



universität
wien

DIPLOMARBEIT / DIPLOMA THESIS

Titel der Diplomarbeit / Title of the Diploma Thesis

„Deutsche Akzente im Französischen und Spanischen – eine
korpusphonologische Studie als Grundlage einer
varietätensensiblen Aussprachedidaktik“

verfasst von / submitted by

Andrea Klingenbrunner

angestrebter akademischer Grad / in partial fulfilment of the requirements for the degree of

Magistra der Philosophie (Mag. phil)

Wien, 2017 / Vienna, 2017

Studienkennzahl lt. Studienblatt /
degree programme code as it appears on
the student record sheet:

A 190 347 353

Studienrichtung lt. Studienblatt /
degree programme as it appears on
the student record sheet:

Lehramtsstudium UF Französisch UF Spanisch

Betreut von / Supervisor:

Univ.-Prof. Dr. habil. Elissa Pustka, Dipl.-Journ. Univ. M.A

Danksagung

Das erste große Dankeschön gilt meiner Betreuerin Prof. Dr. Elissa Pustka, die mich zur Auswahl dieses anspruchsvollen Themas animiert, und immer durch ihre gute Betreuung unterstützt hat und mir zeigen konnte, wozu ich fähig bin.

Außerdem möchte ich mich bei meiner Familie für ihre fortwährende Unterstützung bedanken. Meiner Mutter Daniela, die mich beruhigt und mit ihren eigenen Kenntnissen unterstützt hat, meinem Vater Fritz, der mir selbst in der schwierigsten Phase bedingungslos unter die Arme gegriffen hat, meiner Schwester Lisi, die mich immer zum Lachen bringt, sowie meinem Bruder Alexander und seiner eigenen kleinen Familie, die mich ablenkten, wenn ich es brauchte.

Ebenfalls bedanken möchte ich mich bei allen Beteiligten am Projekt Nordwind und Sonne, das die Grundlage für meine Forschung in vorliegender Arbeit darstellt.

Zu guter Letzt sage ich danke an all meine lieben Freundinnen im In- wie im Ausland, die diese intensive Zeit mit mir durchlebt haben und mir immer zur Seite standen: Franziska, Katharina, Sophia, Laura, Rosa, Lisa R. und Lisa M. Außerdem danke ich zahlreichen hilfsbereiten Studienkolleginnen und -kollegen, sowie vielen anderen Menschen, die mich auf diesem Weg begleitet haben.

Andrea Klingenbrunner

Inhaltsverzeichnis

1. Einleitung	1
2. Stand der Forschung.....	3
2.1 Die Rolle der Aussprache beim Fremdsprachenlernen	3
2.1.1 Interlanguage und fremdsprachliche Akzente	3
2.1.2 Argumente für eine L1-nahe Aussprache als Unterrichtsziel	7
2.2 Phonologische Varietäten des Deutschen als Quellsprache	9
2.2.1 Dialekt-Standard-Kontinuum im deutschen Sprachraum	9
2.2.2 Varietäten und Dialekte.....	11
2.2.2.1 Mittel- und Südbairisch	11
2.2.2.2 Niederdeutsch	11
2.2.2.3 Hochalemannisch	12
2.2.3 Dialekt und Fremdsprachenunterricht	12
2.3 Französisch und Spanisch als Zielsprachen	13
2.3.1 Französische Aussprachenormen.....	14
2.3.2 Spanische Aussprachenormen	15
2.4 Deutsche Akzente im Französischen und Spanischen.....	15
2.4.1 Französische Aussprache deutschsprachiger Lernerinnen und Lerner	16
2.4.2 Spanische Aussprache deutschsprachiger Lernerinnen und Lerner	18
2.5 <i>r</i> -Laute	18
2.5.1 <i>r</i> -Laute in den Sprachen der Welt	18
2.5.2 <i>r</i> -Laute im Deutschen.....	21
2.5.3 Phonotaktische Distribution.....	25
2.5.4 <i>r</i> -Laute im Französischen.....	25
2.5.4.1 Phonotaktische Distribution	27
2.5.4.2 Mögliche Transfers: Deutsch-Französisch.....	27
Exkurs: Aussprache der französischen <i>r</i> -Laute in Lehrbüchern.....	27

2.5.5	<i>r</i> -Laute im Spanischen	28
2.5.5.1	Phonotaktische Distribution	30
2.5.5.2	Mögliche Transfers: Deutsch-Spanisch	30
	Exkurs: Aussprache der spanischen <i>r</i> -Laute in Lehrbüchern	30
3.	Methode	32
3.1	Forschungsfrage und Hypothesen.....	32
3.2	Methode	32
3.2.1	Untersuchungsprotokoll	32
3.2.2	Korpus	34
3.2.3	Analyse	36
3.2.3.1	Codierungs-System	36
3.2.3.2	Statistische Auswertung	49
3.2.4	Vor- und Nachteile der Methode	53
4.	Ergebnisse und Diskussion	55
4.1	Deutsch	56
4.1.1	Vortestungen.....	56
4.1.2	Onset	57
4.1.2.1	verzweigt	58
4.1.2.2	unverzweigt	59
4.1.3	Coda	59
4.1.3.1	verzweigt	60
4.1.3.2	unverzweigt	60
4.2	Französisch	61
4.2.1	Vortestungen.....	61
4.2.2	Onset	62
4.2.2.1	verzweigt	63
4.2.2.2	unverzweigt	63
4.2.3	Coda	63

4.2.3.1	verzweigt	64
4.2.3.2	unverzweigt	64
4.3	Spanisch.....	65
4.3.1	Vortestungen.....	65
4.3.2	Onset	67
4.3.2.1	verzweigt	68
4.3.2.2	unverzweigt	68
4.3.3	Coda	69
4.3.3.1	verzweigt	69
4.3.3.2	unverzweigt	69
5.	Diskussion.....	70
5.1	Deutsch	70
5.2	Französisch	75
5.3	Spanisch.....	80
6.	Schluss	86
	Bibliographie.....	88
	Anhang	96
	Résumé du travail présent en langue française	96
	Zusammenfassung in deutscher Sprache	106
	Abstract in English.....	108
	Eidesstattliche Erklärung.....	110

Abbildungs- und Tabellenverzeichnis

Abb. 1: *r*-Modell nach Lindau (1985)

Abb. 2: *r*-Artikulation im Deutschen nach Göschel (1971)

Abb. 3: HH_02_fr: Der uvulare Vibrant [ʀ] silbeninitial

Abb. 4: [r] in der französischen Aufnahme zwischen zwei Vokalen

Abb. 5: [r] in sowohl postkonsonantischer, als auch intervokalischer Position

Abb. 6: [χ] vor stimmlosem Konsonanten [k] in *qui* [ki]

Abb. 7: [ʁ] im Onset der Silbe, Stimmhaftigkeit am *voice bar* zu erkennen.

Abb. 8: Approximant [ʁ] zwischen zwei Vokalen, es ist deutliche Stimmhaftigkeit erkennbar

Abb. 9: [r] erkennbar, jedoch deutet fehlender *voice bar* auf entstimmten Charakter hin.

Abb. 10: [ɾ] in der letzten Silbe im Onset

Abb. 11: Entstimmter retroflexer Flap [ɾ̥], erkennbar am fehlenden *voice bar*.

Abb. 12: [ɐ] ersetzt konsonantisches *r* in Form eines Monophthongs.

Abb. 13: [ɐ] ersetzt konsonantisches *r* in Form eines Diphthongs

Abb. 14: *r* fällt weg, der davorstehende Vokal [a] wird in die Länge gezogen.

Abb. 15: Schwach erkennbarer Laut [χ] zwischen Vokal [ɔ] und Konsonant [s].

Abb. 16: Aspiration zwischen Tap [ɾ] und Vokal [i], evtl. Einfluss von Plosivkonsonant [p] wortinitial

Abb. 17: Ausschnitt aus Excel-Tabelle der Auswertung

Abb. 18: Variable berechnen in SPSS

Abb. 19: Schwa [ɐ] im deutschen Korpus

Abb. 20: Vokalverlängerung bei *r*-Wegfall im deutschen Korpus

Abb. 21: Informantin HH_03 realisiert verlängerten Vokal [ɔ] bei *r*-Wegfall

Abb. 22: Stimmhafter alveolarer Tap [ɾ] nach Ort

Abb. 23: Häufigste konsonantische Realisierungen im deutschen Korpus nach Ort

Abb. 24: Zwei Informantinnen und deren prozentuelle Realisierungen von [ɐ] im deutschen und französischen Korpus

Abb. 25: Zwei Informantinnen und deren prozentuelle Realisierungen von [ɐ] im deutschen und französischen Korpus

Abb. 26: Auszug aus *r*-Realisierungen im Vergleich

Abb. 27: Schwa [ɐ] im französischen Korpus in Prozent

Abb. 28: Gegenüberstellender Vergleich von [χ] und [ʁ]

Abb. 29: Häufigste konsonantische Realisierungen im französischen Korpus nach Ort

Abb. 30: Häufigkeit stimmhafter alveolarer Tap [r] nach Ort in Prozent

Abb. 31: Vokalverlängerung nach Ort

Abb. 32: HH_03_es realisiert *norte* [no:tɛ]

Abb. 33: Schwa [ɐ] nach Ort

Abb. 34: Häufigste konsonantische Realisierungen im spanischen Korpus nach Ort

Tabelle 1: Finale Auflistung der analysierten Realisierungen von *r*

Tabelle 2: Codes für Variablen

Tabelle 3: Gesamtzahl Okkurrenzen Deutsch

Tabelle 4: Gesamtzahl Okkurrenzen Französisch

Tabelle 5: Gesamtzahl Okkurrenzen Spanisch

1. Einleitung

Neben den klassischen Sparten Grammatik oder Vokabular ist einer der zentralen Aspekte des Erlernens einer Fremdsprache jener einer möglichst authentischen, dem geläufigen Standard der Zielsprache ähnelnden Aussprache. Diese hilft uns in vielerlei Hinsicht bei Kommunikationssituationen in der Zielsprache, und das aus verschiedensten Gründen. Die Art, in der wir sprechen, vermittelt unserem Gegenüber eine Vielzahl an Informationen über uns. Ein gewisser Akzent kann in diesem Zusammenhang schnell verräterisch wirken. Um diese Art der entlarvenden Aussprache so gut wie möglich abzulegen, bedarf es Instruktion und Übung. Doch inwieweit ist es überhaupt realistisch, eine ‚akzentfreie‘, dem Standard möglichst ähnliche Aussprache in der Zielsprache zu erlangen? Gibt es Faktoren, die das Erlernen einer Aussprache, die jener eines *native speakers* gleicht, fördern beziehungsweise behindern?

Relativ einig ist man sich mittlerweile bei der Tatsache, dass die Erstsprache des Lernalters oder der Lernerin eine zentrale Rolle im Lernprozess spielt. Hier gibt es zwar zahlreiche Diskussionen darüber, woher Akzente und Interferenzen rühren, dass es sie gibt und die Konsequenzen mit sich bringen, ist jedoch unumstritten. Wie verhält es sich aber, wenn der Standard besagter Erstsprache oder Quellsprache, und vor allem deren Aussprache, nicht der Realität der gesprochenen Sprache in den Klassenzimmern entspricht? Gerade das Deutsche mit seinen verschiedenen Varietäten und Dialekten sieht sich hier mit der Problematik der unrealistischen Erwartung einer homogenen Quellsprache konfrontiert, auch, oder vor allem, auf der Ebene der Aussprache. Folglich stellt sich die Frage, ob diese dialektale Vielfalt einen wesentlichen Faktor im Lernprozess der Aussprache darstellt. Kann ein Dialektmerkmal einen Vorteil beim Aussprachetraining liefern, ein anderes hingegen nicht? Gibt es Teilaspekte der jeweiligen Dialekte, die für Transfer und Interferenzen verantwortlich sind? Lassen sich diese tatsächlich auf die dialektale Vielfalt zurückführen? Diese Fragen versuche ich im Zuge dieser Forschungsarbeit anhand der Analyse von Sprachaufnahmen deutschsprachiger Informantinnen in Österreich, Deutschland und der Schweiz zu beantworten. Die betreffenden Zielsprachen Französisch und Spanisch werde ich in nicht allzu ferner Zukunft an einer Schule im deutschsprachigen Raum unterrichten, womit auch das eigene Interesse an einer Studie dieser Art gegeben ist.

Um die Analyse einzugrenzen und zu konkretisieren, ist die Variable, an welcher sich die Untersuchung orientiert, die Gruppe der *r*-Laute. Die Auswahl jener Lautklasse geschah nicht rein zufällig. Diese ist aufgrund ihrer gleichzeitigen Einheit in der Schriftsprache und Vielfalt in der Aussprache, nicht nur, aber auch innerhalb des deutschen Sprachgebrauchs, von besonderem Interesse. Und auch deren restriktivere Realisierungsstandards in den Zielsprachen Französisch und Spanisch stellen einen zentralen Aspekt dar. Sowohl das Französische als auch das Spanische verwenden Realisierungen der *r*-Laute, die nicht für alle deutschsprachigen Lernerinnen und Lerner in gleichem Maße einfach zu erlernen oder imitieren sind. Hier lässt sich nun der Dialekt als Schlüsselement vermuten, denn im Deutschen ist ein Großteil der verschiedenen *r*-Laute als Allophone gebräuchlich. Das Französische und das Spanische jedoch grenzen die standardmäßigen Realisierungen stärker ein und haben somit auch eine größere Fehleranfälligkeit bei Lernerinnen und Lernern. Denn zumindest in der Standardaussprache, welche im Fremdsprachenunterricht ja meist den Fokus darstellt, kann nicht einfach eine so große Zahl an Realisierungsmöglichkeiten als Allophon gebraucht werden wie dies im Deutschen der Fall ist. Im Zuge der vorliegenden Arbeit sollen die Dialekte als Schlüsselement im Fremdsprachenlernprozess anhand der Variablen der *r*-Realisierungen untersucht werden. Dazu sollen die erwähnten Sprachaufnahmen aus verschiedenen deutschsprachigen Orten als Basis für eine korpusphonologische Analyse dienen und in weiterer Folge Aufschluss darüber geben, ob die dialektale Vielfalt im deutschen Sprachraum tatsächlich eine signifikante Rolle spielt, was das Erlernen fremdsprachiger Aussprachemerkmale betrifft, und so Teil einer Grundlage für eine dialektensensiblere Gestaltung des Fremdsprachenunterrichts sein. In diesem Sinne soll ein Zitat Zehetners (1983) aus seiner Forschung zu einem angrenzenden Thema eine einleitende Funktion in vorliegende Forschungsarbeit erfüllen:

„Das tatsächliche von Haus aus bei den Schülern vorfindliche Deutsch variiert je nach geographischer und sozialer Herkunft; häufiger ist es dialektal geprägt. Kontrastive Ansätze, die diese Tatsache ignorieren, gehen an der Wirklichkeit der Schule vorbei.“
(Zehetner 1983: 153)

2. Stand der Forschung

2.1 Die Rolle der Aussprache beim Fremdsprachenlernen

Zum erfolgreichen Erlernen einer Fremdsprache zählen viele einzelne Faktoren, darunter auch jene einer möglichst authentischen Aussprache. Jener wichtige Teilaspekt des Sprachenlernens gliedert sich des Weiteren in segmentale und suprasegmentale Elemente der Aussprache, also die Aussprache einzelner Laute sowie Betonung, Pausen und Geschwindigkeit, respektive (vgl. Carbó et al. 2003: 165). Für die vorliegende Forschungsarbeit soll die segmentale Ebene, also die Aussprache einzelner Laute, im Zentrum stehen. Wie eingangs schon erwähnt, ist die Aussprache ein nicht unwichtiger Teil des Erlernens einer Fremdsprache, denn sie verrät eventuellen Gesprächspartnerinnen und -partnern einiges über uns. Das Gegenüber in der Gesprächssituation leitet aus einem etwaigen Akzent beispielsweise Informationen über unsere Herkunft oder unser Niveau in der Fremdsprache ab, macht sich ein Bild von uns und zieht Schlüsse über Dinge wie Sympathie, Intelligenz oder sozialen Status (vgl. Edwards 1999: 2). Im folgenden Kapitel sollen der Aspekt der Aussprache im Kontext des Fremdsprachenlernens beleuchtet, und verschiedene Argumente für dessen Wichtigkeit aufgeführt werden.

2.1.1 Interlanguage und fremdsprachliche Akzente

Lange Zeit war man im Forschungsgebiet des Sprachenlernens der Ansicht, je später eine Person mit dem Erlernen der Fremdsprache beginne, desto unwahrscheinlicher sei es, dass er oder sie die Aussprache auf dem Niveau eines *native speakers* erreiche. Ein zentraler Begriff dieser Argumentationsweise ist jener der kritischen Phase (vgl. Lenneberg 1967). Diese soll einen Schlüsselmoment im Lernprozess beschreiben, nach dessen Verstreichen der Lerner oder die Lernerin in einer mehr oder weniger stark durch einen Akzent geprägten Form der Fremdsprache stehen bleibt. Lenneberg (1967) zieht Parallelen zum Tierreich und vergleicht die kritische Phase etwa mit der Prägung bei Vögeln und spricht von einer „age limited emergence of behavior“ (Lenneberg 1967: 175). Des Weiteren definiert er die Pubertät als jenen Zeitraum, in dem die Fähigkeit, die Erstsprache erfolgreich zu erlernen, sowie sich weitere Sprachen akzentfrei aneignen zu können, zu einem Ende kommen. Dies wird mit Bezug auf die sich verändernde Plastizität des menschlichen Gehirns argumentiert (vgl. Lenneberg 1967: 175-176) und basiert auf der Forschung des kanadischen Neurochirurgen Penfield (1963), eines der ersten, der den „je früher desto besser“-

Ansatz vorantrieb (vgl. Scovel 2000: 214). Folgt man der Hypothese der kritischen Phase (*critical period hypothesis* oder *CPH*), so ergibt sich daraus eine weit verbreitete Denkweise nach der zuvor genannten Devise „je früher, desto besser“ in puncto Zweit- oder Fremdsprachenlernen, die auch die Schul- und Unterrichtsgestaltung bereits stark geprägt hat (vgl. Scovel 2000: 214). Diese Betrachtung des Lernprozesses wurde bereits von zahlreichen weiteren Forscherinnen und Forschern, wie beispielsweise Wolfe (1967) aufgegriffen. Er argumentiert die Unterschiede beim Lernen so, dass Kinder Aussprache unbewusst, Erwachsene jedoch bewusst lernen, jedoch spezifiziert er nicht, was unbewusstes Lernen bedeutet und inwieweit sich dieses auf den Lernprozess überhaupt auswirkt (vgl. Scovel 2000: 214).

Andere Ansätze, wie jener von Geschwind (1966), führen die Diskrepanz zwischen Erwachsenen und Kindern auf die Art der Umstände, wie sie Sprache erlernen, zurück. Ihm zufolge lernen Kinder durch visuelle oder haptische Eindrücke in Verbindung mit akustischem Input, Erwachsene eher durch Übersetzen aus der L1, wie beispielsweise in Kursen (vgl. Scovel 1969: 246). Jedoch berücksichtigt dieser Ansatz nicht, dass Kinder auch in schulischem Umfeld Lernerfolge erzielen können (vgl. Scovel 1969: 246). Es lässt sich also durchaus erkennen, dass die Annahme einer zeitlichen Begrenzung für ein erfolgreiches Erlernen einer akzentfreien Aussprache in der Fremdsprache in der Forschung lange vorherrschte und auch nach wie vor präsent ist, jedoch viele Argumentationsansätze nicht vollständig zufriedenstellend zu sein scheinen.

An die Forschung rund um Lennebergs kritische Phase (1967) knüpften bereits zahlreiche weitere Untersuchungen an, die diese auch zum Teil zu widerlegen oder ergänzen versuchten. So definiert beispielsweise Selinker (1972) die Notion der Interimssprache oder *interlanguage*, einer eigenen Kategorie der Aussprache eines Lernalters oder einer Lernerin, die sich von der Aussprache eines *native speakers* auf verschiedenen Ebenen unterscheidet und eine Art Zwischenstadium zwischen Quell- und Zielsprache darstellt (vgl. Selinker 1972: 214). Auch hier geht man davon aus, dass nur ein geringer Prozentsatz aller Lernerinnen und Lerner Erfolg beim Überwinden jener von Interferenzen geprägten Interimssprache hat (vgl. Selinker 1972: 212), die meisten bleiben an einem Punkt in diesem Stadium dauerhaft „stecken“, man spricht dann von Fossilisierung (vgl. Selinker/Gass 1992: 197).

Die Frage, was eigentlich einen Akzent konstituiert, stellte einige Jahre später auch James Emil Flege (1981). In Bezug auf eine Studie zur englischen Aussprache mit aus Saudi-Arabien stammenden Personen in den USA (vgl. Flege 1980: 130) bemerkt er eine Tendenz der Lernenden, sich mit ähnlichen Lauten aus ihrer Erstsprache in der Fremdsprache zu behelfen, auch wenn diese nicht ident mit jenen des Englischen sind. Außerdem bemerkt er sehr wohl eine Verbesserung der Aussprache bei jenen Menschen, die schon länger im Land der Zielsprache leben. So stellt er sich die Frage, ob die von *native speakers* des Englischen als Akzent wahrgenommene Aussprache zustande kommt, indem versucht wird, nicht den exakten, bekannten Laut aus der Quellsprache zu produzieren, sondern beim Versuch der Produktion des Lautes in der Zielsprache stattdessen eine ähnliche Aussprache produziert wird (vgl. Flege 1981: 444).

Diese Feststellung deckt sich auch mit dem Konzept des positiven und negativen Transfers im Lernprozess. Hierbei geht man davon aus, dass dort, wo sich L1 und L2 ähneln oder gar gleichen, eine dem Lernen zuträgliche Tendenz erkennbar ist, dort wo die beiden sich jedoch unterscheiden, eine hemmende Komponente zu Tage tritt und Fehler passieren (vgl. Gómez Martínez 2011: 1). Transfer kann auf allen sprachlichen Ebenen geschehen, seien dies Syntax, Lexikon, Morphologie, oder gar auf kultureller Ebene, sowie die für vorliegende Arbeit den Fokus darstellende Phonologie (vgl. Major 2001: 1-2). Bereits Weinreich (1953) beobachtet Transfer auf segmentaler sowie suprasegmentaler Ebene und definierte sieben verschiedene Szenarien für dessen Vorkommen. Auf segmentaler Ebene sind dies zu geringe Differenzierung (*Underdifferentiation*) lautlicher Unterschiede, welche geschieht, wenn die L2 distinktive Merkmale besitzt, die die L1 nicht kennt, zu großzügige Differenzierung (*Overdifferentiation*), wenn das Gegenteil der Fall ist, inadäquate Interpretation distinktiver Merkmale (*Reinterpretation of Distinctions*), also die Interpretation eines Merkmals nach dem Schema in der L1, auch wenn sich dieses nicht mit jenem der L2 deckt, sowie der Austausch von Lauten (*Sound Substitution*), also die Tendenz bei der Aussprache eines unbekannten Lautes der L2 den nächstverwandten Laut aus der L1 heranzuziehen. Auf phonotaktischer Ebene stellt er ebenfalls Interferenzmöglichkeiten fest, indem das Lautsystem der L1 eine Anpassung der Silbenstruktur der L2 zur Folge hat (vgl. Weinreich [1953] 2010: 18-19).

Flege stellt in weiterer Folge seiner Erörterungen (1981) die kritische Phase in Frage, da er anmerkt, dass es sowohl Kinder mit bereits erkennbarem Akzent, sowie Erwachsene ohne jeglichen Akzent in der Fremdsprache gibt (vgl. Flege 1981: 448). Er schlägt eine Alternative zu Lennebergs (1967) kritischer Phase vor und widerspricht seiner strikten biologischen Erklärung. Stattdessen zitiert er Catford (1965) mit dem Begriff der *phonological translation* als Hauptgrund für akzentgeprägte Sprache, also die Fähigkeit auf der Basis einer bereits bekannten Sprache, der L1, aus dieser den nächstgelegenen Laut heranzuziehen, um einen Laut in der Zielsprache zu ersetzen (vgl. Catford 1965: 65), so wie es auch Weinreich (1953) beschreibt. Flege (1981) stellt die Hypothese jener *phonological translation* als Alternative zur kritischen Phase vor, und sieht deren Vorkommen sowohl bei Kindern als auch Erwachsenen (vgl. Flege 1981: 452). Major (2001) merkt außerdem an, dass es einerseits Transfer aus der L1 in die L2 gibt, jedoch auch universelle Elemente des Sprachenlernens für Fehler verantwortlich sein können, die keiner Struktur in der L1, sondern vielmehr allgemeinen Prozessen im Laufe des Erlernens der Sprache, er nennt diese *universals*, zugrunde liegen (vgl. Major 2001: 3).

Flege (1987) erkennt, dass es sehr wohl Differenzen zwischen Kindern und Erwachsenen im Bereich des Erlernens der Aussprache gibt, zweifelt jedoch am Konzept der kritischen Phase aufgrund mangelnder Beweise auf neurologischer Ebene (vgl. Flege 1987: 165). Er kritisiert den zu vereinfachten Blickwinkel auf den Sprachenlernprozess, den die Hypothese der kritischen Phase voraussetzt, da Sprache laut ihm nicht nur in einen rein biologischen Kontext eingebettet sein sollte (vgl. Flege 1987: 167). In einer Studie von Moyer (1999) (zitiert von Major 2001: 9) kam man zu dem Schluss, dass das Alter der Lernerin oder des Lerners nicht als unabhängiger Faktor im Prozess des Erlernens der fremdsprachigen Aussprache mitspielt, sondern mit anderen Elementen, wie beispielsweise Motivation, kultureller Empathie, dem Wunsch, wie ein *native speaker* zu klingen, oder der Art des Unterrichts zusammenhängt.

In einer etwas jüngeren Studie zu anglophonen Französischlernerinnen und -lernern konstatiert Birdsong (2007), dass eine von *native speakers* als akzentfrei wahrgenommene Aussprache, oder *nativelike pronunciation* (vgl. Birdsong 2007: 100), möglich ist, vor allem bei jenen Probandinnen und Probanden, die Aussprachetraining genossen hatten und hohe Motivation aufweisen konnten (vgl. Birdsong 2007: 110).

So lässt sich also sagen, dass es, der bisherigen Forschung zufolge, durchaus möglich ist, den verräterischen Akzent abzulegen. Dieser Meinung sind auch Forscherinnen und Forscher in ähnlichen Studien, wie einer vielzitierten von Bongaerts, Planken und Schils (1995) mit niederländischen Probanden und Probandinnen, die Englisch lernen, oder von Bongaerts (1999), welche niederländische Lernerinnen und Lerner des Französischen beobachtet. In beiden dieser Studien gab es unter den Teilnehmenden eine kleine Zahl solcher, deren Aussprache als „akzentfrei“ gewertet wurde (Studien zitiert von Scovel 2000: 217). Also findet man in einer gewissen Zahl an Publikationen immer wieder Probandinnen oder Probanden, denen dies gelingt, jedoch scheinen diese in der Unterzahl zu sein.

Kurz: eine akzentfreie Aussprache in der Zielsprache zu erlangen ist möglich, der Weg zu diesem Ziel ist aber durch Hindernisse wie Interferenzen erschwert und das Erreichen jenes Objekts ist nicht für alle Lernerinnen und Lerner gleichermaßen realisierbar. Gleichzeitig gilt es aber, Hypothesen wie jene der kritischen Phase von Lenneberg (1967) kritisch zu hinterfragen und den breiteren Kontext und die vielfältigen Aspekte des Erlernens von Sprache zu beachten.

2.1.2 Argumente für eine L1-nahe Aussprache als Unterrichtsziel

Beim Erlernen einer neuen Sprache gibt es eine Vielzahl an Komponenten, deren erfolgreiches Erlernen Ziel der Lernerinnen und Lerner ist, unter anderem sind dies beispielsweise Grammatik oder Wortschatz, aber auch die möglichst L1-nahe Aussprache der Zielsprache. Die vorangegangene Erläuterung verschiedener Ansätze zum Erlernen der Aussprache in der Zielsprache lässt erahnen, dass dieses Kapitel im Lernprozess als nicht unwichtig angesehen werden sollte, obwohl den bisherigen Ergebnissen zufolge das Erlernen einer akzentfreien Aussprache ein nur wenigen Lernerinnen und Lernern vorbehaltenes Erfolgserlebnis bleibt. So ist die Frage berechtigt, ob deren Instruktion im Unterricht überhaupt Sinn macht. Welche Argumente gibt es tatsächlich, die für eine akzentfreie oder der Zielsprache möglichst nahe Aussprache als eines der zentralen Lernziele sprechen?

Ein grundlegender Antrieb für das Erlernen einer authentischen Aussprache ist die des Verstandenwerdens in Gesprächen mit *native speakers* der jeweiligen Zielsprache (vgl. Hurtado/Estrada 2010: 74). Um Teil einer Gesellschaft sein zu können, ist es unabdingbar, von ihr verstanden und anerkannt zu werden, die Möglichkeit des Verstandenwerdens ist also ein Schlüsselement zu Partizipation am

gesellschaftlichen Diskurs (vgl. Moyer 2013: 1). Doch dies stellt nur einen von mehreren ebenso zentralen Gründen dar. Sprache ist traditionell ein Faktor, der Information über unseren sozialen Status mitträgt. Ob bewusst oder unbewusst, unser Gegenüber im Gespräch nimmt unweigerlich Merkmale unserer Sprechart wahr. Daraus werden in weiterer Folge Schlüsse über beispielsweise Sympathie, Intelligenz oder sozialen Status, sowie unsere Kompetenz in der Sprache an sich gezogen (vgl. Edwards 1999: 102). Denn selbst ungeschulte Menschen sind sehr gut darin, Akzente in ihrer Erstsprache zu erkennen, auch dann, wenn dies nur relativ kurze Äußerungen sind (vgl. Flege 1984: 694). Der Aspekt der Verständlichkeit wird auch im fünften Kapitel „Phonologische Kompetenz“ des Gemeinsamen Europäischen Referenzrahmen für Sprachen (fortan GERS) behandelt, nach welchem sich der Fremdsprachenunterricht länderübergreifend orientiert.

Bis zum Niveau B1 wird der Fokus auf eines der zuvor genannten Argumente, jenes der Verständlichkeit, gelegt und bedient sich ebenfalls des Terminus *Akzent*. So lautet beispielsweise die Anforderung für phonologische Kompetenz auf Niveau A2: „Die Aussprache ist im Allgemeinen klar genug, um trotz eines merklichen Akzents verstanden zu werden; manchmal wird aber der Gesprächspartner um Wiederholung bitten müssen.“. Auf dem Niveau B1 lautet die Beschreibung wie folgt: „Die Aussprache ist gut verständlich, auch wenn ein fremder Akzent teilweise offensichtlich ist und manchmal etwas falsch ausgesprochen wird.“ (GERS Kapitel 5). Ab Niveau B2 spricht die Beschreibung des GERS von einer „klaren und natürlichen Aussprache und Intonation“ und der Akzentbegriff scheint nicht mehr auf (GERS Kapitel 5). Da für die Untersuchung in vorliegender Arbeit Informantinnen mit etwa zwischen A2 und B1 angesiedelten Kompetenzen in den beiden Fremdsprachen herangezogen wurden, wäre also laut GERS mit einem gewissen Akzent zu rechnen.

Ein weiterer Grund für das Erlernen der Zielsprache ähnlichen Aussprache, der zwar für die vorliegende Arbeit keinen zentralen Aspekt darstellt, jedoch trotzdem Erwähnung finden sollte, ist jener der Motivation bei Lernerinnen und Lernern. Eine Studie zu Motivationsfaktoren beim Fremdsprachenlernen zeigt, dass das Erlangen einer authentischen Aussprache bei Lernenden eines der am höchsten gereihten Ziele ist (vgl. Harlow/Muyskens 1994: 145). Somit erweist sich der Erwerb der einem *native speaker* gleichenden Aussprache aus vielerlei Gründen als erstrebenswert und als zentraler Aspekt des allgemeinen Prozess des Erlernens einer Fremdsprache und

sollte im Unterricht nicht ignoriert werden, obwohl zahlreiche Forscherinnen und Forscher an der Hypothese der kritischen Phase festhalten.

2.2 Phonologische Varietäten des Deutschen als Quellsprache

In diesem Kapitel soll nun die bereits in den vorangegangenen Kapiteln erwähnte Wichtigkeit der Erstsprache ausführlicher behandelt werden. Im Zuge dieser Forschungsarbeit spielen das Deutsche und seine phonologischen Varietäten als L1 die zentrale Rolle. Das Deutsche ist in mehreren europäischen Staaten die offizielle oder eine von mehreren offiziellen Sprachen (vgl. Krech 2009: 1). So wie jede Sprache besitzt auch das Deutsche bestimmte Charakteristika, die seinen Sprecherinnen und Sprechern Probleme sowie Vorteile beim Erlernen weiterer Sprachen bereiten. Diese können überall, beispielsweise in der Syntax oder in der Aneignung des Lexikons auftreten, in der vorliegenden Forschungsarbeit interessieren uns jedoch die Auffälligkeiten beim Erlernen der Aussprache. Diese, wie schon eingangs besprochen, überschattet in Konversationen mit *native speakers* nämlich oftmals bereits sehr hohe Kompetenz in Grammatik und Ausdrucksweise. Im folgenden Kapitel sollen nun das Deutsche und seine dialektale Phonologie als L1 näher beleuchtet werden, sowie die Standards und deren Entwicklung über die vergangenen Jahrzehnte innerhalb eines Dialekt-Standard-Kontinuums im deutschen Sprachraum diskutiert werden.

2.2.1 Dialekt-Standard-Kontinuum im deutschen Sprachraum

Das Definieren sprachlicher Normen ist nichts Ungewöhnliches. So kennt vermutlich der Großteil der deutschsprachigen Welt den Duden als Rechtschreibwörterbuch. Somit ist wahrscheinlich auch der Begriff der Orthographie vielen geläufig. Da aber für die vorliegende Arbeit vorrangig die Regeln und Konventionen der Aussprache im Zentrum stehen, legen wir das Augenmerk auf die Rechtlautung oder Orthoepie, also die ‚korrekte‘ Aussprache des Deutschen. Auch für diese gibt und gab es Versuche einer Normierung. Einer der bekanntesten Vorreiter der deutschen Aussprachenormierung ist Theodor Siebs, welcher im Jahr 1898 den finalen Anstoß zur Niederschrift von Ausspracheregeln für Schauspieler gab (vgl. Ehrlich 2009: 26). Die erste Auflage des Werkes wurde unter dem Titel *Deutsche Bühnenaussprache: Nach den Ergebnissen der Beratungen zur ausgleichenden Regelung der deutschen Bühnenaussprache, die vom 14. bis 16. April 1898 im Apollosaal des Königlichen Schauspielhauses zu Berlin stattgefunden haben* (Albert Ahn Verlag, Berlin, Köln, Leipzig) publiziert.

Nach und nach etablierte sich die Siebs'sche Norm der Bühnenaussprache auch in öffentlichen Domänen wie dem Bildungsbereich. Jedoch wurden die Normen nach Siebs (1898) auch bald als realitätsfremd kritisiert, was kaum verwundert, da sie auf deutliche Aussprache für die Bühne und ein großes Publikum ausgerichtet waren, und nicht auf alltägliche Kommunikationssituationen (vgl. Ehrlich 2009: 27-28). So wurde das Werk nach und nach überarbeitet und beispielsweise in der 19. Auflage 1969 angepasst und um von der Bühnenaussprache getrennte Kategorien wie ‚reine Hochsprache‘ und ‚Alltagslautung‘ erweitert (vgl. Ehrlich 2009: 29). Der Wunsch nach einer möglichst allgemeinen, mundartfreien Aussprachenorm zog sich auch über das 20. Jahrhundert weiter. So fand eine Entwicklung einer neuen Standardaussprache in den 1960er Jahren statt, die sich an einer relativ neuen Gruppe von Vorbildern orientierte. Die neue Norm war geprägt durch die Verbreitung moderner Massenmedien wie Fernsehen. Diese Innovation trug die dort gesprochene Sprache in selbst die entlegensten Gegenden und die Sprecherinnen und Sprecher in jenen neuen Medien nahmen die Vorbildrolle ein (vgl. Ehrlich 2009: 30). Es entstanden Aussprachewörterbücher, die möglichst für den gesamten deutschen Sprachraum Anwendung finden sollten, wie beispielsweise das 1964 erschienene *Wörterbuch der deutschen Aussprache* (Krech).

Durch die an Bedeutung gewinnende Rolle der Plurizentrik des Deutschen wurden neuere Referenzwerke nach und nach den verschiedenen Varietäten des Deutschen gerecht und es entstanden Aussprachewörterbücher für einen deutschen, schweizerischen und österreichischen Standard (vgl. Krech 2009: 2). Jedoch beobachteten Dialektforscherinnen und -forscher nach wie vor eine Einstellung den verschiedenen deutschen Varietäten gegenüber, die das Deutsche, wie es in Deutschland gesprochen wird als Standard, und das österreichische Deutsch beispielsweise als Substandard dessen betrachtet (vgl. Ransmayr 2012: 200). Standardvarietäten des Deutschen wie beispielsweise das österreichische Deutsch oder das Schweizerhochdeutsch werden oftmals mit Dialekten gleichgesetzt (vgl. Ransmayr 2012: 200). Bei großer dialektaler Vielfalt sowie verschiedensten Unterschieden in Register und Kontext gestaltet sich eine klare Abgrenzung zwischen Dialekt, Varietät und Standard schwierig. So bietet sich die Betrachtung des deutschen Sprachgebrauchs als Kontinuum mit den Polen Standardsprache und Mundart an (vgl. Durrell 1998: 22). Merkmale der Aussprache der für vorliegende Forschungsarbeit relevanten Regionen werden im weiteren Verlauf diskutiert werden, sowohl varietäten- als auch dialektspezifische Auffälligkeiten sollen näher betrachtet werden.

2.2.2 Varietäten und Dialekte

Nachdem wir nun festgestellt haben, dass das Deutsche und dessen Aussprache aufgrund seiner Vielfalt nicht der Realität im Klassenzimmer oder Kursraum in Form einer homogenen Quellsprache entsprechen kann, empfiehlt es sich, einen genaueren Blick auf die Varietäten und Dialekte zu richten, insbesondere jene, die im Zuge der vorliegenden Arbeit erforscht wurden. Die drei Varietäten, die sich in weiterer Folge in die verschiedenen Dialekte der Regionen aufgliedern, die anschließend behandelt werden sollen, sind das österreichische Deutsch, das Schweizerdeutsche und das Bundedeutsche (vgl. Muhr 1997: 49). Diese Varietäten gelten jeweils als der akzeptierte Standard in ihrem Sprachgebiet. Da für diese Arbeit jedoch nicht die Varietäten, sondern dialektale Unterschiede den Fokus darstellen, sollen diese im folgenden Abschnitt kurz näher beschrieben werden. Die Besonderheiten der Aussprache der jeweiligen Regionen finden später in den Kapiteln 2.2.3, sowie 2.4 und 2.5 ihren Platz und werden unter Anwendung verschiedener Blickwinkel untersucht. Die diskutierten Aussprachemerkmale finden sich oftmals zwischen dialektaler und varietätenabhängiger Lautung, was nicht zuletzt den Grund hat, dass, wie zuvor erwähnt, sich das Deutsche am besten mithilfe des Begriffs eines Kontinuums darstellen lässt. Selbstverständlich kennt der deutsche Sprachraum noch eine viel größere Zahl an Dialekten als jene, die in dieser Forschungsarbeit untersucht werden. Die auf sechs Orte begrenzte Anzahl von drei Dialektzonen, die für die vorliegende Arbeit relevant sind, beinhalten das Mittel- bis Südbairische, das Niederdeutsche und das Alemannische.

2.2.2.1 Mittel- und Südbairisch

Wirft man einen Blick auf die Grenzen der Dialektzonen, sieht man schnell, dass diese meist nicht mit Staats oder gar Bundeslandgrenzen kongruent sind. So erstreckt sich beispielsweise die südmittelbairische Dialektzone von Nieder- und Oberbayern in Deutschland über große Teile Österreichs, wie beispielsweise Oberösterreich, Niederösterreich, Wien und Teile der Steiermark (vgl. König 2014: 230-231).

2.2.2.2 Niederdeutsch

Das niederdeutsche Dialektgebiet erstreckt sich vor allem im nördlichen Teil Deutschlands, aber auch über die niederländische Grenze hinweg und gliedert sich in verschiedene Mundarten auf (vgl. König 2014: 230-231). Die Informantinnen die für die vorliegende Arbeit herangezogen wurden, stammen aus Hamburg.

2.2.2.3 Hochalemannisch

Die alemannischen Dialekte dehnen sich ebenfalls über eine Grenze hinweg aus. So findet man sie im Westen Österreichs, zum Beispiel in Vorarlberg, in Liechtenstein, dem Elsass, sowie in der deutschsprachigen Schweiz (vgl. König 2014: 230-231). Sie unterscheiden sich von anderen deutschen Dialekten auf den Ebenen der Morphologie, der Syntax und, hier relevant, der Aussprache. Das Alemannische wird in verschiedene Unterkategorien unterteilt. Der in Zürich gesprochene Dialekt, der für die vorliegende Arbeit von Interesse ist, gehört zum Hochalemannischen (vgl. Projekt SynAlm der Universität Konstanz).

2.2.3 Dialekt und Fremdsprachenunterricht

Mit der Rolle der dialektalen Vielfalt in Verbindung mit Fremdsprachenunterricht haben sich bereits viele Forscherinnen und Forscher über die vergangenen Jahrzehnte beschäftigt. Wie schon in den vorangegangenen Kapiteln erwähnt, ergibt sich die Wichtigkeit der Dialektforschung im Zusammenhang mit dem Erlernen von Fremdsprachen dadurch, dass das Deutsche mit einer einheitlichen Aussprache als L1 im Klassenzimmer nicht der Realität entspricht (vgl. Zehetner 1983: 153). So gibt es beispielsweise schon zahlreiche Studien, die den Zusammenhang zwischen deutschen Dialekten und dem Englischen und seiner Aussprache erforscht haben, wie beispielsweise Zehetner (1983), welcher Bairisch sprechende Schüler und deren Schwierigkeiten des Erlernens der englischen Phonologie dokumentierte. Er entwarf für seine Beobachtungen der bairisch gefärbten Aussprache im Englischen den Begriff des *Bavarian English* (vgl. Zehetner 1983: 153), welches sich dadurch auszeichnet, dass die Schüler¹ Laute der Zielsprache durch ihnen ähnliche in der L1, also in diesem Fall Bairisch, ersetzen und ihnen die nun akzentgeprägte Aussprache sogar natürlicher vorkommt, als die „korrekte“, da sie sich in ihrer gewohnten Lautumgebung bewegen (vgl. Zehetner 1983: 156).

Diese Tendenz ist nicht nur in Bayern aufgefallen, sondern scheint in ähnlichen Studien in verschiedensten Gebieten des deutschsprachigen Raums aufzutreten. So bemerken Kettemann und Viereck (1983) in ihrer Studie mit Grazer Schülern² ebenfalls ein „Zurückgreifen“ auf bekannte Laute bei der Aussprache des Englischen, mit für das Steirische typischen Auffälligkeiten, wie beispielsweise stark aspirierte Fortisvarianten

¹ Im Artikel von Zehetner (1983) wird nicht spezifiziert, ob es sich um ausschließlich männliche Teilnehmer oder eine gemischte Informantengruppe handelt, daher hier nur die männliche Form.

² Im Artikel von Viereck (1983) ebenso.

im Silbenauslaut, wie der Kontrast engl. *egg* [ɛg] im Vergleich mit der Aussprache des steirischen Lerner [ɛkʰ] (vgl. Kettemann/Viereck 1983: 80).

Auch mit alemannischen Dialekten als L1 beschäftigte sich die Forschung bereits. So untersucht Schubinger (1983 [1937]) das von ihr so bezeichnete *Alemanic English* auf phonologischer Ebene und stellt Interferenzen an verschiedensten Stellen fest, wie beispielsweise die Entstimmung von stimmhaften Plosiven und Frikativen, die zu Interferenzen wie der Nichtunterscheidung des englischen Minimalpaars *lagging* [læɡɪŋ] ‚verzögernd‘ versus *lacking* [lækɪŋ] ‚fehlend‘ führen, da Sprecherinnen und Sprecher des Alemannischen [g] entstimmt realisieren [ɡ̊] und somit [k] angleichen (vgl. Schubinger 1983: 25). Anhand dieser Beispiele lässt sich schon gut eine gewisse Tendenz erkennen, die den Dialekt als ausschlaggebendes Element über den Standard der Erstsprache stellt, wenn es um den Erwerb der Aussprache in einer Fremdsprache geht. Was jedoch auch in den meisten dieser Werke betont wird, ist, dass diese Auffälligkeiten im Unterricht weitgehend unbeachtet bleiben. In weiterer Folge sollen in vorliegender Arbeit noch weitere solche Beispiele vorkommen, die Interferenzen durch dialektale Merkmale des Deutschen näher beschreiben.

2.3 Französisch und Spanisch als Zielsprachen

Die beiden Zielsprachen, die im Zuge dieser Arbeit erforscht werden, sind das Französische und das Spanische. Beide Sprachen werden in sowohl Deutschland, Österreich als auch der Schweiz in Schulen sowie anderen Institutionen unterrichtet und sind beliebte Fächer, wobei Französisch bereits als Klassiker in den meisten Schulen gilt und Spanisch sich während der letzten Jahre zum Newcomer hocharbeitete, der sich nun einer stetig wachsenden Beliebtheit erfreut. In Deutschland lernen beispielsweise 24% der Schülerinnen und Schüler Französisch und 4% Spanisch (Statista), in Österreich sind dies bei Französisch etwa gleich viele, bei Spanisch momentan nur 2% (Statistik Austria), und in der Schweiz ist Französisch eine von vier Nationalsprachen und somit im Bildungsbereich fest verankert, und Spanisch befindet sich mittlerweile auf Platz 3 der beliebtesten Fremdsprachen (vgl. Elminger/Forster 2005: 28).

Selbstverständlich sind auch das Französische und Spanische plurizentrische Sprachen mit vielen Varietäten und dialektalen Unterschieden. Da es in der vorliegenden Arbeit aber vorrangig um das Erlernen dieser Fremdsprachen geht, soll in weiterer Folge der Standard und dessen Lautung als Ziel herangezogen werden, da sich

Fremdsprachenunterricht vorrangig an diesem orientiert. Einige Aspekte, die Französisch und Spanisch ebenfalls zu plurizentrischen Sprachen machen, werden in den folgenden Kapiteln dennoch ihren Platz finden, da sie zwar in der Schule meist nicht als Lernziel unterrichtet werden, aber deshalb nicht weniger korrekt oder real vorkommend sind.

2.3.1 Französische Aussprachenormen

Das Französische besitzt eine lange Tradition der Sprachnormierung, denn das *français standard* oder *français de référence* wird bereits seit dem 16./17. Jahrhundert systematisch erfasst (vgl. Fagyal 2006: 2). So überrascht es kaum, dass dieses Bestreben zum Teil auch die Aussprache betrifft. Einrichtungen wie die *Académie Française* sind maßgeblich an der Sprachnormierung in den verschiedenen Aspekten der Sprache beteiligt. Das Wörterbuch der *Académie Française*, dessen erste Ausgabe 1694 erschien, hat als Ziel, die Französische Sprache mit Regeln auszustatten, welche es zu befolgen gilt. Aus den Statuten der *Académie* aus dem Gründungsjahr 1635 geht folgendes hervor:

„La principale fonction de l’Académie sera de travailler, avec tout le soin et toute la diligence possibles, à donner des règles certaines à notre langue et à la rendre pure, éloquente et capable de traiter les arts et les sciences.“ (Académie Française : Article XXIV des statuts)

Die *Académie* pflegt also einen präskriptiven Zugang zur Verwendung von Sprache. Wobei jedoch zur Thematik einer dem Standard entsprechenden Aussprache zu bemerken ist, dass bereits im Vorwort des ersten Wörterbuches aus dem 17. Jahrhundert angemerkt wurde, dass die Anschrift der korrekten Aussprache ähnlich viel Sinn mache, als würde man versuchen, einem Blinden, das Konzept von Farben zu erklären (vgl. Rey 2011: 150). Jedoch wird sehr wohl eingeräumt, dass Angaben die Aussprache betreffend dann sinnvoll sind, wenn die lautliche Realisierung weit von ihrer orthographischen Schreibweise abweicht (vgl. Rey 2011: 151). Somit finden sich auch heute in den Wörterbüchern der *Académie* Angaben zur Aussprache, wie Hinweise auf ein *h aspiré* oder bei Einträgen, deren Aussprache nicht auf Anhieb klar ist, wie beispielsweise bei Fremdwörtern (vgl. Dictionnaire de l’Académie Française en ligne). Was die allgemeinen Konventionen der französischen Aussprache angeht, so findet man zahlreiche Referenzwerke, die die Phonologie des Französischen ausführlich und

meist aus deskriptiver Perspektive behandeln, wie beispielsweise spezifisches Lehrmaterial zur französischen Phonologie und Phonetik.

2.3.2 Spanische Aussprachenormen

Der Standard von Rechtschreibung sowie –lautung, oder Orthoepie, wird auch in Spanien stark von einer normierenden Institution beeinflusst, der *Real Academia Española* (RAE) (vgl. Carbó et al. 2003: 164). Auch die RAE gibt Wörterbücher heraus, in welchen die Definition der Orthoepie wie folgt lautet: „(el) arte de pronunciar correctamente“, (Real Academia Española: 2001) also die Kunst des korrekten Aussprechens. Doch auch die normierende RAE gesteht dem Spanischen eine große lautliche Vielfalt zu, die der geografischen Verbreitung der Sprache zuzuschreiben ist, spricht aber auch über eine *base común* und die Wichtigkeit einer gewissen Normierung mit dem Ziel der erfolgreichen Kommunikation, die sich ohne Normen schwierig gestalten würde (vgl. Diccionario panhispánico de dudas: 2005).

2.4 Deutsche Akzente im Französischen und Spanischen

Untersuchungen zur Interphonologie beim Erlernen einer Fremdsprache gibt es bereits zahlreiche. So finden sich beispielsweise einige Studien zur Aussprache englischsprachiger Französischlerner und –lernerinnen (vgl. Birsdson 2007), oder auch englischsprachiger Spanischlerner und –lernerinnen wie jene von Hurtado/Estrada (2010) oder Olsen (2012), welche sich sogar spezifisch mit dem Erlernen der *rhotics*, also der *r*-Laute im Spanischen beschäftigen, die uns in einem der Folgekapitel noch begegnen werden. Geht man einen Schritt weiter, so findet man beispielsweise auch eine Untersuchung, die die Quellsprache in ihrer dialektalen Vielfalt betrachtet, indem sie den Norden und den Süden Englands und dessen dialektale Phonologie im Kontext des Erlernens von Spanisch als Zielsprache beleuchtet (vgl. Herrero de Haro/Andión Herrero 2012). Diese Studie stellte bereits unterschiedliche Aussprachemuster je nach Dialekt fest (vgl. Herrero de Haro/Andión Herrero 2012: 64). Auch Major (2001) erwähnt bereits in seinem einleitenden Kapitel dialektale Diskrepanzen im Aussprachelernprozess bei *native speakers* des Spanischen beim im Englischen vorkommenden [ʃ] wie in engl. *shoe* [ʃu:] ‚Schuh‘, da jener Laut im Standardspanischen nicht vorkommt, in einem Dialekt im mexikanischen Chihuahua jedoch schon (vgl. Major 2001: 3).

Es ist also ersichtlich, dass auch Dialekte als L1 in der bisherigen Forschung von Interesse und Untersuchungen darüber demnach legitim sind. Auch Studien zu

Sprecherinnen und Sprechern deutscher Dialekte und deren Aussprache des Englischen wurden im Laufe vorliegender Arbeit bereits besprochen. Alle dieser Studien gehen von Transfer von der Quell- zur Zielsprache aus, einige davon mit einer Entwicklung zu weniger Interferenzen, je höher das Kompetenzniveau wird (vgl. Olsen 2012: 65) (vgl. Rose 2010: 407), beziehungsweise je länger sie im Land der Zielsprache leben, wie in den Untersuchungen von Birdsong (2007) mit englischsprachigen Menschen, die in Frankreich leben und als Erwachsene Französisch als Fremdsprache lernen (vgl. Birdsong 2007: 114).

Interphonologie ist also etwas Sprachenübergreifendes und wird in verschiedenen Konstellationen von Quell- und Zielsprache beobachtet und erforscht, auch dann, wenn die Quellsprache ein Dialekt ist. Interphonologie und Interferenzen sind aber auch durchaus überwindbar, wie einige Resultate der zuvor genannten Studien zeigen. Um nun wieder zum Thema der Aussprache deutschsprachiger Lernerinnen und Lerner der Sprachen Spanisch und Französisch zurückzukehren, sollen hier einige Eigenheiten in der Aussprache des gesamten deutschen Sprachraums sowie solche der einzelnen hier relevanten Dialektzonen beispielhaft angeführt werden, und wie diese dazu beitragen können, das Erlernen der Aussprache des Spanischen und Französischen zu beeinflussen, beziehungsweise wie diese zu Transfer (sei dieser positiv oder negativ) führen können.

2.4.1 Französische Aussprache deutschsprachiger Lernerinnen und Lerner

Wie eingangs erwähnt, gibt es beim Erlernen der Aussprache in der Fremdsprache unterschiedliche Fehlerquellen. In Kapitel 2.2.3 konnten wir bereits Beispiele dialektal geprägter Interferenzen im Englischen beobachten. Besagte Interferenzen können Aussprachekonventionen des deutschen Sprachraumes als Ganzes zum Ursprung haben, also mit Phonemen, die im Deutschen schlichtweg nicht existieren, oder auf der Lautung eines der vielen Dialekte beruhen. Dieser Abschnitt soll beispielhaft ein paar dieser Transferquellen in Bezug auf die französische Standardaussprache beleuchten. Zuvor soll ein Beispiel für eine mögliche Interferenz aufgrund des deutschen Standards kontrastiv zu Transfer aus den Dialekten angeführt werden. In Bezug auf den gesamten deutschsprachigen Raum haben Lernerinnen und Lerner oft Probleme mit dem im Französischen vorkommenden stimmhaften Konsonanten [ʒ] wie beispielsweise in frz. *rouge* [ʁuʒ] ‚rot‘, da dieser Laut im Deutschen fast nur in (oftmals aus dem Französischen stammenden) Fremd- und Lehnwörtern, wie dt. *Garage*

[gara:fe], zu finden ist (vgl. ADABA – österreichische Aussprachedatenbank) und dort oft seine Stimmhaftigkeit verliert (vgl. Russ 2010: 158) und entstimmt [ʒ] eher [ʃ] wie in dt. *Schuh* [ʃu:] ähnelt. Dies steht im Gegensatz zur französischen Standardaussprache, die den stimmhaften Konsonanten in frz. *garage* [gaʁa:ʒ] vorsieht, da /ʒ/ und /ʃ/ im Französischen Phonemstatus besitzen (vgl. Pustka 2011: 100). Für Sprecherinnen und Sprecher eines bairischen Dialekts (in Süddeutschland wie in Österreich) ist der stimmhafte alveolare Frikativ [z] wie in dt. *Sonne* [zɔnə] im Anlaut ungewöhnlich (vgl. Ehrlich 2009: 92). Dies äußert sich in der Fremdsprache, hier als Beispiel Französisch, oft in Form einer Interferenz, wenn standardmäßig das stimmhafte [z] im Anlaut verlangt wird, wie beispielsweise in frz. *zéro* [zeʁo] ‚null‘, indem Schülerinnen und Schüler den Unterschied zwischen den beiden im Französischen als Phoneme geltenden Lauten /z/ und /s/ (vgl. Pustka 2011: 100) das stimmlose [s] verwenden und dadurch [seʁo] realisieren.

Diese zuletzt genannten Interferenzen lassen sich also bereits auf die regionale Beschaffenheit der L1 zurückführen, und nicht allein auf das Deutsche mit einheitlicher Lautung. Wie wir auch in weiterer Folge in Kapitel 2.5 genauer beleuchten werden, produzieren Sprecherinnen und Sprecher eines alemannischen Dialekts tendenziell auch die *r*-laute in der Coda, während der restliche deutschsprachige Raum auf vokalisierte Formen wie das Lehrer-Schwa [ɐ] zurückgreifen, beziehungsweise den *r*-Laut wegfallen lässt und den Vokal davor verlängern (vgl. Krech 2009: 244). Auch diese Tatsache bereitet eine Grundlage für mögliche Interferenzen, da Lernerinnen und Lerner diese Realisierungsmöglichkeit aus ihrer L1 heranziehen und anstatt eines *r*-Lautes in der Coda wie frz. *fort* [fɔ:r] den Konsonanten meiden und [fɔɐ] aussprechen. Hierfür wird es in Teil III vorliegender Arbeit Beispiele aus dem Korpus geben. Doch auch schweizerische Sprecherinnen und Sprecher des Deutschen haben in der Aussprache des Französischen nicht nur Vorteile. Kolly (2013) beschreibt einige Eigenheiten, die sogar zur Schöpfung des Helvetismus *français fédéral*, also ein durch einen deutschschweizerischen Akzent geprägtes Französisch, führten (vgl. Kolly 2013: 51). Zu diesen Auffälligkeiten zählt Kolly (2013) unter anderem die oft stimmlose Aussprache von [b] in frz. *bien* [bjɛ̃] ‚gut‘ und [ʒ] in frz. *je* [ʒə] ‚ich‘ als [pjɛ̃] und [jə] oder dialektal verursachte Unterschiede in der Vokalqualität, wie beispielsweise deren Verlängerung oder Vokale, die zu offen, tief oder weniger gespannt realisiert werden. Des Weiteren beobachtet sie fallweise eine alveolare Artikulation der *r*-Laute (vgl. Kolly 2013: 52).

2.4.2 Spanische Aussprache deutschsprachiger Lernerinnen und Lerner

In 2.4.1 wurde soeben noch besprochen, dass das auf Französisch vorkommende stimmhafte [z] im Anlaut für Sprecherinnen und Sprecher eines bairischen Dialekts problematisch sein kann. Gegensätzlich dazu begegnen wir norddeutschen Lernerinnen und Lernern, bei denen es zu Interferenzen kommt, wenn es um das stimmlose [s] im Anlaut auf Spanisch geht. So wird aus span. *sol* [sol] oft [zol]. Die Aspiration von Plosiven im Anlaut, beispielsweise in dt. *kalt* [k^halt] entlarvt norddeutsche Lernerinnen und Lerner der Fremdsprache schnell als solche, da diese im Norden Deutschlands gängige Form der Aussprache die Aspiration vorsieht, ebendiese aber in Österreich, sowie im deutschen Süden, nur selten realisiert wird (vgl. Wonka 2015: 75) und in romanischen Sprachen typischerweise auch geringer ausfällt (vgl. Flege 1981: 445). Wie auch schon im vorangegangenen Kapitel zu möglichen Interferenzen deutschsprachiger Lernerinnen und Lerner im Französischen erwähnt, kann die gewohnte Realisierung eines Coda-Schwa [ə] anstatt eines konsonantischen *r*-Lautes auch im Spanischen zu Transfer führen, sodass aus beispielsweise span. *fuerte* [fwerte] ‚stark‘ bei einer deutschsprachigen Lernerin oder einem deutschsprachigen Lerner aus einem nicht-alemannischsprachigen Gebiet [fwæte] wird.

2.5 *r*-Laute

Für dialektale Besonderheiten und deren Verursachen von Transfer in den Fremdsprachen ließen sich noch zahlreiche weitere Beispiele anführen. Im vorangegangenen Kapitel finden sich bereits einige Fälle dialektal geprägter Interferenzen in der Aussprache von Lernerinnen und Lernern. Da es sich jedoch für eine möglichst detaillierte Analyse auf einen enger gefassten Aspekt zu fokussieren gilt, sollen in der vorliegenden Forschungsarbeit daher die *r*-Laute als Dreh- und Angelpunkt im Hinblick auf Transfer im Mittelpunkt stehen. Die Gründe für das vorrangige Interesse an dieser Lautgruppe sollen im folgenden Kapitel beleuchtet werden.

2.5.1 *r*-Laute in den Sprachen der Welt

Die *r*-Laute als Gruppe sind schon seit langem für die sprachwissenschaftliche Forschung interessant und wurden in verschiedensten Zusammenhängen thematisiert. Sie gehören mitunter zu den für Kinder im Phoneminventar ihrer Erstsprache am schwersten zu erlernenden Lauten und sind umgekehrt auch oft die ersten, die bei Aphasie verloren gehen (vgl. Göschel 1971: 85). Die Frage, die sich aber im folgenden

Kapitel stellt, ist jene nach der Vielfalt der *r*-Laute und gleichzeitige Kategorisierung in eine Klasse unter dem Graphem <r>. Es haben sich schon einige Forscherinnen und Forscher, unter anderem Wiese (2001) und Lindau (1980), mit der Kategorisierung dieser Gruppe von Lauten beschäftigt, die sie als verschieden und zugleich einheitlich beschreiben:

[...] while r-sounds seem to vary greatly and in several dimensions, there is yet a unity to these r-sounds which justifies the common usage in phonetics and linguistics to talk about „r“ and the class of „r-sounds“ or „rhotics“. (Wiese 2001: 11)

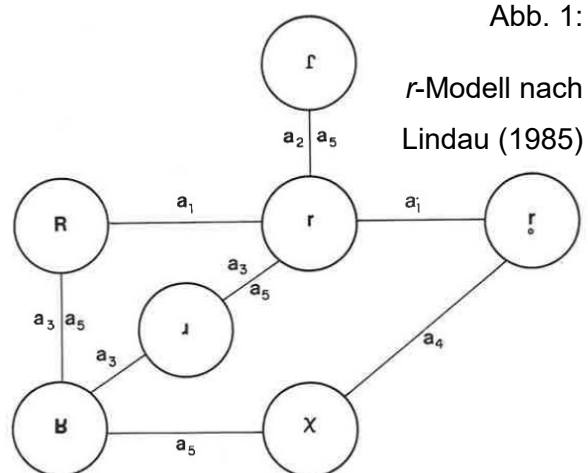
Jenes Zitat von Wiese (2001) beschreibt in Kürze das, womit sich die Forschung schon lange so eingehend beschäftigt. Während <r> als Graphem in der Schriftlichkeit weltweit in verschiedensten Sprachen als ein und dasselbe auftritt, sind die Möglichkeiten seiner lautlichen Realisierung so vielfältig, dass sich die Frage nach Berechtigung der Klassifizierung jener Laute als eine einzelne Gruppe stellt. Auch zu Gründen für jene Gruppierung gibt es zahlreiche Theorien. Die *r*-Laute sind außerdem in den Sprachen der Welt stark verbreitet. Etwa 75% aller Sprachen kennen eine oder mehrere Formen von *r*-Lauten (vgl. Maddieson 1984: 73). Unter dem Graphem <r> sind verschiedenste phonologische Realisierungsmöglichkeiten zusammengefasst. Diese unterscheiden sich in verschiedenen Sprachen, aber auch innerhalb derselben Sprache oftmals relativ stark durch Kriterien wie den verschiedenartigen Ort ihrer Artikulation. So ist es möglich, dass <r> in einer Sprache, wie beispielsweise Spanisch alveolar durch einen stimmhaften Tap [r] wie in span. *pero* [pɛro] ‚aber‘, oder Trill [r] wie in span. *perro* [pɛro] ‚Hund‘ realisiert wird (vgl. Martínez-Celdran/Fernández-Planas/Carrera-Sabaté 2003: 255) und beide den Status eines Phonems innehaben, während das gleiche Graphem von Sprecherinnen und –sprechern des Standard American English als Approximant [ɹ] wie in engl. *rye* [ɹaɪ] ‚Roggen‘ ausgesprochen wird (vgl. International Phonetic Association ¹⁵2014: 41) und im Standard British English durch das so genannte *r-dropping* (vgl. Biersack 2002: 61) an manchen Stellen in der Aussprache elidiert wird, wo es schriftlich nach wie vor zu finden ist, beispielsweise in engl. *car* [ka:] ‚Auto‘.

Doch auch *r*-Laute, die beispielsweise uvular realisiert werden, werden oft unter selbigem Graphem zusammengefasst. So wurden in vorangegangenen Kapiteln bereits die Aussprachemöglichkeiten der französischen *r*-Laute diskutiert, welches den stimmhaften uvularen Frikativ [ʁ] der Standardaussprache zuordnet (vgl. International Phonetic Association ¹⁵2014: 79). Das Phänomen der *r*-Laute beschränkt sich aber nicht nur auf

die Familie der indoeuropäischen Sprachen. Wirft man einen Blick auf die zur afro-asiatischen Sprachfamilie gehörenden Sprache Hausa, die in Westafrika gesprochen wird, so finden sich unter dem Graphem <r> ebenfalls zwei Phoneme, zum einen der stimmhafte alveolare Trill [r] wie in Hausa *bara* [bə̀r̀à] ‚betteln‘, sowie der stimmhafte retroflexe Tap [ɽ] wie in Hausa *bara* [bə̀ɽ̀à] ‚Diener‘ (vgl. International Phonetic Association 152014: 90).

So ist also ersichtlich, dass weder der Artikulationsort noch die Artikulationsart das abgrenzende Merkmal jener Lautgruppe darstellen. Ein auffälliges Merkmal ist die Okkurrenz von *r*-Lauten an ähnlichen Positionen in den Silben verschiedenster Sprachen. So beobachtet Lindau (1985) ihr Vorkommen in Konsonantencluster nahe des Silbenkerns, welcher meist ein Vokal ist. Abgesehen davon beschreibt sie ihre Elision zugunsten eines Vokals oder deren Einfluss auf einen Vokal in ihrer Umgebung (vgl. Lindau 1985: 157). Wiese (2001) vermutet die Legitimation für die Gruppierung der *r*-Lauten als Klasse aufgrund von phonotaktischen Aspekten, indem er ihnen einen eigenen *slot* in der Sonoritätshierarchie, zwischen // und Halbvokalen oder hohen Vokalen zuweist (vgl. Wiese 2001: 13). Lindau (1985) hat, um der Frage nach Gründen für diese Klassifizierung nachzugehen, die *r*-Lauten in vier indoeuropäischen sowie sieben in Westafrika gesprochenen Sprachen untersucht, um deren Gruppierung auf den Grund zu gehen, und kam zu dem Schluss, dass eine Klassifizierung in einer Form von Verwandtschaftsmodell sinnvoll scheint, in dem immer ein Mitglied einem nächsten ähnelt, aber nicht jedes Mitglied allen (vgl. Lindau 1985: 166).

Das Modell von Lindau (1985) in Abbildung 1 beschreibt übersichtlich, wie sich die verschiedenen *r*-Lauten zueinander verhalten. Mitglied *r*1 ähnelt *r*2, welches wiederum *r*3 ähnelt, welches dann wiederum *r*4 ähnelt. So wird schrittweise die Beziehung zwischen den *r*-Lauten hergestellt. Ladefoged und Maddieson (1996) knüpfen in ihrer Abhandlung über die Klasse der *rhotics* an dieses Modell an und schließen mit der Bemerkung ab, dass eben jene familiär anmutenden Beziehungen die Klassifizierung der *r*-Lauten in eine Gruppe legitimieren und, dass die



umfassende Einheitlichkeit in der Schriftsprache auf hauptsächlich historischen Gründen basiert (vgl. Ladefoged/Maddieson 1996: 245). Des Weiteren fassen sie zusammen, dass die Klasse der *r*-Laute zwar vielfältig und variantenreich ist, jedoch alveolare Realisierungsformen wie der Trill [r] und der Tap [ɾ] in den Sprachen der Welt am häufigsten vorkommen (vgl. Ladefoged/Maddieson 1996: 218).

2.5.2 *r*-Laute im Deutschen

Schon bei Siebs (1898) findet man ein eigenes Kapitel, das die ‚Sonoranten‘ behandelt, darunter ein Unterkapitel zu *r*. Er konstatiert,

„dass sich seit dem 17. Jahrhundert neben dem alten deutschen Zungen-*r* das Zäpchen-*r* immer weiter verbreitet (hat), so daß heute beide Formen in der Hochsprache als gleichberechtigt angesehen werden müssen.“ (Siebs [1898] ¹⁸1961: 61)

Als Anmerkung nach jenem Zitat findet man den Verweis auf die Entscheidung des Beraterausschusses 1933, der die Gleichberechtigung der *r*-Laute beschlossen hat (Siebs [1898] ¹⁸1961: 61), also scheint jene Beobachtung des Status der beiden Laute als Allophone nichts allzu Neues zu sein. Die Siebs'sche Norm schreibt des Weiteren vor, dass „*r* vor *t* nicht zum Reibelaut *ch* werden (darf)“, und dass es „im Auslaut nicht zu *a* vokalisiert werden (darf)“ (Siebs [1898] ¹⁸1961: 61-62). Es ist hier anzumerken, dass Siebs (1898) sich nicht auf Dialekte bezieht, sondern, wie schon in 2.2.1 besprochen, auf eine Norm für den gesamten deutschen Sprachraum, noch spezifischer, Normen für die Bühnenaussprache von Schauspielerinnen und Schauspielern. Bekanntlich gibt es in der Realität im deutschen Sprachraum zahlreiche Abweichungen von den Siebs'schen Aussprachenormen, die für die jeweilige Region als richtig angesehen werden (vgl. Ammon et al. 2004: LI). Da die *r*-Laute allein in verschiedenen Sprachen schon so eine große Vielfalt aufweisen, ist es daher auch wenig verwunderlich, dass dies auf regionaler und dialektaler Ebene gleichermaßen der Fall ist. Vielen deutschen Dialekten wird daher jeweils ein anderer *r*-Laut (oder mehrere gleichzeitig) zugeordnet. Die generelle Entwicklung des uvularen *r* im Deutschen zum *r* der Standardaussprache und dessen historischer Hintergrund sind umstritten. Einige sprachwissenschaftliche Historiker und Historikerinnen gehen von einer Verbreitung des Lautes im 17. Jahrhundert aus, als sich das *r* des Pariser Französisch, also der uvulare Frikativ [ʁ], aus Prestige Gründen vor allem in urbanen deutschsprachigen Gegenden durchsetzte und davor ein alveolarer Trill die gängige Aussprache war (vgl. Howell 1987: 318). Dieser Annahme wird aber des Öfteren entgegengesetzt, dass eine uvulare oder

velare Realisierungsmöglichkeit schon vor diesem Zeitpunkt in älteren germanischen Sprachen existiert hatte (vgl. Howell 1987: 319).

Somit sehen wir, dass die Vielfalt der deutschen *r*-Laute kein modernes Phänomen zu sein scheint, jedoch die Entwicklung des als Standard beschriebenen *r*-Lautes eine ist, die viele Deutschsprecherinnen und –sprecher zu Lebzeiten miterlebt haben, als sich während des 20. Jahrhunderts die Konvention von [r] zu einer uvularen Artikulation [ʀ]³ verlagerte (vgl. Wiese 2001: 14). In der Literatur zur deutschen Standardaussprache finden sich verschiedene Beschreibungen des im Deutschen geltenden Standards, verschiedene Aussprachewörterbücher sowie Forschungsartikel und Referenzwerke zur deutschen Phonologie bedienen sich unterschiedlicher Formen der Anschrift und bemerken unterschiedliche Verteilungs- und Häufigkeitsmuster. So beschreibt das Deutsche Aussprachewörterbuch (Krech 2009) ohne Umschweife die gängigsten Aussprachemöglichkeiten nach Varietäten geordnet und legt sich nicht auf einen allgemeingültigen Standard fest. Wiese (2001) sieht den uvularen Approximanten [ʁ] als die meistverbreitete Realisierung an, verweist aber zugleich auf die große dialektale Vielfalt (vgl. Wiese 2001: 14).

Auf der Ebene der Varietäten gibt das Deutsche Aussprachewörterbuch (Krech 2009) zur Distribution der *r*-Laute weiter Aufschluss über Aussprachestandards. Für den in Deutschland geltenden Standard findet man in einer Tabelle zu den Konsonanten nach distinktiven Merkmalen nur [ʁ] als Phonem (vgl. Krech 2009: 29), mit Kommentar dazu, dass das häufig realisierte *r* im deutschen Standard lange Zeit als Vibrant [ʀ] beschrieben wurde, dieser aber laut Krech (2009) nur selten, beispielsweise im Kunstgesang, vorkommt und im Standard am häufigsten der Reibelaut [ʁ] vor akzentuiertem Vokal gesprochen wird (vgl. Krech 2009: 30). Weiters unterteilt sie die Realisierungsmöglichkeiten in konsonantische und vokalische (vgl. Krech 2009: 85), letztere in Form des so genannten ‚Lehrer-Schwa‘ [ɐ] (vgl. Biersack 2002: 66).

Alveolar realisierte *r*-Laute finden zwar Erwähnung, werden jedoch nicht anhand von Beispielen demonstriert, und es wird erneut die Prävalenz des Frikativs [ʁ] betont. Bei der Diskussion der österreichischen Standardaussprache werden sowohl das alveolare „Zungenspitzen-[r]“ (Krech 2009: 244) als auch der uvulare Vibrant [ʀ] als in ganz

³ Die Literatur ist sich in puncto Transkription des deutschen *r* uneins, [ʀ] soll in diesem spezifischen Fall als Repräsentant für verschiedene uvulare Realisierungsmöglichkeiten im Gegensatz zur alveolaren Produktion stehen.

Österreich verbreitet genannt, was ein neueres Phänomen ist. Früher waren die Laute dialektal unterschiedlich verteilt, wobei im Westen in Vorarlberg und im Osten Österreichs die alveolare Realisierung üblich war, und im Süden und in Tirol [ʀ] vorherrschte (vgl. Krech 2009: 244). Den Frikativ [ɣ] stellt Krech (2009) in Österreich nur im Inlaut vor Konsonant und nach Kurzvokalen fest, beispielsweise in *Hirt* [hiʁt]. An der gleichen Stelle nach Langvokalen wird meist zu [ɐ] vokalisiert und *hier* [hi:ɐ] realisiert (vgl. Krech 2009: 244). In der Schweiz beobachtet Krech (2009) den alveolaren Vibrant [r], den uvularen Vibrant [ʀ] sowie den uvularen frikativ [ɣ] als gängige Standardlaute. Sie konstatiert ebenfalls, dass im Dialekt [r] am häufigsten zu finden ist (vgl. Krech 2009: 266).

Wirft man einen weiteren Blick auf die bisherige Forschung zur dialektalen *r*-Lautung im deutschen Sprachgebiet, findet man zwar keine eindeutigen und absoluten Zuordnungen von bestimmten Realisierungen zu bestimmten Dialekten, aber sehr wohl Tendenzen in Richtung einer geografischen Verteilung, wie wir sie auch schon bei Krech (2009) auf Ebene der Varietäten feststellen konnten. In Werken zur dialektalen Lautung findet sich in Kettemann/James (1983) eine Tabelle, die dem deutschen Standard die vier Realisierungsmöglichkeiten [r], [r̥], [ʀ] und [ɣ] zuordnet, während dem Steirischen, welches auch zur bairischen Dialektzone zählt, nur [ʀ], und dem Alemannischen und dem Bairischen [ʀ] und [r̥] zugeschrieben werden, wobei anzumerken ist, dass jene Tabelle als Überblick zu allen im Werk erschienenen Beiträgen zu sehen und daher nicht komplett homogen ist (vgl. Kettemann/James 1983: 15). Nach Simpson (1998) findet man im norddeutschen Sprachgebiet oft den stimmlosen uvularen Frikativ [χ] nach einem initialen Konsonanten wie in *trat* [tχat] (vgl. Simpson 1998: 91). Die wohl umfassendste geografische Zuordnung von *r*-Lauten finden wir in Göschel (1971), dessen Karte (siehe Abbildung 2 unterhalb) verschiedenen Regionen Deutschlands eine Vielzahl an verschiedenen *r*-Lauten zuordnet:

2.5.3 Phonotaktische Distribution

Wie schon weiter oben kurz angesprochen, verhalten sich die *r*-Laute auf Deutsch selbst innerhalb eines Dialektes unterschiedlich, je nachdem an welcher Position in der Silbe sie sich befinden. Die Silbe als Einheit wird unterteilt in Onset, Nukleus und Coda (vgl. Kabatek/Pusch 2009: 68). Im Deutschen kann der Silbenkern aus einem Vokal oder Konsonanten aus der Gruppe der Liquide und der Nasale bestehen (vgl. Müller 2002: 96). Das Deutsche erlaubt in der Silbe ein- bis dreigliedrige Konsonantencuster, sowohl in Onset (vgl. Hall 1992: 65-69) als auch in der Coda (vgl. Hall 1992: 111-116). Die *r*-Laute sind oftmals Teil jener Konsonantenanhäufungen, und man findet sie in mehreren Zusammensetzungen. So trifft man auf die *r*-Laute beispielsweise einzeln im Onset wie in dt. *parat* [para:t], einfach verzweigt wie in dt. *Prinz* [prints] und in dreigliedrigen Onsets, wie in dt. *Sprung* [prʊŋ] (vgl. Hall 1992: 65-69).

Auch in der Coda findet man sie alleine wie in dt. *Narr* [nar], mit einem zweiten Konsonanten verzweigt wie in dt. *Hirn* [hɪrn], oder dreigliedrig, also mit zwei weiteren Konsonanten verzweigt, wie in dt. *entfernt* [ɛntfɛrnt] (vgl. Hall 1992: 111-116). Wie wir aber schon kurz zuvor gesehen haben, werden *r*-Laute in weiten Teilen des deutschsprachigen Raums in der Coda eher zu [ɐ] vokalisiert oder elidiert und nicht als Konsonant realisiert, sodass aus dt. *Hirn* [hɪrn] oft beispielsweise eher [hiɐn] wird. Ganz allgemein lässt sich also sagen, dass das Deutsche und seine Dialekte eine recht breite Palette an *r*-Lauten als Allophone erlauben (vgl. Wiese 2001: 14) und somit die Verständlichkeit der deutschen Aussprache nicht durch die jeweilige Realisierung des *r*-Lautes eingeschränkt wird. Trotz dieser Tatsache ist die Frage nach Distribution der *r*-Laute nach Dialekten von Interesse, denn wenn sie schon nicht für das gegenseitige Verstehen für Deutschsprecherinnen und –sprecher ausschlaggebend ist, so könnte sie sehr wohl einen Einfluss auf diese in den Fremdsprachen ausüben. Jene, die in dieser Forschungsarbeit behandelt werden, besitzen nämlich ein restriktiveres Inventar an standardmäßigen Realisierungen, wie in den folgenden Punkten 2.5.4 und 2.5.5 erläutert werden soll.

2.5.4 *r*-Laute im Französischen

In puncto *r*-Laute kennt das Französische allein eine Vielzahl an Realisierungsmöglichkeiten, die mehr oder weniger in verschiedensten Regionen der Frankophonie zu finden sind. So findet man neben dem stimmhaften uvularen Frikativ [ʁ], der etwa zwi-

schen dem 15. und 16. Jahrhundert begonnen hat, sich von Paris aus in die Frankophonie zu verbreiten, und heute als Standard gilt (vgl. Pustka 2007: 194), zahlreiche andere Möglichkeiten, wie *r* realisiert werden kann. So hört man oft den stimmhaften uvularen trill [ʀ] als Allophon von [ʁ] am Wortanfang (vgl. Pustka 2007: 194), den stimmlosen uvularen Frikativ [χ] als Allophon von [ʁ] in der Nähe eines anderen stimmlosen Konsonanten wie beispielsweise in frz. *présent* (vgl. Tranel 1987: 143) oder gar den alveolaren Trill [r], der heute noch in manchen Regionen Frankreichs, sowie außerhalb des Hexagons, zu hören ist (vgl. Pustka 2007: 194).

Die Distribution der verschiedenen *r*-Laute gestaltet sich jedoch laut Picard (1987) trotz Vielfalt immer gleich, egal in welcher Position sie sich befinden, also erwartet man von einem Sprecher des Standards beispielsweise ein [ʁ] sowohl wortinitial in *renovation* [ʁenovasjõ] ‚Renovierung‘, präkonsonantisch *hambourgeois* [ãbuʁʒwa] ‚hamburgisch‘, postkonsonantisch *matrimonial* [matʁimɔnjal] ‚ehelich‘, intervokalisch *séparation* [sepapasjõ] ‚Trennung‘ und postvokalisch *Hambourg* [ãbuʁ] ‚Hamburg‘ die gleiche Realisierung (vgl. Picard 1987: 52-53). Dies widerspricht der zuvor genannten Erwähnung Tranels (1987), der in Konsonantennähe auch [χ] als Allophon von [ʁ] beobachtet, während Fagyal (2006) wiederum [ʁ] in Konsonantencluster wie *après* [apʁɛ] ‚danach‘ beschreibt und [χ] eher intervokalisch als Ersatz für den uvularen Trill [ʀ] (vgl. Fagyal 2006: 46).

Auch Realisierungen, die durch Assimilationsprozesse zustande kommen, sind beobachtbar (vgl. Fagyal 2006: 46), sodass sich eine große Zahl an verschiedensten *r*-Lauten ergibt, die der Realität vieler Sprecherinnen und Sprecher entspringen. All diese Realisierungsmöglichkeiten als ‚falsch‘ abzutun, entspräche einer stark präskriptiven Auslegung von Sprache. Da sich aber dennoch, wie erwähnt, die vorliegende Forschungsarbeit an den in den Schulen im deutschsprachigen Raum vorherrschenden Normen im Fremdsprachenunterricht orientiert, werden Realisierungsformen wie beispielsweise das alveolare [r] oftmals nicht als ‚richtige‘ Aussprache angesehen, obwohl es in der Realität durchaus vorkommt. Diese Realität soll durch einen kleinen Exkurs mit einem Blick in ein paar gängige Lehrwerke untermauert werden, um zu sehen, auf welche Art und Weise, wenn denn überhaupt, die französische *r*-Lautung im Unterrichtsmaterial behandelt wird.

2.5.4.1 Phonotaktische Distribution

Das Französische kennt sowohl offene als auch geschlossene Silben, wobei erstere den Großteil mit etwa 80% ausmachen (vgl. Lauret 2007: 54). Das Französische erlaubt ebenfalls Konsonantencluster, jedoch findet man insgesamt weniger Fälle von mehrkonsonantischen Silbenelementen (vgl. Léon 2007: 133), oftmals in seltenen Lehnwörtern wie beispielsweise frz. *dextre* [dɛkstʁə] ‚rechte Hand‘ (vgl. Sokol 2007: 92). Im Gegensatz zum Deutschen sind jene laute, die den Silbenkern oder Nukleus bilden können, ausschließlich Vokale (vgl. Pustka 2011: 111).

Die *r*-Laute kommen sowohl in Onset- als auch Codaposition vor und können verzweigt und unverzweigt sein. Ein Beispiel für ein unverzweigtes *r* im Onset ist frz. *rat* [ʁa] ‚Ratte‘, ein verzweigtes *r* im Onset findet man in frz. *drame* [dʁam] ‚Drama‘. Im Französischen existieren ebenfalls Onsets, die aus drei Elementen bestehen, auch hier finden wir *r* als einen Baustein wieder wie in frz. *scribe* [skʁib] ‚Schreiber‘ (vgl. Pustka 2011: 118). Ein unverzweigtes *r* in der Coda lässt sich in frz. *dire* [diʁ] ‚sagen‘ beobachten, verzweigt in frz. *moderne* [mɔdəʁn] ‚modern‘ (letztes Beispiel vgl. Pustka 2011: 118). Auch die Coda lässt noch Kombinationen aus bis zu vier Elementen zu, so findet man in frz. *filtre* [filtʁ] ‚Filter‘ (vgl. Pustka 2011: 118) ebenfalls *r*. Ein Beispiel für eine viergliedrige Coda wurde schon weiter oben in frz. *dextre* erwähnt, auch hier finden wir wieder *r* als Teil ebendieser.

2.5.4.2 Mögliche Transfers: Deutsch-Französisch

Wie schon kurz zuvor in Kapitel 4.2 erwähnt, produzieren Sprecherinnen und Sprecher eines alemannischen Dialekts auf Deutsch tendenziell auch die *r*-Laute in der Coda, während der restliche deutschsprachige Raum auf vokalisierte Formen wie das Lehrer-Schwa [ɐ] zurückgreift, beziehungsweise den *r*-Laut wegfallen lässt und den Vokal davor verlängert. Diese Tatsache bereitet eine Grundlage für mögliche Interferenzen, da Lernerinnen und Lerner diese Realisierungsmöglichkeit aus ihrer Quellsprache heranziehen und anstatt eines *r*-Lautes in der Coda wie frz. *fort* [fɔ:ʁ] den Konsonanten meiden und [fɔɐ] aussprechen, vor allem dann, wenn sie auf Deutsch eher die alveolaren Realisierungsformen von *r* aussprechen und mit uvularen Realisierungen Probleme haben.

Exkurs: Aussprache der französischen *r*-Laute in Lehrbüchern

Da es in weiterer Folge der vorliegenden Forschungsarbeit um die *r*-Laute in sowohl Quell- als auch Zielsprache als Analyseschwerpunkt gehen wird, soll in jeweils einem

Exkurs die Behandlung jener Lautgruppe in Lehrbüchern der Sprache beleuchtet werden.

Wirft man einen Blick in Lehrwerke, die in Schulen oder Französischkursen verwendet werden, stößt man bei der Aussprache der *r*-Laute, sofern ihr überhaupt ein Kapitel oder Abschnitt gewidmet ist, weitgehend auf einen Vergleich des französischen *r* mit dem des Deutschen anhand von Wortbeispielen.

Wenn es im internationalen phonetischen Alphabet transkribiert ist, was selten der Fall ist, so findet man häufig [ʀ], wie beispielsweise im Lehrwerk *Voyages* (vgl. VOYAGES 2013: 12). Im Lehrbuch *À plus! 1* (2007), welches als Lehrwerk für Französisch als dritte Fremdsprache an Gymnasien zu verwenden gedacht ist, findet man vielerlei Hinweise zur Aussprache von Lauten, die es im Französischen zu unterscheiden gilt. Es werden Phoneme wie /z/ und /s/, oder /ʃ/ und /ʒ/ in Gegenüberstellung behandelt, meist auch mit Wort- und Hörbeispielen verdeutlicht (vgl. *À plus! 1* 2007: 27; 43). Zur Aussprache von *r* gibt es keine konkrete Übung oder Darstellung in diesem Lehrbuch, erst im Anhang scheint es in einer Liste der Laute des Französischen in Lautschrift als [ʀ] auf, gefolgt von zwei Wortbeispielen ohne phonetische Transkription, *rue* und *livre* (vgl. *À plus! 1* 2007: 159). Im Lehrbuch *On y va! A2* (2014) findet man eine einzige Hör- und Nachsprechübung zum Laut [ʀ], der an jener Stelle auch als [ʀ] transkribiert erscheint (vgl. *On y va! A2* 2014: 38). Interessant ist hier ebenfalls, dass es sich beim zuletzt genannten Lehrwerk bereits um den zweiten Band handelt, sein Vorgänger für das Niveau A1 (*On y va! A1*: 2014) beinhaltet keinerlei Hinweise zur Aussprache des *r* im Französischen.

2.5.5 *r*-Laute im Spanischen

Wie schon unter Punkt 2.3.2 erwähnt, erhält man Hinweise auf die korrekte Lautung des Spanischen in den Referenzwerken der *Real Academia Española*. So findet man die Konventionen der *RAE* für die Aussprache der *r*-Laute im Spanischen wie folgt vor: Nach Konsonanten wie , <p>, <g>, <k>, <f>, <c>, <d>, <t>, sowie am Silben- oder Wortende ist ein /r/ zu realisieren. Die Anschrift als /r/ ist laut *RAE* eine Vereinfachung für Leserinnen und Leser, denen linguistische Konzepte nicht geläufig sind (*Diccionario en línea de la RAE*). Gemeint ist der stimmhafte alveolare Tap /r/. Für <r> am Wortanfang und nach <n>, <l>, <s>, sowie bei <rr> ist nach *RAE* ein stimmhafter alveolarer Vibrant /r/ auszusprechen, wobei im Wörterbuch wiederholt, aus Gründen der

breiteren Verständlichkeit, eine andere Anschrift, nämlich /rr/ verwendet wird (Diccionario en línea de la RAE). So unterscheiden sich viele Referenzwerke über die spanische Lautung in ihrer Form der Transkription. Quilis (1993) beispielsweise unterscheidet zwischen *vibrante simple* /r/ und *vibrante múltiple* /r̄/ (vgl. Quilis 1993: 330), während Hualde (2005) die Transkription /r/ und /r̄/ des Internationalen phonetischen Alphabets mit Zusatz des diakritischen Zeichens für bessere Unterscheidung von transkribierten Äußerungen verwendet, da er aufgrund der spanischen Orthographie Verwechslungen befürchtet. Er betont, dass die phonetische Transkription für beispielsweise span. *perro* [p̄ero] ‚Hund‘ zu sehr der graphischen Schreibweise von span. *pero* für ‚aber‘ gleicht, und verwendet somit [r̄] für mehr Klarheit (vgl. Hualde 2005: 16).

Die vorliegende Arbeit folgt den aktuellen Konventionen des internationalen phonetischen Alphabets (vgl. International Phonetic Association ¹⁵2014) und verzichtet auf das diakritische Zeichen über [r], um den stimmhaften alveolaren Vibranten zu beschreiben. Die eben genannten beiden Phoneme gelten als der Standard in der spanischen Aussprache (vgl. Hualde 2005: 186). Aber auch in der spanischsprachigen Welt gibt es Variation, was die *r*-Laute betrifft. So findet man in manchen Regionen Fälle von Assibilierung, wie beispielsweise in einigen Dialekten Mexikos, oder einen dem im Englischen Standard verbreiteten [ɹ] ähnlichen Approximanten in Costa Rica sowie Fälle von uvularen Realisierungen, ähnlich dem im Deutschen oft beschriebenen [ʀ], dessen Vorkommen ebenso in Dialekten Costa Ricas beobachtet werden kann (vgl. Hualde 2005: 186-187).

So zeigt sich, dass auch das Spanische in Bezug auf seine *r*-Realisierungen eine große Vielfalt kennt. Jedoch ist im Unterricht, wie auch bei Französisch, die der Norm entsprechende Aussprache geläufig und wird von zahlreichen Referenzwerken sowie den in Schulen verwendeten Lehrwerken als die anzustrebende Ziellautung angesehen. Jemand, der mit der Artikulation der alveolar realisierten *r*-Laute Schwierigkeiten hat und auf seine oder ihre gewohnte Aussprache eines uvularen *r*-Lautes zurückgreift, wird also gemäß den geltenden Standards entsprechend eher als mit Akzent sprechend wahrgenommen werden. Ein kurzer, an das Kapitel über phonotaktische Distribution der *r*-Laute im Spanischen anschließender, Exkurs über die Aussprachehinweise in Lehrbüchern soll auch hier Aufschluss darüber geben, ob und wie Unterrichtsmaterial mit der Instruktion der Standardaussprache umgeht.

2.5.5.1 Phonotaktische Distribution

Das Spanische verfährt bei der Bildung seiner Silben ähnlich wie das Französische. So ist es auch hier der Fall, dass der Silbenkern ausschließlich aus Vokalen bestehen kann (vgl. Hualde 2005: 70). Auch findet man im Spanischen Konsonantencluster, im Onset können diese maximal aus zwei Elementen bestehen, die immer dem Muster der Struktur aus Plosiv oder /f/ gefolgt von /l/ oder /r/ entsprechen, beispielsweise in span. *tres* [tɾɛs] ‚drei‘ (vgl. Hualde 2005: 71). Selbstverständlich findet man auch unverzweigte Onsets im Spanischen, auch in Bezug auf die Gruppe der *r*-Laute. Wenn <r> beispielsweise wortinitial, und somit als Onset, zu finden ist, so ist die Realisierung standardmäßig [r] (vgl. Hualde 2005: 183), wie in span. *rata* [rata] ‚Ratte‘. Unverzweigt in der Coda finden die *r*-Laute oft Verwendung, wie beispielsweise in span. *carta* [karta] ‚Brief‘ oder span. *mar* [mar] ‚Meer‘ (vgl. Hualde 2005: 75). Die Coda erlaubt ebenfalls Cluster, jedoch ist in ihnen kein *r*-Laut erlaubt, sondern ausschließlich /s/ an zweiter Stelle, wie in span. *vals* [vals] ‚Walzer‘.

2.5.5.2 Mögliche Transfers: Deutsch-Spanisch

Wenn wir uns nun der Kategorie der *r*-Laute als Quelle für Transfer und Interferenzen widmen, lassen sich ähnliche Fehlerquellen vermuten wie auf Französisch. Eine Sprecherin, die es gewohnt ist, auf Deutsch ein [ɐ] anstatt eines *r*-Lautes in der Coda zu realisieren und eventuell noch dazu Schwierigkeiten mit der Aussprache des alveolaren Taps [ɾ] und Trills [r] hat, weil sie in der L1 beispielsweise uvular produzierte Laute wie [ʀ] verwendet, könnte span. *fuerte* [fwɛɾtɛ] dann eher wie [fwɛɛtɛ] aussprechen. Eine Spanischlernerin aus der Schweiz, welche einen alemannischen Dialekt spricht, wäre für diese Art von Interferenz aus der Quellsprache demnach weniger anfällig.

Ob sich diese auf dem aktuellen Stand der Forschung beruhenden Annahmen über dialektbasierte Fehlerquellen in der fremdsprachigen Aussprache tatsächlich so manifestieren, soll im weiteren Verlauf dieser Forschungsarbeit nun näher untersucht werden.

Exkurs: Aussprache der spanischen *r*-Laute in Lehrbüchern

Auch auf die *r*-Laute in Lehrbüchern zur spanischen Sprache soll ein Blick in einem kurzen Exkurs geworfen werden. Hier wird ebenfalls eine Auswahl an Lehrwerken darüber informieren, wie diese es mit der Anleitung zur standardmäßigen Aussprache halten. Diese wird oft zu Beginn des Lehrbuches kurz angeschnitten, meist in Form von Wortbeispielen ohne phonetische Transkription oder mit einer Umschreibung, wie

zum Beispiel in einer Übung im Lehrbuch *Caminos neu A1*: „*r* wird am Anfang eines Wortes und wenn es als *rr* geschrieben wird, stark gerollt [...]“. Zur Aussprache des alveolaren Taps [r] gibt es keine Hinweise (*Caminos neu A1* 2009: 11-12). Die Lehrbücher *Descubramos el español – Spanisch interlingual* (2012) und *El nuevo Curso 1* (2001) behandeln beide Phoneme und bedienen sich der Transkription mittels Einzel- und Doppellaut *r* für den alveolaren Tap [r] und *rr* für den Trill [r] (vgl. *Descubramos el español – Spanisch interlingual* 2012: 6), während *El nuevo curso 1* sogar noch einen Schritt weiter geht und den Unterschied zwischen Deutsch und Spanisch in Hinblick auf die Aussprache von *r* am Wortende mit einer Hörverständnisübung hervorhebt (vgl. *El nuevo curso 1* 2001: 37).

3. Methode

3.1 Forschungsfrage und Hypothesen

Die im Stand der Forschung abgehandelten Themen und Erkenntnisse zeigen, dass in Bezug auf das Erlernen der Aussprache in der Fremdsprache der Fokus nicht einzig auf einer standardisierten Normaussprache liegen sollte. Deshalb soll im Forschungsteil der vorliegenden Arbeit eruiert werden, ob dialektale Unterschiede ausschlaggebend für Transfer und Interferenzen auf lautlicher Ebene aus der Quellsprache in die Zielsprache sind. Dazu wurden Informantinnen aus sechs verschiedenen deutschsprachigen Orten in drei Ländern zur Untersuchung ihrer *r*-Lautung herangezogen. Die *r*-Laute wurden ebenfalls im vorangegangenen Teil dieser Arbeit diskutiert und ihre Vielfalt beschrieben, sodass es nahe liegt, dass bei einer solchen Verschiedenheit innerhalb der Lautklasse auch dialektale Aspekte, sowie Besonderheiten, die die fremdsprachige Aussprache betreffen, von Interesse sein könnten.

Die Hypothese lautet folglich, dass durch die dialektale Vielfalt des Deutschen und die daraus resultierende phonetische Varianz für das Erlernen der Aussprache in der Fremdsprache den Ausschlag gibt, nicht aber der deutsche Standard.

Die zweite Hypothese, die in meine Forschung miteinbezogen wird, ist, dass die *r*-Laute aufgrund ihrer Heterogenität als Lautklasse einen nicht zu vernachlässigenden Faktor im Lernprozess spielen, da diese Laute auch in den Zielsprachen verschiedene Realisierungen kennen und somit anfällig für Fehler sind. So führen die bisher behandelten Themen in der Forschung zur Forschungsfrage der vorliegenden Arbeit:

Ist die dialektale, und nicht die standardsprachliche *r*-Lautung in der Aussprache selbiger Lautklasse der Zielsprachen Französisch und Spanisch ausschlaggebend?

3.2 Methode

Dieses Kapitel behandelt das methodische Vorgehen ausgehend von der Planung, beleuchtet die Vorgehensweise meiner Arbeit mit dem Datenmaterial und dessen Auswertung und Testung. Außerdem sollen im Anschluss verschiedene Aspekte meiner Handlungsweise kritisch reflektiert und diskutiert werden, bevor im darauffolgenden Ergebnisteil die Resultate besprochen werden.

3.2.1 Untersuchungsprotokoll

Um deutsche Akzente in den Fremdsprachen Französisch und Spanisch erforschen zu können, bedarf es Tonmaterial. Die Aufnahmen, anhand derer ich meine Analysen

erarbeitet habe stammen aus der Datenbank eines unveröffentlichten Projektes über deutsche Dialekte, deren Phonologie sowie die Ausprägung in den Sprachen Französisch und Spanisch. Im Zuge jenes Projektes wurden in verschiedensten Orten im deutschsprachigen Raum Aufnahmen durchgeführt. Das Projekt wurde von Univ.-Prof. Dr. Elissa Pustka (Universität Wien) und Prof. Dr. Christoph Gabriel (Johannes-Gutenberg-Universität Mainz) ins Leben gerufen. Die Informantinnen nach denen gesucht wurde, sollten in ein bestimmtes Profil passen. Es wird fortan nur mehr von Informantinnen gesprochen, da eine der Anforderungen war, dass nur weibliche Teilnehmerinnen aufgenommen werden sollten. Dies geschah in Hinblick auf die bessere Vergleichbarkeit der Ergebnisse. Des Weiteren beinhaltet das Profil eine Altersspanne von 18-26, sowie die Begebenheit, dass sie im jeweiligen Ort, dessen Dialekt für das Projekt interessant war, geboren wurden und dort wenn möglich so gut wie ihr ganzes bisheriges Leben verbracht haben. Außerdem sollten sie die beiden Fremdsprachen Französisch und Spanisch auf etwa dem Niveau A2-B1 gemäß dem Gemeinsamen Europäischen Referenzrahmen für Sprachen (GERS) beherrschen, wobei das reale Niveau nicht viel höher als das offizielle sein sollte.

Zu den Orten, an denen Aufnahmen durchgeführt wurden, zählen auch jene sechs, die für die vorliegende Arbeit relevant sind. Dies sind in Deutschland Hamburg und München, in Österreich Wien, Krieglach (Steiermark) und Steyr (Oberösterreich) sowie Zürich in der Schweiz. Die Aufnahmen wurden von verschiedenen Personen an jenen Orten durchgeführt. In Hamburg war Jonas Grünke für die Sprachaufnahmen verantwortlich, er ist wissenschaftlicher Mitarbeiter von Prof. Dr. Gabriel. Die Münchnerinnen wurden von der in Wien arbeitenden aber aus München stammenden Universitätsassistentin von Univ.-Prof. Dr. Pustka, Luise Jansen aufgenommen. In der Steiermark sowie in Oberösterreich beschäftigte sich die zum damaligen Zeitpunkt als Studienassistentin für Univ.Prof. Dr. Pustka arbeitende Isabella Rechberger mit der Rekrutierung und Aufnahme von Informantinnen. Die Aufnahmen in Zürich führte Prof. Dr. Stephan Schmid, Lehrender und wissenschaftlicher Mitarbeiter an der Universität Zürich, durch. In Wien war ich persönlich bei einem der Aufnahmetermine in den Räumlichkeiten des Instituts für Schallforschung, sowie an der zuvor notwendigen Rekrutierung der Teilnehmerinnen beteiligt. Dort fanden die Aufnahmen in einem Tonstudio statt, an den anderen Aufnahmeorten meist mit Audio-Aufnahmegeräten, wie beispielsweise jenen der Marke ZOOM, in möglichst geeigneten Räumen wie etwa Besprechungszimmern,

statt. Daraus ergibt sich auch, dass die Tonqualität der verschiedenen Aufnahmen variiert. Selbst in der gleichen Stadt wurden die Aufnahmen nicht immer in den gleichen Räumlichkeiten gemacht, sodass auch innerhalb dieser das Material von unterschiedlicher Qualität ist. Dies hatte auch zur Folge, dass manche Aufnahmen vor deren Analyse noch mit dem Programm *Audacity* bearbeitet wurden, etwa zur Regulierung der Lautstärke, sowie zur Kürzung von nicht zu analysierenden Pausen. Mit Audacity wurde außerdem der Aufnahmetyp von Stereo in Mono umgewandelt, sodass für die spätere Bearbeitung des Materials in PRAAT, einem freien Programm für phonetische Analysen (entwickelt von Paul Boersma und David Weenik), nur eine anstatt von zwei Tonspuren aufschienen.

3.2.2 Korpus

Die Sprachaufnahmen der Informantinnen folgen einem bestimmten Schema, dessen Kern immer gleich ist. Die Informantinnen sollten nacheinander drei Texte vorlesen. Im Falle dieses Projektes wurden diese in den Sprachen Deutsch, Französisch und Spanisch herangezogen. Dabei handelt es sich um den Text „Nordwind und Sonne“, der in den verschiedensten Sprachen existiert. Diese Texte galt es vom Blatt abzulesen:

Der Nordwind und die Sonne

Einst stritten sich Nordwind und Sonne, wer von ihnen beiden wohl der Stärkere wäre, als ein Wanderer, der in einen warmen Mantel gehüllt war, des Weges daherkam. Sie wurden einig, dass derjenige für den Stärkeren gelten sollte, der den Wanderer zwingen würde, seinen Mantel abzunehmen. Der Nordwind blies mit aller Macht, aber je mehr er blies, desto fester hüllte sich der Wanderer in seinen Mantel ein. Endlich gab der Nordwind den Kampf auf. Nun erwärmte die Sonne die Luft mit ihren freundlichen Strahlen, und schon nach wenigen Augenblicken zog der Wanderer seinen Mantel aus. Da musste der Nordwind zugeben, dass die Sonne von ihnen beiden der Stärkere war.

La bise et le soleil

La bise et le soleil se disputaient, chacun assurant qu'il était le plus fort, quand ils ont vu un voyageur qui s'avavançait, enveloppé dans son manteau. Ils sont tombés d'accord que celui qui arriverait le premier à faire ôter son manteau au voyageur serait regardé comme le plus fort. Alors, la bise s'est mise à souffler de toute sa force mais plus elle soufflait, plus le voyageur serrait son manteau autour de lui et à la fin, la bise a renoncé

à le lui faire ôter. Alors le soleil a commencé à briller et au bout d'un moment, le voyageur, réchauffé a ôté son manteau. Ainsi, la bise a du reconnaître que le soleil était le plus fort des deux.

El viento del Norte y el Sol

El viento del norte y el sol estaban disputándose sobre cuál de ellos era el más fuerte, cuando pasó un viajero envuelto en una capa gruesa. Se pusieron de acuerdo en que el primero que lograra que el viajero se quitara la capa sería considerado más fuerte que el otro. Entonces, el viento del norte sopló tan fuerte como pudo, pero cuanto más soplabla más estrechamente se ceñía el viajero la capa al cuerpo hasta que, al fin, el viento del norte se rindió. Luego, el sol brilló calurosamente y de inmediato el viajero se despojó de su capa. Así, el viento del norte se vio obligado a reconocer que el sol era el más valiente de los dos.

Diese drei Texte in ihrer vorgelesenen Form existieren jeweils von allen Informantinnen als separate Tonaufnahmen. Das Dokument, welches den vorzulesenden Text enthält findet sich im digitalen Anhang. Bei manchen Aufnahmeorten stellte man den Informantinnen danach noch die zusätzliche Aufgabe, die Geschichte, die sie zuvor vorgelesen hatten, mündlich zusammenzufassen. Bei den österreichischen Aufnahmetermi-
nen gab es noch eine weitere Aufgabe, die von den Informantinnen verlangte, wie sie Pizza backen würden. Die letzteren beiden Aufgaben dienten dazu, auch Material in Spontansprache zu erhalten, doch dadurch, dass diese nicht aus allen Städten existieren, wurden für die vorliegende Arbeit nur die Aufnahmen der gelesenen Texte für die Analyse verwendet. Jede Informantin füllte vor der Aufnahme eine Einverständniserklärung sowie einen soziodemographischen Fragebogen aus. Da die Daten aber anonymisiert verwendet werden, hat jede Aufnahme der Teilnehmerinnen einen Codenamen zugewiesen bekommen, der über Stadt, eine Nummer für die Informantin, sowie die jeweilige Sprache der Aufnahme informiert. Dieser Code erklärt sich dann beispielsweise für die Wiener Informantin VI_01_de wie folgt:

VI = Wien **01** = Informantin **de** = deutscher Text

Die weiteren Städtekürzel sind *ST* für Krieglach, *OÖ* für Steyr, *ZH* für Zürich, *HH* für Hamburg und *MU* für München. Wie zuvor schon erwähnt, wurden für die vorliegende Arbeit also sechs Orte ausgewählt und pro Ort drei Informantinnen zur Analyse herangezogen. Die Auswahl der einzelnen Informantinnen beruhte vorrangig auf der Niveauevielfalt, sofern diese im engen Rahmen von A2-B1 noch gegeben war. Anschließend

musste auch die Tonqualität miteinbezogen werden, um die akustischen Analysen möglichst genau durchführen zu können. Die Auswahl der bereits oben genannten Orte hatte verschiedene Gründe. Einerseits sollten mehrere Dialektzonen verglichen werden. Mit den sechs untersuchten Orten werden das Niederdeutsche, das Süd- bis-Mittelbairische, sowie das Alemannische abgedeckt. Es soll hier nochmals kurz erwähnt werden, dass Krieglach in der Steiermark sich rein theoretisch bereits im Übergangsgebiet zwischen süd- und mittelbairischer Dialektzone befindet, München, Wien und Steyr hingegen eindeutig dem Mittelbairischen zugeordnet werden (vgl. Wiesinger 1983: 837). Jedoch gestalten sich hier die lautlichen Unterschiede, vor allem die konsonantischen, bereits so gering (vgl. Wiesinger 1983: 839-840), dass eine nochmalige Unterteilung für vorliegende Arbeit nicht als sinnvoll erschien. Deshalb wird Krieglach zusammen mit Wien, München und Steyr in einer Gruppierung des Süd- bis Mittelbairischen erwähnt.

Von jeder Teilnehmerin existieren drei Aufnahmen, jeweils in einer Sprache. Bei den meisten finden sich diese auch in separaten Audiodateien. An manchen Orten und bei manchen Informantinnen wurde die deutsche und die französische Aufnahme in einem Durchgang aufgezeichnet, sodass die Gesamtzahl an Audiodateien, die ich schlussendlich bearbeitete 48 ist.

3.2.3 Analyse

3.2.3.1 Codierungs-System

Die Audiodateien stellen den Kern der Analyse für die vorliegende Arbeit dar. Anhand des Audiomaterials führte ich eine auditiv-akustische Analyse der Daten durch, die sich auf alle Positionen an welchen ein *r*-Laut, beziehungsweise eine Vokalisierung eines graphischen <*r*> möglich ist, konzentriert. Im Inventar der drei Texte gab es *r*-Laute an vielen möglichen Positionen, also silbenfinal- oder intial, im französischen und spanischen Text auch wortinitial, im deutschen Text nicht. Die phonetische Transkription des Materials war der Grundstein für alle weiteren Arbeitsschritte. Mithilfe des Programms PRAAT war es möglich, alle Audiodateien schriftlich im Internationalen Phonetischen Alphabet (IPA) zu transkribieren, den Text innerhalb einer Datei passend zum Tonmaterial in Silben und einzelne Laute zu segmentieren und anhand der Sonagramme Laute aufgrund ihrer Merkmale zu bestimmen. Oftmals erwies sich diese Aufgabe als sehr kompliziert. Wenn ich bei der Bestimmung eines Lautes in PRAAT allein keine eindeutige Realisierung feststellen konnte, behalf ich mir mit einem IPA-

Chart mit Sounds zur Hand um zu vergleichen, wie beispielsweise jenes auf der Website der Universität des Saarlandes, welches die jeweiligen Laute in initialer, zentraler und finaler Position beinhaltet und mit Sonagrammen ergänzt (<http://www.coli.uni-saarland.de/elaut/konsPulm.htm>).

Die an die Transkription anschließende Segmentierung in Silben hatte den Grund, dass für eine möglichst tiefgehende Analyse der Aussprache die Äußerungen, in denen ein *r*-Laut realisiert werden kann, auch beobachtet werden sollte, wie sich die Informantinnen je nach Position des Lautes in der Silbe, also ob jener in Silbenkopf (Onset) oder Silbenreim (Coda) vorkommt, und weiters ob sie mit weiteren Konsonanten verzweigt oder aber unverzweigt sind, verhalten. Hier gibt es nämlich von Sprache zu Sprache Unterschiede und Besonderheiten, wie *r*-Laute realisiert werden können (siehe dazu Kapitel 2.5.4.1 und 2.5.5.1). Da die Aufnahmen aus gelesenen Texten und nicht aus Spontansprache bestehen, ließ sich schon vorab die Zahl der Realisierungen je Sprache einschätzen. Hierzu muss noch angemerkt werden, dass für die Analyse der *r*-Laute die Überschriften des Lesetextes nicht berücksichtigt wurden, da manche Informantinnen sie nicht vorlasen, andere jedoch schon. Somit entschloss ich mich dazu, die Überschrift aus der Bestimmung herauszunehmen um die Anzahl der realisierbaren *r* möglichst homogen zu halten. So kommen für die Analyse folgende Zahlen zustande:

Im deutschen Text sind insgesamt 48 Realisierungen möglich, davon 11 im Onset und 37 in der Coda. Diese Zahl variiert beispielsweise bei manchen Informantinnen, da das Wort „ihren“ im deutschen Text von manchen einsilbig als [iɛn] (vgl. VI_01_de) produziert wird, von anderen aber mehrsilbig als [i:rən] (vgl. ZH_02_de) und somit auch die Position von *r*, beziehungsweise seiner vokalisierten Form in der Silbe unterschiedlich ist.

Im französischen Text sind es insgesamt 27 mögliche Realisierungen, davon 12 im Onset und 15 in der Coda. Bei den Schweizerinnen sind es auf Französisch jedoch insgesamt nur 26 Realisierungsmöglichkeiten, da im Text, den sie vorlesen mussten, kleine Unterschiede vorkamen. So ist in den Texten der Schweizerinnen der Satz „ils sont tombés d'accord que celui qui arriverait le premier à le lui faire ôter“ zu lesen, während in allen anderen Aufnahmeorten die Ersetzung durch das indirekte Objektpronomen *lui* nicht durchgeführt wird und der Wanderer mit dem Wort *voyageur* nochmals

wiederholt wird (vgl. Kapitel 3.2.2). Dadurch kommt in der Schweiz ein Coda-*r* abhanden. Aufgrund dieses Unterschieds ergeben sich daher bei den Schweizerinnen 11 mögliche *r*-Realisierungen im Onset und 15 in der Coda.

Im spanischen Text finden sich 33 mögliche Realisierungen, davon 23 im Onset und 10 in der Coda. Des Weiteren wurde im Spanischen aufgrund seiner Standardlautung davon ausgegangen, dass von 23 Realisierungen im Onset zwei den alveolaren Trill [r] verlangen. Diese beiden Stellen sind die Wörter span. *reconocer* [rɛkonosɛr] ‚erkennen‘ und span. *rindió* [rindjo] ‚gab auf‘ aus dem Text. Die restlichen 21 Möglichkeiten verlangen laut Standard den alveolaren Tap [r] (vgl. Hualde 2005: 185). Diese Annahme basiert auf den für das Spanische geltenden Aussprachenormen, die in den Kapiteln 2.3.2 und 2.5.5. diskutiert wurden.

Diese Anzahl von der anfänglich ausgegangen wurde, variiert selbstverständlich je nach Sprecherin, denn auch in vorgelesenen Texten passieren Fehler oder werden aneinanderstoßende Laute zu einem zusammengefasst, sowie Worte wiederholt gelesen.

Des Weiteren wurden für die anschließende Analyse die Realisierungen nochmals unterteilt, sodass sie schlussendlich in verzweigte und unverzweigte Onset- und Codaposition gegliedert waren. So kam die finale Zahl an Lauten nach Kategorie in Deutsch auf drei verzweigte, sowie neun unverzweigte Onsets, fünf verzweigte, und 30 unverzweigte *r*-Laute in der Coda. Im Französischen Text fanden sich alles in allem drei verzweigte und neun unverzweigte *r*-Laute im Onset, sowie ein verzweigter und 14 unverzweigte *r*-Laute in der Coda. Im Spanischen Text kommt die Unterteilung auf insgesamt sieben verzweigte und 13 unverzweigte *r*-Laute im Onset, keinen verzweigten und zehn unverzweigte *r*-Laute in der Coda. Wie schon vorhin erwähnt, variierte auch hier die Anzahl bei manchen Informantinnen, da durch Lesefehler oder das Zusammenziehen von Lauten nicht immer jede Position gleich realisiert wurde.

Da zusätzlich zur statistischen Auswertung im Anschluss daran noch die reinen Zahlen ohne Miteinbezug der Signifikanztests diskutiert werden, soll auch erklärt werden, wie jene Ergebnisse, meist in Prozent angegeben, zustande kamen. Um einen möglichst anschaulichen Überblick über die verschiedensten Realisierungen nach Aufnahmeorten gewährleisten zu können, errechnete ich die Ergebnisse der nicht-statistischen Auswertung meist anhand der Gesamtzahl an Okkurrenzen pro Ort, nicht wie zuvor

nach Informantin. So wurde im Deutschen mit einer Gesamtzahl an 144 möglichen *r*-Positionen pro Ort, im Französischen 81 beziehungsweise 78 in der Schweiz, sowie im Spanischen mit 99 gerechnet. Wenn für die gesamte Sprache ein Überblick an Realisierungen dargestellt werden sollte, rechnete ich nochmals in die absolute Gesamtanzahl möglicher Okkurrenzen aller Informantinnen pro Sprache um.

Neben den Dateien in PRAAT habe ich die Transkriptionen ebenfalls in Microsoft Word gesammelt, sodass ein schneller Überblick über den gesamten Text in phonetischer Transkription möglich ist. Die fertigen Transkriptionen dienten im Anschluss an die auditiv-akustische Untersuchung als Grundlage für eine gezielte Analyse der *r*-Laute. Vor Beginn der Transkriptionsarbeit stellte ich mich auf sieben Realisierungsmöglichkeiten der *r*-Laute ein, die ich zu erwarten hatte. Dazu zählten der alveolare Tap [ɾ], der alveolare Trill [r], der stimmhafte uvulare Frikativ [ʁ], sowie sein stimmloses Gegenstück [χ], der stimmhafte uvulare Trill [R], sowie Formen der Vokalisierung wie das so genannte Lehrer- oder Tiefschwa [ɐ]⁴ und die Verlängerung von Vokalen vor einem graphischen <r> wie beispielsweise in dt. *war* [a:]. Die Anzahl der vorab erwarteten Laute sollte sich in der Realität mehr als verdoppeln. Nach Abschluss der Transkriptionsarbeit stand eine Auswahl an 15 verschiedenen Realisierungsformen fest. Dies sind der stimmhafte retroflexe Tap bzw. Flap [ɽ], sowie eine entstimmte Realisierung dessen [ɽ̥], der erwartete alveolare Tap, jedoch ebenfalls entstimmt [ɾ̥], der uvulare Approximant [ʁ̥], oder Elisionen, die mit Ø gekennzeichnet werden. Zu den Elisionen zählen einerseits Fälle, die durch eine komplette Auslassung des Wortes passierten, sowie Versprecher oder Lesefehler (z.B. *der Stärkste* anstatt *der Stärkere* ergibt ein <r> weniger, das realisiert wird) und Zusammenziehen von zwei *r*-Lauten an der Wortgrenze wie beispielweise im französischen Text an der Stelle *voyageur rechauffé*.

[alfɪnɛlvjɛntodɛlnoɐtɛsɛ**RinRindij**o] Wiederholung bei OÖ_06_es

[veɛfɒni:nənbaɪdɪvo:ldeɐ**fteekst**ɔvɛɐ] Lesefehler bei OÖ_06_de

[eobudɛmɔmälə**vwajafœxə**ɔfeatoesmāto] Zusammenziehen bei OÖ_04_fr

⁴ Im weiteren Verlauf vorliegender Forschungsarbeit wird der Einfachheit halber nur noch von Schwa [ɐ] gesprochen, da das Hochschwa [ə] keine untersuchte Variable darstellte.

Bei der Auswertung wurden Wiederholungen von *r*-Lauten auf zwei verschiedene Arten gezählt. Wenn es sich um die gleiche Realisierung handelt, wie im vorangegangenen Beispiel OÖ_06_es, kamen die einzelnen [ʀ] auf einen Wert von je 0,5 und wurden so durch 2 mal 0,5 wieder zu einer Okkurrenz mit dem Wert von 1. Kamen in einer partiellen oder kompletten Wortwiederholung jedoch zwei verschiedene Realisierungen vor, wurden die Kommawerte von beispielsweise 0,5 beibehalten und finden sich daher auch in der Auswertung als solche wieder. Bei Zusammenziehen von zwei *r*-Lauten wie im zuvor genannten Beispiel bei OÖ_04_fr wurde zu einer Okkurrenz mit dem Wert 1 zusammengezählt. Somit haben nicht alle Informantinnen immer dieselbe Anzahl an realisierten Lauten, auch wenn bei vorgelesenen Texten die Annahme bestehen könnte, dass dem so sein möge.

Zusätzlich kam es noch zur Eröffnung dreier weiterer Kategorien, die zwar selten vorkamen, jedoch trotzdem markant waren. Eine davon ist die Kategorie ?, beziehungsweise *unbekannt*, für Laute die aus verschiedenen Gründen (beispielsweise Lärm) schlichtweg nicht eindeutig zu bestimmen waren. Eine weitere Kategorie wurde für Realisierungen eingeführt, bei welchen [χ] ansatzweise zu hören und im Sonagramm zu sehen war, jedoch nur sehr schwach. Deshalb bekam diese Kategorie die Bezeichnung [χ] für einen hochgestellten stimmlosen uvularen Frikativ. Die letzte, zwar nur einmal vorkommende, aber dennoch auffallende Realisierung die kategorisiert wurde, ist ein aspirierter alveolarer Tap [rʰ]. Wie es zur Entscheidung kam, diese Kategorien zu eröffnen und wie diese unerwarteten Laute im Sonagramm aussehen, wird im Anschluss an diesen Absatz genauer mit Bildmaterial aus den jeweiligen Aufnahmen mit Sonagramm in PRAAT erläutert. Die finale Auswertung kam auf diese 15 Realisierungen:

[ʀ]	[r]	[r]	[χ]	[ʁ]	[ʁ]	[r]	[r]	[r]	∅	[e]	[a:] ⁵	[χ]	[rʰ]	?
-----	-----	-----	-----	-----	-----	-----	-----	-----	---	-----	-------------------	-----	------	---

Tabelle 1: Finale Auflistung der analysierten Realisierungen von *r*

Eine Erläuterung dazu, wieso die endgültige Entscheidung auf diese 15 Laute gefallen ist, sowie eine kurze Beschreibung, soll nun kurz mit jeweils einem Screenshot einer Aufnahme aus PRAAT zur Veranschaulichung dargestellt werden.

⁵ [a:] steht hier verallgemeinernd für Vokalverlängerungen aller Art, nicht nur für die des Vokals [a]

[R] stimmhafter uvularer Vibrant: Dieser Laut gehört zur Klasse der Trills, da er in der Regel durch mehrere Schläge des hinteren Teils der Zunge am Zäpfchen realisiert wird (vgl. Ladefoged/Maddieson 1996: 225). Das folgende Sonagramm ist ein Screenshot aus der Aufnahme HH_02_fr:

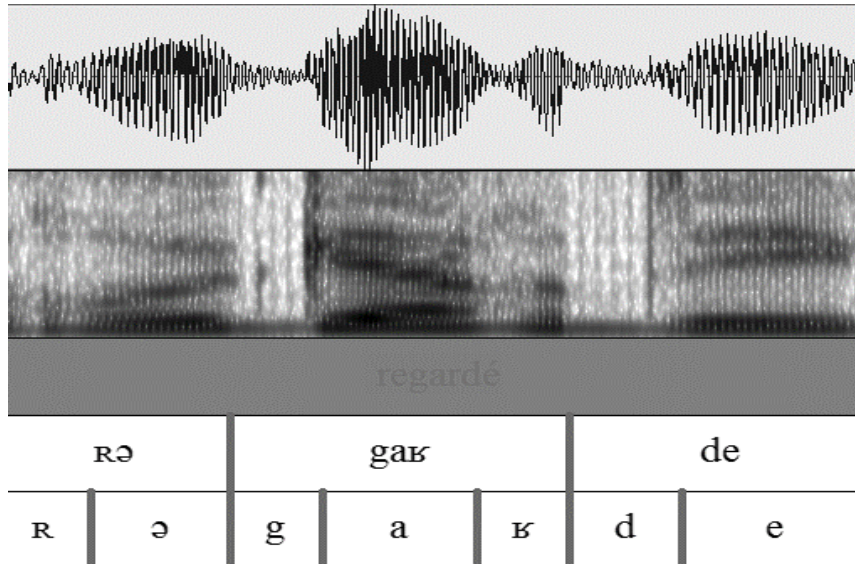


Abb. 3: HH_02_fr: Der uvulare Vibrant [R] silben-initial

HH_02_fr: *regardé* [Rəgaʁde]

[r] stimmhafter alveolarer Vibrant: Dieser Laut gehört ebenfalls zur Klasse der Trills. Er wird normalerweise mit der Zungenspitze an den Alveolen mit mehreren Schlägen realisiert (vgl. Ladefoged/Maddieson 1996: 221). Das Sonagramm stammt aus der Aufnahme ST_02_fr:

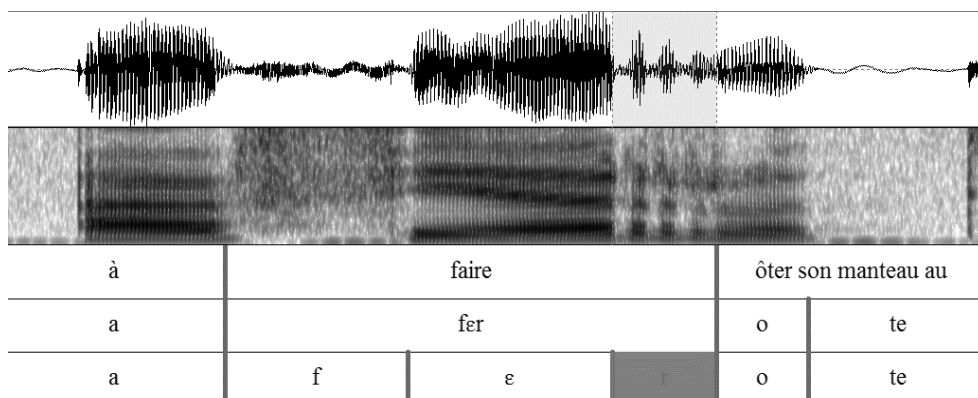


Abb. 4: [r] in der französischen Aufnahme zwischen zwei Vokalen

ST_02_fr: *faire ôter* [færote]

[r] stimmhafter alveolarer Tap: Dieser Laut wird im Gegensatz zu [r] nicht mit mehreren, sondern nur mittels eines Zungenschlags an den Alveolen realisiert (vgl. Göschel 1971: 91). Das folgende Sonagramm in dem man den typischen Schlag erkennt, stammt aus der Aufnahme ZH_02_es:

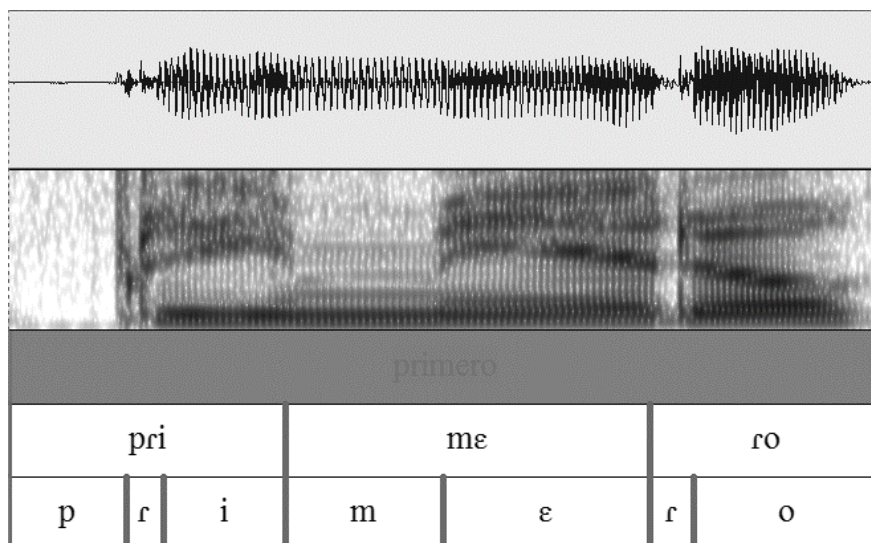


Abb. 5: [r] in sowohl postkonsonantischer, als auch intervokalscher Position

ZH_02_es: *primero* [primεro]

[χ] stimmloser uvularer Frikativ: Dieser Laut wird durch Konstriktion am Zäpfchen sowie den für Reibelaute typischen Luftstrom charakterisiert (vgl. Tranel 1987: 141). Das Beispielfeld stammt aus der Aufnahme HH_02_fr:

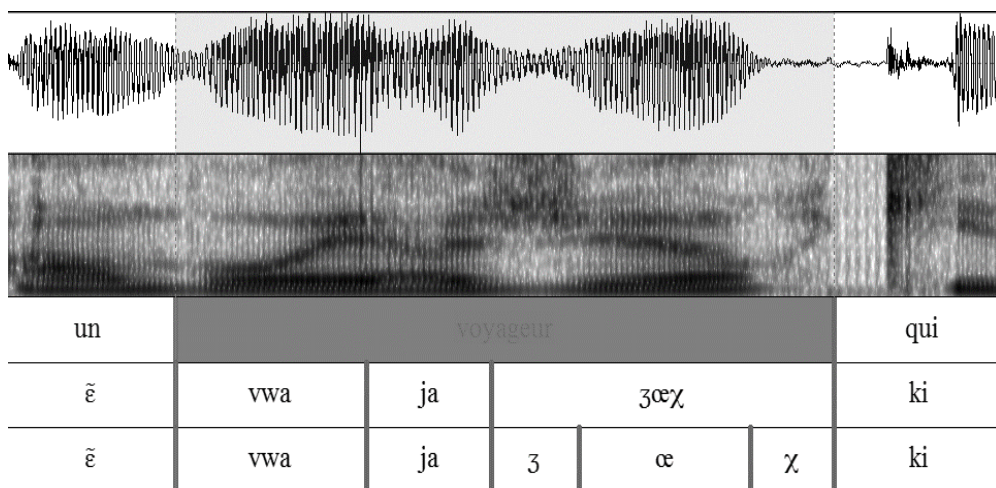


Abb. 6: [χ] vor stimmlosem Konsonanten [k] in *qui* [ki]

HH_02_fr: *voyageur* [vwaja3œχ]

[ʁ] stimmhafter uvularer Frikativ: Dieser Laut wird auf die gleiche Weise wie [χ] produziert, im Gegensatz dazu ist er aber stimmhaft. Das folgende Beispiel stammt aus der Aufnahme HH_03_fr:

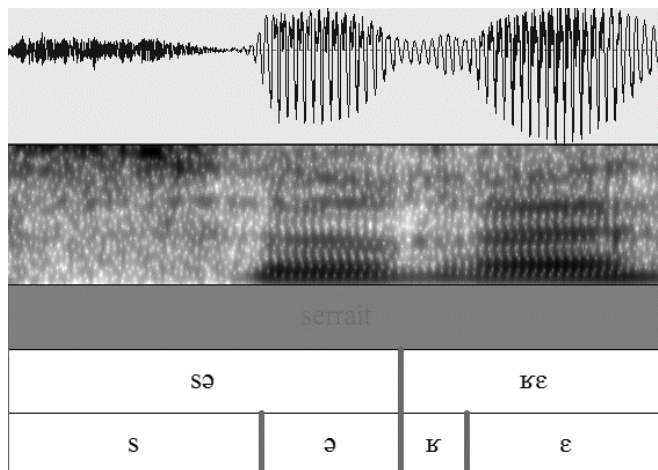


Abb. 7: [ʁ] im Onset der Silbe, Stimmhaftigkeit am *voice bar* (vgl. Hualde 2014: 50) zu erkennen.

HH_03_fr: *serrait* [səʁɛ]

[ɣ] stimmhafter uvularer Approximant: Wie auch [ʁ] und [R] wird dieser Laut uvular realisiert, jedoch findet keine Friktion statt. Diesem Laut wird im Internationalen Phonetischen Alphabet kein eigenes Symbol zugeschrieben, sodass es durch das diakritische Zeichen für gesenkt, ergänzt wird, um ihn zu beschreiben (vgl. International Phonetic Association ¹⁵2014: viii) Im Sonagramm erkennt man [ɣ] im Gegensatz zu [ʁ] vor allem an der fehlenden Friktion (vgl. Thomas 2010: 132) sowie an den Formanten, die jenen der Vokale ähneln, wenngleich weniger stark ausgeprägt (vgl. Hualde 2014: 52) wie im folgenden Beispiel aus HH_04_es zu erkennen ist:

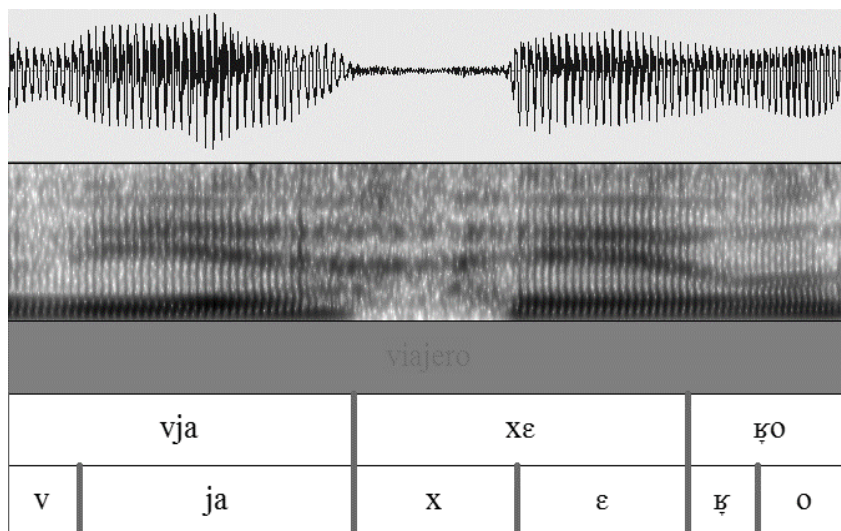


Abb. 8: Approximant [ɣ] zwischen zwei Vokalen, es ist deutliche Stimmhaftigkeit erkennbar

HH_04_es: *viajero* [vjaɣeɾo]

[ɾ] entstimmter alveolarer Tap: Diese Kategorie wurde für alveolare Taps [ɾ] verwendet, die im Sonagramm jedoch als entstimmt erkennbar waren, wie beispielsweise in ST_02_fr:

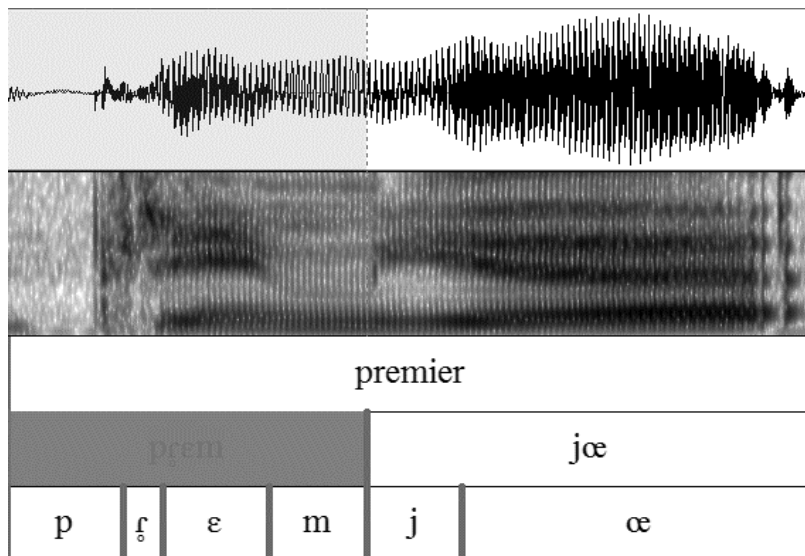


Abb. 9: [ɾ] erkennbar, jedoch deutet fehlender *voice bar* auf entstimmten Charakter hin.

ST_02_fr: *premier* [pɾɛmjœ]

[ɽ] stimmhafter retrofleher Flap: Bei dieser Modifikation des alveolaren Schlaglautes wird die Zungenspitze zurückgebogen und durch die nun den Alveolen angenäherte Zungenunterseite eine Enge gebildet (vgl. Göschel 1971: 95). Im Sonagramm vereint er die verengende Qualität des Approximanten und den Schlag des Tap (The Atlas of North American English), wie im Beispiel aus ST_02_fr zu sehen:

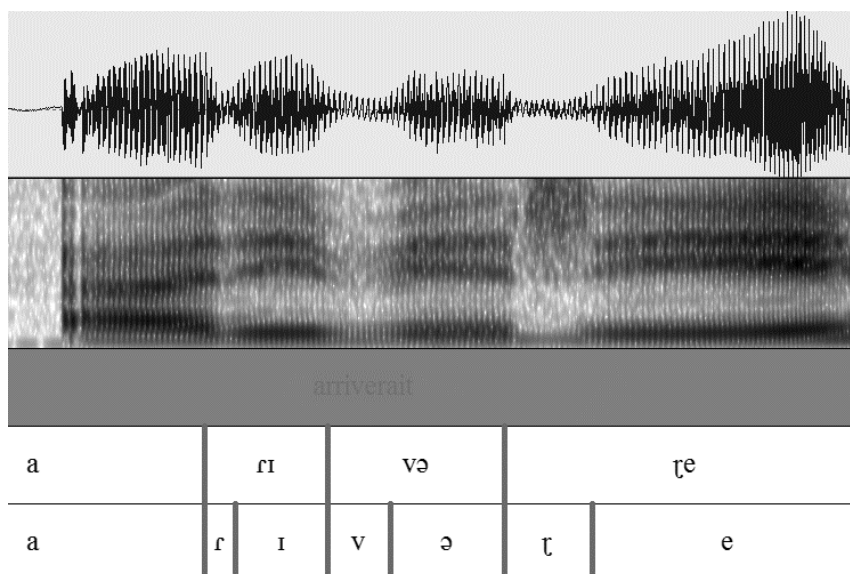


Abb. 10: [ɽ] in der letzten Silbe im Onset

ST_02_fr: *arriverait* [arɪvəɽe]

[ɾ] entstimmter retrofleher Flap: ebenfalls eine Kategorie des retroflexen Flap, der seine Stimmhaftigkeit verloren hat, wie hier in VI_01_es:

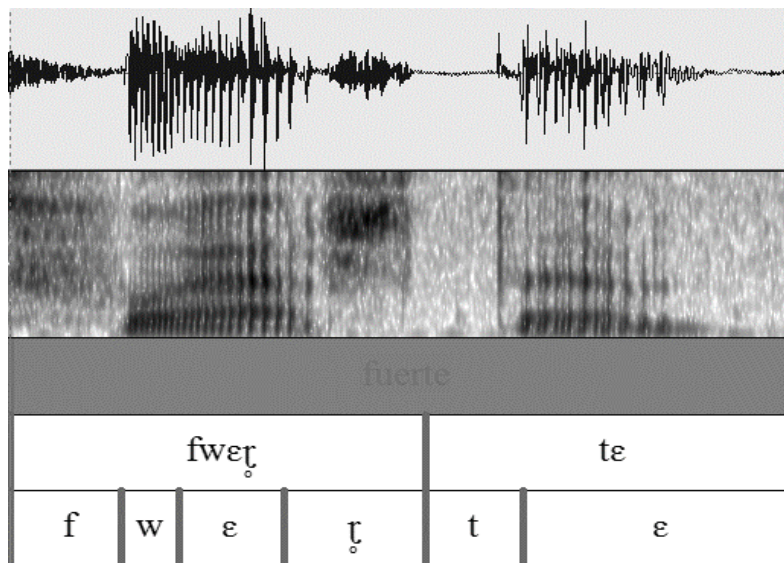


Abb. 11: Entstimmter retrofleher Flap [ɾ], erkennbar am fehlenden *voice bar*.

[ɐ] ‚Lehrer-Schwa‘: ein ungerundeter Zentralvokal, zwischen halber und kompletter Öffnung angesiedelt. Dieser Vokal ersetzt *r*-Laute häufig in Form eines Monophthongs oder in Kombination mit einem davorstehenden Vokal als Diphthong, wie hier in den Beispielen aus VI_02_de zu erkennen:

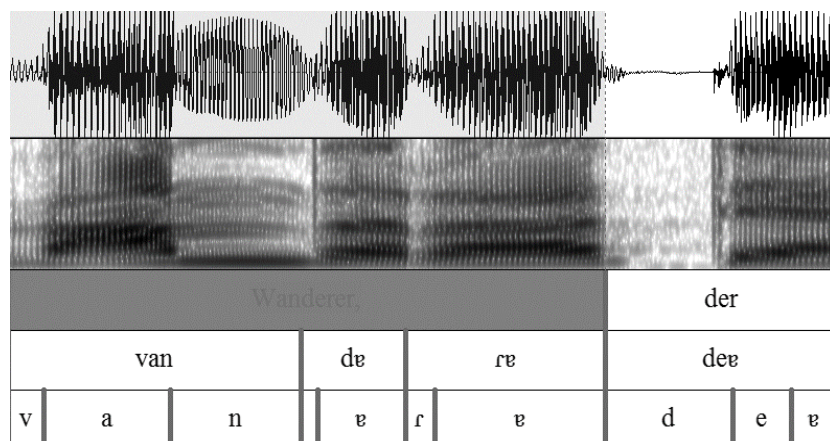


Abb. 12: [ɐ] ersetzt konsonantisches *r* in Form eines Monophthongs.

VI_02_de: *Wanderer* [vandɐrɐ]

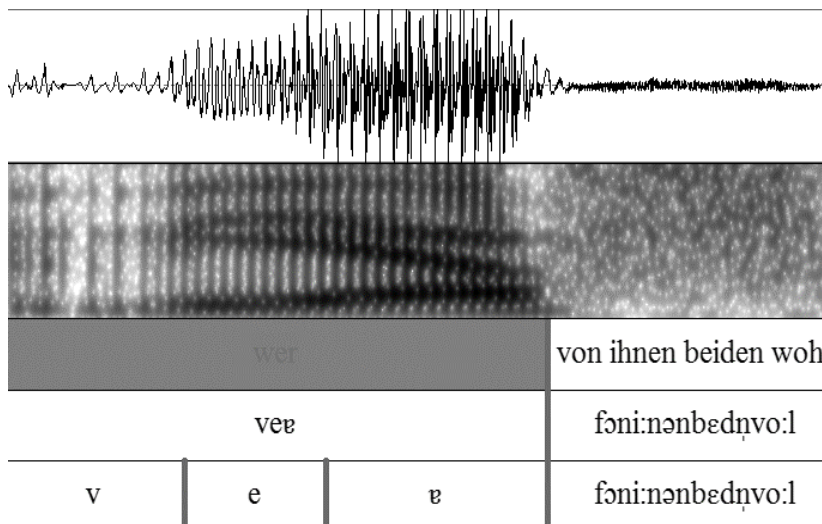


Abb. 13: [ɐ] ersetzt konsonantisches *r* in Form eines Diphthongs

VI_02_de: *wer* [veɐ]

[a:] Verlängerter Vokal: diese Bezeichnung wurde verwendet, wenn der *r*-Laut zur Gänze elidiert und der davorstehende Vokal in die Länge gezogen wurde, wobei es sich nicht ausschließlich um den Vokal [a] handeln muss. Dieser dient hier in der Beschriftung lediglich als Bezeichnung für die Kategorie. Ein Beispiel für diese Realisierungsmöglichkeit liefert uns VI_02_de im folgenden Screenshot:

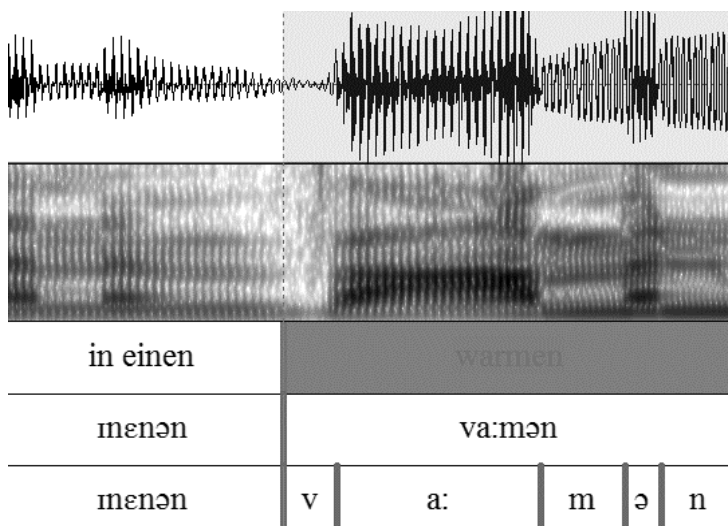


Abb. 14: *r* fällt weg, der davorstehende Vokal [a] wird in die Länge gezogen.

VI_02_de: *warmen* [va:mən]

Ich erwähne hier dezidiert, dass es sich nicht ausschließlich um den Vokal [a] handeln muss, da in den Transkriptionen auch Okkurrenzen eines verlängerten [ɔ] vorkamen.

Diese streiche ich deshalb heraus, da sie in deutschen, aber auch in spanischen Aufnahmen auffielen, und die konventionelle Transkription des Spanischen normalerweise den Vokal [o] anschreibt (vgl. Martínez-Celdran/Fernández-Planas/Carrera-Sabaté 2003: 255). Da ich aber verdeutlichen wollte, dass die Informantinnen in diesem Fall aber eindeutig nicht den gerundeten halbgeschlossenen Hinterzungenvokal [o], sondern den gerundeten halboffenen Hinterzungenvokal [ɔ] realisieren, findet sich dieser in den spanischen Transkriptionen dort, wo er auch so ausgesprochen wurde. Ein Screenshot, der dies veranschaulicht, findet sich später in Kapitel 5.1 und 5.3 im Zuge der Diskussion der Ergebnisse.

[χ] stimmloser uvularer Frikativ hochgestellt: bei dieser Kategorie war ein Laut zu erkennen, jedoch ist dieser akustisch nur sehr schwach ausgeprägt und das Reibegeräusch vermindert (vgl. Vremšak-Richter 2010: 96), sodass die hochgestellte Darstellungsform gewählt wurde. Im Programm PRAAT ist das Hochstellen eines Zeichens nicht möglich, sodass ich dort diesen Laut als $\chi(sup)$ für ‚hochgestellt‘ kennzeichnete, hierzu ein Beispiel aus HH_04_fr:

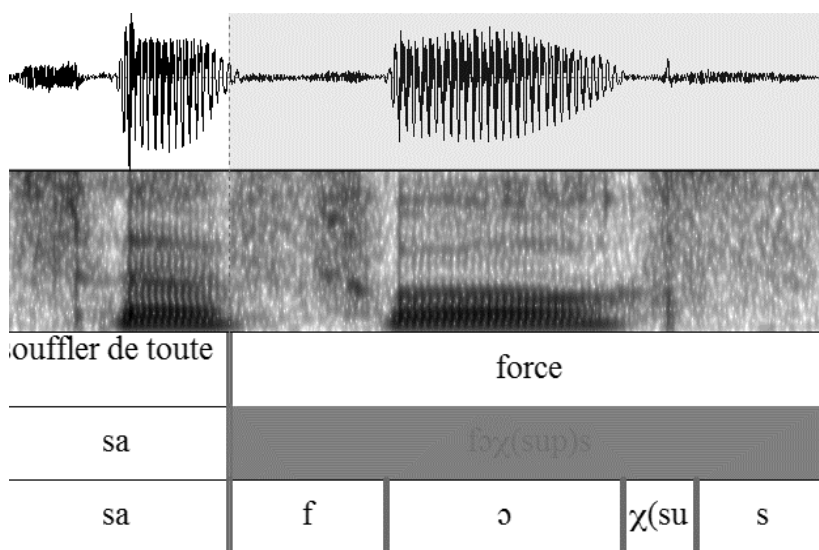


Abb. 15: Schwach erkennbarer Laut [χ] zwischen Vokal [ɔ] und Konsonant [s].

HH_04_fr: force [fɔχs]

[rʰ] stimmhafter alveolarer Tap mit Aspiration: diese Kategorie zählt nur eine einzige Okkurrenz, die trotzdem in die Auswertung aufgenommen wurde. Diese Realisierung fand sich in der Aufnahme...und zeigt eine Aspiration [ʰ] nach [r], die von dem davorstehenden [p] ausgehen könnte:

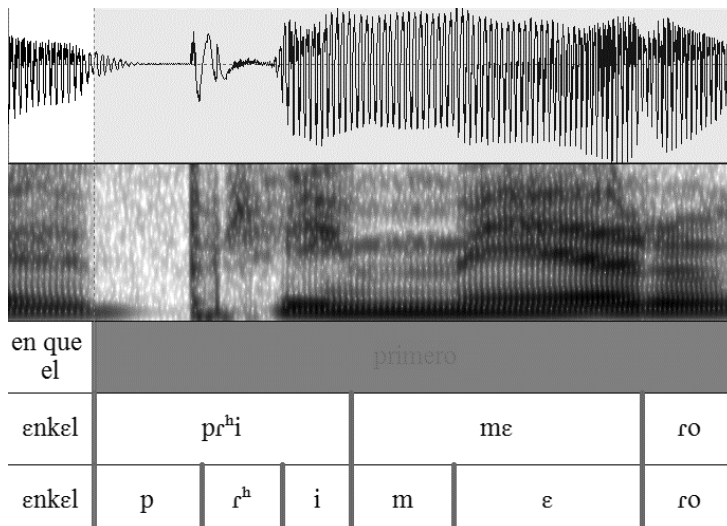


Abb. 16: Aspiration zwischen Tap [r] und Vokal [i], evtl. Einfluss von Plosivkonsonant [p] wortinitial

HH_03_es: *primero* [pr^himɛrɔ]

Ø Elision: meist bedingt durch Lesefehler, in denen Worte falsch vom Blatt gelesen oder zusammengezogen wurden und somit ein *r*-Laut verloren geht.

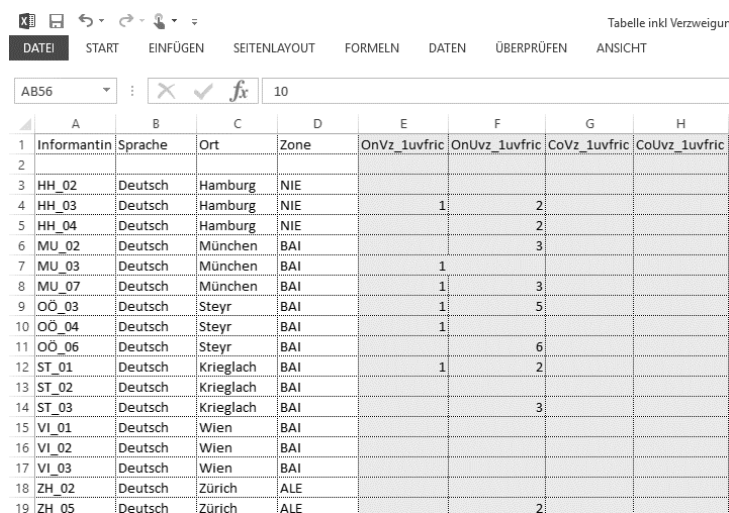
?-Kategorie: Für Laute, die nicht eindeutig zuzuordnen waren, wurde die ?- oder *undefiniert* Kategorie eröffnet.

Um nun all diese Laute aus den Transkriptionen herauszuarbeiten, erstellte ich zuerst Tabellen in Microsoft Excel. Diese beinhalten den Codenamen der Informantin, den Ort, die Sprache der jeweiligen Aufnahme, sowie eine Wortliste und eine Liste mit der phonetischen Transkription des jeweiligen Wortes, die Silbenposition des *r*-Lautes und schlussendlich die Realisierung. Im Anschluss an die Erstellung dieser, der Übersicht über das Gesamtinventar dienenden Tabelle, erstellte ich eine weitere, diesmal schon geordnet nach den schlussendlich 15 verschiedenen Realisierungsformen und deren Vorkommnis je Informantin und Sprache in absoluten Zahlen. Mithilfe jener Tabellen konnten die Ergebnisse im späteren Verlauf auch in Diagrammen visualisiert werden, die verschiedene Realisierungen nach Ort, Dialektzone, Sprache und Silbenposition, sowie Verzweigungen innerhalb dieser anschaulich darstellen. Da ich sowohl Tests unter Miteinbeziehen der Verzweigung, welche die lediglich Onset- und Codaposition, als auch gewisse Laute im Allgemeinen ohne jene Unterscheidung berücksichtigen durchführte, wurden zwei verschiedene Tabellen erstellt, um den Überblick einfacher zu halten. Um zu gewährleisten, dass die erstellten Tabellen in weiterer Folge mit anderen Programmen, wie beispielsweise dem Statistikprogramm SPSS, kompatibel sind, scheinen die Realisierungen im Dokument nicht in ihrer Entsprechung im internationalen phonetischen Alphabet auf, sondern in Form eines von mir selbst erstellten

Codes, welcher darüber informiert, ob der jeweilige Laut stimmhaft (=1) oder stimmlos (=0), oder gar entstimmt (=01) ist, eine Abkürzung von Artikulationsort und –art angibt und ob er sich in Onset- oder Codaposition befindet, sowie ob der Onset oder die Coda eine Verzweigung mit weiteren Konsonanten aufweist. Hierfür wurden die Kürzel *On* und *Co* festgelegt. Des Weiteren entstanden für die Angabe nach etwaiger Verzweigung noch die zusätzlichen Kürzel *Vz* für *verzweigt* und *Uvz* für *unverzweigt*. So steht beispielsweise für den verzweigten stimmhaften uvularen Frikativ [ʁ] in der Coda folgender Code: **CoVz_1uvfric**

Co = Coda Vz = verzweigt 1=stimmhaft uvfric = uvularer Frikativ

In den Tabellen wurde dann die Anzahl an den jeweiligen Realisierungen unter jenem Code in der entsprechenden Zeile und Spalte eingefügt:



	A	B	C	D	E	F	G	H
1	Informantin	Sprache	Ort	Zone	OnVz_1uvfric	OnUvz_1uvfric	CoVz_1uvfric	CoUvz_1uvfric
2								
3	HH_02	Deutsch	Hamburg	NIE				
4	HH_03	Deutsch	Hamburg	NIE	1	2		
5	HH_04	Deutsch	Hamburg	NIE		2		
6	MU_02	Deutsch	München	BAI		3		
7	MU_03	Deutsch	München	BAI	1			
8	MU_07	Deutsch	München	BAI	1	3		
9	OÖ_03	Deutsch	Steyr	BAI	1	5		
10	OÖ_04	Deutsch	Steyr	BAI	1			
11	OÖ_06	Deutsch	Steyr	BAI		6		
12	ST_01	Deutsch	Krieglach	BAI	1	2		
13	ST_02	Deutsch	Krieglach	BAI				
14	ST_03	Deutsch	Krieglach	BAI		3		
15	VI_01	Deutsch	Wien	BAI				
16	VI_02	Deutsch	Wien	BAI				
17	VI_03	Deutsch	Wien	BAI				
18	ZH_02	Deutsch	Zürich	ALE				
19	ZH_05	Deutsch	Zürich	ALE		2		

Abb. 17: Ausschnitt aus-
Excel-Tabelle der Aus-
wertung

3.2.3.2 Statistische Auswertung

Aus dieser Sammlung aller Daten in Zahlen konnten im Anschluss Tests auf etwaige Signifikanz der Ergebnisse ausgeführt werden. Aufgrund der relativ kleinen Datenmenge (drei Informantinnen pro Ort) und der Heterogenität der Ergebnisse eignet sich für die Überprüfung der Signifikanz der Ergebnisse der Chi²-Test nach Pearson, der Abhängigkeiten von Merkmalausprägungen testet (vgl. Duller 2013: 244). So lässt sich beispielsweise überprüfen, ob die Anzahl an Schwa [ə], die von einigen Informantinnen beim Vorlesen des Spanischen Textes realisiert werden, in Abhängigkeit mit dem Herkunftsort der Informantinnen steht oder nicht. Diese Art von Test lässt sich auch schon auf kleinere Datenmengen anwenden. So sollte festgestellt werden, welche Realisie-

rungsformen eventuell tatsächlich mit dem Herkunftsort der Informantin zusammenhängen und bei welchen es laut Testung keinen Zusammenhang gibt. Um diese Überprüfungen durchführen zu können, wurde das Statistikprogramm SPSS verwendet. Die zuvor in Microsoft Excel erstellten Tabellen wurden in SPSS geöffnet um mithilfe des Programms die erwähnten Tests durchzuführen.

Bevor die Chi²-Tests durchgeführt werden konnten, waren in der Tabelle lediglich die Kategorien Informantin, Sprache, Ort, Zone und anschließend die Realisierungsformen in Silbenposition unterteilt. Die Kategorie Zone fasste nochmals die Dialektzonen der Informantinnen in drei neue Gruppen zusammen: *BAI* für Bairisch, *NIE* für Niederdeutsch, *ALE* für Alemannisch. Im Nachhinein testete ich jedoch nicht nochmals dezidiert nach diesen Kategorien, sie dienten vorrangig der Übersicht. Um damit nun aber bestimmten Zusammenhänge nachgehen zu können, musste die Tabelle um einige Variablen erweitert werden. Für die bereits erwähnte Art von Test mussten neben den bereits anfänglich feststehenden Variablen wie beispielsweise dem Ort noch die Kategorien für die Realisierungsvariablen definiert werden. Um sinnvolle Kombinationen testen zu können, war es notwendig, zusätzlich präzisere Variablen zu bestimmen. Dazu erstellte ich Kategorien, in denen ich, abhängig davon was getestet werden sollte, verschiedene Realisierungsformen zu einer größeren Variable zusammenfasste, wie beispielsweise eine für alle gesammelten uvularen Realisierungsmöglichkeiten. Ich beschränkte diese jedoch jeweils auf die Sprache, auf die ich sie testen wollte. Neben Variablen für die jeweiligen Laute und deren Silbenposition versuchte ich ebenfalls, Variablen für eine „akzentfreie“ Ziellautung zu erstellen, in welchen sich die Realisierungsmöglichkeiten befanden, die den jeweiligen Standards des Französischen und Spanischen entsprechen, also [ʁ] und [χ], sowie [r] und [r] (vgl. Kapitel 2.5). Wobei hier ein Element aus dem Ergebnisteil vorgezogen werden muss, um zu erwähnen, dass die Variable für die standardsprachliche Lautung im Spanischen *s.standard* im Endeffekt nur aus den Komponenten *on_1alvtap* und *co_1alvtap* besteht. Dies hat den Grund, dass in keiner der spanischen Aufnahmen im gesamten Korpus der alveolare Trill [r] realisiert wurde und diese Variable in dieser Form sich somit als obsolet erweist.

Alles Weitere konnte jedoch ohne Hindernisse getestet werden, und um spezifisch nach den jeweiligen Sprachen testen zu können, entstanden neue Variablen, die vor

dem Code für den jeweiligen Laut noch ein *f* oder ein *s* für die in Frage stehende Sprache hatten. Auch auf Merkmale der deutschen Sprache wurde getestet, diese Variablen erhielten das Kürzel *d* vor dem Code für den Laut. Wenn keine neu erstellte Variable getestet wurde, sondern lediglich eine der Kategorien die sich bereits in der Tabelle befanden, konnten mittels der in SPSS vorhandenen Funktion ‚Fälle auswählen‘ immer nur auf jene Fälle getestet werden, die der jeweiligen Sprache angehörten.

Dies war meine Vorgehensweise bei der Testung der spezifischen Laute in den jeweiligen Sprachen. Mit dieser Methode wurde auf verschiedenste Realisierungsmöglichkeiten auf ihren Zusammenhang mit der Variable *Ort* auf etwaige Signifikanzen überprüft und durch die Anpassung der Auswahl von Fällen in nur der dann jeweilig relevanten Sprache konnte sichergestellt werden, dass gezielt vorgegangen werden konnte. Für die Überprüfung der verschiedenen Kombinationen an Variablen wurden, wie eben schon erwähnt, Kategorien für neue Variablen erstellt. Dies geschieht im Programm SPSS mit dem Menüpunkt ‚Transformieren‘. Dort wählte ich ‚Variable berechnen‘ aus und so öffnet sich ein Fenster, das einem genau dies ermöglicht. Zuerst gibt man der Zielvariable eine Bezeichnung, beispielsweise *s.1alvtap*, wenn man alle stimmhaften alveolaren Taps, die im Spanischen realisiert wurden, egal in welcher Silbenposition, und ohne Rücksicht auf eventuelle Verzweigungen, untersuchen möchte. Dann gilt es, der Zielvariable einen numerischen Ausdruck zuzuordnen, welcher im Falle dieses Beispiels unser alveolarer Tap [r] ist und fügt diesen aus der Liste der verfügbaren Laute in das Feld ‚Numerischer Ausdruck‘ ein. Wenn gewünscht ist, sowohl Okkurrenzen in Onset als auch in Coda darin zu vereinen, müssen beide, also jener mit der Bezeichnung *ON* und jener mit der Bezeichnung *CO*, aus der Liste ausgewählt, und mit einem + verbunden werden. Im folgenden Beispiel wurde lediglich die Variable für alveolare Taps im Onset im Spanischen erstellt:

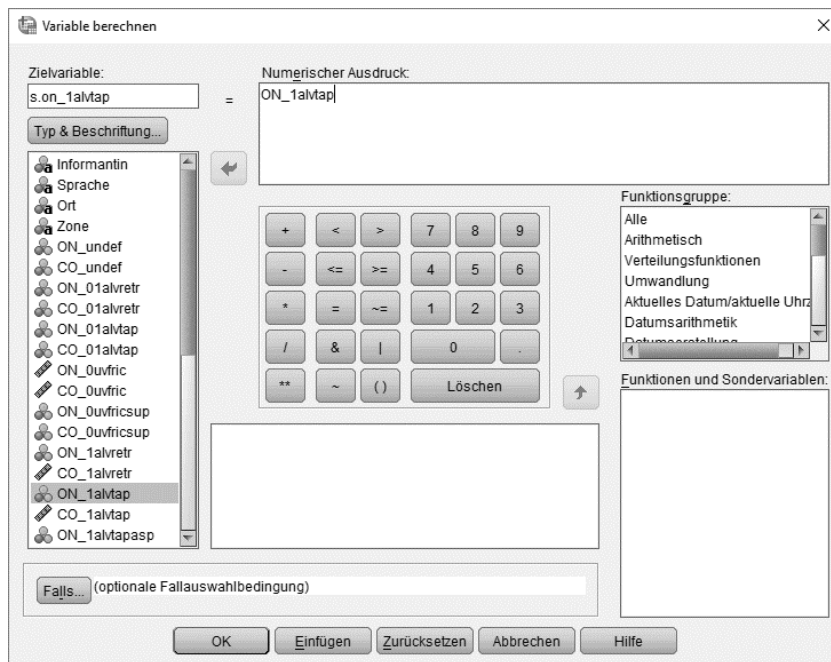


Abb. 18: Variable berechnen in SPSS

Mit einem Klick auf OK entsteht dann eine neue Spalte im Raster, die die ausgewählte Bezeichnung trägt. Wenn man nun jedoch nur die Okkurrenzen im Spanischen testen will, muss man zunächst noch alle anderen etwaig vorkommenden [r] aus der Berechnung entfernen und nur jene Fälle beibehalten, die im Spanischen realisiert wurden, wie bereits erklärt. Auf diese Art und Weise war es möglich, eine Vielzahl an Variablen zu erstellen, die man in weiterer Folge auf ihre Signifikanz in Zusammenhang mit dem Herkunftsort der Informantin testen konnte. Dies geschieht, indem man den Menüpunkt ‚Analysieren‘ auswählt und den Unterpunkt ‚Deskriptive Statistiken‘ auswählt, um in weiterer Folge die Option ‚Kreuztabellen‘ anzuwählen.

In diesem Fenster können nun die Variablen für den Chi²-Test ausgewählt werden. Dazu muss zuvor noch im Menüpunkt ‚Statistiken‘ die Option ‚Chi-Quadrat‘ ausgewählt werden. Als Variablen wurden dann immer ‚Ort‘ und die jeweilige neu definierte oder bereits vorhandene Variable herangezogen.

Eine finale Auflistung der anfänglich enthaltenen Variablen sowie deren „Übersetzung“ befinden sich in folgender Tabelle:

uvfrc = uvularer Frikativ [χ]	alvretr = alveolarer Retroflex [ɾ]
uvfricsup = uvularer Frikativ hochgestellt [χ̥]	schwa = Schwa [ə]
uvapprox = uvularer Approximant [ɣ]	vocv = Vokalverlängerung
uvvibr = uvularer vibrant [ʀ]	elision = Elision

alvvibr = alveolarer Vibrant [r]	? = undefiniert
alvtap = alveolarer Tap [r]	

Tabelle 2: Codes für Variablen

Im digitalen Anhang auf der Daten-CD findet sich das gesammelte Material wieder, das für vorliegende Arbeit herangezogen wurde. Dies beinhaltet Audiodateien, Transkriptionen, TextGrid-Dateien aus PRAAT, sowie die Excel-Tabellen die für die statistische Auswertung verwendet wurden.

3.2.4 Vor- und Nachteile der Methode

Dieser kurze Abschnitt soll der kritischen Betrachtung meines methodischen Vorgehens gewidmet werden. Hier sollen nun einige der bereits zuvor erwähnten Arbeitsschritte reflektiert werden. Zuerst möchte ich die Arbeit mit ausschließlich gelesenen Texten diskutieren. Wie schon kurz zuvor erwähnt, war es aufgrund von nur uneinheitlich vorhandenen Datenmaterials nicht möglich, Spontansprache in die Analyse miteinzubeziehen. Um eine möglichst abgerundete Betrachtung dialektaler Merkmale zu gewährleisten, bräuchte es solches Material, da gelesene Texte nur einen Bruchteil der tatsächlichen Sprechweise der Informantinnen abbilden können (vgl. Rothstein 2011: 83).

Auch die Rekrutierung der Informantinnen erweist sich aufgrund der sehr eng gefassten Anforderungen an das Profil als sehr schwierig, weshalb der Datenpool sehr klein ist. Als nächstes möchte ich die Auswahl der Informantinnen aus dem bereits existierenden Datenpool besprechen. Bei jeder der sechs ausgewählten Städte existieren zwischen vier und sieben Informantinnen, aus denen ich jeweils drei auswählte, da eine Analyse von allen Informantinnen den Rahmen der Arbeit gesprengt hätte. Mein Ziel war es, innerhalb der Städte auch möglichst variierende Sprecherinnen auszuwählen. Dabei ignorierte ich die Tonqualität mancher Aufnahmen und begann mit der Analyse, um dann noch einmal zur Neuauswahl zurückkehren zu müssen, da ein paar wenige Aufnahmen ein Echo aufwiesen, das es, wie sich herausstellte, nicht möglich machte, Laute eindeutig zu bestimmen.

Um dies zu verhindern, hätten schon bei meiner Auswahl der Informantinnen aus dem Aufnahmekorpus zu Beginn mehrere Faktoren berücksichtigt werden müssen, da die Aufnahmen selbstverständlich nicht immer an den gleichen Orten stattfinden konnten und somit die Aufnahmequalität variiert. In puncto akustische Analyse empfiehlt es sich

des Weiteren, die phonetischen Transkriptionen von Beginn an parallel mit PRAAT zu machen, da ich bei den ersten Aufnahmen zuerst nur akustisch direkt in Word gearbeitet habe und den Text anschließend in PRAAT übertragen habe, um ihn abzugleichen. Dort wurde dann schnell klar, dass ein Großteil des Materials überarbeitet werden musste. Danach begann ich, PRAAT parallel zur Word-Transkription geöffnet zu haben und arbeitete mit dem Sound sowie der visuellen Hilfe der Sonagramme und der leichteren Steuerbarkeit durch die Segmentierung weitaus effizienter. Auch die Hilfestellung, die online in Form von IPA-Charts mit Sound zu finden ist, hätte ich von Beginn an wahrnehmen sollen, um Zeit zu sparen.

Ein weiterer Aspekt, der in Hinblick auf eine statistisch aussagekräftige Auswertung der Ergebnisse problematisch ist, ist die Unterteilung der auszuwertenden Kategorien. Durch die Aufspaltung der Ergebnisse auf Onset und Coda inklusive Verzweigung entstehen viele Variablen auf die sich dann in weiterer Folge die Anzahl der realisierten Laute aufteilt. Somit gibt es pro Kategorie nur noch wenige Okkurrenzen, was signifikante Ergebnisse nochmals weniger wahrscheinlich macht. Jedoch entschied ich mich trotzdem für jene Unterteilung, da sie für die Betrachtung der *r*-Laute abseits von statistischen Tests dennoch interessant ist.

Um nicht nur Nachteile zu nennen, soll auch ein vorteilhafter Aspekt der gewählten Methode hervorgehoben werden. Die Tatsache, dass vorgelesene Texte verwendet wurden, hatte zwar, wie schon erwähnt, auch negative Seiten, jedoch erleichtert sie die Herausbildung von relativ einheitlichen Kategorien, die in weiterer Folge zu besserer Vergleichbarkeit der Ergebnisse führen. Dadurch, dass die Informantinnen alle die gleichen Texte vorzulesen hatten, war eine voraussehbare Kategorisierung möglich. Diese Tatsache ermöglichte auch vorab die Unterteilung und Schaffung der Variablen in Onset und Coda sowie des Weiteren deren etwaige Verzweigung, da die vorgelesenen Texte, bis auf wenige Abweichungen, ident sind.

Im nun folgenden Teil der vorliegenden Arbeit sollen die Ergebnisse, die auf Basis des eben dargelegten methodischen Vorgehens entstanden sind, diskutiert, interpretiert, sowie graphisch visualisiert werden.

4. Ergebnisse und Diskussion

Wie bereits im Methodenteil erwähnt, wurden die aus den Aufnahmen gesammelten und bestimmten Daten von mir quantifiziert und im Anschluss auf Zusammenhänge und deren Signifikanz getestet. Dies soll im nun folgenden Ergebnisteil eingehend aufgeschlüsselt und diskutiert werden. Für eine möglichst gute Übersicht sind Präsentation und Diskussion nach den drei untersuchten Sprachen geordnet, beginnend mit Deutsch und seinen Dialekten, danach Französisch und im Anschluss daran Spanisch. Jede der Variablen, auf die in Verbindung mit dem Herkunftsort der Informantin getestet wurde, soll hier schrittweise beleuchtet und erklärt werden. Als erstes wurden jene Realisierungsformen getestet, die am ehesten dem jeweiligen Standard der Aussprache der Zielsprache entsprechen, anschließend wurden unerwartete Laute wie beispielsweise der entstimmte alveolare Tap [ɾ] und mögliche Interferenzen wie [ɐ] in Französisch und Spanisch in der Coda näher betrachtet. Außerdem versuchte ich, Kategorien zu erstellen, die auf die standardmäßige Ziellautung in den Fremdsprachen zugeschnitten sind und demnach mehrere Laute beinhalten.

Ein Ergebnis ist dann signifikant, wenn der Wert von $p < 0,05$ ist (vgl. Gries 2008: 155). Nicht signifikant ist, was sich darüber befindet, doch kann bei $p < 0,1$ von einer Tendenz zur Signifikanz gesprochen werden (vgl. Zöfel 2003: 95), jedoch nimmt die Wahrscheinlichkeit eines Zusammenhangs in Richtung $p < 0,1$ bereits stark ab (Methodological Education for the Social Sciences). Wenn das Ergebnis einen p-Wert unter 0,1 aufweist, wird dieser in weiterer Folge mit seinem genauen Ergebnis von mir angeführt. Jene Ergebnisse, deren p-Wert über der Grenze zu Signifikanz oder Tendenz $p > 0,1$ liegt, werden zur Vereinfachung und Beibehaltung der Übersicht lediglich mit $p > 0,05$ aufgeführt, um zu zeigen, dass sie in keiner Weise signifikant sind. Diese folgenden drei Kapitel sollen des Weiteren anhand ihrer Silbenstruktur unterteilt werden, also verzweigte und unverzweigte Onset- und Codapositionen. Zuvor werden noch Tests mit größer gefassten Kategorien besprochen. Dies hat den Grund, dass durch die Kategorisierung in verzweigte Onset- und Codaposition die Zahlen pro Kategorie noch niedriger waren und somit die Chance auf signifikante Ergebnisse abnahm (vgl. Zöfel 2003: 185). Um dies auszugleichen, sollte nochmals mit etwas größeren Kategorien gearbeitet werden. Wenn es in einer Kategorie von vorne herein zu wenig Okkurrenzen gab, konnte diese nicht getestet werden (vgl. Zöfel 2003: 185).

Im Anschluss an die Aufschlüsselung der Analyse und Signifikanztests sollen auch noch jene Ergebnisse nochmals genauer beleuchtet werden, die bei der Testung keine signifikanten Werte aufweisen konnten, bei denen es aber bei Betrachtung der absoluten Zahlen Auffälligkeiten gibt, die mir diskussionswürdig scheinen. Dazu sollen Diagramme zur besseren Veranschaulichung hinzugezogen werden.

Für eine anschließende Analyse der Daten werden die absoluten Zahlen aus den Ergebnissen herangezogen, um sehen zu können, wo diese Ergebnisse ihren Ursprung haben. Danach beleuchten wir erneut die nicht signifikanten Resultate, ziehen die absoluten Zahlen dieser heran, um zu diskutieren, ob trotz negativer Signifikanz Interesse besteht, ihr Vorkommen noch einmal näher zu betrachten.

4.1 Deutsch

Das Deutsche beziehungsweise die Daten der Dialektzonen sollen hier als erstes dargestellt werden. Sie stellen zwar für eine Analyse der fremdsprachlichen Lautung nur einen Teilaspekt dar, aber da sie, wie im Stand der Forschung ausführlich besprochen wurde, eine der Grundlagen für Transfer und mögliche Interferenzen bilden und im Falle der hier durchgeführten Forschung das Deutsche die Erstsprache der Informantinnen ist, sollen sie hier auch in der Struktur der Forschungsarbeit den ersten Abschnitt bilden. Um die verschiedenen *r*-Laute auf signifikante Zusammenhänge mit dem Herkunftsort der Informantin zu testen, wurden in SPSS eigene Kategorien geschaffen, nach dem gleichen Schema, wie es in 3.2.3.2 bereits ausführlich beschrieben wurde. Zuerst sollen allgemeine Okkurrenzen ohne Unterteilung Platz finden, danach sind die Ergebnisse nach Onset und Coda geordnet, bevor in weiterer Folge Ergebnisse von Tests mit Rücksicht auf Verzweigungen behandelt werden. Diese Kategorien und deren Ergebnisse sowie eine kurze Erläuterung folgen:

4.1.1 Vortestungen

Dieses Unterkapitel dient dem Zweck, auch Tests Raum zu geben, die nicht in die anschließenden Unterkategorien mit Silbenposition und Verzweigung aufgegliedert wurden. Dies beschloss ich aus dem Grund, dass durch die Gliederung in immer kleinere Einheiten auch die Stichprobengröße kleiner und somit anfälliger für nicht signifikante Ergebnisse war. Diese Vortestungen ohne Rücksicht auf Silbenposition und Verzweigung beinhalten gröber gefasste Kategorien.

- Da der p-Wert von Schwa [ə] auf Deutsch in Zusammenhang mit dem Ort der Informantin unter $p=0,1$ liegt, kann von einer **Tendenz zur Signifikanz** gesprochen werden:

$$\text{d.schwa: } \chi^2 (40) = 53,64; p=0,073$$

- Da der p-Wert von einer Verlängerung des Vokals in Zusammenhang mit dem Ort bei $p=0,014$ liegt, kann man hier von einem **signifikanten Ergebnis** sprechen:

$$\text{d.vocv: } \chi^2 (20) = 36,364; p=0.014$$

- Da der p-Wert des stimmhaften alveolaren Taps [ɾ] in Zusammenhang mit dem Ort der Informantin $p>0,05$ ist, liegt **kein signifikantes Ergebnis** vor:

$$\text{d.1alvtap: } \chi^2 (16) = 20 ; p>0,05$$

- Da der p-Wert für die Variable bestehend aus den gesammelten uvularen Realisierungsmöglichkeiten ([ʁ], [χ], [R], [ʀ] sowie [x]) in Zusammenhang mit dem Ort der Informantin $p>0,05$ ist, liegt **kein signifikantes Ergebnis** vor:

$$\text{d.uvular: } \chi^2 (4) = 5 ; p>0,05$$

4.1.2 Onset

- Da der p-Wert des stimmhaften alveolaren Taps [ɾ] im Onset in Zusammenhang mit dem Ort der Informantin $p>0,05$ ist, liegt **kein signifikantes Ergebnis** vor:

$$\text{d.On_1alvtap: } \chi^2 (6) = 7,5 ; p>0,05$$

- Da der p-Wert des stimmhaften uvularen Frikativs [ʁ] im Onset in Zusammenhang mit dem Ort der Informantin $p>0,05$ ist, liegt **kein signifikantes Ergebnis** vor:

$$\text{d.On_1uvfric: } \chi^2 (16) = 19,467 ; p>0,05$$

- Da der p-Wert des stimmlosen uvularen Frikativs [χ] im Onset in Zusammenhang mit dem Ort der Informantin $p > 0,05$ ist, liegt **kein signifikantes Ergebnis** vor:

$$\text{d.On_0uvfric: } \chi^2 (10) = 11,762 ; p > 0,05$$

- Da der p-Wert des stimmhaften uvularen Approximanten [ʁ] im Onset in Zusammenhang mit dem Ort der Informantin $p > 0,05$ ist, liegt **kein signifikantes Ergebnis** vor:

$$\text{d.On_1uvapprox: } \chi^2 (30) = 31 ; p > 0,05$$

- Da der p-Wert des stimmhaften uvularen Vibranten [R] im Onset in Zusammenhang mit dem Ort der Informantin $p > 0,05$ ist, liegt **kein signifikantes Ergebnis** vor:

$$\text{d.On_1uvvibr: } \chi^2 (4) = 5 ; p > 0,05$$

4.1.2.1 verzweigt

Auf die Realisierungsformen [ʁ], [ʁ̥] und [R] konnte nicht getestet werden, da zu wenige Okkurrenzen vorhanden waren. Die restlichen interessanten Laute ergaben folgende Werte:

- Da der p-Wert des stimmlosen uvularen Frikativs [χ] im verzweigten Onset in Zusammenhang mit dem Ort der Informantin $p > 0,05$ ist, liegt **kein signifikantes Ergebnis** vor:

$$\text{d.OnVz_0uvfric: } \chi^2 (25) = 19,907 ; p > 0,05$$

- Da der p-Wert des stimmhaften alveolaren Taps [ɾ] im verzweigten Onset in Zusammenhang mit dem Ort der Informantin $p > 0,05$ ist, liegt **kein signifikantes Ergebnis** vor:

$$\text{d.OnVz_1alvtap: } \chi^2 (4) = 6,667 ; p > 0,05$$

4.1.2.2 unverzweigt

- Da der p-Wert des stimmhaften uvularen Frikativs [ʁ] im unverzweigten Onset in Zusammenhang mit dem Ort der Informantin $p > 0,05$ ist, liegt **kein signifikantes Ergebnis** vor:

$$\text{d.OnUvz_1uvfric: } \chi^2 (12) = 15,375 ; p > 0,05$$

- Da der p-Wert des stimmlosen uvularen Frikativs [χ] im unverzweigten Onset in Zusammenhang mit dem Ort der Informantin $p > 0,05$ ist, liegt **kein signifikantes Ergebnis** vor:

$$\text{d.OnUvz_0uvfric: } \chi^2 (45) = 40,368 ; p > 0,05$$

- Da der p-Wert des stimmhaften uvularen Approximanten [ʁ̥] im unverzweigten Onset in Zusammenhang mit dem Ort der Informantin $p > 0,05$ ist, liegt **kein signifikantes Ergebnis** vor:

$$\text{d.OnUvz_1uvapprox: } \chi^2 (30) = 31 ; p > 0,05$$

- Da der p-Wert des stimmhaften uvularen Vibranten [ʀ] im unverzweigten Onset in Zusammenhang mit dem Ort der Informantin $p > 0,05$ ist, liegt **kein signifikantes Ergebnis** vor:

$$\text{d.OnUvz_1uvvibr: } \chi^2 (2) = 3 ; p > 0,05$$

- Da der p-Wert des stimmhaften alveolaren Taps [ɾ] im unverzweigten Onset in Zusammenhang mit dem Ort der Informantin $p > 0,05$ ist, liegt **kein signifikantes Ergebnis** vor:

$$\text{d.OnUvz_1alvtap: } \chi^2 (2) = 5 ; p > 0,05$$

4.1.3 Coda

Die konsonantischen Realisierungsformen [ʀ], [ʁ], [χ] und [ʁ̥] in der deutschen Coda waren zahlenmäßig zu begrenzt, um Tests durchführen zu können (vgl. Zöfel 2003: 185). Die Ergebnisse jener Überprüfungen, die möglich waren, lauten wie folgt:

- Da der p-Wert des stimmhaften alveolaren Taps [r] in der Coda in Zusammenhang mit dem Ort der Informantin $p > 0,05$ ist, liegt **kein signifikantes Ergebnis** vor:

$$\text{d.Co_1alvtap: } \chi^2 (4) = 7 ; p > 0,05$$

4.1.3.1 verzweigt

Auch in der verzweigten Coda konnte auf bestimmte lautliche Realisierungen nicht getestet werden, da die Okkurrenzen fehlten. Dazu zählen hier die Laute [ɸ], [ɸ̥], [ʀ] und [χ].

- Da der p-Wert von des stimmhaften alveolaren Taps [r] in der verzweigten Coda in Zusammenhang mit dem Ort der Informantin $p > 0,05$ ist, liegt **kein signifikantes Ergebnis** vor:

$$\text{d.CoVz_1alvtap: } \chi^2 (4) = 7 ; p > 0,05$$

- Da der p-Wert von Schwa [ə] in der verzweigten Coda in Zusammenhang mit dem Ort der Informantin $p = 0,065$ ist, kann man von einer **Tendenz zur Signifikanz** sprechen:

$$\text{d.CoVz_schwa: } \chi^2 (20) = 30,324 ; p = 0,065$$

- Da der p-Wert von einer Verlängerung des Vokals in der verzweigten Coda in Zusammenhang mit dem Ort der Informantin $p > 0,05$ ist, liegt **kein signifikantes Ergebnis** vor:

$$\text{d.CoVz_vocv: } \chi^2 (2) = 1,333 ; p > 0,05$$

4.1.3.2 unverzweigt

In der unverzweigten Coda wurde ebenfalls nur auf jene lautlichen Realisierungen getestet, die genügend Okkurrenzen aufweisen konnten. Zu wenige beziehungsweise keine Okkurrenzen kamen bei [ɸ], [ɸ̥], [ʀ] und [χ] vor.

- Da der p-Wert von Schwa [ɐ] in der unverzweigten Coda in Zusammenhang mit dem Ort der Informantin $p > 0,05$ ist, liegt **kein signifikantes Ergebnis** vor:

$$\text{d.CoUvz_schwa: } \chi^2 (35) = 42,815 ; p > 0,05$$

- Da der p-Wert von einer Verlängerung des Vokals in der unverzweigten Coda in Zusammenhang mit dem Ort der Informantin $p > 0,05$ ist, liegt **kein signifikantes Ergebnis** vor:

$$\text{d.CoUvz_vocv: } \chi^2 (5) = 7,248 ; p > 0,05$$

4.2 Französisch

Die Tests, die für die Ermittlung etwaiger Signifikanzen zwischen Dialekt und Französischer *r*-Lautung durchgeführt wurden, verliefen ebenfalls nach dem Schema, das bereits in Kapitel 4.1 besprochen wurde. Nachdem auf allgemeine Okkurrenzen ohne Verzweigung, sowie die Kategorie *f.standard* getestet wurde, gab es noch den Versuch, in die Richtung einzelner Laute nach Silbenposition und Verzweigung Untersuchungen anzustellen. Wie schon im Methodenkapitel erwähnt, beinhaltet die Kategorie *f.standard* den stimmhaften uvularen Frikativ [ʁ] und den stimmlosen uvularen Frikativ [χ], da [ʁ] als Standard verbreitet ist und [χ] sehr häufig vorkommt (siehe Kapitel 2.5.4.). Die Bezeichnung durch ‚Standard‘ wurde mit Blick auf möglichst einfache Benennung der Variablen kreiert.

4.2.1 Vortestungen

- Da der p-Wert von Schwa [ɐ] in Zusammenhang mit dem Ort der Informantin $p > 0,05$ ist, liegt **kein signifikantes Ergebnis** vor:

$$\text{f.schwa: } \chi^2 (35) = 32,139 ; p > 0,05$$

- Da der p-Wert des stimmhaften uvularen Frikativs [ʁ] in Zusammenhang mit dem Ort der Informantin $p > 0,05$ ist, liegt **kein signifikantes Ergebnis** vor:

$$\text{f.1uvfric: } \chi^2 (30) = 34,125 ; p > 0,05$$

- Da der p-Wert des stimmlosen uvularen Frikativs [χ] in Zusammenhang mit dem Ort der Informantin $p > 0,05$ ist, liegt **kein signifikantes Ergebnis** vor:

$$\text{f.0uvfric: } \chi^2 (60) = 69,333 ; p > 0,05$$

- Da der p-Wert von einer Verlängerung des Vokals in Zusammenhang mit dem Ort der Informantin $p > 0,05$ ist, liegt **kein signifikantes Ergebnis** vor:

$$\text{f.vocv: } \chi^2 (10) = 12,5 ; p > 0,05$$

- Da der p-Wert für die Variable *f.standard*, bestehend aus stimmhaftem und stimmlosem uvularen Frikativ [ʁ] und [χ], in Zusammenhang mit dem Ort der Informantin $p > 0,05$ ist, liegt **kein signifikantes Ergebnis** vor:

$$\text{f.standard: } \chi^2 (60) = 62,5 ; p > 0,05$$

4.2.2 Onset

- Da der p-Wert des stimmhaften uvularen Frikativs [ʁ] im Onset in Zusammenhang mit dem Ort der Informantin $p > 0,05$ ist, liegt **kein signifikanten Ergebnis** vor.

$$\text{f.On_1uvfric: } \chi^2 (40) = 39,825 ; p > 0,05$$

- Da der p-Wert des stimmlosen uvularen Frikativs [χ] im Onset in Zusammenhang mit dem Ort der Informantin $p > 0,05$, liegt **kein signifikanten Ergebnis** vor.

$$\text{f.On_0uvfric: } \chi^2 (50) = 50,22 ; p > 0,05$$

- Da der p-Wert des stimmhaften uvularen Approximanten [ʁ] im Onset in Zusammenhang mit dem Ort der Informantin $p > 0,05$ ist, liegt **kein signifikanten Ergebnis** vor.

$$\text{f.On_1uvapprox: } \chi^2 (15) = 8,85 ; p > 0,05$$

4.2.2.1 verzweigt

Für die in dieser Kategorie interessanten Realisierungsformen gab es durch die Unterteilung ebenfalls zu wenige beziehungsweise keine Okkurrenzen, somit erwiesen sich Überprüfungen nach Zusammenhängen der Laute [ʁ] [χ] und [ʁ̥] mit dem Ort der Informantinnen als nicht durchführbar.

- Da der p-Wert des stimmhaften alveolaren Taps [r] im verzweigten Onset in Zusammenhang mit dem Ort der Informantin $p > 0,05$ ist, , **liegt kein signifikantes Ergebnis** vor:

$$\text{f.OnVz_1alvtap: } \chi^2 (1) = 2 ; p > 0,05$$

4.2.2.2 unverzweigt

- Da der p-Wert des stimmhaften uvularen Frikativs [ʁ] im unverzweigten Onset in Zusammenhang mit dem Ort der Informantin $p = 0,039$ ist, kann von einem **signifikanten Ergebnis** gesprochen werden.

$$\text{f.OnUvz_1uvfric: } \chi^2 (30) = 45 ; p = 0,039$$

- Da der p-Wert des stimmlosen uvularen Frikativs [χ] im unverzweigten Onset in Zusammenhang mit dem Ort der Informantin $p > 0,05$ ist, **liegt kein signifikantes Ergebnis** vor:

$$\text{f.OnUvz_0uvfric: } \chi^2 (45) = 43,167 ; p > 0,05$$

- Da der p-Wert des stimmhaften uvularen Approximanten [ʁ̥] im unverzweigten Onset in Zusammenhang mit dem Ort der Informantin $p > 0,05$ ist, liegt **kein signifikantes Ergebnis** vor:

$$\text{f.OnUvz_1uvapprox: } \chi^2 (15) = 8,85 ; p > 0,05$$

4.2.3 Coda

- Da der p-Wert von Schwa [ə] in der Coda in Zusammenhang mit dem Ort der Informantin $p > 0,05$ ist, liegt **kein signifikantes Ergebnis** vor:

$$f.co_schwa: \chi^2 (35) = 32,139 \text{ } p>0,05$$

- Da der p-Wert des stimmhaften uvularen Frikativs [ʁ] in der Coda in Zusammenhang mit dem Ort der Informantin $p>0,05$ ist, liegt **kein signifikantes Ergebnis** vor:

$$f.co_1uvfric: \chi^2 (30) = 35,96 \text{ ; } p>0,05$$

- Da der p-Wert des stimmlosen uvularen Frikativs [χ] in der Coda in Zusammenhang mit dem Ort der Informantin $p>0,05$ ist, liegt **kein signifikantes Ergebnis** vor:

$$f.co_0uvfric: \chi^2 (20) = 28,007 \text{ ; } p>0,05$$

- Da der p-Wert des stimmhaften uvularen Approximanten [ʁ̥] in der Coda in Zusammenhang mit dem Ort der Informantin $p>0,05$ ist, liegt **kein signifikantes Ergebnis** vor:

$$f.co_1uvapprox: \chi^2 (20) = 28,067 \text{ ; } p>0,05$$

4.2.3.1 verzweigt

In der verzweigten Coda begegnen wir ebenfalls der Problematik von zu wenigen beziehungsweise nicht vorhandenen Okkurrenzen für die Laute [ʁ], [χ], [ʁ̥] sowie [ə]. Daher waren Tests nach diesen Lauten auch hier nicht durchführbar.

4.2.3.2 unverzweigt

- Da der p-Wert des stimmhaften uvularen Frikativs [ʁ] in der unverzweigten Coda in Zusammenhang mit dem Ort der Informantin $p>0,05$ ist, liegt **kein signifikantes Ergebnis** vor:

$$f.CoUvz_1uvfric: \chi^2 (16) = 17,008 \text{ ; } p>0,05$$

- Da der p-Wert des stimmlosen uvularen Frikativs [χ] in der unverzweigten Coda in Zusammenhang mit dem Ort der Informantin $p>0,05$ ist, liegt **kein signifikantes Ergebnis** vor:

f.CoUvz_0uvfric: $\chi^2 (40) = 45$; $p > 0,05$

- Da der p-Wert des stimmhaften uvularen Approximanten [ɣ] in der unverzweigten Coda in Zusammenhang mit dem Ort der Informantin $p > 0,05$ ist, liegt **kein signifikantes Ergebnis** vor:

f.CoUvz_1uvapprox: $\chi^2 (12) = 12,933$; $p > 0,05$

- Da der p-Wert von Schwa [ə] in der unverzweigten Coda in Zusammenhang mit dem Ort der Informantin $p > 0,05$ ist, liegt **kein signifikantes Ergebnis** vor.

f.CoUvz_schwa: $\chi^2 (40) = 40,806$; $p > 0,05$

- Da der p-Wert der Verlängerung des Vokals in der unverzweigten Coda in Zusammenhang mit dem Ort der Informantin $p > 0,05$ ist, liegt **kein signifikantes Ergebnis** vor.

f.CoUvz_vocv: $\chi^2 (10) = 12,5$; $p > 0,05$

4.3 Spanisch

Auch bei den spanischen Ergebnissen wurde nach dem gleichen Schema getestet, so gab es hier einen Versuch, die Kategorie ‚Spanisch Standard‘ mit der Variablenbezeichnung *s.standard* zu generieren, welche die beiden Phoneme /r/ und /r/ beinhaltet. Wie aber schon im Methodenskapitel unter 3.2.3.2. erwähnt, wurde in keiner der spanischen Aufnahmen [r] realisiert und somit beschränkt sich diese Kategorie auf eine Testung nach [r] in allen Positionen. Danach wurde, wie bei den Resultaten der Französischen Datensammlung, auf einzelne Laute getestet. Die Ergebnisse der Signifikanztests mit χ^2 lauten wie folgt:

4.3.1 Vortestungen

Die Variable *s.standard* wurde aufgrund fehlender Okkurrenz des alveolaren Trills [r] nicht eigens getestet. Es wurden daher lediglich Zusammenhänge zwischen Herkunftsort der Informantin und Realisierung des alveolaren Taps [r], sowie weitere interessante Realisierungen überprüft.

- Da der p-Wert des stimmhaften alveolaren Taps [ɾ] in Zusammenhang mit dem Ort der Informantin $p=0,077$ ist, kann man von einer **Tendenz zur Signifikanz** sprechen:

$$\text{s.1alvtap: } \chi^2 (40) = 53,333 ; p=0,077$$

- Da der p-Wert von Schwa [ə] in Zusammenhang mit dem Ort der Informantin $p=0,008$ ist, kann man hier von einem **signifikanten Ergebnis** sprechen:

$$\text{s.schwa: } \chi^2 (30) = 52 ; p=0,008$$

- Da der p-Wert der Verlängerung des Vokals in Zusammenhang mit dem Ort der Informantin $p>0,05$ ist, liegt **kein signifikantes Ergebnis** vor:

$$\text{s.vocv: } \chi^2 (2) = 5 ; p>0,05$$

- Da der p-Wert des stimmhaften uvularen Frikativs [ʁ] in Zusammenhang mit dem Ort der Informantin $p>0,05$ ist, liegt **kein signifikantes Ergebnis** vor:

$$\text{s.1uvfric: } \chi^2 (40) = 46,2 ; p>0,05$$

- Da der p-Wert des stimmlosen uvularen Frikativs [χ] in Zusammenhang mit dem Ort der Informantin $p>0,05$ ist, liegt **kein signifikantes Ergebnis** vor:

$$\text{s.0uvfric: } \chi^2 (15) = 18,083 ; p>0,05$$

- Da der p-Wert des stimmhaften uvularen Vibranten [ʀ] in Zusammenhang mit dem Ort der Informantin $p>0,05$ ist, liegt **kein signifikantes Ergebnis** vor:

$$\text{s.1uvvibr: } \chi^2 (3) = 4 ; p>0,05$$

- Da der p-Wert des stimmhaften uvularen Approximanten [ɣ] in Zusammenhang mit dem Ort der Informantin $p > 0,05$ ist, liegt **kein signifikantes Ergebnis** vor:

$$\text{s.1uvapprox: } \chi^2 (15) = 18,08 ; p > 0,05$$

4.3.2 Onset

- Da der p-Wert des stimmhaften alveolaren Taps [ɾ] im Onset in Zusammenhang mit dem Ort der Informantin $p > 0,05$ ist, liegt **kein signifikantes Ergebnis** vor:

$$\text{s.on_1alvtap: } \chi^2 (15) = 20,1 ; p > 0,05$$

- Da der p-Wert des stimmhaften uvularen Vibranten [ʀ] im Onset in Zusammenhang mit dem Ort der Informantin $p > 0,05$ ist, liegt **kein signifikantes Ergebnis** vor:

$$\text{s.on_1uvvibr: } \chi^2 (8) = 7,5 ; p > 0,05$$

- Da der p-Wert des stimmhaften uvularen Frikativs [ʁ] im Onset in Zusammenhang mit dem Ort der Informantin $p > 0,05$ ist, liegt **kein signifikantes Ergebnis** vor:

$$\text{s.on_1uvfric: } \chi^2 (18) = 17,25 ; p > 0,05$$

- Da der p-Wert des stimmlosen uvularen Frikativs [χ] im Onset in Zusammenhang mit dem Ort der Informantin $p > 0,05$ ist, liegt **kein signifikantes Ergebnis** vor:

$$\text{s.on_0uvfric: } \chi^2 (18) = 23,25 ; p > 0,05$$

- Da der p-Wert des stimmhaften uvularen Approximanten [ɣ] im Onset in Zusammenhang mit dem Ort der Informantin $p > 0,05$ ist, liegt **kein signifikantes Ergebnis** vor:

$$\text{s.on_1uvapprox: } \chi^2 (15) = 20,25 ; p > 0,05$$

4.3.2.1 verzweigt

Für die Überprüfung auf Zusammenhänge zwischen Ort und den Lauten [ɸ], [ɸ̥] sowie [ʀ] fehlte auch in diesem Abschnitt wieder eine ausreichende Anzahl an Okkurrenzen. Die übrigen Tests ergaben folgende Werte:

- Da der p-Wert des stimmhaften alveolaren Taps [r] im verzweigten Onset in Zusammenhang mit dem Ort der Informantin $p > 0,05$ ist, liegt **kein signifikantes Ergebnis** vor:

s.OnVz_1alvtap: $\chi^2 (12) = 15,714$; $p > 0,05$

- Da der p-Wert des stimmlosen uvularen Frikativs [χ] im verzweigten Onset in Zusammenhang mit dem Ort der Informantin $p > 0,05$ ist, liegt **kein signifikantes Ergebnis** vor:

s.OnVz_0uvfric: $\chi^2 (9) = 11,86$; $p > 0,05$

4.3.2.2 unverzweigt

- Da der p-Wert des stimmhaften alveolaren Taps [r] im unverzweigten Onset im unverzweigten Onset in Zusammenhang mit dem Ort der Informantin $p = 0,022$ ist, kann von einem **signifikanten Ergebnis** gesprochen werden.

s.OnUvz_1alvtap: $\chi^2 (10) = 20,857$; $p = 0,022$

- Da der p-Wert des stimmhaften uvularen Vibranten [ʀ] im verzweigten Onset in Zusammenhang mit dem Ort der Informantin $p > 0,05$ ist, liegt **kein signifikantes Ergebnis** vor:

s.OnUvz_1uvvibr: $\chi^2 (3) = 4$; $p > 0,05$

- Da der p-Wert des stimmhaften uvularen Approximanten [ɹ] im verzweigten Onset in Zusammenhang mit dem Ort der Informantin $p > 0,05$ ist, liegt **kein signifikantes Ergebnis** vor:

s.OnUvz_1uvapprox: $\chi^2 (12) = 18,083$; $p > 0,05$

4.3.3 Coda

Aufgrund zu geringen Vorkommens der Laute [ɸ], [ɸ̥] und [ʀ] wurden diese auch in der spanischen Coda nicht getestet. Die Laute, deren Zusammenhang überprüft wurde, lieferten folgende Ergebnisse:

- Da der p-Wert des stimmhaften alveolaren Taps [r] in der Coda in Zusammenhang mit dem Ort der Informantin $p > 0,05$ ist, liegt **kein signifikantes Ergebnis** vor:

$$\text{s.Co_1alvtap: } \chi^2 (45) = 55 ; p > 0,05$$

- Da der p-Wert des stimmlosen uvularen Frikativs [χ] in der Coda in Zusammenhang mit dem Ort der Informantin $p > 0,05$ ist, liegt **kein signifikantes Ergebnis** vor:

$$\text{s.Co_0uvfric: } \chi^2 (6) = 5 ; p > 0,05$$

4.3.3.1 verzweigt

Im spanischen Text, der von den Informantinnen vorzulesen war, gab es keine Okkurrenzen verzweigter *r*-Laute in der Coda. Somit konnten mit dieser Kategorie für Spanische keine Tests durchgeführt werden.

4.3.3.2 unverzweigt

- Da der p-Wert des stimmhaften alveolaren Taps [r] in der unverzweigten Coda in Zusammenhang mit dem Ort der Informantin $p > 0,05$ ist, liegt **kein signifikantes Ergebnis** vor:

$$\text{s.CoUvz_1alvtap: } \chi^2 (32) = 39,417 ; p > 0,05$$

- Da der p-Wert von Schwa [ə] in der unverzweigten Coda in Zusammenhang mit dem Ort der Informantin mit $p = 0,055$ knapp über 0,05 liegt, kann von einer **Tendenz zur Signifikanz** gesprochen werden.

$$\text{s.CoUvz_schwa: } \chi^2 (20) = 31 ; p = 0,055$$

- Da der p-Wert der Verlängerung des Vokals in der unverzweigten Coda in Zusammenhang mit dem Ort der Informantin bei $p=0,05$ liegt, kann von einem **signifikanten Ergebnis** gesprochen werden.

$$\text{s.CoUvz_vocv: } \chi^2 (2) = 6 ; p=0,050$$

- Da der p-Wert des stimmlosen uvularen Frikativs [χ] in der unverzweigten Coda in Zusammenhang mit dem Ort der Informantin $p>0,05$ ist, liegt **kein signifikantes Ergebnis** vor:

$$\text{s.CoUvz_0uvfric: } \chi^2 (4) = 4; p>0,05$$

5. Diskussion

Wie in Kapitel 4. ausführlich dargeboten wurde, stellte sich eine signifikante statistische Analyse als schwierig heraus. Dies lag nicht zuletzt an der relativ kleinen Datenmenge, die durch die Aufspaltung in einzelne Kategorien wie verzweigte und unverzweigte Silbenposition nochmals in kleinere Einheiten aufgesplittet wurde. Somit halten sich statistisch aussagekräftige Ergebnisse zwar in Grenzen, jedoch lässt sich bei genauerer Betrachtung der Ergebnisse in Zahlen feststellen, dass durchaus interessante Verteilungen der *r*-Laute zu sehen sind. In diesem Kapitel, welches die Auswertung der Ergebnisse nochmals abseits von statistischen Tests darstellen und diskutieren soll, finden nun jene Beobachtungen Platz, sowie eine nochmalige genauere Erläuterung jener Ergebnisse, die einen signifikanten Zusammenhang aufweisen konnten.

5.1 Deutsch

Im Deutschen ließen sich vor allem interessante Beobachtungen anstellen, die sich auf die Vokalisierung von *r*-Lauten beziehen. So konnte ich beispielsweise feststellen, dass die aus der Schweiz stammenden Informantinnen weniger dazu tendierten, [ɐ] anstelle eines *r*-Lautes in der Coda zu realisieren und stattdessen ein konsonantisches *r* aussprachen. Dies spiegelt sich sogar in den bereits besprochenen signifikanten Ergebnissen unter Kapitel 4.1.1 in Form einer Tendenz zur Signifikanz wieder. In Zahlen äußert sich diese Beobachtung wie folgt:

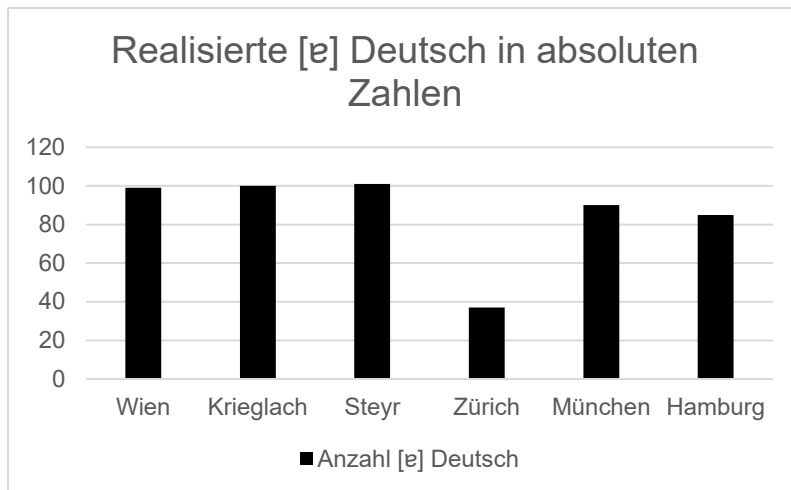


Abb. 19: Schwa [ɐ] im deutschen Korpus

Dieser Unterschied deckt sich mit der bereits diskutierten Literatur zur Aussprache von *r*-Lauten im deutschen Sprachraum, die ebenfalls eine geographische Verteilung der [ɐ]-Realisierung anstatt eines Konsonanten feststellt (vgl. Panizzolo 1982: 30), (vgl. Simpson 1998: 92). Während in den bairischen Dialektzonen auf Deutsch jeweils um die 100 [ɐ] pro Ort in der Coda realisiert wurden und in Hamburg ebenfalls über 80, wurden in der Schweiz weniger als die Hälfte davon, insgesamt nur 37 produziert. Diese Zahlen beziehen sich hier immer auf die gesammelten Realisierungen aller Informantinnen pro Ort.

Eine weitere Beobachtung betrifft die Vokalverlängerung in Kombination mit Wegfall eines konsonantischen *r*. Das signifikante Ergebnis hierzu aus 4.1.1 lässt sich in absoluten Zahlen so begründen, dass bei den Informantinnen aus Hamburg häufiger zu dieser Realisierung tendiert wurde, als bei den restlichen Sprecherinnen:

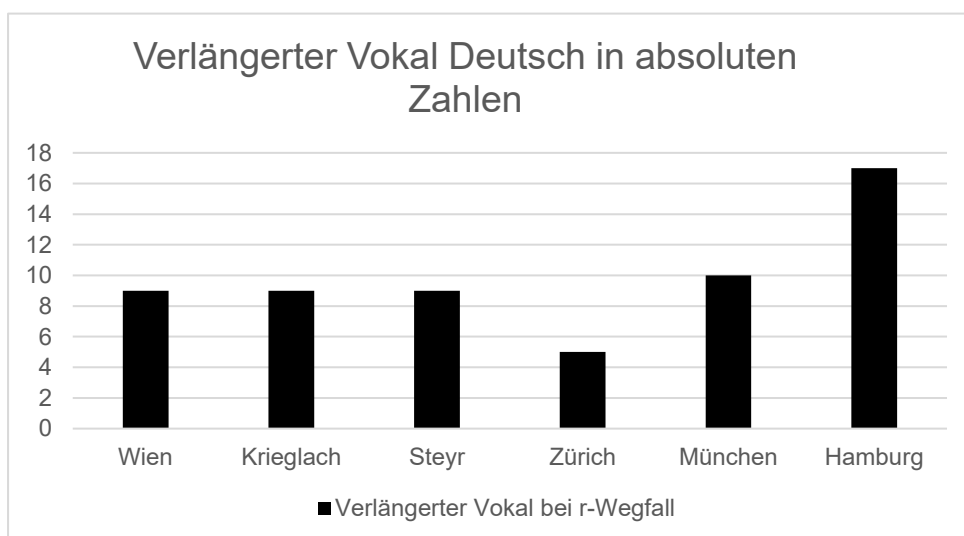


Abb. 20: Vokalverlängerung bei *r*-Wegfall im deutschen Korpus

In Hamburg wurde beispielsweise fast doppelt so oft auf genannte Realisierungsmöglichkeit zurückgegriffen, wie jeweils in den drei österreichischen Orten. Dies fand vor

allem in verzweigter Codaposition statt, denn das Wort *Nordwind*, welches *r* in verzweigter Silbencoda beinhaltet, kam im Text insgesamt vier Mal vor. Die meisten Informantinnen tendierten hier entweder zu [ɐ] oder, wie beispielsweise bei den Schweizerinnen beobachtbar, zu [r] oder einer anderen konsonantischen Realisierung von *r*, während die Hamburgerinnen öfter weder Konsonant noch Schwa [ɐ] realisierten und stattdessen den vorangehenden Vokal [ɔ] in die Länge zogen. Ein Beispiel liefert uns die Informantin HH_03 aus der deutschen Aufnahme in folgender Abbildung:

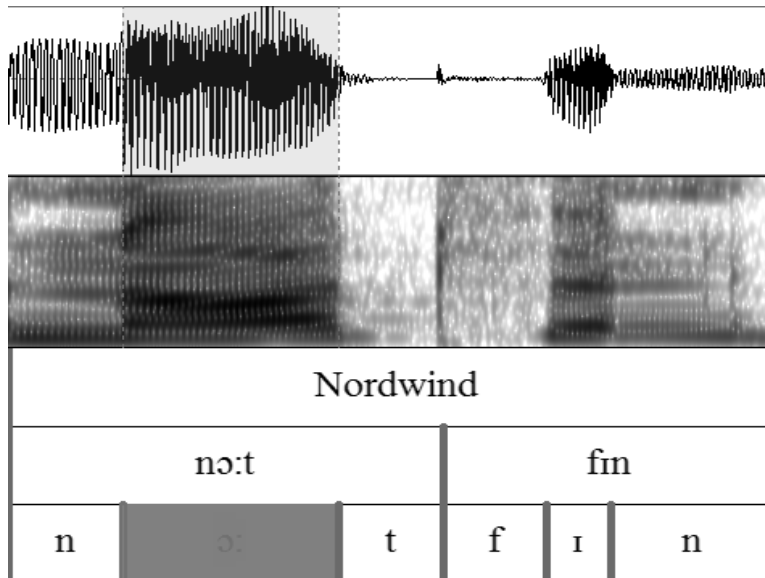


Abb. 21: Informantin HH_03 realisiert verlängerten Vokal [ɔ] bei *r*-Wegfall

Zu leider keinem signifikanten Ergebnis kam die Testung nach Zusammenhang von Realisierung des alveolaren Tap [r] mit dem Ort oder der Dialektzone der Informantin. Jedoch kann man auch hier wieder anhand der Zahlen eine gewisse Verteilung dieser Realisierungsform beobachten:

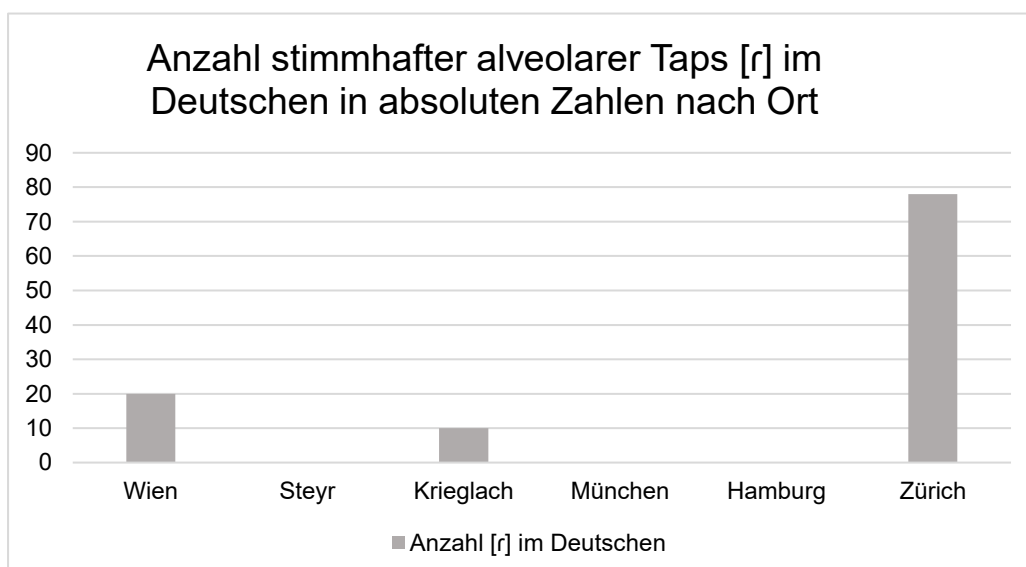


Abb. 22: Stimmhafter alveolarer Tap [r] nach Ort

Es lässt sich anhand der Grafik in Abbildung 22 erkennen, dass in Zürich, sowie Wien und der Steiermark auf Deutsch die höchste Anzahl an alveolaren Taps [ɾ] realisiert wurde, wohingegen in Oberösterreich, München und Hamburg keine einzige Okkurrenz zu finden war. Die höchste Anzahl an realisierten alveolaren Taps [ɾ] findet sich wenig überraschend in der Schweiz, da, wie schon weiter oben behandelt, die Sprecherinnen in diesem Ort auch mehr dazu tendierten, anstatt eines Schwa [ə] in der Coda ebenfalls ihre gewohnte konsonantische *r*-Realisierung auszusprechen. Oberösterreich und München fallen aus der Reihe, da sie der Literatur zur deutschen Dialektlautung zufolge eigentlich ebenfalls zur bairischen Dialektzone gehören, und somit zum alveolaren Tap [ɾ] tendieren sollten (vgl. Kapitel 2.5.2), dies aber bei den Informantinnen im Korpus der in dieser Forschungsarbeit verwendet wurde, nicht der Fall war.

Eine weitere Auffälligkeit im Deutschen betrifft das Vorkommen alveolarer *r*-Laute. In der Literatur werden sie als Allophone oft in Verbindung mit bestimmten Dialekten angeführt, wie wir schon in Kapitel 2.5.2 sehen konnten. Interessant ist hierbei, dass in der Literatur über deutsche *r*-Laute häufiger der alveolare Trill [r] angesprochen wird als der alveolare Tap [ɾ]. Wirft man nun allerdings einen Blick auf den deutschen Korpus der „Nordwind und Sonne“ – Aufnahmen, fällt schnell auf, dass jene Informantinnen, die zur alveolaren *r*-Realisierung greifen, ausschließlich den stimmhaften Tap [ɾ], die entstimmte Version [ɾ̥] oder einen alveolaren Flap [ɾ̥] sowie dessen entstimmte Variante [ɾ̥̥] aussprechen, niemals den mehrschlägigen Trill [r].

Generell wurde auch bei uvularen Realisierungsformen wenig auf Trills wie beispielsweise [ʀ] zurückgegriffen, jener Laut kommt in den deutschen Aufnahmen insgesamt nur zehn Mal vor, das immer im Onset. Wie schon in Kapitel 2.5.2 erwähnt, wird beispielsweise im Artikel von Wiese (2001) der stimmhafte uvulare Approximant [ʁ] als der deutsche Standard beschrieben. Dieser konnte im Korpus tatsächlich relativ häufig beobachtet werden. Am häufigsten produzierten den Approximanten [ʁ] die Münchnerinnen mit insgesamt 13 Okkurrenzen, anschließend daran die Hamburgerinnen mit gesamt zwölf. Im Anschluss daran reihen sich Wien mit acht, Krieglach mit sieben, Steyr mit sechs und Zürich mit fünf Okkurrenzen. Außerdem fand sich im deutschen Korpus dieser Laut bis auf eine Ausnahme ausschließlich im unverzweigten Silbenonset. Die Ausnahme machte eine Okkurrenz in Zürich im verzweigten Onset, die restlichen vier wie in den anderen Aufnahmeorten aber unverzweigt.

Um nun abschließend noch einen Überblick über die Verteilung der verschiedensten Realisierungen im deutschen Korpus geben zu können, zeigt folgender Überblick die häufigsten konsonantischen sowie vokalischen *r*-Realisierungsmöglichkeiten, die zu realisieren waren, in absoluten Zahlen:

Laut	Okkurrenzen
[ə]	512
[r]	108
Vokalverlängerung	62
[ʁ]	55
[x]	39
[ʁ]	34
Elision	25
[R]	9
[r̥]	6
[X]	4
[r̥]	1
?	1
[r]	0
[r̥]	0
[r ^h]	0

Tabelle 3: Gesamtzahl Okkurrenzen Deutsch

Die häufigsten konsonantischen Realisierungsformen, also exklusive Schwa [ə] und Vokalverlängerung, finden sich in folgender Karte nach Orten kategorisiert:

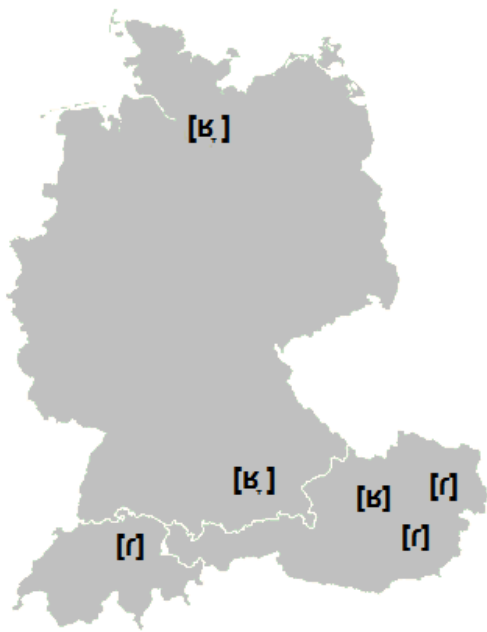


Abb. 23 : Häufigste konsonantische Realisierungen im deutschen Korpus nach Ort

Anhand dieser grundlegenden Information über das Verhalten der Informantinnen in ihrer L1 können in den nun folgenden Kapiteln 5.2 und 5.3 auch weiterführende Beobachtungen in den Fremdsprachen angestellt werden.

5.2 Französisch

Auch im Französischen konnten anhand der vorhandenen Ergebnisse in Zahlen einige interessante Beobachtungen angestellt werden. Eine interessante Tendenz zur *r*-Lautung deutschsprachiger Lernerinnen im Französischen war dort zu beobachten, wo in der Quellsprache Deutsch eher alveolare Realisierungen wie der Tap [r] stattfanden. Ein Teil jener Informantinnen tendierte nämlich im Französischen dazu, ähnlich wie im Deutschen, *r*-Laute in der Coda als Schwa [ə] zu realisieren, anstatt ihren gewohnten konsonantischen *r*-Laut auch an dieser Stelle zu übernehmen. Diese Beobachtung schließt jedoch die Informantinnen aus der Schweiz aus, da sie, wie wir schon im Stand der Forschung besprochen und in 5.1 anhand von Beispielen gesehen haben, nicht zur Coda-Realisierung mit [ə] tendieren, sondern ebenfalls ihren entsprechenden konsonantischen *r*-Laut produzieren. Informantinnen, die auf Deutsch eher uvulare Realisierungen produzieren, führten diese Substitution durch Schwa [ə] seltener durch, wie

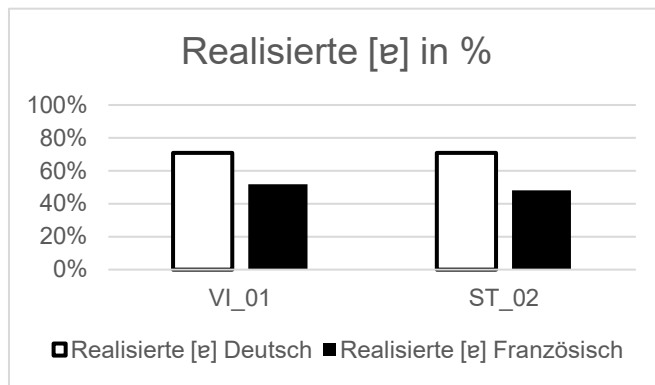


Abb. 24: Zwei Informantinnen und deren prozentuelle Realisierungen von [ɐ] im deutschen und französischen Korpus

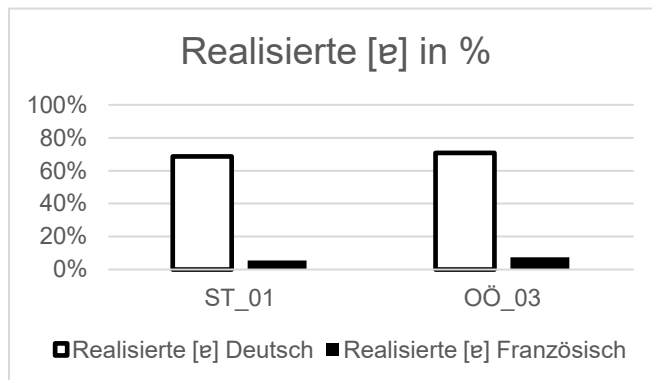


Abb. 25: Zwei Informantinnen und deren prozentuelle Realisierungen von [ɐ] im deutschen und französischen Korpus

Taps [r] aussprachen. Jene, die zu einer uvularen Realisierungsform tendieren, weisen auch weniger Schwas [ɐ] in der französischen Coda auf. Eine Visualisierung jener oben beschriebenen Informantinnen und deren Aussprachemuster im Deutschen findet sich in folgender beispielhafter Grafik:

die Zahlen in folgenden vergleichenden Grafiken in den Abbildungen 24 und 25 verraten:

Die Informantinnen VI_01 und ST_02 realisierten, wie im Diagramm ersichtlich, auch im Französischen häufig Schwa [ɐ] in der Coda, von allen realisierbaren *r*-Lauten nämlich zu über 51% und über 48% respektive. Nun realisierten jedoch beinahe alle Informantinnen auf Deutsch eine ähnliche Anzahl an [ɐ], ausgenommen jene aus der Schweiz, deren am häufigsten produzierte Form das konsonantische *r* in Gestalt des alveolaren Taps [r] in sowohl Onset als auch Coda war. Die Tendenz zur höheren Anzahl an [ɐ] in der Französischen Coda findet sich, wie oben kurz erwähnt, oftmals dort, wo Informantinnen in der Erstsprache als konsonantische Form am ehesten die alveolare Realisierungsform des

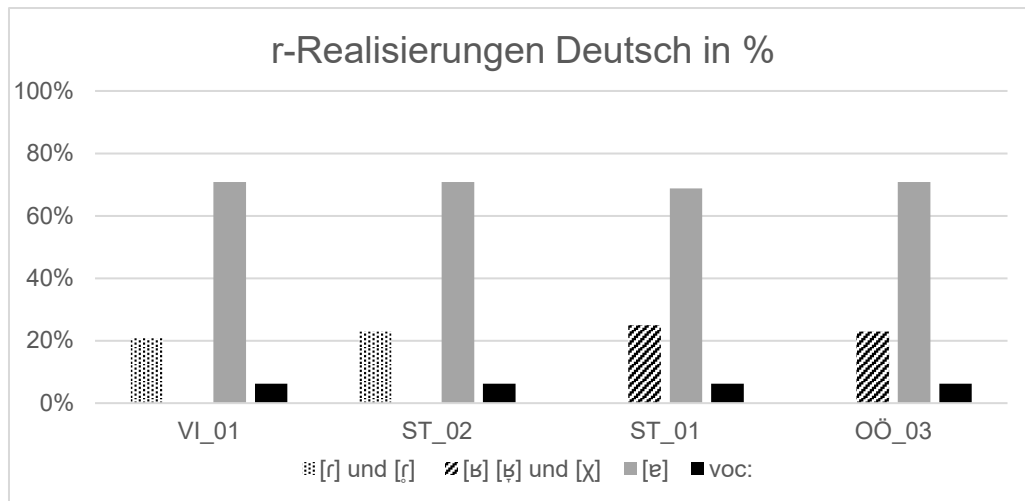


Abb. 26: Auszug aus *r*-Realisierungen im Vergleich

Vergleicht man nun die Werte aus den Grafiken der Abbildungen 24, 25 und 26 sieht man einen deutlichen Unterschied im Verhalten der Informantinnen im Französischen. Während Informantinnen ST_01 und OÖ_03 am ehesten uvulare Realisierungen von *r* produzierten, wie hier in einem Balken zusammengefasst, den stimmhaften und den stimmlosen uvularen Frikativ [ʁ] sowie [χ] und den Approximanten [ʁ̥], tendierten VI_01 und ST_02 zum stimmhaften oder entstimmten alveolaren Tap [r] und [ʀ].

Um nicht nur diese Fälle herauszuheben, soll folgende Grafik noch einen Gesamtüberblick über realisierte Schwa [ɐ] bei allen Informantinnen geben:

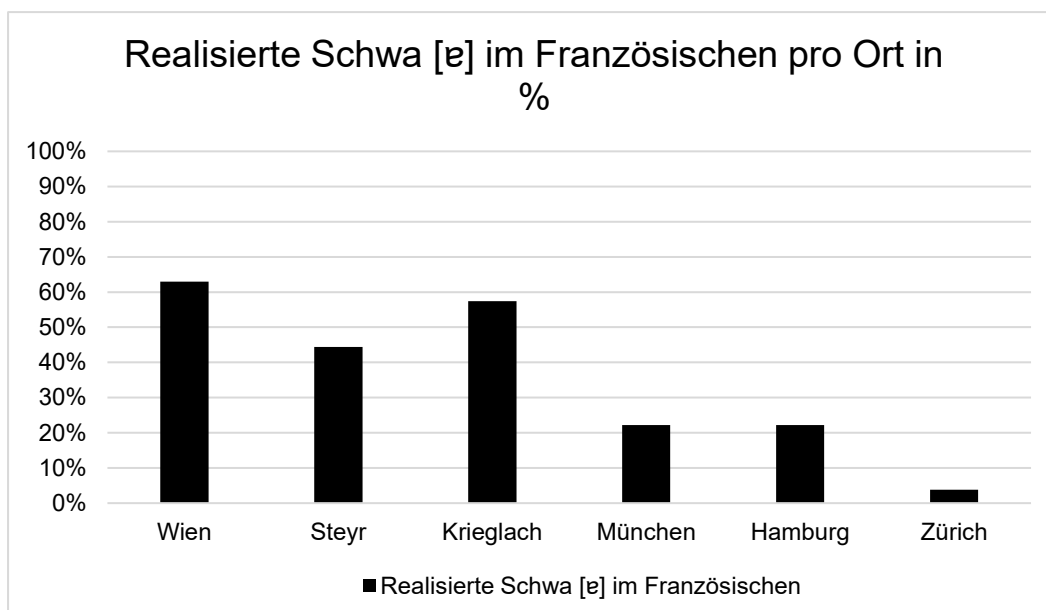


Abb. 27:
Schwa [ɐ]
im franzö-
sischen
Korpus in
Prozent

Wie in der Grafik ersichtlich, wurde in den österreichischen Aufnahmeorten die höchste Anzahl an Schwa [ɐ] in der französischen Coda produziert. Berechnet wurde die Anzahl an Schwa [ɐ] als Anteil von 27 möglichen Okkurrenzen von *r* in Österreich und

Deutschland, beziehungsweise 26 in der Schweiz, in Prozent. Die Schweiz bildet mit einer einzigen Realisierung die Gruppe, in welcher diese Realisierungsform mit nur 3,84% am wenigsten ausgeprägt ist.

Bewegen wir uns nun zu konsonantischen Realisierungsformen im Französischen, so können wir erkennen, dass die Informantinnen tendenziell eher zum stimmlosen uvularen Frikativ [χ] als zum stimmhaften [ʁ] tendieren, auch unabhängig davon, ob sich das auszusprechende *r* in Onset oder Coda befindet oder mit weiteren Konsonanten verzweigt ist. Obwohl die Diskrepanz zwischen realisierten stimmlosen [χ] und stimmhaften [ʁ] von Ort zu Ort variiert, ist erstere Realisierung trotzdem zahlenmäßig immer letzterer voraus:

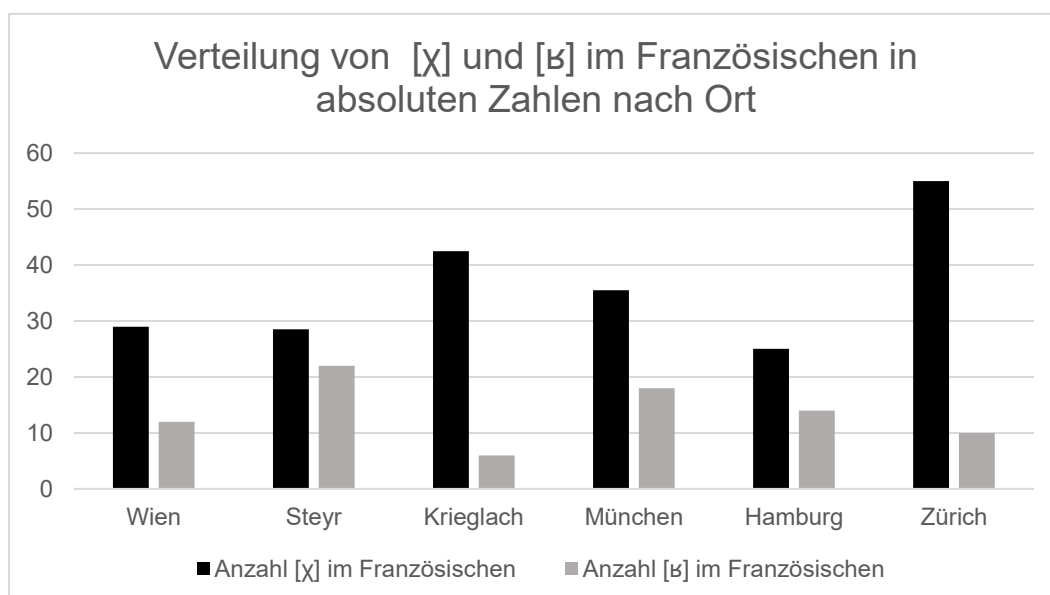


Abb. 28:
Gegen-
überstel-
lender
Vergleich
von [χ]
und [ʁ]

Der bereits im Kapitel zum Deutschen beschriebene uvulare Approximant [ʁ] verteilt sich auch im Französischen unterschiedlich. An einem Ort im Korpus wurde besonders oft auf diese Realisierungsform zurückgegriffen. Die Hamburger Informantinnen realisierten insgesamt 19 Mal den uvularen Approximanten [ʁ] im Französischen, zehn Okkurrenzen mehr als die Gruppe der Münchnerinnen, welche sich an zweiter Stelle mit neun realisierten Approximanten [ʁ] befindet. In Hamburg verteilt sich die Anzahl an realisierten Approximanten [ʁ] fast gleichmäßig auf Onset und Coda, mit elf sowie acht Okkurrenzen respektive. Auffällig dabei ist noch, dass sich der uvulare Approximant [ʁ] bei den Hamburgerinnen mit Ausnahme einer Okkurrenz immer in unverzweigter Onset- oder Codaposition befindet. Auch im restlichen Korpus, wo dieser Laut zahlenmäßig weniger oft vertreten ist, findet sich das Muster dieser Verteilung auf unverzweigte Onset- und Codaposition wieder.

Wie bereits in Kapitel 3. über mein methodisches Vorgehen erwähnt wurde, richtete ich eigens eine Kategorie für einen hochgestellten stimmlosen uvularen Frikativ [χ] ein, da es an manchen Stellen durch geringe akustische Ausprägung nicht passend schien, [χ] zu transkribieren, aber doch genügend, um die Einführung jener Kategorie für die Auswertung als legitim zu erachten. Diese Form der Realisierung kam zwar nicht sehr häufig vor, dennoch möchte ich auf die Verteilung dieser Realisierung auf ihre Position in der Silbe eingehen. Insgesamt im französischen Korpus drei Mal wurde dieser Laut im Onset realisiert, neun Mal in der Coda, am häufigsten bei den Schweizer Sprecherinnen mit insgesamt fünf realisierten [χ]. Wirft man einen Blick auf die Transkriptionen der jeweiligen Informantinnen, findet sich diese Realisierung am häufigsten in der Äußerung *fort* sowie jeweils einmal bei *force* [fɔχs] (vgl. HH_04_fr) und *voyageur* [vwajaʒœχ] (vgl. HH_02_fr). Im Onset finden wir diesen Laut beispielsweise bei Informantin HH_04 an der Stelle *premier* [pχɛmje] (vgl. HH_04_fr).

In der Gesamtanzahl von allen realisierten Lauten im Französischen verteilen sich die Okkurrenzen wie folgt:

Laut	Okkurrenzen
[χ]	215,5
[ʁ]	82
[e]	57,5
[ɸ]	50
[r]	18
Elision	15,5
Vokalverlängerung	15
[X]	12
[R]	8,5
[r̥]	4
[ɾ]	2
?	2
[r]	1
[r̃]	0
[rʰ]	0

Tabelle 4: Gesamtzahl Okkurrenzen Französisch

Wie auch zuvor in Kapitel 5.1 zeigt folgende Karte des deutschsprachigen Raums auch hier die häufigsten konsonantischen Realisierungsformen im französischen Korpus nach Ort:



Abb. 29: Häufigste konsonantische Realisierungen im französischen Korpus nach Ort

5.3 Spanisch

Für das Spanische sind wir in Kapitel 4.3 auf einige wenige signifikante Ergebnisse gestoßen, deren Vorkommen nun vor jenen, die nicht statistisch auffällig waren, hier genauer erläutert werden soll.

Als erstes fanden wir einen signifikanten Zusammenhang zwischen dem Ort der Informantin und der Anzahl an realisierten alveolaren Taps [r] vor. Dies lässt sich anhand der Zahlen so erklären, dass die Schweizer und Wiener Informantinnen die höchste Anzahl jenes Lautes im Spanischen realisierten, beziehungsweise in seltenen Fällen auf den entstimmten Tap [ɾ] oder entstimmten retroflexen alveolaren Flap [ɽ] zurückgriffen. Die Informantinnen ZH_02 und ZH_06 kamen von 33 möglichen *r*-Realisierungsmöglichkeiten auf jeweils 33 stimmhafte alveolare Taps, die Sprecherin ZH_05 auf 29 stimmhafte alveolare Taps [r], drei entstimmte [ɾ] sowie einen entstimmten alveolaren Flap [ɽ]. In Wien realisierten VI_02 und VI_03 jeweils 31 alveolare Taps [r] und VI_01 insgesamt 27. Die restlichen Realisierungen verteilten sich in Wien auch wieder auf entstimmte Taps [ɾ] und Flaps [ɽ] sowie, anders als in Zürich, auch Schwa [e]:

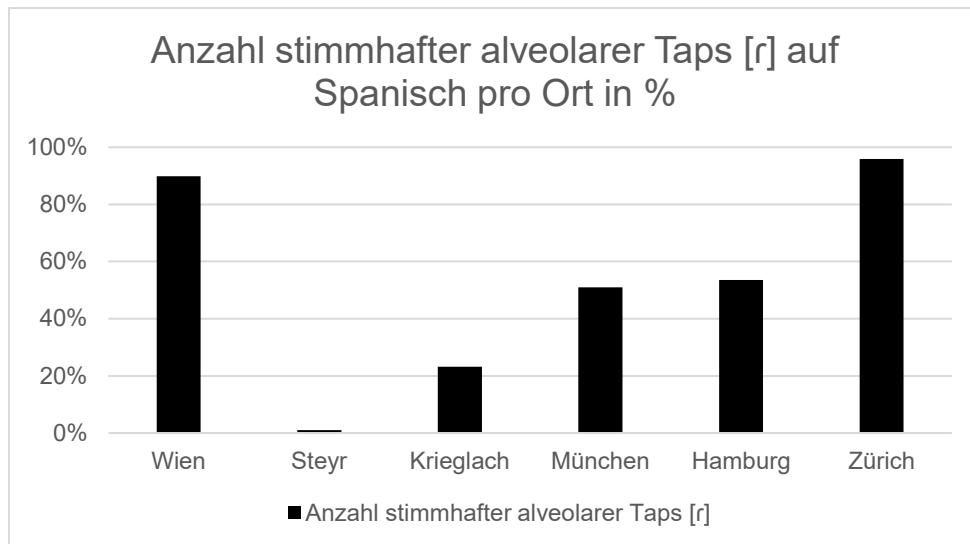


Abb. 30: Häufigkeit stimmhafter alveolarer Tap [r] nach Ort in Prozent

Ein weiteres diskussionswürdiges Ergebnis aus 4.3. ist die Vokalverlängerung im Spanischen anstatt eines konsonantischen *r*-Lautes. Auch dieses Phänomen scheint laut Test mit dem Ort der Informantinnen zusammenzuhängen. Den Grund für dieses Ergebnis können wir anhand der vorhandenen Zahlen herleiten. Wie auch schon im Zuge der Analyse der deutschen Aufnahmen erwähnt, wurde in einem Ort zahlenmäßig öfter auf jene Realisierungsmöglichkeit zurückgegriffen als in den anderen. Dieser Ort war Hamburg. Bei Betrachtung der Werte im Spanischen Aufnahmekorpus fiel diese Art der *r*-Realisierung wieder in Hamburg, aber auch in München häufiger, je Ort insgesamt sechs Mal, auf, eine weitere Informantin aus Oberösterreich bediente sich ebenfalls zwei Mal dieser Realisierung. Sonst kam im spanischen Korpus in keinem anderen Ort eine Okkurrenz jener Vokalverlängerung von [ɔ] vor. Folgende Grafik illustriert die Anzahl jener Realisierungsform in absoluten Zahlen für jene Orte, wo diese vorkam:

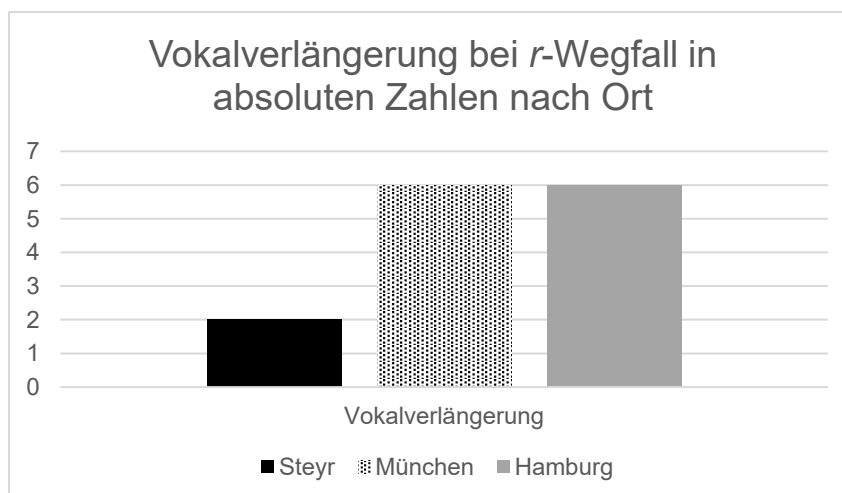


Abb. 31: Vokalverlängerung nach Ort

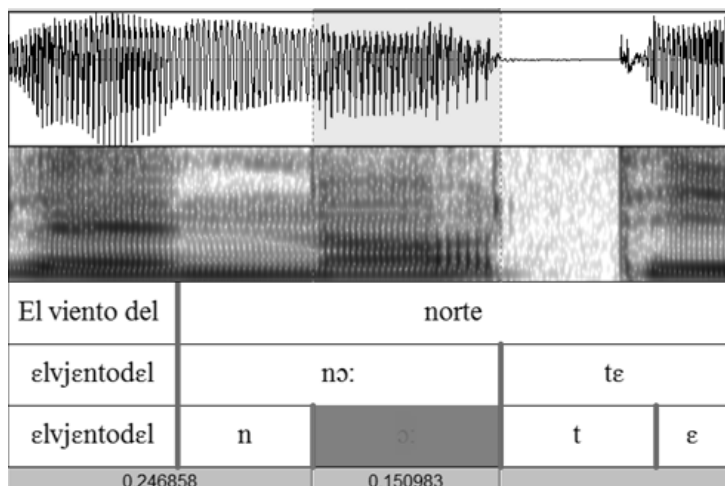


Abb. 32: HH_03_es realisiert *norte* [no:tɛ]

Ein weiteres interessantes Merkmal ließ sich bei der Realisierung von Schwa [ə] in der spanischen Coda feststellen. Wie im vorangegangenen Punkt 4.2. zu Französisch kam es auch im spanischen Korpus zu Interferenzen auf diesem Gebiet. Viele Schwa [ə] wurden ebenfalls im Spanischen Korpus dort realisiert, wo ein Großteil der Informantinnen sie auch auf Deutsch bilden würde, also in der Coda anstatt eines konsonantischen *r*. In den Betrachtungen der französischen Aussprache der Informantinnen konnten wir außerdem bereits eine Tendenz zu einer höheren Anzahl an Schwa [ə] dort erkennen, wo sich die Lautung in der Erstsprache der Informantinnen nicht mit jener des gängigen Standards in der Zielsprache deckte. Dies kann man auch im Spanischen beobachten, denn es wurden tendenziell bei jenen Informantinnen Schwa [ə] realisiert, deren *r*-Lautung im Deutschen eher uvular stattfand. Außerdem kam es auch in den Spanischen Aufnahmen in Zürich zu keiner einzigen Okkurrenz von Schwa [ə] in der Coda, also noch weniger als im französischen Korpus. Deshalb scheint jener Ort auch nicht unterhalb in Abbildung 33 auf. Die folgende Grafik soll die Anzahl an realisierten Schwa im Allgemeinen pro Ort in absoluten Zahlen illustrieren, anschließend daran werden beispielhaft Informantinnen und deren bevorzugte *r*-Lautung in der L1 zum Vergleich herangezogen.

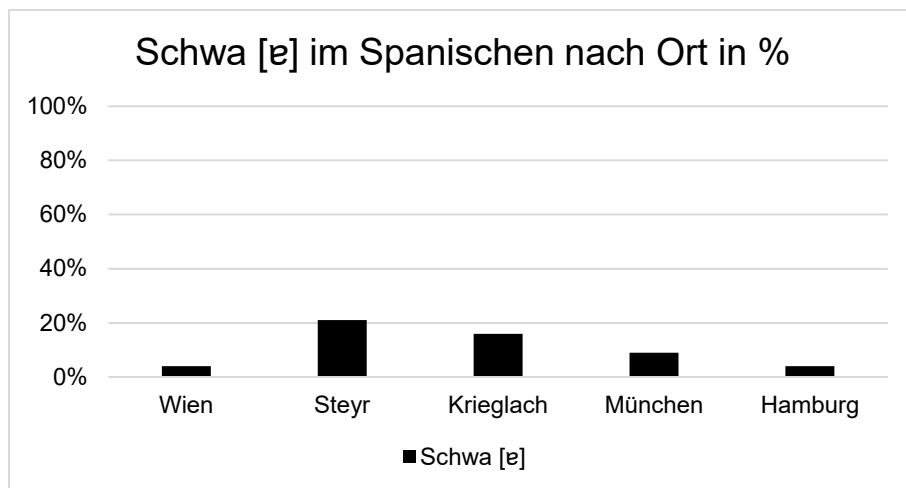


Abb. 33: Schwa [ə] nach Ort

Wirft man nun beispielsweise einen Blick auf die bevorzugte konsonantische *r*-Realisierung im Deutschen fällt auch hier wieder auf, dass die Lautung in der L1 zu Interferenzen führen kann. Die höchste Anzahl an Schwa [ə] im spanischen Korpus wurde in Oberösterreich realisiert. Im deutschen Korpus tendierten jene Informantinnen am ehesten zu uvularen Realisierungsformen wie dem stimmlosen uvularen Frikativ [χ]. Dort wo in der L1 häufiger zu alveolaren Realisierungen tendiert wurde, wie beispielsweise in Wien, kam es auch im Spanischen zu weniger Realisierungen von Schwa [ə]. Eine Beobachtung zu produzierten Schwa [ə] muss jedoch noch zusätzlich angestellt werden, da sich die oben genannten Ergebnisse nicht vollständig nach dem Schema verhalten, nach welchem die Interferenzen dort passieren, wo sich quell- und zielsprachliche Lautung unterscheiden. So gibt es beispielsweise Informantinnen, die auf Deutsch als gängigste konsonantische Form von *r* den stimmhaften alveolaren Tap [r] realisieren, diesen auch im Spanischen produzieren, jedoch trotzdem stellenweise ein Schwa [ə] in der Coda anstatt eines Konsonanten aussprechen. Dies ist beispielsweise bei den Informantinnen VI_01, VI_02 und ST_02 der Fall. Die ersten beiden produzierten jeweils zwei Mal Schwa [ə] anstatt eines konsonantischen *r*-Lautes in der Coda, Informantin ST_02 insgesamt sogar neun Mal. In Hamburg wurden im spanischen Korpus ebenfalls nur wenige Okkurrenzen von Schwa [ə] gezählt, jedoch tendierten diese Informantinnen in der L1 zu uvular realisierten Konsonanten. Da die Realisierung von [ə] auf Deutsch in der Coda keine Seltenheit ist und in den jeweiligen Standardvarietäten Österreichs und Deutschlands beschrieben wird (vgl. Kapitel 2.5.2), kann man bei dessen Übertragung ins Spanische von einer möglichen Interferenz ausgehen. Jedoch hängt diese nicht zwangsläufig damit zusammen, ob die Informantin generell

dazu fähig ist, ein konsonantisches *r* alveolar zu artikulieren, wie uns die zuletzt genannten Informantinnen aus Wien, der Steiermark und Hamburg zeigen.

Als nächstes möchte ich noch auf jene Laute eingehen, die zwar zahlenmäßig nur in kleinen Mengen vorkommen, jedoch für mich vor Beginn der akustischen Analyse unerwartet waren. Dazu zählen der entstimmte alveolare Tap [ɾ] sowie der stimmhafte alveolare Flap [ɹ] und dessen entstimmte Variante [ɾ̥]. Zunächst möchte ich anmerken, dass der entstimmte alveolare Tap [ɾ] sowohl bei jenen Informantinnen vorkam, die in ihrem deutschen Dialekt als konsonantisches Laut am häufigsten den stimmhaften alveolaren Tap [r] produzierten, als auch bei jenen, die auf Deutsch eher zu uvularen Realisierungsformen tendierten. So besteht Grund zur Annahme, dass diese Art von Okkurrenz nicht eindeutig mit dialektaler Lautung in Verbindung gebracht werden kann. Somit bietet es sich an, nach anderen Verteilungsmustern zu suchen.

Wirft man einen Blick auf das Vorkommen von [ɾ] in Silbenposition und etwaiger Verzweigung, so findet man diese Realisierung am häufigsten im verzweigten Onset sowie in der unverzweigten Coda. Um nun sehen zu können, welche Stellen im Text jene sind, an denen wir diesen Laut häufig finden, empfiehlt sich ein Blick in die Transkriptionen. Dort verteilen sich entstimmter Tap [ɾ] sowie stimmhafter Flap [ɹ] und entstimmter Flap [ɾ̥] in folgenden Textstellen unterschiedlich häufig: Bei den Textstellen *norte* und *fuerte* fanden sich diese Laute am häufigsten wieder, also in unverzweigter Codaposition. Da aber auf die Silben *nor-* und *fuer-* die nächste Silbe mit einem konsonantisches Onset in Form eines stimmlosen Konsonanten [t] folgt, könnte man hier die Produktion entstimmter Konsonanten auf einen regressiven Assimilationsprozess die Stimmhaftigkeit betreffend (vgl. Hualde 2014: 99) zurückführen. Ein Assimilationsprozess, der die Stimmhaftigkeit des *r*-Lautes beeinflussen könnte, ließe sich demnach auch bei Textstellen wie *primero* oder *estrechamente* beobachten, hier jedoch in progressiver Form, von den stimmlosen Konsonanten [p] oder [t] in Richtung des darauffolgenden *r*-Lautes.

Abschließend sollen auch hier die gesamten Okkurrenzen im spanischen Korpus in einer Tabelle aufgelistet sowie die häufigste Konsonantische Form nach Ort in der darauffolgenden Karte illustriert werden:

Laut	Okkurrenzen
[r]	311,5
[χ]	70
[e]	54
[ʁ]	48,5
[ʁ̥]	36
[R]	16,5
[ɾ]	15
Vokalverlängerung	14
[X]	8
Elision	6
[t̪]	5
[ɾ]	4,5
?	4
[r ^h]	1
[r]	0

Tabelle 5: Gesamtzahl Okkurrenzen Spanisch

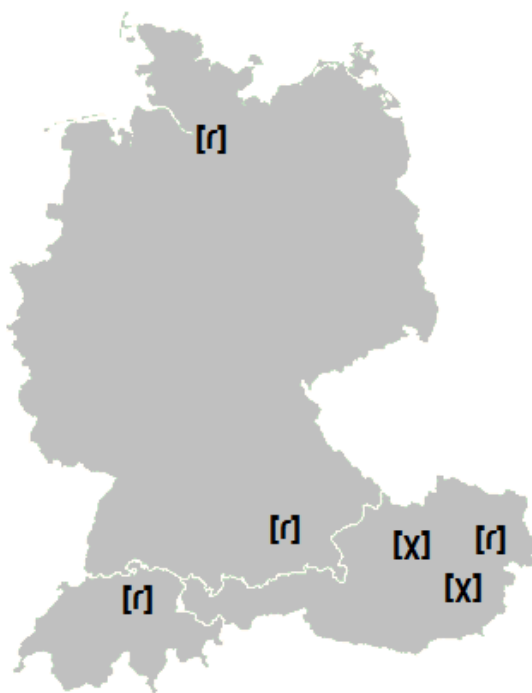


Abb. 33: Häufigste konsonantische Realisierungen im spanischen Korpus nach Ort

6. Schluss

Aus der bisherigen Forschung zum Erlernen von Fremdsprachen und im Speziellen deren Aussprache geht hervor, dass es durchaus von Interesse ist, dieser den Platz zuzugestehen, den sie benötigt. Wir haben gesehen, dass eine akzentgeprägte Aussprache in der Zielsprache viel über uns verrät und in so manchen sozialen Kontexten unser Gegenüber zu voreiligen Schlüssen über unsere Sprachkompetenz verleiten kann. Auch haben wir gesehen, dass es durchaus legitim ist, den Dialekten als Quellsprache weiter besonderes Augenmerk zukommen zu lassen. Der Großteil der bisher in diesem Forschungsgebiet durchgeführten Untersuchungen beschreibt vorrangig die Standardsprache, wenn sie die Quellsprache als beeinflussenden Faktor auf die Zielsprache im Lernprozess beleuchtet. Doch haben wir sowohl im Stand der Forschung als auch im praktischen Teil dieser Arbeit gesehen, dass das Miteinbeziehen von Dialekten und eine plurizentrische Betrachtung von Sprache notwendige Faktoren sind, wenn man im Bereich des Erlernens der Aussprache an ein gewünschtes Ziel gelangen möchte.

Auch wenn die Ergebnisse der im Zuge der vorliegenden Arbeit durchgeführten Tests nur in seltenen Fällen signifikante Zahlen aufweisen konnten, erkennt man dennoch bei näherer Betrachtung der verschiedenartigen Lautung Zusammenhänge, die nicht allein auf Deutsch als Standardsprache zurückgeführt werden können. Dies wurde vor allem bei Betrachtung der Zahlen in Kapitel 5. klar, wo die Verteilung gewisser Aussprachemerkmale oftmals mit dialektaler Lautung nahezu kongruent war. Viele dieser Ergebnisse, die bereits bei Informantinnen in Aufnahmen auf Deutsch auffielen, zogen sich weiter in die fremdsprachliche Lautung. Weiterführende Forschung zu diesem Thema hat also durchaus eine gewisse Daseinsberechtigung, wobei es in jedem Fall größerer Mengen an Datenmaterial bedarf.

Ein weiterer Forschungsaspekt, der von Interesse für die Zukunft ist, wäre die Untersuchung von Spontansprache, die im Zuge der vorliegenden Arbeit aufgrund nicht einheitlich vorhandenen Materials nicht möglich war. Selbstverständlich war es auch im Stand der Forschung sowie im Ergebnisteil vorliegender Arbeit schon ersichtlich, dass die Dialektzone allein nicht zur Gänze als Erklärung für bestimmte Aussprachemerkmale herangezogen werden kann. Oftmals stellt der Dialekt nur einen Teilaspekt für das Vorkommen von Transfer dar. Auch die Frage nach anderen Aussprachemerkmalen, abseits der Klasse der *r*-Laute, könnte neue Einblicke in die Thematik liefern. In

einer größer angelegten Studie können Faktoren wie Unterrichtsmethode in den Fremdsprachen, eventuelle Auslandsaufenthalte und andere soziodemographische Fakten miteinbezogen werden, um noch genauer zu differenzieren, welche Einflüsse die Aussprache in den Zielsprachen formen. Jedoch bietet diese Forschungsarbeit, so hoffe ich, einen Teil der Grundlage für ein fortbestehendes Interesse an besprochener Thematik und einer dialektensensibleren Ausrichtung der Fremdsprachendidaktik in Klassenzimmern und Kursräumen.

Bibliographie

- Ammon, Ulrich et al. (2004): *Variantenwörterbuch des Deutschen. Die Standardsprache in Österreich, der Schweiz und Deutschland sowie in Liechtenstein, Luxemburg, Ostbelgien und Südtirol*, Berlin, Boston: De Gruyter.
- Bächle, Hans et al. (2007): *À Plus! 1, Méthode intensive*. Berlin: Cornelsen.
- Biersack, Sonja (2002): „Systematische Aussprachefehler deutscher Muttersprachler im Englischen. Eine phonetisch-phonologische Bestandsaufnahme." *Forschungsberichte des Instituts für Phonetik und sprachliche Kommunikation der Universität München (FIPKM)* 39, 37-130.
- Birdsong, David (2007): „Nativelike pronunciation among late learners of French as a second language", in: *Bohn, Ocke-Schwen/Munro, Murray (Hrsg.): Language experience in second language speech learning*, Amsterdam: John Benjamins, 99-116.
- Boersma, Paul/Weenik, David (2017): *Praat: Doing phonetics by computer*. Aktuelle Version: März 2017 (www.praat.org)
- Carbó, Carme et al. (2003): „Estándar oral y enseñanza de la pronunciación del español como primera lengua y como lengua extranjera“, in: *ELUA, Estudios de Lingüística de la Universidad de Alicante* 17, 161-180.
- Catford, John Cunnison (1965): *A Linguistic Theory of Translation*, London: Oxford University Press.
- Duller, Christine (2013): *Einführung in die Statistik mit EXCEL und SPSS: Ein anwendungsorientiertes Lehr-und Arbeitsbuch*, Berlin: Springer.
- Durrell, Martin (1998): „Zum Problem des sprachlichen Kontinuums im Deutschen“, in: *Zeitschrift für Germanistische Linguistik*, 26.1, 17-30.
- Edwards, John (1999): „Refining Our Understanding of Language Attitudes“, in: *Journal of Language and Social Psychology* 18.1, 101-110.
- Ehrlich, Karoline (2009): *Die Aussprache des österreichischen Standarddeutsch – umfassende Sprech- und Sprachstandserhebung der österreichischen Orthoepie* (Dissertation, Universität Wien).

- Elminger, Daniel/Forster, Simone (2005) : „*La Suisse face à ses langues: Histoire et politique du plurilinguisme – situation actuelle de l’enseignement des langues*“, Neuchâtel : Institut de recherche et de documentation pédagogique.
- Fagyal, Zsuzsanna/Kibbee, Douglas/Jenkins, Fred (2006): *French: A linguistic introduction*, Cambridge: Cambridge University Press.
- Flege, James Emil (1980): „Phonetic Approximation in Second Language Acquisition“. in: *Language Learning* 30, 117–134.
- Flege, James Emil (1981): „The Phonological Basis of Foreign Accent: A Hypothesis“, in: *TESOL Quarterly* 15.4, 443-455.
- Flege, James Emil (1984): „The Detection of French Accent by American Listeners“, in: *The Journal of the Acoustical Society of America* 76, 692-707.
- Flege, James Emil (1987): „A Critical Period for Learning to Pronounce Foreign Languages?“, in: *Applied Linguistics* 8, 162-177.
- Gómez de Olea, Lourdes, López Pernía, Palmira, Heydel, Marlies (2001) : *El nuevo curso 1*. München : Langenscheidt.
- Gómez Martínez, Susana (2011): „The syllable structure: understanding spanish speakers pronunciation of english as a L2“, in: *RAEL: revista electrónica de lingüística aplicada* 10, 1-7.
- Görrissen, Margarita (2009): *Caminos neu A1*. Stuttgart : Klett.
- Göschel, Joachim (1971): „Artikulation und Distribution der sogenannten Liquida r in den europäischen Sprachen“, in: *Indogermanische Forschungen* 76, 84-126.
- Gries, Stefan Th. (2008): *Statistik für Sprachwissenschaftler*. Göttingen: Vandenhoeck & Ruprecht.
- Hall, Allan Tracy (1992): *Syllable Structure and Syllable-Related Processes in German*. Berlin: De Gruyter.
- Harlow, Linda/Muyskens, Judith (1994): „Priorities for Intermediate-Level Language Instruction“, in: *The Modern Language Journal* 78, 141-154.
- Herrero de Haro, Alfredo/Andión Herrero, Antonieta (2012): „La fonética de aprendices de español del sur y del norte de Inglaterra. Consideraciones acerca de su enseñanza.“, in: *Revista Española de Lingüística Aplicada* 25, 45-68.

- Holzinger, Gabriele et al. (2012) *Descubramos el español. Spanisch interlingual*. Wien: HTP.
- Howell, Robert B. (1987): „Tracing the Origin of Uvular *R* in the Germanic Languages“, in: *Folia Linguistica Historica*. 7.2, 317–350.
- Hualde, José Ignacio (2005): *The sounds of Spanish*. Cambridge: Cambridge University Press.
- Hualde, José Ignacio (2014): *Los sonidos del español*. Cambridge: Cambridge University Press.
- Hurtado, Luz/Estrada, Chelsea (2010) „Factors influencing the second language acquisition of Spanish vibrants.“ in: *The Modern Language Journal* 94.1, 74–86.
- International Phonetic Association [1999] (¹⁵2014): *Handbook of the International Phonetic Association*, Cambridge: Cambridge University Press.
- Jambon, Krystelle (2013): *Voyages 1 – Französisch für Erwachsene*. Stuttgart: Klett.
- Kabatek, Johannes/Pusch, Claus D. (2009): *Spanische Sprachwissenschaft*. Tübingen: Narr.
- Kettemann, Bernhard/Viereck, Wolfgang (1983): „Muttersprachiger Dialekt und Fremdsprachenerwerb. Dialekttransfer und Interferenz: Steirisch-Englisch“ in: James, Allan/Kettemann, Bernhard (Hrsg.): *Dialektphonologie und Fremdsprachenerwerb, Dialect Phonology and Foreign Language Acquisition*, Tübingen: Narr.
- Kolly, Maria-José (2013): „Akzent auf die Standardsprachen: Regionale Spuren in "Français fédéral" und "Schweizerhochdeutsch"“, in: *Linguistik Online*, 58.1, 37-76.
- König, Werner (2014): *dtv-Atlas deutsche Sprache*. München: dtv.
- Krech, Eva-Maria (2009): *Deutsches Aussprachewörterbuch*. Berlin: De Gruyter.
- Ladefoged, Peter/Maddieson, Ian (1996): *The Sounds of the World's Languages*. Oxford: Blackwell.
- Laudut, Nicole/ Patte-Möllmann, Catherine (2014): *On y va! A2*. Ismaning: Hueber.
- Lauret, Bertrand (2007): *Enseigner la prononciation du français: questions et outils*.

- Paris: Hachette.
- Lenneberg, Eric (1967): *The biological foundations of language*, New York: Wiley.
- Léon, Pierre Roger (2007): *Phonétisme et prononciations du français*. Paris: Armand Colin.
- Lindau, Mona (1985): „The story of /r/“, in: Fromkin, Victoria (Hrsg.): *Phonetic Linguistics. Essays in Honor of Peter Ladefoged*, Academic Press.
- Maddieson, Ian (1984): *Patterns of Sounds*. Cambridge: Cambridge University Press.
- Major, Roy Coleman (2001): *Foreign Accent: The Ontogeny and Phylogeny of Second Language Phonology*. Mahwah: Erlbaum.
- Martínez-Celdrán, Eugenio/Fernández-Planas, Ana María/Carrera-Sabaté, Josefina (2003): „Castilian Spanish“, in: *Journal of the International Phonetic Association*, 33(2), 255–259.
- Moyer, Alison (2013): *Foreign Accent: The Phenomenon of Non-native Speech*. Cambridge: Cambridge University Press.
- Muhr, Rudolf (1997): „Zur Terminologie und Methode der Beschreibung plurizentrischer Sprachen und deren Varietäten am Beispiel des Deutschen.“ In: Muhr, Rudolf/Schrodt, Richard (Hrsg.): *Österreichisches Deutsch und andere nationale Varietäten plurizentrische Sprachen in Europa*, Wien: Holder-Pichler-Tempsky. 40-66.
- Müller, Horst M. (2002): *Arbeitsbuch Linguistik*. Paderborn: Ferdinand Schöningh.
- Olsen, Michael K. (2012): „The L2 Acquisition of Spanish Rhotics by L1 English Speakers: The Effect of L1 Articulatory Routines and Phonetic Context for Allophonic Variation“, in: *Hispania*, 95.1, 65-82.
- Panizzolo, Paola (1982): *Die schweizerische Variante des Hochdeutschen*, N.G. Elwert: Marburg.
- Picard, Marc (1987): *An Introduction to the Comparative Phonetics of English and French in North America*. Amsterdam: John Benjamins Publishing Company.
- Pustka, Elissa (2007): *Phonologie et variétés en contact : Aveyronnais et Guadeloupéens à Paris*. Tübingen: Narr.
- Pustka, Elissa (2011): *Einführung in die Phonetik und Phonologie des Französischen*. Berlin: Erich Schmidt Verlag.
- Quilis, Antonio (1993): *Tratado de fonología y fonética españolas*. Madrid: Gredos

- Ransmayr, Jutta (2012): „Zur Wahrnehmung der Varietäten des Deutschen im Unterricht für Deutsch als Muttersprache und Deutsch als Fremdsprache“, in: *Der Sprachdienst*, 5, 198-207.
- Rey, Christophe (2011): „Les recommandations sur la prononciation dans le Dictionnaire de l'Académie française“ in: *Langue commune et changements de norme*. Honoré Champion, 149-162.
- Rose, Marda (2010): „Intervocalic Tap and Trill Production in the Acquisition of Spanish as a Second Language“, *Studies in Hispanic and Lusophone Linguistics*, 3.2, 379-420.
- Rothstein, Björn (2011): *Wissenschaftliches Arbeiten für Linguisten*. Tübingen: Narr.
- Russ, Charles V. J. (2010): *The sounds of German*. Cambridge, Cambridge University Press.
- Scovel, Thomas (1969): „Foreign Accents, Language Acquisition and Cerebral Dominance“, in: *Language Learning*, 19. 3-4, 245-253.
- Scovel, Thomas (2000): „A Critical Review of the Critical Period Research“, in: *Annual Review of Applied Linguistics*, 20, 213-223.
- Selinker, Larry (1972): „Interlanguage“, in: *International Review of Applied Linguistics* 10, 209-231.
- Selinker, Larry/Gass, Susan (1992): *Language Transfer in Language Learning*. Amsterdam : John Benjamins Publishing Company.
- Siebs, Theodor [1898] (¹⁸1961): *Deutsche Bühnenaussprache*. Köln: Ahn.
- Simpson, Adrian (1998): „Accounting for the Phonetics of German *r* without process“, in: *ZAS Working Papers in Linguistics*, 11, 91-104.
- Sokol, Monika (2007): *Französische Sprachwissenschaft. Ein Arbeitsbuch mit thematischem Reader*. Tübingen: Narr.
- Thomas, Erik (2011): *Sociophonetics: An Introduction*. Basingstoke: Palgrave Macmillan.
- Tranel, Bernard (1987): *The Sounds of French*, Cambridge: Cambridge University Press.

- Vremšak-Richter, Vanda (2010): *Unterschiede zwischen den deutschen Aussprachewörterbüchern: Grosses Wörterbuch der deutschen Aussprache (1982), Duden. Aussprachewörterbuch (2005), Deutsches Aussprachewörterbuch (2009)*. Ljubljana University Press, Faculty of Arts.
- Weinreich, Uriel [1953] (⁹2010): *Languages in Contact. Findings and Problems*. Berlin, Boston: De Gruyter Mouton.
- Wiese, Richard (2001): „The unity and variation of (German) /r/“. in: Hans van de Velde/Roeland van Hout (Hrsg.): *r-atics: sociolinguistic, phonetic and phonological characteristics of /r/*. Brüssel, Institut des Langues Vivantes et de Phonétique, 11-26.
- Wiesinger, Peter (1983): „Die Einteilung der deutschen Dialekte“. in: Besch, W. u. a. (Hrsg.): *Dialektologie*, Band 2. Berlin/New York: De Gruyter, 807-899.
- Wonka, Lisa (2015) *Merkmale des gesprochenen österreichischen Deutsch anhand der Analyse von ORF-Sendungen*, Wien: Präsens.
- Zehetner, Ludwig (1983) „'Bairisches Englisch'. Muttersprachiger Dialekttransfer im Fremdsprachenerwerb am Beispiel regionalspezifischer Schwierigkeiten und Möglichkeiten für den Englischunterricht im bairischen Dialektraum“, in James, Allan/Kettemann, Bernhard (Hrsg.): *Dialektphonologie und Fremdsprachenerwerb, Dialect Phonology and Foreign Language Acquisition*. Tübingen: Narr, 153-172.
- Zöfel, Peter (2003): *Statistik für Psychologen: im Klartext*. München: Pearson Studium.

Internetquellen:

Académie Française: Dictionnaire de l'Académie Française en ligne:

<http://atilf.atilf.fr/dendien/scripts/generic/form.exe?54;s=367666650;>
[20.12.2016].

Académie Française: Statutenliste: http://www.academie-francaise.fr/sites/academie-francaise.fr/files/statuts_af_0.pdf [20.12.2016]

ADABA: Österreichische Aussprachedatenbank: <http://www.aussprache.at/>
[30.11.2016].

eLAUT: Projekt der Fachrichtung der Computerlinguistik und Phonetik der Universität des Saarlandes <http://www.coli.uni-saarland.de/elaut/konsPulm.htm>
[4.12.2016].

GERS: Gemeinsamer Europäischer Referenzrahmen für Sprachen
http://www.oesz.at/download/publikationen/Broschuere_interaktiv.pdf
[17.1.2016]

MESOS: Methodological education for the Social Sciences
http://www.mesosworld.ch/lerninhalte/Biv_Chi/de/html/index.html [23.1.2016]

Real Academia Española: Diccionario de la lengua española en línea:
<http://www.rae.es/diccionario-panhispanico-de-dudas/representacion-de-sonidos> [2.1.2017].

Real Academia Española: Diccionario panhispánico de dudas en línea:
<http://www.rae.es/diccionario-panhispanico-de-dudas/que-es> [2.1.2017].

StatistikAustria:

http://www.statistik.at/web_de/services/publikationen/5/index.html?includePage=detailView§ionName=Bildung%2C+Kultur&pubId=722 [15.12.2016].

Statista: <https://de.statista.com/statistik/daten/studie/161316/umfrage/anteil-der-schueler-die-franzoesisch-als-fremdsprache-lernen-in-deutschland-und-der-eu/>
[15.12.2016].

SynAlm: Projekt der Universität Konstanz: <https://cms.uni-konstanz.de/ling/syntax-des-alemannischen/projekt/alemannisch/> [2.1.2017]

Bildquellen:

Stumme Karte des deutschsprachigen Raums:

<http://www.voith.com/de/karriere/unsere-standorte-291.html> [27.3.2017]

Anhang

Résumé du travail présent en langue française

État actuel des recherches

Dans le cadre de l'apprentissage d'une langue étrangère, l'apprenant doit considérer une multitude d'aspects en ce qui concerne la maîtrise de la langue en question. Le vocabulaire, la grammaire, la syntaxe, ainsi que la prononciation ne sont que quelques chapitres qu'il faut prendre en compte en s'entraînant. Le dernier aspect, celui de la prononciation, est souvent négligé dans les cours de langue étrangère, que ce soit à l'école ou à l'université. Ainsi, il faut se poser la question si ce classement secondaire de la prononciation de l'apprenant est justifié quand on considère la multitude de raisons pour lesquelles sa maîtrise est importante.

Une des raisons principales est que l'apprentissage d'une prononciation qui s'approche de celle d'un locuteur natif est qu'elle est nécessaire pour garantir la compréhension mutuelle entre deux locuteurs (cf. Estrada/Hurtado 2010: 74). De plus, la prononciation devrait occuper une place dans le cadre éducatif parce que les apprenants eux-mêmes souhaitent obtenir un niveau de *nativelike pronunciation* (cf. Birdsong 2007: 110), donc il s'agit aussi d'un facteur qui joue un rôle en ce qui concerne la motivation des élèves. Une étude de Harlow et Muyskens (1994) a même montré que l'obtention d'une prononciation authentique est un des buts principaux des apprenants. Finalement, il ne faut pas oublier le rôle de la prononciation dans le discours social. Notre façon de parler permet aux interlocuteurs de tirer des conclusions sur notre statut social, notre intelligence et bien sûr notre compétence dans la langue étrangère, même si l'apprenant a déjà obtenu un niveau assez haut dans les autres champs d'apprentissage, comme la grammaire ou le lexique (cf. Edwards 1999: 2).

Nous voyons que la prononciation et son apprentissage sont des chapitres qu'il faut prendre en considération. Cependant, suite à cette observation, vient alors la question concernant la possibilité d'atteindre le niveau d'un locuteur natif. Beaucoup d'études et une multitude de scientifiques s'intéressent au sujet de la capacité de l'être humain de maîtriser la prononciation de la langue apprise. Pendant longtemps, beaucoup de scientifiques étaient d'accord avec l'hypothèse d'une limite liée à l'âge de l'apprenant

concernant la capacité d'obtenir la dite *nativelike pronunciation*. Lenneberg (1967) a fait avancer la notion de la période critique, après laquelle l'apprenant n'est plus capable d'arriver au niveau souhaité de la prononciation dans la langue étrangère. Selon Lenneberg (1967), cette limite est due à la plasticité diminuante du cerveau humain (cf. Lenneberg 1967: 175-176) et il place cet instant de changement cérébral vers l'âge de la puberté. Après avoir surmonté cette limite, selon Lenneberg (1967), l'apprenant parle avec un certain accent et sera donc incapable de s'en débarrasser. Ce genre d'argumentation a déjà beaucoup influencé non seulement la linguistique, mais aussi la didactique des langues étrangères. La question de l'accent est un sujet auquel se sont intéressés plusieurs linguistes après Lenneberg, parmi eux Selinker (1972) qui a proposé la notion d'*interlanguage*, ou interlangue, un concept qui décrit l'état d'un apprenant entre la prononciation de sa L1 et de sa L2. Cette interlangue est marquée par un accent plus ou moins fort et Selinker (1972) constate que seulement un pourcentage assez faible d'apprenants sera capable de surmonter cet état (cf. Selinker 1972: 212). La plupart des apprenants restent dans ce stade d'interlangue et continuent à parler avec un accent, un concept que Selinker appelle *fossilisation* (vgl. Selinker/Gass 1992: 197).

La question de ce que constitue un accent étranger dans la deuxième langue est aussi une question qui se pose à plusieurs chercheurs après, mais aussi déjà avant Lenneberg et la période critique (1967). Weinreich (1953) a observé des interférences linguistiques au niveau segmental ainsi qu'au niveau de la prosodie. Il a défini sept scénarios différents dans lesquels le transfert linguistique peut apparaître. Tous se réfèrent à l'interaction entre la première et la deuxième langue de l'apprenant. La notion du transfert linguistique est également présente dans les travaux de Flege (1981), qui constate qu'un accent étranger apparaît quand l'apprenant est confronté à un aspect phonétique de la L2 qui diffère de la L1. Alors, en essayant de produire le nouveau son dans la L2, il prend le son le plus proche de sa L1 au lieu de produire le son inconnu, ainsi prononçant un son entre celui de la L1 et de la L2. Cette observation est congruente avec la notion de Catford (1965) de *phonological translation*, qui a proposé un angle d'approche semblable à celui de Flege (1981). Flege (1981) préfère la notion de *phonological translation* à celle de la période critique de Lenneberg (1967), car la dernière lui paraît trop simplifiée quant à la complexité du processus de l'apprentissage d'une langue et il critique l'approche exclusivement biologique proposée par la période critique (cf. Flege 1987: 167). Cependant, Flege (1987) admet aussi qu'il existe bien

des différences entre de jeunes apprenants et des adultes, mais il critique un manque de preuves au niveau neurologique ainsi qu'une approche plus complexe qui inclut d'autres facteurs comme la motivation, l'empathie culturelle ou la méthode didactique.

Et si l'obtention d'une prononciation sans accent dans la langue étrangère n'est pas possible après l'âge de la puberté, est-il raisonnable d'inclure son instruction dans les curricula? Certains linguistes argumentent que l'apprentissage d'une prononciation authentique, sans accent, est possible dans des circonstances favorables. Dans une étude plus récente de Birdsong (2007), il a découvert des succès au niveau de l'obtention d'une *nativelike pronunciation* parmi les apprenants qui montrent une très haute motivation ainsi que ceux qui ont profité d'une instruction spécifiquement ciblée à entraîner la prononciation (cf. Birdsong 2007: 110). Dans des études semblables, comme celles de Bongaerts, Planken et Schils (1995) ou Bongaerts (1999), on trouve des résultats similaires. On peut donc conclure que l'obtention d'une prononciation authentique dans la L2 est possible, mais les travaux existants suggèrent que la plupart des apprenants ne vont pas atteindre ce but.

On a observé que la prononciation et les accents étrangers sont souvent basés sur l'influence de la L1. Un aspect qui est souvent négligé est celui de la diversité de dite L1. Dans le cas de ce travail, la L1 en question est l'allemand, une langue avec plusieurs variétés et dialectes. Il faut maintenant se poser la question de quelle langue maternelle nous parlons quand nous nous référons à la L1 et au transfert linguistique. Un apprenant hambourgeois produit-il les mêmes interférences qu'un apprenant de Zurich ou de Vienne? Il y a un tel florilège de dialectes dans le monde germanophone qu'il paraît impossible de parler d'une L1 homogène. Plusieurs linguistes se sont déjà posés cette question de diversité dialectale dans le processus de l'apprentissage d'une L2.

Au niveau de la prononciation, on peut citer Zehetner (1983), qui a observé l'influence dialectale dans le cadre d'une étude avec des apprenants d'une région qui parle un dialecte bavarois et leurs problèmes avec la prononciation anglaise. Il a fait les mêmes observations que celles décrites auparavant, à savoir que les apprenants se servent d'un son de leur langue maternelle au lieu de produire le son inconnu anglais. Une étude semblable de Kettemann et Viereck (1983) constate le même phénomène avec des élèves de Graz en Autriche. Schubinger (1983 [1937]) a fait des recherches similaires dans une région autrichienne où les apprenants parlent un dialecte alémanique.

Dans le cadre du travail présent, nous allons considérer trois régions dialectales germanophones différentes dans, au total, six villes. Les dialectes sont l'alémanique, le bavarois et le bas-allemand. Les six villes sont Zurich, où on parle un dialecte alémanique, Hambourg avec son dialecte bas-allemand, et Munich, Vienne, Steyr (Haute Autriche), ainsi que Krieglach (Styrie), des villes qui appartiennent à la zone dialectale bavaroise. Les langues étrangères dont nous allons observer la prononciation sont le français et l'espagnol.

Sur la base de l'état actuel des recherches sur les interférences dialectales au niveau phonétique, nous pouvons faire des hypothèses sur les aspects qui pourraient poser des problèmes pour les apprenants en prenant en compte leur provenance. Pour donner des exemples, je vais présenter quelques interférences possibles basées sur des aspects dialectales en français ainsi qu'en espagnol.

Les apprenants qui parlent un dialecte bavarois normalement n'utilisent pas la consonne voisée [z] en position initiale du mot, comme le font les locuteurs du bas-allemand (cf. Ehrlich 2009: 92). Ces derniers prononceraient le mot allemand *Sonne* [zɔnə] 'soleil', contrairement aux locuteurs bavarois, qui le prononceraient plutôt avec la consonne sourde [s], le résultat étant la prononciation comme [sɔnə] ou désonorisée comme [zɔnə]. Ce fait peut causer une interférence en français, par exemple dans un mot qui exige le [z] voisée initiale, comme fr. *zéro* [zɛʁo]. Un apprenant qui parle un allemand bavarois pourrait le prononcer comme [sɛʁo]. En revanche, les locuteurs d'un dialecte bas-allemand ont un problème pareil avec le même groupe de consonnes, à savoir [s] et [z], en espagnol. Dans le cas de cette langue, on ne trouve pas de [z] initiale et les apprenants d'une région parlant un dialecte bas-allemand sont plus susceptibles de prononcer le mot espagnol *sol* [sol] 'soleil' comme [zol], et il est peu probable que les apprenants qui parlent un dialecte bavarois fassent ce type d'interférence. Kolly (2013) décrit des interférences observées chez les locuteurs des dialectes suisse allemands qui ont même provoqués la création de l'helvétisme du *français fédéral*, un français marqué par un accent suisse allemand, notamment la prononciation désonorisée ou sourde de [b] dans *bien* [bjɛ̃] comme [pjɛ̃] ou [ʒ] dans *je* [ʒə] comme [jə] (cf. Kolly 2013: 52).

Donc nous voyons déjà que les interférences causées par les caractéristiques dialectales ne sont pas rares. Pour permettre une analyse plus précise des interférences

dialectales entre l'allemand et les langues étrangères dans ce travail, il fallait se concentrer sur une caractéristique plus spécifique. Le groupe de sons sur lequel j'ai travaillé est celui des consonnes rhotiques. Beaucoup de linguistes ont déjà travaillé sur ce groupe consonantique à cause de sa diversité et à la fois, son uniformité. Le graphème <r> existe dans la plupart des langues du monde, mais la manière de le prononcer diverge souvent selon la langue ou le dialecte observé. Wiese (2001), dans un article sur les consonnes rhotiques dans la langue allemande, résume dans la citation suivante la question de dite uniformité et diversité simultanée:

[...] while r-sounds seem to vary greatly and in several dimensions, there is yet a unity to these r-sounds which justifies the common usage in phonetics and linguistics to talk about „r“ and the class of „r-sounds“ or „rhotics“. (Wiese 2001: 11)

Dans cette observation, il aborde la question de la justification de regrouper des sons tellement diverses sous un seul graphème <r>. La manière de prononcer ce graphème, qui est très répandu dans cette forme dans les langues du monde, est très variée. Quand on regarde la prononciation de *r* dans les langues du monde, on trouve un florilège de possibilités de le prononcer. Par exemple en espagnol, en fonction d'où se trouve le *r* dans le mot, on peut le prononcer de deux façons. En position initiale, on prononce la consonne roulée alvéolaire voisée [r] comme dans esp. *río* [rio] ,fleuve', et en position intervocalique on prononce la consonne battue alvéolaire voisée [ɾ], comme dans esp. *pero* [pɛro] ,mais'. En anglais américain standardisé, le même graphème <r> est prononcé comme l'approximant alvéolaire voisé [ɹ], comme dans angl. *rye* [ɹaɪ] ,seigle', et en anglais britannique standard, parfois on ne le prononce pas du tout à cause du phénomène du *r-dropping* (cf. Biersack 2002: 61), comme dans le mot angl. *car* [ka:] ,voiture'. Mais nous trouvons le graphème <r> aussi dans les langues qui se servent plutôt d'une prononciation uvulaire, comme l'allemand ou le français, où on peut observer les consonnes [ʀ], [R] ou [χ] et beaucoup plus.

Néanmoins, le phénomène du groupe des consonnes rhotiques ne se limite pas aux langues indo-européennes. En hausa, une langue appartenant à la famille linguistique des langues chamito-sémitiques, on trouve <r> dans l'orthographe, qui peut désigner deux phonèmes, à savoir la consonne roulée alvéolaire sonore [r], comme dans hausa *bara* [bɛrɔ] ,mendier', ou la consonne battue rétroflexe voisée [ɽ], comme dans hausa *bara* [bɛɽɔ] ,serviteur' (cf. International Phonetic Association ¹⁵2014: 90).

Il est donc évident que ni le point, ni le mode d'articulation ne sont des éléments cruciaux pour classer tous ces sons dans un groupe. Une caractéristique commune observée dans ce champ de recherche est leur occurrence en mêmes positions dans la syllabe dans des langues différentes. Lindau (1985) observe les rhotiques dans les groupes consonantiques près du noyau syllabique, ainsi que leur tendance à être élidés en faveur d'une voyelle à proximité, ou leur influence sur des voyelles dans l'environnement phonétique (cf. Lindau 1985: 157). Wiese (2001) soupçonne la légitimation du groupement des rhotiques dans une classe de sons à cause de leur comportement phonotactique, et leur donne une propre place dans le classement sur l'échelle de sonorité, entre // et les semi-voyelles, ainsi que les voyelles fermées (cf. Wiese 2001: 13). Lindau (1985) a donc proposé de grouper les rhotiques dans un modèle familial, dans lequel une manière de prononcer est toujours apparentée à une autre, qui est à la suite apparentée à encore une autre, qui peut nous donner une solution au problème de l'hétérogénéité de cette classe consonantique (cf. Lindau 1985: 166).

Dans les langues sur lesquelles j'ai travaillé, nous sommes encore une fois confrontés à la diversité dans les trois langues observées, ainsi que la diversité dans la même langue. En français et en espagnol standardisé nous trouvons des manières très différentes de prononcer le *r*. En français standard, on utilise la consonne fricative uvulaire voisée [ʁ] et l'allophone sourd très souvent utilisé, [χ], en espagnol on trouve les phonèmes [r] et [r̄]. En allemand standard, la littérature diverge en ce qui concerne le standard répandu. Beaucoup de manuels désignent la consonne roulée uvulaire voisée [ʀ] comme le standard, pendant que d'autres observent la prévalence de la consonne fricative uvulaire voisée [ʁ] ou même l'approximant uvulaire voisée [ɐ̯] (cf. Wiese 2001: 14). Dans la littérature on trouve aussi la consonne roulée [r] ou battue [r̄] comme assez répandue en allemand, par exemple dans les dialectes bavarois (cf. Krech 2009: 266). Un allophone de *r* en allemand peut aussi être une voyelle, dans notre cas, cette voyelle est [ɐ], normalement elle remplace un *r* consonantique dans la coda de la syllabe dans la plupart des variétés et des dialectes de l'allemand, à l'exception du suisse allemand et les dialectes alémaniques (cf. Panizzolo 1982: 30).

On peut donc constater, en prenant en compte tout ce que nous avons appris sur le potentiel pour des interférences qui pose la L1 ou ses dialectes, que le caractère hétérogène de l'allemand pourrait mener aux interférences spécifiques dans les langues

cibles. Dans la partie suivante je vais décrire la méthodologie appliquée pour faire des recherches sur ce sujet, ainsi que les résultats.

Méthode

Pour pouvoir faire des recherches sur la prononciation dialectale il faut des participants. Les locutrices desquelles je pouvais utiliser le matériel audio avaient fait partie d'un projet inédit sur la phonétique dialectale et l'apprentissage des langues étrangères, mené par Dr. Elissa Pustka (Universität Wien) et Prof. Dr. Christoph Gabriel (Johannes-Gutenberg-Universität Mainz). Les enregistrements suivaient toujours un certain schéma. Les locutrices devaient toutes être des jeunes femmes entre 18-16 ans, devaient venir de la région dialectale en question et devaient toutes parler le français et l'espagnol au niveau A2-B1. Les enregistrements ont été faits dans les villes mentionnées auparavant, par Isabella Rechberger (à Steyr et à Krieglach), Luise Jansen (à Munich), Jonas Grünke (à Hambourg), Prof. Dr. Stephan Schmid (à Zurich), et à Vienne je faisais partie de l'équipe qui s'occupait du recrutement et des enregistrements.

Les locutrices devaient toujours lire à voix haute un texte, d'abord en allemand, ensuite en français, puis en espagnol. Le texte en question est „La bise et le soleil“, un texte qui existe dans plusieurs langues et qui est souvent utilisé dans la recherche linguistique (cf. International Phonetic Association ¹⁵2014). À partir de ces enregistrements, j'ai préparé des transcriptions phonétiques à l'aide du programme PRAAT, pour pouvoir voir les sonagrammes et pour y mettre les transcriptions phonétiques, encore segmentées en syllabes et en sons individuels. Après j'ai fait mes analyses des consonnes rhotiques sur la base du matériel audio. J'ai identifié toutes les occurrences dans les textes où le *r* peut être prononcé. Pour effectuer des tests statistiques plus spécifiques, j'ai segmenté les occurrences de *r* dans sa position dans la syllabe, c'est-à-dire si le *r* se trouve dans l'attaque ou la coda, et s'il est ramifié avec une autre consonne ou non.

Après, j'ai tout rassemblé dans un tableau Microsoft Excel, classé par locutrice, langue et manière de réalisation du *r*. À la suite de cela, j'ai transféré ce tableau dans le programme SPSS pour pouvoir effectuer un test de signification χ^2 . Ce genre de test peut montrer s'il y a une corrélation statistique entre deux éléments dans un tableau, dans notre cas c'est le nombre d'une certaine manière de prononcer *r* en corrélation avec la région dialectale de laquelle vient la locutrice. Je n'ai pas seulement effectué

des tests avec les catégories assez restreintes, c'est à dire les catégories qui prennent en compte la position dans la syllabe et une ramification éventuelle, mais aussi avec les catégories plus larges, par exemple seulement la prononciation d'un certain *r* sans respecter la position syllabique, ou la création d'une nouvelle catégorie qui incluait plusieurs manières de prononcer le *r*.

Résultats

Comme je l'ai déjà mentionné dans la partie sur mon approche méthodique, les catégories qui incluent une différenciation très spécifique sont assez petites, c'est à dire qu'elles comptent très peu d'occurrences par locutrice. Ainsi il n'est pas très surprenant que les résultats significatifs se trouvent plutôt dans les catégories plus larges, même si elles sont également assez petites et les résultats significatifs sont rares.

Il y a néanmoins des résultats intéressants quand on regarde les nombres absolus d'occurrences et leur distribution sur les régions dialectales. Il était, par exemple, possible de voir une distribution de schwa [ə] dans les enregistrements allemands qui est congruente à ce que nous dit la littérature, à savoir que ce son se trouve en grand nombre chez les locutrices allemandes et autrichiennes, et on ne le trouve quasiment pas chez les locutrices suisses. En chiffres, cela veut dire que dans les régions bava-roises on trouvait 100 réalisations de [ə] par ville, à Hambourg on pouvait observer plus de 80, mais à Zurich je n'ai compté que 37 occurrences. Une tendance intéressante concerne une autre manière de vocaliser *r* dans la coda en allemand, à savoir la prolongation de la voyelle avant *r* et de laisser tomber complètement la consonne. Cette tendance était plus prévalente à Hambourg que dans les autres villes. Malheureusement, les tests dans SPSS ne montraient pas de résultats significatifs concernant la prévalence d'une certaine manière de prononcer le *r* comme consonne en relation avec la provenance de la locutrice, mais de même on peut observer une tendance en regardant les chiffres. Ce qui n'est pas surprenant est la prévalence plus fréquente de la consonne battue alvéolaire voisée [r] en Suisse, car les locutrices la prononcent également dans la coda d'une syllabe, contrairement aux autres locutrices de l'Allemagne ou de l'Autriche.

À partir de ces observations dans la L1 on peut maintenant continuer à regarder les résultats dans les langues étrangères, commençant par le français, suivi par l'espagnol.

Dans le corpus français il n'y avait pas de résultats significatifs en ce qui concerne la prononciation des *r* consonantiques en corrélation avec l'origine de la locutrice. Néanmoins, on peut voir des tendances aux interférences en regardant certains schémas. Par exemple, les locutrices qui, dans leur L1, prononcent plutôt la consonne [r] ont aussi plutôt tendance à réaliser [ə] dans la coda en français. Celles qui prononcent une consonne uvulaire, quoi qu'elle en soit, en allemand, ne font pas cette substitution aussi fréquemment. L'exception de cette observation concerne les locutrices suisses, dont la plupart prononcent [r] en allemand, mais elles n'ont pas tendance à substituer par [ə] dans la coda.

Une observation générale à faire est la tendance de toutes les locutrices à prononcer la consonne fricative uvulaire sourde [χ] à la place de la consonne fricative uvulaire voisée [ʁ]. Une autre observation intéressante est la prononciation assez fréquente de l'approximant uvulaire voisée [ʁ], surtout à Hambourg, où cette consonne apparaît le plus fréquemment, aussi dans le corpus allemand.

Passons à l'espagnol, où les résultats significatifs sont également rares, mais où nous pouvons aussi voir des tendances dans les chiffres. D'abord il faut mentionner un résultat significatif, à savoir celui concernant le nombre des [r] prononcés en corrélation avec l'origine de la locutrice. La signification de cette corrélation peut être expliquée en regardant les chiffres de Vienne et de Zurich, où les locutrices ont prononcé le plus grand nombre de consonnes battues alvéolaires [r], ainsi qu'un petit nombre de la variante désonorisée [ɾ], ou la consonne battue rétroflexe désonorisée [ɽ]. Un autre résultat significatif concerne la prolongation de la voyelle à la place de *r* consonantique. Ici on trouve également une corrélation. Cette corrélation peut être due au fait que les locutrices hambourgeoises ont tendance à substituer la consonne par la voyelle prolongée plus fréquemment que les autres.

En ce qui concerne la corrélation entre le nombre de schwa [ə] au lieu de *r* consonantique nous n'avons pas de résultat significatif, mais nous trouvons encore une fois des détails intéressants dans les chiffres absolus. Encore une fois, la tendance à prononcer [ə] dans la coda se laisse observer plus fréquemment chez les locutrices qui, dans leur L1 prononcent plutôt des consonnes uvulaires. Il faut néanmoins constater que ce cas d'interférence n'est pas aussi congruent que dans le corpus français, parce qu'on peut observer quelques locutrices qui prononcent des consonnes alvéolaires en allemand et, malgré ce fait, le substituent par schwa [ə] en espagnol.

Une dernière observation que j'aimerais mentionner est l'occurrence des éléments désonorisés. Ces occurrences ne semblent avoir rien à voir avec l'origine de la locutrice dans le corpus mais plutôt avec la prévalence des consonnes normalement voisées désonorisées dans un environnement d'autres consonnes sourdes, comme dans les parties du texte espagnol *fuerte*, *estrechamente*, ou *norte*. Une explication peut venir de l'assimilation régressive, c'est-à-dire l'influence des consonnes sourdes sur les consonnes voisées dans leur environnement, comme c'est le cas dans les mots mentionnés.

Conclusion

Pour conclure, il reste à dire que les dialectes semblent jouer un rôle qu'on ne devrait pas négliger dans le contexte de l'apprentissage de la prononciation d'une langue étrangère. Nous avons vu que certaines caractéristiques linguistiques régionales se manifestent également dans la prononciation de la L2, causant des interférences. Clairement, la L1 n'est pas le seul facteur qui influence notre prononciation dans la langue étrangère, comme nous avons observé à la fin du dernier chapitre, mais selon les résultats et l'état de la recherche actuel on peut constater que des recherches plus approfondies sont tout à fait légitimes.

Zusammenfassung in deutscher Sprache

Die vorliegende Diplomarbeit beschäftigt sich mit der Rolle verschiedener deutscher Dialekte beim Erlernen der Aussprache des Französischen und Spanischen. Zunächst werfen wir einen Blick auf den Stand der Forschung, um zu sehen, ob das Erlernen einer akzentfreien Aussprache in der L2 überhaupt möglich ist und deren gezielte Instruktion somit sinnvoll. Hierzu gibt es bereits zahlreiche Studien, von denen ein beachtlicher Teil zu dem Schluss kam, dass akzentfreies Sprechen nach dem Verstreichen einer sogenannten kritischen Phase (vgl. Lenneberg 1967) nicht mehr, oder nur sehr schwer möglich zu erlernen ist. Da der Lerner oder die Lernerin bei Konfrontation mit einem unbekannten Laut in der Zielsprache dazu neigt, sich des nächstverwandten Lautes aus der L1 zu bedienen und in weiterer Folge einen Laut „dazwischen“ zu produzieren, entsteht ein Akzent. Selinker (1972) nennt dieses Zwischenstadium *interlanguage*, also Interrimssprache. Lernerinnen und Lerner können in diesem Stadium stecken bleiben, dann spricht man von Fossilisierung (vgl. Selinker/ Gass 1992: 197). Andere Studien, wie beispielsweise jene von Birdsong (2007), stellten jedoch fest, dass das Erlernen einer von ihm so bezeichneten *nativelike pronunciation* (vgl. Birdsong 2007: 100) sehr wohl ein erreichbares Ziel sein kann, wenn die Lernenden entsprechende Instruktion erhalten sowie hohe Motivation aufweisen können (vgl. Birdsong 2007: 110). Nun haben wir festgestellt, dass Interferenzen aus unserer L1 zu einem Akzent führen können, jedoch werden wir mit der Problematik konfrontiert, dass der Begriff der L1 oftmals auf Standardvarietäten der Sprache beschränkt ist und Varietäten sowie Dialekte außer Acht lässt. Die Sprachen der Welt unterscheiden sich in ihrer dialektalen Form auf verschiedensten Ebenen, sowie im Bereich der Aussprache. So drängt sich die Annahme auf, dass dialektale Phonologie für so manche Interferenzen verantwortlich sein könnte und die Instruktion der fremdsprachigen Lautung somit regional spezifiziert werden müsste. Um dies zu testen, fokussierte ich mich auf die Gruppe der *r*-Laute oder *rhotics*, die als Lautklasse bereits viel Aufmerksamkeit in der Forschung erhalten haben und aufgrund ihrer Vielfalt, oftmals sogar innerhalb der gleichen Sprache, von Interesse sind. Um dialektale Unterschiede in den verschiedenen Sprachen testen zu können, untersuchte ich Informantinnen aus sechs deutschsprachigen Orten. Dies sind Hamburg, München, Wien, Steyr (Oberösterreich), Krieglach (Steiermark) und Zürich. Von jeder Informantin standen mir drei vorgelesene Texte in Form von Aufnahmen zur Verfügung, jeweils einmal auf Deutsch, Französisch und

Spanisch. Aus diesen Aufnahmen arbeitete ich mittels des Programms PRAAT (Borsma / Weenik) die verschiedenen Realisierungen von *r* heraus und überprüfte anschließend deren Okkurrenzen in Silbenposition sowie etwaige Verzweigung in Zusammenhang mit der Herkunft der Informantin.

Die Überprüfungen ergaben durchaus ein paar interessante Merkmale in der Aussprache, die auf den Dialekt der Sprecherin zurückzuführen waren. Dadurch, dass der Datensatz relativ klein war, kam es zwar zu kaum statistisch relevanten Ergebnissen, jedoch lässt sich eine Tendenz erkennen, die weitere Forschung in Richtung der Entwicklung einer varietätensensibleren Aussprachedidaktik legitimiert.

Abstract in English

This diploma thesis focuses on the process of learning the pronunciation of a foreign language. Mastering the aspects of nativelike pronunciation in a foreign language comes with a variety of difficulties for the learner, some linguists even argue that, after having reached a certain age, also called the *critical period* (see Lenneberg 1967), reaching proficiency is impossible. If that is the case, why do we even bother with the instruction of pronunciation? What research has shown us, is that being able to speak a foreign language like a native speaker comes with a wide range of benefits for the learner. Among said benefits we find that nativelike pronunciation leads to being perceived as more sympathetic and intelligent (see Edwards 1999: 2), thus being accepted into a certain social group. Additionally, ridding themselves of a foreign accent also ranks as one of the highest goals of learners when they are questioned about their objectives (see Harlow/Muyskens 1994: 145). When talking about what constitutes a foreign accent we often hear that the L1 of the learner is responsible for most of the interferences that happen on the segmental and suprasegmental level. The learner hears a sound in their target language that is new to them and tries to reproduce it by taking whatever sound in their L1 inventory is closest to the new sound, hence creating a sound that is neither exactly like the one in their L1, nor like the new L2 sound. This type of accent-marked speech is called interlanguage (see Selinker 1972: 212) and it is argued that most learners will not be able to surpass this stage completely. Some studies that focused on foreign accents showed that a certain amount of learners *do* get past the stage of interlanguage, when they received specific instruction and showed high motivation (see Birdsong 2007). What I attempted to show in this diploma thesis, is that L1 interferences can and do happen frequently, but that we might be looking at too narrow of a definition of what the learner's L1 actually is. Most languages of the world know many varieties and dialects, thus making it almost impossible to focus on the standard of the L1 when talking about interferences on a phonological level. My hypothesis is predicated on the assumption that what constitutes a foreign accent is oftentimes a result of the learner's dialect phonology, more so than the standard pronunciation. My research focused on German as a native language and French and Spanish as target languages. To be able to analyze learners' pronunciation, I centered my research on a very specific set of sounds, namely *rhotics* or *r-sounds*. As a group, rhotics are very diverse and tend to be interference-prone for learners. The re-

search for this thesis was done by analyzing the pronunciation of female native German speakers that were learning both French and Spanish from six different cities or towns in Austria, Switzerland and Germany.

My findings showed that there are indeed some differences in *r*-pronunciation depending on the learner's origin. The relatively small sample size made an in-depth analysis difficult and statistically significant results were thus rare. Nonetheless, I was able to show certain patterns in regional differences among some of the speakers and laid what I hope could be the groundwork for further research in this field. The aim of future studies could be to optimize foreign language pronunciation teaching with a focus on regional adaptation of certain aspects.

Eidesstattliche Erklärung

Ich erkläre hiermit, dass ich vorliegende Arbeit selbstständig verfasst habe und nur Hilfsmittel benutzt habe, die von mir durch Belege gekennzeichnet wurden. Ich habe keinerlei unerlaubte Hilfe in Anspruch genommen. Alle Quellen, sowohl in Textform als auch online, die für diese Arbeit herangezogen wurden, wurden mit einer genauen Angabe ihrer Herkunft als Kurzbeleg sowie in der Bibliographie genannt.

Ich habe die Arbeit bislang weder im In- noch im Ausland in irgendeiner Form als Prüfungsarbeit vorgelegt und nicht veröffentlicht.

Wien, 2017