in saying their sentence (“The CATS have CHASED MICE”).

⁶⁶ All print materials for the Day 7 class may be found in Appendix C.

⁶⁷ From Celce-Murcia et al. (2010), p. 210.

The students with the other numbers successively followed suit until everyone was saying their own sentence, but following the same rhythmic pattern. The other sentences were, “The CATS will CHASE the MICE”, “The CATS have been CHASing the MICE”, and “The CATS could have been CHASing the MICE”. Although each sentence is successively longer, the same amount of time is taken to say the sentence. The purpose of this activity was to review the importance of reduction in order to maintain the rhythm of English.

After this activity on rhythm, content words and function words were briefly reviewed. This led into the students practicing changing focus words for disagreeing and correcting one another. For this they had a sheet with mini dialogues in which person A made a statement and person B corrected or disagreed with it, for example: Speaker A: “I buy books at the **L**ibrary”, Speaker B: “No, you **B**orrow books at the library”⁶⁸. Then, the students were asked to make up their own statements and their partners would correct or disagree with what they had said. This controlled-to-free practice activity helped enable the students to appreciate the importance of applying stress to the proper focus word in speaking situations where there are mistakes or disagreements. By having the students make up their own erroneous sentences, which required correction from their partners, meant that the students had to correctly hear and produce the focus words in order to maintain the proper exchange of information and avoid miscommunication. All in all, the students had no problems applying the focussing techniques on the right words.

The next activity was a fluency practice one in which we pretended we were having a party and everyone had to say what they were going to bring. I started out by saying, “We’re having a party and I’m bringing a bottle of wine”. Then, the first student had to say that we were

⁶⁸ From Gilbert (2005a), p. 63.

having a party and that I was bringing a bottle of wine and that he or she was bringing something else. The following student had to repeat what both of us were bringing and so it continued with the rest of the students. For this activity the students were instructed beforehand to pay attention to thought groups, focus words, stress, peak and connected speech and linking. The students had a lot of fun with this activity, and in addition to sometimes being creative with their answers, for the most part, they were able to accurately produce these different areas of suprasegmental pronunciation together.

The final activity was also a fluency practice one and for this the students were shown different kinds of photographs (see Appendix C) and instructed to just talk about them, that is, what they saw, what they thought the person (or animal) was thinking, how it would feel to be in that position etc. Basically, they were just asked to talk about whatever came to their mind when they saw the images. Of course, beforehand, they were instructed to pay attention to the different pronunciation features that had been covered in class. To help make this a more relaxed activity and to show my appreciation to the students for taking part in the investigation, I passed out maple syrup cookies for them to enjoy while they talked about the photographs. Finally, I thanked them for participating in the course, and I finished the class off by playing a parody of Gotye's song, "Somebody that I used to know" – "Some study that I used to know", by College Humor. I gave the students the lyrics, played the video and encouraged them to sing along. I did this for both the songs and the no-songs group. Gotye's song was very popular at the time, and so the students had fun with the parody.

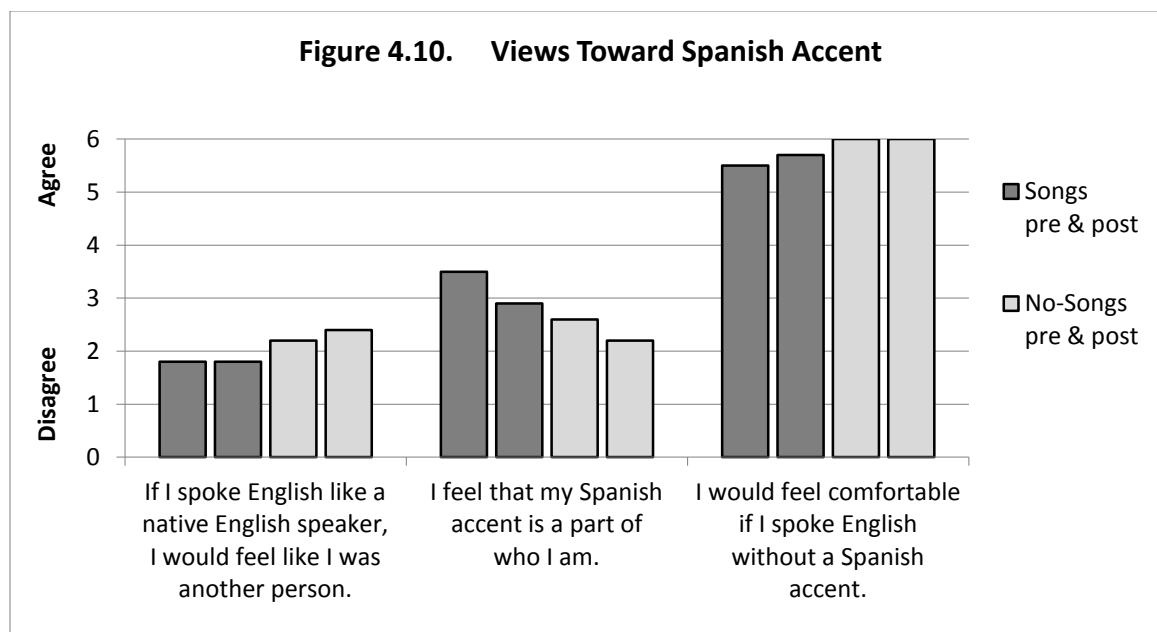
As mentioned earlier in this chapter, after completing the course, all of the participants completed the post-course listening and speaking tests, which were identical to the pre-course

ones. The songs and no-songs groups, however, also completed the Final English Pronunciation Questionnaire.

Data Analysis Techniques

After the initial and final questionnaires and listening and speaking tests were completed by the students, they were analysed. The data were recorded and compared, both pre- and post-course, among the groups. The results are presented and discussed, and the research questions answered in Chapter Four.

Questionnaires. All of the responses in the surveys were put into an Excel document, with the questions and statements along the top horizontal axis and the participants, separated into groups, along the left vertical axis, with room allotted for the answers to both the initial and final questionnaires. With regard to the statements, they were developed in order to know how the songs and no-songs groups felt before and after the course with regard to matters such as pronunciation and accent, for example. For these I drew up charts in order to illustrate the responses that the participants gave. The charts show on average how each group responded to particular statements based on a six-point scale ranging from *strongly disagree* to *strongly agree* (i.e. 1=*strongly disagree*, 2=*disagree*, 3=*slightly disagree*, 4=*slightly agree*, 5=*agree*, 6=*strongly agree*). The following figure, which is explained in the next chapter, is an example of one of the charts.



In the final questionnaire, there was also a space provided to allow the participants to make any comments they wished. This qualitative data was included in the discussion of the results of the listening and speaking tests as a way of incorporating the participants' views of the course and how they felt it helped them.

Listening tests. First of all, for Listening Test # 1, I marked the initial and final tests for all groups and then put the total into an Excel document out of a score of ten. In addition, how each participant scored pre- and post-course on Listening Test # 1 in thought groups, focus words, word stress, and contractions / reductions was also entered into the spreadsheet. Each section was graded out of ten, with only one permitting the allocation of half points because two words were involved in achieving a correct answer in this section (contractions / reductions). Marking the tests was a straightforward task and there were no complications involved. That is, answers were either correct or not.

For Listening Test # 2, which dealt only with linking and connected speech, an Excel spreadsheet was set up with the answer key on the horizontal axis along with the allotment of points for each sentence. As with Listening Test # 1, the total score as well as each pronunciation area were out of ten. On the vertical axis, the participants were indicated according to their code (e.g. C1, S5, NS2). The pre- and post-course results for each student were grouped vertically. While marking the tests, if the answer was fully correct, the (full) points were simply indicated in the cell. If there were errors, then what the student wrote was typed into the corresponding cell and the points allotted were indicated. The participants received a point if they got the element of connected speech correct. Normally this coincided with writing the correct phrase, but not always. For example a participant wrote “Yes, sure” instead of “Yes, sir”. In cases such as this, in which the answer provided was meaningful and included an element of connected speech, the participant was awarded the point. The mistakes that were evaluated included only the areas of connected speech and linking that were covered in the course. Finally, if a participant did not write anything, or in the case of C2 and C5 who did not take the test, the cell was left blank and zero points were awarded. The following table is a snapshot of how the results were noted in the Excel document.

Table 3.2.

Listening Test # 2 Marking Example

		d + y and t+ y and elision of sound				s + s	
		9				10	
	PH	Could you make sure that you don't speed?	dz	ty	elision	Yes sir!	link
		3				1	
C	1	3	1	1	1	... - 0	0
		3	1	1	1	1	1
C	2		0	0	0		0
			0	0	0		0
C	3	3	1	1	1	1	1
		Can you make sure that you don't speed? - 2	0	1	1	1	1
C	4	Could you make sure you don't speak - 2	1	0	1	1	1
		Could you make sure you don't speak - 2	1	0	1	1	1
C	5		0	0	0		0
			0	0	0		0

Dörnyei (2007) indicates that descriptive statistics are helpful for describing what learners have attained and so to better understand how the participants performed from pre- to post-course, descriptive statistics were run. In order to do this, SPSS was the program used to calculate the statistical results for this study.⁶⁹ Therefore, all of the results for both listening tests were transferred to an SPSS spreadsheet.

Firstly, descriptive statistics tests were run on the pre- and post-course scores, initially with the participants as a whole and then separated or “split” by group. These tests generated a variety of information, but for the purposes of this study, it was important to look at the mean,

⁶⁹ Dörnyei (2007) mentions that SPSS is typically the statistics program used for studies in Applied Linguistics.

standard deviation (SD), minimum and maximum scores, as well as extreme scores (“fringeliers”) and outliers. In addition, the percentage difference from pre- to post-course for both listening tests was computed as a new variable in the SPSS spreadsheet. Again, descriptive tests were run both with the participants as a whole and then separated or “split” by group. Generating this information enabled the researcher to determine whether there were potential outliers who would need to be eliminated.

All participants took Listening Test # 1 pre- and post-course; however, this was not the case for Listening Test # 2. C2 and C5 did not take the pre-course Listening Test # 2. Therefore, it was necessary to exclude these two participants from the calculations for Listening Test # 2. Additionally, C8 was removed from the Listening Test # 2 results because she was determined to be an outlier who had to be eliminated for both statistical and logical reasons. This participant scored third highest in pre- and post-course Listening Test # 1 and post-course Listening Test # 2, but scored second lowest on the pre-course Listening # 2 test. Such a score appeared to be out-of-character for this participant’s listening ability. Furthermore, the fact that her score increased by 100% from pre to post-course in Listening Test # 2, an increase not achieved by any participant in any of the tests, and flagged statistically as an outlier, the decision was made to eliminate C8 from the Listening Test # 2 results. Therefore, the Listening Test # 2 results for the control group are based on five participants, C1, C3, C4, C6, and C7.

With regard to the songs group, there was one participant, S5, who was the only one of all three groups that performed worse on the post-course test. In fact, his performance was so bad that his score dropped 20.09% from pre- to post-course. However, when he was analysed against all of the participants, he was not statistically flagged as an outlier, but rather was a borderline outlier or “fringelier”. Nevertheless, when analysed as part of the songs group only, he was

statistically an outlier in the post-course test. Logically, the researcher felt that the participant should be eliminated from Listening Test # 2 because his post-course score was not consistent with his performance. In Listening Test # 1, this participant's percentage improvement was slightly above average for the group, even though he was absent for two days in a row. Furthermore, this participant was going to drop out of the study and then decided not to. After talking to him about this, it was clear that the participant had personal or health issues. Therefore, because of somewhat conflicting statistical, quantitative, and qualitative data, the researcher decided to run two sets of statistical tests for Listening Test # 2, one with S5 and one without S5. Initially, the quantitative results for the songs group are shown using both scenarios, i.e. with S5 and without S5. However, when the discussion of results moves to the individual pronunciation phenomena, S5 is not included as part of the songs group for any of the statistical tests related to Listening Test # 2.

Once the outliers were eliminated, then a descriptive statistics test which split the data of the three groups was run in order to determine the low and high score, mean, and standard deviation for each group in each listening and speaking test. After this, the percentage improvement from pre- to post-course was calculated for each listening and speaking test. The formula used was $(\text{post minus pre}) / \text{pre} * 100^{70}$. To calculate the mean and percentage difference of the individual pronunciation phenomena covered in the course, the above statistical tests were run using the pre- and post-course scores of each pronunciation area as the variables and then the percentage difference was calculated for each of the pronunciation areas.

⁷⁰ When a pre-course score was zero, which sometimes occurred in the speaking tests, the program could not calculate a percentage difference score. When this was the case, the percentage decrease was instead considered and the value of zero was entered. In this way, all participants would be included in the ensuing statistical calculations.

After running the descriptive statistics tests described above, inferential statistics tests were run. Independent Samples T-Tests were run when comparing the control group to the treatment group (songs + no-songs) for research questions 1a and 2a, and one-way analysis of variance (ANOVA) tests were run for research questions 1b and 2b when the control, songs, and no-songs groups were being compared. These tests allowed us to compare the means of the percentage improvement overall for the listening and speaking tests as well as the individual pronunciation areas in order to see whether there was any significant difference between the performance of the groups. Due to the small sample sizes, Games-Howell post hoc tests were used (see Newsom, 2006).

Finally, after the quantitative analysis of the listening tests was done, qualitative information that was obtained on members of the songs group from the Final English Pronunciation Questionnaire and Speaking Test # 2 was first presented and then discussed in relation to the actual results of the Listening Tests, percentage of improvement from pre- to post-course and how they did in comparison to their group as well as across the groups.

Speaking tests. The complete analysis of Speaking Test # 1 occurred in stages. In what we can refer to as stage one, an Excel spreadsheet was set up. The participants were listed on the vertical axis according to their identification code (e.g. C1, S4, NS2) while the speech phenomena that were included in the course were listed along the horizontal axis. Results could not be calculated out of a total score as was the case with the listening tests. Unlike the listening tests in which the students either heard the pronunciation phenomena or they did not, the speaking tests are not so straightforward. The flexible nature of thought groups, for example, means that the number of thought groups in any one text will vary and consequently affect the

number of focus words, in addition to instances of connected speech and linking⁷¹. Thus, it was necessary to consider instead the number of errors that the participants made⁷². To do this, I listened to each recording and when a participant made an error in any of the areas taught, it was indicated on the paper copy of the test by the error being circled and labeled according to the type of error made. For example, when a participant made a focus word error, the (incorrect) word that was stressed was circled and “focus word” was written above it. A separate copy of “The Rainbow Passage” was used for every test that was analysed, thus allowing for the place and type of errors to be clearly visible.⁷³ A thorough analysis of the reading of the text required listening to each recording a number of times. To help ensure that all relevant errors were caught, the participants’ initial and final recordings were analysed back to back, and it was often necessary to go back and double-check the analysis of the initial recording against the final one and vice-versa. By comparing the two utterances that corresponded to the same part of the text, I was able to distinguish errors that were sometimes difficult to perceive, such as the schwa. Once I had analysed all of the recordings, I noted the number of errors in the excel document.

Stage two involved having the speaking tests analysed by my research assistant⁷⁴. Before he began, I provided him with the audios of the speaking tests, a blank copy of the Excel

⁷¹ These issues are discussed in further detail below.

⁷² Ellis (2008) discusses the use of error analysis in L2 production. Couper (2003) used error analysis in his study of pronunciation by summing up the errors and averaging them. Derwing and Rossiter (2003) categorized and counted errors while Malmeier (2014), in a case study, analysed different types of errors made by the participants at two different times.

⁷³ While certainly an option, I decided against phonetically transcribing the speech according to IPA conventions because doing so would not have allowed for the same amount of clarity. The sheer number of symbols and diacritics used in narrow IPA transcription would have obscured the select few types of errors that we were looking for, not to mention there not being a direct correspondence between elements of the Prosody Pyramid and IPA symbols.

⁷⁴ My research assistant is a native Portuguese speaker from Brazil. His proficiency in English is at a low advanced level. He completed the English Intensive Program at the University of Ottawa and was the Valedictorian for the bridging level. In addition he was the top student in all three pronunciation workshops that I offered during the three terms that he studied ESL at the University. His demonstrated skill in perceiving and producing English pronunciation and his familiarity with the course content taught in Chile, made him an appropriate candidate as a

spreadsheet for noting the errors, as well as photocopies of the reading text, “The Rainbow Passage”. It should be noted that at no time did my assistant have access to my results of the speaking tests. What follows are the steps he took while analysing the speaking tests.

All the recordings were listened to multiple times in order to identify the errors made by the participants. That is, for each pronunciation area (thought groups, focus words, word stress, schwa, and linking), the recording was listened to at least once and only that pronunciation area was focussed on. Then the recording was listened to again and a different pronunciation area was then concentrated on. This process was repeated until all of the pronunciation areas were covered. As the pronunciation errors were found, they were indicated on the copy of “The Rainbow Passage” that pertained to the participant’s recording that was being evaluated, e.g. NS2 initial. Then the total number of errors that each speaker made in each area was indicated for both the pre- and post-course tests in the Excel spreadsheet.

After completing the analysis, my assistant noted some of the difficulties he faced. He found that the analysis of thought groups was sometimes challenging, because it was difficult to determine whether a hesitation should be considered a thought group or a pronunciation error⁷⁵. Since there are different ways to arrange thought groups in some sentences, it was difficult for him, as a non-native speaker, to determine when there were too few or too many thought groups. After discussing the analysis of thought groups with him, he felt that identifying where there is a thought group error is much easier and almost intuitive for those who speak English as a first language. Analysing word stress and schwa problems were also sometimes quite challenging. In

research assistant. He is currently working on an interdisciplinary degree in science and technology. He plans on completing another undergraduate degree in neuroscience, specializing in neurolinguistics.

⁷⁵ Kormos and Dénes (2004) say that “unfilled pauses shorter than three seconds are generally regarded as articulation pauses and not as hesitation phenomena” (p. 151). Therefore, only pauses longer than 3 seconds were considered to be errors if there were at a thought group boundary.

many tests, he found it difficult to determine whether or not both pronunciation errors occurred in one word. The pronunciation of the word *horizon* is one of the best examples. The majority of the participants tended to pronounce this word as ['ho.ɪzən] instead of the correct way [hə 'jaɪzən]. From the transcriptions of the word into the International Phonetic Alphabet, it can be seen that there are two errors⁷⁶: a word stress and a schwa error. The phenomenon described can be best explained by the fact that a word stress error can cause a schwa error, because a schwa cannot be the vowel sound in a syllable that has primary stress. In words where the two pronunciation errors occurred, the major difficulty was identifying which error was the main one and if it caused the second error. In the case of the word *horizon*, the word stress error was considered the main problem, and since it caused the accompanying schwa error in that syllable, the only error that was recorded was the word stress error. Regarding focus words and linking, my assistant had no problems of analysis.

Stage three was when I compared my results in stage one with those of my assistant in stage two, and then performed a final double reanalysis of each speaking test in order to resolve discrepancies in our two sets of results. First of all, I compared our results to see how similar and different they were. Since there were a number of differences, I went through and analysed each speaking test twice again: once by listening and checking my original working copy and a second time by listening and checking my assistant's working copy. By doing this I was able to first determine where our results differed and then clarify the discrepancies. For example, there were times when my assistant noted a thought group error where there was not one. At other times, he noted linking errors that I had missed. Going through each test twice allowed me to

⁷⁶ There are actually three errors, however, since the schwa was the only vowel sound that was covered in the course, other segmental errors were not evaluated or counted as errors.

confirm (or not) where my assistant and I had noted identical errors, and for the errors in which there was not consensus, determine whether or not there was an error and if so, which kind. Then, my research assistant and I got together to listen to the recordings again, so that we could discuss any questionable discrepancies between our sets of results. Having a research assistant and following these stages of analysis allowed for more accurate results than would have been possible with only one person completing the analysis.

Finally, it is important to mention that it was advantageous to have a research assistant who was a non-native speaker of English. Although he tended to perceive thought group errors where there were none or failed to perceive schwa errors when there were ones, he was able to perceive linking errors that I as a native speaker had missed. With regards to focus words and word stress, those areas were the easiest for both of us to analyse and consequently were the two areas where our results tended to agree most often.

Both for myself and for my research assistant, the analysis of some pronunciation areas posed a few problems. First of all, although thought groups seem to be generally quite straightforward, if a speaker made an abnormally long pause (i.e. more than three seconds), it was considered a thought group error. Since the participants had been provided with the pronunciation of a few individual words and had been allowed to practice the reading before recording it, they should not have made abnormally long pauses. In fact, most of them did not. However, there were a few that did and if it was obvious that a participant was stumbling, this was taken into consideration, and such stumbles (pauses) were not considered to be errors.

Furthermore, it was sometimes difficult to determine whether a slight pause should be considered a thought group error or a linking error. If the pause occurred very close to a previous

thought group boundary, e.g. between the verb and the object or adverbial phrase when there had been a thought group boundary between the subject and the verb, then it was considered to be a linking error. For example, “When the sunlight strikes / raindropsin the air, / they / act* as a prism...”. If there was a pause between “act” and “as”, then it was counted as a linking error, because even though it could have been reasonable to make a thought group boundary there, there was already one after the subject, so there should not have been one before the adverbial phrase.

Another aspect to consider was that as the thought group boundaries change, so might the focus words. For example, the following two ways of saying the same phrase are correct. Note that thought group divisions are indicated with slashes and focus words are underlined.

“...one showing mainly red and yellow, / with little or no / green or blue.”

“...one showing mainly red and yellow, / with little or no green / or blue.”

The most difficult speech phenomenon to analyse was the schwa. Sometimes it was obvious when a participant said a full vowel and sometimes it was obvious when they used a schwa. However, there were times when it seemed that they used a sound that was in between a schwa and the full vowel. Such cases required listening to the segment over and over again until we were able to decide whether the sound the participants emitted was closer to a schwa or the full vowel.

All in all, the process involved in analysing the speech in terms of the elements of the Prosody Pyramid, proved to be much more interesting and challenging than initially anticipated. Having a non-native speaker assist in the analysis helped provide a great deal of insight into speech analysis that a native speaker alone might not consider. Furthermore, while the elements

of the Prosody Pyramid are presented in the textbooks as occurring in generally predictable ways (that correspond to grammatical conventions), we discovered that this is not always the case.

For Speaking Test # 2, the same analytical steps that were taken in Speaking Test # 1 were applied. An Excel spreadsheet, much like the one for Speaking Test # 1 was set up. The participants were listed on the vertical axis while the speech phenomena that were included in the course were listed along the horizontal axis. Then, after determining the total amount of time that each participant spoke for, I selected a one-minute segment of their recording, being careful to choose a segment that was in the middle. Then each segment of the recording was listened to and transcribed into written English. Doing this allowed me to visually note the pronunciation errors on the transcription. The following is an excerpt from one of the participants before analysing it for errors:

S5 – PRE Total time: 0:0 – 2:17 min/sec Analysed: :30 – 1:30 min/sec

The whole book eh it's situated in in the middle of this road where the two men are waiting and they are only talking about ehm how how how they their lives have passed and and they are mainly talking about eh their lives in a very in a very complicated way I mean they talk about their relationship eh through the years as they have been together for 15 years wondering around the world wo I guess and the thing is that they they while they are waiting they met a a man with his slave and they talk a lot and they get to to realize

As in Speaking Test # 1, each recording was listened to multiple times until all pronunciation areas covered in the course were evaluated and agreed upon by both my research assistant and myself. It is important to stress that only the pronunciation areas covered in the

course were evaluated. While interesting (and salient in Speaking Test # 2), fluency phenomena were not evaluated because to do so would have been beyond the scope of this thesis.

When looking at the average total number of errors between Speaking Tests # 1 and 2, one will notice that the participants had a higher number of errors on the Speaking Test # 1, which was a test of reading and therefore a controlled activity, than they did on Speaking Test # 2, which was a free speaking activity. This difference in number of errors is partly due to the fact that the reading passage was longer than the free speech samples. More importantly, though, the difference in errors also stems from the fact that the two speaking tests are by nature different. Speaking freely is natural speech and something that people do on a regular, daily basis. Reading out loud, though, for most people is not a common speaking activity (Levis and Barriuso, 2012). Furthermore, read speech is not considered natural or “spontaneous” speech (Llisterri, 1992). The point is not that reading is a bad measure of speaking, but rather that reading aloud and spontaneous speech are different speaking styles which may lead to different results. Indeed, Levis and Barriuso (2012) indicate that the types of errors may not be the same especially with L2 speakers because English spelling and pronunciation have an indirect correspondence. Therefore, although the average number of errors between the two tests should not be compared, the types of errors made in the two tests should. Doing so may provide insight into whether or not controlled and free speaking tests result in similar types of errors.

With regard to the data analysis of the speaking tests, the same SPSS procedures that were done with the listening tests were completed with both speaking tests. That is, the data were entered on the spreadsheet and then descriptive and inferential statistics were run. After that,

qualitative data from the Final English Pronunciation Questionnaire, Speaking Test # 2⁷⁷, and my personal observations during the classes were included in the discussion of the results. Doing this allowed for the results to be viewed from different perspectives⁷⁸.

Summary

In this chapter we described and discussed the methods and procedures that were followed in the design and implementation of an English pronunciation course that was given over a two-week period to Chilean undergraduate students studying English in Santiago. It began with a discussion of the research, why and how it was designed in the way that it was. The criteria regarding the participants, the number of groups, and the course length, content, and approach were explained before covering the steps that were involved in designing and implementing the intervention. The various tasks that each group of participants had to do were explained as well as how the pronunciation classes were carried out. Information was provided regarding the questionnaires and tests that were given to the students before and after the treatment period. Finally, we discussed how these data sources were evaluated.

⁷⁷ For the students who chose to talk about the pronunciation course in the post-course Speaking Test # 2, all of what they said pertaining to this topic was transcribed.

⁷⁸ Duff (2008) mentions that the use of multiple perspectives or triangulation is common to case studies.

Chapter Four

Research Findings

Introduction

The following discussion of results begins with an inquiry as to whether the pronunciation course produced any effect, independent of the type of materials used, on the participants' ability to perceive pronunciation phenomena. How the control group and treatment group performed on the listening tests, overall as well as in the specific pronunciation areas taught, are considered. Then, the extent to which songs versus other materials contributed to this effect are looked at in a similar fashion, except for this purpose the treatment group is divided into a songs and a no-songs group, with the control group kept unchanged for purposes of comparison. Next, further data obtained is discussed in light of the results of the listening tests. Then, the discussion turns to the efficacy of the course with regard to the production of the pronunciation phenomena covered, in which the results of the speaking tests are discussed. The same analytical steps that were taken for the listening tests are followed for the speaking tests, including a discussion of further data obtained. Finally, additional information contained in the final pronunciation questionnaire is discussed as separate, additional findings.

Analyses of Research Questions

Research Question # 1a

Listening Tests. The first research question asks if teaching suprasegmental phenomena in a two-week pronunciation course can enable L2 learners to better perceive suprasegmental phenomena of the target language. In order to determine this, two separate listening tests were applied pre- and post-course. Listening Test # 1, taken in part from Gilbert (2005b), was a five-part test of decontextualized / controlled speech that specifically covered areas of the Prosody Pyramid, namely syllables, word stress, focus words, contractions and reductions, and thought groups. Listening Test # 2 was a dictation exercise of a dialogue which contained instances of linking and connected speech that were covered in the course⁷⁹. The performance of the control group and the treatment group (i.e. the songs and no-songs groups together) were compared for both Listening Test # 1 and Listening Test # 2, and, as we will see, although the treatment group had a greater mean percentage increase from pre- to post-course on both listening tests, the increase was not statistically significant. Nevertheless, when each pronunciation phenomenon was considered separately, there were statistically significant⁸⁰ results in the area of word stress, which was measured in Listening Test # 1.

To provide an initial summary of the findings, we applied descriptive statistics, namely, measures of central tendency and measures of variability. Tables 4.1 and 4.2 show the number of participants, the mean scores (out of a reduced value of 10), and the standard deviation of the control and treatment groups for Listening Test # 1 and Listening Test # 2 pre- and post-course.

⁷⁹ For the purposes of this investigation, although linking is by nature a feature of connected speech, it is considered as a separate pronunciation phenomenon from connected speech, which (here) involves actual segmental change, such as assimilation and reduction.

⁸⁰ For all statistical tests, an alpha level of .05 was used. That is, a *p*-value above .05 we classify as not significant and anything below, we classify as significant.

Table 4.1.

Listening Test # 1 Descriptive Statistics for the Control and Treatment Groups

	Group			
	Control		Treatment	
	Pre-	Post-	Pre-	Post-
<i>Number of Participants (n)</i>	8	8	15	15
<i>Mean (M)</i>	8.41	8.75	7.92	8.62
<i>Standard Deviation (SD)</i>	1.10	.67	.60	.72

Table 4.2.

Listening Test # 2 Descriptive Statistics for the Control and Treatment Groups

	Group					
	Control		Treatment (with S5) ⁸¹		Treatment (without S5)	
	Pre-	Post-	Pre-	Post-	Pre-	Post-
<i>Number of Participants (n)</i>	5	5	15	15	14	14
<i>Mean (M)</i>	6.40	7.53	7.09	8.42	7.12	8.64
<i>Standard Deviation (SD)</i>	2.28	1.98	1.61	1.21	1.67	.89

In order to know how the control and treatment groups did in comparison to each other, we will look at their average (i.e. mean) pre- and post-course scores and the degree to which they improved in each listening test. Table 4.3 shows the mean of the pre-course scores, the mean of the post-course scores, and the mean percentage increase from pre- to post-course. The mean percentage increase is the sum of the percentage increases for each participant in a group divided by the number of participants in that group. It is not simply the percentage increase, which is the percentage difference between the pre-course mean and post-course mean of the group as a whole. This is important to make clear because sometimes the percentage increase

⁸¹ Chapter Three discusses the dilemma involving the inclusion or not of S5 in the results for Listening Test # 2.

and mean percentage increase can be quite different⁸². We present the mean percentage increase here because it is a more accurate measure of each group's improvement, taking into account the variability within the group, and it is thus a measure used in statistics.

Table 4.3.

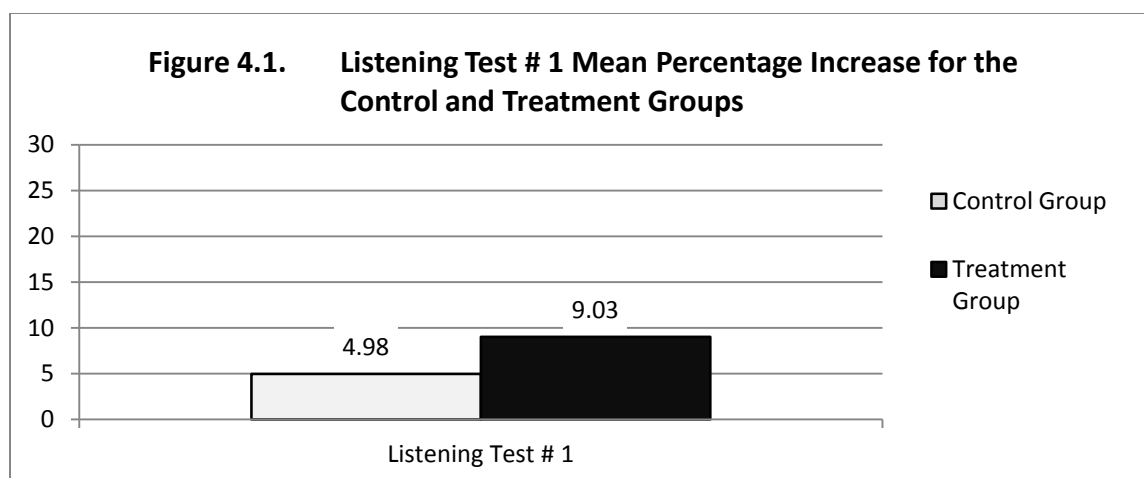
Listening Test # 1 Mean Scores and Mean Percentage Increase for the Control and Treatment Groups

	Group					
	Control			Treatment		
	Pre	Post	%	Pre	Post	%
Listening Test # 1	8.41	8.75	4.98	7.92	8.62	9.03

As we can see in Table 4.3, both groups had a high pre-course test score in Listening Test # 1, which left little room for improvement, since the highest score possible was ten. The control group had a pre-course mean of 8.41 and a post-course one of 8.75, which translated into an increase of 4.98%. While it was not expected that the members of the control group would better their score, their slight improvement could be due to test familiarity or simply that they are part of an English language programme and they are learning the language. As the chart indicates, the treatment group had a 9.03% increase from their pre-course mean of 7.92 to post-course, 8.62.

Figure 4.1 shows the mean percentage increase for the two groups.

⁸² If we have group mean pre-test and post-test scores of 14 and 20 respectively, then there is an overall group increase of 6, which translates into a percentage difference of $6/14 = 43\%$. With five participants in the group whose percentage differences from pre- to post-course were 50%, -43%, -36%, 100% and 100% respectively. The sum of these percentages divided by 5 (the number of participants) gives us a mean percentage difference of 34%, lower than the group percentage difference, since it takes into account two actual percentage decreases (negative numbers) which pull down the mean percentage increase. It can be seen from this example that the mean percentage difference takes into account variability between participants, whereas the group percentage difference does not.



In order to determine whether this increase was statistically significant, an independent-samples t-test was conducted to compare the mean percentage increase in Listening Test # 1 for the control and treatment groups. There was a not a significant difference in the scores for the control group ($M = 4.98$, $SD = 9.16$) and the treatment group ($M = 9.03$, $SD = 8.24$) conditions; $t(21) = 1.08$, $p = .29$. These results suggest that for our case, the pronunciation course had no clear effect on the ability of the participants to better perceive pronunciation phenomena contained in the Prosody Pyramid.

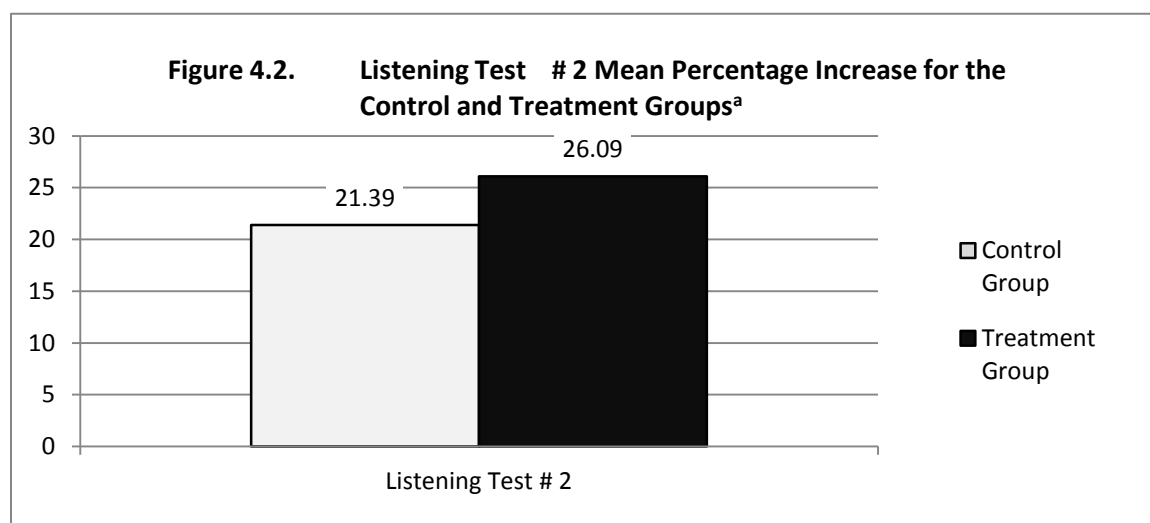
For Listening Test # 2, Table 4.4 indicates the pre- and post-course scores and mean percentage increase.

Table 4.4.

Listening Test # 2 Mean Scores and Mean Percentage Increase for the Control and Treatment Groups

	Control Group			Treatment Group (with S5)			Treatment Group (without S5)		
	Pre	Post	%	Pre	Post	%	Pre	Post	%
Listening Test # 2	6.40	7.53	21.39	7.09	8.42	23.01	7.12	8.64	26.09

In Listening Test # 2, the groups scored lower in the pre-course test, which indicates that this test was more challenging for the participants than Listening Test # 1. Because their scores were initially lower, it allowed for more improvement to be made, which was in fact the case. The control group, for some reason had a 21.39% improvement from pre- to post-course, which is difficult to explain. The results of the treatment group differed depending on whether the participant S5 was included or not. With S5 the treatment group improved on average by 23.01%; however, without said participant⁸³, the improvement was 26.09%.



^aWithout S5.

In order to determine whether this increase was statistically significant, an independent-samples t-test was conducted to compare the mean percentage increase in Listening Test # 2 for the control and treatment groups. There was a not a significant difference between the control group ($M = 21.39$, $SD = 14.73$) and the treatment group ($M = 26.09$, $SD = 23.65$) conditions; $t(17) = .41$, $p = .69$. These results suggest that for our case, the pronunciation course did not have an

⁸³ All discussion from here on of the treatment group's performance in Listening Test # 2 excludes S5.

effect on the ability of the participants to better perceive linking and connected speech in free speech.

Now that the results of the two listening tests as a whole have been analysed, it is important to see what kind of effect the pronunciation course might have had in the individual pronunciation areas contained in the tests. Therefore, we will consider how the groups performed in the different pronunciation areas of the tests. For Listening Test # 1, these are: thought groups, focus words, word stress, and contractions / reductions. For Listening Test # 2, they are linking and connected speech. Table 4.5 combines the results of both tests.

Table 4.5.

Listening Test # 1 and 2 Mean Scores and Mean Percentage Increase in Each Pronunciation Area for the Control and Treatment Groups

Pronunciation Area	Group					
	Control			Treatment		
	Pre	Post	%	Pre	Post	%
Thought Groups	9.50	9.75	3.09	9.73	9.73	.33
Focus Words	6.75	8.13	36.67	6.07	7.47	30.31
Word Stress	8.88	8.50	-3.99	7.27	8.40	22.52
Contractions & Reductions	8.50	8.63	2.37	8.60	8.87	3.35
Linking	5.53	7.18	39.84	6.77	8.40	32.67
Connected Speech	7.54	8.00	7.34	7.58	8.96	21.66

Table 4.5 shows that in the area of thought groups, the treatment group actually had a lower mean percentage increase (.33%) than did the control group (3.09%). It is important to mention, though, that in the case of both groups, there was very little room for any increase given that the scores were out of a possible maximum of ten. The high pre-course scores that both the control (9.50) and treatment (9.73) achieved are an indication that the students were already competent in the perception of thought groups.

In focus words, the control group also had a higher mean percentage increase (36.67%) than the treatment group (30.31%). Judging from the pre-course scores, focus words is an area in which there is room for improvement, although it is not clear whether pronunciation perception instruction could help. Furthermore, given that the control group experienced such a high percentage increase, it is possible that test familiarity positively affected their performance. For both thought groups and focus words, there were no statistically significant differences in the performance of the two groups of participants.

In the area of word stress, however, if we look at the mean percentage increase for the two groups, we can see that the control group had a 3.99% decrease from pre- to post-course whereas the treatment group had a 22.52% increase. An independent-samples *t*-test was conducted to compare the mean percentage difference in word stress in Listening Test # 1 for the control and treatment groups. There was a significant difference in the scores for the control group ($M = -3.99$, $SD = 10.48$) and the treatment group ($M = 22.52$, $SD = 43.85$) conditions; $t(6.35) = 2.23$, $p = .04$. The effect size (eta squared) was .23, which was large⁸⁴, which means that the significance is important, i.e. not ignorable.

In contractions and reductions, both the control group and treatment group had high pre-course scores, 8.50 and 8.60, respectively. While such high scores would suggest that the participants were already competent in this area, it is important to remember that this pronunciation area corresponds to Listening Test # 1, in which the speech was controlled and decontextualized. Therefore, it should not be assumed that the participants would perform as well had the test contained natural speech.

⁸⁴ Dörnyei, Z. (2007) indicates that effect sizes (eta squared) are usually interpreted in the following manner: “.01 = small effect, .06 = moderate effect, and .14 = large effect” (p. 217).

Linking, like focus words, is a weaker area for the participants considering their pre-course test scores. The fact that the control group improved by 39.84% and the treatment group by 32.67% suggests either test familiarity and/or a lack of effect of the pronunciation course on the perception of linking.

Finally, with regard to connected speech, we can see that the mean percentage increase for the treatment group (21.66%) is much higher than that of the control group (7.34%). Nevertheless, there was not a significant difference between the control group ($M = 7.34$, $SD = 7.21$) and the treatment group ($M = 21.66$, $SD = 23.26$) conditions; $t(17) = 1.33$, $p = .051$, $\eta^2 = .09$. Although the percentage increase for the treatment group was not found to be significant, the fact that it is so close to being significant as well as having close to a large effect size is reason to consider conducting further research in this pronunciation area.

In sum, the results in Table 4.5 suggest that for our case, the pronunciation course had a significant effect on the ability of the participants to better perceive word stress. Another pronunciation area that might benefit from pronunciation instruction is connected speech. For other phenomena, no significant improvement was observed.

Summary of Research Question # 1a. On whether the pronunciation course benefited the treatment group's ability to better perceive pronunciation phenomena, the overall results suggest that it did not. However, after analysing the results for each pronunciation area, we can see that the perception of word stress was significantly helped. The analysis of connected speech did not show a significant difference, because p was .001 away from being significant. In light of such a result, though, we cannot clearly say that the course did not benefit the treatment group; in fact, it seems that it has. In the other pronunciation areas, namely, thought groups, focus words,

contractions and reductions, and linking, the analysis of the results indicates that the course had no clear effect on the participants' perception of these areas.

Research Question # 1b

Listening Tests. The second part of the first research question asks whether songs as a material are beneficial for perceiving the suprasegmental phenomena taught in a pronunciation course over a period of two weeks. In order to determine this, the treatment group was divided into two separate groups: the songs group, who received instruction using songs and the no-songs group, whose instruction did not include songs. We will see that when we look at how the three groups of participants performed on Listening Test # 1 and Listening Test # 2, it appears that songs may enable language learners to better perceive some pronunciation phenomena – but the results are inconclusive. In Listening Test # 1, the songs group had a higher mean percentage improvement in their pre- to post-course scores than did the control group, but this was not the case in Listening Test # 2. In both tests, although the no-songs group had the greatest mean percentage increase from pre- to post- course, the increases were not significant. With regard to individual scores, there were specific songs group participants who performed higher than average across the groups and who commented that songs were indeed beneficial for improving their (pronunciation) listening skills⁸⁵.

To provide an initial summary of the findings, we have applied descriptive statistics namely, measures of central tendency and measures of variability. In Tables 4.6 and 4.7, we can

⁸⁵ In a few pages, when we look at the qualitative data, more detail will be provided on the participants and how they felt the course helped them.

see the number of participants, the mean scores, and the standard deviation in the pre- and post-course scores for Listening Test # 1 and Listening Test # 2.

Table 4.6.

Listening Test # 1 Descriptive Statistics for Each Group

	Group					
	Control		Songs		No-Songs	
	Pre-	Post-	Pre-	Post-	Pre-	Post-
<i>Number of Participants (n)</i>	8	8	10	10	5	5
<i>Mean (M)</i>	8.41	8.75	7.92	8.58	7.93	8.70
<i>Standard Deviation (SD)</i>	1.10	.67	.70	.84	.38	.47

Table 4.7.

Listening Test # 2 Descriptive Statistics for Each Group

	Control		Songs (with S5)		Songs (without S5)		No-Songs	
	Pre-	Post-	Pre-	Post-	Pre-	Post-	Pre-	Post-
<i>Number of Participants (n)</i>	5	5	10	10	9	9	5	5
<i>Mean (M)</i>	6.40	7.53	7.50	8.53	7.59	8.89	6.27	8.20
<i>Standard Deviation (SD)</i>	2.28	1.98	1.47	1.36	1.53	.82	1.72	.93

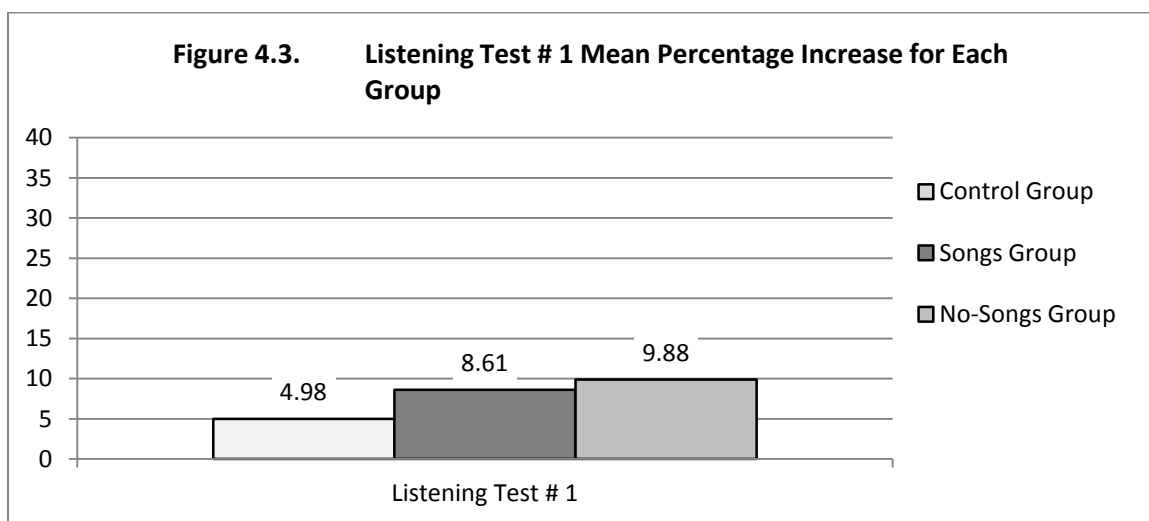
In order to know how the three groups did in comparison to each other, we will look at their average (i.e. mean) pre- and post-course scores and the degree to which they improved in each listening test.

Table 4.8.

Listening Test # 1 Mean Scores and Mean Percentage Increase for Each Group

	Control Group			Songs Group			No-Songs Group		
	Pre	Post	%	Pre	Post	%	Pre	Post	%
Listening Test # 1	8.41	8.75	4.98	7.92	8.58	8.61	7.93	8.70	9.88

In Listening Test # 1, the pre-course average score for the control group was 8.41 and their post-course average was 8.75, which translates into a mean increase of 4.98% in the Listening Test # 1 post-course test. The songs group had a pre-course average score of 7.92 in Listening Test # 1 and finished with an 8.58 in the post-course test, which is a mean increase of 8.61%. This indicates that the improvement of the songs group was almost twice (i.e. 1.73) that of the control group. Finally, with regard to the no-songs group, in Listening Test # 1, the pre-course average was 7.93 and their post-course average was 8.70, a mean increase of 9.88%. Their increase was almost twice that of the control group (i.e. 1.98). Figure 4.3 shows the mean percentage increase for each of the three groups.



Although, the results of Listening Test # 1 suggest an improvement with pronunciation instruction, irrespective of the materials used (songs versus non-song texts), it was important to determine whether or not this improvement was statistically significant. Table 4.9 shows the results of the one-way ANOVA that was run for the mean percentage difference in scores from pre- to post-course for Listening Test # 1.

Table 4.9.

Listening Test # 1 One-Way Analysis of Variance

Descriptives								
	N	Mean	Std. Deviation	Std. Error	95% Confidence Interval for Mean		Minimum	Maximum
					Lower Bound	Upper Bound		
Control	8	4.9827	9.15916	3.23825	-2.6746	12.6399	-6.53	19.71
Songs	10	8.6118	9.39360	2.97052	1.8921	15.3316	-5.71	23.55
No Songs	5	9.8751	6.13446	2.74341	2.2582	17.4920	.00	16.94
Total	23	7.6241	8.58864	1.79086	3.9101	11.3381	-6.53	23.55

ANOVA						
	Sum of Squares	df	Mean Square	F	Sig.	
Between Groups	90.908	2	45.454	.593	.562	
Within Groups	1531.916	20	76.596			
Total	1622.824	22				

As can be seen in the above table, there were no statistically significant differences in the mean percentage difference between the three groups according to the one-way ANOVA $F(2, 20) = .593, p = .562$. These results suggest that the pronunciation course did not have an effect on the participants' ability to better perceive pronunciation phenomena contained in the Prosody Pyramid.

In Listening Test # 2, the results are rather different. The following table, Table 4.10, indicates the pre- and post-course scores and mean percentage increase for Listening Test # 2.

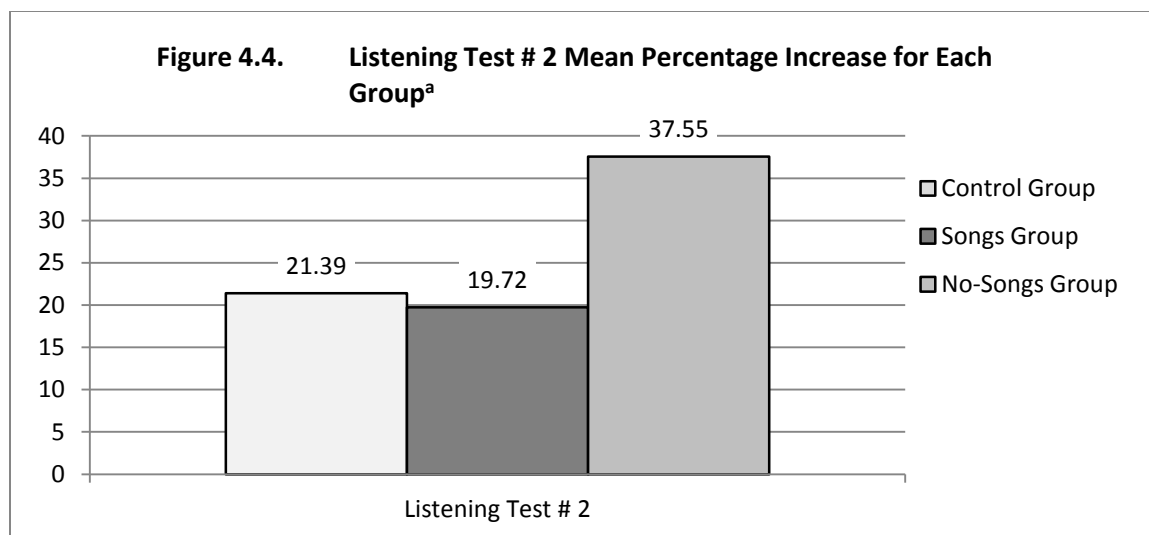
Table 4.10.

Listening Test # 2 Mean Scores and Mean Percentage Increase for Each Group

	Control Group			Songs Group (with S5)			Songs Group (without S5)			No-Songs Group		
	Pre	Post	%	Pre	Post	%	Pre	Post	%	Pre	Post	%
Listening Test # 2	6.40	7.53	21.39	7.50	8.53	15.74	7.59	8.89	19.72	6.27	8.20	37.55

In this test, the control group had an average pre-course score of 6.40 and a post-course score of 7.53, which translates into a mean percentage increase of 21.39%. An increase of this magnitude is unexpected for a control group. In Listening Test # 2, in which S5 was included, the songs group's mean percentage increase was 15.74%, which was lower than that of the control group. To reiterate what was discussed in Chapter Three, it was in this test, Listening Test # 2, that the post-course performance of S5 is in question, because this participant was the only one out of all of the groups to score lower in the post- than in the pre-test. In fact, this student had a 20.09% *decrease* from pre- to post-course, which was inconsistent considering how the participant performed on other tests. As a result, the researcher felt that his post-course score was not reliable and so the decision was made to eliminate S5 from the songs group and consider the scenario without him⁸⁶. Nevertheless, in the scenario in which S5 was not included, although the songs group's mean percentage increase jumped to 19.72%, it was still lower than that of the control group. The no-songs group, on the other hand, improved by 37.55% from pre- to post-course, which was 1.76 times that of the control group and 2.39 times that of the songs group (without S5). The following figure, Figure 4.4, displays the improvement that each group made from pre- to post-course.

⁸⁶ All discussion from here on of the songs group performance in Listening Test # 2 excludes S5.



^aWithout S5.

As we can see in Figure 4.4, all groups had higher mean percentage increases in their post-course test scores in Listening Test # 2 than they had in Listening Test # 1 (which were 4.98% for the control group, 8.61% for the songs, and 9.88 for the no-songs). It is important to mention, though, that because all of the groups scored lower in the pre-course Listening Test # 2 than they did in Listening Test # 1, the possibility of achieving a higher mean percentage increase was in fact greater. Nevertheless, in Listening Test # 2, the no-songs group had the highest increase of the three groups, followed by the control group. In order to know whether the differences in the increases among the groups were significant, a one-way ANOVA was run. It indicated that there was no significant difference between any of the groups, $F(2, 16) = 1.240$, $p = .316$. Table 4.11 shows these results.

Table 4.11.

Listening Test # 2 One-Way Analysis of Variance

Descriptives								
	N	Mean	Std. Deviation	Std. Error	95% Confidence Interval for Mean		Minimum	Maximum
					Lower Bound	Upper Bound		
Control	5	21.3936	14.73111	6.58795	3.1025	39.6846	3.64	36.24
Songs	9	19.7196	16.38778	5.46259	7.1228	32.3164	3.41	46.91
No Songs	5	37.5520	32.02067	14.32008	-2.2069	77.3109	4.30	73.40
Total	19	24.8529	21.37121	4.90289	14.5523	35.1534	3.41	73.40

ANOVA						
	Sum of Squares	df	Mean Square	F	Sig.	
Between Groups	1103.328	2	551.664	1.240	.316	
Within Groups	7117.788	16	444.862			
Total	8221.116	18				

After looking at the overall results of Listening Test # 1 and 2, which present no clear evidence that the course was helpful for improving the students' perception of pronunciation phenomena, it is necessary to look at how each group performed in each pronunciation area covered in the course.

Table 4.12.

Listening Test # 1 and 2 Mean Scores and Mean Percentage Increase in Each Pronunciation Area for Each Group

Pronunciation Area	Group								
	Control			Songs			No-Songs		
	Pre	Post	%	Pre	Post	%	Pre	Post	%
Thought Groups	9.50	9.75	3.09	9.80	9.80	.50	9.60	9.60	.00
Focus Words	6.75	8.13	36.67	5.90	7.40	32.14	6.40	7.60	26.67
Word Stress	8.88	8.50	-3.99	7.30	8.30	18.11	7.20	8.60	31.35
Contractions & Reductions	8.50	8.63	2.37	8.65	8.80	1.81	8.50	9.00	6.43
Linking	5.53	7.18	39.84	7.19	8.63	25.27	6.00	8.00	45.98
Connected Speech	7.54	8.00	7.34	8.12	9.23	16.38	6.61	8.46	31.16

In Table 4.12 we can see that all of the groups performed the best in the area of thought groups. Since all of the groups had an average pre-course score of 9.50 or higher, it is an indication that thought groups did not pose a problem for these participants. Contractions and reductions, much like thought groups, was an area in which the students also initially did well. The mean percentage improvement of the groups in these two areas was fairly low, because they all scored quite high on the pre-course test. With regard to focus words, the control group had, inexplicably, the highest increase (36.67%) followed by the songs group (32.14%) and no-songs group (26.67%). The no-songs group, however, had the highest mean percentage increase in the areas of word stress (31.35%), linking (45.98%) and connected speech (31.16%). Finally, the songs group, in addition to focus words, had the second highest increase in word stress (18.11) and connected speech (16.38). Overall, the no-songs group tended to have a higher mean percentage of improvement from pre- to post-course than the songs group in the individual pronunciation areas.

Finally, with regard to how the students performed in the individual pronunciation areas on the listening tests, although both the songs and no-songs groups generally had a higher percentage improvement than the control group, these improvements were not statistically significant according to the ANOVAs that were performed on the mean percentage increases in all of the pronunciation areas. Further results, though, are nevertheless, interesting. That is, the researcher did not set out to use statistical significance as the only measure for determining the usefulness of the course. This study also takes into consideration how the participants themselves felt about whether the use of songs could help in the perception of L2 pronunciation phenomena. Therefore, we turn now to an examination of additional data which, after reviewing, will be discussed in light of the results obtained.

Comments in Speaking Test # 2. The post-course Speaking Test # 2 (which was a test of free speech) revealed some information about what the students thought about songs in a pronunciation course and whether they can enable L2 learners to better perceive suprasegmental phenomena of the target language. Although, they were allowed to speak on any topic, seven of the ten members of the songs group decided to talk about the pronunciation course. Therefore the comments that they made were particularly revealing because they were not prompted to talk about the topics that they did. The following excerpts show that there were some students who indicated that the course helped them better perceive English suprasegmental phenomena.

“It has been very helpful for for us or for me especially, because when you’re learning a second language is very difficult to acknowledge the the the connected speech that uhm a native speaker has so uhm she has gave us uhm she have she had she have gave us uhm many useful tips for improvin’ our pronunciation skills” (S1)

“when listening a song you try to imitate the sounds and that helps you to pronounce better yourself” (S7)

“as your ear gets more used to it, to hearing the sounds isolating them and actually following truly what’s being said, you get more used to certain ways of saying certain phrases, and then you actually start understanding a whole lot more, and I think that has to do with like, the focus groups and the thought groups, and the kinda things that our class did together and they make sense when all is said together in a certain specific ways. So, I think that’s really really useful to like, know, in a more specific way now how all that works so that we can focus a little bit more on it and pay more attention” (S8)

“If you hear music in that language over and over will help to improve your listening skill” (S9)

Questionnaire results for usefulness of songs. The Initial and Final English Pronunciation Questionnaires had a variety of general statements related to how the participants felt about their accent, pronunciation and learning with songs. For these questions they answered on a six-point scale ranging from *strongly disagree* (1) to *strongly agree* (6). As well, at the end of the final questionnaires, the students were given the opportunity to provide additional comments. With regard to the statement, “Taking the pronunciation course helped improve my listening comprehension”, eight out of ten participants in the songs group strongly agreed (S2-S4, S6-S10), and two agreed (S1 and S5)⁸⁷. Another statement was, “Listening to English music can help me to learn English”. To this sentence, seven members of the songs group strongly agreed both before and after the course (S2-S4, S6-S9), two agreed (S1 and S10), while one (S5) changed his opinion from *agree* to *strongly agree* after the course. What this information reveals

⁸⁷ Incidentally, in the no-songs group, all five participants strongly agreed that taking the pronunciation course helped them with their listening comprehension.

is that *all* of the participants in the songs group believe that the pronunciation course helped their listening comprehension and that listening to songs are beneficial for learning the language. That is, that after having taken a course involving the use of songs, the participants felt that they had learned. In the comments section of the questionnaire, two of the participants referred to songs as being helpful for perceiving pronunciation phenomena and those comments are as follows:

“I learned quite a lot of things that I find to be really useful when it comes to pronunciation and noticing how native speakers talk.” (S6)

“I’d love [sic] to take this course because it helped me to understand that music can help you in the improvement of listening skill.” (S9)

Consolidation of listening data for the songs group. Since we have an idea as to how the participants in the songs group feel about the course, we will look at their individual performance on the listening tests in relation to the averages of the three groups and consider this in light of their comments.

If we look at how each one of the songs participants, and in particular those who gave their opinion, i.e. S1, S6, S7, S8, S9, performed on the listening tests, we will see that these students often performed above average for their group, and at times, above average for all of the groups. In the two tables below, Table 4.13 reminds us of the mean percentage improvement for each group while Table 4.14 indicates the improvement in each pronunciation area for the individual participants of the songs group.

Table 4.13.

Listening Tests Mean Percentage Increase for Each Group in Each Pronunciation Area

Pronunciation Area	Group		
	Control	Songs	No-Songs
Thought Groups	3.09	.50	.00
Focus Words	36.67	32.14	26.67
Word Stress	-3.99	18.11	31.35
Contractions & Reductions	2.37	1.81	6.43
Linking	39.84	25.27	45.98
Connected Speech	7.34	16.38	31.16

Table 4.14.

Listening Tests Percentage Increase in Each Pronunciation Area for the Songs Group

Pronunciation Area	Songs Group Participants									
	S1	S2	S3	S4	S5	S6	S7	S8	S9	S10
Thought Groups	.00	25.00	-20.00	.00	.00	.00	.00	.00	.00	.00
Focus Words	.00	.00	.00	28.57	80.00	133.33	42.86	16.67	-20.00	40.00
Word Stress	100.00	.00	.00	11.11	-11.00	50.00	11.11	42.86	20.00	-42.86
Contractions & Reductions	5.88	5.88	.00	.00	.00	-5.88	6.67	5.56	.00	.00
Linking	24.93	7.71	.00	6.27	-	.00	74.95	30.10	33.46	50.00
Connected Speech	.00	18.20	18.20	.00	-	8.34	20.03	71.56	11.13	.00

Above average for songs group

Above average for all groups

In Table 4.14, we can see how the members of the songs group did in each pronunciation area as well as which participants in the songs group had an above average increase for their group and for all three groups in each pronunciation area.

In thought groups, only one member (S2) performed above average for all of the groups; however, in focus words, four members (S5, S6, S7, and S10) had an above average percentage increase across the groups. In word stress, S1, S6 and S8 were above average across the groups, while S9 had an above average increase for the songs group. In contractions and reductions, only one member, S7 had a percentage increase that was above the three groups, although S1, S2, and S8 had a percentage increase above that of the songs group. In linking, S7 and S10 had percentage increases that were above average across the groups, while S8 and S9's increases were above average for the songs group. In connected speech, three members (S2, S3, and S7) were above average for the songs group and S8 had a percentage increase that was above average across the groups.

To have a clearer idea of how each songs group participant performed, we will now consider this in light of their comments as well as any class absences that appear to have negatively affected their performance.

S1, who indicated that the course gave her “many useful tips”, performed well above average for all groups in word stress and above average for the songs group in contractions and reductions. This participant, however, performed below average in the areas of connected speech and linking. Because she was absent from the first class in which these pronunciation areas were taught, her performance in these areas could be attributed to this.

Participants S2-S5 did not provide any comments that specifically related to how the pronunciation course might have affected their listening skills. Nevertheless, S2 had an above average increase across the groups in thought groups, and an above average increase for the songs group in contractions and reductions and connected speech. However, this person was

below average in linking and was absent on one of the class days in which linking was taught. However, since she did perform above average in connected speech, which was also introduced in that same class, we cannot say for sure that her absence directly affected her performance in linking.

S3, like S2, was above average for the songs group in connected speech but below it in linking. Also, although she missed two classes in which connected speech and listening were taught, it is not clear that this could have affected her performance.

S4 was not above average in any area. Although she did not miss any classes, she was quite sick for a number of them. This could have affected her performance in the listening tests.

S5 had an above average percentage increase for all groups in focus words even though he was absent from the class in which focus words were reviewed. He was also absent for the first class of connected speech and linking. However, although he was considered to be an outlier in Listening Test # 2 because he performed worse on the post-test and was the only one to do so, his performance cannot be directly attributable to his absences. Rather, as mentioned earlier in this thesis, this participant showed signs of personal or health issues.

Participant S6, who mentioned having learned “quite a lot of things” that she found useful with regard to pronunciation and “noticing how native speakers talk”, had the highest percentage improvement of all participants in focus words and was above average for all groups in word stress.

S7, who indicated that we try to imitate sounds when we listen to a song, had an above average increase for all groups focus words, contractions and reductions, and linking, in addition

to being above average for the songs group in connected speech. Clearly this person put into practice what she said. Incidentally, she was absent the day that word stress was covered and she performed below average for the songs group in word stress.

S8, a strong advocate for songs, who appreciated knowing how English suprasegmentals function together, achieved the highest increase overall in connected speech and had an above average improvement for all groups in word stress, as well as an above average improvement for the songs group in contractions and reductions and linking. Although he was absent for a class on connected speech and linking, it did not appear to affect his performance in these areas.

S9, who mentioned twice that songs can help improve one's (pronunciation) listening skills, had an above average improvement for the songs group in word stress and linking. However, this participant missed three classes, one in which word stress was taught, one in which linking was taught, and one in which focus words were taught. Since she performed below average in focus words, one might be inclined to attribute this to her absence, but doing so would not explain her above average improvement in stress and linking.

Finally, S10, who did not make any comments, had an above average percentage increase across the groups in focus words and linking. With regard to her absences, she was away for the class on focus words and one of the ones for linking. She also missed the class in which word stress was taught, and she performed the worst overall in word stress. As was the case with S9, while we might be inclined to attribute poor performance to an absence, it is difficult to explain how a participant can excel in other areas after having missed classes.

In summary, the participants who chose to talk about how they felt about how songs can help in the perception of pronunciation area, were aware that their pronunciation listening skills

had improved as a result of the course (even though at no time were they ever made aware of their test results). It is interesting to note that those who did not comment on the usefulness of the course and the songs, performed less well than those who did (with the exception of S2); S3 and S5, for instance, only performed above average in one area, and S4⁸⁸ in none. With regard to student absences, it was not clear that they had a detrimental effect on performance.

Summary of Research Question # 1b. What we can conclude from this comparison of data is that although the statistical analyses did not show that using songs in a pronunciation course can help learners to better perceive the pronunciation areas taught, we should keep in mind that statistics alone cannot always provide an accurate and comprehensive understanding of whether or not students have learned. In our case, by considering how the participants viewed the course and comparing their performance with group averages, we saw indications that the course and songs in particular, *were* beneficial for many members of the songs group. That is, the songs participants were clearly aware that songs were a material that they personally found beneficial for improving their ability to perceive pronunciation phenomena in English and this was apparent in their performance on the two listening tests.

Research Question # 2a

Speaking Tests. The second research question asks if teaching suprasegmental phenomena in a two-week pronunciation course can enable L2 learners to better produce suprasegmental phenomena of the target language. In order to determine this, two separate speaking tests were applied pre- and post-course. Speaking Test # 1 was a controlled speaking

⁸⁸ S4, incidentally, was quite sick during the course, and although she did not miss any classes, being sick could have had a negative impact on her learning.

test in which the participants read a one-page text called The Rainbow Passage, which is used to test connected speech in general. Speaking Test # 2 was a test of spontaneous speech in which the participants spoke freely on a particular topic. The performance of the control group and the treatment group (i.e. the songs and no-songs groups together) were compared for both Speaking Test # 1 and Speaking Test # 2. It is important to remember that with the speaking tests, a positive value means that the participants had an increase in the number of pronunciation errors on the post-course tests and that a negative value indicates a reduction of errors, i.e. an improvement in pronunciation. As we will see, overall, in Speaking Test #1, the treatment group had a significant mean percentage decrease in errors compared to the control group whereas for Speaking Test # 2 it was not significant. With regard to the individual pronunciation areas, the treatment group significantly reduced its mean percentage of errors in the schwa in Speaking Test # 1, and in thought groups in Speaking Test # 2.

To provide an initial summary of the results, descriptive statistics were applied in order to show the number of participants, the mean number of errors, and the standard deviation of the control and treatment group (songs + no-songs) for Speaking Test # 1 and Speaking Test # 2 pre- and post-course.

Table 4.15.

Speaking Test # 1 Descriptive Statistics for the Control and Treatment Groups

	Groups			
	Control		Treatment	
	Pre-	Post-	Pre-	Post-
<i>Number of Participants (n)</i>	8	8	15	15
<i>Mean (M)</i>	27.00	27.63	24.47	17.67
<i>Standard Deviation (SD)</i>	7.86	7.56	8.50	7.58

Table 4.16.

Speaking Test # 2 Descriptive Statistics for the Control and Treatment Groups

	Groups			
	Control		Treatment	
	Pre-	Post-	Pre-	Post-
<i>Number of Participants (n)</i>	8	8	15	15
<i>Mean (M)</i>	15.00	13.13	15.27	12.33
<i>Standard Deviation (SD)</i>	4.41	4.64	4.86	2.94

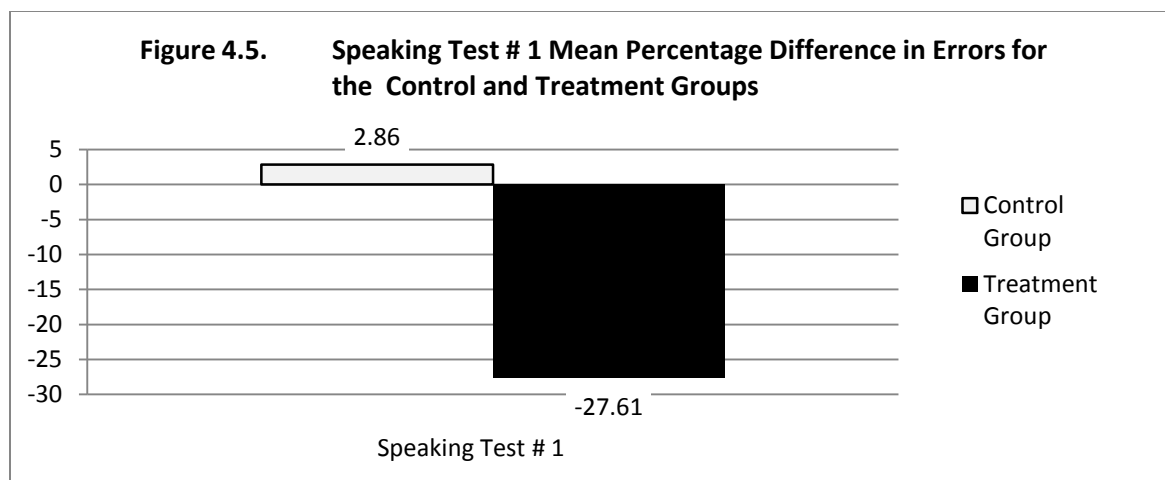
In order to know how the control and treatment groups did in comparison to each other, we will look at their average (i.e. mean) pre- and post-course errors and the degree to which they differed in each speaking test.

Table 4.17.

Speaking Test # 1 Mean Errors and Mean Percentage Difference for the Control and Treatment Groups

	Group					
	Control			Treatment		
	Pre	Post	%	Pre	Post	%
Speaking Test # 1	27.00	27.63	+2.86	24.47	17.67	-27.61

In Speaking Test # 1, the control group made more errors in the post-course test (27.63) than they did pre-course (27.00), resulting in a 2.86% increase in errors. The treatment group, on the other hand, reduced their errors from 24.47 pre-course to 17.67 post-course, which is a 27.61% decrease in the number of errors. Figure 4.5 shows these mean percentage differences for the two groups.



In order to determine whether the mean difference in errors was statistically significant, an independent-samples t-test was conducted to compare the mean percentage difference in Speaking Test # 1 for the control and treatment groups. In this case, there was a significant difference in the errors for the control group ($M = 2.86$, $SD = 6.56$) and the treatment group ($M = -27.61$, $SD = 13.51$) conditions; $t(21) = -5.97$, $p < .001$. The effect size was .63, which suggests that the significance is important. These results indicate that the pronunciation course had an effect on the ability of the participants to better produce the pronunciation phenomena when reading out loud.

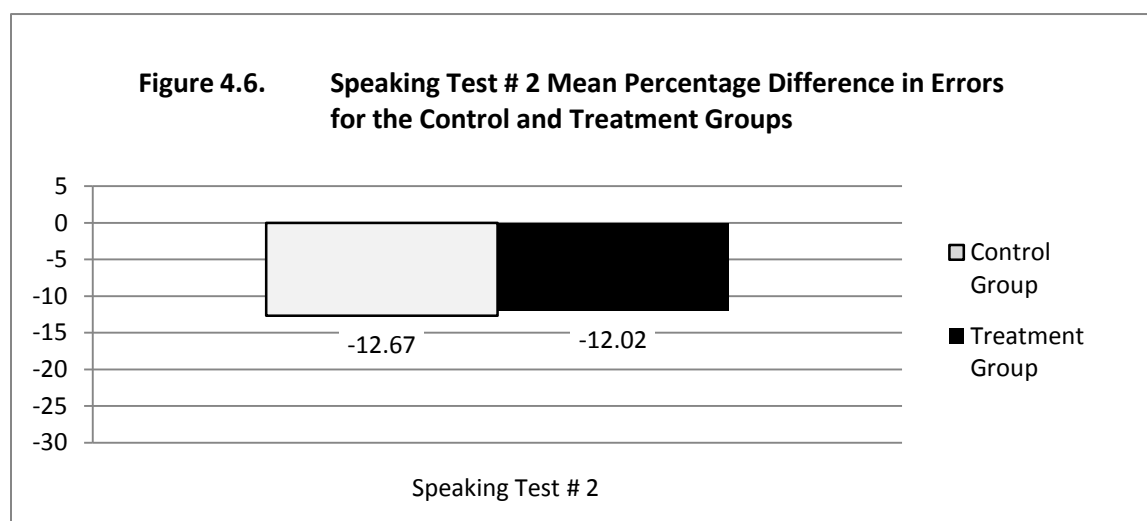
For Speaking Test # 2, Table 4.18 indicates the mean pre- and post-course errors and mean percentage difference for the two groups.

Table 4.18.

Speaking Test # 2 Mean Errors and Mean Percentage Difference for the Control and Treatment Groups

	Group					
	Control			Treatment		
	Pre	Post	%	Pre	Post	%
Speaking Test # 2	15.00	13.13	-12.67	15.27	12.33	-12.02

As we can see in Table 4.18, the control and treatment groups made a similar number of errors in both the pre- and post-course tests and thus, had similar mean percentage decreases in errors, i.e. -12.67% for the control group and -12.02% for the treatment group. Clearly this difference is not statistically significant.



Now that we have considered the results of the two speaking tests as a whole, it is necessary to see what effect the pronunciation course had in how the participants produced the individual pronunciation areas covered in the course. For both speaking tests, these are: thought groups, focus words, word stress, schwa, linking and connected speech. The following two tables show the mean percentage difference in errors in each pronunciation area for Speaking Test # 1 and Speaking Test # 2.

Table 4.19.

Speaking Test # 1 Mean Errors and Mean Percentage Difference in Each Pronunciation Area for the Control and Treatment Groups

Pronunciation Area	Group					
	Control			Treatment		
	Pre	Post	%	Pre	Post	%
Thought Groups	5.50	5.38	+13.31	3.00	2.07	-14.24
Focus Words	2.13	2.75	+79.17	1.93	1.20	-23.56
Word Stress	.75	.88	.00	1.47	1.27	-8.89
Schwa	14.25	15.00	+7.52	13.07	10.07	-18.92
Linking	4.38	3.63	-6.53	5.00	3.07	-45.80

As we can see in Table 4.19, there are pronunciation areas, e.g. focus words and word stress, in which the participants made very few mistakes in both the pre- and post-course tests. What this tells us is that focus words and word stress are areas that did not pose problems for our participants.

In thought groups, the participants also did not make a lot of errors, and although the mean percentage difference in errors from pre- to post-course for the treatment group (-14.24) was considerably different than the control group, who had a 13.31% increase in errors, the difference was not statistically significant.

With regard to the schwa, the area in which the participants made the most mistakes, the difference in the percentage of pre- to post-course errors between the two groups was significant. That is, an independent-samples t-test was conducted to compare the mean percentage difference in schwa errors in Speaking Test # 1 for the control and treatment groups. There was a significant difference in the percentage for the control group ($M = 7.52$, $SD = 13.73$) and the treatment group ($M = -18.92$, $SD = 22.32$) conditions; $t(21) = -3.04$, $p = .006$. The effect size

was .30, indicating that the magnitude of the difference was large. What these results tell us is that the pronunciation course helped the participants pronounce the schwa while reading aloud.

Finally, with respect to linking, although there was a large difference between the control groups' mean percentage decrease (-6.53) and that of the treatment group (-45.80), the difference was shown to be not significant when an independent-samples t-test was conducted. That is, for the control group ($M = -6.53$, $SD = 49.08$) and the treatment group ($M = -45.80$, $SD = 44.35$) conditions; $t(21) = -1.95$, $p = .868$.

In sum, the results of Speaking Test # 1, in our case, indicate that there were suprasegmental pronunciation areas in which the students were already competent pre-course. Nevertheless, in the area in which they were less competent, that is the schwa, the course helped the students improve their pronunciation of it.

Table 4.20

Speaking Test # 2 Mean Errors and Mean Percentage Difference in Each Pronunciation Area for the Control and Treatment Groups

Pronunciation Area	Group					
	Control			Treatment		
	Pre	Post	%	Pre	Post	%
Thought Groups	4.50	5.25	16.15	5.87	4.53	-19.05
Focus Words	.25	.00	-12.50	.47	.20	-20.00
Word Stress	.25	.13	-6.25	.13	.07	-6.67
Schwa	5.50	4.50	-13.81	5.00	4.60	.27
Linking	.75	.38	-21.88	1.53	.80	-26.67
Connected Speech	3.75	2.88	-14.14	2.27	2.13	61.86

In Table 4.20 the mean pre- and post-course errors and mean percentage difference for Speaking Test # 2 are given for the pronunciation areas covered in the course. As we can see, the

participants made very few errors in focus words, word stress, and linking, and none of the percentage differences in errors were meaningful or significant.

If we look at thought groups, we can see that the control group had an increase in errors post-course while the treatment group made fewer errors post-course. For this, an independent-samples t-test was conducted to compare the mean percentage difference in thought group errors in Speaking Test # 2. There was a significant difference in the percentage for the control group ($M = 16.15$, $SD = 24.84$) and the treatment group ($M = -19.05$, $SD = 33.34$) conditions; $t(21) = -2.61$, $p = .016$. The effect size is .24. What this means is that in our study, the pronunciation course helped the students make fewer errors in thought groups.

Regarding the schwa, despite the treatment group on average making fewer errors post-course (4.60) than pre-course (5.00), the mean percentage difference indicates that there was actually a slight increase (.27). Also, while the control group had a 13.81% mean decrease in errors, this difference between the two groups was not significant according to an independent-samples t-test. For connected speech, the situation was similar. The control group had a 14.14% mean decrease in errors post-course whereas the treatment group had an increase in errors (61.86%), even though the post-course mean (2.13) was slightly lower than the pre-course mean (2.27). Nevertheless, there was not a significant difference in the mean percentage difference of errors in the area of connected speech for the control group ($M = -14.14$, $SD = 61.35$) and the treatment group ($M = 61.86$, $SD = 151.92$) conditions; $t(21) = 1.35$, $p = .069$.

Summary of Research Question # 2a. To sum up, the overall results for Speaking Test # 1 indicated that the pronunciation course helped the students better produce the pronunciation phenomena taught when reading. The tests conducted on the individual pronunciation areas

showed that the only pronunciation area that the course seemed to help in was the schwa. The two areas in which the students were already particularly proficient at pre-course were focus words and word stress. For Speaking Test # 2, the overall percentage difference in errors between the control and treatment group suggested that the course did not benefit the participants' pronunciation in free speech. However, after analysing each pronunciation area, the results indicated that the course helped the learners in thought groups. In focus words, word stress, and linking, very few errors were made pre-course, suggesting that the participants were already strong in these areas.

Research Question # 2b

Speaking Tests. Part two of the second research question asks whether songs as a material are beneficial for producing the suprasegmental phenomena taught in a pronunciation course over a two-week period. When we look at how the three groups of participants performed on Speaking Test # 1 and Speaking Test # 2, we will see that the songs group had the largest mean percentage decrease in errors of the three groups, and in Speaking Test # 1 this decrease was significant compared to the control group. In the individual pronunciation areas, the songs group made significant gains in the schwa. In Speaking Test # 2, the improvement was significant in the area of thought groups.

An initial summary of the findings can be found in Tables 4.21 and 4.22 where we can see the number, mean, and standard deviation in the pre- and post-course errors for Speaking Test # 1 and Speaking Test # 2.

Table 4.21.

Speaking Test # 1 Descriptive Statistics for Each Group

	Group					
	Control		Songs		No-Songs	
	Pre-	Post-	Pre-	Post-	Pre-	Post-
<i>Number of Participants (n)</i>	8	8	10	10	5	5
<i>Mean (M)</i>	27.00	27.63	24.30	16.80	24.80	19.40
<i>Standard Deviation (SD)</i>	7.86	7.56	7.09	4.21	11.82	12.48

Table 4.22.

Speaking Test # 2 Descriptive Statistics for Each Group

	Group					
	Control		Songs		No-Songs	
	Pre-	Post-	Pre-	Post-	Pre-	Post-
<i>Number of Participants (n)</i>	8	8	10	10	5	5
<i>Mean (M)</i>	15.00	13.13	15.50	12.30	14.80	12.40
<i>Standard Deviation (SD)</i>	4.41	4.64	3.31	2.67	7.60	3.78

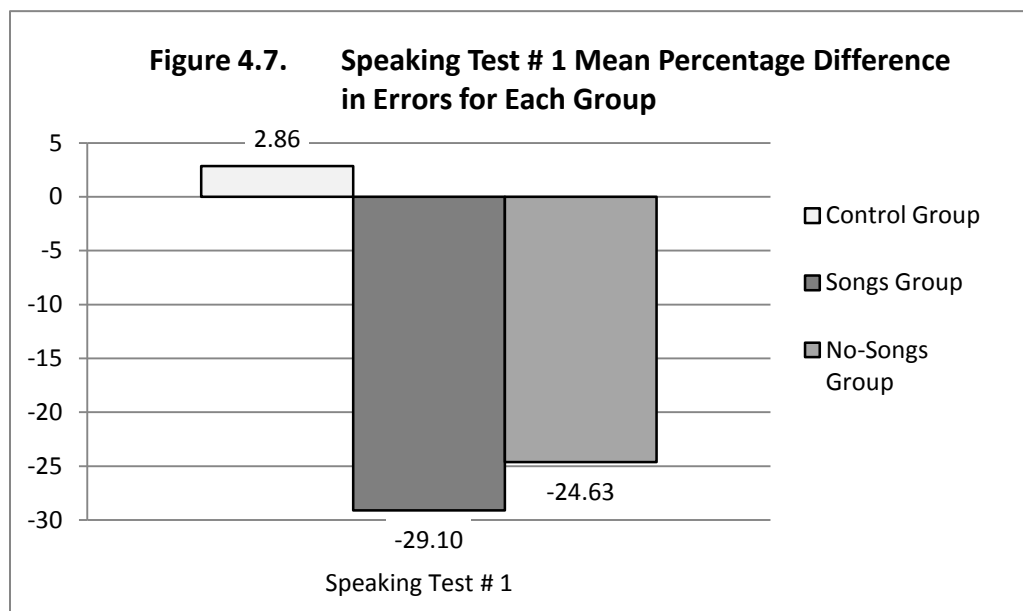
In order to know how the three groups did in comparison to each other, we will look at their average (i.e. mean) pre- and post-course errors and the degree to which they improved in each speaking test, which is what Table 4.23 shows.

Table 4.23.

Speaking Test # 1 Mean Errors and Mean Percentage Difference for Each Group

	Group								
	Control			Songs			No-Songs		
	Pre	Post	%	Pre	Post	%	Pre	Post	%
Speaking Test # 1	27.00	27.63	+2.86	24.30	16.80	-29.10	24.80	19.40	-24.63

Table 4.23 illustrates how the groups did in comparison to each other before and after the course with respect to the number of errors in Speaking Test # 1. For the control group, the pre-course average number of speaking errors was 27.00 and their post-course average was 27.63, which means they had a mean increase of 2.86% in errors post-course. The songs group's average errors dropped from 24.30 to 16.80 for a 29.10% decrease. Finally, the no-songs group had a 24.63% average decrease in the number of errors from pre-course (24.80) to post-course (19.40). The mean percentage differences are shown in Figure 4.7.



In order to know whether these differences in errors were significant, a one-way analysis of variance (ANOVA) was run and the results found in the following tables, Table 4.24, which shows the descriptives and the ANOVA values and Table 4.25, which provides the post hoc figures.

Table 4.24.

Speaking Test # 1 One-Way Analysis of Variance

Descriptives								
	N	Mean	Std. Deviation	Std. Error	95% Confidence Interval for Mean		Minimum	Maximum
					Lower Bound	Upper Bound		
Control	8	2.8610	6.55929	2.31906	-2.6227	8.3447	-5.41	12.50
Songs	10	-29.0983	14.02216	4.43420	-39.1292	-19.0675	-50.00	-7.69
No Songs	5	-24.6347	13.41369	5.99879	-41.2900	-7.9794	-38.46	-2.44
Total	23	-17.0117	18.70984	3.90127	-25.1024	-8.9209	-50.00	12.50

ANOVA						
	Sum of Squares	df	Mean Square	F	Sig.	
Between Groups	4910.812	2	2455.406	17.599	.000	
Within Groups	2790.468	20	139.523			
Total	7701.281	22				

Table 4.25.

Speaking Test # 1 Post Hoc Tests

Multiple Comparisons						
Dependent Variable: Speaking Test # 1 % decrease in errors						
Games-Howell						
(I) Group	(J) Group	Mean Difference (I-J)	Std. Error	Sig.	95% Confidence Interval	
					Lower Bound	Upper Bound
Control	Songs	31.95934*	5.00401	.000	18.7852	45.1335
	No Songs	27.49569*	6.43144	.016	6.8777	48.1136
Songs	Control	-31.95934*	5.00401	.000	-45.1335	-18.7852
	No Songs	-4.46365	7.45973	.825	-25.5459	16.6185
No Songs	Control	-27.49569*	6.43144	.016	-48.1136	-6.8777
	Songs	4.46365	7.45973	.825	-16.6185	25.5459

*. The mean difference is significant at the 0.05 level.

As can be seen in the above tables, there was a significant difference in the percentage decrease of errors between the control group ($M = 2.86$, $SD = 6.56$) and the songs group ($M = -29.10$, $SD = 14.02$) and the no-songs group ($M = -24.63$, $SD = 13.41$), $F(2, 20) = 17.60$, $p < .001$. The effect size was large (eta squared = .64). Games-Howell post hoc tests showed that the songs group had a significant decrease in errors compared to the control group, $p < .001$ as did the no-songs group compared to the control group, $p < .016$. The songs group and the no-songs group did not differ from each other significantly.

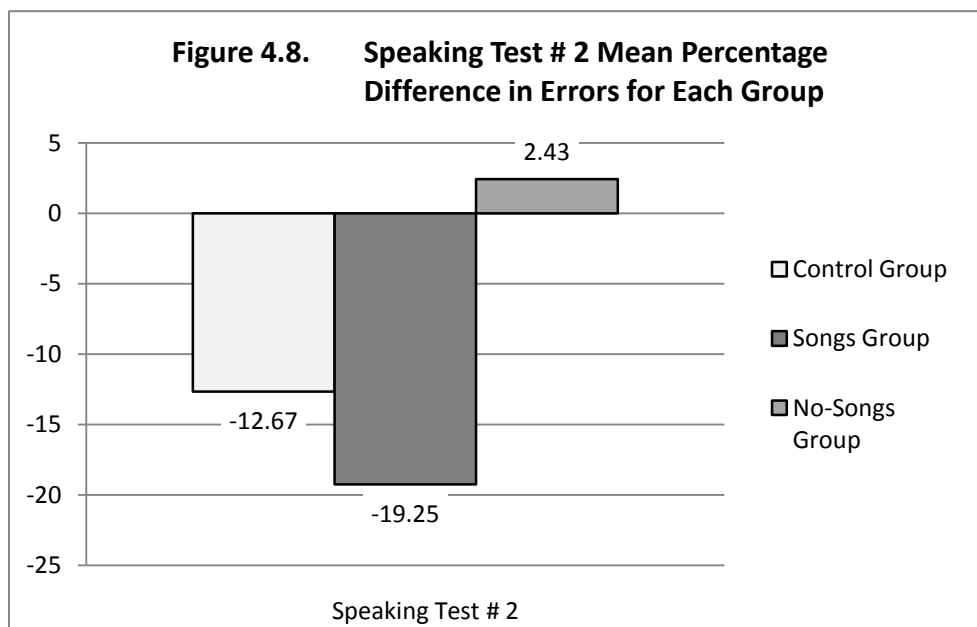
In Speaking Test # 2, there was also a difference in how the groups did in comparison to each other before and after the course with respect to the percentage of speaking errors.

Table 4.26.

Speaking Test # 2 Mean Errors and Mean Percentage Difference for Each Group

	Group								
	Control			Songs			No-Songs		
	Pre	Post	%	Pre	Post	%	Pre	Post	%
Speaking Test # 2	15.00	13.13	-12.67	15.50	12.30	-19.25	14.80	12.40	+2.43

In Speaking Test # 2, the control group's average number of errors dropped by 12.67% from 15.00 in the pre-course test to 13.13 post-course, while the songs group mean drop in errors was 19.25%, from 15.50 pre-course to 12.30 post-course. As for the no-songs group, although their average number of errors dropped from 14.80 pre-course to 12.40 post-course, this actually resulted in a mean percent *increase* of 2.43. Figure 4.8 provides a visual representation of the percent differences between the groups.



In order to know whether the differences in the percentage of errors were significant, a one-way analysis of variance (ANOVA) was run. Despite what the figure suggests, there were no

statistically significant differences in the mean percentage difference of errors between the control group ($M = -12.67$, $SD = 16.82$) and the songs group ($M = -19.25$, $SD = 15.29$) and the no-songs group ($M = 2.43$, $SD = 64.52$), $F(2, 20) = .756$, $p = .482$.

To summarize how the participants performed overall on Speaking Tests # 1 and 2, we know that in Speaking Test # 1, both the songs and no-songs group had a significant decrease in errors, whereas in Speaking Test # 2, although the songs group had the greatest decrease of the three groups (as they did in Speaking Test # 1), the amount was not significant.

Now that we have an indication as to how the three groups performed in general on the speaking tests, we will consider how they performed in the different pronunciation areas of the tests. For both speaking tests, these are: thought groups, focus words, word stress, schwa, linking and connected speech. The table below shows the average pre- and post-course errors and mean decrease in errors in each pronunciation area for Speaking Test # 1.

Table 4.27

Speaking Test # 1 Mean Errors and Mean Percentage Difference in Each Pronunciation Area for All Groups

Pronunciation Area	Group								
	Control			Songs			No-Songs		
	Pre	Post	%	Pre	Post	%	Pre	Post	%
Thought Groups	5.50	5.38	+13.31	2.00	1.20	-14.00	5.00	3.80	-14.70
Focus Words	2.13	2.75	+79.17	1.60	1.10	-23.33	2.60	1.40	-24.00
Word Stress	.75	.88	.00	1.50	1.20	-18.33	1.40	1.40	+10.00
Schwa	14.25	15.00	+7.52	14.00	10.40	-25.95	11.20	9.40	-4.85
Linking	4.38	3.63	-6.53	5.20	2.90	-40.52	4.60	3.40	-56.36

If we look at the above table, we can see that with the exception of the schwa, the groups, on average, did not make a lot of speaking errors in the different pronunciation areas. In word stress, for example, each group made an average of less than two errors. Despite this low incidence of errors, the songs group had an 18.33% reduction in word stress errors (although this difference was not statistically significant). The no-songs group had the same number of errors post-course and the control group had a slight increase in errors. Focus words, similarly, did not pose a problem for the participants as each group's average number of errors pre-course was less than three. Nevertheless, both the songs and the no-songs groups reduced their focus word errors by 23.33% and 24.00%, respectively, whereas the control group had a 79.17% increase in focus word errors. Again, the mean differences in errors between the groups were not significant.

Regarding the areas in which the students made more errors, we can see that in thought groups, the songs group had a 14.00% reduction in errors post-course and the no-songs group, a 14.70% reduction while the control group had a 13.31% increase in errors. In linking all of the groups had a reduction in errors post-course with the no-songs group having the highest mean percentage reduction in errors (56.36) followed by the songs group (40.52) and the control group (6.53). Finally, for the schwa, the songs group had the highest mean percentage decrease in errors (25.95%), followed by the no-songs group at 4.85%, while the control group had a 7.52% increase in errors.

In order to know whether these percentage decreases in errors are significant, a one-way ANOVA was run in the pronunciation areas where the participants made the most errors, namely the schwa, thought groups, and linking. The following tables, Table 4.28 and 4.29 show the descriptives and ANOVA, and the post hoc tests for the mean percentage decrease in errors for the schwa in Speaking Test # 1.

Table 4.28.

Speaking Test # 1 One-way Analysis of Variance for the Schwa

Descriptives								
	N	Mean	Std. Deviation	Std. Error	95% Confidence Interval for Mean		Minimum	Maximum
					Lower Bound	Upper Bound		
Control	8	7.5191	13.73233	4.85511	-3.9614	18.9996	-9.52	25.00
Songs	10	-25.9514	15.61089	4.93660	-37.1188	-14.7841	-50.00	.00
No Songs	5	-4.8485	28.71602	12.84219	-40.5041	30.8072	-33.33	33.33
Total	23	-9.7219	23.29895	4.85817	-19.7971	.3533	-50.00	33.33

ANOVA						
	Sum of Squares	df	Mean Square	F	Sig.	
Between Groups	5130.734	2	2565.367	7.532	.004	
Within Groups	6811.775	20	340.589			
Total	11942.509	22				

Table 4.29.

Speaking Test # 1 Post Hoc Tests for the Schwa

Multiple Comparisons						
Dependent Variable: Speaking Test # 1 schwa % difference in errors						
Games-Howell						
(I) Group	(J) Group	Mean Difference (I-J)	Std. Error	Sig.	95% Confidence Interval	
					Lower Bound	Upper Bound
Control	Songs	33.47050*	6.92402	.001	15.5831	51.3579
	No Songs	12.36756	13.72931	.662	-31.8013	56.5364
Songs	Control	-33.47050*	6.92402	.001	-51.3579	-15.5831
	No Songs	-21.10293	13.75834	.350	-65.2081	23.0022
No Songs	Control	-12.36756	13.72931	.662	-56.5364	31.8013
	Songs	21.10293	13.75834	.350	-23.0022	65.2081

*. The mean difference is significant at the 0.05 level.

As can be seen in the above tables, there was a significant difference in the mean percentage decrease of schwa errors between the control group ($M = 7.52$, $SD = 13.73$) and the songs group ($M = -25.95$, $SD = 15.61$), $F(2, 20) = 7.53$, $p = .004$. The effect size was large (eta squared = .43). Games-Howell post hoc tests showed that the songs group had a significant mean decrease in schwa errors compared to the control group, $p = .001$. The songs and the no-songs group did not differ from each other significantly nor did the control group and the no-songs group.

The results of the tests in thought groups and linking were different from the schwa. For thought groups, there were no statistically significant differences in the mean percentage decrease of errors between the three groups according to the one-way ANOVA $F(2, 20) = .769$, $p = .477$. Similarly, in linking there were no statistically significant differences in the mean percentage decrease of linking errors between the three groups according to the one-way ANOVA $F(2, 20) = 2.039$, $p = .156$.

Moving to Speaking Test # 2, the following table indicates the performance of the three groups in the various pronunciation areas.

Table 4.30.

Speaking Test # 2 Mean Errors and Mean Percentage Difference in Each Pronunciation Area for All Groups

Pronunciation Area	Group								
	Control			Songs			No-Songs		
	Pre	Post	%	Pre	Post	%	Pre	Post	%
Thought Groups	4.50	5.25	16.15	6.70	4.70	-30.90	4.20	4.20	4.67
Focus Words	.25	.00	-12.50	.40	.00	-30.00	.60	.60	.00
Word Stress	.25	.13	-6.25	.20	.10	-10.00	.00	.00	.00
Schwa	5.50	4.50	-13.81	5.50	5.10	-3.67	4.00	3.60	8.14
Linking	.75	.38	-21.88	1.20	1.00	-20.00	2.2	.40	-40.00
Connected Speech	3.75	2.88	-14.14	1.50	1.40	28.50	3.80	3.60	128.57

Looking at Table 4.30, we can see that the three groups made very few errors in most of the areas namely focus words, word stress, linking, and connected speech. What this means is that the participants in our study were already quite adept at these pronunciation areas when speaking freely. With regard to thought groups, the area in which each group made the most errors, only the songs group had a mean percentage decrease which was 30.90%. Both the no-songs group and the control group had an increase in errors post-course. With regard to the schwa, although the songs group did have a 3.67% mean decrease in errors, for some reason the control group had a higher decrease post-course (-13.81%).

As we know from Figure 4.8. above, the songs group had the highest overall mean percentage decrease in errors in Speaking Test # 2. The individual pronunciation areas in which the songs group had the greatest mean percentage decrease in errors were thought groups

(-30.90%), focus words (-30.00%) and word stress (-10.00%). Despite the small number of errors in focus words and word stress, a one-way analysis of variance was run for each area; however, the mean percentage difference in errors between the groups was not significant for either pronunciation area. Because it was clear in the table that the songs group did not outperform the other two groups in schwa, linking and connected speech, no ANOVA results are reported⁸⁹. In thought groups, though, a one-way ANOVA was run and the following tables, Table 4.31. and 4.32., show the descriptives and ANOVA, as well as the post hoc tests for the mean percentage decrease in errors for thought groups in Speaking Test # 2.

Table 4.31.

Speaking Test # 2 One-way Analysis of Variance for Thought Groups

Descriptives								
	N	Mean	Std. Deviation	Std. Error	95% Confidence Interval for Mean		Minimum	Maximum
					Lower Bound	Upper Bound		
Control	8	16.1458	24.84451	8.78386	-4.6247	36.9164	-33.33	50.00
Songs	10	-30.9048	21.71699	6.86751	-46.4402	-15.3694	-66.67	.00
No Songs	5	4.6667	42.13734	18.84439	-47.6538	56.9871	-60.00	50.00
Total	23	-6.8064	34.60657	7.21597	-21.7714	8.1586	-66.67	50.00
ANOVA								
		Sum of Squares	df	Mean Square	F	Sig.		
Between Groups		10679.907	2	5339.953	6.817	.006		
Within Groups		15667.617	20	783.381				
Total		26347.524	22					

⁸⁹ ANOVA tests were indeed run, though, to determine whether the control or no-songs group had a significant percentage decrease in errors post-course. In all three tests there was no significant difference between any of the groups.

Table 4.32.

Speaking Test # 2 Post Hoc Tests for Thought Groups

Multiple Comparisons						
Dependent Variable: Speaking Test # 2 thought groups % difference in errors						
Games-Howell						
(I) Group	(J) Group	Mean Difference (I-J)	Std. Error	Sig.	95% Confidence Interval	
					Lower Bound	Upper Bound
Control	Songs	47.05060*	11.14984	.002	17.8875	76.2137
	No Songs	11.47917	20.79104	.849	-53.0521	76.0104
Songs	Control	-47.05060*	11.14984	.002	-76.2137	-17.8875
	No Songs	-35.57143	20.05677	.268	-100.4102	29.2673
No Songs	Control	-11.47917	20.79104	.849	-76.0104	53.0521
	Songs	35.57143	20.05677	.268	-29.2673	100.4102

*. The mean difference is significant at the 0.05 level.

As we can see, there was a significant difference in the mean percentage decrease of thought group errors between the control group ($M = 16.14$, $SD = 24.84$) and the songs group ($M = -30.90$, $SD = 21.72$), $F(2, 20) = 6.82$, $p = .006$. The effect size was large (eta squared = .41). Games-Howell post hoc tests showed that the songs group had a significant mean percentage decrease in schwa errors compared to the control group, $p = .002$.

Finally, to summarize how the songs group did in Speaking Test # 1, we can say that overall, they had the largest mean percentage reduction of speaking errors post-course of the three groups and this reduction was found to be statistically significant compared to the control group. In the specific pronunciation areas that were taught in the course, the songs group had the greatest mean percentage reduction of errors in word stress and schwa. This reduction was statistically significant for the schwa but not for word stress. In thought groups and linking, although the songs group differed considerably from the control group in the mean percentage difference in post-course errors, the differences were not statistically significant. In Speaking

Test # 2, although overall the songs group also had the largest mean percentage reduction in speaking errors post-course, this decrease was not statistically significant. They had the largest mean percentage decrease in errors in thought groups, focus words, and word stress and although in focus words and word stress this difference was not significant, it was in thought groups. With regard to the schwa, linking, and connected speech, the songs group did not perform as well as the other two groups on the post-course test, but the mean percentage difference in errors between the groups was not significant in these areas. Now that we have a clear idea of the test results, we will now consider additional data which entails the participants' viewpoints on whether the use of songs could help in the production of L2 pronunciation phenomena.

Comments in Speaking Test # 2. In the post-course Speaking Test # 2, some of the participants in the songs group talked about how they felt the pronunciation course had affected their pronunciation skills. What follows are excerpts from their comments that relate to this.

“the first time we were mentioned this this project I wasn’t completely sure about if it was going to work or not, and I had my doubts about it, and then I came here and I realized that it works a lot... then I came here and I realized it’s really effective” (S2)

“I really liked the course and I think it’s gonna help me a lot” (S3)

“I think this is going to help me a lot when I’m trying to learn a song, so it helps me with the pitches and where to put the voice and so on, well not where to put the voice, but in which words I have to be very careful so they sound just right... The thing is, I like to sing, I like this course, it’s going to work well, it’s going to help me with singing” (S4)

“throughout this course I learned to pronounce better and especially how to emphasize certain words when when speaking or reading and personally I think the songs helped a lot... I think this course will help me eh in my major too because uhm because my English will be more fluent and clear and I will express myself better for example in the in the oral exams” (S7)

“this investigation proved to me that that if I that music is a really good option if you want to uhm be more fluent in the language... it’s more than that you you are always singing and always trying to uhm imitate the way that the singer sings” (S9)

Questionnaire results for usefulness of songs. To reiterate what was mentioned earlier, the pronunciation questionnaires had a variety of general statements related to how the participants felt about the course, their pronunciation, and learning with songs, as well as some other topics related to language learning. For these questions, the participants read a variety of statements and then answered on a six-point scale ranging from *strongly disagree* (1) to *strongly agree* (6).

There were some questions that were included to determine whether or not the participants viewed the course and songs as being useful for learning pronunciation. To the statement, “Taking the pronunciation course helped me improve my English pronunciation” five members of the songs group (S1, S3, S5, S8, and S9) agreed that it had and five (S2, S4, S6, S7, and S10) strongly agreed that it had. When asked, “I think my English pronunciation will continue to improve because I took this course”, five members (S1, S3, S8, S9, and S10) agreed to the statement and five (S2 and S4-S7) strongly agreed. As for the statement, “Listening to English music can help me improve my English pronunciation”, all of the songs group members initially, that is, pre-course, either agreed or strongly agreed. After the course, however, S4, S5,

and S6 changed their response from *agree* to *strongly agree*. What this information tells us is that the songs participants *unanimously* believed that the course helped and would continue to help their pronunciation and that music can assist in that process.

Other questions were asked to determine the participants' level of enjoyment related to the course as well as their motivation to improve their pronunciation. To the statement "I enjoyed the pronunciation course", all of them strongly agreed (with the exception of S5 who simply agreed). Similarly, after the course, nine of the songs group participants strongly agreed to the statement, "I am interested in improving my English pronunciation". (S8 simply agreed.) Finally, to the statement, "I think it is possible to improve my English accent", nine of the participants strongly agreed, while one (S8) simply agreed. All in all, the participants' answers to these statements indicate that after the course their level of motivation was still high and their enjoyment of the course would have contributed to that. Therefore, whether or not the results of the speaking tests showed that significant improvement had been made, the experience did serve to maintain the students' very high level of motivation to learn.

The opinions expressed above are reiterated in the comments section of the Final Pronunciation Questionnaire, in which a number of the participants mentioned that the pronunciation course was helpful for them. What follows are their specific comments.

"I received many useful tips to improve pronunciation." (S1)

"Me gustó mucho el curso y siento que en dos semanas he aprendido cosas útiles" (S2)

“it’s quite refreshing to speak with an American accent when you’ve been forced to use a British one. We were taught things I already knew but didn’t think were actual rules, I just thought I wasn’t speaking correctly.” (S3)

“I think that the workshop was very helpful” (S5)

“I learned quite a lot of things that I find to be really useful when it comes to pronunciation” (S6)

“I think this course helped me a lot with my pronunciation” (S7)

Consolidation of speaking data for the songs group. Now that we know how the participants feel, we will compare how their comments on the Final English Pronunciation Questionnaire and Speaking Test # 2 relate to their actual performance on the post-course speaking tests. That is, we will look at how their mean percentage decrease in errors related to the group’s average. Since the songs group had the highest mean percentage decrease in errors, participants who had an above average decrease in errors for the songs group, had an above average decrease across all groups.

Table 4.33.

Speaking Tests # 1 and 2 Mean Percentage Decrease in Errors for the Songs Group

	S1	S2	S3	S4	S5	S6	S7	S8	S9	S10	Group Average
Speaking Test # 1	-50.00	-43.75	-26.67	-15.79	-19.05	-16.00	-37.04	-37.50	-7.69	-37.50	-29.10
Speaking Test # 2	-16.67	-28.57	0.00	-7.69	-35.29	0.00	-36.36	-9.09	-17.65	-41.18	-19.25

Above average percentage decrease in errors across all groups

We can see in Table 4.33 that S1, S2, and S7, S8 and S10 performed above average for their group in Speaking Test # 1. In Speaking Test # 2, participants S2, S5, S7 and S10 had an above average percentage decrease in errors. Of these participants, S1, S2, S5, and S7 specifically commented that the course was helpful for them. This leads us to believe that these participants were aware that there was improvement in their pronunciation skills (even though at no time were they ever made aware of their test results). Furthermore, since all of the members of the songs group either agreed or strongly agreed that the pronunciation course helped them to improve their pronunciation, this is consistent with the fact that all of them showed some degree of improvement on the speaking tests. That is, everyone made fewer errors in Speaking Test # 1 and all but two (S3 and S6) made fewer errors in Speaking Test # 2.

Now that we have a clearer idea as to how the participants of the songs group performed on average in the speaking tests, we will look at how they did in each pronunciation area and compare this to their comments. Table 4.34 reminds us of the mean percentage difference in errors for each pronunciation area from pre- to post-course for all of the groups, while Table 4.35 provides the percentage difference in errors for the members of the songs group. The pronunciation areas in which they performed above average for their group as well as all groups is indicated through shading.

Table 4.34

Speaking Test # 1 Mean Percentage Difference in Errors in Each Pronunciation Area for All Groups

Pronunciation Area	Group		
	Control	Songs	No-Songs
Thought Groups	+13.31	-14.00	-14.70
Focus Words	+79.17	-23.33	-24.00
Word Stress	.00	-18.33	+10.00
Schwa	+7.52	-25.95	-4.85
Linking	-6.53	-40.52	-56.36

Table 4.35.

Speaking Test # 1 Percentage Decrease in Errors in Each Pronunciation Area for the Songs Group

Pronunciation Area	Songs Group Participants									
	S1	S2	S3	S4	S5	S6	S7	S8	S9	S10
Thought Groups	-100.00	-40.00	.00	.00	.00	100.00	.00	.00	.00	-100.00
Focus Words	-50.00	.00	.00	-33.33	.00	-50.00	.00	-50.00	.00	-50.00
Word Stress	-33.33	-50.00	.00	.00	.00	.00	-100.00	.00	.00	.00
Schwa	-21.05	-50.00	-33.33	-16.67	-41.67	-26.67	-25.00	-38.46	-6.67	.00
Linking	-90.91	-40.00	-100.00	.00	.00	.00	-100.00	.00	-14.29	-60.00

Above average percentage decrease in errors than both the control and no-songs groups

Above average percentage decrease in errors for both the songs group and across all groups

Table 4.35 provides for an understanding of how the members of the songs group performed in Speaking Test # 1 that simple group averages cannot. The cells that are lightly shaded indicate that the members had a greater decrease in errors than the control and no-songs group averages. The cells that are more darkly shaded indicate that the members had a greater

than average percentage decrease in errors across the groups. That is, not only did they outperform their group on average, but by default, all of the groups due to the fact that the songs group had on average the greatest mean percentage decrease in errors of all the groups.

In the area of thought groups, three members (S1, S2, and S10) had a higher than average percentage reduction in errors across the groups. In fact, two of them (S1 and S10) had a 100% reduction in their thought group errors. S4 did not make any thought group errors. Other members (S3, S7, S8, and S9) did not have a change in thought group errors from pre- to post-course, while S5 actually made one error post-course, as did S6.

In focus words, five members (S1, S4, S6, S8 and S10) had a higher than average percentage decrease in errors across all groups, while the other five (S2, S3, S5, S7, and S9) did not have a change in errors.

In word stress, S1, S2, and S7 had higher than average percentage decreases across the groups. Six members (S3-S6 and S9-S10) had no change in word stress errors from pre- to post-course, while S8 did not make any errors pre- or post-course.

For the schwa, we can see that all but one member had a greater percentage decrease in errors across all groups. This member, S10, did not have a change in errors whereas the other nine had a decrease in errors. The ones who had an above average decrease in errors for the songs group were: S2, S3, S5, S6, and S8.

For linking, S1, S3, S7 and S10 had an above average percentage decrease in errors, two (S2 and S9) had a below average decrease in errors for their group and the no-songs group. S4, S5, and S6 had no change in errors from pre- to post-course, while and S8 had no errors in the pre- or post-course test.

Finally, by looking strictly at the shading in Table 4.35, we can see that S1 had an above average percentage decrease in errors in all five pronunciation areas. For S2, S7, and S10 it was in three out of five areas, whereas S3, S4, S6 and S8 were above average in two. S5 and S9 were only above average in one area.

Before discussing these results in light of the comments the students made, it is necessary to see how the individual members of the songs group did in Speaking Test # 2. Table 4.36 reminds us how the three groups performed in the pronunciation areas and Table 4.37 shows, via shading, how the songs group participants did compared to the other groups.

Table 4.36.

Speaking Test # 2 Mean Percentage Difference in Errors for Each Pronunciation Area for All Groups

	Group		
Pronunciation Area	Control	Songs	No-Songs
Thought Groups	16.15	-30.90	4.67
Focus Words	-12.50	-30.00	.00
Word Stress	-6.25	-10.00	.00
Schwa	-13.81	-3.67	8.14
Linking	-21.88	-20.00	-40.00
Connected Speech	-14.14	28.50	128.57

Table 4.37.

Speaking Test # 2 Percentage Decrease in Errors in Each Pronunciation Area for the Songs Group

Songs Group Participants										
Pronunciation Area	S1	S2	S3	S4	S5	S6	S7	S8	S9	S10
Thought Groups	.00	-20.00	-42.86	-50.00	-42.86	-30.00	.00	-66.67	-40.00	-16.67
Focus Words	.00	.00	.00	.00	.00	-100.00	-100.00	.00	.00	-100.00
Word Stress	.00	-100.00	.00	.00	.00	.00	.00	.00	.00	.00
Schwa	-42.86	-50.00	75.00	.00	20.00	60.00	-12.50	-12.50	-16.67	-57.14
Linking	.00	.00	-100.00	-100.00	-100.00	.00	-100.00	.00	200.00	.00
Connected Speech	.00	.00	.00	400.00	-75.00	.00	.00	.00	-40.00	.00
Above average percentage decrease in errors than both the control and no-songs groups										
Above average percentage decrease in errors for both the songs group and across all groups										

In Table 4.37, we can see that in thought groups, five members (S3, S4, S5, S8, and S9) had an above average percentage decrease in errors for the songs group, which had the greatest percentage decrease of all of the groups. The other five members (S1, S2, S6, S7, and S10) performed better than the average of the control and no-songs groups.

In focus words, seven of the songs group members (S1-S5, S8-S9) made no errors pre- or post-course. However, S6, S7, and S10 had a 100% decrease in errors, which was greater than the highest percentage decrease across the groups.

For word stress, eight members (S1, S4-S10) did not have any errors in the pre- or post-course test. The one who did, S2, had a 100% reduction in errors, which was greater than the

highest percentage decrease across the groups. S3, actually made an error post-course, and thus did not have a percentage decrease in errors⁹⁰.

In the schwa, an area in which the songs group did not have on average the greatest percentage decrease in errors, S1, S2, S9, and S10 had an above average decrease in errors across the groups. S7 and S8 had an above average decrease for the songs group. Three members actually made more errors (S3, S5, and S6) post-course while one (S4) had no change in errors.

In linking, one member (S9) had an increase in errors post course. S3, S4, S5, and S7 reduced their errors by 100%, which was an above average percentage decrease across all of the groups. Three members (S1, S2, and S6) had one error in the post-course test, one (S8) had two, whereas S10 had no change in errors from pre- to post course.

In connected speech, S5 and S9 had an above average decrease in errors across the groups, S3, S6, and S7 had no change in errors from pre- to post-course, which was above average for the songs group. Similarly, four members (S1, S2, S8, and S10) made no errors in the pre- or post-course tests. S4, actually made more errors in the post-course test.

Finally, by looking at the shading in Table 4.37, it is clear that *everyone* in the songs group were above average in thought groups. In addition, we can see that S5 and S9 had an above average decrease in errors across the groups in three of the six pronunciation areas. S7 had an above average decrease across the groups in two areas, above the control and no-songs group in one, and (although not shown with shading) was above average for the songs group in the schwa. S3, had an above average decrease across the groups in two, and was above average for the songs group in connected speech. S2, S4, and S10 had an above average decrease in two

⁹⁰ The percentage increase cannot be calculated with a pre-course error count of zero.

areas across the groups, while S1 only had one. S6 and S8 were above average in one area across the groups and above average for the songs group in one as well, for S8 this was the schwa.

In order to have a better understanding of how the participants' performance on the speaking tests related to their opinions and comments, we will discuss the information participant by participant. Since all members of the songs group agreed or strongly agreed that the pronunciation course help their pronunciation and that songs can assist in that, this information will not be included in the ensuing discussion. What will be added, however, are any class absences that the students had that corresponded with a less-than-average performance in the pronunciation areas covered in those classes.

S1 mentioned having received "many useful tips" for improving her pronunciation. Overall in Speaking Test # 1, this participant had an above average decrease in errors across the groups, with the decrease being above average across the groups in thought groups, focus words, word stress, and linking, and above average for the control and no-songs group in the schwa. In Speaking Test # 2, she had an above average decrease in errors in thought groups and the schwa. This participant's recognition and application of the "tips" clearly resulted in an improvement in her pronunciation.

S2 said that the "project" worked and that it's "really effective" and added that she felt that in two weeks, she had learned a lot of useful things. In both speaking tests, she had an above average decrease in errors across the groups. In both speaking tests, she was above average in thought groups, word stress, and schwa.

S3 said that she liked the course and thought that it was going to help her a lot. She also indicated that she was pleased to be able to speak with an American accent and have the patterns of her accent validated. Overall, this participant did not have an above average decrease in errors in the two speaking tests, although in Speaking Test # 1, she did have an above average decrease in errors in the schwa and linking. For Speaking Test # 2, it was in thought groups and linking, and in connected speech her performance was above average for the songs group.

S4 felt that the course would help her when singing and specifically with the pronunciation of certain words. Although she did not have an above average decrease in errors overall in either of the speaking tests, she did have an above average decrease in focus words and the schwa in Speaking Test # 1, and thought groups and linking in Speaking Test # 2. It is interesting to notice that her performance in focus words directly relate to her comments.

S5 simply commented that “the workshop was very helpful”. Overall in Speaking Test # 2, he had an above average decrease in errors, which translated to an above average decrease in thought groups, linking, and connected speech. In Speaking Test # 1, he had an above average decrease in errors in the schwa, but in Speaking Test # 2, he actually had an increase in schwa errors. Because his two absences corresponded with the classes in which the schwa was taught and reviewed, his lack of improvement of the use of schwa in spontaneous speech could be attributed to this; however, it would then be difficult to understand how his use of schwa improved in Speaking Test # 1.

S6 said that she had learned “quite a lot” of useful things for pronunciation. Overall, this participant did not perform above average on the two tests, but she did have an above average

decrease in errors in focus words and the schwa in Speaking Test # 1, and in thought groups and focus words in Speaking Test # 2.

S7 said that the course helped her a lot with the pronunciation and “especially how to emphasize certain words”. She also felt that her English would be “more fluent and clear”. Overall, S7 had an above average decrease in errors in both speaking tests. In Speaking Test # 1, this translated into performing above average in word stress, schwa and linking. In Speaking Test # 2, she was above average across the groups in focus words and linking, and above average for the songs group in schwa. It is interesting to note how her performance in the tests is consistent with her comments, especially with regard to focus words and linking.

S8 did not make any specific comments about how the course might specifically affect pronunciation. This person did, however, have an above average decrease in errors in Speaking Test # 1, which specifically were in focus words and the schwa. In Speaking Test # 2, his decrease in errors in thought groups was above average across the groups, while in the schwa it was above average for the songs group.

S9 mentioned that music can help you become more “fluent”, and although overall she did not have an above average decrease in errors, she did have it in the schwa in Speaking Test # 1 and in thought groups, schwa, and connected speech in Speaking Test # 2.

S10 also did not comment on the course. Nevertheless, overall she had an above average decrease in errors in both speaking tests. In Speaking Test # 1, this above average decrease was in thought groups, focus words, and linking; in Speaking Test # 2, it was in thought groups, focus words and the schwa.

Summary of Research Question # 2b. This comparison of data for the songs group allowed us to better understand how the individual participants performed in relation to the other two groups and how their feelings about the pronunciation course related to their performance. While descriptive and inferential statistical analyses were able to provide a general idea of how the songs group performed in comparison to the control and no-songs group, namely that they improved significantly in the schwa and thought groups, a detailed look at the performance of the songs group participants indicated that many of them improved in other areas beyond that of the average for the other groups. Finally, except in the case of possibly one participant, the songs group's absences from class did not appear to negatively affect their pronunciation in the areas taught on those days.

Additional Findings

Although this study was most interested in understanding how songs might benefit students who are learning pronunciation, the performance and opinions of the no-songs group were also collected.

No-Songs Group Results and Opinions

Tables 4.38 and 4.39 show how the members of the no-songs group improved their receptive and productive pronunciation skills from pre- to post-course. As we saw above, the no-songs group on average performed better post-course than the songs group in listening, but not as well as the songs group or control group in speaking.

Table 4.38

Listening Tests # 1 and 2 Mean Percentage Increase for the No-Songs Group

	NS1	NS2	NS3	NS4	NS5	Group Average
Listening Test # 1	0.00	11.35	10.38	10.70	16.94	9.88
Listening Test # 2	4.30	8.38	34.93	66.75	73.40	37.55

Above average for the control and songs group

Above average for all groups

In Table 4.38 we can see that in Listening Test # 1, three of the no-songs group members performed above average for all groups, while in Listening Test # 2, this was the case with two participants, with one participant performing above average for just the no-songs group.

On the Final Pronunciation Questionnaire, everyone in the no-songs group chose “strongly agree” for the statement, “Taking the pronunciation course helped improve my listening comprehension”. As we can see, this is generally consistent with their actual performance in both Listening Test # 1 and 2.

Table 4.39

Speaking Tests # 1 and 2 Mean Percentage Decrease in Errors No-Songs Group

	NS1	NS2	NS3	NS4	NS5	Group Average
Speaking Test # 1	-30.00	-25.00	-27.27	-2.44	-38.46	-24.63
Speaking Test # 2	-27.27	116.67	-35.71	-11.11	-30.43	+2.43

Above average for control and no-songs group

Above average for all groups

In Table 4.39, we can see that in Speaking Test # 1, two members had a greater percentage decrease in errors across the groups, while two had this for their group only. In Speaking Test # 2, three members performed above average across the groups while one was above average for just the no-songs group. The table also indicates that the very high percentage increase of errors that participant NS2 made in Speaking Test # 2 greatly affected the group average.

With regard to how the students felt about their pronunciation, four members of the group strongly agreed with the statement, “Taking the pronunciation course helped me improve my English pronunciation”, while one simply agreed. Again, these responses are in line with how the students actually performed on the tests, with the possible exception of NS2 in Speaking Test # 2.⁹¹ Moreover, all of the participants strongly agreed to the statement, “I think my English pronunciation will continue to improve because I took this course”. Furthermore, all of the members of the no-songs group strongly agreed to the statement “I enjoyed the pronunciation course” and “I am interested in improving my English pronunciation”. What these responses tell us is that the students’ pre-course (high) level of motivation was maintained.

In the post-course Speaking Test # 2, the only member of the no-songs group to talk about the course was NS4. This participant performed above average across the groups in both of the post-course listening tests and above average for the no-songs group in Speaking Test # 2. Although the participant does not specifically comment on her personal performance, she does indicate that she felt that she would have learned more if she had been part of the songs group.

⁹¹ While it appears that NS2 performed poorly on Speaking Test # 2, with a 116.67% increase in errors, this figure is a little bit misleading. This participant had the lowest number of errors (six) of all of the participants in the study in the pre-course test. Post-course she made 13 errors, which was only slightly above average for the no-songs group who made 12.4 errors on average.

Nevertheless, as was the case with the members of the songs group, this person's comments about the course are also positive:

"I think it was very interesting to me because I really like to see the different sound uh that exist in different languages uh and I like to compare the native speakers and non-native speakers uhm and how we have to do the sounds correctly uh for them to understand what I am saying so uh I would have loved being in the songs group because I love singing and I think I would understand better a lot of things but I enjoyed this course as well uh I loved eh knowing how to contract words and and the elision of the sounds that native speakers do and uh and uh the different intonations and the differences between between uhm different uh intonations I mean the differences between a native speaker from the US and from Canada the example of Toronto and Torono I think it was very fun and I really enjoyed it a lot, and I think it is very good for for us that we are studying English and I think it is very very important for uhm for us to to learn how to how to make these thought groups for native people understand what we are actually saying"
(NS4)

With regard to the comments that the members of the no-songs group made on the final questionnaires, as with the songs group, their comments also are all positive.

"Even though I did not participate in the "Songs Group" I do believe that I've improved a lot in these last two weeks. The activities were great, funny and really helpful. I must admit that, at first, I wanted to be in the "S. G." But now I can say that I enjoyed so much this course, it was useful and amazing. Thank you for the opportunity to be part of it, and to realize that there are more ways of learning English that don't specifically use songs." (NS 1)

“I really enjoyed this course, I think it was really helpful because now I’m able to identify what are the mistakes I make when I speak, and I learnt some techniques to sound more like an american native speaker. Thank you very much for this opportunity, you are a really nice person and I hope everything goes great with your investigation ☺” (NS2)

“I really enjoyed the course. It was fun, and I liked talking to you! haha I hope everything comes out well and I’m really looking forward to know the results so when you finish analyzing send us the results. Thanks for everything!! It was helpful ☺” (NS3)

“Siento que este curso me ayudó mucho a entender cómo un native speaker puede entendernos mejor cuando hablamos, a través de thought groups and different intonations. También me gustó trabajar con Karen porque es muy valioso para nosotros que un native speaker nos corrija y nos ayude a mejorar nuestra pronunciation.” (NS4)

“It was an enjoyable and really useful course. Maybe I didn’t learn or improve that much, but now I am aware of things I wasn’t before and I will try to use them from now on. Thank you! ☺” (NS5)

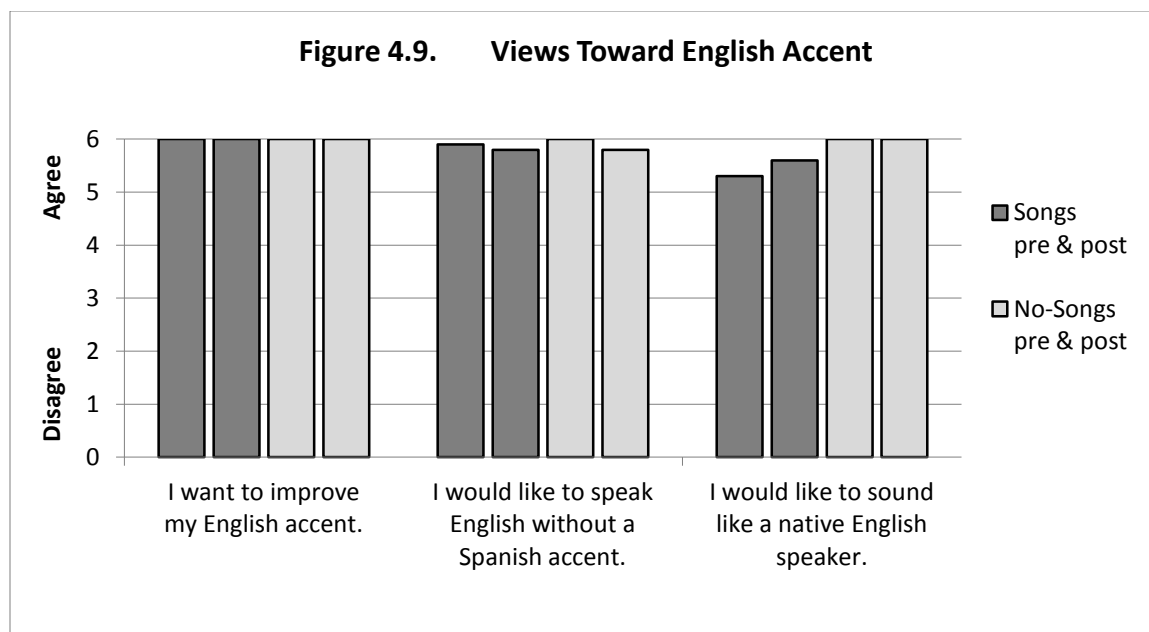
In summary, after looking at how the no-songs group felt about the pronunciation course and how well they performed in general in the listening and speaking tests, we can say that taking the course was beneficial for them.

Questionnaire Results

The Initial and Final English Pronunciation Questionnaires had a series of statements related to accent and musical aptitude. As mentioned earlier in the thesis, the responses to these

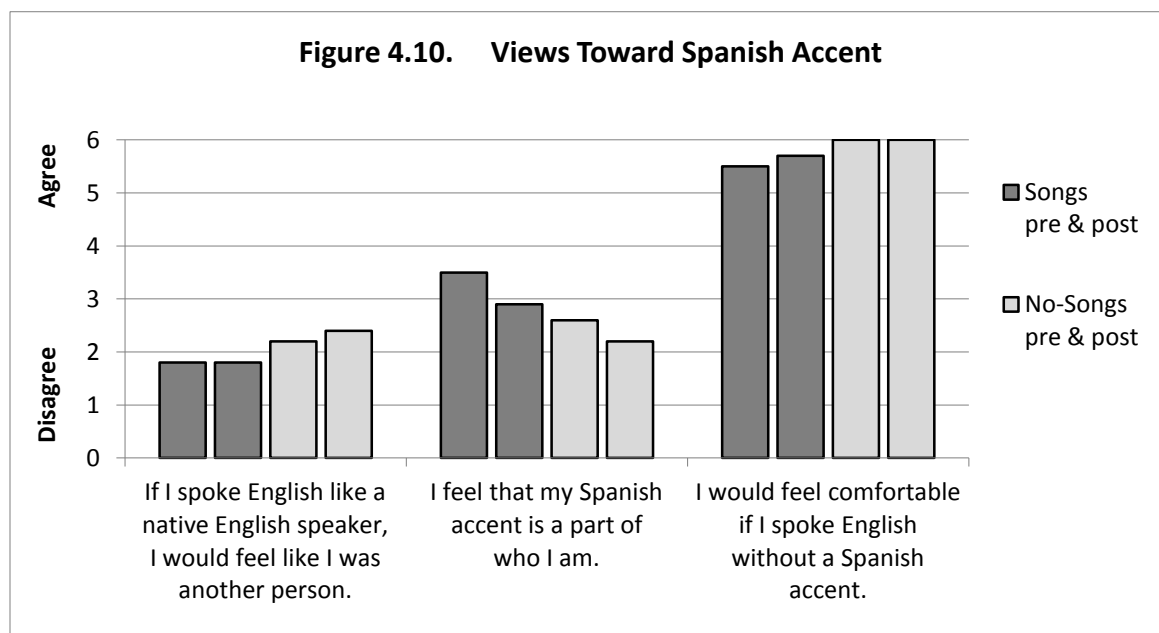
statements were used to help distribute the participants among the groups, but they also serve to reveal their personal views and goals both before and after the course.

Participants' Opinions on Accent. There were some statements in the English Pronunciation Questionnaires that served to find out how the participants felt about the English accent. One of the statements was “I want to improve my English accent”, to which all of the songs and no-songs participants strongly agreed both pre- and post-course. Another statement was “I would like to speak English without a Spanish accent”. To this, most of the songs group participants strongly agreed both pre- and post-course, however, two changed their answer to *agree* post-course while one changed it to *strongly agree* post-course. In the end, eight songs group participants strongly agreed to this statement. For the no-songs group, all participants strongly agreed both pre- and post-course, with the exception of one who changed to *agree* post-course. To the statement “I would like to sound like a native English speaker”, at pre-course six members of the songs group strongly agreed, three agreed and one slightly agreed. Post-course, two went from *agree* to *strongly agree* and one from *slightly agree* to *agree*, while another member went from *strongly agree* to *slightly agree*. The no-songs group unanimously strongly agreed both pre- and post-course. Overall, what the results to these three questions tell us is that, after the course, the participants still want to improve their English accents. As well, they still want to speak English without a Spanish accent, although they are slightly less convinced of this than they were at pre-course. At the same time, though, they want to sound like a native English speaker. Figure 4.9 displays the pre- and post-course average results to these questions for the songs and no-songs groups.

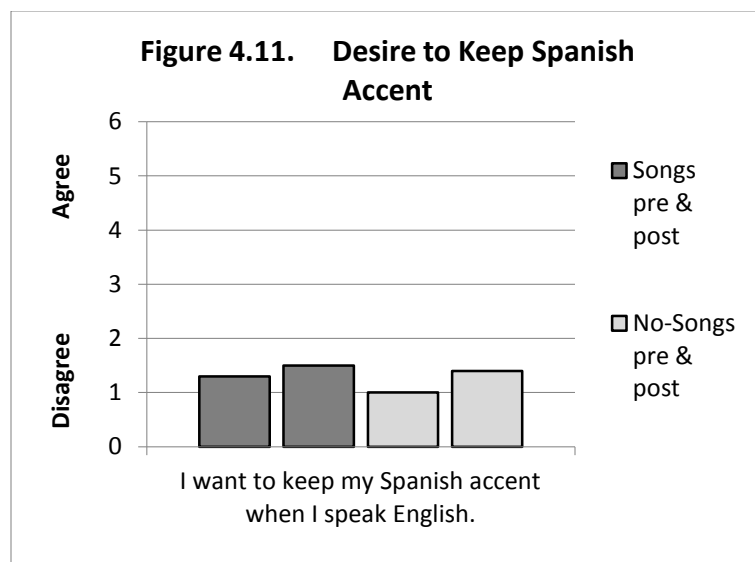


Another series of statements were concerned with what kind of influence accent might have on the participants. One statement was “If I spoke English like a native English speaker, I would feel like I was another person”. Of the members of the songs group, nine either strongly disagreed or disagreed pre- and post-course. The no-songs group responses were more varied both pre- and post-course. (The only notable changes in opinion were that one member went from *slightly agree* to *strongly disagree* after the course, while another went from *disagree* to *agree*.) What their responses to this statement suggest is that overall, post-course, the students (with the exception of one) do not feel that their accents are part of their identity. However, the participants’ answers to “I feel that my Spanish accent is part of who I am” indicate that they do feel somewhat tied to it, because their responses are varied. In the songs group, five members disagreed or slightly disagreed pre-course, while the other five agreed to one degree or another. Post-course, three either slightly agreed or agreed, whereas seven disagreed to one degree or another. The no-songs group, on the other hand, had four members who slightly disagreed pre-

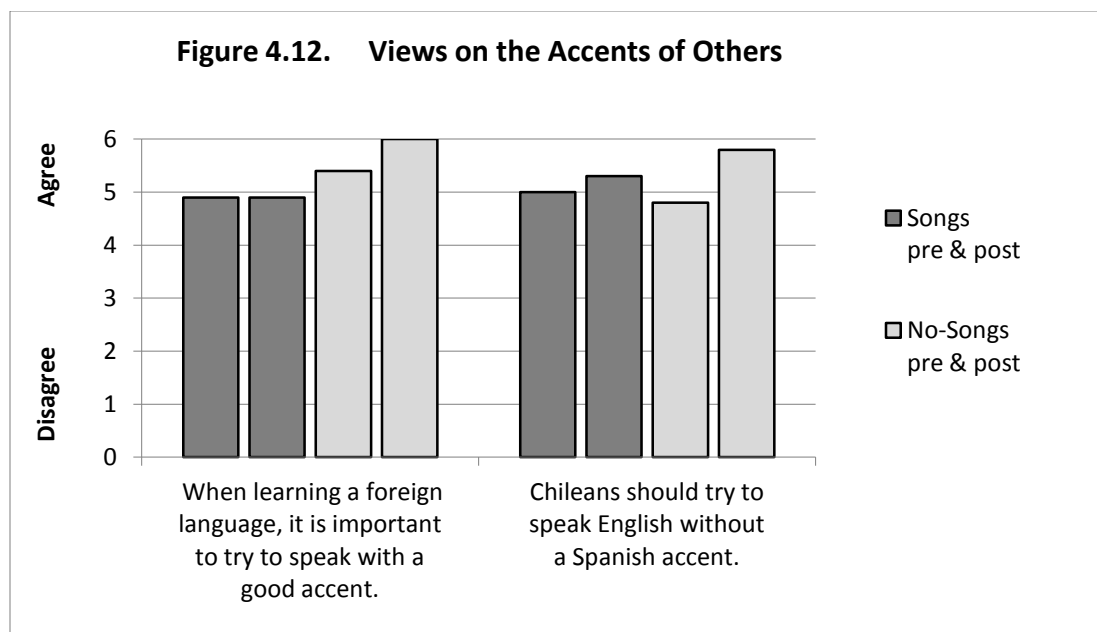
course while one strongly disagreed. Post-course, the notable changes were that one member went from *slightly disagree* to *agree*, while two went from *slightly disagree* to *strongly disagree*. Overall, these results suggest that after the course neither group felt quite as strongly about their Spanish accent being a part of them. The participants' responses to the statement "I would feel comfortable if I spoke English without a Spanish accent" seem to support this view. Pre-course six members of the songs group strongly agreed to the statement, whereas three agreed and one slightly agreed. Post-course eight strongly agreed, one agreed and one slightly agreed. All of the members of the no-songs group strongly agreed both pre- and post-course. In general, the groups' responses to these statements suggest that the participants do not feel that their Spanish accent is particularly tied to their identity. Figure 4.10 depicts the average group responses to these questions.



Finally, the responses to the following statement “I want to keep my Spanish accent” indicate that the participants do not want to have a Spanish accent when speaking English. That is, pre- and post-course, all of the songs participants disagree to one extent or another. The no-songs group unanimously strongly disagree pre-course, while post-course one member changed to *slightly disagree*.



Finally, there were two statements which were the following: “When learning a foreign language, it is important to try to speak with a good accent” and “Chileans should try to speak English without a Spanish accent”. These were included help determine whether the participants felt differently about the accent of people other than themselves. To the first question all but two members of the treatment group agreed more post-course. To the second question, all but one either maintained or increased their level of agreement post-course. Therefore, even with respect to other people, the participants unanimously feel that accent is important.



Participants and Musical Aptitude. There were a few statements in the Initial English Pronunciation Questionnaire that required the participants to assess themselves regarding their musical aptitude. The statements were “I am a good singer”, “I consider myself a musical person” and “I play one or more musical instruments”. With regard to the statements, we know that it is possible for people to consider themselves musical and play an instrument even though they are poor singers. Similarly, people who can sing well would consider themselves musical even though they cannot play instruments. While each of these scenarios involves separate musical abilities, in both cases, auditory skill is one of them. Therefore, in order to calculate a musical aptitude score, the two statements with the highest ratings were averaged and rounded off. Then the musical aptitude level was determined based on the score. Scores that were between one and two were considered *low*, those between three and four were *mid*, and those between five to six were *high*. The participants’ statement ratings, musical aptitude score, and musical aptitude level can be seen in Table 4.40.

Table 4.40.

Musical Aptitude

Participant	I am a good singer.	I consider myself a musical person.	I play one or more musical instruments.	Musical Aptitude Score	Musical Aptitude Level
S1	1	5	1	3	mid
S2	4	6	1	5	high
S3	5	6	1	6	high
S4	5	5	2	5	high
S5	1	1	1	1	low
S6	1	5	2	4	mid
S7	2	6	1	4	mid
S8	1	2	1	2	low
S9	4	2	1	3	mid
S10	4	4	4	4	mid
NS1	2	5	4	5	high
NS2	3	3	1	3	mid
NS3	4	4	5	5	high
NS4	6	6	5	6	high
NS5	1	1	1	1	low

In order to know how the participants' musical aptitude related to their performance on the listening and speaking tests, the percentage difference between their pre- and post-course tests for both of the listening and speaking tests were averaged and then ranked from 1 to 15, the total number of participants in the treatment group. Then the ranking was converted to a

descriptive term, i.e. *low*, *mid*, *high*. That is, participants ranked between one and five, were *high*, those between six and ten were *mid*, and those between eleven and fifteen were *low*. In this way we were able to compare their performance with their self professed musical aptitude level. Table 4.41 displays this information.

Table 4.41.

Participant Ranking and Musical Aptitude

Participant	Listening		Speaking		Musical Aptitude Level
	Rank	Level	Rank	Level	
S1	6	mid	5	high	mid
S2	10	mid	3	high	high
S3	15	low	10	mid	high
S4	12	low	12	low	high
S5	13	low	8	mid	low
S6	7	mid	13	low	mid
S7	4	high	2	high	mid
S8	3	high	9	mid	low
S9	8	mid	11	low	mid
S10	9	mid	1	high	mid
NS1	14	low	7	mid	high
NS2	11	low	15	low	mid
NS3	5	high	6	mid	high
NS4	2	high	14	low	high
NS5	1	high	4	high	low

As we can see in Table 4.41, there does not appear to a clear connection between self-professed musical aptitude and success in learning receptive and productive pronunciation skills. For example, NS5 had low musical aptitude, yet this participant had the highest percentage improvement on the listening tests, and ranked fourth (high) for improvement in speaking. S7, who had mid musical aptitude, was a high performer in both listening and speaking. S4, who had high musical aptitude, was a low performer in both listening and speaking. Nevertheless, there is some degree of consistency between musical aptitude and listening: S1, S5, S6, S9, S10, NS3, and NS4 all had a listening performance level and musical aptitude that matched. This did not carry over into speaking performance, though: only S2 had a speaking performance and musical aptitude that matched. The fact that musical aptitude seems to relate somewhat to listening skills but not to speaking skills is potentially of interest, but overall we cannot conclusively determine that there was a relation between our participants' self professed musical aptitude and their pronunciation performance.

Summary of the Findings

This chapter began by determining whether the pronunciation course had an effect on the participants' ability to perceive pronunciation phenomena irrespective of the materials used. For this, the songs and no-songs groups were considered together as the treatment group. Overall, with respect to listening, the treatment group performed better than the control group on both listening tests; however, the percentage increase was not statistically significant. When each pronunciation area was evaluated separately, it was found that the treatment group definitely

performed significantly better in word stress, and almost significantly in connected speech. Then the treatment group was split into the songs and no-songs groups to enable us to determine whether songs as a material can better enable students to perceive pronunciation phenomena. The results were inconclusive. On the one hand, the no-songs group performed on average better than the songs group on the listening tests. Nevertheless, there were some songs group participants who performed higher than average across the groups in different pronunciation areas, especially focus words and word stress. Next we set out to determine whether the pronunciation course had an effect on the participants' ability to produce pronunciation phenomena irrespective of the materials used. Again, the songs and no-songs groups were combined as the treatment group. The results showed that for Speaking Test # 1, the treatment group performed significantly better overall; the same was the case with the schwa. In Speaking Test # 2, although the treatment group's percentage decrease in errors was greater than that of the control group, overall, the difference was not significant. It was significant, however, in the area of thought groups. After this, once again the treatment group was separated into the songs and no-songs groups to see the effect that songs might have in the pronunciation of the participants. Overall, the songs group performed significantly better than the control group in Speaking Test # 1. Their percentage decrease in errors for the schwa was significant in Speaking Test # 1, whereas the decrease in the area of thought groups was significant in Speaking Test # 2. As was the case with the listening tests, there were songs group members who performed well above average, particularly in focus words, schwa, and linking in Speaking Test # 1, and in thought groups, schwa, and linking in Speaking Test # 2. In addition to the quantitative results, the students made comments regarding the course and the songs group felt very strongly that the pronunciation course helped them and that songs played a role. Finally, there were some additional findings that, while not part of the research questions themselves, were, nevertheless, interesting to report on. It was found that the

no-songs group also felt that they benefitted from the course and most of them performed above average in the tests. Lastly, there were some statements from the Initial and Final Pronunciation Questionnaires that sought to better understand how the songs and no-songs participants felt about their pronunciation, accent, and musical aptitude. Both pre- and post-course both groups felt strongly about wanting to improve their accent, speaking without a Spanish accent, and sounding like a native English speaker. Furthermore, they indicated that they would feel comfortable with this, with results suggesting that their accent is not strongly tied to their identity, and that they do not want to speak English with a Spanish accent. With regard to musical aptitude, our results did not show a clear link between self-professed musical aptitude and better results in learning pronunciation, although there was some consistency between musical aptitude and listening performance.

Chapter Five

Summary, Discussion, and Implications

Introduction

Our study set out to explore L2 pronunciation. We saw that pronunciation with regard to research and teaching had been neglected for a number of years and we looked at reasons for this neglect. The challenge for L2 learners to find help for their pronunciation shortcomings along with consequences that they may face as a result of their pronunciation problems were also discussed. By recognising the importance of pronunciation along with a general lack of knowledge regarding the teaching of it, we introduced the idea of using songs as a material to be used in the L2 classroom. Given the fact that songs are only one of many materials that may be utilized in pronunciation teaching, we wanted to know if in general a group of Spanish speakers could learn suprasegmental aspects of English pronunciation through a nine-hour two-week course and if so, whether songs could help with this. In the discussion that follows, the answers to the research questions will be considered in light of both the quantitative and qualitative results. That is, we will consider the participants' performance on the tests, their answers and comments in the questionnaires, comments they made in the second post-course speaking test, as well as observations that I made in class. After that, we consider how our study contributes to the existing knowledge in the field, and make a number of suggestions for future study. We end with a brief summary and a reminder to teachers and researchers about the importance of continuing to learn so that we might best help our students. First, though, we must briefly consider the challenges we faced in our empirical research, and the limitations that these might imply.

Limitations of the Study

In this study, we wanted to find out whether teaching suprasegmental phenomena over a two week period could help L2 learners to better perceive and produce the pronunciation areas taught as well as whether the use of songs as a material would be beneficial for this. In order to achieve these objectives, we conducted a short-term case study of a group of 23 Chilean university students studying English Language and Literature at the Pontificia Universidad Católica de Chile in Santiago, Chile. As indicated in the Introduction chapter of this thesis, a nine-hour course over a period of two weeks is a very short period of time. However, because it was important to have a homogeneous sample of Spanish speakers, I was only able to have the participants over a two-week period because that was the maximum amount of time that the receiving institution could allow the students to not attend their regular classes and instead be available for mine.

Another limitation, which was a direct result of the methodological choice to conduct research in the classroom, was the uneven number of participants. Out of an initial pool of 27 interested students, I created three even-numbered groups. However, some participants withdrew from the study and I was forced to do some shuffling between groups. Even so, I was still unable to even out the groups due to scheduling conflicts. Thus, the no-songs group was left with only five members. This is, however, still a sufficient number according to Duff (2008) for a mixed methods case study like ours.

Furthermore, there was a terrible flu going around which resulted in a number of class absences; only three participants attended all classes. The rest, on average, missed 1.8 classes,

that is, between one and three. In order to mitigate the effect of class absences, I offered to give the students make-up classes; however, only one participant⁹² took advantage of this option.

Methodologically, it would have been beneficial to have had an assistant before and during the intervention period. Because I was the researcher and the teacher, it was not possible to take detailed notes during the pronunciation classes. Instead, I had to rely on what I had observed, which was understandably less detailed than would have been the observations of an assistant free to focus on making detailed notes. Having such an assistant in the class taking observation notes would have provided a more robust body of qualitative data from another perspective. On the other hand, though, having another person in the classroom (or even a camera) to observe the students might have made the participants feel uncomfortable which in turn could have had a negative effect on their learning.

While time constraints, uneven numbers of participants, illnesses and absences and a teacher working alone are less-than-desirable circumstances, they are, according to Rossiter (2001) a reality of classroom research:

Faced with developing constraints related to non-equivalent groups, student and teacher participants, data collection, data analysis, task differences, and ethical considerations, the temptation for many classroom-oriented researchers in my position might be to curtail or even abandon their study. I maintain, however, that what are often perceived as problems by researchers are in fact the daily realities of the contexts in which most teachers practise. The limitations in these research settings may frustrate investigators and pose possible threats to the reliability and validity of quasi-experimental findings; they are, however, part and parcel of the classroom context. (p. 36)

⁹² NS3 missed three classes and took two make-up classes. Whether or not this had a bearing on her performance, we do not know.

With regard to the tests, Listening Test # 1 (see Appendix D) proved to be too easy for the participants because the speaking rate was generally unnaturally slow. It was impossible, however, to have known ahead of time that the test was too easy, because it was designed for intermediate and high intermediate students (Gilbert, 2005b), and the participants in this study were identified as low to mid-intermediate by the institution they were attending. In order to have a clearer idea of how participants' listening skills progress from pre- to post-course, a longer and more challenging listening test would have to be developed, i.e. one with a potential for a greater number of errors to be made and one in which the speech is spoken at a natural pace throughout. Listening Test # 2 was more of a true test of listening because it used natural speech spoken at a normal pace. The problem, however, was that it was a dialogue spoken by only one person. Because of the unavailability of other native speakers for recording purposes, I had no choice but to be that speaker. Although the speech that the participants listened to was authentic, it would have been preferable to have two native speaker voices in the recording.

In order to eliminate the possibility of test familiarity as a factor in elevated post-course scores, which may have occurred in this study, the listening tests could have been further improved. Developing completely different pre- and post-course tests while preserving the type and point value of the pronunciation areas being tested, as well as the level of difficulty, would be good for short-term studies. Nevertheless, an alternative could be to use similar test items, but alter them. Muller Levis and Levis (2012) provide an example in their study of focus words. In their pre- and post-course listening tests, they used the same sentences, but they changed the focus words in the post-test so that the word that was stressed in the post-test was not the same one as in the pre-test.

With regard to speaking tests, the design of Speaking Test # 1 should ideally be rethought. This test (see Appendix D) required the participants to read *The Rainbow Passage*, a standard reading text used to evaluate the production of connected speech. Unfortunately, however, a few of the (very common) connected speech features that were covered in the pronunciation classes did not surface in the reading. Careful selection or design of a new reading text would have to ensure that all areas taught in the course would better occur in the reading and with a higher degree of frequency. The design of such a test, while difficult, would not be impossible. Again, though, ideally different pre- and post-tests should be used.

Our Speaking Test # 2 required the students to use free speech rather than the controlled style typical of reading, for example. However, because it was a free speaking test, it was difficult to compare the results obtained before and after the pronunciation course because the content was naturally not identical. However, in any tests of true spontaneous speech, the content will always be different. Choosing a picture story (a sequenced set of drawings depicting a story) might have helped make the pre and post course results easier to compare because it might have led the participants to use, to a certain extent, the same lexical and grammatical structures. However, using such a speaking prompt would provide the participants with fixed content and thus reduce the cognitive load (Kormos & Dénes, 2004). As such, it could not be said to be a true test of free or spontaneous speech. Therefore, since the ideal test of pronunciation is through free speech, the application of effective pre- and post-course speaking tests is not easily accomplished. In studies with small numbers of participants, it could be done with the help of experienced spoken language proficiency interviewers⁹³ who were adept at eliciting the required pronunciation phenomena. A problem with this, and one that is difficult to

⁹³ Interviewers who have experience carrying out the oral portions of the CanTEST, IELTS, or Canada Language Proficiency Test, for example, would be able to conduct effective interviews.

escape, is Labov's *Observer Paradox*, which states, "The aim of linguistic research in the community must be to find out how people talk when they are not being systematically observed; yet we can only obtain these data by systematic observation" (1972, p. 209). Certainly when people know that their speech is being recorded or evaluated, they are more likely to produce it more carefully.

Assessment of Research Questions

Research Question # 1a

Our first research question was whether teaching suprasegmental phenomena in a two-week pronunciation course can enable L2 learners to better perceive suprasegmental phenomena of the target language. When we compared the mean percentage difference from pre- to post-course for the control and treatment groups, although the treatment group had a higher mean percentage improvement than the control group, the increase was not statistically significant. Nevertheless, when each pronunciation area was analysed separately, the treatment group showed significant improvement over the control group in the area of word stress and a result extremely close to being significant (i.e. $p = .051$ as opposed to $p = .05$) in connected speech. In the other pronunciation areas, namely, thought groups, focus words, contractions and reductions, and linking, quantitatively, the results indicate that the course had no clear effect on the participants' perception of these areas. However, this could be due to the test not being difficult enough, or it could simply mean that the participants were already competent in their ability to perceive those pronunciation phenomena. Of the few studies that have evaluated pronunciation perception (Abe, 2009, 2011; Couper, 2006; Miller, 2012; Muller Levis and Levis, 2012;

Pennington and Ellis, 2000; and Wang, Spence, Jongman, and Sereno, 1999), in the short-term ones such as ours, we can find some similarities in the results. Couper (2006), with regard to perception, was also unable to draw any conclusions from his research which was on (vowel) epenthesis and (consonant) absence and wondered whether perception was more difficult to change than production. Wang et al. (1999) noted a significant improvement in the perception of Mandarin tones after a two-week course in which the participants had received five hours of training. Since (word) stress involves pitch changes, as do tones, it is important to recognise that both studies showed that the participant's ability to perceive pitch changes improved. Similarly Abe (2009), in a study involving connected speech phenomena, namely, rhythm, linking, assimilation, and elision, found that after two hours of instruction over a period of five weeks, significant improvement had been made. In addition, Abe (2011) noticed significant improvement in the perception of weak forms after instruction for a month. With this in mind, perhaps the ability to perceive connected speech and contractions, reductions, and the schwa require more attention and time for assimilation than our course was able to provide. A study that showed different results to ours was that of Pennington and Ellis (2000) in which there was a significant improvement in the perception of focus words; in our study, this was not the case.

The following results for this research question, which pertain to only those of the no-songs group⁹⁴, showed that the course was beneficial. All of the participants strongly agreed to the questionnaire statement, "Taking the pronunciation course helped improve my listening comprehension". In the comments section of the questionnaire, with regard to listening specifically, the participants did not comment. Instead, they referred to the course in general and said that it was enjoyable and helpful. From my perspective as the teacher, I noticed that the

⁹⁴ Because the results of the songs group are discussed in questions 1b and 2b, for purposes of clarity, only the results of the no-songs group are discussed under questions 1a and 2a, with the obvious exception of the listening and speaking tests in which the two groups are combined.

students were engaged in the activities and paid attention during the listening activities. I think because the tasks started out easy and progressed into more difficult ones, the participants were able to approach the activities with confidence. The only exception to this was the first time they attempted the repetition exercises with the movie clips. These exercises required the students to repeat lines of an authentic spoken dialogue without time to reflect on the meaning of what was said. For this exercise, they simply needed reassurance and positive feedback to allow them to go along with the activity.

In summary, with regard to whether the course helped the students' pronunciation listening skills, the test results show that it did, especially in the areas of word stress and connected speech. Furthermore, the participants themselves felt that their listening skills had improved, while I as the teacher noted their ability to handle progressively more difficult listening tasks.

Research Question # 1b

As an extension of our first research question, we wanted to know whether songs as a material could be beneficial for the process of learning to perceive suprasegmental phenomena in the target language. A comparison of the mean percentage improvement for the control, songs, and no-songs groups overall as well as in the individual pronunciation areas, did not indicate any significant improvement of the songs group compared to the other two groups. However, we did notice that in the area of word stress and focus words, three and four of the songs participants, respectively, had a percentage increase above that of the other groups. Similarly, in contractions and reductions, linking, and connected speech, four songs participants performed above average for their group, and in some cases, above average of the other groups. What this means is that

with the exception of thought groups (the pronunciation area which the participants were already adept at), most of the songs group participants (seven out of ten) in one or more pronunciation areas, improved more than the average of any of the groups. Therefore, although the statistics tests did not indicate significant results for teaching perception of pronunciation with songs, clearly the course helped most of the songs group participants in at least one pronunciation area. As mentioned in Chapter Two: A Review of the Literature, there have been very few studies that have used songs for teaching perception of pronunciation phenomena (Karimer, 1984 and Schön et al., 2008) and only Schön et al. has studied whether songs can help with the perception of suprasegmental pronunciation phenomena, which in their study was the perception of word boundaries. As a result we cannot make any quantitative comparisons between our results and those of the other studies. Qualitatively, however, it is possible. The songs participants, as a group, felt strongly that the pronunciation course had improved their listening comprehension and that listening to English music could help them to learn English. Specific comments made by the participants echo what is said in the literature. For example, two of them mentioned learning to notice how native speakers talk (compare to Richards, 1993; Schmidt, 2010; Van den Berg, 2011) while another specifically said that listening to music can help improve one's listening skills (compare to Avery and Ehrlich, 1992; Borland, 2012; Grant, 2007; Van den Berg, 2011). Another participant said that listening to songs results made her to imitate the sounds (compare to Celce-Murcia et al., 2010; Lake, 2011). Finally, one of the participants alluded to how songs could help with language awareness (compare to Bolitho et al., 2003). The student referred to the process involved and said that by training the ear and practicing the sounds, one not only gets more used to saying things differently, but also begins to understand how the suprasegmental areas work together which helps one to understand better what is being said,

which is what other authors have also indicated (compare to Celce-Murcia et al., 2010; Chunxuan, 2009).

In summary, with regard to whether teaching with songs can help learners to better perceive suprasegmental pronunciation phenomena, there is evidence that they do, despite the lack of statistically significant results. Individually, all of the songs group participants indicated that the pronunciation course had helped them improve their listening skills with some of them commenting on exactly how they felt the course had helped them learn. That the course did benefit them is supported by individual post-course test results that show that most of the participants in the songs group made gains that were above average of that of the other groups in one or more pronunciation areas.

Research Question # 2a

Our second research question was whether teaching suprasegmental phenomena in a two-week pronunciation course can enable L2 learners to better *produce* suprasegmental phenomena of the target language. As we saw, in controlled speaking, that is in the reading of a text, the treatment group showed a significant improvement overall as well as in the schwa. In spontaneous speech, there was significant improvement in the production of thought groups. In the literature on pronunciation teaching, of the studies that measured pronunciation production (Abe, 2009, 2011; Couper, 2003, 2006; Derwing, Munro, and Wiebe, 1997, 1998; Derwing & Rossiter, 2003; Lord, 2005, 2008; Macdonald, Yule, and Powers, 1994; Muller Levis and Levis, 2012; Saito, 2011; Saito & Lyster, 2012; Sardegna, 2011; Sardegna & McGregor, 2013; Sturm, 2013) only a few (Abe, 2009, 2011; Couper, 2003; Muller Levis and Levis, 2012; Sturm, 2013; Sardegna, 2011; Sardegna & McGregor, 2013) measured some of the same pronunciation areas

as our study. As in these studies, our treatment group also improved significantly overall when reading even though in general our teaching course was shorter, though Muller Levis and Levis (2012) and Sturm (2013) also gave short courses. None of the above studies targeted specifically the use of the schwa; however, since its use corresponds to the production of weak forms, reduction, and elision, we can say that the results of our study (i.e. a significant improvement in the use of schwa in a reading test) coincide with the results of Couper (2003), Sardegna and McGregor (2013), and Abe (2009, 2011). As mentioned above, in spontaneous speech, our study found that the treatment group improved in the area of thought groups. With regard to other studies, the only one to measure the production of specific suprasegmental areas in spontaneous speech was Couper (2003); however, the suprasegmental areas that he measured were weak forms, linking, word stress, and focus words.

The rest of the results for this research question also show that the course was beneficial. As a group, the no-songs participants not only felt strongly that the course had helped improve their pronunciation, but also that their pronunciation would continue to improve in the future. There was no loss of motivation for them from pre- to post-course and, as we saw earlier, they all indicated that they enjoyed the course and found it helpful. One of the members talked about the value of the course in that by helping them improve their pronunciation, it would mean that native speakers would be able to understand them better. Another participant indicated being able to identify her mistakes, while someone else mentioned a heightened sense of awareness of pronunciation and the intention to put to practice what was learned in class. While two members expressed regret for not being a part of the songs group, they both recognised the value of the course even though no songs were involved. From my perspective as the teacher, I noticed that the students readily took to the speaking tasks and paid attention to the areas that were covered in

the classes. They appeared to all be serious students and were fully engaged in the tasks and activities.

To sum up, with regard to whether the course helped the students' improve their English pronunciation, the test results show that it did, especially in controlled speech. When reading there was significant improvement in their use of the schwa; in spontaneous speech, significant improvement was noted in thought groups. The students themselves felt strongly not only that their pronunciation had improved, but also that it would continue to improve in the future. From my perspective as the teacher, the students applied themselves and showed a genuine desire to learn.

Research Question # 2b

Finally, the second part of research question number two asked whether songs as a material could be beneficial for the process of learning to produce the suprasegmental phenomena. A comparison of the pre- to post-course performance for the control, songs, and no-songs groups showed that the songs group had the largest mean percentage decrease in errors of the three groups in both speaking tests and this decrease was significant in controlled speech. The individual pronunciation areas in which the songs group did significantly better than the other groups was in the schwa when reading, and in thought groups when speaking freely. With regard to the studies that have used songs for teaching pronunciation production (speaking), namely, Chunxuan (2009), Fischler (2009), and Rengifo (2009), while all of them showed that improvement in pronunciation had been made, we cannot further compare their results with ours. This is because it is either not clear which suprasegmental pronunciation areas were included or, in the case of Fischler (2009), only intelligibility ratings were measured.

When we looked at how individual songs group members performed compared to group averages, five of them had an above average decrease in errors across the groups when reading and four when speaking freely. If we look at the individual pronunciation areas, we will see that a number of the members outperformed the other groups in terms of their percentage decrease in errors. When reading a text, this was the case for nine members with the schwa, for five with focus words, four with linking, and three with both thought groups and word stress. What this tells us is that learning pronunciation with songs helped all of the songs group participants in one or more areas when reading. When speaking freely, all ten songs group participants had a greater percentage decrease in errors than the mean averages of the other groups in thought groups. In the schwa and linking, this was the case for four participants. For focus words, connected speech and word stress, three, two, and one, respectively, outperformed the other groups. With regard to focus words and word stress, the other participants did not show improvement because they were already skilled in those pronunciation areas, with most of them not making any errors. Therefore, if we consider the statistical results as well as the individual results for the songs group, it is clear that teaching pronunciation with songs helped improve the pronunciation of the participants.

The other data also supports this. In the questionnaires, all of the songs group participants indicated that the course had helped them improve their English pronunciation and all of them felt that it would continue to improve because they took the course. Furthermore, although before the course all of the participants felt that listening to English music can help with learning pronunciation, after the course, three of the participants felt more strongly that it can help. When we look at the comments that the participants made about the course, we can see that there are common themes that relate to the classroom dimensions.

With regard to the linguistic aspect, almost all of the participants commented in one way or another that the course was helpful or useful for learning pronunciation, which is what Jolly (1975) also found. One of the participants who provided more detail said that she felt her English would be “more fluent and clear” and that she would express herself better, essentially echoing Chunxuan (2009) when he says that songs help learners speak more fluently and correctly. As well, some of them mentioned gaining an awareness of the pronunciation of English, which is what Bolitho et al. (2003) indicate about using songs in language teaching. The comments that relate to the affective aspect revealed that the students really enjoyed the course, found it to be a lot of fun, and were grateful for having been a part of it. These comments are similar to those that Adkins (1997) received from her students after teaching them with songs. In addition, the fact that the songs group participants maintained a high level of motivation, with some of them mentioning the intention to keep working on their pronunciation, lends credence to the motivating influence of songs, as indicated by a number of authors (Adkins, 1997; Baoan, 2008; Borland, 2012; Chunxuan, 2009; Ebong & Sabbadini, 2006; Lê, 1999; Murphey, 1995; Stansell, 2005; Van den Berg, 2011). Evidence of the songs having a bearing on the social / cultural aspect could be seen in comments that referred to me as the teacher. There were students that wished me well and said I was very “nice”. Personally, I was surprised that I was able to gain such a strong rapport with the students after only two weeks of classes, even though, as we have seen, Lê (1999) mentions this as a benefit of using songs in the classroom. Regarding the cognitive aspect, a student mentioned that she had experienced the Song-Stuck-in-my-Head (SSIMH) phenomenon (Murphey, 1990) when she says, “I really liked it and the music was really good, although I was like singing after it a while”. This brief summary of the comments made by the songs participants not only illustrates some of the different ways that learning with

songs benefitted them, but it also lends more credibility to what language teachers who work with songs have been saying all along, even though they might not have carried out quantitative research on the use of songs. Finally, from my perspective as the teacher, I noticed that the students were interested, engaged, and enjoying themselves during the activities. Therefore, these observations along with the results of the speaking tests and what the students themselves have indicated, definitely show that songs were beneficial for learning pronunciation.

General Discussion of the Study

After looking at the answers to our research questions, we now turn to a more general discussion of our study. Using songs within a communicative approach, first of all, constitutes a novel way to teach pronunciation. We have shown, in our lesson plans, how songs can be used to teach learners how to perceive and produce suprasegmental phenomena. Careful sequencing of activities coupled with progressively more challenging tasks within an overall theme, show how songs and pronunciation teaching can be incorporated into a communicative language class.

Secondly, by choosing to have L1 Spanish speakers as our participants, we have gained a better understanding of which pronunciation areas of theirs might be better helped with songs as well as other materials. That is, we discovered that a short pronunciation course improved our learners' ability to perceive and produce suprasegmental phenomena. In perception, our participants made significant progress with word stress. Moreover, their perception of connected speech, focus words, contractions and reductions, and linking was also positively affected. In production, our participants made significant progress in their use of the schwa and thought groups. As well, progress in focus words and linking was apparent. With this knowledge, EFL

teachers and researchers in Spanish-speaking countries (especially) will be able to focus more meaningfully on precise areas of pronunciation in their teaching and research.

Furthermore, the lesson plans, including the handouts that I developed, are a concrete model for teachers wishing to incorporate more pronunciation teaching into their classrooms. Those that already use The Prosody Pyramid can easily pick and choose the songs and activities they wish to provide for extra practice. Whether or not teachers follow a textbook, the materials developed for this study will provide them with ideas for creating similar materials of their own.

Finally, the in-depth analysis of the individual songs participants provided more information than statistical tests alone could do. We saw that almost all of the songs group participants achieved an above average improvement in at least one suprasegmental area, with many of them having such improvement in various areas. Furthermore, the participants only had positive things to say about the course. Understanding how songs affected the students' performance and learning experience helped me to appreciate even more the value of songs. Just as learning outcomes are important, so are learning experiences. Songs really can enhance the classroom experience in many ways. The dimensions of the classroom are real and ways of teaching that keep these dimensions in mind can benefit both students and teachers. If our students are happy, then that makes our jobs as teachers much more pleasant and gratifying. Therefore, I hope that our study in this respect has lent some legitimacy to what has been said about songs – but rarely studied. I also hope that more teachers and researchers will do more studies involving songs. It is possible that once one has learned how to listen for pronunciation features in music then one will always do so. That is, our ear may be forever changed once it has been educated through song.

Implications for Future Study

Even though more studies are beginning to be done on the effects of teaching on L2 pronunciation, even more research is needed in our quest to integrate pronunciation teaching appropriately into the L2 curriculum. We know that teaching can have a positive effect on a learner's pronunciation, but it is still not clear to what extent. We are still not certain as to which pronunciation areas may be more or less affected depending on the teaching methodology and materials used. Clearly studies that have been done in pronunciation teaching can guide us in new studies and ours is no exception to that. As we saw in Chapter Four: Research Findings, there were some pronunciation areas that our participants were already strong in and others in which they were weaker, and these areas were different depending on whether pronunciation perception or production was being measured. Furthermore, regardless of how skilled the participants were in the pronunciation areas, some areas appeared to be more or less resistant to change. These observations contribute to the following reasons for conducting future studies in pronunciation teaching.

In listening, the participants appeared to be already proficient in the ability to identify thought groups because they scored very high in the pre-course test. It is important to remember, though, that it is possible that the test was a little too easy for the students, even though in theory it was designed for their proficiency level. Nevertheless, although there was a little room for improvement, almost no improvement was made post-course. With this in mind it would be advantageous to further test Spanish speakers using both this listening test as well as a more challenging one with, as mentioned above, slightly altered pre- and post-course tests. By doing this, we would know with more certainty whether Spanish speakers require pronunciation

training in the perception of thought groups and whether they are resistant to improved perception. With regard to our participants' production of thought groups, in both controlled and spontaneous speech, they were less proficient. Although we noticed gains in reading irrespective of the materials used, the fact that the songs group had a significant improvement in the use of thought groups in spontaneous speech is a good reason to see whether these results can be duplicated in other studies. Furthermore, if the studies also target thought group perception, then we will have a better understanding not only of what kinds of materials might be most useful, but also more about the relationship between the perception and production of thought groups.

In the perception of word stress, our participants were strong but slightly less proficient than in thought groups, and both the songs and no-songs groups showed improvement post-course whereas the control group did not. This suggests that improved perception of word stress is possible through teaching. Nevertheless, in production, it was the songs group who made the most gains in both controlled and spontaneous speech. Because there have been so few studies on teaching word stress (Couper, 2003; Sardegna and McGregor, 2013) and only one involving music (Fischler, 2009), it would be worth conducting more in this area to see what combinations of teaching methods and materials might be most effective.

Our participants, as Spanish speakers, naturally, had a few difficulties in focus words. All groups showed improvement in the perception of focus words post-course, *especially* the control group, which lead us to wonder whether test familiarity plays a role. As mentioned above, Muller Levis and Levis (2012) and Pennington and Ellis (2000) both conducted studies on focus words, and while Pennington and Ellis found significant results in the perception of focus words, Muller Levis and Levis did not. With regard to production, our participants made more errors when reading and fewer when speaking freely, and although in both cases fewer

errors were made by the songs group post-course, the results were not significant, which is in contrast to Muller Levis and Levis who found significant results. Since we have three studies which produced different results in the perception and production of focus words, clearly more research is needed.

In the area of contractions and reductions, while quite proficient in perception, our participants proved to be less so in production. In listening we saw that only a little improvement was made (although, as mentioned, this could be due to the test being too easy); in speaking, significant progress in the production of the schwa was made when reading, but this did not carry over to free speech. While we know that controlled and spontaneous speech will not necessarily show the same kinds of pronunciation errors (Levis and Barriuso, 2012), further research on possible teaching methods and materials for bridging this gap should be done. Because the schwa is the most common sound in English and the one that results in syllables being reduced, it is a fundamental pronunciation area not just for Spanish-speaking learners, but for other L2 learners of English with different L1s..

In addition, more investigations into teaching methods that target linking and connected speech are needed. In linking, our participants were better at producing it than perceiving it, and their linking was better when speaking freely than when reading. Abe (2009) noted significant improvement in the perception of linking and since in our study we did not, it would be worth trying to replicate his study with Spanish speakers. In connected speech, Abe (2009) noted significant improvement in both perception and production, whereas in our study, there was clear improvement in perception, but not production, because our participants were already proficient in producing connected speech phenomena.

Another idea for future research has to do with the supposed connection between language learning and musical aptitude. While we did not specifically set out to determine whether a connection exists, as we saw when we compared self-professed musical aptitude and success in learning receptive and productive pronunciation skills, there was some indication of a possible link between musical aptitude and listening, but not with speaking. This differs somewhat from Stansell (2005), who says that musical people are better at both perceiving and producing pronunciation phenomena.

Finally, research into how learners feel about their L2 accent in relation to their identity is definitely needed. Our participants wanted to sound like a native English speaker, did not want to keep their Spanish accent, and felt that they would feel comfortable doing so. They disagreed that they would feel like another person if they spoke like a native English speaker and tended to disagree that their Spanish accent is part of who they are. These opinions seem reasonable and are mutually compatible. However, some have expressed the idea that the loss of an L1 accent results in the loss of the L1 identity (see Setter & Jenkins, 2005). Therefore, given the complex nature of, and relationship between, identity and accent, more studies that consider the opinions of learners of varying proficiencies are necessary.

With regards to the use of songs as a material, the field is especially wide open in all pronunciation areas. As we have seen, our study is one of the very few that have used songs to teach pronunciation. We targeted a number of suprasegmental areas in a short period of time in order to find out if songs could have an effect on the pronunciation skills of a group of Spanish speakers. Since we found significant results in the schwa and thought groups when speaking, a good place to start with further classroom studies could be these areas with L1 Spanish speakers as well as speakers of other L1s. As we saw, a deeper analysis of the performance and opinions

of the songs group participants showed that songs can be helpful, not only in different pronunciation areas, but also in the overall learning experience of the students. With this in mind, more mixed-method studies need to be done with more participants and for a longer intervention period in order to see if the results can achieve significance. Studying a combination of pronunciation areas, as we did, is important because each one is an integral component of a larger system which contributes to effective communication. While it is important to gain a better understanding of how songs might be especially useful for improving certain pronunciation areas, we also need to know if they *fail* to help other areas, so that alternative materials can be explored.

We had a small relatively homogeneous sample of highly motivated participants who received nine hours of classes. Studies with a longer intervention period and more participants with varied levels of motivation will shed further light on the power of songs for teaching pronunciation. It would also be advantageous to conduct studies with groups of participants who are beginners. The extent to which pronunciation training influences both perception and production might benefit L2 beginners in the long-term would contribute to L2 language acquisition studies in general. In addition, pronunciation studies with songs that target segmentals would provide further insight on how songs might help improve L2 pronunciation.

Finally, learning more about how spontaneous speech can be affected through pronunciation training is necessary. As we have seen, many pronunciation studies have measured controlled speech (Abe, 2009, 2011; Couper, 2003, 2006; Derwing, Munro, and Wiebe, 1997, 1998; Fischler, 2009; Lord, 2005, 2008; Muller Levis and Levis, 2012; Saito, 2011; Saito and Lyster, 2012; Sardegna, 2011; Sardegna and McGregor, 2013; Sturm, 2013) with fewer measuring spontaneous speech (Couper, 2003; Derwing, Munro, and Wiebe, 1998; Derwing and

Rossiter, 2003; Fischler, 2009; Lord, 2008; Saito, 2011; Saito and Lyster, 2012). However, if the ultimate goal of most language learners is to be able to *speak* effectively (as opposed to *reading* effectively), then researchers should focus more on the effects of pronunciation training on spontaneous speech.

On the basis of the results of our study and how they compare to those of others, there is much that needs to be done before we can have a clearer understanding of how L2 pronunciation can improve through teaching.

Summary

Studies, including this one, have shown that significant progress in L2 pronunciation can be made. We have several guiding principles and now, an under-utilized and under-explored material – songs – for applying those principles. That is not to say, however, that songs are the *only* material that can help learners. In our study, we have seen that a combination of other materials can also be effective. Songs, though, uniquely permeate the different dimensions of the classroom, and thus contribute to a positive learning – and teaching – environment.

With more studies, we will get closer to discovering how a person's L2 pronunciation perception and production might improve. Although we know that improvement can be made, we still do not know to what extent. With this in mind, although it is important for L2 speakers to be intelligible, we must be careful not to limit them to a goal of mere intelligibility if they consider that that is not enough for them. Learners have different aptitudes, opportunities, and goals. As teachers, both the needs and desires of our students should always be kept in mind.

While many learners might only need to be intelligible, there are others who want to achieve more. They may have professional aspirations that require an excellent command of the language or they may simply have a strong personal desire to sound like a native speaker. Since we are still discovering new and better ways to teach pronunciation and have yet to discover the extent to which effective pronunciation teaching can help, we should not rule out the possibility of learners achieving proficient or even native-like speech if they have ample desire, goals, aptitude, and opportunities. We should keep in mind the words of Marinova-Todd et al. (2012) who say in reference to teachers: “they can do much to influence a student's learning strategies, motivation, and learning environment. Thus, such teachers are justified in holding high expectations for their students and can give their motivated students research-based information about how to improve their own chances for learning to a high level” (p. 30). This is not to say that we should be imposing impossibly higher standards or giving our students false hope of achieving native-like pronunciation. But nor should we impose lower standards or limit them by suggesting that mere intelligibility is a good enough goal when it might not be. We need to be careful when telling learners what their needs should be. Not only is it unfair to apply a blanket approach to students’ needs, but we should also question whether or not we have the right to impose upon them standards of speaking which they may see as insufficient. Imposing a lower standard of pronunciation on L2 learners might also have racist overtones in a multiethnic context, blocking immigrant learners of different races, for instance, from achieving future social and professional success. We have seen that racism is alive and well, especially with regard to speaking, and advising language learners to only reach for intelligibility may not be enough. After all, given that we are still learning about pronunciation, we need to be careful not to limit our learners by imposing our (partial) knowledge upon them. What is best for our students

should not be a matter of choosing between intelligibility or nativelikeness. In this complex matter, there must be a middle ground.

Research is a quest for knowledge and as such researchers should be careful not to be such strong adherents to their beliefs that they ignore or try to discredit the research of adherents to other beliefs. (see Marinova-Todd et al., 2012). Considering that pronunciation teaching and research has only in the last decade or so begun to be studied more systematically, we should not jump to premature conclusions. We still do not know if some areas of pronunciation might be more resistant to perception than production (Couper, 2006) even though there is the notion that pronunciation perception must precede production (Cheng, 1998). We are only beginning to learn that fossilisation, for example, may be overcome (Couper, 2006; Derwing, Munro, & Wiebe, 1997), but we are a long way from being able to say that fossilisation is not a permanent affliction for some learners. We still do not know the potential of effective teaching of pronunciation for L2 learners. We must keep an open mind. It is only in this way that we will keep learning.

REFERENCES

- Abe, H. (2009). The effect of interactive input enhancement on the acquisition of the English connected speech by Japanese college students. *Phonetics Teaching & Learning Proceedings, University College London*. Retrieved from http://www.phon.ucl.ac.uk/ptlc/proceedings/ptlcpaper_29e.pdf.
- Abe, A. (2011). Effects of form-focused instruction on the acquisition of English weak forms by Japanese EFL learners. In W.-S. Lee & E. Zee (Eds.), *Proceedings of the 17th International Congress of Phonetic Sciences, Hong Kong* (pp. 184–187). Retrieved from <http://www.docstoc.com/docs/171830653/effects-of-form-focused-instruction-on-the-acquisition-of-weak-forms#>
- Abrahamsson, N., & Hylenstam, K. (2009). Age of onset and nativelikeness in a second language: Listener perception versus linguistic scrutiny. *Language Learning*, 59(2), 249-306. doi: 10.1111/j.1467-9922.2009.00507.x
- Adamowski, E. (1997). *The ESL songbook*. Toronto: Oxford University Press.
- Ajioka, M. (2010). Grammar, pronunciation, or something else? Native Japanese speakers' judgments of "native-like" speech. *Issues in Applied Linguistics*, 12(8), 233-250. Retrieved from <http://escholarship.org/uc/item/50w3991n>
- Anton, R. (1990). Combining singing and psychology. *Hispania*, 73(4), 1166-1170. Retrieved from http://www.cervantesvirtual.com/s3/BVMC_OBRAS/027/564/ea8/2b2/11d/fac/c70/021/85c/e60/64/mimes/p0000029.htm

- Avery, P., & Ehrlich, S. (1992). *Teaching American English pronunciation*. New York: Oxford University Press.
- Baluja, T. (2011). The pressure to smooth an accent in the workplace. *The Globe and Mail*. Retrieved from <http://www.theglobeandmail.com/report-on-business/careers/career-advice/the-pressure-to-smooth-an-accent-in-the-workplace/article587505/>
- Bambury, B. (2011). How music therapy soothed the bullet-damaged brain – Canada - CBC News. *CBC.ca - Canadian News Sports Entertainment Kids Docs Radio TV*. Retrieved from <http://www.cbc.ca/news/canada/story/2011/11/18/f-vp-bambury.html>
- Baoan, W. (2008). Application of popular English songs in EFL classroom teaching. Humanising Language Teaching (HLT) Free Online EFL Magazine for Teachers of English. Retrieved from <http://www.hltmag.co.uk/jun08/less03.htm>
- Bareither, T. (2007). *Tongue twisters, rhymes, & songs to improve your English pronunciation*. Bloomington: Authorhouse.
- Bhatt, M. (2013). Losing an accent could be a good thing for newcomers. *Canadian Immigrant*. Retrieved from <http://canadianimmigrant.ca/settling-in-canada/losing-an-accent-could-be-a-good-thing-for-newcomers-bhatt>
- Birdsong, D. (2005). Nativelikeness and non-nativelikeness in L2A research. *IRAL - International Review of Applied Linguistics in Language Teaching*, 43(4), 319-328. doi:10.1515/iral.2005.43.4.319

- Bitchener, J. (2010). *Writing an applied linguistics thesis or dissertation: A guide to presenting empirical research*. New York: Palgrave Macmillan.
- Bolitho, R., Carter, R., Hughes, R., Ivanic, R., Masuhara, H., & Tomlinson, B. (2003). Ten questions about language awareness. *ELT Journal*, 57(3), 251-259. doi:10.1093/elt/57.3.251
- Bongaerts, T., Summeren, C., Planken, B., & Schils, E. (1997). Age and ultimate attainment in the pronunciation of a foreign language. *Studies in Second Language Aquisition*, 19, 447-465. Retrieved from http://resolver.scholarsportal.info/resolve/02722631/v19i0004/447_aauaitpoaf
- Bongaerts, T., Susan M., & Slik, F. (2000). Authenticity of pronunciation in naturalistic second language acquisition: The case of very advanced late learners of Dutch as a second language. *Studia Linguistica*, 54(2), 298–308. doi: 10.1111/1467-9582.00069
- Borland, K. (2012). Teaching English and Spanish through songs. *Memoria Conferencia Annual 2011. American Association of Teachers of Spanish and Portuguese Capítulo Ontario (AATSP-ON)* (January 2012). <https://docs.google.com/viewer?a=v&pid=sites&srcid=ZGVmYXVsdGRvbWFpbXNhYXZlcG9ub3R0YXdhY29uZmVyZW5jZXxneDo3MDg5ZDE0NWZhMDk4NWE4>
- Breitkreutz, J., Derwing, T., & Rossiter, M. (2001). Pronunciation teaching practices in Canada. *TESL Canada Journal*, 19, 51–61.
- Brown, H. (2000). *Principles of language learning and teaching* (4th ed.). New York: Longman.

- Carlson, M. (Producer) (2012). 247. *A stolen butterfinger*. Retrieved from <http://www.eslyes.com/nyc/nyc/nycs247.htm>
- Cauldwell, R. (2013). Models, metaphors, and the evidence of spontaneous speech: A new relationship for pronunciation and listening. *Pre Conference Workshop. Pronunciation in the Language Teaching Curriculum Pronunciation in Second Language Learning and Teaching PSLLT 5th Annual Conference*.
- Celce-Murcia, M., Brinton, D., & Goodwin, J. (2010). *Teaching pronunciation: A course book and reference guide* (2nd ed.). New York: Cambridge University Press.
- Chan, T. (1999) Teaching pronunciation using song lyrics. *Temple University, Japan Campus*, 27. Retrieved from http://www.tuj.ac.jp/newsite/main/tesol/publications/studies/vol_27/labrenz.html
- Cheng, F. (1998). The teaching of pronunciation to Chinese students of English. *English Teaching Forum, Jan-Mar, 1998*, 37-39. Retrieved from <http://dosfan.lib.uic.edu/usia/E-USIA/forum/vols/vol36/no1/p37.htm>
- Cho, S., & Reich, G. (2008). New immigrants, new challenges: High school social studies teachers and English language learner instruction. *The Social Studies*, 99(6), 235-242.
- Christiner, M., & Reiterer, S. (2013). Song and speech: Examining the link between singing talent and speech imitation ability. *Frontiers in Psychology*, 4, 45-55.
- Christison, M. (1996). Teaching and learning languages through multiple intelligences. *TESOL Journal*, 6(1), 10-14.

- Chunxuan, S. (2009). Using English songs: An enjoyable and effective approach to ELT. *English Language Teaching*, 2(1), 88-94. Retrieved from <http://www.ccsenet.org/journal/index.php/elt/article/view/341>
- Cliff, J. (Performer). (2011, August 31). I can see clearly now [Web Video]. Retrieved from http://www.youtube.com/watch?v=4P_vX31VA_Y
- Cole, H. (Performer). (2012, April 12). I can see clearly now [Web Video]. Retrieved from http://www.youtube.com/watch?v=-a_Iy3VoWJM
- College Humor. (Performer). (2012, June 06). Some study that I used to know [Web Video]. Retrieved from <http://www.youtube.com/watch?v=VxkHM4DUDKM>
- Compas. (2009). Best practices: Employers and IEPs speak about strategies for integrating internationally educated professionals into the Canadian labour force. Retrieved December 14, 2012, from <http://www.compas.ca/090311-IEP-B.html>.
- Couper, G. (2002). ESOL course evaluation through a survey of post-course experiences. *TESOLANZ Journal*, 10, 36–51.
- Couper, G. (2003). The value of an explicit pronunciation syllabus in ESOL teaching. *Prospect*, 18, 53-70.
- Couper, G. (2006). The short and long-term effects of pronunciation instruction. *Prospect*, 2 (1), 46-66.
- Crane, B. (Performer). (2012, May 21). Whatcha gonna do? [Web Video]. Retrieved from <http://www.youtube.com/watch?v=QW4RUIjOp9k>

- Cummings, B. (Composer). (2011, April 28). Burton Cummings - Break it to them gently [Web Video]. Retrieved from <http://www.youtube.com/watch?v=8AReU2zUTgs>
- Dale, P., & Poms, L. (2005). *English pronunciation made simple*. White Plains, Pearson Education.
- Darcy, I., Ewert, D., & Lidster, R. (2012). Bringing pronunciation instruction back into the classroom: An ESL teachers' pronunciation "toolbox". In J. Levis & K. LeVelle (Eds.). *Proceedings of the 3rd Pronunciation in Second Language Learning and Teaching Conference*, Sept. 2011. (pp. 93-108). Ames, IA: Iowa State University.
- Deng, J., Holtby, A., Howden-Weaver, L., Nessim, L., Nicholas, B., Nickle, K., Pannekoek, C., Stephan, S., & Sun, M. (2009). English pronunciation research: The neglected orphan of second language acquisition studies? (*WP 05-09*). Edmonton, AB: Prairie Metropolis Centre.
- Derwing, T. (2003). What do ESL students say about their accents? *Canadian Modern Language Review*, 59, 547–566.
- Derwing, T. (2010). Utopian goals for pronunciation teaching. In J. Levis & K. LeVelle (Eds.), *Proceedings of the 1st Pronunciation in Second Language Learning and Teaching Conference*, Iowa State University, Sept. 2009. (pp. 24-37), Ames, IA: Iowa State University.
- Derwing, T., Munro, M., & Wiebe, G. (1997). Pronunciation instruction for "fossilized" learners: Can it help? *Applied Language Learning*, 13, 1-17.

- Derwing, T., Munro, M., & Wiebe, G. (1998). Evidence in favour of a broad framework for pronunciation instruction. *Language Learning*, 48, 393-410.
- Derwing, T., & Munro, M. (2005). Second language accent and pronunciation teaching: A research-based approach. *TESOL Quarterly*, 39, 379-397.
- Derwing, T., & Munro, M. (2006). How pronunciation research can inform teaching practices. *Paper presented at TESL Canada Conference, 2006*. Retrieved from www.sfu.ca/~mjmunro/Info%20for%20Teachers_files/TCHdt.pdf
- Derwing, T., & Rossiter, M. (2002). ESL learners' perceptions of their pronunciation needs and strategies. *System*, 30, 155-166.
- Derwing, T., & Rossiter, M. (2003). The effect of pronunciation instruction on the accuracy, fluency and complexity of L2 accented speech. *Applied Language Learning*, 13(1), 1-17.
- Derwing, T., Rossiter, M., & Munro, M. (2002). Teaching native speakers to listen to foreign-accented speech. *Journal of Multilingual and Multicultural Development*, 23(4), 245-259.
- Dörnyei, Z. & Ottó, I. (1998). Motivation in action. *Working Papers in Applied Linguistics*, 4, 43-69.
- Dörnyei, Z. (2007). *Research methods in applied linguistics: Quantitative, qualitative and mixed methodologies*. Oxford: Oxford University Press.
- Dörnyei, Z., & Ushioda, E. (2011). *Teaching and researching motivation* (2nd ed.). Harlow: Longman.
- Duff, P. (2008). *Case study research in applied linguistics*. New York: Routledge.

- Ebong, B., & Sabbadini, M. (2006). Developing pronunciation through songs *British Council BBC*. Retrieved from <http://www.teachingenglish.org.uk/articles/developing-pronunciation-through-songs>
- Ellis, R. (2008). *The Study of Second Language Acquisition, 2nd ed.* Oxford: Oxford University Press.
- Engh, D. (2013a). Why use music in English language learning? A survey of the literature. *English Language Teaching. Canadian Centre of Science and Education*, 6(2).
- Engh, D. (2013b). Effective use of music in language learning: A needs analysis. *Humanizing Language Teaching*, 15(5), Retrieved from <http://www.hltmag.co.uk/oct13/mart03.htm>.
- Erickson, B., & Nosanchuk, T. (1979). *Understanding data: An introduction to exploratory and confirmatory data analysis for students in the social sciences*. Milton Keynes: Open University Press.
- Fairbanks, G. (2009, February 22). *The rainbow passage*. Retrieved from [http://everything2.com/title/The Rainbow Passage](http://everything2.com/title/The+Rainbow+Passage)
- Field, J. (2005). Intelligibility and the listener: The role of lexical stress. *TESOL Quarterly*, 39, 399-423.
- Fischler, Janelle. (2009). The rap on stress: Teaching stress patterns to English language learners through rap music. *Minnesota and Wisconsin Teachers of English to Speakers of Other Languages*. Retrieved from the University of Minnesota Digital Conservancy, <http://purl.umn.edu/109937>.

- Flege, J., Munro, M. & Mackay, I. (1995). Factors affecting strength of perceived foreign accent in a second language. *Journal of the Acoustical Society of America*, 97, 3125-3134.
- Foote, J., Holtby, A., & Derwing, T. (2011). 2010 survey of pronunciation teaching in adult ESL programs in Canada. *TESL Canada Journal*.
- Francis, J. (2010). A few notes on the use of singing and music in the ESL classroom. *Otsuma Journal of Social Information Studies*, 19, 157-164. Retrieved from <http://ci.nii.ac.jp/naid/110008427693/>
- Fraser, H. (2000). Coordinating improvements in pronunciation teaching for adult learners of English as a second language. Canberra: DETYA.
- Gardner, Howard. (1993) *Frames of mind: The theory of multiple intelligences* (2nd ed.). London: Fontana Press.
- Geranpayeh, A. (2003). A quick review of the English Quick Placement Test. *Research Notes*, 12, 8-10. Retrieved from http://www.uniss.it/documenti/lingue/what_is_the_QPT.pdf
- Giffords, G. (2011, November). A message from Rep. Gabrielle Giffords by Rep Giffords on SoundCloud - Create, record and share your sounds for free. *SoundCloud - share your sounds*. Retrieved December 20, 2011, from <http://soundcloud.com/user8083112/a-message-from-rep-gabrielle>
- Gilbert, J. (2005a). *Clear speech: Pronunciation and listening comprehension in North American English: student's book* (3rd ed.) Cambridge: Cambridge University Press.

- Gilbert, J. (2005b). *Clear speech: Pronunciation and listening comprehension in North American English: Teacher's resource book* (3rd ed.) Cambridge: Cambridge University Press.
- Gilbert, J. (2008). *Teaching pronunciation using the prosody pyramid*. Cambridge: Cambridge University Press.
- Gilbert, J. (2010). Pronunciation as orphan: What can be done? *As We Speak, Newsletter of TESOL*. Retrieved from http://www.cambridge.org/other_files/downloads/esl/clearspeech/orphan.pdf.
- Gordon, J. (2012). Extra-linguistic factors in the teaching and learning of pronunciation in an ESL class. In J. Levis & K. LeVelle (Eds.). *Proceedings of the 3rd Pronunciation in Second Language Learning and Teaching Conference*, Sept. 2011. (pp. 134-146). Ames, IA: Iowa State University.
- Graham, C. (1978). *Jazz chants*. New York: Oxford University Press.
- Grant, L. (2007). *Well Said Intro: Pronunciation for clear communication*. Boston: Heinle.
- Grant, L. (2010). *Well said: Pronunciation for clear communication* (3rd ed.). Boston: Heinle.
- Hadfield, J. & Dörnyei, Z. (2013). *Motivating learning*. Harlow: Longman.
- Hahn, L. (2004). Primary stress and intelligibility: Research to motivate the teaching of suprasegmentals. *TESOL Quarterly*, 38, 201-223.
- Hancock, M. (1995). *Pronunciation games*. Cambridge: Cambridge University Press.

- Hatch E. and Farhady H. (1982). *Research design and statistics for applied linguistics*. Rowley: Newbury House Publishers, Inc.
- Henderson, A., Frost, D., Tergujeff, E., Kautzsch, A., Murphy, D., Kirkova-Naskova, A., Waniek-Klimczak, E., Levey, D., Cunningham, U. & Curnick, L. (2012). The English pronunciation teaching in Europe survey: Selected results. *Research in Language*, 10.1, 5–27.
- Herr, K., & Anderson, G. (2005). *The action research dissertation: A guide for students and faculty*. Thousand Oaks, Calif.: SAGE Publications. Retrieved from http://www.ciae.uchile.cl/download.php?file=otros/tallerinvestigacion/anderson_y_kerr_actionresearchdiss.pdf
- Immen, W. (2011, July 15). Maybe your English isn't as good as you think. *The Globe and Mail*. Retrieved August 4, 2011, from <http://www.theglobeandmail.com/report-on-business/maybe-your-english-isnt-as-good-as-you-think/article187885/>
- Ioup, G., Boustagui, E., El Tigi, M., & Moselle, M. (1994). Reexamining the critical period hypothesis: A case study of successful adult SLA in a naturalistic environment. *Studies in Second Language Acquisition*, 16(1), 73-98.
- Isaacs, T. (2009). Integrating form and meaning in L2 pronunciation instruction. *TESL Canada Journal*, 27(1), 1–12.
- Jenkins, J. (2000). *The phonology of English as an international language*. Oxford: Oxford University Press.

- Jenkins, J. (2004). Research in teaching pronunciation and intonation. *Annual review of applied linguistics*, 24, 109-125.
- Jobs, S. (Performer). (2009, September 14). Steve Jobs explains the rules for success [Web Video]. Retrieved from <http://www.youtube.com/watch?v=KuNQgln6TL0>
- Jolly, Y. (1975). The use of songs in teaching foreign language. *The Modern Language Journal*, 59, 11-14.
- Karimer, L. (1984). Can Southeast Asian students learn to discriminate between English phonemes more quickly with the aid of music and rhythm? *Illinois: Annual state Convention of the Illinois Teachers of English to Speakers of other Languages/Bilingual Education*. (ERIC Document Reproduction Service No. ED 263 783).
- Kelly, L. (1969). *25 centuries of language teaching*. Rowley, MA: Newbury House.
- Kennedy, R., & Scott, A. (2005). The effects of music therapy intervention on middle school students' ESL skills. *Journal of Music Therapy*, 42(4), 244-261.
- Kirkova-Naskova, A., Tergujeff, E., Frost, D., Henderson, A., Kautzsch, A., Levey, D., Murphy, D., Waniek- Klimczak, E. (2013). Teachers' views on their professional training and assessment practices: Selected results from the English Pronunciation Teaching in Europe survey. In J. Levis & K. LeVelle (Eds.). *Proceedings of the 4th Pronunciation in Second Language Learning and Teaching Conference*. Aug. 2012. (pp. 29-42). Ames, IA: Iowa State University.

- Kjellin, O. (1999). *Accent addition: Prosody and perception facilitates second language learning*. In O. Fujimura, B. D. Joseph, & B. Palek (Eds.), *Proceedings of LP'98 (Linguistics and Phonetics Conference) at Ohio State University, Columbus, Ohio, September 1998* (Vol. 2, pp. 373-398). Prague: The Karolinum Press. Retrieved from <http://olle-kjellin.com/SpeechDoctor/ProcLP98.html>
- K'Naan. (Performer). (2009, February 06). Wavin' flag [Web Video]. Retrieved from <http://www.youtube.com/watch?v=TxmEd9lcn0k>
- Kormos, J., & Dénes, M. (2004). Exploring measures and perceptions of fluency in the speech of second language learners. *System*, 32(2), 145-164.
- Labov, W. (1972). *Sociolinguistic patterns*. Philadelphia, PA: University of Pennsylvania Press.
- Laerd Statistics. (2013). One-way ANOVA. Retrieved from <https://statistics.laerd.com/statistical-guides/one-way-anova-statistical-guide-4.php>
- Lake, R. (2002-03). Enhancing acquisition through music. *The Journal of the Imagination in Language Learning*, 7. Retrieved from <http://www.njcu.edu/cill/journal-index.html>.
- Lane, L. (2013a). *Focus on pronunciation 1* (3rd ed.). New York: Pearson Education.
- Lane, L. (2013b). *Focus on pronunciation 2* (3rd ed.). New York: Pearson Education.
- Lane, L. (2013c). *Focus on pronunciation 3* (3rd ed.). New York: Pearson Education.
- Lê, M. (1999). The role of music in second language learning: A Vietnamese perspective. *AARE*. Retrieved from <http://www.aare.edu.au/99pap/le99034.htm>

- Lems, K. (1996). For a song: Music across the ESL curriculum. Paper presented at *Annual Meeting of the Teachers of English to Speakers of Other Languages*, 1-18. ED 396524. Retrieved from <http://files.eric.ed.gov/fulltext/ED396524.pdf>
- Lennon, P. (1990). Investigating fluency in EFL: A quantitative approach. *Language Learning*, 40(3): 387-417.
- Lev-Ari, S., & Keysar, B. (2010). Why don't we believe non-native speakers? The influence of accent on credibility. *Journal of Experimental Social Psychology* 46(6), 1093-1096.
- Levis, J. (2005). Changing contexts and shifting paradigms in pronunciation teaching. *TESOL Quarterly*, 39(3), 369-377.
- Levis, J., & Barriuso, T. A. (2012). Nonnative speakers' pronunciation errors in spoken and read English. In J. Levis & K. LeVelle (Eds.). *Proceedings of the 3rd Pronunciation in Second Language Learning and Teaching Conference*, Sept. 2011. (pp. 187-194). Ames, IA: Iowa State University.
- Levitin, D. (2006). *This is your brain on music: The science of a human obsession*. New York: Plume.
- Levitin, D. (2010, December 8) Music and the brain. *Home TVO Main*. Retrieved December 12, 2011, from <http://ww3.tvo.org/video/163937/daniel-levitin-music-and-brain>
- Levitin, D. (2011, December 23) Why music moves us. *Home TVO Main*. Retrieved December 12, 2011, from <http://ww3.tvo.org/video/168528/daniel-levitin-why-music-moves-us>

- Lima, E. F. (2012). A comparative study of the perception of ITAs by native and nonnative undergraduate students. In. J. Levis & K. LeVelle (Eds.). *Proceedings of the 3rd Pronunciation in Second Language Learning and Teaching Conference*, Sept. 2011. (pp. 54-64). Ames, IA: Iowa State University.
- Lin, H., Fan, C., & Chen, C. (1995). *Teaching pronunciation in the learner-centered classroom*. (ERIC Document Reproduction Service No. ED393292)
- Lippi-Green, R. (2012). *English with an accent: Language, ideology and discrimination in the United States* (2nd ed.) New York: Routledge.
- Llisterri, J. (1992). Speaking styles in speech research, *ELSNET/ESCA/SALT Workshop on Integrating Speech and Natural Language*. Dublin, Ireland, 15-17 July 1992. http://liceu.uab.es/~joaquim/publicacions/SpeakingStyles_92.pdf
- Long, M. (1990). Maturational constraints on language development. *Studies in Second Language Acquisition*, 12(3), 251-285.
- Lord, G. (2005). (How) can we teach foreign language pronunciation? On the effects of a Spanish phonetics course. *Hispania*, 88, 557–567.
- Lord, G. (2008). Podcasting communities and second language pronunciation. *Foreign Language Annals*, 41, 374–389.
- Macdonald, D., Yule, G., & Powers, M. (1994). Attempts to improve English L2 pronunciation: The variable effects of different types of instruction. *Language Learning*, 44, 75–100.

- MacIntyre, P. (2007). Willingness to communicate in the second language: Understanding the decision to speak as a volitional process. *The Modern Language Journal*, 91(4), 564-576. Retrieved from http://faculty.cbu.ca/pmacintyre/research_pages/journals/macintyre,2007.pdf
- Makarova, V., & Ryan, S. (2000). Language teaching attitudes from learners' perspectives: a cross-cultural approach. *Speech Communication Education*, 23, 135-165.
- Malmeier, B. (2014). A comparison of linguistic errors of Iranian EFL learners at two different levels of proficiency. *Theory and Practice in Language Studies*, Vol. 4, No. 4, pp. 850-856. doi:10.4304/tpls.4.4.850-856
- Marinova-Todd, S., Marshall, D., & Snow, C. (2012). Three misconceptions about age and L2 learning. *TESOL Quarterly*, 34(1), 9-34.
- Martin, S. (Performer). (2008, July 04). Steve Martin English lesson [Web Video]. Retrieved from <http://www.youtube.com/watch?v=2NOow2ks5RE>
- McFerrin, B., Faculty M., & Donaghy, J. (Performer). (2009, February 01). I can see clearly now [Web Video]. Retrieved from <http://www.youtube.com/watch?v=cXeY696gZW4>
- Miller, J. S. (2012). Teaching pronunciation with phonetics in a beginner French course: Impact of sound perception. In J. Levis & K. LeVelle (Eds.). *Proceedings of the 3rd Pronunciation in Second Language Learning and Teaching Conference*, Sept. 2011. (pp. 109-123). Ames, IA: Iowa State University.
- Miyake, S. (2004, June 22). Pronunciation and music. *Sophia Junior College*. Japan. Web. Retrieved September 13, 2011, from <http://www.jrc.sophia.ac.jp/kiyou/ki24/miyake.pdf>.

- Mora, C. (2000). Foreign language acquisition and melody singing. *ELT Journal*, 54(2), 146-52.
- Morley, J. (1991) The pronunciation component in teaching English to speakers of other languages. *TESOL Quarterly*, 25(1), 51-74.
- Muller Levis, G. & Levis, J. (2012). Learning to produce contrastive focus: A study of advanced learners of English. In. J. Levis & K. LeVelle (Eds.). *Proceedings of the 3rd Pronunciation in Second Language Learning and Teaching Conference*, Sept. 2011. (pp. 124-133). Ames, IA: Iowa State University
- Munro, M. (2003). A primer on accent discrimination in the Canadian context. *TESL Canada Journal*, 20(2), 38-51.
- Munro, M. (2011). Intelligibility: Buzzword or buzzworthy?. In. J. Levis & K. LeVelle (Eds.). *Proceedings of the 2nd Pronunciation in Second Language Learning and Teaching Conference*, Sept. 2010. (pp.7-16), Ames, IA: Iowa State University.
- Munro, M., & Derwing, T. (1999). Foreign accent, comprehensibility, and intelligibility in the speech of second language learners. Article reprinted in *Language Learning* , 49, *Supplement 1*, 285-310. [Originally published as Munro, M., & Derwing, T. (1995). *Language Learning*, 45, 73-97.]
- Munro, M., Derwing, T., & Sato, K. (2006). Salient accents, covert attitudes: Consciousness-raising for preservice second language teachers. *Prospect: An Australian Journal of TESOL*, 21, 65-77.
- Murphey, T. (1985). A pop song register: The motherese of adolescents as affective foreigner talk. *TESOL Quarterly*, 793-796.

- Murphey, T. (1990). The song stuck in my head phenomenon: A melodic din in the lad? *System*, 18(1), 53–64. Retrieved from http://www.kuis.ac.jp/~murphey-t/Tim_Murphey/Articles_-_Music_and_Song_files/SongStuckIMHP.pdf
- Murphey, T. (1992). The discourse of pop songs. *TESOL Quarterly*, 26(4), 770-774.
- Murphey, T. (1995). *Music & song* (4. impr.) Oxford: Oxford University Press.
- Murray, K. (2005, June 22). Learning a second language through music. *The Free Library*. Retrieved from [http://www.thefreelibrary.com/Learning a second language through music.-a0136071099](http://www.thefreelibrary.com/Learning+a+second+language+through+music.-a0136071099)
- Newsom. (2006). Post hoc tests. Retrieved from http://www.upa.pdx.edu/IOA/newsom/dal/ho_posthoc.doc
- Osorio, I. (Performer). (2012, May 23). Ingles gratis online pronunciacion schwa [Web Video]. Retrieved from http://www.youtube.com/watch?v=o9bkby_hAD8
- Pardede, P. (2010). The role of pronunciation in a foreign language program. Universitas Kristen Indonesia. Web. Retrieved October 24, 2011, from <http://parlindunganpardede.wordpress.com/2010/10/07/349/>
- Pennington, M. & Ellis, N. (2000). Cantonese speakers' memory for English sentences with prosodic cues. *The Modern Language Journal*, 84, 372-389.
- Piñera, S. (2014). Estrategia nacional de inglés 2014-2030. Retrieved from [http://gestion2010-2014.cumplimiento.gob.cl/wp-content/uploads/2014/03/Estrategia Nacional-de-Ingles-2014-2030.pdf](http://gestion2010-2014.cumplimiento.gob.cl/wp-content/uploads/2014/03/Estrategia+Nacional-de-Ingles-2014-2030.pdf)

- Pullen, E. (2012). Cultural identity, pronunciation, and attitudes of Turkish speakers of English: Language identity in an EFL context. In J. Levis & K. LeVelle (Eds.). *Proceedings of the 3rd Pronunciation in Second Language Learning and Teaching Conference*, Sept. 2011. (pp. 65-83). Ames, IA: Iowa State University.
- Rachel. (Designer). (2010, September 15). American English imitation exercise: People change [Web Video]. Retrieved from <http://www.youtube.com/watch?v=TebNvIrUvXk>
- Rachel. (Designer). (2011, July 12). American English imitation exercise: Peggy's new office [Web Video]. Retrieved from <http://www.youtube.com/watch?v=84BCcZhgpIc>
- Rajagopalan, K. (2010). The soft ideological underbelly of the notion of intelligibility in discussions about 'world Englishes' *Applied Linguistics*, 31(3), 465-470. *first published online April 25, 2010 doi:10.1093/applin/amq014*
- Rengifo, R. (2009). Improving pronunciation through the use of karaoke in an adult English class. *Profile issues in teachers' professional development*, 11, 91-105. Retrieved from <http://www.revistas.unal.edu.co/index.php/profile/article/viewFile/10547/11010>
- Richards, J. (2006). *Communicative language teaching today*. New York: Cambridge University Press. Retrieved from http://www.cambridge.org/other_files/downloads/esl/booklets/Richards-Communicative-Language.pdf

- Richards, J., & Rodgers, T. (2001). *Approaches and methods in language teaching: A description and analysis* (2nd ed.). New York: Cambridge University Press.
- Richards, R. (1993). Music and rhythm in the classroom. *Learn: Playful techniques to accelerated learning*, 109-113. (ERIC Document Reproduction Service No. ED379071)
- Richards & Schmidt. (2002) *Longman dictionary of language teaching and applied linguistics* (3rd ed.). Toronto: Pearson Education Ltd.
- Rossiter, M. (2001). The challenges of classroom-based SLA research. *Applied Language Learning*, 12(1), 31-44.
- Rubin, D. (1992). Nonlanguage factors affecting undergraduates' judgments of nonnative English-speaking teaching assistants. *Research in Higher Education*, 33(4), 511-531.
- Rubin, D. (2012). The power of prejudice in accent perception: Reverse linguistic stereotyping and its impact on listener judgments and decisions. In. J. Levis & K. LeVelle (Eds.). *Proceedings of the 3rd Pronunciation in Second Language Learning and Teaching Conference*, Sept. 2011. (pp. 11- 17). Ames, IA: Iowa State University.
- Saglam, E., Kayaoglu, M., & Aydinli, J. (2010). *Music, language and second language acquisition*. Saarbrücken: Lap Lambert Academic Publishing.
- Saito, K. (2011). Examining the role of explicit phonetic instruction in native-like and comprehensible pronunciation development: An instructed SLA approach to L2 phonology. *Language Awareness*, 20, 45–59.
- Saito, K., & Lyster, R. (2012). Effects of form-focused instruction and corrective

- feedback on L2 pronunciation development of /ɹ/ by Japanese learners of English. *Language Learning*, 62, 595–633. doi:10.1111/j.1467-9922.2011.00639.x
- Salcedo, C. (2010). The effects of songs in the foreign language classroom on text recall, delayed text recall and involuntary mental rehearsal. *Journal of College Teaching & Learning (TLC)*. The Clute Institute. Retrieved from <http://journals.cluteonline.com/index.php/TLC/article/view/126/123>
- Sardegna, V. G. (2011). Pronunciation learning strategies that improve ESL learners' linking. In J. Levis & K. LeVelle (Eds.). *Proceedings of the 2nd Pronunciation in Second Language Learning and Teaching Conference*, Sept. 2010. (pp. 105-121), Ames, IA: Iowa State University.
- Sardegna, V. G., & McGregor, A. (2012). Principles for teaching pronunciation to international teaching assistants. *TESOL SPLIS Newsletter*, 8(1).
- Sardegna, V. G. & McGregor, A. (2013). Scaffolding students' self-regulated efforts for effective pronunciation practice. In J. Levis & K. LeVelle (Eds.). *Proceedings of the 4th Pronunciation in Second Language Learning and Teaching Conference*, Aug. 2012. (pp. 182-193). Ames, IA: Iowa State University.
- Saussure, F. (1959). *Course in general linguistics*. New York: Philosophical Library.
- Schmidt, A., & Sullivan, S. (2003). Clinical training in foreign accent modification: a national survey. *Contemporary Issues in Communication Science and Disorders*, 30, 127-135.

- Schmidt, R. (2010). Attention, awareness, and individual differences in language learning. In W. M. Chan, S. Chi, K. N. Cin, J. Istanto, M. Nagami, J. W. Sew, T. Suthiwan, & I. Walker, *Proceedings of CLaSIC 2010*, Singapore, December 2-4 (pp. 721-737). Singapore: National University of Singapore, Centre for Language Studies.
- Schoepp, K. (2001). Reasons for using songs in the ESL/EFL classroom. *Internet TESL Journal (For ESL/EFL Teachers)*. Retrieved February 9, 2011, from <http://iteslj.org/Articles/Schoepp-Songs.htm>
- Schön, D., Boyer, M., Moreno, S., Besson, M., Peretz, I., & Kolinsky, R. (2008). Songs as an aid for language acquisition. *Cognition*, 106(2), 975-983. *Science Direct*. doi:10.1016/j.cognition.2007.03.005
- Scovel, T. (1978). The effect of affect on foreign language learning: A review of the anxiety research. *Language Learning*, 28, 129–142.
- Setter, J. & Jenkins, J. (2005). Pronunciation. *Language Teaching* 38, 1-17.
- Silver, H. F., Strong, R. W., & Perini, M. J. (2000). *So each may learn: Integrating learning styles and multiple intelligences*. Alexandria, VA: ASCD.
- Smith, M. (2002, 2008). Howard Gardner and multiple intelligences. *The encyclopedia of informal education*, <http://www.infed.org/thinkers/gardner.htm>.
- Sposet, B. (2008) *The role of music in second language acquisition: A bibliographical review of seventy years of research, 1937-2007*. New York: E. Mellen Press.

- Stansell, J. (2005, September 14). The use of music in learning languages. *Home Page, MSTE, University of Illinois*. Web. Retrieved September 29, 2011, from <http://mste.illinois.edu/courses/ci407su02/students/stansell/Literature%20Review%201.htm>
- Stevens, C. (Performer). (2007, November 28). Father and son [Web Video]. Retrieved from <http://www.youtube.com/watch?v=Q29YR5-t3gg>
- Sturm, J. (2013). Liaison in L2 French: The effects of instruction. In J. Levis & K. LeVelle (Eds.). *Proceedings of the 4th Pronunciation in Second Language Learning and Teaching Conference*, Aug. 2013. (pp. 157-166). Ames, IA: Iowa State University.
- Tan, D., & Woodworth, J. (2011, May). *Perfect Pronunciation: Say What? Steps to Teach Integrated Pronunciation Strategies*. Paper presented at *TESL Canada Conference Spring 2011*, Halifax.
- Tate, C. (Performer) (2009, November 28). *Catherine Tate interpreter* [Web Video]. Retrieved from <http://www.youtube.com/watch?v=QNKn5ykP9PU>
- Thomson, R. I. (2013). ESL teachers' beliefs and practices in pronunciation teaching: Confidently right or confidently wrong? In J. Levis & K. LeVelle (Eds.). *Proceedings of the 4th Pronunciation in Second Language Learning and Teaching Conference*. Aug. 2012. (pp. 224-233). Ames, IA: Iowa State University.
- Thomson, R. I., & Derwing, T. M. (2015). The effectiveness of L2 pronunciation instruction: A narrative review. *Applied Linguistics*, 36(3), 326-344. doi:10.1093/applin/amu076

- Trofimovich, P. & Baker, W. (2006). Learning second language suprasegmentals: Effect of L2 experience on prosody and fluency characteristics of L2 speech. *Studies in Second Language Acquisition*, 28(1), 1-30. DOI: 10.1017/S0272263106060013
- University of Chicago (2010, July 20). Foreign accents make speakers seem less truthful to listeners, study finds. *ScienceDaily*. Retrieved December 13, 2010, from <http://www.sciencedaily.com/releases/2010/07/100719164002.htm>
- University of Ottawa (n.d.). *Undergraduate admissions: Language requirements*. Retrieved August 29, 2015, from <http://www.registrar.uottawa.ca/Default.aspx?tabid=4539>
- Van den Berg, L. (2011, January 7). The English rhythm: The beneficial use of music in the acquisition of English stress patterns. BA thesis. Retrieved from [igitur-archive.library.uu.nl/student-theses/2011-0810-200820/The%20English%20Rhythm.pdf](http://digitar.archive.library.uu.nl/student-theses/2011-0810-200820/The%20English%20Rhythm.pdf)
- Vandergrift, L. (2004). Listening to learn or learning to listen? *Annual Review of Applied Linguistics*, 24, 3-25.
- Vandergrift, L. (2007). Recent developments in second and foreign language listening comprehension research. *Language Teaching*, 40, 191-210.
- Vann, B. (2011-2012). *Dissertation and case study handbook*. Williamsburg: University of the Cumberlands, Cumberlands College. Retrieved from www.ucumberlands.edu/academics/.../DissertationHandbook12.pdf
- Véliz-Campos, M. (2011). A critical interrogation of the prevailing teaching model(s) of English pronunciation at teacher-training college level: A Chilean evidence-based study. *Literatura y Lingüística*, 23, 213-236.

- Wang, X., Spence, M. M., Jongman, A., & Sereno, J. A. (1999). Training American listeners to perceive Mandarin tones. *Journal of the Acoustical Society of America*, 106, 3649-3658.
- Wei, M. (2006). A literature review on strategies for teaching pronunciation. *ERIC Education Information Resources Center*. Web. Retrieved from eric.ed.gov/PDFS/ED491566.pdf.
- Weinstein, N. (2001). Whaddaya say: Guided practice in relaxed speech. *Longman*, 119.
- Weinstein, N. (2013, October 19). Whaddaya Say? An interview with Nina Weinstein. *Global English for Business*. Retrieved from <http://globalenglishforbusiness.wordpress.com/>.
- Wilcox, W. (1995). Music cues from classroom singing for second language acquisition: Prosodic memory for pronunciation of target vocabulary by adult non-native English speakers. (Doctoral dissertation, University of Kansas, 1995). *Dissertation Abstracts International* 45, 332.
- Williams, M., & Burden, R. (1997). *Psychology for language teachers: A social constructivist approach*. Cambridge: Cambridge University Press.
- Willing, K. (1998). *Learning styles in adult migrant education*. Adelaide: National Curriculum Resource Centre.
- Wong, R. (1993). Pronunciation myths and facts. *English Teaching Forum*, 45-46.
- Young, N. (Performer). (2009, January 16). Neil Young - Old man [Web Video]. Retrieved from http://www.youtube.com/watch?v=An2a1_Do_fc
- Young Artists for Haiti. (Performer). (2010, March 09). Young Artists for Haiti –

Wavin' flag [Web Video]. Retrieved from
<http://www.youtube.com/watch?v=nB7L1BIDELc>

Zielinski, B. (2012). The social impact of pronunciation difficulties: Confidence and willingness to speak. In. J. Levis & K. LeVelle (Eds.). *Proceedings of the 3rd Pronunciation in Second Language Learning and Teaching Conference*, Sept. 2011. (pp.18-26). Ames, IA: Iowa State University.

Appendices

Appendix A

File Number: 06-12-20

Date (mm/dd/yyyy): 08/28/2012



Université d'Ottawa
Bureau d'éthique et d'intégrité de la recherche

University of Ottawa
Office of Research Ethics and Integrity

Ethics Approval Notice **Social Science and Humanities REB**

Principal Investigator / Supervisor / Co-investigator(s) / Student(s)

<u>First Name</u>	<u>Last Name</u>	<u>Affiliation</u>	<u>Role</u>
Karen	Borland	Arts / Modern Languages and Literatures	Principal Investigator
Rodney	Williamson	Arts / Modern Languages and Literatures	Supervisor

File Number: 06-12-20

Type of Project: PhD Thesis

Title: The Use of Songs in the ESL/EFL Classroom as a Means of Teaching Pronunciation: A Case Study of Chilean University Students

Approval Date (mm/dd/yyyy)	Expiry Date (mm/dd/yyyy)	Approval Type
08/28/2012	08/27/2013	Ia
(Ia: Approval, Ib: Approval for initial stage only)		

Special Conditions / Comments:
N/A



Université d'Ottawa **University of Ottawa**
Bureau d'éthique et d'intégrité de la recherche Office of Research Ethics and Integrity

This is to confirm that the University of Ottawa Research Ethics Board identified above, which operates in accordance with the Tri-Council Policy Statement and other applicable laws and regulations in Ontario, has examined and approved the application for ethical approval for the above named research project as of the Ethics Approval Date indicated for the period above and subject to the conditions listed the section above entitled "Special Conditions / Comments".

During the course of the study the protocol may not be modified without prior written approval from the REB except when necessary to remove subjects from immediate endangerment or when the modification(s) pertain to only administrative or logistical components of the study (e.g. change of telephone number). Investigators must also promptly alert the REB of any changes which increase the risk to participant(s), any changes which considerably affect the conduct of the project, all unanticipated and harmful events that occur, and new information that may negatively affect the conduct of the project and safety of the participant(s). Modifications to the project, information/consent documentation, and/or recruitment documentation, should be submitted to this office for approval using the "Modification to research project" form available at:
<http://www.research.uottawa.ca/ethics/forms.html>

Please submit an annual status report to the Protocol Officer four weeks before the above-referenced expiry date to either close the file or request a renewal of ethics approval. This document can be found at:
<http://www.research.uottawa.ca/ethics/forms.html>

If you have any questions, please do not hesitate to contact the Ethics Office at extension 5387 or by e-mail at: ethics@uOttawa.ca.

Signature:

Protocol Officer for Ethics in Research
For Barbara Graves, Chair of the Social Sciences and Humanities REB



PONTIFICIA UNIVERSIDAD CATÓLICA DE CHILE
FACULTAD DE LETRAS

CARTA AUTORIZACIÓN



José Luis Samaniego Aldazábal, Decano de la Facultad de Letras de la Pontificia Universidad Católica de Chile, autoriza a la señora Karen Borland para que realice la investigación correspondiente a su proyecto de tesis doctoral en esta Facultad.

La profesora Borland podrá acceder a los estudiantes que requiera para su trabajo, de acuerdo con lo conversado con nuestra académica la profesora Dra. Miriam Cid Uribe. Asimismo, dispondrá del espacio físico, acceso a la biblioteca y demás facilidades para el logro de objetivos.

Atentamente,

Santiago, 23 de agosto de 2012

Recruitment Text for Potential Participants

Dear Students,

The purpose of this letter is to invite you to participate in a study that I am conducting in the area of English language teaching. My study aims to assess the influence of pronunciation instruction on correct pronunciation and whether using songs helps this process. In order to gather the right information, I will have to divide participants into three groups of students, two receiving a short course of pronunciation instruction and one not. All groups would be given short listening and speaking tests (which will be recorded) at the beginning and the end of the study, as well as an initial questionnaire and, in the case of the groups receiving instruction, a final questionnaire. Participants will be assigned to groups in random fashion, so which group you are assigned to has nothing to do with your merits, capabilities or level in English. Most importantly, I want to assure you that the input of ALL participants is EQUALLY valuable, so thank you in advance for your help! I'm looking forward to obtaining some very useful information from you, and if that helps me to help you out in your future studies, I shall be very happy indeed! But please be assured too that at no time are you under any obligation to participate, even after signing the consent form. That is, you may withdraw from this study at any time if you wish to do so. Prior to beginning this, Dr. Miriam Cid Uribe and I will meet with you to answer any questions or concerns that you may have. Collectively, it is our goal to obtain results that will be useful to you as language learners and I will be in contact with her and the Facultad de Letras, PUC for active follow-up after the investigation has ended and the results have been processed.

Many thanks,

Karen Borland
Ph.D. candidate
Department of Modern Languages and Literatures
University of Ottawa
Canada

Texto de reclutamiento de posibles participantes

Queridos alumnos:

El propósito de la presente es invitarlos a participar en un estudio que estoy llevando a cabo en el área de la enseñanza de la lengua inglesa. En mi estudio me propongo investigar la influencia de la enseñanza de la pronunciación sobre la adquisición de la pronunciación correcta, y asimismo si el uso de canciones sirve para apoyar este proceso. Para recabar los datos necesarios, tendré que formar tres grupos de participantes, dos de los cuales recibirán un breve curso de enseñanza fonética, y otro que sólo me ayudará con unas breves pruebas y un cuestionario. Todos los grupos harán pruebas de escucha y de habla, las cuales serán grabadas, al principio y al final de la investigación, y completarán cuestionarios, uno inicial y otro final en el caso de los grupos que realizan el curso, y uno inicial solamente en el caso del grupo que no recibe instrucción. Los participantes se repartirán de manera aleatoria entre los diferentes grupos, de manera que el grupo al que les asignamos no tiene nada que ver con sus méritos, sus capacidades ni su nivel de inglés. Lo importante para mí es que quiero asegurarles que la contribución de TODOS los participantes es IGUAL de valiosa, de modo que les quiero agradecer desde ya su muy apreciada ayuda. Con su ayuda espero poder recabar datos de gran utilidad, y si éstos me llevan luego a poder ayudarlos en sus futuros estudios, ¡me sentiré sumamente feliz! Pero también quiero decirles que en ningún momento deben sentirse obligados a participar, incluso después de firmar el formulario de aceptación. Es decir, pueden retirarse del estudio en cualquier momento si quieren. Antes de comenzar con la investigación, la Dra. Miriam Cid Uribe y yo nos reuniremos con ustedes para responder a cualquier pregunta o duda que puedan tener. Juntas esperamos poder obtener datos que sean útiles para su aprendizaje del inglés, y por mi parte mantendré contacto con ella y con la Facultad de Letras de la PUC para un seguimiento activo después de concluir la investigación y procesar los datos.

Muchas gracias,

Karen Borland
Candidata al doctorado
Departamento de Lenguas y Literaturas Modernas
Universidad de Ottawa
Canadá

Consent Forms



Université d'Ottawa • University of Ottawa

Faculté des arts
Langues et littératures modernes

Faculty of Arts
Modern Languages and Literatures

Consent Form

Title of the study: The Use of Songs in the ESL / EFL Classroom as a Means of Teaching
Pronunciation: A Case Study of Chilean University Students

Name of researcher: Karen Borland, Department of Modern Languages and Literatures, Faculty of Arts, University of Ottawa

Telephone: (613) 562-5800 ext. 3746

Name of Supervisor: Professor Rodney Williamson, Department of Modern Languages and Literatures, Faculty of Arts, University of Ottawa

Invitation to Participate: We would like to invite you to participate in the research study mentioned above.

Purpose of the Study: The purpose of the study is to find out the influence of pronunciation classes on the pronunciation of second or foreign language learners, and whether or not songs can help with this.

Participation: Your participation will consist of taking both a listening and speaking test and completing an initial questionnaire. Then you will be randomly assigned to one of three groups. Please know that being assigned to one group or another bears absolutely no relation to your merits, language level, or

abilities. As well, it is important that you know that THE CONTRIBUTION OF ALL GROUPS IS EQUALLY VALUABLE. If you are assigned to group One, you will only be required to take another listening and speaking test at the end of the project period (about two weeks later). If you are assigned to group Two or Three, you will attend and participate in a short, two-week long, pronunciation course, in addition to completing a final questionnaire and final listening and speaking test. The questionnaires will take 15-30 minutes each to complete, and the tests will take 30-40 minutes to complete. Please note that the speaking tests will be recorded and collected as data. The pronunciation classes will be held at the faculty and will be 1 hour and 20 minutes long, possibly during your regular class schedule. (Dates and times to be determined.) During the pronunciation classes there will be various listening, speaking, reading and writing activities.

Risks: Your participation in this study will mean that you participate actively in the pronunciation classes (if you are in group Two or Three) as you would in any other university foreign-language classes. Therefore, there are no potential perceived risks. Participating in this study might mean that you would need to spend extra time at the university if it is not possible to schedule the pronunciation classes and tests during your regular university program classes.

Benefits: Your participation in this study means that you might receive some pronunciation training during your first year of English studies. This additional training early on in the program could be of help in preparing you for your third-year English Phonetics and Phonology course. In addition, you may be exposed to additional English culture and language, which will also support and reinforce the material that you are learning in your other English-language courses. Since Chile is a country whose education system is dedicated to the teaching and learning of English with the goal of becoming a bilingual country, we hope that this study will be useful because it aims to further knowledge with respect to pronunciation teaching and the role that songs can play in this.

Confidentiality and anonymity: Rest assured that all information that you share with us will remain strictly confidential. The contents will be used only for academic purposes and your confidentiality will be protected. Your name will never be used nor will any personally identifiable characteristics. Your anonymity will be protected in the following manner: Participants will be labeled as opposed to identified. That is, based on the initial questionnaire results, you will be placed into one of three groups. Within each group, you will be labeled using letters and numbers. Only the researcher, supervisor and Dr. Miriam Cid Uribe will have access to your identity and it will not be revealed in publications.

Conservation of data: The data collected, that is, the questionnaires, the listening tests and the recordings of the speaking tests will be kept in locked filing cabinets in locked offices. Electronic material will be stored in computers equipped with functioning initial entry passwords, firewalls and antivirus software. All data will be conserved for a period of 5 years after the thesis has been completed and only the researcher, supervisor and Dr. Miriam Cid Uribe will have access to it.

Voluntary Participation: You are under no obligation to participate and if you choose to participate, you can withdraw from the study at any time and/or refuse to answer any questions without suffering any negative consequences. If you choose to withdraw, all data gathered until the time of withdrawal will be eliminated from the study and securely discarded by means of shredding and/or secure electronic deletion.

Acceptance: Your signature below indicates that you agree to participate in the above research study conducted by Karen Borland of the Department of Modern Languages and Literatures, Faculty of Arts, University of Ottawa, whose research is under the supervision of Professor Rodney Williamson. If you have any questions about the study, please feel free to contact the researcher or her supervisor.

If you have any questions regarding the ethical conduct of this study, you may contact the Protocol Officer for Ethics in Research, University of Ottawa, Tabaret Hall, 550 Cumberland Street, Room 154, Ottawa, ON K1N 6N5. Tel.: (613) 562-5387. Email: ethics@uottawa.ca

There are two copies of the consent form, one of which is yours to keep.

Participant's signature: *(Signature)* Date:

Researcher's signature: *(Signature)* Date:



Université d'Ottawa • University of Ottawa

Faculté des arts
Langues et littératures modernes

Faculty of Arts
Modern Languages and Literatures

Formulario de aceptación

Título del estudio: El uso de canciones como medio para enseñar la pronunciación en los cursos de inglés como segunda lengua o como lengua extranjera: un estudio de caso con alumnos universitarios chilenos

Nombre de la investigadora:

Karen Borland, Departamento de Lenguas y Literaturas Modernas,

Facultad de Letras, Universidad de Ottawa

Teléfono: 1-613-562-5800 ext. 3746

Nombre del director de tesis:

Dr. Rodney Williamson, Departamento de Lenguas y Literaturas Modernas,

Facultad de Letras, Universidad de Ottawa

Invitación a participar: Queremos invitarte a participar como informante en la investigación mencionada arriba.

Propósito del estudio: El propósito de este estudio es evaluar el efecto de las clases de enseñanza fonética sobre la pronunciación de estudiantes del inglés como segunda lengua o como lengua extranjera, y si el uso de canciones sirve para apoyar este proceso.

Participación: Tu participación consistirá en hacer pruebas de habla y escucha, y completar un cuestionario inicial. Después serás asignado/a uno de tres grupos. Como los participantes se repartirán aleatoriamente entre los tres grupos, el grupo asignado no tiene nada que ver con tus méritos, tus capacidades ni tu nivel de inglés. Además, es importante que sepas que LA CONTRIBUCIÓN DE TODOS LOS GRUPOS ES IGUAL DE VALIOSA. Si te asignamos al grupo uno, sólo tendrás que hacer otra prueba de habla y escucha al final del período de investigación (aproximadamente dos semanas más tarde). Si te asignamos al grupo dos o al tres, te pediremos seguir un breve curso de dos semanas de pronunciación inglesa, además de completar un cuestionario final y las pruebas finales de habla y escucha. Completar los cuestionarios les llevará de 15 a 30 minutos, y las pruebas de 30 a 40 minutos. Las pruebas de habla serán grabadas y las grabaciones serán conservadas como datos del estudio. El curso de enseñanza fonética se llevará a cabo en la Facultad, y cada clase durará 1 hora y 20 minutos. Intentaremos integrarlo a su programa normal de cursos (fechas y horario por determinarse). Las clases de pronunciación inglesa incluirán diversas tareas de habla, escucha, lectura y escritura.

Riesgos: Tu participación en este estudio sólo implica hacer una pruebas y seguir un curso de pronunciación inglesa (si te asignamos al grupo dos o al tres) comparable con cualquier otro curso universitario de lengua extranjera. Por lo tanto, no corren ningún riesgo previsible. El único inconveniente posible sería tener que pasar unas horas extra en la Universidad si no lográramos integrar las clases y pruebas de pronunciación en tu horario normal de cursos.

Beneficios: Tu participación en este estudio podría darte la ventaja de un poco de enseñanza adicional para mejorar tu pronunciación durante tu primer año de estudios de lengua inglesa. Estas clases adicionales podrían ayudarte a la hora de prepararte para tu curso de fonética y fonología del inglés en el tercer año. Por otra parte, te presentarán nuevos elementos de cultura y lengua inglesas, los cuales podrían ayudar a apoyar y reforzar los materiales de aprendizaje de otros cursos que estás tomando. Ya que el sistema educativo de Chile tiene la meta de enseñar el inglés a todos los ciudadanos para que el país sea bilingüe, esperamos que este estudio haga una contribución útil en este contexto porque busca mejorar nuestro conocimiento de la eficacia de la enseñanza fonética y del papel que juegan las canciones en el proceso de mejorar la pronunciación.

Confidencialidad y anonimidad: Puedes tener la seguridad de que mantendremos la confidencialidad de toda la información que compartas con nosotras. Los datos del proyecto sólo se emplearán para fines académicos, y tus datos personales serán protegidos. Nunca emplearemos tu nombre ni otros datos de identidad personal. La anonimidad de tus datos se garantizará de la siguiente manera: a cada participante se identificará mediante un código de letras y números de acuerdo con el grupo asignado. Sólo la investigadora, su director de tesis y la profesora responsable de la PUC, la Dra. Miriam Cid Uribe, tendrán acceso a sus datos personales y éstos nunca se emplearán en publicaciones.

Conservación de los datos: Los datos recabados, es decir, los cuestionarios, las pruebas de escucha y las grabaciones de las pruebas de habla se conservarán bajo llave en archiveros en oficinas cerradas. Los materiales electrónicos se almacenarán en computadores protegidos con contraseña, cortafuegos y programas antivirus. Los datos serán conservados por un período de cinco años después de completar la tesis, y sólo la investigadora, el director de tesis y la profesora responsable de la PUC, la Dra. Miriam Cid Uribe, tendrán acceso a ellos.

Participación libre y voluntaria: No tienes ninguna obligación de participar en el estudio y puedes suspender tu participación en cualquier momento o negarte a contestar cualquier pregunta sin sufrir consecuencias negativas. Si decides retirarte del estudio, todos los datos tuyos que se hayan recabado hasta ese momento serán eliminados, destruidos o borrados utilizando métodos físicos y electrónicos seguros.

Aceptación: Firmando abajo indicarás que aceptas participar en la investigación de referencia llevada a cabo por Karen Borland del Departamento de Lenguas y Literaturas Modernas de la Facultad de Letras de la Universidad de Ottawa, bajo la dirección del Dr. Rodney Williamson. Si tienes preguntas sobre el estudio, te invitamos a ponerte en contacto con nosotros, sea con la investigadora misma, sea con su director.

Si tienes preguntas sobre la ética de este estudio, puedes ponerte en contacto con el Protocol Officer for Ethics in Research, University of Ottawa, Tabaret Hall, 550 Cumberland Street, Room 154, Ottawa, ON K1N 6N5. Tel.: (613) 562-5387. Correo electrónico: ethics@uottawa.ca

El presente formulario de aceptación se prepara con dos copias, una para cada uno de los signatarios.

Firma del/ de la participante: *(Firma)* Fecha:

Firma de la investigadora: *(Firma)* Fecha:

**INITIAL ENGLISH PRONUNCIATION QUESTIONNAIRE / CUESTIONARIO INICIAL DE
LA PRONUNCIACIÓN INGLESA**

DATE / FECHA: _____

Please help us by taking the time to carefully answer the following questions about learning English. This questionnaire is part of a research project conducted at the Department of Modern Languages and Literatures, University of Ottawa. Please be honest when you answer the questions and please answer every question, as this will help with the success of the investigation. Thank you very much for your help.
/ Le agradecería mucho que tomara el tiempo para contestar cuidadosamente las siguientes preguntas acerca del aprendizaje del inglés. Este cuestionario forma parte de una investigación llevada a cabo en el Department of Modern Languages and Literatures de la University of Ottawa. Por favor, conteste cada una de las preguntas con sinceridad, porque de esta manera contribuirá al éxito de la investigación. Muchísimas gracias por su ayuda.

PART ONE / PRIMERA PARTE

1. Name / Nombre: _____

2. Age / Edad: _____

3. Sex / Género: ☐ Male / Masculino ☐ Female / Feminino

4. Where were you born? / ¿Dónde nació?

City / Ciudad _____ Country / País _____

5. Where have you lived for at least five years / ¿Dónde ha vivido por un mínimo de cinco años?

City / Ciudad _____ Country / País _____ Years / Años _____

City / Ciudad _____ Country / País _____ Years / Años _____

City / Ciudad _____ Country / País _____ Years / Años _____

6. Have you ever lived in another country? / ¿Alguna vez ha vivido en otro país?

☐ Yes / Sí ☐ No

a) If yes, what country did you live in? / Si contestó afirmativamente, ¿en cuál país vivió?

b) How long did you live there? / ¿Por cuánto tiempo vivió allá? e.g. 2.5 years / años

c) How old were you? / ¿Cuántos años tenía usted? e.g. 2-4 years old / años

7. Is Spanish your native language? / ¿Es el español su lengua materna? Yes / Sí ☐ No

8. Do you speak another language other than English or Spanish? / ¿Habla usted otro idioma que no sea el inglés o español? Yes / Sí ☐ No

a) If yes, indicate the language(s) and check your proficiency level. / Si contestó afirmativamente, indica la(s) lengua(s) junto con su nivel de competencia en ellas.

Language / Lengua: _____

- ☐ Advanced / Avanzado
- ☐ Intermediate / Intermedio
- ☐ Pre-intermediate / Pre-intermedio
- ☐ Beginner / Básico
- ☐ Only a few words and phrases / Solamente algunas palabras y frases

Language / Lengua: _____

- ☐ Advanced / Avanzado
- ☐ Intermediate / Intermedio
- ☐ Beginner / Básico
- ☐ Only a few words and phrases / Solamente algunas palabras y frases

9. Why are you learning English? Check all that apply. / ¿Por qué está aprendiendo el inglés?
Marque todas las que sean pertinentes.

- ☐ I like English. / Me gusta el inglés.
- ☐ I want to travel to an English-speaking country. / Quiero viajar a un país donde se habla inglés.
- ☐ I will need it for my job. / Lo voy a necesitar en mi trabajo.
- ☐ I need it for my future studies. / Lo voy a necesitar para mis estudios en el futuro.
- ☐ _____ Otra razón (Especificar) _____
- ☐ _____ Otra razón (Especificar) _____

10. What do you want to do with your English in the future? / ¿Qué es lo que quiere hacer con su inglés en el futuro?

11. Have you ever taken classes in English pronunciation? ☐ Yes / Sí ☐ No

a) If yes, indicate where and for how long. / Si contestó afirmativamente, indica dónde y por cuánto tiempo.

Where? / ¿Dónde? _____

For how long? / ¿Por cuánto tiempo? _____

12. Please check any English accents that you would be happy to learn. / Por favor, indique los acentos que estaría contento aprender.

☐ General American / Acento norte americano	☐ Received Pronunciation / Pronunciación
--	---

general	estándar del inglés británico
☐ Australian / Australiano	☐ Scottish / Escocés
☐ Canadian / Canadiense	☐ Irish / Irlandés
☐ South African / Sudafricano	☐ New Zealand / Nueva Zelandia
☐ Other / Otro _____	☐ Other / Otro _____

13. Please check any English accents you would not like to learn. / Por favor, indique los acentos que no le gustaría aprender.

☐ General American / Acento norte americano general	☐ Received Pronunciation / Pronunciación estándar del inglés británico
☐ Australian / Australiano	☐ Scottish / Escocés
☐ Canadian / Canadiense	☐ Irish / Irlandés
☐ South African / Sudafricano	☐ New Zealand / Nueva Zelandia
☐ Other / Otro _____	☐ Other / Otro _____

14. Do you like the music of English-speaking countries? / ¿Le gusta la música de los países ingleses? ☐ Yes / Sí ☐ No ☐ I don't know. / No sé.

PART TWO / SEGUNDA PARTE

In this part, please indicate how much you agree or disagree with the following statements. / En esta parte, indique, por favor, cuánto está de acuerdo o desacuerdo con las siguientes declaraciones.

Strongly disagree / Totalmente en desacuerdo	Disagree / En desacuerdo	Slightly disagree / Ligeramente en desacuerdo	Slightly agree / Ligeramente de acuerdo	Agree / De acuerdo	Strongly agree / Totalmente de acuerdo
1	2	3	4	5	6

1. When learning a foreign language, it is important to try to speak with a good accent. / Cuando uno aprende una lengua extranjera, es importante tratar de hablar con un buen acento.	1	2	3	4	5	6
2. I am a good singer. / Canto bien.	1	2	3	4	5	6
3. I listen to English music. / Escucho la música inglesa.	1	2	3	4	5	6
4. I would like to speak English without a Spanish accent. / Me gustaría hablar el inglés sin acento hispanico.	1	2	3	4	5	6
5. I am comfortable singing in front of other people. / Me siento cómodo cantando frente a los demás.	1	2	3	4	5	6

6. Chileans should try to speak English without a Spanish accent. / Los chilenos deberían tratar de hablar inglés sin acento hispanico.	1	2	3	4	5	6
7. I consider myself a musical person. / Me consider una persona musical.	1	2	3	4	5	6
8. I want to improve my English accent. / Quiero mejorar mi acento en inglés.	1	2	3	4	5	6
9. I would feel comfortable if I spoke English without a Spanish accent. / Me sentiría cómodo si hablara inglés sin acento hispanico.	1	2	3	4	5	6
10. Listening to English music can help me to learn English. / Escuchar la música inglesa me puede ayudar a aprender inglés.	1	2	3	4	5	6
11. I feel that my Spanish accent is a part of who I am. / Siento que mi acento hispanico es parte de mí.	1	2	3	4	5	6
12. It is OK if I speak English with a Spanish accent. / Está bien si hablo inglés con un acento hispanico.	1	2	3	4	5	6
13. I want native English speakers to understand me when I speak English. / Quiero que los hablantes nativos del inglés me entiendan cuando hablo inglés.	1	2	3	4	5	6
14. I play one or more musical instruments. / Toco uno o varios instrumentos musicales.	1	2	3	4	5	6
15. I want to keep my Spanish accent when I speak English. / Quiero mantener mi acento hispanico cuando hablo inglés.	1	2	3	4	5	6
16. I think it is possible to improve my English accent. / Creo que es possible mejorar mi acento en inglés.	1	2	3	4	5	6
17. I would like to sound like a native English speaker. / Me gustaría pronunciar como un hablante nativo del inglés.	1	2	3	4	5	6
18. I can understand the lyrics when I listen to English songs. / Entiendo las letras cuando escucho las canciones inglesas.	1	2	3	4	5	6
19. I am interested in improving my English pronunciation. / Me interesa mejorar mi pronunciación inglesa.	1	2	3	4	5	6
20. I am confident speaking English to a native speaker. / Me siento seguro/a de mí mismo/a cuando hablo inglés con un hablante nativo.	1	2	3	4	5	6
21. If I spoke English like a native English speaker, I would feel like I was another person. / Si hablara inglés como un hablante nativo del inglés, me sentiría como si fuera otra persona.	1	2	3	4	5	6
22. Listening to English music can help me improve my English pronunciation. /	1	2	3	4	5	6

Escuchar la música inglesa me puede ayudar a mejorar mi pronunciación del inglés.						
---	--	--	--	--	--	--

PART THREE / TERCERA PARTE

1. Please list up to 10 of your favourite English songs and the artist or group who sings them. / Por favor, apunte hasta 10 de sus canciones inglesas favoritas y el cantante o grupo musical que las cantan.

Name of Song / Nombre de la canción

Group / Artist / Grupo / Artista musical

2. Would you be willing to attend a two-week English pronunciation course in October? It would be offered at the faculty outside of your regular class hours. Actual schedule to be determined. / ¿Estaría dispuesto a asistir a un curso de la pronunciación inglesa que duraría dos semanas en octubre? Se lo ofrecería donde la facultad fuera del horario de sus clases regulares. Horario por determinar.

☐ Yes / Sí ☐ No ☐ I don't know. / No sé.

THANK YOU! / ¡GRACIAS!



If you have any questions, I can be contacted at the following email address:

Su usted tiene cualquier duda, me puede contactar a la siguiente dirección electrónica:



Université d'Ottawa • University of Ottawa

Faculté des arts
Langues et littératures modernes

Faculty of Arts
Modern Languages and Literatures

Information Form for Group One

**Title of the study: The Use of Songs in the ESL / EFL Classroom as a Means of Teaching
Pronunciation: A Case Study of Chilean University Students**

Name of researcher: Karen Borland, Department of Modern Languages and Literatures, Faculty of Arts, University of Ottawa - Telephone: 1-613-562-5800 ext. 3746

Name of assisting PUC professor and member of thesis committee:

Dr. Miriam Cid Uribe, M.A., Ph.D., Associate Professor, Director of M.A.
Programme in Applied Linguistics, PUC de Chile

Dear Student,

First of all, I wish to thank you for agreeing to take part in this study! Your participation as a member of one of the groups is valuable and will contribute to a greater understanding of the effect of pronunciation classes on the listening and speaking skills of adult language learners. As a member of one of the groups, your participation will consist of the following:

- ❖ Filling out an initial questionnaire (15-30 minutes) – Where: _____ When: _____
- ❖ Taking a short (initial) listening test (15-20 minutes) – Where: _____ When: _____
- ❖ Taking a short (initial) speaking test (15-20 minutes) – Where: _____ When: _____
- ❖ Taking a short (final) listening test (15-20 minutes) – Where: _____ When: _____
- ❖ Taking a short (final) speaking test (15-20 minutes) – Where: _____ When: _____
- ❖ Attending a follow-up Question and Answer / Information session (30-45 minutes)
 - Where: _____ When: _____

Your input is especially valuable to us so please perform to the best of your ability on the tests and answer the questionnaire as honestly as possible.

Also, please remember that even though you have signed the consent form, you are free to discontinue your participation at any time if you are unable or unwilling to continue to be a part of this study.

Sincerely,

Karen Borland

Hoja informativa para el grupo uno

Título del estudio: El uso de canciones como medio para enseñar la pronunciación en los cursos de inglés como segunda lengua o como lengua extranjera

Nombre de la investigadora: Karen Borland, Departamento de Lenguas y Literaturas Modernas,
Facultad de Letras, Universidad de Ottawa - Teléfono: 1-613-562-5800 ext. 3746

Nombre de la profesora responsable en la PUC (miembro del comité de tesis doctoral)

Dr. Miriam Cid Uribe, M.A., Ph.D., Profesora de Estado en Inglés, Directora del
Programa de Master en Lingüística Aplicada, PUC de Chile

Querido/a alumno/a:

En primer lugar, quiero darte las gracias por haber aceptado colaborar en este estudio. Tu contribución será muy valiosa y nos ayudará a comprender mejor el efecto de la enseñanza fonética sobre la competencia oral y auditiva de estudiantes adultos que aprenden el inglés. Como participante de uno de los grupos, tu contribución consistirá en lo siguiente:

- | | | |
|--|--------------|-------------|
| ❖ Completar un cuestionario inicial (15-30 minutos) | Lugar: _____ | Hora: _____ |
| ❖ Hacer una breve prueba auditiva inicial (15-20 minutos) | Lugar: _____ | Hora: _____ |
| ❖ Hacer una breve prueba oral inicial (15-20 minutos) | Lugar: _____ | Hora: _____ |
| ❖ Hacer una breve prueba auditiva final (15-20 minutos) | Lugar: _____ | Hora: _____ |
| ❖ Hacer una breve prueba oral final (15-20 minutos) | Lugar: _____ | Hora: _____ |
| ❖ Asistir a una sesión final de información (preguntas y respuestas) (30-45 minutos) | | |

Lugar: _____ Hora:

Tu contribución es sumamente valiosa, por lo cual te pedimos que hagas las pruebas lo mejor que puedas y que contestes las preguntas del cuestionario con toda sinceridad.

Te recuerdo también que, aun después de haber firmado el formulario de aceptación, tienes el derecho de suspender tu participación en este estudio en cualquier momento, si no quieres o no puedes continuar.

Cordialmente,

Karen Borland



Université d'Ottawa • University of Ottawa

Faculté des arts
Langues et littératures modernes

Faculty of Arts
Modern Languages and Literatures

Information Form for Groups Two and Three

**Title of the study: The Use of Songs in the ESL / EFL Classroom as a Means of Teaching
Pronunciation: A Case Study of Chilean University Students**

Name of researcher: Karen Borland, Department of Modern Languages and Literatures, Faculty of Arts, University of Ottawa - Telephone: 1-613-562-5800 ext. 3746

Name of assisting PUC professor and member of thesis committee:

Dr. Miriam Cid Uribe, M.A., Ph.D., Associate Professor, Director of M.A.
Programme in Applied Linguistics, PUC de Chile

Dear Student,

First of all, I wish to thank you for agreeing to take part in this study! Your participation is valuable and will contribute to a greater understanding of the effect of pronunciation classes on the listening and speaking skills of adult English language learners. As a member of one of the pronunciation instruction groups, your participation will consist of the following:

- ❖ Filling out an initial questionnaire (15-30 minutes) – Where: _____ When: _____
- ❖ Taking a short (initial) listening test (15-20 minutes) – Where: _____ When: _____

- ❖ Taking a short (initial) speaking test (15-20 minutes) – Where: _____ When: _____
- ❖ Attending English pronunciation instruction classes over a two-week period (10-13 hours in total). In the pronunciation classes, which will be 1 hour and 20 minutes in length, you will be actively participating in various listening, speaking, reading and writing tasks which will vary from class to class.

– Where: _____ When: _____

- ❖ Taking a short (final) listening test (15-20 minutes) – Where: _____ When: _____
- ❖ Taking a short (final) speaking test (15-20 minutes) – Where: _____ When: _____
- ❖ Filling out a final questionnaire (15-20 minutes) – Where: _____ When: _____
- ❖ Attending a follow-up Question and Answer / Information session (30-45 minutes)

– Where: _____ When: _____

Your input is especially valuable to us so please perform to the best of your ability on the tests and answer the questionnaires as honestly as possible.

Also, please remember that even though you have signed the consent form, you are free to discontinue your participation at any time if you are unable or unwilling to continue to be a part of this study.

Sincerely,

Karen Borland

Hoja informativa para los grupos dos y tres

Título del estudio: El uso de canciones como medio para enseñar la pronunciación en los cursos de inglés como segunda lengua o como lengua extranjera

Nombre de la investigadora: Karen Borland, Departamento de Lenguas y Literaturas Modernas,
Facultad de Letras, Universidad de Ottawa - Teléfono: 1-613-562-5800 ext. 3746

Nombre de la profesora responsable en la PUC (miembro del comité de tesis doctoral)

Dr. Miriam Cid Uribe, M.A., Ph.D., Profesora de Estado en Inglés, Directora del
Programa de Master en Lingüística Aplicada, PUC de Chile

Querido/a alumno/a:

En primer lugar, quiero darte las gracias por haber aceptado colaborar en este estudio. Tu contribución será muy valiosa y nos ayudará a comprender mejor el efecto de la enseñanza fonética sobre la competencia oral y auditiva de estudiantes adultos que aprenden el inglés. Como participante de uno de los grupos que seguirán el curso de enseñanza fonética, tu contribución consistirá en lo siguiente:

- ❖ Completar un cuestionario inicial (15-30 minutos) Lugar: _____ Hora: _____
- ❖ Hacer una breve prueba auditiva inicial (15-20 minutos) Lugar: _____ Hora: _____
- ❖ Hacer una breve prueba oral inicial (15-20 minutos) Lugar: _____ Hora: _____
- ❖ Asistir a clases de pronunciación inglesa durante un período de dos semanas (10-13 horas en total). En las clases de pronunciación, cada una de una duración de 1 hora y 20 minutos, participarás activamente en distintas tareas de escucha, habla, lectura y escritura que variarán de una clase a otra. Lugar: _____ Hora: _____
- ❖ Hacer una breve prueba auditiva final (15-20 minutos) Lugar: _____ Hora: _____

- ❖ Hacer una breve prueba oral final (15-20 minutos) Lugar: _____ Hora: _____
 - ❖ Completar un cuestionario final (15-20 minutos) Lugar: _____ Hora: _____
 - ❖ Asistir a una sesión final de información (preguntas y respuestas) (30-45 minutos) Lugar: _____ Hora: _____
- _____

Tu contribución es sumamente valiosa, por lo cual te pedimos que hagas las pruebas lo mejor que puedas y que contestes las preguntas de los cuestionarios con toda sinceridad.

Te recuerdo también que, aun después de haber firmado el formulario de aceptación, tienes el derecho de suspender tu participación en este estudio en cualquier momento, si no quieres o no puedes continuar.

Cordialmente,

Karen Borland

Final English Pronunciation Questionnaire

FINAL ENGLISH PRONUNCIATION QUESTIONNAIRE / CUESTIONARIO FINAL DE LA PRONUNCIACIÓN INGLESA

Now that you have finished taking the English pronunciation course, please take the time to carefully answer the following questions. This questionnaire is part of a research project conducted at the Department of Modern Languages and Literatures, University of Ottawa. Please answer every question honestly, as this will help with the success of the investigation. Thank you very much for your help. / Ya que usted ha completado el curso de pronunciación inglesa, le agradecería mucho que tomara el tiempo para contestar cuidadosamente las siguientes preguntas acerca del aprendizaje del inglés. Este cuestionario forma parte de una investigación llevada a cabo en el Department of Modern Languages and Literatures de la University of Ottawa. Por favor, conteste cada una de las preguntas con sinceridad, porque de esta manera contribuirá al éxito de la investigación. Muchísimas gracias por su ayuda.

Name / Nombre: _____

Please indicate how much you agree or disagree with the following statements. / Indique, por favor, cuánto está de acuerdo o desacuerdo con las siguientes declaraciones.

Strongly disagree / Totalmente en desacuerdo	Disagree / En desacuerdo	Slightly disagree / Ligeramente de desacuerdo	Slightly agree / Ligeramente de acuerdo	Agree / De acuerdo	Strongly agree / Totalmente de acuerdo
1	2	3	4	5	6

1. When learning a foreign language, it is important to try to speak with a good accent. / Cuando uno aprende una lengua extranjera, es importante tratar de hablar con un buen acento.	1	2	3	4	5	6
2. I enjoyed the pronunciation course. / Me gustó el curso de pronunciación.	1	2	3	4	5	6
3. Taking the pronunciation course helped improve my English pronunciation. / El curso de pronunciación me ayudó a mejorar mi pronunciación del inglés.	1	2	3	4	5	6
4. I would like to speak English without a Spanish accent. / Me gustaría hablar el inglés sin acento hispanico.	1	2	3	4	5	6
5. I am comfortable singing in front of other people. / Me siento cómodo cantando frente a los demás.	1	2	3	4	5	6
6. Chileans should try to speak English without a Spanish accent. / Los chilenos	1	2	3	4	5	6

deberían tratar de hablar inglés sin acento hispanico						
7. Taking the pronunciation course helped improve my listening comprehension. / El curso de pronunciación me ayudó a mejorar mi comprensión del inglés.	1	2	3	4	5	6
8. I want to improve my English accent. / Quiero mejorar mi acento en inglés.	1	2	3	4	5	6
9. I would feel comfortable if I spoke English without a Spanish accent. / Me sentiría cómodo si hablara inglés sin acento hispanico.	1	2	3	4	5	6
10. Listening to English music can help me to learn English. / Escuchar la música inglesa me puede ayudar a aprender inglés.	1	2	3	4	5	6
11. I feel that my Spanish accent is a part of who I am. / Siento que mi acento hispanico es parte de mí identidad.	1	2	3	4	5	6
12. It is OK if I speak English with a Spanish accent. / Está bien si hablo inglés con un acento hispanico.	1	2	3	4	5	6
13. I want native English speakers to understand me when I speak English. / Quiero que los hablantes nativos del inglés me entiendan cuando hablo inglés.	1	2	3	4	5	6
14. I think my English pronunciation will continue to improve because I took this course. / Creo que seguiré mejorando mi pronunciación del inglés por haber tomado este curso.	1	2	3	4	5	6
15. I want to keep my Spanish accent when I speak English. / Quiero mantener mi acento hispanico cuando hablo inglés.	1	2	3	4	5	6
16. I think it is possible to improve my English accent. / Creo que es posible mejorar mi acento en inglés.	1	2	3	4	5	6
17. I would like to sound like a native English speaker. / Me gustaría pronunciar como un hablante nativo del inglés.	1	2	3	4	5	6
18. I can understand the lyrics when I listen to English songs. / Entiendo las letras cuando escucho las canciones inglesas.	1	2	3	4	5	6
19. I am interested in improving my English pronunciation. / Me interesa mejorar mi pronunciación en inglés.	1	2	3	4	5	6
20. I am confident speaking English to a native speaker. / Me siento segura de mí misma cuando hablo inglés con un hablante nativo.	1	2	3	4	5	6
21. If I spoke English like a native English speaker, I would feel like I was another person. / Si hablara inglés como un hablante nativo del inglés, me sentiría como si fuera otra persona.	1	2	3	4	5	6

22. Listening to English music can help me improve my English pronunciation. / Escuchar la música inglesa me puede ayudar a mejorar mi pronunciación inglesa.	1	2	3	4	5	6
--	---	---	---	---	---	---

COMMENTS. / COMENTARIOS.

Do you grant permission for your comments to be quoted in the final thesis? Your name will not be mentioned. / ¿Me da el permiso de citar sus comentarios en la tesis final? No voy a mencionar su nombre. ☐ Yes / Sí ☐ No

THANK YOU! / ¡GRACIAS!



If you have any questions, I can be contacted at the following email address:

Su usted tiene cualquier duda, me puede contactar al siguiente dirección electrónica:

LESSON PLAN - DAY # 1 - “I can do that”	Date: October 1 st , 2012
Pronunciation Areas: Introduction and raising awareness; Thought Groups	
Anticipated Problems for Students: - Ss might be shy. - Ss might have problems with reverse accent mimicry.	Solution: - Encourage ss to relax, have fun, and not worry about making mistakes, and that we learn more sometimes by making mistakes. - If they can’t say full sentences, have them think of words. Or, they can join a pair who is getting the concept.
Anticipated Problems for Teacher: - Not knowing the names of the students yet. - Technical problems	Solution: - Take name cards to class. Distribute and then collect for next class. - Check the functioning of the computer and speakers the week before, if possible.
Goals: 1. To raise ss awareness of how different languages differ with respect to their sounds and delivery. 2. To teach what causes foreign accents. 3. To introduce the Prosody Pyramid and teach Thought Groups. Objectives: 1. Ss will produce stereotypical English sounds. 2. Ss will mimic North Americans speaking	Complete List of Materials • Lesson Plan • Attendance form • Class copies of the course outline • Name cards • Lesson 1 - PPP • Steve Martin – Pink Panther video • Catherine Tate video - The Interpreter • (S) Cat Stevens – Father and Son (audio, lyrics, answer key) • (NS) Thought Groups and Focus Words, Chapter 8 (both audios, pp. 103-104, answer key)

Spanish with a thick English accent. 3. Ss will produce a phrase with automaticity through Quality Repetition. 4. Ss will separate Thought Groups while listening to and while reading a SONG/TEXT.	
Homework: Observe and note North Americans speaking English. Can you check a movie clip, a TV show or remember things that you have noticed before? Are North Americans loud? Do they stand far apart? Do they use their hands a lot? Do they make faces or move their heads a lot? What characteristic sounds do they make? Etc. Write down at least 5 things that you notice about North Americans when they speak. Please hand this in next class.	Feedback:

AT START OF CLASS:

- Before class: Put up the first slide of the PPP which shows the plan for the day.
- Greet the students, pass out the name cards, take attendance.

INTRODUCTION:

- Ice breaker: Play Pink Panther video to show students how the class will *not* be.
- Distribute Course Outline and go over.
- Ask ss to have fun and laugh during the course; tell them that they will feel strange making new sounds but that the strangeness will disappear
- Tell them that the focus is not on learning new vocabulary or grammar, although they might learn some. The focus is on hearing how English sounds and how to say its sounds and sound combinations. The course is probably not going to be difficult, but it will be effective.

Activity 1: *What does North American English sound like?*

Suggested Time:

- How different languages sound: Show Catherine Tate video.
- Then, in small groups, have ss ‘pretend’ that they are speaking American English, like the woman in the video did for other languages.
- Brainstorm ss on the sounds that they are making. Write these in one column on the board. Tell ss that these are some of the characteristic sounds of English.

Activity 2: *How does an English-speaker sound when they are trying to speak Spanish?*

Suggested Time:

- ss practice mimicking/making fun of an American speaking Spanish (Give ss an example of this, “*Joe nou keyero nada. Eso es toe doe.*”)
- Go around and listen in and help.
- Ask ss to say what sounds they are making. Write these in a second column on the board.

Presentation: *The Nature of Foreign Accents*

- Show nature of foreign accents PPP, terminology.

Activity 3: *Quality Repetition phrase: “I can do THAT”*

Suggested Time: 5 minutes

- Don’t write on board. Explain the concept of Quality Repetition, i.e. helps them learn chunks of language with all of the ‘musical’ pronunciation features that we will learn in this course. Tell ss that there will be one of these every day.
- Ask ss to think back to the Catherine Tate video. Do they remember what she said when her boss said he needed someone who could translate into seven different languages? (“*I can do that.*”)
- Say the phrase a number of times and have ss just listen first. Then continue saying and repeating it while they say and repeat it as well. (Make sure they reduce the vowel in ‘can’ and that ‘that’ is stressed.)
- Randomly check to see if ss got it.

Presentation: *The Prosody Pyramid and Thought Groups*

Suggested Time:

- Go through PPP on the Prosody Pyramid and Thought Groups.

- (S) Play Cat Stevens – Father and Son and have ss separate the thought groups using slashes.
- Take up answers and then have ss take turns reading the song to each other.

- (NS) Go through ch. 8 handouts and have ss separate the thought groups using slashes.
- Take up answers and then have ss take turns reading the text to each other.

Last Activity:

- Brainstorm ss on what they can do to improve their pronunciation during and outside of class.
(Listen to native speakers speaking. Don't worry about understanding the message, but rather listen to the sounds and sound of English; Practice the sounds you learn when speaking and at home; Talk to yourself out loud.)

Homework:

- Show instructions on the PPP.
- Observe and note North Americans speaking English. Can you check a movie clip, a TV show or remember things that you have noticed before? Are North Americans loud? Do they stand far apart? Do they use their hands a lot? Do they make faces or move their heads a lot? What characteristic sounds do they make? Etc. Write down at least 5 things that you notice about North Americans when they speak. Please hand this in next class.

Slide 1

LESSON 1 - Introduction and Raising Awareness; Thought Groups



- Introduction to the course
- How different languages sound
- Making fun of how North American English speakers sound
- The nature of foreign accents
- Quality Repetition phrase of the day
- The Prosody Pyramid and Thought Groups
- Ways to improve your pronunciation
- Homework

Image:

http://sandbox.yoyogames.com/extras/image/name/san2/197/420197/original/Pink-Panther_1_.jpg

Slide 2

The Nature of Foreign Accents



Images:

<http://img.over-blog.com/223x198/0/55/44/49/crackers/xom.jpg>
<http://i.ytimg.com/vi/BmUPnei30c0/0.jpg>
http://api.ning.com/files/AyXbjZP6TRj9P*UCbmNH474jmJSLnTioJlq8LwUJMhiSVzxmP2PkWKPyWWfKry-JppF6535-ufunV92K1xKs2F2gMz*cFtU9/710261928.jpg

Slide 3

Why do we have an accent when we speak a foreign language?

- Languages may have different consonant and vowel sounds. That is, sounds that exist in one language might not exist in another. For example: /r/ does not exist in North American English
- Languages may have similar consonant and vowel sounds, but they may be made in a different way or a different place in our mouths. For example: the Spanish *r* and *d* are made with the tongue against the teeth, but in English, the tongue touches the alveolar ridge. Compare Sp. "todo" and Eng. "to do"

Slide 4

- Two languages might have the same sounds, but they might not correspond to the same letters. For example: /ð/ corresponds to the letter "d" in Spanish, as in "*lado*". In English, this sound corresponds to "th", as in "*the*".
- Languages have different patterns of stress, rhythm and intonation. These contribute to the *music* of the language. If any of these areas are different, then the music of the language is different.

Slide 5

So...

- If we speak a foreign language using the sounds and the music of our native language, then we will have an accent.
- The accent will be strong or weak depending on which and how many pronunciation aspects of our native language we use when speaking the foreign language.
- If our accent is strong, it may be very hard for native speakers or speakers of other languages to understand us.

Slide 6

Terminology

Some words you should know for this course:

- **Segmentals** – Consonant sounds and vowel sounds.
- **Suprasegmentals** – The pronunciation units that extend over more than one sound in an utterance, that is, *rhythm, stress and intonation – the prosody*.

Slide 7

- **Rhythm** – The beat of a language.
- **Stress** – A more energetic or forceful pronunciation of a syllable or word than those around it. A stressed syllable or word may be louder, higher or lower in pitch, or longer than the ones next to it.
- **Intonation** – Patterns of changes in pitch, loudness and rhythm that convey information over and above that which is expressed by the individual words in an utterance.

Slide 8

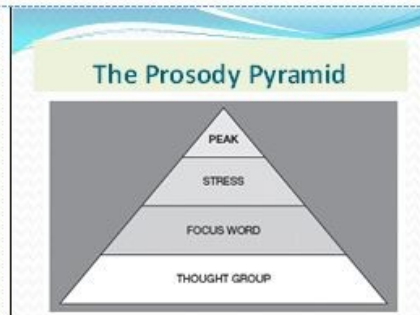
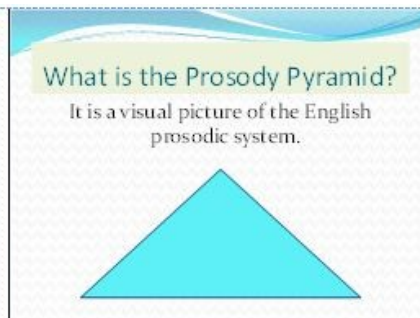


Image:
http://www.teachers2010.xpg.com.br/index_arquivos/image3581.gif

Slide 9



Slide 10

What is a Thought Group?

- It is a group of words (phrase, clause, short sentence).
- English speakers organize their thoughts into these groups of words or chunks when they speak.
- Separating what you say into chunks makes it easier for a listener to understand what you are saying.

Slide 11

Note!

- If you speak word by word, it is actually harder to understand what you are saying.
- Therefore, you need to speak thought by thought.

Slide 12

- In writing, punctuation markers help readers separate thought groups.
- In speaking, prosodic markers are used to separate thought groups. These are:
 - A pause
 - A drop in pitch
 - Lengthening of the last stressed syllable

Slide 13

HOMEWORK

- Observe and note North Americans speaking English. Can you check a movie clip, a TV show or remember things that you have noticed before? Are North Americans loud? Do they stand far apart? Do they use their hands a lot? Do they make faces or move their heads a lot? What characteristic sounds do they make? Etc.
- Write down at least 5 things that you notice about North Americans when they speak. Please hand this in next class.

Instructions: *Listen to the following song and separate the lyrics into thought groups by putting a slash where one thought group ends and the next one begins. The first few have been done for you. (Be careful! Thought groups do not always correspond to written punctuation markers.)*

CAT STEVENS/YUSUF ISLAM (Father and Son)

It's not time / to make a change,/
Just relax,/ take it easy./
You're still young,/ that's your fault,/
There's so much you have to know./
...

The rest of the lyrics to the song can be found at:

<http://www.azlyrics.com/lyrics/catstevens/fatherandson.html>

Instructions: *Listen to the following song and separate the lyrics into thought groups by putting a slash where one thought group ends and the next one begins. The first few have been done for you. (Be careful! Thought groups do not always correspond to written punctuation markers.)*

CAT STEVENS/YUSUF ISLAM (Father and Son)

It's not time / to make a change,/
Just relax,/ take it easy./
You're still young,/ that's your fault,/
There's so much you have to know./
Find a girl, /settle down,/
If you want /you can marry./
Look at me, /I am old,/ but I'm happy./

I was once like you are now, /and I know that it's not easy,/
To be calm /when you've found /something going on./
But take your time,/ think a lot,/
Why, think of everything you've got.
For you will still be here tomorrow/ but your dreams may not./

How can I try to explain,/ when I do /he turns away again./
It's always been the same,/ same old story./
From the moment I could talk /I was ordered to listen.
Now there's a way/ and I know /that I have to go away./
I know/I have to go./

It's not time/ to make a change,/
Just sit down,/ take it slowly./
You're still young,/ that's your fault,/
There's so much you have to go through./
Find a girl, /settle down,/
if you want /you can marry.
Look at me,/ I am old,/ but I'm happy./

All the times /that I cried,/ keeping all the things I knew inside,
It's hard, /but it's harder to ignore it./
If they were right, /I'd agree,/ but it's them they know /not me./
Now there's a way /and I know /that I have to go away.
I know/I have to go./

Thought Groups and Focus Words

Fluent speakers organize their speech into phrases or **thought groups**.

They speak thought by thought.

They. Do. Not. Speak. Word. By. Word.

Thought groups make information easier for the listener to understand. Read these examples. Pause briefly at the end of each thought group.

Phone number: 202 / 555 / 1212

Sentence: Some of my best friends / are people I've met in class.

Every thought group has a **focus word**: one key word with more emphasis than the others.



Some of my **BEST FRIENDS**/are **PEO**ple I've **MET** in **CLASS**.

You might hear other stressed words in a thought group, but generally you will hear only one focus word.

In this chapter, you will learn about dividing speech into thought groups. You will also learn about highlighting the key word or focus word in each thought group.

Listen!

Listening Activity 1



Track 46

Listen to your teacher or the speaker on audio say the phrases. If you hear one thought group, circle 1. If you hear two thought groups, circle 2. The first one has been done.

1. a. seven-week-long vacations
- b. seven / week-long vacations



- | | |
|-------------------------------|---|
| 2. a. three-hour-long tests | ☝ |
| b. three / hour-long tests | ☝ |
| 3. a. thirty-nine-cent stamps | ☝ |
| b. thirty / nine-cent stamps | ☝ |
| 4. a. I don't know George. | ☝ |
| b. I don't know, / George. | ☝ |
| 5. a. Who's hiring Julia? | ☝ |
| b. Who's hiring, / Julia? | ☝ |

Check your answers with your teacher.

Listening Activity 2



Part A: Close your book and listen to this radio advertisement. Then open your book and listen again for the brief pauses or breaks. Mark the end of each break or thought group with a slanted line (/). The first two have been marked.

"Unlike other copier companies, / Mita doesn't make cameras, /
or televisions, or calculators, or DVD players, or answering machines,
or vacuum cleaners, or dishwashers, or cell phones, or laptops. The fact
is, Mita doesn't make anything but great copiers. After all, we didn't
become the fastest growing copier company for the last five years by
selling microwave ovens. Mita. All we make are great copiers."

Check your answers with your teacher.

2. ☒ a. three-hour-long tests 1
- b. three / hour-long tests 2
3. a. thirty-nine-cent stamps 1
- ☒ b. thirty / nine-cent stamps 2
4. a. I don't know George. 1
- ☒ b. I don't know, / George. 2
5. ☒ a. Who's hiring Julia? 1
- b. Who's hiring, / Julia? 2

Check your answers with your teacher.

Listening Activity 2



Part A: Close your book and listen to this radio advertisement. Then open your book and listen again for the brief pauses or breaks. Mark the end of each break or thought group with a slanted line (/). The first two have been marked.

"Unlike other copier companies, / Mita doesn't make cameras, /
or televisions, / or calculators, / or DVD players, / or answering machines, /
or vacuum cleaners, / or dishwashers, / or cell phones, / or laptops. The fact
is, / Mita doesn't make anything, / but great copiers. / After all, / we didn't
become the fastest growing copier company, / for the last five years, / by
selling microwave ovens. / Mita, / All we make, / are great copiers." /

Check your answers with your teacher.

LESSON PLAN - DAY # 2 -	Date: October 2 nd , 2012
Pronunciation Areas: Review Thought Groups; Focus Words	
Anticipated Problems for Students: - Ss might not understand intonation	Solution: - Let them know that it is when the pitch of the voice goes up or down. Can give them an example, even of a yes/no question in Spanish, so that they don't just associate intonation with questions.
Anticipated Problems for Teacher: - Ss might be shy to sing along.	Solution: - Have them read the song to each other if this is the case.
Goals: 1. To review Thought Groups. 2. To teach Focus Words Objectives: 1. Ss will separate thought groups when listening, reading and speaking. 2. Ss will identify focus words in song/speech 3. Ss will practice using focus words in a scripted dialogue	Complete List of Materials • Lesson Plan • Attendance form • Name cards • Lesson 2 – PPP • Holly Cole Trio – I Can See Clearly Now video • Class copies of I Can See Clearly Now (thought groups) • Answer Key – I Can See Clearly Now (thought groups) • 23 Unit 15 Task M • Class copies of Difficult Children - student copy (thought groups) • Difficult Children - answer key (thought groups) • Class copies of I Can See Clearly Now (focus words) • I Can See Clearly Now - answer key (focus words) • Focus Words - Dialogue - Teacher's Copy (no audio, T reads from this. Is also the answer key) • Class copies of Focus Words - Dialogue - Student Copy • Class copies of Focus Words - Dialogue - Teacher's Copy

Homework: Listen for and practice using focus words. - (S) Group: in songs and speech - (NS) Group: in audios and speech	Feedback:

AT START OF CLASS:

- Before class: Put up the first slide of the PPP which shows the plan for the day and

- For (s) group, play Holly Cole – I Can See Clearly Now
--

- For (ns) group don't play anything

- Take attendance; set up name cards (should know names after today's class)

Warm-up: *Review thought groups*

Suggested Time: 10 minutes

- Remind ss that we looked at thought groups yesterday. Ask them what they are and why we use them when we speak. (*groups of words; easier for listeners to understand what we say*).
- Tell them that we use thought groups in different circumstances, like when we say our phone numbers. For example: 555 / 26 76 or 56 / 2 / 417 / 85 85
- Ask ss to take turns giving their phone numbers to the people on either side of them. Then, read back the phone numbers you were given, using thought groups.

Activity 1: *Thought Group recognition and controlled practice*

Suggested Time: 15 minutes

(S) group: distribute class copies of <i>I Can See Clearly Now</i>
--

- | |
|---|
| <ul style="list-style-type: none"> - Ask ss to listen to the song and mark the thought groups, the first few have been done already - Read the lyrics and have ss check their thought group boundaries. While doing so, for “<i>it's going to be</i>” say “<i>it's gonna be</i>” as an inductive way of introducing connected speech. - Have ss compare their answers and then play the song and have ss sing along. |
|---|

(Ns) group: distribute student copies of *Difficult Children*

- Play the audio and ask the ss to separate it into thought groups.
- have ss compare their answers and then read the dialogue to each other.

Activity 2

Suggested Time: 10 minutes

- Now tell the ss that you'd like them to talk to their partner using thought groups.
- Ask ss to – using thought groups - take a couple of minutes to tell their partner what they did last night. Have them ask each other three questions about their evening so that it is a little conversation. For example: *Last night / I went to the mall. / I looked around / and then I bought / some jeans. / After that, / I had a cup of coffee. / Finally, / I went home, / cooked dinner, / and went to bed.* Now, ask the ss to ask me 3 questions, using thought groups. Exaggerate the thought group boundaries a bit when speaking so that the ss notice them.
- Go around and listen in and talk to the ss.

Presentation: Focus Words

Suggested Time: 15 minutes

- Show the Survey PPP and have the ss take the survey.
- Then move on to What is The Focus Word PPP and stop at Which butterfly is easiest to see?

Activity 3: *Quality Repetition phrase “It’s gonna be a cold day.”*

Suggested Time: 2 minutes

- Tell ss that now we’re going to do a little bit of quality repetition.
- Direct them to hear [ɪts gɒnə biːə kɒld:ei].
- Then continue saying and repeating it while they say and repeat it as well.
- Randomly check to see if ss got it.

Activity 4: *Identifying focus words*

Suggested Time: 15 minutes

- Now tell the ss that they will be identifying the focus words.

(S) group:

- Distribute the I Can See Clearly Now – Student Copy - Focus Words
- Tell ss that you will play the song twice. Play the Bobby McFerrin McNally Smith Faculty Judi Donaghy version.
- Ss listen to the song and circle the focus words, that is, the words that are the really clear ones. (The first few have been done and the thought groups have already been separated.)
- Ask ss what they noticed about how the singer sang the song. (some words she either didn't say or didn't say clearly – that's exactly the point!)
- Have ss compare their answers together and go around and see how they did.
- Ask them if they notice any pattern as to where the focus words are located within a thought group. (a: often the last word in a thought group)

- (Ns) group: - Distribute Focus Words – Dialogue – Student Copy
- Tell the ss that they will hear the dialogue twice
- Read the dialogue using the Teacher's Copy (there is no audio for this)
- Have the ss listen and circle the focus words that they hear.
- Have ss compare their answers together and go around and see how they did.
- Ask them if they notice any pattern as to where the focus words are located within a thought group. (a: often the last word in a thought group)

- Now ask ss if there were words that they found were not pronounced clearly.

- (S) Group ss will say: "going to" is pronounced "gonna". The (NS) group should mention the contractions.

- Tell them this is exactly right and that this is the opposite of focussing, so-to-speak. That is, we have to de-emphasize some words so that other words can be focus words. Remind them of the slide on the PPP and encourage them to start using these kinds of speech forms when they talk.

Activity 5: *Practicing focus words*

Suggested Time: 10 minutes

- (S) Group: Distribute the Teacher's Copy of the Focus Words – Dialogue
- (NS) Group: They have their Student Copy of this.

- Ask ss to work in pairs practicing the dialogue out loud and switch roles.

If Time:

- (S) Group: Play another version of I Can See Clearly Now and have ss sing along.
- (NS) Group: Being sure to use focus words, have ss take turns reading Difficult Children – student copy (thought groups) using focus words. Play the audio again for them, have them underline the focus words and then practice and even act out.

Homework:

- Listen for and practice using focus words.

- (S) Group: in songs and speech

- (NS) Group: in audios and speech

Slide 1

LESSON 2 – The Focus Word

- Reviewing Thought Groups and practicing them
- Survey
- Quality Repetition phrase of the day
- The Focus Word
- Identifying Focus Words
- Practicing Focus Words

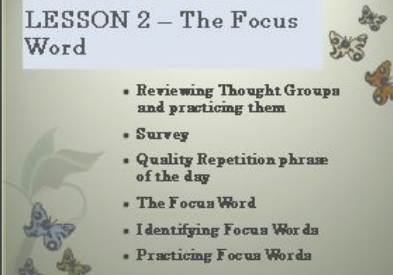


Image:

Slide 2

Survey

Read and decide whether the following statements are true or false.

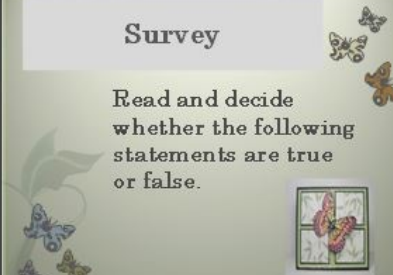
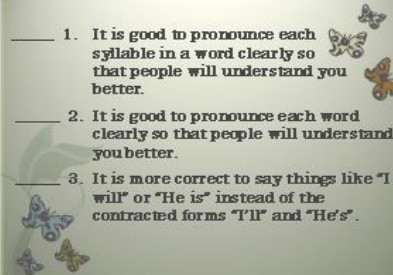


Image:

Butterfly: <https://encrypted-tbn0.google.com/images?q=tbn:ANd9GcQ16zNdWiEqnciwPvYdc96o-KqUUCOL3Mpor4w3f0Q44vFNcbgncA>

Slide 3

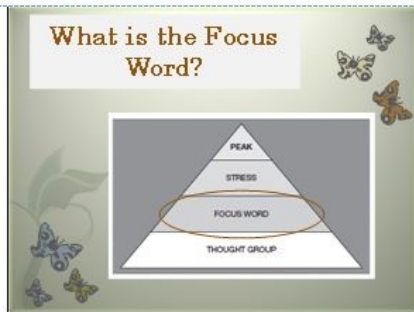
1. It is good to pronounce each syllable in a word clearly so that people will understand you better.
2. It is good to pronounce each word clearly so that people will understand you better.
3. It is more correct to say things like "I will" or "He is" instead of the contracted forms "I'll" and "He's".



Slide 4

- F** 1. It is good to pronounce each syllable in a word clearly so that people will understand you better.
- F** 2. It is good to pronounce each word clearly so that people will understand you better.
- F** 3. It is more correct to say things like "I will" or "He is" instead of the contracted forms "I'll" and "He's".

Slide 5



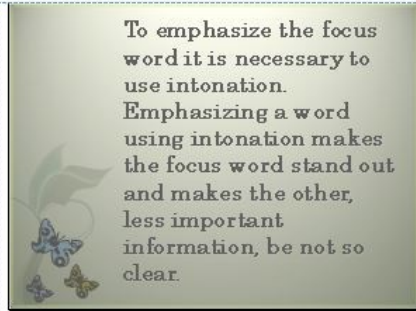
Slide 6

The Focus Word is the most important word in a Thought Group. "It is the word that the speaker wants the listener to notice most, and it is therefore emphasized."

Gilbert, Judy. © 2009 Teaching Pronunciation Using The Prosody Pyramid, P. 12. Cambridge University Press.

Taken from Judy Gilbert, pg. 12 Teaching Pronunciation Using The Prosody Pyramid (2008)

Slide 7



Taken from Judy Gilbert, pg. 12-13.
Teaching Pronunciation Using The Prosody Pyramid (2008)

Slide 8

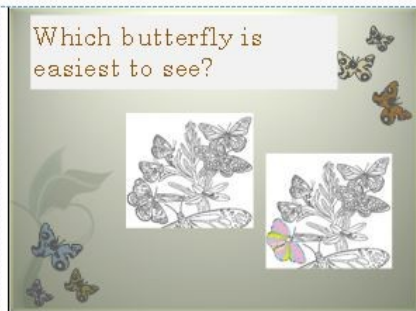
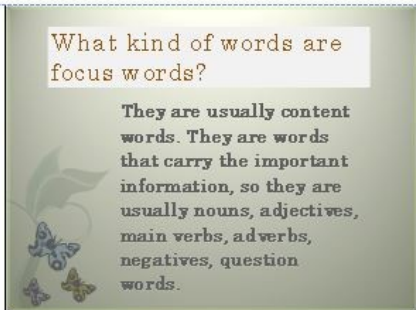


Image: <http://www.coloring-pictures.net/drawings/butterfly/many-butterflies.gif>
http://www.teachers2010.xpg.com.br/index_arquivos/image3581.gif

Slide 9



Slide 10

Homework



Listen for
and practice
using focus
words.

I CAN SEE CLEARLY NOW (HOLLY COLE)

STUDENT COPY

I can see clearly now, /

The rain is gone. /

I can see all obstacles / in my way. /

Gone are the dark clouds / that had me blind. /

...

The rest of the lyrics to the song can be found at:

http://www.lyricsmode.com/lyrics/h/holly_cole/i_can_see_clearly_now.html

I CAN SEE CLEARLY NOW (HOLLY COLE) ANSWER KEY (THOUGHT GROUPS)

I can see clearly now, /
The rain is gone. /
I can see all obstacles / in my way. /
Gone are the dark clouds / that had me blind. /
It's going to be a bright /
It's going to be bright / sunshiny day. /
I think I can make it now, /
The pain is gone. /
All of my bad feelings have / disappeared. /
There is the rainbow / I've been praying for. /
It's going to be a bright /
It's going to be bright / sunshiny day. /
Look all around /
Nothing but blue sky. /
Look straight ahead /
Nothing but blue sky. /
Look all around /
Nothing but blue sky. /
Just look straight ahead /
Nothing but blue sky. /
I can see clearly now /
The rain is gone. /
I can see all obstacles / in my way. /
Gone is the dark cloud / that had me blind. /
It's going to be a bright /
Going to be bright / sunshiny day. /
Going to be a bright, / gonna be bright, /
It's going to be bright / sunshiny day. /

Thought Groups

Listen to the following dialogue, and make a slash (/) at the end of each thought group. Then practice the dialogue with a partner. Use pauses, pitch, and syllable lengthening to make the thought groups clear.

Difficult Children

Mother: We want a turkey and cheese sandwich, / and two tuna sandwiches.

Server: On white, whole wheat, or rye?

Mother: The turkey and cheese on rye, and the other two on whole wheat.

1st child: No! No! I want white bread!

Mother: Whole wheat's good for you.

2nd child: I want peanut butter and jelly, not tuna!

Mother: OK. One turkey and cheese on rye, one tuna on white, and one peanut butter and jelly.

Server: What would you like to drink?

Mother: One iced tea, and two glasses of milk.

1st child: No milk! Lemonade!

Mother: Three sandwiches, one iced tea, and two glasses of water.

Source: Gilbert, Judy. (2005) *Clear Speech*, 3rd ed. Cambridge University Press: New York.
Page 135

Thought Groups

Listen to the following dialogue, and make a slash (/) at the end of each thought group. Then practice the dialogue with a partner. Use pauses, pitch, and syllable lengthening to make the thought groups clear.

Difficult Children

Mother: We want a turkey and cheese sandwich, / and two tuna sandwiches./

Server: On white, / whole wheat,/ or rye?/

Mother: The turkey and cheese on rye, / and the other two on whole wheat./

1st child: No!/ No! / I want white bread!/

Mother: Whole wheat's good for you./

2nd child: I want peanut butter and jelly, / not tuna!/

Mother: OK. / One turkey and cheese on rye, / one tuna on white, / and one peanut butter and jelly./

Server: What would you like to drink?/

Mother: One iced tea, / and two glasses of milk./

1st child: No milk! / Lemonade!/

Mother: Three sandwiches, / one iced tea, / and two glasses of water./

Source: Gilbert, Judy. (2005) *Clear Speech*, 3rd ed. Cambridge University Press: New York.
Page 135

I CAN SEE CLEARLY NOW (BOBBY MCFERRIN, JUDI DONAGHY) STUDENT COPY

I can see clearly now, /

The rain is gone. /

I can see all obstacles / in my way. /

Gone are the dark clouds / that had me blind. /

It's going to be a bright, / bright / sunshiny day./

It's going to be bright / bright / sunshiny day./

I think I can make it now, /

The pain is gone./

All of the bad feelings have disappeared./

Here is that rainbow / I've been praying for.

It's going to be a bright, / bright / sunshiny day./

It's going to be bright, / bright /sunshiny day./

Look all around /

Nothing but blue sky./

Look straight ahead /

Nothing but blue sky./

I can see clearly now./

The rain is gone./

Said I can see all obstacles / in my way. /

Gone are the clouds / that had me blind./

It's going to be a bright, / bright / sunshiny day./

Oh looks like a bright, / bright, / sunshiny day. /

How can you fix my ?????????????? /

Going to be a bright, / bright /sunshiny day.

I CAN SEE CLEARLY NOW (BOBBY MCFERRIN, JUDI DONAGHY) ANSWER KEY

I can see clearly now, /
The rain is gone. /
I can see all obstacles / in my way. /
Gone are the dark clouds / that had me blind. /
It's going to be a bright, / bright / sunshiny day. /
It's going to be bright / bright / sunshiny day. /
I think I can make it now, /
The pain is gone. /
All of the bad feelings have disappeared. /
Here is that rainbow / I've been praying for.
It's going to be a bright, / bright / sunshiny day. /
It's going to be bright, / bright / sunshiny day. /
Look all around /
Nothing but blue sky. /
Look straight ahead /
Nothing but blue sky. /
I can see clearly now. /
The rain is gone. /
Said I can see all obstacles / in my way. /
Gone are the clouds / that had me blind. /
It's going to be a bright, / bright / sunshiny day. /
Oh looks like a bright, / bright, / sunshiny day. /
How can you fix my ?????????????? /
Going to be a bright, / bright / sunshiny day.

FOCUS WORDS

Instructions: *Listen to the dialogue. Circle the focus words.*

John: Anna, who was on the phone?

Anna: My old friend Mary.

John: Mary Jones?

Anna: No. Mary Hall.

John: I don't know Mary Hall. Where is she from?

Anna: She's from Washington.

John: Washington the state or Washington the city?

Anna: Washington, D.C., the capital of the United States.

John: Is that where she lives?

Anna: Yes, she still lives in the white house.

John: The White house? With the president?

Anna: No, silly. The white house on First Street.

John: What did she want?

Anna: She wants to come here.

John: Come here? When?

Anna: In a week. She's bringing her black bird, her collie, her snakes, her...

John: Stop! She's bringing a zoo to our house?

Anna: No, John. She's opening a pet store here in town.

Taken from pages 93-94 of Dale & Poms (2005) *English Pronunciation Made Simple*

FOCUS WORDS

John: ANNA, who was on the PHONE?

Anna: My old friend MARY.

John: Mary JONES?

Anna: NO. Mary HALL.

John: I don't know Mary HALL. Where is she FROM?

Anna: She's from WASHINGTON.

John: Washington the STATE or Washington the CITY?

Anna: Washington, D.C., the CAPITAL of the United States.

John: Is that where she LIVES?

Anna: YES, she still lives in the white HOUSE.

John: The WHITE house? With the PRESIDENT?

Anna: No, SILLY. The white HOUSE on FIRST Street.

John: What did she WANT?

Anna: She wants to COME here.

John: Come HERE? WHEN?

Anna: In a WEEK. She's bringing her black BIRD, her COLLIE, her SNAKES, her...

John: STOP! She's bringing a ZOO to OUR house?

Anna: NO, John. She's opening a PET store here in TOWN.

Taken from pages 93-94 of Dale & Poms (2005) *English Pronunciation Made Simple*

LESSON PLAN - DAY # 3	Date: October 3 rd , 2012
Pronunciation Areas: Stress; Peak; Syllables and Word Stress; Schwa	
Anticipated Problems for Students: - Ss might have a hard time hearing the schwa	Solution: - Remind them that it is very hard to hear, that maybe if they notice that they don't hear a full vowel, they will understand that what they are hearing is the schwa.
Anticipated Problems for Teacher: - Ss getting frustrated if they don't hear or say the stressed syllable properly	Solution: - Have them practice stressing different syllables in order to hear the difference and work out which one sounds right.
Goals: 1. To finish teaching the last two general aspects of the Prosody Pyramid: Stress & Peak 2. To teach Syllables and Word Stress 3. To introduce the Schwa Objectives: 1. Ss will use focus words in sentences and will disagree with or correct what was said. 2. Ss will identify the peak syllables in focus words and then practice them orally. 3. Ss identify and orally practice the schwa in connected speech. 4. Ss will imitate connected speech without a text.	Complete List of Materials • Lesson Plan • Lesson 3 – PPP • Class copies of <i>Pair work: Listening for the focus word</i> • 'Wavin' Flag' by K'Naan (video) • Class copies of K'Naan - Wavin' Flag - student copy (peak) • If by Rudyard Kipling (audio clip) • class copies of Rudyard Kipling - If - student copy (peak) • Pronunciación el sonido schwa (video) • Basic Rhythm – Reduced Words (audios, tracks 45, 46, 47) • Class copies of Basic Rhythm – Reduced Words • Imitation exercise: Explanation and Peggy's New Office (video)
Homework: • Imitation exercise: Explanation and Peggy's New Office (video)	Feedback:

AT START OF CLASS:

- Before class: Put up the first slide of the PPP which shows the plan for the day.

Warm-up and Review: *Review Thought Groups and Focus Words*

Suggested Time: 10 minutes

- Show Review slide for *Thought Groups and Focus Words*.
- Go through it and then model the example and point out the areas indicated.
- Distribute *Pair work: Listening for the focus word* handout. Model the example with a student and then pair them up to practice. If there are weak ss, pair them with stronger ones. Circulate, listen and help.

Presentation: *Stress and Peak Syllables*

Suggested Time: 15 minutes

- Show the *What is Peak and Stress* slide.
- Ask ss to try to answer this question.
- Then go through the rest of the slides on this. Emphasize the stressed syllables and point out that they are clear because we (almost) don't say the vowels in the crossed out syllables.
- Stop at this point because the next activity will (hopefully) help the ss to hear not only the peak syllables, but also the reduced ones. (Syllable stress and the schwa will be introduced after this.)

Activity 1: *Recognizing the peak and reduced syllables*

Suggested Time: 20 minutes

(S) Group:

- Distribute class copies of K'Naan - Wavin' Flag - student copy (peak).
- Play K'naan's Wavin' Flag and ask the ss to follow the instructions at the top of the page.
- Have ss compare their answers. Circulate and check while they do this.
- Then go back to the PPP to the slide with those focus words – slide #9. Ask the ss to listen to the sound of the vowel in the non-peak syllables. What do they notice? (They should say that they all sound the same.) Point out that they all have different vowel letters, but one sound that is very hard to hear: the schwa.

(NS) Group:

- Distribute class copies of Rudyard Kipling - If - student copy (peak).
- Play If – Rudyard Kipling audio clip a couple of times and ask the ss to **circle the peak syllable in the focus words that have been underlined.**
- Have ss compare their answers. Circulate and check while they do this.
- Then go back to the PPP to the slide with those focus words – slide #10. Ask the ss to listen to the sound of the vowel in the non-peak syllables. What do they notice? (They should say that they all sound the same.) Point out that they all have different vowel letters, but one sound that is very hard to hear: the schwa.

Presentation: *The Schwa*

Suggested Time: 15 minutes (tops)

- Play the video clip, Pronunciación el sonido schwa.
- The show the PPP slides on schwa and model and have ss practice saying it.

Activity 2

Suggested Time: 15 minutes

- Distribute class copies of Basic Rhythm – Reduced Words
- Go through the exercises with the class and play the associated audios
- In pairs have ss practice the alternatives on page 90.

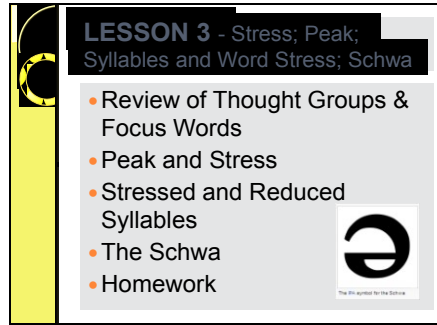
If time left over:

- Ask ss to practice reading their handouts out loud, being sure to use the schwa

Homework:

- Imitation exercise: Explanation and Peggy's New Office
- Give URL (and is on PP) in case there is time in class to begin or practice.

Slide 1



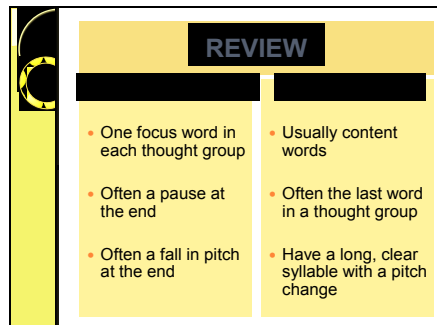
LESSON 3 - Stress; Peak; Syllables and Word Stress; Schwa

- Review of Thought Groups & Focus Words
- Peak and Stress
- Stressed and Reduced Syllables
- The Schwa
- Homework



http://www.translationdirectory.com/images_articles/wikipedia/Schwa.jpg

Slide 2

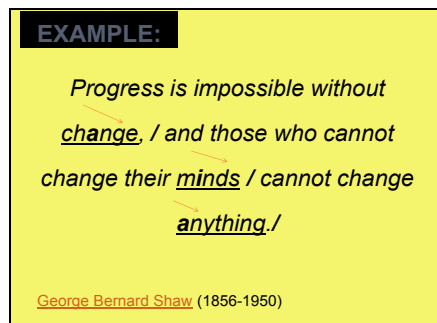


REVIEW

- One focus word in each thought group
- Often a pause at the end
- Often a fall in pitch at the end
- Usually content words
- Often the last word in a thought group
- Have a long, clear syllable with a pitch change

Content taken and adapted from: Gilbert, Judy. (2008) Teaching Pronunciation Using The Prosody Pyramid. Cambridge University Press. Page 45.

Slide 3

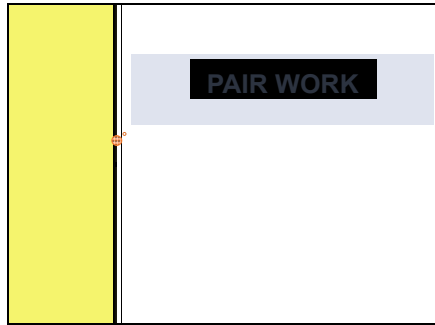


EXAMPLE:

*Progress is impossible without
change. / and those who cannot
change their minds / cannot change
anything. /*

George Bernard Shaw (1856-1950)

Slide 4



Slide 5

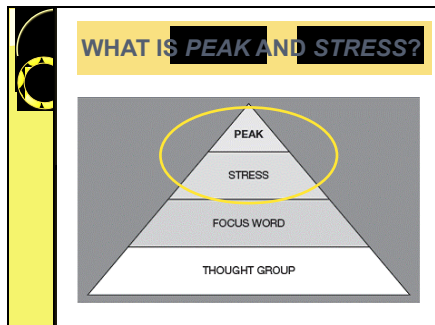


Chart taken and adapted from: Gilbert, Judy. (2008) Teaching Pronunciation Using The Prosody Pyramid. Cambridge University Press. Page 10.

Slide 6

- The most important syllable in a focus word:
 - Receives the most stress
 - Indicates the peak of information in a thought group

Because this syllable is the most important one in the most important word of a thought group, it must be heard the most clearly.

Content taken and adapted from: Gilbert, Judy. (2008) Teaching Pronunciation Using The Prosody Pyramid. Cambridge University Press. Page 14.

Slide 7



<http://chrismaddencartoons.files.wordpress.com/2011/09/athletics-olympics-winners-podium-cartoon-cjmadden.jpg?w=630>
<http://www.flagpics.net/i/chile-flag-924.gif>

Slide 8



Circle the stressed syllables in the following focus (*italicized*) words.

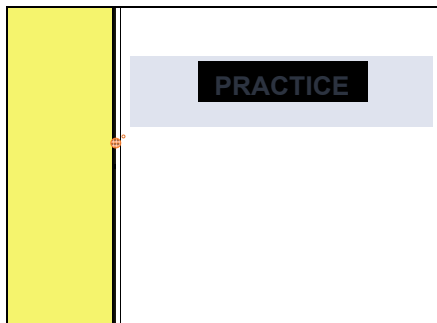
"When we lose the right to be *different*, we lose the *privilege* to be *free*."
Charles Evans Hughes

"Next in importance to freedom and *justice* is popular *education*, without which neither freedom *nor* justice can be *permanently* maintained."
James A. Garfield

Sources:
<http://www.quotationspage.com/subjects/freedom/>
<http://www.quotationspage.com/subjects/education/>

(accessed Sep. 11, 2012)

Slide 9



Slide 10

STRESSED AND REDUCED SYLLABLES	
DIFF erent PRIV ilege FREE	JUS tice edu CAT ion NOR PER manently

Slide 11

STRESSED AND REDUCED SYLLABLES	
al LOWance dis AS ter <i>pitch-and-TOSS</i>	<i>Min</i> ute <i>SE</i> conds

Slide 12

THE SCHWA


Image: http://danienglish.com.br/wp-content/uploads/2012/03/homer-simpson_schwa.jpg (accessed Sep. 12, 2012)

Slide 13



- The schwa can appear in one-syllable words when the vowel is followed by an /ɹ/:
 - *sir, fur, her, heard,*
- It can appear in two-syllable words in the first or second syllable:
 - *advice, escape, disease, tonight, support*
 - *palace, college, tulip, purpose, minute*

Slide 14



- It can appear more than once and in any syllable in three or more syllable words:
 - *completion, calendar, instrument*
- It can be the sound in a function word
 - *a, the, can, you, and*

Slide 15



HOMEWORK

Imitation exercise:

http://www.rachelsenglish.com/im_peggys_new_office

(URL accessed Sep. 12, 2012)

G *Pair work: Listening for the focus word*

Student A: Say sentence **a** or **b**.

Student B: Listen closely for the focus word, and say the matching response.

Example

Student A: "It's a **big** dog."

Student B: "No, it's really more medium-sized."

OR

Student A: "It's a big **dog**."

Student B: "No, it's a wolf."

One way to make this exercise more fun is to hum the sentence. Or you could use a kazoo (toy humming instrument). When Student A hums the sentence, Student B listens closely to the pitch pattern and then says the response.

- | | |
|--|--|
| 1. a. It's a big dog . | No, it's a <u>wolf</u> . |
| b. It's a big dog. | No, it's really more <u>medium-sized</u> . |
| 2. a. But we asked for two <u>coffees</u> ! | Oh, I thought you wanted <u>tea</u> . |
| b. But we asked for <u>two</u> coffees! | Oh, I thought you wanted <u>one</u> . |
| 3. a. I thought you bought a big <u>car</u> . | No, it was a <u>motorcycle</u> . |
| b. I thought you bought a big car. | No, it was a <u>little</u> one. |
| 4. a. Is that a silver <u>watch</u> ? | No, it's a <u>bracelet</u> . |
| b. Is that a <u>silver</u> watch? | No, it's <u>platinum</u> . |
| 5. a. I prefer beef <u>soup</u> . | Not <u>stew</u> ? |
| b. I prefer <u>beef</u> soup. | Not <u>chicken</u> ? |
| 6. a. Is there milk in the <u>refrigerator</u> ? | No, it's on the <u>table</u> . |
| b. Is there <u>milk</u> in the refrigerator? | No, but there's <u>juice</u> . |

Instructions: *While you listen to the song, circle the peak syllable in the focus words that have been underlined for you. The first one has been done. Then, in the second column, begin singing along to the song. Raise your hands every time you sing a focus word.*

WAVIN' FLAG (K'NAAN)

When I get <u>older</u> I will be <u>stronger</u> They'll call me <u>freedom</u>, just like a <u>wavin'</u> flag Born to a throne, stronger than Rome A violent prone, poor people zone But it's my home, all I have known Where I got grown, streets we would roam Out of the <u>darkness</u>, I came the <u>farthest</u> Among the <u>hardest</u> survival Learn from these streets, it can be bleak Accept no <u>defeat</u>, surrender, <u>retreat</u> So we <u>struggling</u>, fighting to eat And we <u>wondering</u> when we'll be free So we patiently wait for that fateful day It's not far <u>away</u>, but for now we say When I get <u>older</u> I will be <u>stronger</u> They'll call me <u>freedom</u> just like a <u>wavin'</u> flag And then it goes back, and then it goes back And then it goes back, and then it goes Repeat So many wars, settling scores Bringing us <u>promises</u>, leaving us poor I heard them say 'love is the way' 'Love is the <u>answer</u>,' that's what they say But look how they treat us, make us <u>believers</u> We fight their <u>battles</u>, then they <u>deceive</u> us Try to <u>control</u> us, they couldn't hold us 'Cause we just move <u>forward</u> like <u>Buffalo</u> Soldiers	But we <u>struggling</u>, fighting to eat And we <u>wondering</u>, when we'll be free So we patiently wait for that faithful day It's not far <u>away</u> but for now we say When I get <u>older</u> I will be <u>stronger</u> They'll call me <u>freedom</u> just like a <u>wavin'</u> flag And then it goes back, and then it goes back And then it goes back, and then it goes ... The rest of the lyrics can be found at: http://www.azlyrics.com/lyrics/knaan/wavinflag.html
---	--

IF (Rudyard Kipling)

Student copy

If you can keep your *head* when all about you
Are losing *theirs* and blaming it on *you*,
If you can trust *yourself* when all men *doubt* you,
But make *allowance* for their *doubting too*;
If you can *wait* and not be tired by *waiting*,
Or being *lied* about, don't deal in *lies*,
Or being *hated*, don't give way to *hating*,
And yet don't look *too* good, nor talk too *wise*:

If you can *dream* - and not make dreams your *master*;
If you can *think* - and not make thoughts your *aim*;
If you can meet with Triumph and *Disaster*
And treat those two impostors just the *same*;
If you can *bear* to hear the truth you've *spoken*
Twisted by *knaves* to make a trap for *fools*,
Or watch the things you gave your *life* to, *broken*,
And *stoop* and build 'em *up* with worn-out *tools*:

If you can make one *heap* of all your *winnings*
And risk it on one *turn* of *pitch-and-toss*,
And *lose*, and start *again* at your *beginnings*
And never *breathe* a word about your *loss*;
If you can force your *heart* and *nerve* and *sinew*
To serve your *turn* long after they are *gone*,
And so hold *on* when there is nothing *in* you
Except the *Will* which says to them: '*Hold on!*'

If you can talk with *crowds* and keep your *virtue*,
' Or walk with *Kings* - nor lose the common *touch*,
if neither *foes* nor loving friends can *hurt* you,
If all men *count* with you, but none too *much*;
If you can *fill* the unforgiving *minute*
With sixty *seconds'* worth of distance *run*,
Yours is the *Earth* and everything that's *in* it,
And - which is *more* - you'll be a *Man*, my *son*!

Poem: http://www.kipling.org.uk/poems_if.htm

Audio: http://ia600300.us.archive.org/23/items/if_kipling_librivox/if_kipling_sw.mp3

CHAPTER Basic Rhythm—Reduced Words 10

Function words are sometimes hard to hear. They are *reduced* so that the important content words stand out.

You've SHUT the DOOR on my HAND!

In this chapter you will learn . . .

- What reduced words sound like.
- How reduced words are weakened.

Get Set!



Listen to one-half of a telephone conversation. Circle the words in parentheses that you hear. Check your answers. Or work with a partner as follows.



Student A: You are on the phone. You are planning a party for your friend, Sam. Turn to page 98. Read your half of the conversation.

Student B: You are listening to Student A talk on the phone. You can hear only his half of the conversation. Circle the words in parentheses that you hear.

The party'll be on Sunday in the Terrace Apartments.

That's right . . . (427, 4 to 7).

I (can, can't) take you.

I (can, can't) pick you up before noon.

We're having salad (in, and) sandwiches.

I'll ask (for, four) volunteers to help clean up.

Victor (can, can't) come, but he (can, can't) get the gift.

Don't worry. Sam'll love it. Snakes are quiet. And you only have to feed them once a month!

Listen!



Listening Activity 1 The teacher or the speaker on the audio will say sentence a. or b. Circle the one you hear.

Examples: a. I'll ASK for volunTEERS to help.
 (b.) I'll ASK FOUR volunTEERS to help.

1. a. It's 4 to 7 (FOUR to SEVEN – *time*).
 b. It's 4 2 7 (FOUR TWO SEVEN – *street number*).
2. a. We can TALK.
 b. We CAN'T TALK.
3. a. It's for EYES.
 b. It's FOUR EYES.
4. a. This is to CLEAN.
 b. This is TOO CLEAN!
5. a. What's H to O?
 b. What's H₂O? (water)

Check your answers.



Now listen to both a. and b. Can you hear a difference in rhythm?



Listening Hint

Do native speakers speak too fast? Many students think so. The problem isn't just speed, however. Native speakers also reduce function words. Function words get lost in conversations. Students get frustrated.

As speech gets more casual, function words get harder to hear.

<i>Formal</i>	We should <u>have</u> called him.
↓	We should <u>əv</u> called <u>əm</u> .
<i>Casual</i>	We should <u>ə</u> called <u>əm</u> .

It is not poor English to use reduced words. They are a part of English rhythm. When you listen to English, don't try to hear every word clearly. Pay more attention to the stressed words. You will be a more effective listener!

LESSON PLAN - DAY # 4	Date: October, 4th
Pronunciation Areas: Connected Speech and Linking 1	
Anticipated Problems for Students: - Ss will think some of this is “bad English” and they won’t feel comfortable making the assimilations and contractions etc.	Solution: - Tell them that it is perfectly correct spoken English and that they need to use them if they want to sound natural when they speak.
Anticipated Problems for Teacher: - Ss sometimes get angry about the language not being phonetic.	Solution: - Try to reassure them that once they know the different connected speech phenomena, they will find English so much easier to understand.
Goals: 1. To finish covering any missed areas when reviewing schwa; syllable and word stress. 2. To teach three connected speech phenomena Objectives: 1. Ss will orally use and physically show word stress 2. Ss will perceive, write, and say the connected speech phenomena taught.	Complete List of Materials • Lesson Plan • Lesson 4 – PPP • Audio of Tears in Heaven – Eric Clapton • Class copies of Tears in Heaven – Eric Clapton (student copy) • Class copies of Would you rather...? (student copy) • Class copies of Student Worksheet 1 Linking consonant to vowel
Homework: - Ss study and complete Student Worksheet 1: Linking consonant to vowel at home. We will practice this at the beginning of next class.- Imitation exercise: People Change	Feedback:

AT START OF CLASS:

Suggested Time: 5 minutes

- Before class: Put up the first slide of the PPP which shows the plan for the day.
- Talk to the ss to see how they liked the imitation exercise. Can they remember any phrases that they said/learned? Ask them to share them with the class. Answer any questions they might have.

Warm-up: Tongue twister

Suggested Time: 10 minutes

- Tell the ss that we're going to begin today with a tongue twister. (Emphasize the /t/ in the words *tongue* and *twister*. Briefly mention how the English /t/ is made and different from the Spanish /t/. This is not, the point of the activity but the tongue twister uses the schwa in the function word *to*, so there no harm in briefly teaching the /t/ sound.
- Just say the tongue twister first, slowly and then a little faster showing the rhythm with some hand movements or tapping on the desk.
- Tongue twister: "A tutor who tooted the **flute**, tried to tutor two tooters to **toot**. Said the two to their tutor, "is it harder to **toot**, or to tutor two tooters to **toot**?"
- Ask students to try repeating this exactly the way I say it.
- Point out the picture depicting the situation on the first slide of the PPP, then show the tongue twister slide and let the ss read it so that they understand it.
- Ask the ss to identify the words/syllables with a schwa sound.
- Show the answers on the next slide.

Review + Presentation: Schwa; Syllables and Word Stress

Suggested Time: 15 minutes

- Show the slides on Schwa and go over.
- Continue with the slides that deal with word stress.

Activity 1: Stress Moves (a game)

Suggested Time: 15 minutes

(to review see page 15 of Pronunciation Games)

- The *stress move* is to open and raise a hand when saying a stressed syllable

- Model the stress move with the words *concentrate, photograph, telephone*
- Ask ss to give and model some examples of their own to make sure they know how to do it.
- Put ss in a circle and give them each two cards and time to practice the correct pronunciation to themselves.
- Tell ss that the first person shows one of their cards to the others and pronounces it with the stress move. Then the next person does that with one of their own cards as well as the word the previous person said and so on.
- Start fresh when the list of words gets too long. Continue until all words are done.
- Then, if time, have each student make up and say a meaningful (even silly) sentence using both of their words and anyone else's words if they wish.

Activity 2

Suggested Time: 20 minutes

- (S) Group: Tears in Heaven (Eric Clapton)
- Tell ss that now we're going to listen to a song by Eric Clapton. Ask if anyone knows his music.
- Play the song once and just have ss listen to see how much they understand. Tell them not to worry if they don't understand much.
- Distribute class copies of Tears in Heaven – Eric Clapton (student copy)
- Tell ss that there are a lot of words missing in the lyrics and that their job is to try to fill in the blanks with the missing words.
- Remind them that English doesn't always sound the way it is spelled.
- Tell them that *some* of the blanks will have reduced function words.
- Play the song 2-3 times to allow them to fill in as many spaces as possible.
- Take up the answers by showing them on the PPP
- Ask ss what they noticed about some of the blanks (they should say: "He say's woodja, 'cause, I fie etc.")
- Tell them that this is exactly right and move on to the PPP slides on Connected Speech 1

- *Note: Only d + y = /dʒ/, C + V and contractions are covered in the PPP. Others will be covered in the next couple of classes.

- Go back to the song. Students will want to hear it again. Encourage them to sing along using the pronunciation covered– if possible or necessary, play it loud enough so that they can sing without feeling embarrassed.

- (NS) Group: Tell the ss that now they are going to listen to some hypothetical situations in which they have two choices.

- Read the Would you rather...? scenarios and ask the ss to just listen. Tell them not to worry if they don't understand much.

- Distribute the class copies of Would you rather...? (Student Copy)

- Tell ss that there are lots of words missing and that their job is to try to fill in the blanks with the missing words.

- Remind them that English doesn't always sound the way it is spelled.

- Tell them that *some* of the blanks will have reduced function words.

- Read the handout (or play the audio***if I have time to record someone) 2-3 times to allow them to fill in as many spaces as possible.

- Take up the answers by showing them on the PPP (***be sure to quickly jump over the Eric Clapton answer key PPP slides)

- Ask ss what they noticed about some of the blanks (they should say: *woodja, lotsa* etc.)

- Tell them that this is exactly right and move on to the PPP slides on Connected Speech 1

- *Note: Only d + y = /dʒ/, C + V and contractions are covered in the PPP. Others will be covered in the next couple of classes.

- Go back to the handout. Ask ss to pair up and read through these questions using the pronunciation covered.

Activity 3: *Fluency practice*

Suggested Time: 10 minutes

- (S) Group: Have ss get together in groups of four and ask each other questions and discuss the situation in the song. What would you do in Heaven? Who would you talk to? What would you say? Etc. Members of the group should try to be creative with their answers.

- (NS) Group: Have ss get together in groups of four and ask each other some of the questions in the handout as well as some new ones. Members of the group should answer and explain their answers.

Homework:

- Ss study and complete Student Worksheet 1: Linking consonant to vowel at home. We will practice this at the beginning of next class.
- Imitation exercise: People Change

Slide 1

**LESSON 4 - CONNECTED SPEECH
AND LINKING 1**



- A tongue twister
- Review of schwa, syllables and word stress
- Connected Speech and Linking 1
- Homework

<http://www.gutenberg.org/files/24560/24560-h/images/img5.jpg> (Accessed Sep. 12th, 2012)

Slide 2

A tutor who tooted the **flute**, tried to tutor two tooters to **toot**.

Said the two to their tutor, "is it harder to **toot**, or to tutor two tooters to **toot**?"

Slide 3

A tutor who tooted the flute, tried to tutor two tooters to toot.

Said the two to their tutor, "is it harder to toot, or to tutor two tooters to toot?"

Slide 4

SCHWA

- The schwa is the most common sound in English. It is a very short and unclear sound. That is, it is very hard to hear.
- All of the vowel sounds in English can be reduced to a schwa. That is, it can correspond to any vowel letter.
- It is usually used in function words.
- The schwa is very important in the English stress system.

Slide 5

FUNCTION WORDS

Usually unstressed, unless in final position or when used emphatically

articles	(a, an, the)
auxiliary verbs	(can, will, do, is)
personal pronouns	(I, you)
possessive adjectives	(my, your)
demonstrative adjectives	(this, that)
prepositions	(in, on, at)
conjunctions	(and, or, but)

Taken and adapted from Celce-Murcia et al (2010) Teaching Pronunciation. Page 210.

Slide 6

- Function (or *structure*) words can be pronounced in their strong form or their weak form.
- The strong form is used when they are pronounced in isolation or if they need to be the focus word as a way of emphasizing, clarifying or correcting information.
- Normally, however, they are pronounced in their weak form, which uses the schwa. The following slide shows the strong and weak pronunciation of some common function words.

Slide 7

FUNCTION WORD	STRONG FORM	WEAK FORM
“can”	/kæn/	/kən/
“will”	/wɪl/	/wəl/ , /əl/
“have”	/hæv/	/əv/ , /v/
“to”	/tʰu:/	/tʰə/
“them”	/ðem/	/əm/ , /əm/
“and”	/ænd/	/ən/ , /n/

Slide 8

WHY IS IT IMPORTANT TO HAVE SOME WORDS THAT ARE FOCUS WORDS AND OTHERS THAT ARE REDUCED?

- The important information is easier to hear (remember contrast is important)
- This is the basic rhythm of English

Slide 9

WORD STRESS

- There are three types of syllables in English:
 1. Stressed (primary stress)
 2. Unstressed (secondary stress)
 3. Reduced (unstressed)
- *The terms in brackets are ones that some texts will use.

Slide 10

- The vowel in the stressed syllable of a word is extra long and extra clear.
concentrate photograph telephone
(Only the vowel in a stressed syllable can be the peak of a focus word!)
- The vowel in the unstressed syllable is short and clear.
concentrate photograph telephone
- The vowel in the reduced syllable is very short and unclear. That is, its sound is the schwa.
concentrate photograph telephone

Taken and adapted from: Gilbert, Judy.
(2005) Clear Speech. Pages 27-28.

Slide 11

TEARS IN HEAVEN (ERIC CLAPTON) Answer Key

Would you know my name if I saw you in Heaven
Would it be the same if I saw you in Heaven
 I must be strong and carry on
Because I know I don't belong here in Heaven
Would you hold my hand if I saw you in Heaven
Would you help me stand if I saw you in Heaven
I'll find my way through night and day
Because I know I just can't stay here in Heaven

Slide 12

Time can bring you down, time can bend your knees
 Time can break your heart have you begging, please
 Begging, please
 Beyond the door, there's peace I'm sure
 And I know there'll be no more tears in Heaven
Would you know my name if I saw you in Heaven
Would it be the same if I saw you in Heaven
 I must be strong and carry on
Because I know I don't belong here in Heaven
Because I know I don't belong here in Heaven

Slide 13

4. What is better, being rich and miserable or poor and happy?
5. Would you rather have a nice teacher who is bad at teaching or a mean one who is good at teaching?
6. Would you rather date a rich workaholic or someone poor who has lots of time for you?
7. Would you rather know if there is a God or if there is a devil?
8. What is better, having only one friend who is honest, or lots of friends who lie to you?

Slide 14

- WOULD YOU RATHER...?** Answer Key
1. If you or your partner could not have kids, would you rather adopt or live without children?
 2. Would you rather have to sew all of your clothes or grow all of your own food?
 3. If your partner was unfaithful to you, which would you prefer: To be sad because you found out or to be happy because you didn't?

Slide 15

CONNECTED SPEECH AND LINKING 1

would you
w u (dʒ) j u:

Jew?

wadizzit? w ɔ t i z i t e b a u t
What is it about?

Images:

<http://www.tipsforenglish.com/wp-content/uploads/2011/08/goodTime1-300x218.jpg>

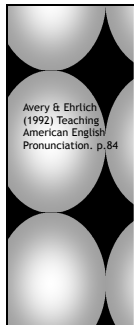
<http://www.tipsforenglish.com/wp-content/uploads/2011/08/oneNight-300x218.jpg>

<http://www.tipsforenglish.com/wp-content/uploads/2011/08/cantYou-300x218.jpg>

<http://www.tipsforenglish.com/wp-content/uploads/2011/08/youAre1-300x218.jpg>

<http://www.tipsforenglish.com/wp-content/uploads/2011/08/whatIs1-300x218.jpg>

Slide 16



Avery & Ehrlich
(1992) Teaching
American English
Pronunciation, p.84

"In connected speech, words within the same phrase or sentence often blend together. Connecting groups of words together is referred to as linking. When words are properly linked, there is a smooth transition from one word to the next."

Slide 17

Linking a consonant (sound) to a vowel (sound)	Examples:
	<i>stop_it</i>
	<i>come_in</i>
	<i>pass_out</i>
	<i>stayed_at</i>

When a word that ends in a consonant comes before a word that begins with a vowel, it sounds like the second word begins with the consonant.

Slide 18

/d/ + /j/ = /dʒ/ Examples:

When a word that ends in a /d/ comes before a word that begins with a /j/, the two sounds merge and become a different sound.

would you
did you
could you
should you

Slide 19

Contractions Examples:

Auxiliary verbs are often contracted in speech. Remember, they are function words and therefore should be unstressed or reduced.

I'm
He'll
We'd
There's

Slide 20

HOMEWORK

- Study and complete Student Worksheet 1: Linking consonant to vowel at home. We will practice this at the beginning of next class.
- Practice the following imitation exercise: People Change
<http://www.youtube.com/watch?v=TeBNvIrUvXk&feature=relmfu>

_____ know my name _____ saw you in Heaven

_____ be the same _____ saw you in Heaven

I must be strong _____ carry on

_____ I know I _____ belong _____ Heaven

_____ hold my _____ saw you in Heaven

_____ help me _____ saw you in Heaven

_____ find my way through _____ day

_____ I know I just _____ stay _____ Heaven

Time _____ bring _____ down, time _____ bend _____ knees

Time _____ break your heart have _____ begging, please Begging,
please

Beyond _____ door, _____ peace _____ sure

And I know _____ be no more tears in Heaven

_____ know my name _____ saw you in Heaven

_____ be the same _____ saw you in Heaven

I must be strong _____ carry on

_____ I know I _____ belong _____ Heaven

_____ I know I _____ belong _____ Heaven

TEARS IN HEAVEN (ERIC CLAPTON)

Answer Key

Would you know my name if I saw you in Heaven

Would it be the same if I saw you in Heaven

I must be strong and carry on

Because I know I don't belong here in Heaven

Would you hold my hand if I saw you in Heaven

Would you help me stand if I saw you in Heaven

I'll find my way through night and day

Because I know I just can't stay here in Heaven

Time can bring you down, time can bend your knees

Time can break your heart have you begging, please

Begging, please

Beyond the door, there's peace I'm sure

And I know there'll be no more tears in Heaven

Would you know my name if I saw you in Heaven

Would it be the same if I saw you in Heaven

I must be strong and carry on

Because I know I don't belong here in Heaven

Because I know I don't belong here in Heaven

WOULD YOU RATHER...?

Student copy

1. _____ or your partner _____ have kids, _____
rather adopt or live without children?
2. _____ rather _____ sew _____ your clothes or
grow _____ your own food?
3. _____ partner was unfaithful _____, which _____
_____ prefer: To be sad because you _____ or to be happy because you
_____?
4. _____ better, being rich _____ miserable or poor _____ happy?
5. _____ rather have a nice teacher _____ bad at teaching or a
mean one _____ good at teaching?
6. _____ rather _____ rich workaholic or someone poor who has
_____ time for you?
7. _____ rather know if _____ a God or if _____ a
devil?
8. _____ better, having only one friend _____ honest, or _____
_____ friends who lie to you?

Content taken and adapted from: <http://www.rrrather.com/>

DESCRIPTION AND ANALYSIS

When introducing connected speech in the ESL or EFL classroom, you can illustrate the most common patterns using worksheets like those presented in Figures 5.6–5.9.

STUDENT WORKSHEET I	
Linking consonant to vowel	
<i>Rule 1:</i> When a word ends in two consonants and the next begins with a vowel, the final consonant sounds like the initial consonant of the following word. <i>Send it</i> sounds like <i>sen-dit</i> <i>Camp out</i> sounds like <i>cam-pout</i>	
<i>Rule 2:</i> When a word ends in a single consonant and the next begins with a vowel, the consonant straddles the two syllables.	
push_up	stop_it
come_in	take_off
Practice: Repeat the following phrases, paying attention to linking.	
Two consonants + vowel	Single consonant + vowel
hol/d_on	is_it?
lef/t_it	keep_up
fin/d_out	gone_in
Your examples:	
Think of verbs ending in consonants that would complete the following phrases. Write one verb in each blank. Then practice saying the phrases with your partner.	
_____ it in.	_____ at me.
_____ it down.	_____ out.
_____ up.	_____ on it.

Figure 5.6 Student worksheet for linking consonant to vowel

LESSON PLAN - DAY # 5	Date: Oct. 8 th , 2012
Pronunciation Areas: Connected Speech and Linking 2	
Anticipated Problems for Students: - Ss might find the last activity difficult.	Solution: - Remind them that they are new to this and that it will get easier. Tell them that doing today's homework will give them plenty of practice with this.
Anticipated Problems for Teacher:	Solution:
Goals: 1. To review and recycle forms covered in Connected Speech and Linking 1 2. To teach additional forms of connected speech Objectives: 1. Ss will mimic connected speech with the use of lyrics to a rap 2. Ss will orally practice connected speech phenomena covered in the last class 3. Ss will write phrases containing a variety of forms of connected speech	Complete List of Materials • Lesson Plan – Lesson 5 • Lesson 5 – PPP • Whatcha Gonna Do - Burton Crane (video) • Whatcha Gonna Do - Burton Crane (lyrics) • Gonna Wanna Hafta 45 Track 45 (audio clip) • Gonna Wanna Hafta Jazz Chant (Student Copy) • Neil Young OLD MAN (video) • Neil Young - Old Man (student copy) • Neil Young - Old Man (Answer Key) • Steve Jobs explains the rules for success – YouTube • Steve Jobs - The Rules for Success (Student Copy) • Steve Jobs - The Rules for Success (Answer Key)
Homework: - Go to YouTube and look for videos similar to the last one done in class. Listen a number of times, try to understand and repeat/mimic what you can.	Feedback:

AT START OF CLASS:

Warm-up: *Whatcha Gonna Do?* Rap from American Idol.

Suggested Time: 10-15 minutes

- Tell ss that we're going to do a rap today and that it was one that a retired teacher performed on American Idol.

- (S) Group: Play video and have ss just listen first. Discuss briefly, make sure they know what *whatcha gonna do* means. Distribute lyrics, discuss the part in italics.

- Play the video again and have ss sing along to what they can while trying to mimic the singer.

- (NS) Group: Read the song and have the ss just listen first. Discuss briefly, make sure they know what *whatcha gonna do* means. Distribute lyrics, discuss the part in italics.

- Read and ask ss read along out loud.

Review : *Using connected speech*

Suggested Time: 10 minutes

- Ask students to take out their homework: Worksheet 1 Linking consonant to vowel.

- Take up their answers and give ss time to practice together in pairs.

Presentation: Connected Speech and Linking 2

Suggested Time: 15 minutes

- Show the slide Connected Speech and Linking 2 and talk about the two areas.

- The first area is like what we saw last class and the ss should have gotten this after the rap.

- For the third area, see below and mention teacher talk phenomena and textbook audios, which are like teacher talk.

- Be sure to discuss that these changes are important for having a natural authentic-sounding accent. At the same time, discuss the fact that there are different levels of formality and that for the auxiliary verbs + to, these changes might not occur in formal speaking situations. For this area, also mention that these forms are not usually written, except maybe in songs, text messages, or when the writer wants to show the pronunciation (for whatever reason).

Activity 1: Jazz Chant

Suggested Time: 5 minutes

- Distribute class copies of the jazz chant from page 48.
- Go through the instructions with the ss, play the audio, and then give them time to practice the chant in pairs.

Activity 2: Hearing and practicing connected speech

Suggested Time: 30 minutes

- (S) Group: Tell ss that now they will be listening for much of the connected speech phenomena that we've seen so far.
- Play Neil Young's Old Man and just ask the ss to listen. Then question the class on what they understood.
- Distribute the class copies of Neil Young - Old Man (student copy). When the ss see the number of blanks, tell them not to panic and that we'll listen to the song a few times, so that they will have a chance to fill in the blanks. Let them know that many of the blanks are repeated lyrics.
- Play the song 2-3 times for the ss. (start at around 40 seconds in order to skip the intro).
- Give ss a chance to compare their answers in a group of four, so they can discuss the pronunciation.
- Then show the PPP slides with the full lyrics.
- Have ss read the lyrics into the computer, recording themselves and then checking their recording to see if they are using the connected speech phenomena.

- (NS) Group: Tell ss that now they will be listening for much of the connected speech phenomena that we've seen so far.
- Play Steve Jobs – Rule for Success and just ask the ss to listen. Then question the class on what they understood.
- Distribute the class copies of Steve Jobs – Rule for Success (student copy). When the ss see the number of blanks, tell them not to panic and that we'll listen to the audio a few times, so that they will have a chance to fill in the blanks. Let them know that many of the blanks are repeated words.
- Play the audio a few times for the ss. (May need to be 4 times or so.)

- Give ss a chance to compare their answers in a group of four, so they can discuss the pronunciation.
- Then read the script using the Steve Jobs – Rule for Success (Answer Key)
- Have ss read the script into the computer, recording themselves and then checking their recording to see if they are using the connected speech phenomena.

Homework:

- Go to YouTube and look for videos similar to the last one done in class. Listen a number of times, try to understand and repeat/mimic what you can.

Slide 1	LESSON 5: Connected Speech and Linking 2 ❖ A rap ❖ Review of Connected Speech and Linking 1 ❖ Connected Speech and Linking 2 ❖ Homework can't you k ɑ: n tʃ j u:	Image: http://www.tipsforenglish.com/wp-content/uploads/2011/08/cantYou-300x218.jpg
Slide 2	CONNECTED SPEECH AND LINKING 2 ▪ /t/ + /y/ = /tʃ/ not yet, what you ▪ some auxiliary verbs + "to" got to, going to, have to, want to	
Slide 3	OLD MAN (NEIL YOUNG) Old man look at my life, I'm a lot like you were. Old man look at my life, I'm a lot like you were. Old man look at my life, Twenty four and there's so much more Live alone in a paradise That makes me think of two. Love lost, such a cost, Give me things that don't get lost. Like a coin that won't get tossed Rolling home to you. Old man take a look at my life I'm a lot like you I need someone to love me the whole day through Ah, one look in my eyes and you can tell that's true.	

Slide 4	Lullabies, look in your eyes, Run around the same old town. Doesn't mean that much to me To mean that much to you. I've been first and last Look at how the time goes past. But I'm all alone at last. Rolling home to you. Old man take a look at my life I'm a lot like you I need someone to love me the whole day through Ah, one look in my eyes and you can tell that's true. Old man look at my life, I'm a lot like you were. Old man look at my life, I'm a lot like you were.	
Slide 5	HOMEWORK Go to YouTube and look for videos similar to the last one done in class. Listen a number of times, try to understand and repeat/mimic what you can.	

Whatcha Gonna Do? (Burton Crane)

Here are the lyrics. Note that the grammar will be different than what you have been taught. It is not grammar that you should use. In addition, there is some slang that you should understand, but not use yourself.

Momma says, “You go to school”.

She don’t know that it’s a tool.

I’d rather cut¹ and sip some poo² (¹ not go, ²brew=beer)

If she was me, she’s do it too.

Whatcha gonna do? Whatcha gonna do?

Whatcha gonna do? Whatcha gonna do?

“You better wake up”, Momma said,

“And straighten out what’s in your head”.

“Time that you got in some bread³ (³made some money)

A fine new girl who you will wed”

Chorus

I went to a dance and met this chick⁴ (⁴girl, woman)

She gave me her number, we seemed to click

I rang her up and took her to a flick⁵. (⁵movie)

She left me cold thinking she’s sick.

Chorus

Everything that came my way was a fake.

Things got so bad I didn’t want to wake.

‘Till from the oven I smelled the cake.

Momma had baked me a birthday cake.

Whatcha gonna do? Whatcha gonna do?

Whatcha gonna do?

Eat it!

**Part 1**

- ① Listen to the tape or your teacher say the following pronunciation chant. Snap your fingers to the rhythm. Pay attention to the pronunciation of the reductions.

A: Do you **wanna** get a prize?
Do you **wanna** get a prize?

B: Yes, I **wanna** get a prize.
Yes, I **wanna** get a prize.

A: First you **hafta** send the money.
First you **hafta** send the money.

B: I don't **wanna** send the money.
I don't **wanna** send the money.

A: You **hafta** send it now.
You **hafta** send it now.

B: I'm **gonna** call the cops.
I'm **gonna** call the cops.

- ② Work in pairs. Student A is a con artist. Student B is a victim. Read the chant aloud three times. Start slowly and speed up each time. Remember to reduce the boldfaced words. Switch roles and repeat the chant three more times.

- ③ Work in two groups. Group A is the con artist. Group B is the victim. As a group, read the chant one time. Then switch roles and read it again.

(NEIL YOUNG)

(Student Copy)

_____ my life,
_____ were.

_____ my life,
_____ were.

_____ my life,
Twenty four

_____ much more
_____ paradise
That makes me think of two.

Love _____ a cost,
Give me things
_____ lost.
Like a coin _____ tossed
Rolling home to you.

_____ look at my life
_____ like you
I need someone _____ love me
the whole day through
Ah, one look in my eyes
and you can tell _____ .

Lullabies, _____ your eyes,
_____ the same _____ .
_____ to me
_____ to you.

I've been _____ last
_____ how the time goes past.
_____ all alone _____ .
Rolling home _____ you.

_____ look at my life
_____ like you
I need someone _____ love me
the whole day through
Ah, one look in my eyes
and you can tell _____ .

_____ my life,
_____ were.
_____ my life,
_____ were.

OLD MAN (NEIL YOUNG)

(Answer Key)

Old man look at my life,

I'm a lot like you were.

Old man look at my life,

I'm a lot like you were.

...

The rest of the lyrics may be found at:

<http://www.azlyrics.com/lyrics/neilyoung/oldman.html>

Steve Jobs Explains the Rules for Success

(Student Copy)

“People say you _____ have _____ passion for _____ doing and it’s totally true. _____ the reason is because it’s so hard _____ you _____, any rational person would give up. It’s really hard. And you _____ do it over _____ time. So if you don’t love it, if you’re not _____ fun doing it, you don’t really love it, (uhh) _____ . And that’s what happens to most people, actually. If you really look at the ones that (uhh) ended up, you know, being successful (unquote) in the eyes of society and the ones that _____, _____, it’s the ones who were successful loved what they did so they could persevere, you know, when it _____ really tough. And the ones that _____ love it quit _____ they’re sane, right? Who would _____ with this stuff if you _____ love it?

So it’s a lot of hard work and it’s a lot of worrying constantly _____ (uhh) if you _____ love it, _____ fail. So _____ love it and _____ got to have passion and I think that’s the high-order bit.

The second thing is, (uhm) _____ be a really good _____ because no matter how smart _____, you need a team _____ great people and _____ figure out how to (how to) size people up fairly quickly, make decisions without knowing people too well and hire them and, you

know, see how _____ do and refine your intuition and be able _____ help, you

know, build _____ that can eventually just, you know, _____

_____ you need _____ you.”

Steve Jobs Explains the Rules for Success

(Answer key)

“People say you have to have a lot of passion for what you’re doing and it’s totally true. And the reason is because it’s so hard that if you don’t, any rational person would give up. It’s really hard. And you have to do it over a sustained period of time. So if you don’t love it, if you’re not having fun doing it, you don’t really love it, you’re going to give up. And that’s what happens to most people, actually. If you really look at the ones that ended up, you know, being “successful” in the eyes of society and the ones that didn’t, oftentimes, it’s the ones who were successful loved what they did so they could persevere, you know, when it got really tough. And the ones that didn’t love it quit because they’re sane, right? Who would want to put up with this stuff if you don’t love it?

So it’s a lot of hard work and it’s a lot of worrying constantly and if you don’t love it, you’re going to fail. So you’ve got to love it and you’ve got to have passion and I think that’s the high-order bit.

The second thing is, you’ve got to be a really good talent scout because no matter how smart you are, you need a team of great people and you’ve got to figure out how to size people up fairly quickly, make decisions without knowing people too well and hire them and, you know, see how you do and refine your intuition and be able to help, you know, build an organization that can eventually just, you know, build itself because you need great people around you.”

LESSON PLAN - DAY # 6	Date: Oct. 9 th , 2012
Pronunciation Areas: Connected Speech and Linking 3	
Anticipated Problems for Students:	Solution:
Anticipated Problems for Teacher:	Solution:
Goals: 1. To review and recycle forms covered in Connected Speech and Linking 2 2. To teach additional forms of connected speech Objectives: 1. Ss will orally practice connected speech phenomena covered in the last class 2. Ss will write phrases containing a variety of forms of connected speech based on what they hear 3. Ss will read connected speech that they have written	Complete List of Materials • Lesson Plan – Lesson 6 • Lesson 6 - PPP • Class copies of Do and Don't Sentences Practice • Class copies of Appendix 2: Common Words with Omitted Syllables • Burton Cummings- Break It To Them Gently (Audio) • Break It To Them Gently - Burton Cummings - (Student Copy) • Break It To Them Gently - Burton Cummings - (Answer Key) • Link to audio of A Stolen Butterfinger • A stolen Butterfinger (Student Copy) • A stolen Butterfinger (Answer Key)
Homework: Review material and practice either on your own or with others.	Feedback:

AT START OF CLASS:

Suggested Time: 10 minutes

- Take some time to ask ss about any problems they might be having noticing or saying connected speech phenomena.

- Ask them if they have noticed anything we haven't seen yet in class.

Review and Warm up:

Suggested Time: 10 minutes

- Distribute class copies of Do and Don't Sentences Practice.
- Go through the instructions with the ss and then give them time to practice. Go around and listen in and help.

Presentation:

Suggested Time: 15 minutes

- For the first area, words ending in *-ing* and said *-in'*, this will be easy to teach and most likely will have been noticed by the students and or mentioned in a teachable moment. Remember, some of the blanks we examples of this articulation.
- For the second area, mention that this happens a lot in function words (can't, don't, shouldn't etc.) and especially with the final /t/, which is a weak sound in General American English (it even changes to /r/ when it is intervocalic). Also mention that this happens in word-final consonant clusters before another word or syllable that begins with a consonant. Show the slide with examples.

Activity 1: *Talking with omitted syllables*

Suggested Time: 15 minutes

- Distribute class copies of Appendix 2: Common Words with Omitted Syllables.
- Read through the handout to model the pronunciation and then ask ss to make up an impromptu dialogue using phrases in the handout.

Activity 2: Hearing and practicing connected speech

Suggested Time: 30 minutes

- (S) Group: Tell ss that now they will be listening for much of the connected speech phenomena that we've seen so far.
- Play Burton Cummings - Break it to Them Gently and just ask the ss to listen. Then question the class on what they understood.
- Distribute the class copies of Burton Cummings - Break it to Them Gently (student copy). When the ss see the number of blanks, tell them not to panic and that we'll listen to the song a few times, so that they will have a chance to fill in the blanks. Let them know that many of the blanks are repeated lyrics.
- Draw their attention, though, to the fact that with this song, they don't know how many words each blank has (which is different from the previous fill-in-the-blanks activities).

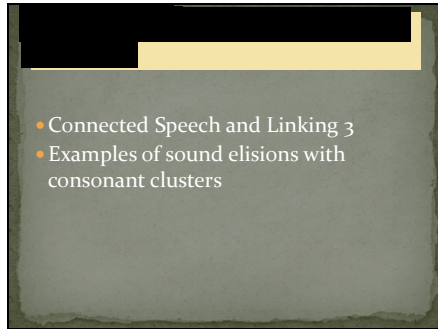
- Play the song 2-3 times for the ss.
- Give ss a chance to compare their answers in a group of four, so they can discuss the pronunciation.
- Then show the PPP slides with the full lyrics.
- In groups of 2-3, have ss change the lyrics and create their own and then practice speaking them into the computer, recording themselves and then checking their recording to see if they are using the connected speech phenomena.

- (NS) Group: Tell ss that now they will be listening for much of the connected speech phenomena that we've seen so far.
- Play audio of A Stolen Butterfinger and just ask the ss to listen. Then question the class on what they understood.
- Distribute the class copies of A Stolen Butterfinger (student copy).
- Draw their attention to the fact that with this text, they don't know how many words each blank has (which is different from the previous fill-in-the-blanks activities).
- Play the story twice for the ss.
- Give ss a chance to compare their answers in a group of four, so they can discuss the pronunciation.
- Then read the text slowly so they can verify their answers.
- In groups of 2-3, have ss change the ending and create their own and then practice speaking this into the computer, recording themselves and then checking their recording to see if they are using the connected speech phenomena.

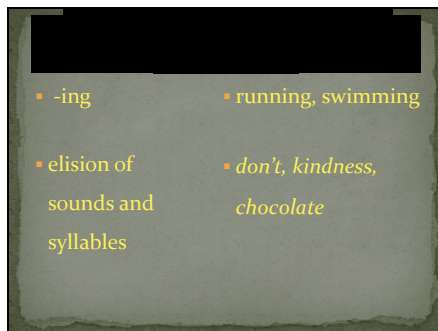
Homework:

Review material and practice either on your own or with others.

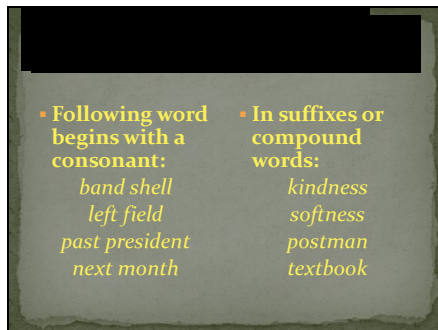
Slide 1



Slide 2



Slide 3



Examples taken from: Avery & Ehrlich (1992) Teaching American English Pronunciation. p.86.

Slide 4

BREAK IT TO THEM GENTLY (BURTON CUMMINGS)
ANSWER KEY

Break it to them gently when you tell my mom and dad
When you see my baby sister be as kind as you can
Break it to my grandma, who said "That boy's wild and bad"
Break it to them gently when you tell them that I won't be coming home again

Because I'm running with a gun and it isn't any fun as a fugitive
Fighting for my life and I don't know if I'll make it alone
Running with a gun and it isn't any fun as a fugitive
God I want to go home
Lord I wish I was home

When you see my lady with the twinkle in her eyes
Tell it to her softly and hold her if she cries
Tell her that I love her and I will until the day I die
Tell it to her gently when you tell her that I won't be coming home again

I got in too deep with strangers
Thinking they could help me find my way
Nobody warned me of the dangers
And it's always the young and foolish that have to pay

Slide 5

So break it to them gently when you tell my mom and dad
Thank them for the good years and all the loving that I had
And break it to my grandma, who said "the boy is wild and bad"
Break it to them gently when you tell them that I won't be coming home again

Running with a gun and it isn't any fun as a fugitive
Fighting for my life and I don't know if I'll make it alone
Running with a gun and it isn't any fun as a fugitive
Lord I want to go home
Lord I want to go home

You got to break it to them gently, break it to them gently
You got to break it to them gently
Got to really try to roll them
Got to break it to them gently
Got to really try to soothe them, got to really try to soothe them
Got to really try to roll them
You got to roll it to my mother
Got to roll it to my grandma
Got to roll the old lady
Roll it to my mother, roll it to my mother and roll the old lady
Roll it to my grandma, she's damn near eighty
Roll the old lady

DO AND DON'T SENTENCES

Practice:

It's important to relax and maintain good health. Write do and don't sentences using the following nouns and verbs and then practice saying them with your partner.

Verbs: *eat, miss, forget, lose, get, count*

Nouns: *sleep, exercise, temper, rest, vitamins, calories, vegetables*

DO

Eat _____ your vegetables.

vitamins.

_____ your _____.

_____ your _____.

_____ your _____.

DON'T

Don't _____ forget your

Don't _____ your _____.

Don't _____ your _____.

Don't _____ your _____.

Taken from Celce-Murcia et al (2010) Teaching Pronunciation. Page 179.

Appendix 2: Common Words with Omitted Syllables

Native speakers usually omit one of the syllables in these words. They omit the unstressed syllable that comes after the syllable with primary stress.

1. EV-ning	every evening Saturday evening Good evening!	6. INT-res ting	It's interesting. very interesting quite interesting
2. CHOC-ate	chocolate cake chocolate sauce chocolate milk	7. DIFF-erent	different from completely different different ways
3. CAM-era	video camera security camera digital camera	8. FAV-rite	favorite one favorite book favorite teacher
4. BUS-ness	small business family business business associate	9. REST-aurant	local restaurant popular restaurant Chinese restaurant
5. VEG-eta ble	bowl of vegetable soup fresh vegetables green vegetables	10. FAM-ly	the whole family family friends family income

BREAK IT TO THEM GENTLY (BURTON CUMMINGS)

Break it to them gently when you tell my mom _____ dad
When you see my baby sister _____ kind as _____ can
Break it to my grandma, who said " _____ "
Break it to them gently when you tell _____ that I _____ home again
_____ I'm _____ with a gun _____ it _____ any _____ a fugitive
_____ for my life and I _____ if I'll _____ alone
_____ with a gun and it _____ any _____ a fugitive
God I _____ go home
Lord I wish I was home

When you see my lady with the twinkle in _____ eyes
Tell it to _____ softly and hold _____ if she cries
_____ that I love _____ and I will _____ the day I die
Tell it to her gently when you tell _____ that I _____ home again
I _____ too deep with strangers
_____ they could help me find my way
Nobody warned me of the dangers
And it's always the young and foolish _____ pay

So break it to them gently when you tell my mom _____ dad
Thank them for the good years and all the _____ had
_____ break it to my grandma, who said " _____ "
Break it to them gently when you tell _____ that I _____ home again
_____ with a gun and it _____ any _____ a fugitive
_____ for my life and I _____ if I'll _____ alone
_____ with a gun and it _____ any _____ a fugitive
Lord I _____ go home
Lord I _____ go home

You _____ break it to them gently, break it to them gently
You _____ break it to them gently
_____ really try to _____
_____ break it to them gently
_____ really try to soothe them, _____ really try _____ soothe
them
_____ really try to _____
You _____ to my mother
_____ to my grandma
_____ old lady
_____ to my mother, _____ to my mother _____ old lady
_____ grandma, she's damn near eighty

BREAK IT TO THEM GENTLY (BURTON CUMMINGS)

ANSWER KEY

Break it to them gently when you tell my mom and dad
When you see my baby sister be as kind as you can
Break it to my grandma, who said "That boy's wild and bad"
Break it to them gently when you tell them that I won't be coming home again

Because I'm running with a gun and it isn't any fun as a fugitive
Fighting for my life and I don't know if I'll make it alone
Running with a gun and it isn't any fun as a fugitive
God I want to go home
Lord I wish I was home

When you see my lady with the twinkle in her eyes
Tell it to her softly and hold her if she cries
Tell her that I love her and I will until the day I die
Tell it to her gently when you tell her that I won't be coming home again

I got in too deep with strangers
Thinking they could help me find my way
Nobody warned me of the dangers
And it's always the young and foolish that have to pay

So break it to them gently when you tell my mom and dad
Thank them for the good years and all the loving that I had
And break it to my grandma, who said "the boy is wild and bad"
Break it to them gently when you tell them that I won't be coming home again

Running with a gun and it isn't any fun as a fugitive
Fighting for my life and I don't know if I'll make it alone
Running with a gun and it isn't any fun as a fugitive
Lord I want to go home
Lord I want to go home

You got to break it to them gently, break it to them gently
You got to break it to them gently
Got to really try to roll them
Got to break it to them gently
Got to really try to soothe them, got to really try to soothe them
Got to really try to roll them
You got to roll it to my mother
Got to roll it to my grandma
Got to roll the old lady
Roll it to my mother, roll it to my mother and roll the old lady
Roll it to my grandma, she's damn near eighty
Roll the old lady

A STOLEN BUTTERFINGER

Wade, 12, _____ a corner market. He _____.
He _____ anybody. He _____ candy bar into
his jacket pocket. It was a Butterfinger. He loved Butterfingers.
He _____ again. The coast was clear. He walked
_____ clerk. He walked toward the
_____. The clerk said, "Stop, please." Wade ignored
the clerk. The clerk shouted, "Stop, _____ the
police!" Wade stopped. He walked back to the clerk. The clerk
said, "_____ that mirror?" Wade _____
near the ceiling. There was a big, convex mirror. The candy bar
section was reflected in the mirror. The clerk had
_____. "_____ the candy bar," the clerk
said. Wade dug _____ pocket. He
_____ the clerk. "_____!" the clerk said.
Wade said, "_____ candy bar." The clerk said,
"_____."



<http://www.butterfinger.com/images/products/bf-detail.jpg>

A STOLEN BUTTERFINGER (answer key)

Wade, 12, was in a corner market. He looked around. He couldn't see anybody. He slipped a candy bar into his jacket pocket. It was a Butterfinger. He loved Butterfingers. He looked around again. The coast was clear. He walked past the clerk. He walked toward the front door. The clerk said, "Stop, please." Wade ignored the clerk. The clerk shouted, "Stop, or I'll call the police!" Wade stopped. He walked back to the clerk. The clerk said, "Do you see that mirror?" Wade looked up near the ceiling. There was a big, convex mirror. The candy bar section was reflected in the mirror. The clerk had seen everything. "Give me the candy bar," the clerk said. Wade dug it out of his pocket. He gave it to the clerk. "Shame on you!" the clerk said. Wade said, "It's just a candy bar." The clerk said, "Get out of here."

<http://www.eslyes.com/nyc/nyc/nycs247.htm>

LESSON PLAN - DAY # 7	Date: Oct. 10 th , 2012
Pronunciation Areas: All areas covered	
Anticipated Problems for Students:	Solution:
Anticipated Problems for Teacher:	Solution:
Goals: 1. To review the different areas taught. Objectives: 1. Ss will use the suprasegmental features in controlled and free speaking contexts.	Complete List of Materials • Lesson Plan – Lesson 7 • Class copies of Sentence Stress and Rhythm • Class copies of Disagreeing and Correcting • Photos for Discussion – PPP
Homework: - Practice speaking using the materials from class. Go over activities, listen to the audios, talk with classmates.	Feedback:

AT START OF CLASS:

Warm-up: Connected speech and linking practice

Suggested Time: 10 minutes

- Ask ss to think of the names of different movies or books and come up to the board, write them down and show and then say the linking. For example: *Lost in Translation*.

- Have all ss practice saying what's on the board.

Review:

Activity 1: Brainstorm

Suggested Time: 10 minutes

- In groups, have ss brainstorm on the different pronunciation areas that we've covered in the course. (thought groups, focus words, syllable, stress, peak, schwa, connected speech and linking).

- Then have everyone say one these out to the class explain what each one is and give an example.

Activity 2: Reviewing sentence stress and rhythm

Suggested Time: 15 minutes

- Distribute class copies of sentence stress and rhythm.

Go through chart 1 with the class and then, in pairs, have the ss change some of the words in the phrases, but keep the same stress patterns. Provide an example for them first. Go around and help and listen in.

- Then do chart 2 as a class chant with each pair taking on one sentence. Start with pair 1 saying their sentence, then after two times saying it, pair 2 enters saying their sentence and so on.
- Then, do it again, but completely switching the sentences that the pairs have.
- Then briefly review content and function words.

Activity 3: Practicing changing focus words

Suggested Time: 15 minutes

- Tell ss that now we're going to spend a little bit of time on using focus words to disagree and correct. We saw this a bit on Day 3, but this activity is to reinforce and expand on that.
- Distribute class copies of Disagreeing and Correcting.
- Play the audio and then have ss practice the examples in pairs and then switch roles.
- Then have the pairs take turns disagreeing and correcting things that they've made up. For example, say, "Chile is in North **America**." (Ask a student to correct me. He or she should say, "No, Chile is in **South** America.")

Activity 4: A party (fluency practice)

Suggested Time: 10 minutes

- Tell ss that we're going to pretend to plan a party. Ask them what some different things you take to a party are.
- Write these on the board in columns according to the number of syllables
- Instructions: Have each student say what they are bringing, but starting with what the people before said they are bringing. That is, "We're having a party and I'm bringing a bottle of wine. We're having a party and Karen's bringing a bottle of wine and I'm bringing potato chips. Etc."
- Remind ss to use thought groups, focus words, correct stress etc.

Activity 5: Photos for discussion

Suggested Time: 20 minutes

- Tell ss that now we're going to look at some different images from the internet and their job is to say something about them using the different pronunciation features we've covered.
- Pass out cookies and make it a relaxed activity that students do in 2-3 groups.

Homework:

- Practice speaking using the materials from class. Go over activities, listen to the audios, talk with classmates.

End of class:

- Thank the ss for participating etc.
- Remind them of tomorrow's listening and speaking tests and that they would also have a short questionnaire to answer.
- Play the Gotye parody, Some Study That I Used To Know (Ss will know the original.)

SENTENCE STRESS AND RHYTHM

Chart 1: Rhythm practice comparing word and sentence stress

<u>Rhythmic Pattern</u>	● •	• ●
word	mother	attend
sentence	Do it.	You did?
	Pay them.	It hurts.
<u>Rhythmic Pattern</u>	• ● •	• • ●
word	abandon	guarantee
sentence	I saw you.	Have some cake.
	We found it.	Where's the beef?
<u>Rhythmic Pattern</u>	• • ● •	• • ● • •
word	education	nationality
sentence	Mary saw it.	Come to Canada.
	John's a lawyer.	Where's your bicycle?
<u>Rhythmic Pattern</u>	• • • ● •	• • • • ● •
word	communication	electrification
sentence	I want a soda.	We took a vacation.
	I think he's got it.	I went to the station.

Celce-Murcia et al (2010, p. 209)



<http://us.cdn4.123rf.com>

Celce-Murcia et al
(2010, p. 210)

Chart 2: “Cats Chase Mice” rhythm drill

1.	CATS		CHASE	MICE
2.	The CATS	have	CHASED	MICE
3.	The CATS	will	CHASE	the MICE
4.	The CATS	have been	CHASing	the MICE
5.	The CATS	could have been	CHASing	the MICE

CONTENT WORDS	FUNCTION WORDS
- Often stressed	- Usually unstressed, unless in final position or when used emphatically
nouns	articles
main verbs	auxiliary verbs
adjectives	personal pronouns
possessive pronouns	possessive adjectives
demonstrative pronouns	demonstrative adjectives
interrogatives	prepositions
not / negative contractions	conjunctions
adverbs	
adverbial particles	

Taken and adapted from Celce-Murcia et al (2010, p. 210)

F *Pair work: Disagreeing and correcting*

- 1** Listen. Notice how the focus word in the second sentence below (“month”) is a correction for the word “week” in the first sentence.

A: He was in Spain for a week.

B: No, he was in Spain for a month.

- 2** Listen. In the second sentence below, the word “France” is a correction for the word “Spain” in the first sentence.

A: He was in Spain for a week.

B: No, he was in France for a week.

Focus Rule 6

When there is a disagreement or a correction, the word that corrects the information from the previous statement is the new focus word.

- 3** Practice saying these dialogues with a partner. Emphasize the underlined focus words. Take turns as Speaker A and Speaker B.

1. A: I buy books at the library.

B: No, you borrow books at the library.

2. A: I buy books at the library.

B: No, you buy books at the bookstore.

3. A: Madrid is the capital of Germany.

B: No, it's the capital of Spain.

4. A: Madrid is the capital of Germany.

B: No, Berlin is the capital of Germany.

5. A: “Actual” means “in the present time.”

B: No, “actual” means “real.”

6. A: A ship is smaller than a boat.

B: I don't think so. A ship is bigger than a boat.

7. A: Is Dallas in California?

B: No, it's in Texas.

8. A: Is Dallas in California?

B: No, but San Francisco is in California.

Slide 1

PHOTOS FOR DISCUSSION

Slide 2



http://farm2.static.flickr.com/1377/3352909583_17c8de6888_o.jpg

Slide 3

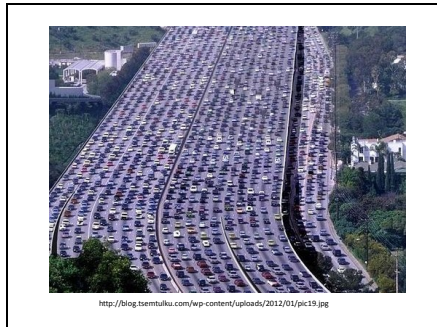


<http://media.smashingmagazine.com/wp-content/uploads/2010/05/awardwinningphotos30.jpg>

Slide 4



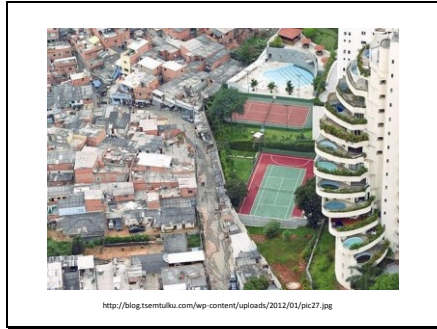
Slide 5



Slide 6



Slide 7



Slide 8



Slide 9



Slide 10



Slide 11



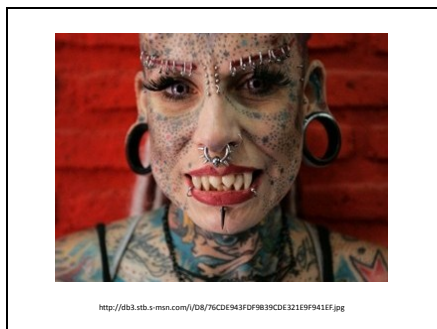
Slide 12



Slide 13



Slide 14



Slide 15



Slide 16



SOME STUDY THAT I USED TO KNOW

(COLLEGE HUMOR)

Now and then I think of what I learned in high school
Like AP Bio an-d British Literature
Is that igneous or metamorphic?
I don't need to write in Iambic.
And I'll admit I don't know shit about Millard Fillmore

...

The rest of the lyrics may be found at

<http://lybio.net/college-humor-some-study-that-i-used-to-know-gotye/parody/>

Clear listening test

How you hear English is closely connected with how you speak English.

Part 1 **Consonants**

[10 points]



Listen. You will hear either sentence **a** or sentence **b**. Circle the letter of the sentence you hear.

1. a. Do you want everything?
 (b) Do you wash everything?
2. a. They saved old bottles.
 b. They save old bottles.
3. a. She loves each child.
 b. She loved each child.
4. a. We'll put it away.
 b. We've put it away.
5. a. He spills everything.
 b. He spilled everything.
6. a. Does she bring her card every day?
 b. Does she bring her car every day?
7. a. What does "leave" mean?
 b. What does "leaf" mean?
8. a. Who'll ask you?
 b. Who'd ask you?
9. a. We wash all of them.
 b. We watch all of them.
10. a. He put the tickets away.
 b. He put the ticket away.
11. a. Is this the long road?
 b. Is this the wrong road?

Part 2 *Vowels*

[10 points]



Listen. You will hear either sentence **a** or sentence **b**. Circle the letter of the sentence you hear.

1. (a.) Did you bring the bat?
b. Did you bring the bait?
2. a. I prefer this test.
b. I prefer this taste.
3. a. It's a good bet.
b. It's a good bit.
4. a. It's on the track.
b. It's on the truck.
5. a. The men worked hard.
b. The man worked hard.
6. a. How do you spell "scene"?
b. How do you spell "sin"?
7. a. How do you spell "luck"?
b. How do you spell "lock"?
8. a. We used a map.
b. We used a mop.
9. a. Is John coming?
b. Is Joan coming?
10. a. Everybody left.
b. Everybody laughed.
11. a. I ran to school every day.
b. I run to school every day.

Part 3 *Syllables*

[10 points]



Listen and write the number of syllables in each word.

- | | | | |
|---------------|-------|-----------------|-------|
| 1. easy | 2 | 7. opened | |
| 2. closet | | 8. first | |
| 3. sport | | 9. caused | |
| 4. clothes | | 10. Wednesday | |
| 5. simplify | | 11. arrangement | |
| 6. frightened | | | |

Part 4 *Word stress*

[10 points]



Listen. In each word, one syllable is stressed more than the others. Underline the stressed syllable in each word.

- | | |
|----------------|-------------------|
| 1. arrangement | 7. Europe |
| 2. political | 8. information |
| 3. photograph | 9. economy |
| 4. photography | 10. economic |
| 5. Canadian | 11. participating |
| 6. geography | |

Part 5 *Emphasizing focus words*

[20 points]



Listen to the following dialogue. In each sentence, one word is emphasized more than the others. Underline the emphasized word in each sentence.

A: Do you think food in this country is expensive?

B: Not really.

A: Well, I think it's expensive.

B: That's because you eat in restaurants.

A: Where do you eat?

B: At home.

A: You must like to cook.

B: Actually, I never cook.

A: So what do you eat?

B: Usually, just cheese.

A: That's awful!

Part 6 *De-emphasizing with contractions and reductions*

[20 points]



Listen. You will hear each sentence two times. Write the missing words in the blanks.

1. Do you think *she's* OR *she is* in her room?
2. you ask?
3. work good?
4. Please the information.
5. want food?

6. How you been here?
7. Matt done lately?
8. Why come so early?
9. they gone?
10. We'd like some vegetables.
11. They'll need glasses.

Part 7 *Thought groups*

[20 points]



Listen. You will hear sentence **a** or sentence **b**. After you hear the sentence two times, answer the question that follows.

1. a. John said, "My father is in the kitchen."
b. "John," said my father, "is in the kitchen."

Question: Who was speaking? *my father*

2. a. The president shouted, "That reporter is lying!"
b. "The president," shouted that reporter, "is lying!"

Question: Who shouted?

3. a. She wants pineapples.
b. She wants pie and apples.

Question: What does she want?

4. a. Would you like a Super Salad?
b. Would you like a soup or salad?

Question: What were you offered?

5. a. We used wooden matches to start the fire.
b. We used wood and matches to start the fire.

Question: What was used to start the fire?

6. a. He sold his houseboat and car.
b. He sold his house, boat, and car.

Question: How many things did he sell?

Clear listening test

How you hear English is closely connected with how you speak English.

Part 1 **Consonants**

[10 points]



Listen. You will hear either sentence **a** or sentence **b**. Circle the letter of the sentence you hear.

Teacher: Use the Class Audio Program or read the sentences in bold type.

1. a. Do you want everything?
 (b.) Do you wash everything?
2. a. They saved old bottles.
 (b.) They save old bottles.
3. a. She loves each child.
 (b.) She loved each child.
4. a. We'll put it away.
 (b.) We've put it away.
5. **(a.) He spills everything.**
 b. He spilled everything.
6. **(a.) Does she bring her card every day?**
 b. Does she bring her car every day?
7. **(a.) What does "leave" mean?**
 b. What does "leaf" mean?
8. **(a.) Who'll ask you?**
 b. Who'd ask you?
9. a. We wash all of them.
 (b.) We watch all of them.
10. **(a.) He put the tickets away.**
 b. He put the ticker away.
11. a. Is this the long road?
 (b.) Is this the wrong road?

Part 2 Vowels

[10 points]



Listen. You will hear either sentence **a** or sentence **b**. Circle the letter of the sentence you hear.

Teacher: Use the Class Audio Program or read the sentences in bold type.

1. (a.) Did you bring the bat?
b. Did you bring the bait?
2. (a.) I prefer this test.
b. I prefer this taste.
3. a. It's a good bet.
(b.) It's a good bit.
4. a. It's on the track.
(b.) It's on the truck.
5. (a.) The men worked hard.
b. The man worked hard.
6. (a.) How do you spell "scene"?
b. How do you spell "sin"?
7. (a.) How do you spell "luck"?
b. How do you spell "lock"?
8. (a.) We used a map.
b. We used a mop.
9. (a.) Is John coming?
b. Is Joan coming?
10. a. Everybody left.
(b.) Everybody laughed.
11. (a.) I ran to school every day.
b. I run to school every day.

Part 3 Syllables

[10 points]



Listen and write the number of syllables in each word.

Teacher: Use the Class Audio Program or read each word.

- | | | | |
|---------------|---|-----------------|---|
| 1. easy | 2 | 7. opened | 2 |
| 2. closet | 2 | 8. first | 1 |
| 3. sport | 1 | 9. caused | 1 |
| 4. clothes | 1 | 10. Wednesday | 2 |
| 5. simplify | 3 | 11. arrangement | 3 |
| 6. frightened | 2 | | |

Part 4 *Word stress*

[10 points]



Listen. In each word, one syllable is stressed more than the others. Underline the stressed syllable in each word.

Teacher: Use the Class Audio Program or read each word, stressing the underlined syllable.

- | | |
|------------------------|---------------------------|
| 1. <u>arr</u> angement | 7. <u>E</u> urope |
| 2. <u>pol</u> itical | 8. <u>in</u> formation |
| 3. <u>phot</u> ograph | 9. <u>e</u> conomy |
| 4. <u>phot</u> ography | 10. <u>e</u> conomic |
| 5. <u>Can</u> adian | 11. <u>part</u> icipating |
| 6. <u>geogr</u> aphy | |

Part 5 *Emphasizing focus words*

[20 points]



Listen to the following dialogue. In each sentence, one word is emphasized more than the others. Underline the emphasized word in each sentence.

Teacher: Use the Class Audio Program or read the dialogue, emphasizing the underlined words.

A: Do you think food in this country is expensive?

B: Not really.

A: Well, I think it's expensive.

B: That's because you eat in restaurants.

A: Where do you eat?

B: At home.

A: You must like to cook.

B: Actually, I never cook.

A: So what do you eat?

B: Usually, just cheese.

A: That's awful!

Part 6 *De-emphasizing with contractions and reductions*

[20 points]



Listen. You will hear each sentence two times. Write the missing words in the blanks.

Teacher: Use the Class Audio Program or read the sentences in an informal style using contractions and reductions. Accept contractions or full forms in your students' answers. Do not subtract points for spelling errors if you can recognize the words.

1. Do you think she's OR she is in her room?
2. Why'd OR Why did you ask?
3. Is her work good?

4. Please *give him* the information.
5. *Does he* want food?
6. How *long have* you been here?
7. *What's* OR *What has* Matt done lately?
8. Why *did he* come so early?
9. *Where have* they gone?
10. We'd like some *fish and* vegetables.
11. They'll need *cups or* glasses.

Part 7 **Thought groups**

[20 points]



Listen. You will hear sentence **a** or sentence **b**. After you hear the sentence two times, answer the question that follows.

Teacher: Use the Class Audio Program or read the sentences in bold type and the questions that follow.

1. a. John said, "My father is in the kitchen."
b. "John," said my father, "is in the kitchen."
Question: Who was speaking? *my father*
2. a. The president shouted, "That reporter is lying!"
b. "The president," shouted that reporter, "is lying!"
Question: Who shouted? *the president*
3. a. She wants pineapples.
b. She wants **pie and apples**.
Question: What does she want? *pie and apples*
4. a. Would you like a Super Salad?
b. Would you like a **soup or salad**?
Question: What were you offered? *soup or salad*
5. a. We used **wooden matches** to start the fire.
b. We used wood and matches to start the fire.
Question: What was used to start the fire? *wooden matches*
6. a. He sold his houseboat and car.
b. He sold his house, **boat, and car**.
Question: How many things did he sell? *three*

Listening Test # 2

Name: _____

LISTENING TEST #2

You will hear a dialogue and each line will be said twice. Write what you hear in correct written English. You can use acceptable contractions.

1.

2.

3.

4.

5.

6.

7.

8.

9.

10.

LISTENING TEST #2

(1-5 are from *Well Said Intro*, pg. 133.)

Play the recording of the following. Repeat.

1. Did you get it? - 2
2. Got it! - 1
3. What a relief! / Bet it made your day. - 3
4. Sure did. / Want a ride home? - 2
5. Thanks anyway. / I think I'll walk. - 2
6. Why? / What are you afraid of? / I'm not going to get in a bad car crash. - 7
7. That's what you think. / I think you're an accident / waiting to happen. - 3
8. Fine. / Suit yourself. / Don't trip or slip on your way back home. - 5
9. I won't. / Could you make sure that you don't speed? - 4
10. Yes sir! - 1

Marking: *Correct the word mistakes, but only evaluate connected speech.*

Total points: 30

Connected Speech: /30 (minus 1 for each mistake)

Speaking Test # 1

Instructions: *Read the following text to yourself. When you are ready, record yourself by reading it out loud into the computer. Return this piece of paper to your teacher after you have finished recording your speech. You may write on this sheet.*

The Rainbow Passage

When the sunlight strikes raindrops in the air, they act as a prism and form a rainbow. The rainbow is a division of white light into many beautiful colors. These take the shape of a long round arch, with its path high above, and its two ends apparently beyond the horizon.

There is, according to legend, a boiling pot of gold at one end. People look, but no one ever finds it. When a man looks for something beyond his reach, his friends say he is looking for the pot of gold at the end of the rainbow.

Throughout the centuries people have explained the rainbow in various ways. Some have accepted it as a miracle without physical explanation. To the Hebrews it was a token that there would be no more universal floods. The Greeks used to imagine that it was a sign from the gods to foretell war or heavy rain. The Norsemen considered the rainbow as a bridge over which the gods passed from earth to their home in the sky.

Others have tried to explain the phenomenon physically. Aristotle thought that the rainbow was caused by reflection of the sun's rays by the rain. Since then physicists have found that it is not reflection, but refraction by the raindrops which causes the rainbows.

Many complicated ideas about the rainbow have been formed. The difference in the rainbow depends considerably upon the size of the drops; the width of the colored band increases as the size of the drops increases. The actual primary rainbow observed is said to be the effect of a super-imposition of a number of bows. If the red of the second bow falls upon the green of the first, the result is to give a bow with an abnormally wide yellow band, since red and green light when mixed form yellow. This is a very common type of bow, one showing mainly red and yellow, with little or no green or blue.

Sample Transcriptions of Speaking Test # 2

C4 – INITIAL Total time: 2:09 – 4:36 Analyzed: 3:00 – 4:00

A lot of uh more realistic uh books I guess. Um I've read a lot of books in Spanish which were written by uhm Spanish or mostly um Argentinean authors. Um I haven't rich I haven't read sorry a lot of these books um in English um other than the ones in our literature courses now. Uhm a really good book is uhm el Tunel by Ernesto Sabato it's ah a very short book it's about a very crazy character who goes through kind of a mental breakdown. Uh It's a really good book to read so I would recommend it to anyone. Ah In English I have read the Sofi Ginsela I guess about uhm confessions of a shopaholic. It's a really funny book. It's really entertaining. It's a light read.

C5 – FINAL Total time: 0 – 2:35 Analyzed: 1:00 – 2:00

She is from Peru and she always rent like a summer house there in in a in a in a beach ah in the north near Tumancura which is a very famous like town like summer town it's town we will be there for 10 days we'll spend with my friend and family and me and then we will go down to Lima which is the capital of Peru and we will uhm rent uh rooms in a hotel and we will visit places we will walk around the city get to know the city uhm we will be there like 5 days no more and after that we will come back to Santiago so practically I will be there for 15 days and then on the other hand I'm planning with my friends two friends of mine and I are planning to go to Bolivia which is another uh country which is

S1 – FINAL Total time: 2:19 – 4:08 Analyzed: 2:45 – 3:45

the connected speech that uhm a native speaker has so ahm she has give us uhm she have she had she have gave us uhm many useful tips for improvin' our pronunciation skills and ahh however it's very difficult to apply them immediately in our in our in our speech so I hope that uhm my pronunciation will keep improving till I get ahm similar accent to those that are native speakers so I wanted to thank Karen for ah the patience and dedication that she had towards us in

S4 – INITIAL Total time: 2:34 – 6:15 Analyzed: 4:00 – 5:00

I dreamed a dream or my favorite things from actually I don't know what the song where the song is from but ehm it's it's kind of hard to remember because I've actually never seen that musical but anyway I first got obsessed with musicals when I started watching Glee but I don't watch Glee anymore I have actually I actually have the soundtracks of a lot of musicals I haven't even seen I don't know even I don't even know what they are about but I lo I love them and ehm the problem is that none of my friends love them er well some in school love them but not anymore here in the university so when I try to sing the songs or whatever they usually usually uhm talk to me like rudely

S5 – INITIAL Total time: 0 – 2:17 Analyzed: :30 – 1:30

The whole book eh it's situated in in the middle of this road where the two men are waiting and they are only talking about ehm how how how they their lives have passed and and they are mainly talking about eh their lives in a very in a very complicated way I mean they talk about their relationship eh through the years as they have been together for 15 years wondering around the world wo I guess and the thing is that they they while they are waiting they met a a man with his slave and they talk a lot and they get to to realize

S9 – INITIAL Total time: 0 – 2:35 Analyzed: :30 – 1:30

And work in another country and live in an English speaker country. That's it's my plan for the future I will like to go abroad and get a get opportunity in the field and live not in Chile tsk uhm I'm trying to find some by the time that I will have graduate I want to uhm get a scholarship and go to study to England or tsk uhm United States I'm thinking also in Canada but I don't know what I li what I'm thinking that I like to do is screen writing I want to work in the film industry I'm not completely sure if that's the path that I will choose

NS2 – INITIAL Total time: 0 – 2:06 Analyzed: :30 – 1:30

So I'm gonna visit her and I'm gonna give her the present I bought her which is a book by Franz Kafka it's a very nice book and it has like really pretty illustrations that's why I bought it to her uhm so that's what I'm going to do then I have like a huge window until like two oclock and then I can go to my home this is going to be kind of a nice week cause I don't have many tests actually I just have one test which is a grammar test and I'm not very prepared I hate grammar and uhm I'll make sure that I'll study like uhm during the week a whatever uhm I'm gonna talk about uhm I think it's gonna be a good week because I only have this test

NS3 – FINAL Total time: 0 – 2:41 Analyzed: :30 – 1:30

And ah the story it's not when they are young but it's it's divided in three parts when they are young then when they are like 18 or 19 and then in their adult life 24 or 35 approximately it's a a really good story it's very well written because the the author just keeps the reader in the dark for the entire novel and it's not only the readers who are in the dark but it's the protagonists too so there's there's empathy produced but from the readers to the protagonists it's like you are friends of them and it's narrated for from Kaheach which is one of the girls and she tells the story

NS5 – FINAL

Total time: 0 – 2:50 Analyzed: 1:00 – 2:00

They don't have to worry about it being such a ??? and everything which is really funny because actually it's other way around and here you don't have to worry about you have like lots of supermarkets and and all eh the institutions are here so if something that's going to happen the the answers and help and everything is going to be right away you don't have to wait like two or three days like they have to wait the last time there was a huge earthquake so I'm really sure that

people from Santiago are really good at exaggerating everything from a little earthquake to I don't know what else it's really funny and the worst part of that is that all the cell phone lines are not working and they should be fixed by now but it's not