

MASTERARBEIT / MASTER'S THESIS

Titel der Masterarbeit / Title of the Master's Thesis

**“Attitudes toward Original Pronunciation
Performances of Shakespeare”**

verfasst von / submitted by

Mag. Sergio Piaia, BEd

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

Master of Education (MEd)

Wien, 2018 / Vienna 2018

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

A 199 507 517 2

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

Masterstudium Lehramt Sek (AB) Lehrverbund
UF Englisch Lehrverbund
UF Italienisch Lehrverbund

Betreut von / Supervisor:

Univ.-Prof. Mag. Dr. Nikolaus Ritt

Attitudes toward Original Pronunciation Performances of Shakespeare

Sergio Piaia

Master's Thesis, University of Vienna

Abstract

Plays by Shakespeare staged in reconstructed Elizabethan speech are known as “Original Pronunciation” (OP) performances. Theater practitioners and scholars, such as David Crystal, argue that Original Pronunciation can bring modern audiences “closer” to Shakespeare. To validate these claims empirically, a language attitudes study was conducted by means of the matched guise technique. With an online questionnaire composed of both semantic differential rating scales and direct questions, participants were asked to evaluate four different recordings of Shakespeare’s *Sonnet 116*, recited in General American, Contemporary RP, Conservative RP, and Original Pronunciation. Responses were obtained from an equal number of native speakers and non-native speakers of English, for comparative purposes. Analysis of the data shows that reconstructed Elizabethan speech was rated lowest along the evaluative dimension *prestige* (social status) and was also deemed less suitable for performances of Shakespeare than the modern British standard varieties. Original Pronunciation was most closely associated with Irish English by a majority of the respondents and very clearly evaluated as the least intelligible variety. Nevertheless, the high evaluations of Original Pronunciation along the *accessibility* (social attractiveness) dimension—particularly by British English native speakers—point to a “covert prestige” that this non-standard variety enjoys as a medium for performances of Shakespeare. Although the results of the study suggest that a widespread adoption of Original Pronunciation may not take hold, they nevertheless appear to encourage the proselytizing efforts of the “OP movement.”

Keywords: language attitudes, Original Pronunciation, Shakespeare

Speak the speech, I pray you, as I pronounced
it to you—trippingly on the tongue; but if you mouth
it, as many of your players do, I had as lief the town-
crier had spoke my lines.

Hamlet (3.2.1–4)

Acknowledgments

Words cannot express my gratitude to Professor Nikolaus Ritt. This thesis would not have been conceivable without his invaluable input and enthusiastic support.

My particular appreciation goes to all the unnamed participants out there who patiently responded to the online questionnaire.

The unwavering support of my family made all of it possible. Again.

Table of Contents

List of Tables.....	iii
List of Figures.....	iv
1. Introduction: “Original Pronunciation”.....	1
2. Shakespearean Phonology and Original Pronunciation Performances.....	4
2.1. The Phonological Reconstruction of Shakespearean Pronunciation and Original Pronunciation Performances	4
2.2. The Evidence for the Reconstruction of Shakespearean Phonology	7
2.3. Distinctive Phonological Features of Original Pronunciation.....	10
2.3.1. <i>Short vowels</i>	10
2.3.2. <i>Long vowels</i>	11
2.3.3. <i>Diphthongs</i>	12
2.3.4. <i>Consonants</i>	13
2.4. Implications of Original Pronunciation for Performances of Shakespeare	15
3. Language Attitudes	19
3.1. Defining Language Attitudes	19
3.1.1. <i>Attitudes</i>	19
3.1.2. <i>Language attitudes: social categorization and stereotyping</i>	20
3.2. The Measurement of Language Attitudes.....	22
3.2.1. <i>Main approaches</i>	22
3.2.2. <i>Matched and verbal guise studies: the “speaker evaluation paradigm”</i>	24
3.2.3. <i>Rating scales</i>	27
3.3. Attitudes toward Shakespeare and Shakespearean Language	29
4. The Research Design of the Study	32
4.1. Aims and Hypotheses	32
4.2. The Research Instrument	34
4.3. The Audio Stimuli.....	35
4.3.1. <i>Speaker A: General American</i>	35
4.3.2. <i>Speaker B: Contemporary RP</i>	36
4.3.3. <i>Speaker C: Conservative RP</i>	36
4.3.4. <i>Speaker D: Original Pronunciation</i>	36

4.4. The Questionnaire	39
4.4.1. <i>Piloting phase</i>	39
4.4.2. <i>Final questionnaire</i>	43
4.5. Questionnaire Administration and Data Collection	45
5. Results and Discussion	47
5.1. Description of Participants	47
5.1.1. <i>Entire sample</i>	47
5.1.2. <i>Native speaker respondents</i>	48
5.1.2. <i>Non-native speaker respondents</i>	50
5.2. Analysis of the Data	51
5.3. Main Results of the Analysis	53
5.3.1. <i>Native speaker respondents</i>	53
5.3.2. <i>Non-native speaker respondents</i>	56
5.4. The Speech Variety Variable	59
5.4.1. <i>Native speaker respondents</i>	59
5.4.2. <i>Non-native speaker respondents</i>	61
5.5. The Tertiary Education Variable	64
5.6. The Personal Relevance of Shakespeare Variable	66
5.7. Discussion of Results	68
5.7.1. <i>Prestige</i>	68
5.7.2. <i>Accessibility</i>	69
5.7.3. <i>Intelligibility</i>	71
5.7.4. <i>Suitability for Shakespeare</i>	72
5.7.5. <i>Ranking order of preferences</i>	73
5.7.6. <i>Modern-day speech varieties associated with Original Pronunciation</i>	74
5.8. Limitations of the Study	75
6. Summary and Conclusion	77
References	81
Appendix	87
Zusammenfassung in deutscher Sprache	93

List of Tables

Table 1. Individual Components of Evaluative Dimensions.....	51
Table 2. Mean Scores of Individual Traits Expressed by NSR.....	54
Table 3. Mean Scores of Individual Traits Expressed by Non-NSR	57
Table 4. Evaluative Dimensions by Tertiary English Degree of NSR	64
Table 5. Evaluative Dimensions by Tertiary English Degree of Non-NSR.....	65
Table 6. Evaluative Dimensions by Personal Relevance of Shakespeare for NSR	66
Table 7. Evaluative Dimensions by Personal Relevance of Shakespeare for Non-NSR ..	67

List of Figures

Figure 1. Full IPA chart.....	14
Figure 2. Pilot questionnaire: semantic differential scales.....	40
Figure 3. Pilot questionnaire: suitability for Shakespeare.....	41
Figure 4. Pilot questionnaire: speaker provenance	42
Figure 5. Final questionnaire: revised semantic differential scales	43
Figure 6. Final questionnaire: intelligibility of the speaker	43
Figure 7. Final questionnaire: background variables.....	44
Figure 8. Age range distribution of the study population	47
Figure 9. Level of tertiary English degree of participants.....	48
Figure 10. Personal relevance of Shakespeare for participants	48
Figure 11. Age range distribution of NSR	49
Figure 12. Speech varieties of NSR.....	49
Figure 13. Age range distribution of Non-NSR.....	50
Figure 14. Model speech varieties of Non-NSR.....	50
Figure 15. Evaluative dimensions expressed by NSR (entire sample).....	54
Figure 16. Rank ordering of speakers expressed by NSR (entire sample).....	55

Figure 17. Modern English varieties associated with Original Pronunciation by NSR (entire sample)	55
Figure 18. Evaluative dimensions expressed by Non-NSR (entire sample)	57
Figure 19. Rank ordering of speakers expressed by Non-NSR (entire sample)	58
Figure 20. Modern English varieties associated with Original Pronunciation by Non-NSR (entire sample)	58
Figure 21. Evaluative dimensions expressed by American English NSR.....	59
Figure 22. Evaluative dimensions expressed by British English NSR.....	59
Figure 23. Rank ordering distribution expressed by American English Non-NSR.....	60
Figure 24. Rank ordering distribution expressed by British English Non-NSR.....	60
Figure 25. Evaluative dimensions expressed by American English Non-NSR	62
Figure 26. Evaluative dimensions expressed by British English Non-NSR	62
Figure 27. Rank ordering distribution expressed by American English Non-NSR.....	63
Figure 28. Rank ordering distribution expressed by British English Non-NSR.....	63

1. Introduction: “Original Pronunciation”

In June of 2004, Shakespeare’s Globe in London staged three performances of *Romeo and Juliet* using reconstructed Elizabethan English pronunciation. It was the first time in half a century that Shakespeare’s words were spoken in the theater in such way as his contemporary audiences presumably would have heard them uttered during the Early Modern period. The event was carried out in close collaboration with the renowned linguist David Crystal, who subsequently recounted that experience in his book *Pronouncing Shakespeare: The Globe Experiment* (2005). The success of these performances and the positive feedback the initiators received from members of the audience during talkback sessions subsequently sparked a minor “Original Pronunciation” movement—or “OP,” as it is referred to by its proponents, chiefly among whom David Crystal’s own son, the actor Ben Crystal.

Advocates of Original Pronunciation adduce several reasons for staging Shakespeare in reconstructed Elizabethan speech: As a form of historically informed practice (Barrett 2013: 15), it offers actors previously untapped veins for the exploration of Shakespeare’s characters and their interactions; it also opens new avenues of interpretation for directors, dramaturgs, and audiences alike (Crystal 2016: x); above all, Original Pronunciation provides a unique opportunity of getting “closer to Shakespeare,” in the opinion of its proponents. They argue that the sound changes affecting the English language in the four centuries since Shakespeare’s death in 1616 have significantly altered the stage effects on which the Bard relied in his plays. These sound shifts have introduced irregularities in rhythm and meter, obscured puns and word-play quintessential to Shakespeare’s craft, and muddled many rhymes to the point where they do not work anymore in modern pronunciation. All of which, it is claimed, can be restored using Original Pronunciation. Significantly, reconstructed Elizabethan speech also offers “a new auditory aesthetic” for audiences because of its marked contrast with British Received Pronunciation (RP), the modern standard English variety in which many anglophone audiences, particularly in the United Kingdom, will most likely have experienced Shakespeare (Crystal 2016: x):

Those who speak with an accent other than RP (which in the UK comprises most of the population) say that OP reaches out to them in a way that RP does not,

primarily because they recognize in it echoes of the way they themselves speak. [. . .] OP thus offers a new kind of ‘ownership’ of Shakespeare.

Nevertheless, there are also some who balk at the notion of performing Shakespeare in reconstructed Early Modern English pronunciation, dismissing the entire concept of historically informed practice—which also includes the use of original costumes and staging techniques—as “nonsense” (Lodewyck 2013: 41). A chief objection raised against Original Pronunciation is the manifest impossibility of recovering authentic Early Modern English phonology in its entirety—a point that linguists gladly concede, stressing instead the “plausibility” of the reconstructed sound system (Crystal 2016: xxxix). There are also concerns that modern audiences will experience even more difficulty understanding Shakespeare’s language in the theater. In his review of Shakespeare’s Globe’s 2005 Original Pronunciation performance of *Troilus and Cressida*, theater critic John Lahr (2005: 98–99), for instance, estimated that he was able to comprehend “only about thirty per cent of the production” and that he had to read along to get the full meaning. Lahr’s critique, in which he maintained that “the waves of words produce a mesmerizing static, sort of like listening to poetry underwater,” also exemplifies the apprehension of many literati “that Shakespeare’s ‘sacred’ poetry” might be “sullied by a delivery unworthy of his ‘soaring’ words” (Lodewyck 2013: 44). Indeed, one of the touted strengths of Original Pronunciation—that it allows for “a new kind of ownership of Shakespeare”—may also harbor potential downsides: the echoes of contemporary nonstandard English varieties people hear in Original Pronunciation might weigh on spectators’ judgments with an entirely different load of associations. Tynan (1961: 29), for instance, noted of an Original Pronunciation performance of *Macbeth* in the 1950s that the “reconstructed dialect bears a violent likeness to the Stage Rustic, with a touch of Stage Dublin” and that “the poetry, in these accents, rings false.”

Thus far, no empirical study has been conducted to investigate contemporary audiences’ perception of Original Pronunciation (Lodewyck 2013: 55) as well as on its suitability as a performance medium in relation to the “normal” style of stage pronunciation (Barret 2013: 12). I would therefore like to contribute to the debate on Original Pronunciation with an empirical investigation of people’s attitudes toward performances of Shakespeare in reconstructed Elizabethan speech, which I accomplished with the matched guise technique, a research tool drawn from “classic” language attitudes

studies. In order to collect data from participants, I have devised an online questionnaire with a combination of direct questions and semantic differential scales in response to audio stimuli and obtained responses from an equal number of native speakers and non-native speakers of English, for comparative purposes.

In this thesis, then, I will first briefly outline the history of the reconstruction of Early Modern English Phonology and of the attempts to revive it for theatrical performances of Shakespeare. In particular, I will focus on the distinctive phonological features of Original Pronunciation as they differ from both present-day British Received Pronunciation and General American pronunciation. Then I will discuss the implications of using Original Pronunciation for performances of Shakespeare in more detail, based on selected examples from Shakespeare's plays and sonnets. Next, I will outline the theoretical concepts of language attitudes and the research instruments employed to investigate them, before presenting the research design of the study conducted for this thesis. After elaborating on the research questions and hypotheses generated for this study, I will dwell on the questionnaire, in particular on its key component the semantic differential scales. Lastly, I will present the results of the analysis of the data collected with the questionnaire and discuss those results in detail, as well as the conclusions that can be drawn from them.

2. Shakespearean Phonology and Original Pronunciation Performances

2.1. The Phonological Reconstruction of Shakespearean Pronunciation and Original Pronunciation Performances

Interest in reconstructing the phonology of Early Modern English has always been tied to a desire to recover the original pronunciation of Shakespearean language. Nevertheless, as Crystal (2013: 5) notes, Shakespearean phonology has been remarkably neglected by scholarly linguistics. (When not otherwise stated, this section and the following one are based on Crystal 2013 and 2016.) It is therefore hardly surprising that the study of Original Pronunciation first began among literary circles, in the form of philological attention to individual puns, rhymes, and metrical idiosyncrasies of the Bard as well as the prosody of his verss. A proper reconstruction of Elizabethan English pronunciation, however, was not attempted until the second half of the nineteenth century, when the American literary critic Richard Grant White (1865) published an appendix to an essay on Shakespeare in which he provided evidence of Early Modern English pronunciation through phonetic respellings of rhymes and puns. Curiously, White (1865: 438) also foresaw the typical reactions of modern audiences to the unfamiliar sounds of Shakespearean phonology:

They will imagine *Hamlet* exclaiming:– ‘A *baste* that wants *discoorse* of *rayson* / Would *haive moorn’d* longer!’ [. . .] and, overcome by the astonishing effect of the passages thus spoken, they will refuse to believe that they were ever thus pronounced out of Ireland.

Several publications that appeared in the subsequent decades testify to a nascent interest in the reconstruction of Early Modern English phonology on both sides of the Atlantic. In an exhaustive article for *North American Review*, for instance, Charles S. Peirce and J. B. Noyes (1864) surveyed the work of George Craik, Richard Grant White, and George P. Marsh on the subject. The latter also delivered a lecture at Columbia College in the early 1860s, titled “Orthoepical change in English.” In Britain, the major contributor to the field was the early phonetician Alexander Ellis, whose extensive tome *On Early English Pronunciation* (1869–1874) spanned over thousand pages, focusing in particular on Chaucer and Shakespeare. The fundamental novelty of Ellis’s work was the use of a phonetic alphabet termed “palaeotype symbols,” which he derived from Melville Bell’s

“Visible Speech” system. Ellis’s account, “unprecedented in its detail” (Crystal 2013: 8), would not be surpassed until the work of Eric J. Dobson almost a century later. But although his transcription system was a major step forward, Ellis’s work was still severely limited by the absence of a true phonetic alphabet, which would only be introduced in the 1880s by the International Phonetic Association. As interest in Early Modern English phonology continued to grow, especially in Germany, the introduction of the international phonetic alphabet would indeed prove to be a watershed. In 1906, for instance, Wilhelm Viëtor published two books on the subject using the new transcription system. His work, too, exhibited some limitations; for instance, he adopted a simplified transcription that did not distinguish between strong and weak vowels and did not account for many variant pronunciations. Nevertheless, it was strongly influential on subsequent scholars, including Daniel Jones, who first became a phonetician while studying German at the University of Marburg, where Viëtor was teaching.

Jones would go on to stage the first known public performance of Shakespeare in Original Pronunciation since the Early Modern period. With “Scenes from Shakespeare in the original pronunciation,” at University College of London in 1909, he himself took over the roles of Prospero from *The Tempest* and of Andrew Aguecheek from *Twelfth Night*. Though the event was so successful that Jones repeated it in December of the same year, plans for a shortened version of *Twelfth Night* in Original Pronunciation, however, never materialized. In the subsequent decades, Jones produced several radio programs on Shakespearean pronunciation for the BBC, confidently stating in 1949 that “we now have a pretty accurate picture of the way in which English pronunciation has developed from Anglo-Saxon times to the present day” (cited in Crystal 2013: 12). By the early 1950s however, a veritable dispute erupted when Kökeritz (1953: 6), whose opinion was that “Shakespeare’s pronunciation strongly resembled modern English,” heavily criticized the work of Viëtor and Jones. Kökeritz’s view was, in turn, strongly criticized by Dobson (1957) and later quite comprehensively refuted by Cercignani (1981: 28), who asserted that “the types of speech reflected in Shakespeare’s works and in those of contemporary writers on orthography and pronunciation reveal considerable discrepancies between Elizabethan and present-day standard usage”. To this day, Cercignani’s detailed study *Shakespeare’s Works and Elizabethan Pronunciation* (1981) remains “the latest and

fullest attempt to review all the evidence of rhymes, puns, spellings and metrics in the Shakespeare corpus” Crystal (2013: 13).

As Barrett (2013: 134) notes, “the early British experiments with OP in theatre and radio paved the way for more ambitious projects involving university departments and professional theatre.” In 1951, Bernard Miles staged the first scene of *Hamlet* in Original Pronunciation at the Mermaid Theatre in London; a year later, the first known complete Original Pronunciation performance of a play by Shakespeare in the modern era, *Julius Caesar*, was performed by the Marlow Society at the Arts Theatre in Cambridge (Barrett 2013: 137–139). This was followed, two years later, by the first Original Pronunciation production of Shakespeare in North America, staged at the Yale Shakespeare Festival by the university’s English department, for which Kökeritz served as a voice-coach (Barrett 2013: 140–143). After this remarkable surge in interest, however, there were no more Original Pronunciation productions until the end of the century, when the Chamberlain’s men’s historical playhouse was recreated on the south bank of the Thames, just a few hundred yards away from its original site.

Shakespeare’s Globe, the brainchild of the American theater impresario Sam Wanamaker, opened to the public in 1997 and has been devoted, from its inception, to what is referred to as “historically informed performances.” In the context of Shakespeare’s Globe, this term describes theatrical practices informed by the historical conditions of the Elizabethan stage, such as open-air setting, daylight performances, original costumes, and use of music. Curiously, for a theater devoted to recreating the Elizabethan stage experience, an production using Original Pronunciation was not attempted until 2004, when director Tim Carrol staged *Romeo and Juliet* simultaneously in modern English and in Original Pronunciation. The following year, Shakespeare’s Globe staged its first production performed solely in reconstructed Elizabethan speech. *Troilus and Cressida*, directed by Giles Block, enjoyed a considerably longer run than *Romeo and Juliet* had and was also greeted with overwhelmingly positive acclaim (Barrett 2013: 84). Despite the relative commercial success, however, Shakespeare’s Globe has to this date not repeated these experiments with Original Pronunciation. That mantle was taken up chiefly by academic circles in North America. In 2010, for instance, British-born Paul Meier staged a production of *A Midsummer Night’s Dream* at the University of

Kansas, after researching Original Pronunciation for a year. It was followed by a production of *Hamlet* at the University of Nevada in 2011, directed by Rob Gander. These North American academic experiments, too, heavily relied on the linguistic expertise of Crystal for the transcription and coaching process (Barrett 2013: 145). In the next part, I will briefly detail the three major sources of evidence from which linguists such as Crystal (2005, 2016)—on whom the following section is based, when not otherwise stated—draw upon for their reconstruction of Shakespearean phonology.

2.2. The Evidence for the Reconstruction of Shakespearean Phonology

As Nevalainen (2006: 118) notes, “[i]t is much easier to describe the grammar of Early Modern English than to account for its pronunciation.” Not surprisingly, Shakespeare’s works constitute a prime source of evidence for scholars seeking to reconstruct the sounds of the Elizabethan period. The current reconstruction of Early Modern English phonology is based on four kinds of evidence—painstakingly gathered, among others, by Kökeritz (1953), Dobson (1957), Barber (1976), and Cercignani (1981)—namely spellings, rhymes, puns, and observations by contemporary writers and orthoepists. Spelling was not yet standardized in the Early Modern period; therefore, the choices made by the various writers, scribes, and typesetters “often provide pointers as to how a word was pronounced”, and, “[w]ith no agreed spelling for a word, the way it was said was likely to influence the way it was spelled” (Crystal 2016: xxi). Nevertheless, the preeminent scholar on Elizabethan phonology, Cercignani (1981: 17), cautions that “[a] correct interpretation of orthographic evidence requires an unequivocal acceptance of the limitations of spellings as a source of information.” Because of the nature of their graphemic-phonemic correspondence, “spellings cannot be expected to offer any clue to the actual realization of a phoneme” and “can only be interpreted from a phonemic point of view.” Indeed,

we do not even know how Shakespeare pronounced his own surname, since the first part of his signature (*Shakspere* and *Shakspeare*) may imply either the antecedent of present *shake* or a variant pronounced like *shack*, while both *-spere* and *-speare* may conceal either the antecedent of present *spear* or a form rhyming with *pear*.

Nevertheless, there are several examples in which the spelling of a word unequivocally points toward a particular pronunciation. In the first act of *Romeo and Juliet* (1.4.64), for instance, where Mercutio describes Queen Mab as having a whip with “a lash of film,” the Folio and Quarto spellings of *philome*, for *film*, indicate a bisyllabic pronunciation, “*fillum*,” as is still occurs in modern Irish English (Meier 2011: 211).

Rhymes arguably deliver the most compelling evidence for original Elizabethan phonology. Shakespeare’s sonnets and narrative poems are a particularly reliable indicator of how words were most likely pronounced, because of their transparent rhyme scheme. But his plays, too, are a vital source of information. Fifty-five percent of all verse scenes (using the *Oxford Shakespeare* scene division) end in a rhyming couplet or have one close by, such as the following lines from *Romeo and Juliet* (2.2.93–94):

ROMEO

O, let us hence! I stand on sudden haste.

FRIAR LAURENCE

Wisely and slow. They stumble that run fast.

These instances when a rhyme ostensibly fails in Modern English have in the past frequently been dismissed as inexact rhyming and often drew the label “half rhymes,” “eye rhymes,” or even “traditional rhymes” that are “phonetically inexact in terms of Shakespeare’s pronunciation,” most notably by Kökeritz (1953: 246). Theatrical plays of the period, however, were primarily written for the stage and thus meant to be listened to rather than read. Indeed, as many scholars have demonstrated, Elizabethans had a unique “aural sensitivity” (Miola 2000: 2) and had an “intimate relation to an oral tradition of language” (Smith 1999: 13). Moreover, the large audiences flocking to the original Globe Theatre would have mostly been illiterate and presumably not much cared for “eye rhymes.” Cercignani (1981)—and, more recently, Crystal (2016)—have thus strongly and convincingly argued in favor of rhymes as internal evidence for Early Modern English pronunciation, particularly through cross-referencing and comparison of different spelling variants. The above mentioned rhyming couplet from *Romeo and Juliet* is an indicative example: in the First Folio Edition of Shakespeare’s works, *haste* is spelled as *hast*, which was the most common orthographical rendition of the word in that period, whereas *fast* rhymes with *last* in *Venus and Adonis* and with *past* in *The Rape of Lucrece*. Because there are many instances in Shakespeare’s works of words rhyming with *haste* and *fast*, it

can thus be safely concluded that both words would likely have been pronounced with short [a], similar to modern Received Pronunciation *fast*.

Shakespeare's fondness of punning, as Cercignani (1981: 12) notes, reflects the spirit of the Elizabethan age and can therefore offer valuable information on Original Pronunciation. In open dispute with Kökeritz (1953), who considered wordplay the veritable "be-all and end-all" of phonological evidence, Cercignani (1981: 13), however, cautions that only in those instances when the context can "reasonably assumed to *require* some kind of quibbling," can puns be safely adduced as phonological evidence, as Crystal also concurs. For instance, in the following exchange from *The Two Gentlemen of Verona* (2.1.1–2),

SPEED (*offering Valentine a glove*)
Sir, your glove.
VALENTINE Not mine. My gloves are on.
SPEED
Why then, this may be yours, for this is but one.

the context itself provides unequivocal proof of a pun, revealing that the pronunciation of *one* and *on* must have been at least similar enough to have warranted the witticism. This particular homophony is further confirmed, for instance, by a rhyming couplet in *A Midsummer Night's Dream* (3.2.118–119), where Puck quips

Then will two at once woo one
That must needs be sport alone;

The most reliable source of evidence for Early Modern English phonology, however, is found in the many commentaries on pronunciation and correct usage by Early Modern writers and orthoepists, in works such as John Hart's *Orthographie* (1569) and Ben Johnson's *English Grammar* (1640). Johnson (cited in Crystal 2016: xx), for example, describes the letter *o* as "[i]n the short time more flat, and akin to u; as [. . .] brother, love, prove," providing a clear indication that in the very frequently occurring rhyme *prove* and *love* the pronunciation of the former matches the present-day pronunciation of the latter. (This for modern ears quite startling pronunciation will also feature prominently in the stimulus material employed in the empirical study conducted for this thesis.)

Nevertheless, there is one significant roadblock to unequivocally establishing the one true "original" pronunciation of the English spoken during Shakespeare's lifetime. As Barber (1976: 103) notes, "[i]n the Early Modern period, there was considerable variety of

pronunciation, and probably more toleration of variety than in the couple of centuries that followed.” Although Elizabethan English pronunciation “clearly had many stable features,” it was also still undergoing a series of sound changes (Nevalainen 2006: 119), which accounts for this extraordinary diversity in speech. Crystal (2016), for his part, also stresses that the sound system of Shakespearean English would have been realized in a variety of different accents. Nevertheless a “common core” of shared pronunciation during Shakespeare’s lifetime clearly existed, and these reconstructed features now form the basis of Original Pronunciation phonology. In the next section (which is based on Crystal 2005 and 2016, when not otherwise stated), I will briefly outline the main distinctive sounds of Shakespearean phonology with regard to how they differ from both present-day British Received Pronunciation (RP) and present-day General American pronunciation (GenAm). To that end, I will refer to the standard lexical sets for modern English introduced by Wells (1982), as well as to the International Phonetic Alphabet. (A full chart of its most recent version is displayed in Figure 1 at the end of the section.)

2.3. Distinctive Phonological Features of Original Pronunciation

2.3.1. Short vowels

Most of the short vowels in Original Pronunciation exhibit only minor phonetic differences with Modern English, so that entire sentences in Shakespeare’s works, such as “Is she gone to the King?” (*All’s Well That Ends Well*, 2.5.19), would have been virtually undistinguishable in Elizabethan speech. Nevertheless, there are several short vowels that presumably had a more open quality than their modern equivalents. For instance, the vowel in the KIT set (RP: /ɪ/, GenAm: /ɪ/, as in *sit*, *will*, *prince*) was likely realized more in the direction of [e]. On the other hand, the vowel in the DRESS set (RP: /e/, GenAm: /ɛ/, as in *met*, *tell*, *hen*), seems to have been pronounced in the same manner as in present-day RP, although some speakers might have used a more open variant, closer to GenAm. Additionally, it appears likely that in Early Modern English both the TRAP set vowel (RP: /æ/, GenAm: /æ/, as in *cap*, *fat*, *lamb*) and the BATH set vowel (RP: /ɑ/, GenAm: /æ/, as in *fast* and *father*) were pronounced as [a], akin to present-day northern British dialects—thus marking a strong contrast with long [ɑ:], which is a feature of present-day RP in words such as *bath* and *past*. Words including *haste*, *taste*, and *any*

were also pronounced with this vowel, as is evidenced by rhymes with *fast* and *blast*. Whereas the vowel in the STRUT set (RP: /ʌ/, GenAm: /ʌ/, as in *cup*, *stuff*, *drum*) was realized more similar to [ɤ]—that is, further back and more closed than its contemporary counterpart sound in RP and GenAm: “Opinions vary as to how far back it would have been, with values proposed between [ə] and unrounded cardinal 7 [ɤ]” (Crystal 2016: xliii). Crystal deems the latter more likely, since the *u* spelling was the norm for this vowel, but notes that there were several instances of overlap with *o*. By contrast, the FOOT set vowel (RP: /ʊ/, GenAm: /ʊ/, as in *put*, *look*, *foot*, and also in *fool* and *tooth*) presumably had the same phonetic value as in modern-day Conservative RP, although there is a degree of uncertainty regarding the extent to which it was used as a variant pronunciation in words with long /u:/. The back vowel in the LOT set (RP: /ɒ/, GenAm: /ɑ/, as in *hot*, *fog*, *bog*) may have been more open than its present-day RP equivalent [ɒ]. Nonetheless, it most probably featured a certain degree of lip-rounding and was pronounced fairly short, which would differentiate it from the [ɑ] sound in modern GenAm words such as *hot*. The use of this vowel in words such as *tongue* and *youth* is evidenced by rhymes and occasional puns, and was, moreover, “an important metrical alternative in such words as *satisfaction*: ,satis'faksɪən, -sɪən” (Crystal 2016: xliv). On the other hand, in the Early Modern period a short open back [ɑ] vowel followed a /w/ or /ʌ/ consonant in words such as *what* or *watch*, rather than the rounded [ɒ] of present-day RP, or the or the more open [ʌ] of contemporary GenAm. A similar effect can be observed in /ɑ/ vowels followed by velar /l/, in words such as *festival*, *special*, or *burial*, which in Modern English exhibit a schwa-like quality, whereas in Elizabethan English they were realized with a short [ɑ] sound.

2.3.2. Long vowels

The so-called “meet-meat” merger, by which Middle English /ɛ:/ and /e:/ both became /i:/ in Modern English (known as the FLEECE lexical set), had been completed in standard English accents by the end of the seventeenth century (Wells 1982: 195). Although “[i]t is not clear just how far a merger would have taken place by the end of the sixteenth century, or which words would have been affected,” there is “a consensus that the gradual rising in this part of the vowel-space still had some way to go before reaching the present-day

value of /i:/” (Crystal 2016: xlv). Gimson (1962) has shown that this particular sound presumably was realized nearer the cardinal vowel /e/ during Shakespeare’s period. The FACE set vowel (RP: /eɪ/, GenAm: /eɪ/, as in *day, place, make*) was pronounced with the monophthongal quality of the present-day SQUARE set (RP: /ɛ:/, GenAm: /ɛ:/) in words such as *fair, hare, and there*, as well as in words that would evolve to /i:/, such as *reason* and *season*; as a variant pronunciation, [ɛ:] also occurred in such words as *here*, which is evidenced by various rhymes and puns in Shakespeare. The NURSE set (RP: /ɜ:/, GenAm: /ɜr/, as in *bird, mercy, sir*) was pronounced /ɛ:/ before /ɹ/, a sound combination that is still retained in many regional British and North American accents. The vowel in the NORTH set (RP: /ɔ:/, GenAm: /ɔ:r/, as in *war, warm*) and THOUGHT set (RP: /ɔ:/, GenAm: /ɔ/, as in *wall, all, fall*) was pronounced with a long open back /ɑ:/. For this particular sound, Crystal (2016: xlv) notes that

[i]t must have been a noticeable feature of OP as Jonson, among others, pays special attention to it, contrasting it with the normal use of *a* (‘pronounced less than the French à’): ‘when it comes before l, in the end of a syllabe, it obtaineth the full French sound, and is uttered with the mouth and tongue wide opened, the tongue bent back from the teeth’. He gives *all, small, salt, calm* among his examples.

Words in the GOAT set (RP: /əʊ/, GenAm: /o/, as in *go, soul, moan*) were pronounced with the long monophthong /o:/ used in many present-day English accents, most notably of the Celtic area: “Rhymes show its use as a variant in words that later would have more open vowels, such as *one / throne, none / bone*” (Crystal 2016: xlv). The value in the GOOSE set (RP: /u:/, GenAm: /u/, as in *do, shoe, spoon, new, cure*), on the other hand, seems to have been identical to present-day conservative RP, although several words that today are pronounced with /u:/, such as *fool*, could be shortened.

2.3.3. Diphthongs

A particularly distinguishing feature of Early Modern English phonology, compared with modern standard pronunciations, is the centralized diphthong /ɔɪ/ in words of both the PRICE set (RP: /aɪ/, GenAm: /aɪ/, as in *my, sigh, fright, mile*) and the CHOICE set (RP: /ɔɪ/, GenAm: /ɔɪ/, as in *joy, boy, enjoy*). The identical pronunciation of these two diphthongs in Elizabethan English—which are clearly contrasted in present-day English in words such as *voice / vice*, or *boil / bile*—was an important source of puns for

Shakespeare. Furthermore, a most prominent feature of Shakespearean phonology compared with present-day English is the occurrence of /əɪ/ at the end of words in the so-called “happyY” lexical set (RP: /i/, GenAm: /i/, as in *lively*, *ready*, *chastity*); though the diphthong in this case is an unstressed syllable and was spoken very rapidly. The same raising and centralizing also applied to the diphthong in the MOUTH set (RP: /aʊ/, GenAm: /aʊ/, as in *now*, *brown*, *house*, *allow*) where the vowel was pronounced /əʊ/, akin to the contemporary phonological process known as “Canadian Raising.”

2.3.4. Consonants

Most of the consonants in Early Modern English had the same phonetic value as they do in Modern English; nevertheless, there are some notable minor differences to Modern English. In relation to present-day RP, and southern British speech in general, the chief distinguishing feature of Original Pronunciation is its rhoticity. The exact phonetic quality of /r/ after vowels is unclear, however. From Ben Johnson we know that “[i]t is sounded firm in the beginning of the words, and more *liquid* in the middle and ends; as in *rarer*, *riper*” (cited in Crystal 2016: xlvi). Crystal interprets this to mean a continuant [ɹ], akin to the sound found in the West Country of England and in most accents of North America, although he cautions that it is impossible to determine whether the focus of the articulation had been post-alveolar or retroflex. Another differentiating feature of Early Modern English phonology was the aspirated /ʍ/ sound in words such as *where*, *whale*, and *what*, *which*, contrasting the unaspirated /w/ sound; a distinction that is still present in some modern regional accents of English. Additionally, several words pronounced with voiceless fricative /θ/ in present-day English were pronounced /t/ in Elizabethan English, as is evidenced by spellings such as *fift* and *sixt* for *fifth* and *sixth*, or by rhymes such as *nothing* / *a-doting*. Other distinguishing consonantal features of Elizabethan speech were: the use of /sɪ/ instead of present-day /ʃ/ in words such as *suspicious*, *pensioner*, and *musician*; the variant pronunciation /n/ for /ŋ/ in the *-ing* verb inflection and adjectival uses; the use of /t/ instead of present-day /tʃ/ in words such as *nature*, *lecture*, or *tempestuous*. Lastly, elisions of consonants, reductions of words, and, in particular, *h*-dropping at the beginning of words were very common in the Early Modern period.

THE INTERNATIONAL PHONETIC ALPHABET (revised to 2015)

CONSONANTS (PULMONIC)

© 2015 IPA

	Bilabial	Labiodental	Dental	Alveolar	Postalveolar	Retroflex	Palatal	Velar	Uvular	Pharyngeal	Glottal
Plosive	p b		t d			ʈ ɖ	c ɟ	k ɡ	q ɢ		ʔ
Nasal	m	ɱ	n			ɳ	ɲ	ŋ	ɴ		
Trill	ʙ		r						ʀ		
Tap or Flap		ⱱ	ɾ			ɽ					
Fricative	ɸ β	f v	θ ð	s z	ʃ ʒ	ʂ ʐ	ç ʝ	x ɣ	χ ʁ	ħ ʕ	h ɦ
Lateral fricative			ɬ ɮ								
Approximant		ʋ	ɹ			ɻ	j	ɰ			
Lateral approximant			l			ɭ	ʎ	ʟ			

Symbols to the right in a cell are voiced, to the left are voiceless. Shaded areas denote articulations judged impossible.

CONSONANTS (NON-PULMONIC)

Clicks	Voiced implosives	Ejectives
◌ ɸ Bilabial	ɓ Bilabial	ʼ Examples:
◌ ɗ Dental	ɗ Dental/alveolar	pʼ Bilabial
◌ ʄ (Post)alveolar	ɸ Palatal	tʼ Dental/alveolar
◌ ɥ Palatoalveolar	ɡ Velar	kʼ Velar
◌ ɮ Alveolar lateral	ɠ Uvular	sʼ Alveolar fricative

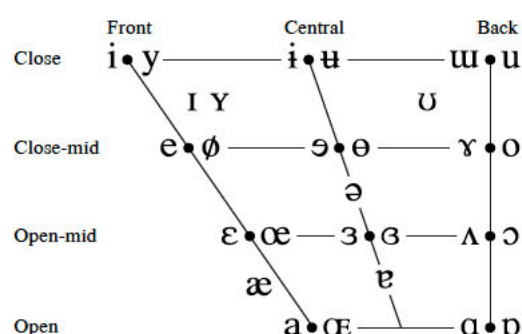
OTHER SYMBOLS

ɱ Voiceless labial-velar fricative	ɕ ʑ Alveolo-palatal fricatives
ʋ Voiced labial-velar approximant	ɭ Voiced alveolar lateral flap
ɰ Voiced labial-palatal approximant	ɥ Simultaneous ʃ and x
ħ Voiceless epiglottal fricative	Affricates and double articulations can be represented by two symbols joined by a tie bar if necessary.
ʕ Voiced epiglottal fricative	
ʡ Epiglottal plosive	

DIACRITICS Some diacritics may be placed above a symbol with a descender, e.g. ɲ̥̊

◌ ɹ Voiceless	◌ ɹ Breathy voiced	◌ ɹ Dental
◌ ɹ Voiced	◌ ɹ Creaky voiced	◌ ɹ Apical
◌ ɹ Aspirated	◌ ɹ Linguolabial	◌ ɹ Laminal
◌ ɹ More rounded	◌ ɹ Labialized	◌ ɹ Nasalized
◌ ɹ Less rounded	◌ ɹ Palatalized	◌ ɹ Nasal release
◌ ɹ Advanced	◌ ɹ Velarized	◌ ɹ Lateral release
◌ ɹ Retracted	◌ ɹ Pharyngealized	◌ ɹ No audible release
◌ ɹ Centralized	◌ ɹ Velarized or pharyngealized	
◌ ɹ Mid-centralized	◌ ɹ Raised	
◌ ɹ Syllabic	◌ ɹ Lowered	
◌ ɹ Non-syllabic	◌ ɹ Advanced Tongue Root	
◌ ɹ Rhoticity	◌ ɹ Retracted Tongue Root	

VOWELS



Where symbols appear in pairs, the one to the right represents a rounded vowel.

SUPRASEGMENTALS

ˈ Primary stress	ˈfounəˈtʃən
ˌ Secondary stress	
ː Long	eː
ˑ Half-long	eˑ
◌ ɹ Extra-short	ẽ
◌ ɹ Minor (foot) group	
◌ ɹ Major (intonation) group	
◌ ɹ Syllable break	ˌi.ækt
◌ ɹ Linking (absence of a break)	

TONES AND WORD ACCENTS

LEVEL	CONTOUR
ẽ or ˩ Extra high	ẽ or ˩ Rising
é ˨ High	ê ˨ Falling
ē ˨ Mid	ẽ ˨ High rising
è ˨ Low	ẽ ˨ Low rising
ẽ ˨ Extra low	ẽ ˨ Rising-falling
˩ Downstep	↗ Global rise
˩ Upstep	↘ Global fall

Figure 1. Full IPA chart

2.4. Implications of Original Pronunciation for Performances of Shakespeare

As Barrett (2013: 15) suggests, “[t]he use of original pronunciation (OP) in modern productions may be viewed as a transfer of culture from one historical period to another” and, in this regard, “it is analogous to the transfer of other cultural elements, such as costume, and practices, such as gesture and blocking.” This cultural transfer in historically informed practice does not, however, signify historical “authenticity,” which, as Lodewyck (2013: 44) suggests, is “neither possible nor particularly advisable.” Similarly, Crystal (2013: 16) considers this notion to be a fallacy:

[A]uthenticity is not the aim, just as it is not when describing the experience of sitting in the reconstructed Globe theatre. Rather, the operative word is ‘plausible’. [. . .] The excitement comes from exploring these possibilities to see what insights into meaning, aesthetic, or performance can be obtained.

Indeed, the value of Original Pronunciation “may be found [. . .] in the potential it holds to bring something new to the performance” (Barrett 2013: 17). In practical terms, as mentioned before, performing Shakespeare in reconstructed Elizabethan speech allows the restoration of the rhymes, puns, and wordplay that have been lost as the sound system of English gradually changed over time. Recovered rhymes are arguably the most immediately recognizable benefit for modern audiences experiencing Original Pronunciation. As noted before, rhyming couplets are frequently a marker of scene endings in Shakespeare’s works. Crystal (2016: xx) calculates that twelve percent of these pairings fail to rhyme in Modern English. It is thus particularly jarring when this happens at the ending of a play. The closing monologue of *Macbeth* (5.11.40–41), for instance,

So thanks to all at once, and to each one,
Whom we invite so see us crowned at Scone.

has posed a head-scratcher for “generations of actors” who “have tried and failed to make something of *one* rhyming with *Scone*—a rhyme that only works in OP” (Crystal 2016: xxii). Likewise, the final lines of *King Lear* ([Folio Text], 5.3.301–302), which cap one of the most stirring moments in the entire Shakespearean canon, are arguably rendered even more potent by the restoration of their rhyming in Original Pronunciation:

The oldest hath borne most; we that are young,
Shall never see so much, nor live so long.

Failed rhymes pose an even greater problem for Shakespeare’s poetry. In his sonnets, which are obviously built around a rigid scheme, thirteen percent of all lines fail to rhyme

in Modern English (Crystal 2016: xxii); in two instances, this is true for most of the line-pairs—four out of seven—as is exemplified by *Sonnet 154*:

The fairest votary took up that fire
Which many legions of true hearts had warmed,
And so the general of hot desire
Was sleeping by a virgin hand disarmed.
This brand she quenched in a cool well by,
Which from love's fire took heat perpetual,
Growing a bath and healthful remedy
For men diseased; but I, my mistress' thrall,
Came there for cure; and this by that I prove:
Love's fire heats water, water cools not love.

Restoring rhymes using reconstructed Elizabethan speech additionally offers “a bold demonstration of how the sound, feeling, and fluidity of words in Renaissance English influence the performative experience and the meaning made by the play” (Lodewyck 2013: 45). In the following passage of *A Midsummer Night's Dream* (3.2.102–109), for example, Original Pronunciation uncovers a series of repeated /aɪ/-endings:

Flower of this purple dye,
Hit with Cupid's archery,
Sink in apple of his eye.
When his love he doth espy,
Let her shine as gloriously
As the Venus of the sky.
When thou wak'st, if she be by,
Beg of her for remedy.

As Meier (2011: 211) observes, this for modern ears highly unusual sequence of sounds produces “a dreamy hypnotic effect, highly appropriate to what [Oberon] is doing at this point,” that is completely lost on modern audiences because of a “phonaesthetic clash” between the /aɪ/ and /i/ endings in present-day English.

Rhymes and assonances are also the main engine driving punning and wordplay, which, as Kökeritz (1954: 54) remarks, seems to have reached the zenith of its popularity during the Elizabethan period and “was not only a source of much merriment but also a greatly cherished intellectual exercise,” being “as much part of sophisticated conversation as it was a stock ingredient of contemporary comedy” and not disdained even by pulpit oratory. As already noted, Shakespeare had a penchant for puns, many of which, especially those of the subtler variety, are lost in Modern English. A striking example of

such a pun recovered through Original Pronunciation is found at the beginning of *Romeo and Juliet* (Prologue 5–6):

From forth the fatal loins of these two foes
A pair of star-crossed lovers take their life,
Whose misadventured piteous overthrows
Doth with their death bury their parents' strife.

Indeed, as Crystal (2016: xxvi) notes, once we are aware that the two modern diphthongs /ɔɪ/ and /aɪ/ were both pronounced /əɪ/ in Early Modern English, “then we are presented with the possibility that there is a genealogical nuance (*lines*) which can be added to the physical sense of *loins*.” But Shakespeare was also guilty of far bawdier wordplay that would undoubtedly have inflamed the famously raucous Elizabethan theater crowds. An especially notorious example is found in the following passage from *As You Like It* (2.7.23–28), in which Jaques relates to Duke Senior what he has gleaned from eavesdropping on Touchstone:

‘Thus we may say’, quoth he, ‘how the world wags:
'Tis but an hour ago since it was nine,
And after one hour more 'twill be eleven,
And so from hour to hour we ripe, and ripe,
And then from hour to hour we rot, and rot,
And thereby hangs a tale.’

These lines prompt Jaques to “laugh sans intermission / An hour by his dial” (*As You Like It* 2.7.32–33), but modern audiences are at a complete loss regarding what exactly provoked this intense hilarity. Only if contemporary spectators are aware of the homophony between *hour* and *whore*—pronounced both /o:ɪ/ in Early Modern English—are they able to understand Jaques’s merriment (Crystal 2016: xxvi).

Practitioners of Original Pronunciation relate other interesting secondary effects of performing Shakespeare using reconstructed Elizabethan speech, both for actors and for spectators. Tom Cornford (cited in Crystal 2005: 144), who served as assistant director for Shakespeare’s Globe production of *Romeo and Juliet*, for instance noticed that Original Pronunciation resonated differently in the actors’ physicality compared with the concurrent Received Pronunciation version of the performance:

I was fascinated by the effect on the actors’ bodies. Capulet’s second line is a good example, where Montague ‘flourishes his blade in spite of me’. In OP, blade sits lower and wider in the body than the RP version, and in sounding dangerous (the

RP equivalent sounds very correct and polite) it makes the actor look and feel dangerous.

The faster and more casual delivery in Original Pronunciation also results in a shorter playtime, so that the 2004 *Romeo and Juliet* production, for instance, ran ten minutes shorter than the RP version (Crystal 2005: 65), thus coming closer to the “two hours traffic of our stage” mentioned in its prologue.

Aside from simply speeding up the performance, the less formal delivery in Original Pronunciation also has profound consequences for audiences regarding their expectations of Shakespearean language. As Barrett (2013: 73) rightly notes, “Shakespearean actors have always been expected to enunciate clearly, pronouncing every consonant.” Over the ages, Shakespeare has thus increasingly been associated with the cultural elites and the artifice of refined Received Pronunciation, particularly in Britain (Lodewyck 2013: 43–44). Shakespeare himself, however, seems to have held diametrically opposite views on how his words should be pronounced, as evidenced by Hamlet’s intimation to his players that they speak “trippingly on the tongue” and not “mouth” the speech—that is, enunciate the words exaggeratedly (*Hamlet* 3.2.2). As already noted, audiences’ responses to the unfamiliar sounds of Original Pronunciation, as well as their expectations of what constitutes “proper” pronunciation of Shakespearean language, have so far never been empirically investigated. Before moving on to the study conducted for this thesis, in the following chapter I will first briefly elucidate the theory of language attitudes and outline the research instruments with which these can be measured.

3. Language Attitudes

3.1. Defining Language Attitudes

3.1.1. *Attitudes*

The notion of “attitude” is one of the most distinctive and crucial concepts in social psychology (Perloff 1993: 26) and as Garrett, Coupland, and Williams (2003: 2) observe, a fundamental concept in sociolinguistics since at least Labov’s (1966) seminal study on the social stratification of speech communities in New York City. Nevertheless, theorizing attitudes is neither easy nor straightforward, because they are “a latent psychological constant” (Agheyisi & Fishman 1970: 138) and “are never directly observed” (Allport 1935: 839). Consequently, various definitions have been attempted. Allport (1954: 19), for instance, defines attitudes as “a learned disposition to think, feel and behave toward a person (or object) in a particular way.” On this point, there seems to be universal agreement among scholars, as well as that “they are not momentary but relatively ‘enduring’” (Agheyisi & Fishman 1970: 139). Others, such as Sarnoff (1970: 279) suggest that attitudes are “a disposition to react favourably or unfavourably to a class of objects.” More recently, this has been expanded to include the notion of “construct,” in line with structuralist tendencies in social sciences. Oppenheim (1982: 39), for instance, describes attitude as

an abstraction which cannot be directly apprehended. It is an inner component of mental life which expresses itself, directly or indirectly, through such more obvious processes as stereotypes, beliefs, verbal statements or reactions, ideas and opinions, selective recall, anger or satisfaction or some other emotion and in various other aspects of behaviour.

A further approach defines attitude as “an evaluative reaction – a judgment regarding one’s liking or disliking of a person, event, or other aspect of the environment” that, as such, is “a non-neutral position” and “can range in its intensity” (Weber 1992: 117). For this thesis, then, I will adopt the following working definition proposed by Garret et. al. (2003: 3), who “take it as axiomatic” that

an attitude is an evaluative orientation to a social object of some sort, but that, being a ‘disposition’, an attitude is at least potentially an evaluative stance that is sufficiently stable to allow it to be identified and in some sense measured.

A consensus among scholars in the field stipulates that attitudes have a “tripartite” structure, consisting of a cognitive, an affective, and a conative (behavioral) component (Edwards 1982). There is considerable disagreement, however, concerning the function of these three components in relation to attitudes, with recent research cautioning against equating them with attitudes themselves. According to this view, cognition, affect, and behavior should rather be regarded as causes and triggers of attitudes (Garrett 2010: 23). Considerable doubts have also been expressed about the interplay of these components, arguing that attitudes are complex phenomena that exhibit many facets and manifestations (Garret et. al. 2003: 7). A particular thorny issue, then, is the relationship between attitudes and behavior. A strong commonsensical view exists that if one is able to change people’s attitudes toward something, their behavior is also subject to modification as a consequence; and, vice versa, that it is straightforward to infer people’s attitude from the way in which they behave—an assumption on which much of advertising and marketing is based (Garret et. al. 2003: 7). Nevertheless, there are strong indications that attitude and behavior are in fact not so neatly aligned. Ajzen and Fishbein’s (1980) “theory of reasoned action,” for example, emphasizes the social context within which any individual person operates, thereby diminishing the relative importance of private attitudes.

3.1.2. Language attitudes: social categorization and stereotyping

According to Dragojevic (2017)—on whom this section is based on, when not otherwise stated—attitudes toward languages can thus be defined as evaluative reactions to different language varieties. The latter can, in turn, can be described as “set[s] of linguistic items with similar social distribution” (Hudson 1996: 22). Empirical research of language attitudes began in the first half of the twentieth century and has thus far primarily focused on the cognitive component. Among the pioneering studies in this field is Pear’s (1931) survey in which he asked British radio listeners to compile personality profiles of different voices, discovering that various forms of accents and dialects caused fundamental changes in the respondents’ perception of the speakers. In the subsequent decades, much research has been conducted to determine whether voice parameters reflect people’s actual persona—that is, their dispositional states. The resulting findings concluded that there was only very modest overlap between respondents’ ratings of speakers’ vocal features and peer-ratings of the same persons’ personalities, as Giles and Billings (2005: 188) also

note. Nevertheless, there was remarkable consistency in respondents' wrongful inferences about speakers' vocation, extraversion, and other traits, suggesting that these judgements were based on vocal stereotypes.

Several theoretical models have been proposed in order to explain the social meanings associated with language varieties; nonetheless, there seems to be an agreement that language attitudes are, at least partly, a product of two sequential cognitive processes known as social categorization and stereotyping. According to this concept, listeners first infer which social group a speaker belongs to, based on language cues such as her or his accent, then they proceed to stereotype the speaker—that is, attribute to her or him the stereotypic characteristics associated with her or his inferred group membership. In the same vein as other social stereotypes, language-based stereotypes, too, are organized along evaluative dimensions (Giles and Billings 2005:190). According to McKenzie (2010: 47), a number of researchers (among which Giles & Coupland 1991; Edwards 1994; Dalton-Puffer, Kaltenböck & Smit 1997; Garrett et al. 2003, and Lindemann 2003) have clearly demonstrated that these dimensions can be condensed into two particularly salient evaluational categories that account for most of the attitude variance: *social status* (or competence) and *social attractiveness* (or solidarity). *Social status* comprises traits such as a speaker's education, intelligence, and success, whereas *social attractiveness* encompasses traits such as a speaker's friendliness, pleasantness, and honesty. Different language varieties are associated with different stereotypes along these dimensions, though the dimensions themselves might be labelled differently by different researchers.

Past language attitudes studies have primarily focused on investigating the stereotypes associated with standard and nonstandard language varieties, evincing that they elicit significantly different evaluations: speakers of standard varieties typically rate higher on the *social status* dimension than speakers of nonstandard varieties. Research has also unearthed a strong correlation between the degree to which a person's speech deviates from the standard variety and the level of status she or he is attributed. According to some scholars, this is reflective of a so-called "standard language ideology" prevalent in many societies, such as those where the widely diffused languages English, French, and Spanish are spoken. These "powerful ideological positions" are usually not conscious to the speakers, who "believe their attitudes to language to be common sense and assume

that virtually everyone agrees with them” (Millroy 2016: 133). Nonstandard varieties, conversely, tend to be evaluated higher on the *social attractiveness* dimension, particularly by members of the same linguistic community as a result of so-called “in-group loyalty.” Therefore, nonstandard varieties can indeed possess what is referred to as “covert prestige.” The propensity to attribute standard language varieties more *social status*—and, vice versa, assign nonstandard varieties more *social attractiveness*—has been shown to occur uniformly across different strata of a particular society, and, most notably, has been demonstrated cross-culturally and worldwide. There has consequently been an explosion of language attitude research in different parts of the world, beginning in the 1960s, which has demonstrated that people express definite and consistent attitudes toward speakers who use a particular language variety (Garrett et al. 2003: 12). This type of research can be accomplished by employing different instruments.

3.2. The Measurement of Language Attitudes

3.2.1. Main approaches

There are, broadly speaking, three main approaches to researching language attitudes: analysis of the societal treatment of language varieties, direct measures, and indirect measures (Ryan, Giles & Hewstone 2005). The societal treatment approach generally involves content analysis of the “treatment” given to languages, language varieties, and to their speakers within societies, hence it is also referred to as “content analysis” (Knops and van Hout 1988: 6). Studies that are grouped under this label are considered to be unobtrusive, because researchers typically infer peoples’ attitudes by conducting participant observation, by carrying out observational and ethnographic studies, or by performing analyses of a variety of sources in the public domain. Consequently, many of these studies are largely qualitative in nature. As Garrett (2010: 51) points out, the societal treatment approach is often overlooked in discussions of language attitudes research, though not for dearth of such work. The leading concern among many scholars is that much of the research in this mold is unsuited for statistical analysis and does not permit generalizations. Hence, some hold the view that these types of studies may best serve as preliminary investigation to more rigorous scholarly inquiry (Garrett et al. 2003: 16).

The direct approach, on the other hand, is characterized by a high degree of obtrusiveness: with questionnaires and interviews, respondents are asked direct questions about language evaluation, personal preferences, and other similar queries. In the direct approach it is thus the respondents themselves who report their attitudes, rather than researchers drawing inferences from their observations, as Knops and van Hout (1988: 7) remark. (One of the earliest and most famous studies conducted in this vein is Labov's (1966) previously mentioned investigation in which he asked New York City respondents to choose between alternative pronunciations of words with postvocalic /r/ and to express their views on which variant they deemed to be the "correct" one.) Data collection within the direct approach typically occurs either with so-called "word-of-mouth" techniques (interviews, surveys, and polls), or by means of written response methods such as questionnaires.

Garrett et al. (2003: 27–31) point to a number of potential pitfalls affecting direct approach studies. These are: asking hypothetical questions (how people *would* react to a particular situation), which are typically poor predictors of people's actual reaction or behavior; asking strongly slanted questions (containing "loaded" items, such as "Nazi," "healthy," or "natural") that tend to goad people into answering a particular way; asking multiple questions (where a positive answer may refer to more than one component), or double negative questions (to which answering negatively would be ambiguous); an inherent social-desirability bias (stemming from the tendency for people to provide socially appropriate responses to questions); an inherent acquiescence bias (which prompts respondents to agree with a question, regardless of its content); characteristics of the researcher (whereby responses may be affected by the gender and ethnicity of both interviewer and interviewee); effects of prior discussion (whether a discussion of the questions is allowed before respondents complete the questionnaires). A further issue seen as underlying this approach is "whether subjects' verbal statements of their attitudes and their behavioural reactions in concrete situations can indeed both be interpreted as manifestations of the same underlying dispositions" (Knops & van Hout 1988: 7). To counter the latter, in particular, so-called indirect measurement techniques have been developed.

The indirect approach to researching language attitudes entails more subtle techniques that have frequently been termed “deceptive.” As Garrett et al. (2003: 51) note, this approach is more or less synonymous with the so-called “matched guise technique,” developed by Lambert et al. (1960) in their seminal investigation of inter-ethnic language attitudes in Montreal. Lambert sought to find out how French Canadian and English Canadian speakers perceive each other, but he distrusted people’s overt responses that are elicited through direct approaches as a true reflection of their privately held views. As a result, Lambert developed the matched guise technique in order to detect covert attitudes toward speakers of different language varieties.

3.2.2. Matched and verbal guise studies: the “speaker evaluation paradigm”

Studies using the matched guise technique typically ask respondents to listen to a number of recorded voices each representing a different language variety—referred to as audio “stimuli”—and to rate each speaker on various character traits. As Giles and Billings (2005: 189) note (and on whom this paragraph is based), “[t]he procedure is built on the assumption that speech style triggers certain social categorizations that will lead to a set of group-related trait-inferences.” (In his original application of the technique, Lambert for instance supposed that hearing a voice that is classified as “French Canadian” would predispose listeners, based on their own group-membership, to infer that the speaker has a particular set of personality attributes.) The voices on the audio recordings are usually produced using so-called “balanced” bilinguals—that is, persons with equal facility in two, or more, languages or language varieties. The speakers then read the same passage of an “ethically neutral” text in different languages (for instance, French and English in Lambert’s original study) or language varieties, which form so-called “guises.” To avoid these guises being identified as produced by the same speaker, each version is interspersed with recordings of so-called “filler” voices of other speakers reading the same text. The crucial aspect in employing the matched guise technique is framing the test administration in such a way that respondents believe they are listening to a series of *different* speakers—from whence stems an inevitable element of deception, as has already been mentioned. As McKenzie (2010: 48) notes, considerable care and effort are thus expended on issues of stimulus control,

ensuring that prosodic and paralinguistic features of voice such as pitch, speech rate, voice quality and hesitations remain constant. Attention is also paid to minimising differences in features of reading style and expressiveness and ensuring that the recordings are perceived by the listener-judges as authentic.

Using the same speaker to render the different guises confers the matched guise technique “the benefit of controlling for other extraneous vocal characteristics that may naturally vary between speakers and influence listeners’ responses” (Dragojevic 2017).

Consequently, it is argued that response differences are attributable primarily to social expectations toward the guises being contrasted. These social expectations are, in turn, based on language cues (Giles and Billings 2005: 189). In Lambert’s (1960) Montreal study, for example, the results showed that Anglo-Canadian respondents rated the English guises more positively than the French guises on seven out of fourteen traits. Remarkably, however, Franco-Canadians also preferred the English guises and evaluated them even more favorably than Anglo-Canadian respondents on ten out of fourteen traits.

The matched guise technique quickly began to dominate language attitudes research—so much so that it came to be known as the “speaker evaluation paradigm” over the subsequent decades (Giles & Billings 2005: 189). As Garret et al. note (2003: 54), it is a “rigorous and elegant” instrument for investigating people’s privately held attitudes—as opposed to the direct questioning of respondents, which is far less likely to elicit covert attitudes and more liable to yield expressions of attitudes that are considered socially acceptable or even socially desirable by respondents. Another key advantage of this technique is the ability to collect data that lends itself optimally to statistical analysis (McKenzie 2010: 47). It has also led to a detailed and convincing demonstration of the role that language code and style choice play in impression formation. Furthermore, it has generated a considerable number of studies internationally—particularly in bilingual, bi-ethnic, multilingual, and multi-ethnic contexts—with a reasonable level of comparability, which has made a cumulative development of theory possible. Finally, it has laid the foundations for a vital cross-disciplinary field at the intersection of social psychology, language, and sociolinguistics (Giles & Billings 2005: 190).

Nevertheless, it has also come under criticism for some of its perceived shortcomings. Edwards (1982) for instance, claims that the matched guise technique does not so much measure language attitudes as attitudes toward representative speakers of

languages and language varieties. Another criticism that is frequently levelled against studies performed with the matched guise technique is that the use of decontextualized speech samples undermines their validity. The speech stimuli are often considered to be contrived and artificial, because “they do not take into account the social meaning of people’s ability to modify their speech styles in different contexts and at different linguistic levels” (Garret et al. 2003: 53). Other issues identified by Garret et al. (2003: 58–61) are, in their terms, the “saliency problem,” the “perception problem,” the “neutrality problem,” as well as several issues concerning authenticity, such as mimicking accent and style. Ensuring saliency is especially a problem when using the matched guise technique, because the routine of a repeated message content—particularly if there is a great number of speakers—may lead to respondents placing too much emphasis on vocal variations. Variations in speech and language become thus more salient than they would otherwise be, outside the experimental environment. Likewise, researchers have no means of ascertaining just how reliably respondents perceive the manipulated variables. For instance, a nonstandard accent variety might be judged by some test participants as “bad” or “incorrect” language (Garret et al. 2003: 58). This is especially an issue in studies that involve dialect areas, for which it is important to assure that respondents identify the provenance of the speakers. Furthermore, stimulus control—the biggest strength of the matched guise technique—also tends to level other speech characteristics, such as intonation or discourse patterning, that usually co-vary with accent varieties, leading to issues of authenticity regarding the use of these varieties. (Reading stimulus material aloud, especially, leads to a marked verbal style with distinctive prosodic and sequential phonological features). Finally, producing a “factually neutral” text is considered highly doubtful (Garret et al. 2003: 60). For instance, Giles et al. (1990) found it impossible to generate a text that was “age neutral” in a cross-generational study, because respondents were found to interpret the same passage of text differently according to the speaker’s perceived age.

As an alternative to these perceived shortcomings, efforts have been made to modify the matched guise technique. Some researchers have used different speakers to produce the stimuli, which has come to be known as “verbal guise technique.” This design is also sometimes employed out of necessity when matched guise recordings are not

feasible; for instance, if a single balanced bilingual or bidialectal speaker is not available or unable to produce the different varieties required for the guises. The use of different speakers is claimed to counter the accusation of artificiality, though as Garret et al. (2003: 54) remark, the notion of “authenticity” remains problematic in these cases. Some verbal guise studies, such as McKenzie (2008), have thus even resorted to use “natural speech,” by asking different speakers to speak “spontaneously” on the recordings. But regardless of whether a study employs matched or verbal guise recordings, both techniques predominantly make use of questionnaires with rating scales asking respondents to evaluate speakers according to specific traits or attributes, as introduced by Lambert et al. (1960). In the next section, I will outline these proceedings in more detail.

3.2.3. Rating scales

There are three types of rating scales employed in language attitude research: Thurstone-type scales, Likert-type scales, and semantic differential scales. Thurstone-type scales, though considered to be fairly reliable, are scarcely used nowadays because they are fairly laborious to devise, requiring a pool of relevant attitude statements derived from scholarly literature as well as from pilot studies (Garret et al. 2003: 39). Likert-type scales, on the other hand, are far easier to prepare and are considered to be more reliable (Oppenheim 1992: 200). Questionnaires with Likert-type scales generally contain a certain number of statements about the attitudes that the researcher wishes to measure, phrased for instance as “Language X is . . .”. Respondents are then asked whether—and to what extent—they agree or disagree with these statements. Because this technique captures overt attitudes toward languages, Likert-type scales are more often used in studies following the direct approach. Responses are scored on a numerical scale, very often ranging from 1 to 5—as proposed by Likert (1931) himself—whereby 1 indicates the least favorable response (usually labelled as “strongly disagree”) and 5 indicates the most favorable response (usually labelled as “strongly agree”). Other forms of Likert-type scales include expanded odd-numbered scales and even-numbered scales, usually ranging from 1 to 6.

An odd-numbered scale allows for a “neutral” mid-point, which is why most language-attitude researchers prefer five-point or seven-point scales (Garret et al. 2003: 51). Nevertheless, as Oppenheim (1992: 200) points out, scores in the middle of a

scale are often ambiguous: although they might indicate a much-considered weighing of attitudes, they may also reflect uncertainty and that respondents have little involvement in the issue. Hence, some researchers prefer to use even-numbered scales, sacrificing the option for respondents to express a “considered commitment” to a mid-point. Some scholars (for example, Henerson et al. 1987: 86), claim that Likert-type scales are especially useful for measuring the intensity of attitudes, which is to say that ratings on either end of the scale are indicative of a higher degree of attitude intensity than if a respondent selects the mid-point of the scale. This same claim is also made for so-called semantic differential scales, which originate from the work of Osgood (Osgood, Suci & Tannenbaum 1957).

Although Likert-type scales are sometimes used to assess the cognitive components of language attitudes in a more direct fashion (Zahn and Hopper 1985: 113), semantic differential scales “are the type of attitude-rating scales most typically associated with the matched guise technique” (Garret 2010: 55). Contrary to Likert-type scales, which ask respondents to agree to statements, semantic differential scales measure the connotative meaning of attributes in response to a prompt typically formulated as “Speaker X sounds. . .”. The design of this type of scales consists of opposing semantic labels—usually antonyms such as *friendly* / *unfriendly*, or *pleasant* / *unpleasant*—representing the attributes on which respondents are asked to evaluate the speakers. The labels are placed on either end of a gradient spaced by equidistant numbers, or by points frequently labeled as “somewhat,” “rather,” or “very,” which offer respondents a verbal indication of the extent to which they agree or disagree with the connotative description of the speaker. A key advantage of semantic differential scales over Likert-type scales is that they lend themselves to more rapid completion by eliciting snap judgements from participants and minimizing opportunities for their mental processing, thus compelling respondents to rely more on their “gut feeling” toward a particular speaker. Because they only require respondents to describe speakers with the aid of connotative meanings, rather than forcing them to agree or disagree with statements about speakers, it is generally thought that semantic differential scales greatly reduce the risk of social desirability and acquiescence bias (Garret 2010: 56).

Empirical research on attitudes toward languages has thus far, and for obvious reasons, only been conducted on “living” languages and language varieties. To this date, there has not been an attempt to investigate perceptions of and attitudes toward a “dead” language variety or an anachronistic mode of speech used solely for artistic purposes, such as Original Shakespearean Pronunciation. Therefore, before moving on to outline the research design of the present study in greater detail, I will briefly touch on attitudes toward Shakespeare in general, and Shakespearean Language in particular.

3.3. Attitudes toward Shakespeare and Shakespearean Language

As Lanier (2015: 440) notes (and on whom this paragraph is based), “Shakespeare” frequently serves as a metonym for charged concepts such as literature, classical theatre, intellectualism, and highbrow culture in general. Consequently, he is often treated as a veritable “repository of ‘universal’ truths,” invoked to elevate the cultural register of a work. Popular perceptions of the Bard, particularly in Britain, are often tightly linked with those cultural institutions that use Shakespeare to purvey a notion of “proper” culture and who reproduce and regulate his high-cultural status, thereby protecting the class privilege of their own members. It thus comes as no surprise that the “class-coded canons of stylistic decorum” underlying “authentic” Shakespeare are frequently accompanied by a whiff of what Lanier calls the “taint of elitism and antiquarianism.” These canons of taste are closely linked for modern audiences with the “quaint, quasi-scriptural ring of Shakespearian language.”

It is no surprise, then, that Shakespeare has come to be almost inextricably associated with Received Pronunciation, the speech variety of the British political and cultural elite. Received Pronunciation, as Pensalfini (2009: 146) observes, “was the accent of choice for stage performances – particularly for portraying upper and middle-class characters – throughout much of the twentieth century.” There have been efforts in the United Kingdom to promote regional English accents and varieties in the theater, spearheaded in the 1950s by so-called “kitchen sink dramas” such as John Osborne’s *Look Back in Anger* (1956). As Barrett (2013: 119) notes, “[t]his era was the start of a long slow process of evolution, where regional accents slowly gained acceptance and the social stigma associated with their use gradually diminished”—a process that is far from

completed, because “some theatre-goers still frown at Shakespeare in dialect.” Recent productions in the United Kingdom, for instance by Shakespeare’s Globe and the National Theatre, have featured performances by actors with a variety of British regional and non-British accents. Nevertheless, most major roles in contemporary enactments of Shakespeare are still overwhelmingly performed in Received Pronunciation, while use of regional varieties is relegated to minor and supporting roles. There are, however, strong indications that, at long last, the appeal of the more archaic variety of Received Pronunciation—nowadays referred to as “Conservative RP”—is slowly but surely declining in favor of its more modern incarnation “Contemporary RP.” This shift in tides was never more evident than in the recent, highly publicized National Theatre production of *Hamlet* at the Barbican Centre in London, which I personally attended in 2015:

Whereas older thespians—the likes of Ciarán Hinds in the role of Claudius, Anastasia Hille as Gertrude and, most prominently, Karl Johnson as the ghost of Hamlet’s father—all indulged in Conservative RP, the younger actors in the play, such as Benedict Cumberbatch as Hamlet and Siân Brooke as Ophelia, performed in a decidedly modern Received Pronunciation—which for Siân Brooke could even be said to display tendencies toward so-called “Estuary English.” Regional British and non-British accents, on the other hand, were relegated to supporting roles, such as that of the gravedigger.

No empirical study has ever been conducted on whether contemporary audiences indeed regard Received Pronunciation as the “proper” language variety for performances of Shakespeare, nor have attitudes toward other varieties used in the context of Shakespeare performances ever been investigated. In his account of Shakespeare’s Globe’s *Romeo & Juliet* production, Crystal (2005) relates his own informal impressions of audience reactions to Original Pronunciation. Although they are anecdotal in nature—and, as Lodewyck (2013: 55) remarks, very likely biased by his investment in the project—they nevertheless provide an interesting starting point for the empirical investigation of people’s attitudes toward performances of Shakespeare in reconstructed Elizabethan pronunciation. From various talkback sessions, Crystal (2005: 142) gained the impression that audiences had “got into” the language by the end of the first couple of scenes and that, in their perception, the characters were more “down to earth” when speaking Original Pronunciation. Further reactions by the members of the audience

characterized the speech as being “rustic,” which Crystal understands to mean “honest,” “open,” and “direct.” He thus sums up his impression of the spectators’ auditory experience attending *Romeo and Juliet*:

[Original Pronunciation] reduced the psychological distance between speaker and listener, and to that extent presented a more immediate opportunity to access the speaker’s thought. [. . .] Far from pulling the mind away from the moment, OP seemed to help focus it.

Younger respondents in particular, who in Crystal’s (2005: 137) view are “the litmus test” for audience engagement, reported that they felt “closer to the characters” when they spoke in reconstructed Elizabethan speech. Lodewyck (2013: 55), for her part, interprets these comments as Original Pronunciation “speak[ing] truth to the teens because it is relatable, while the ‘posh’ nature of Received Pronunciation is a kind of artifice.” Echoing Lanier (2015), Lodewyck (2013: 55) further writes that “the connotations of ‘posh’” associated with Received Pronunciation are “steeped in ideas about the cultural capital of Shakespeare and theatre itself,” while Original Pronunciations seemingly grants younger and less privileged spectators access to this high status commodity. Bearing these anecdotal reports and considerations in mind, I proceeded to design a matched guise test with which to collect quantitative data in order to empirically investigate people’s covert attitudes toward Original Pronunciation performances of Shakespeare. The following chapter outlines the aims and hypotheses that were formulated for this study and presents the research design thereof in greater detail.

4. The Research Design of the Study

4.1. Aims and Hypotheses

Because no quantitative inquiry in this mold has thus far been attempted, the work for this thesis constitutes a somewhat unique precedent. Hence, one of the intentions underlying this study is to demonstrate that the tried-and-true empirical research instruments of sociolinguistics and social psychology can be effectively put to practice on topics crossing over into the field of aesthetics and the arts. Nevertheless, the principal aim of the study is to investigate people's attitudes toward Original Pronunciation—the plausibly reconstructed Early Modern English spoken during Shakespeare's own lifetime—and to establish whether it is deemed an acceptable speech variety for contemporary performances of Shakespeare, as compared with the three most dominant and widespread contemporary standard varieties of English: General American, Contemporary RP, and Conservative RP. To that end, the following research questions have been formulated:

- (i) What evaluative reactions do people have toward performances of Shakespeare in the different modern standard varieties of English and in Original Pronunciation?
- (ii) Do people have more difficulties understanding Original Pronunciation than the modern standard varieties of English in performances of Shakespeare?
- (iii) Do people deem Original Pronunciation less suitable for performances of Shakespeare than the modern standard varieties of English?
- (iv) How do people rank Original Pronunciation in their preference for performances of Shakespeare compared to the modern standard varieties of English?
- (v) Which modern English variety do people associate with Original Pronunciation?
- (vi) Are there significant differences in the perception of Original Pronunciation between native speakers and non-native speakers of English and do further background variables influence these perceptions?

Because of the near absence of previous research on the subject, it was very difficult to a priori predict the direction of the responses to the research questions outlined above. Nonetheless, based on the results of previous language attitudes studies, which have

evidenced that speakers of standard varieties typically rate higher on *social status* than speakers of nonstandard varieties, it was safe to assume that Original Pronunciation would follow this pattern and rate low along the same evaluative dimension. Far more difficult, in that regard, was to foresee evaluations along the *social attractiveness* dimension, because, not being familiar with Original Pronunciation, respondents presumably would not experience the same “in-group loyalty” that previous studies have shown is regularly expressed toward nonstandard language varieties. Altogether, taken Crystal’s (2005, 2013) informal reports on audience reactions to Original Pronunciation performances into account, as well as Barrett’s (2013) and Gibbons’s (2011) remarks on the dominant role of Received Pronunciation regarding theater performances of Shakespeare, the following predictions were made:

- (1) Original Pronunciation will rate low on *social status*, whereas the modern British and American standard varieties of English will rate high on the same evaluative dimension.
- (2) Original Pronunciation will rate high on *social attractiveness*, whereas the modern British and American standard varieties of English will rate low on the same evaluative dimension.
- (3) Original Pronunciation will not prove more difficult to understand than the modern British and American standard varieties of English.
- (4) Original Pronunciation will not be deemed less suitable for performances of Shakespeare than the modern British and American standard varieties of English.
- (5) Original Pronunciation will not rank lower in preference for performances of Shakespeare than the modern British and American standard varieties of English.
- (6) Irish English will be the modern-day speech variety most closely associated with Original Pronunciation.

For research question (vi) it was not considered appropriate to formulate a hypothesis, because of the lack of research on attitudes of non-native English speakers toward non-standard varieties of English. As McKenzie (2010) has shown with his seminal investigation of Japanese speakers’ attitudes toward different native and non-native

speech varieties of English, the field of inquiry of English(es) in the global context is an area that undoubtedly merits further exploration.

4.2. The Research Instrument

Because of its unobtrusive nature and ability to detect covert attitudes toward languages, an indirect approach by means of the matched guise technique was deemed the most appropriate method for investigating the research questions formulated in this study. The most crucial decision in the preparation of the questionnaire concerned the text passage to be used as audio stimulus to evaluate the different varieties. Producing the required audio material for the questionnaire raised a significant problem, however. Because stimulus control is of the utmost importance within the speaker evaluation paradigm, the decision was made to select a sonnet, rather than a passage from a play, as audio material.

Individual interpretation of dramatic texts is arguably one of the most crucial aspects of theatrical performances, therefore it is virtually impossible to ensure that extraneous factors such as pitch, speech rate, voice quality and hesitations remain constant in renditions of the same passage by different actors. With recitations of sonnets by different people, on the other hand, it is much easier to minimize differences in reading style and expressiveness and therefore ensure sufficient comparability between the speakers.

Owing to the lack of available resources, it was not feasible to recruit a native speaker who could convincingly recite a sonnet by Shakespeare in Original Pronunciation. The choice of stimulus material was thus severely restricted by the availability of professional audio. At the time of this study, the only recording of Original Pronunciation available to the public was the CD *Shakespeare's Original Pronunciation: Speeches and Scenes Performed as Shakespeare Would Have Heard Them* (Crystal et al. 2012), published by British Library Recordings. It contains a variety of short extracts from Shakespeare's plays, as well as from a few select sonnets, interpreted in Original Pronunciation by various actors. One track, *Sonnet 116*, stood out because it had been recorded both in Original Pronunciation and in Contemporary RP by the same actor, Ben Crystal. It was therefore regarded as the ideal choice for stimulus material in this study. From a phonological perspective, a further consideration in favor of using *Sonnet 116* was that it showcases some of the most salient phonological features of Original

Pronunciation, such as the exact rhyming of *love* with *prove*, monophthong vowels in *shaken* and *taken*, or a long /e:/ sound in *cheeks* and *weeks*:

Let me not to the marriage of true minds
Admit impediments. Love is not love
Which alters when it alteration finds,
Or bends with the remover to remove.
O no, it is an ever-fixèd mark
That looks on tempests and is never shaken;
It is the star to every wand'ring barque,
Whose worth's unknown, although his height be taken.
Love's not time's fool, though rosy lips and cheeks
Within his bending sickle's compass come;
Love alters not with his brief hours and weeks,
But bears it out even to the edge of doom.
If this be error and upon me proved,
I never writ, nor no man ever loved.

Given the great popularity of this sonnet (which has gained some notoriety for its preponderance at wedding ceremonies) and the availability of a wide variety of different recordings in the public domain, it was not difficult to retrieve two interpretations—in General American and in Conservative RP—that would be as close as possible to Ben Crystal's renditions in pitch, speed, and intonation. In the following section, I will describe the phonological features of each audio stimulus in more detail.

4.3. The Audio Stimuli

4.3.1. Speaker A: General American

The recording of *Sonnet 116* in General American (Brewster-Geisz 2006) was obtained through the Librivox project online platform. Phonologically, this rendition is chiefly distinguished by the postalveolar approximant [ɹ] occurring postvocalic in *alters*, *alteration*, *or*, *remover*, *ever*, *mark*, *never*, *star*, *barque*, *worth*, *hours*, *bears*, *error*, and *never*. The /ɑ/ vowel in *not*, a characteristic feature of General American, is curiously rendered by the speaker somewhat closer to its Received Pronunciation counterpart /ɒ/, but without having quite the same frontal realization. This is perhaps because the speaker, consciously or unconsciously, wished to suppress the most marked feature of his American pronunciation—presumably because he deemed it more “appropriate” for a Shakespearean sonnet. A similar observation can be made about the instances where the

speaker does not render intervocalic /t/ as an alveolar “tap” or “flap,” such as in the word combinations *It is* and *it out*. Again, the speaker presumably must have intended to suppress the most evident signs of his American accent in “obsequiousness” to Shakespeare. The other prominent feature that characterizes the General American recording is the short und rounded [ɔ] in *alters* and *alteration*, compared to the longer and rounded [ɔ:] occurring in both Received Pronunciation renditions.

4.3.2. Speaker B: Contemporary RP

The reading on the Contemporary RP recording (Crystal et al. 2012) was also performed by Ben Crystal, as previously mentioned. Compared with the General American version, the most distinguishing phonological features of this rendition are the absence of postvocalic /r/, as well as the presence of an /ɒ/ vowel in *not* and a longer and more rounded rendition of /ɔ:/ in *alters* and *alteration*.

4.3.3. Speaker C: Conservative RP

The Conservative RP rendition was retrieved from the online video platform YouTube (Selwyn 2009) and is notable for several phonological differences compared to both the General American and Contemporary RP recordings. In the very first word of the sonnet, *Let*, the phonetic realization of the /e/ phoneme is more closed, [e], than either General American or Contemporary RP, where it is more open [ɛ]. A pronounced feature of this rendition is intervocalic /r/ rendered as alveolar flap [ɾ] in *marriage*, *alteration*, *error*, and as alveolar trill [r] in initial /r/ in *remover*, *remove*, *every*, *wand’ring*, *rosy*, *brief*. There is also notable aspiration of voiceless plosives [t^h] in *let*, *not*, and of [k^h] in *mark*, *shaken*, *bark*, *taken*. The /ʌ/ vowel in *love* is further open and back, toward the [ɑ] position. The /eɪ/ diphthong in *taken* and *shaken* is more closed, [eɪ], compared with the [ɛɪ] diphthong in the General American and Contemporary RP version. Furthermore, final /d/ in *loved* and *proved* is fully voiced, as opposed to the General American and Contemporary RP versions where it is devoiced and realized as [t̚].

4.3.4. Speaker D: Original Pronunciation

As noted before, the Original Pronunciation recording of *Sonnet 116* was taken from Crystal et al. (2012) and recited by Ben Crystal. As does the General American version,

the Original Pronunciation rendition, too, features postvocalic /ɹ/ realized as postalveolar approximant [ɹ̠]. An additional idiosyncratic feature of this version is the presence of a retroflex approximant [ɻ] in postvocalic preconsonantal position in the words *mark* and *barque*. Because it is the most phonologically distinct version among the four audio stimuli, a full phonemic transcription of the Original Pronunciation recording is provided below:

lɛt mɪ nɑt tə ðə mæɹjɑz ə tɹu məɪndz
 ɛdmit ɪmpədəmənts | hʌv ɪz nɑt hʌv
 ʌɪf altəɪz ʌn ɪt altəɪɜːsɪən fəɪndz
 əɪ bɛnds wɪ ðə ɹɪmʌvəɪ tə ɹɪmʌv
 o noː | ɪt ɪz ən ɛvəɪ fɪksɛd maɪk
 ðæt lʊks ʌn tɛmpəsts ɛnd ɪz nɛvəɪ ʃɛːkən
 ɪt ɪz ðə stɑɪ tu ɛv.ɪ wɑnd.ɪn baɪk
 hʌz wɛɪθs ɹɪnoːn ʌldoː hɪz hɛɪt bɪ tɛːkən
 hʌvz nɑt təɪmz fʊl ðo ɹoːzɪ lɪps n ʃɛːks
 wɪðɪn ɪz bɛndɪn sɪklz kɹɪmpəs kɹɪm
 hʌv altəɪz nɑt wɪð ɪz brɛf oːɪz n wɛːks
 bæt bɛːɪz ɪt əʊt ɛːn tə ðɪ ɛdʒ ə dɹɪm
 ɪf ðɪs bɪ ɛɪəɪ ən əpɑn mɪ pɹɪvɪd
 əɪ nɛvəɪ ɹɪt nɛɪ noː mæn ɛvəɪ hɹɪvɪd

The first line of the sonnet in Original Pronunciation already displays considerable differences with respect to the other three renditions. Though *Let* is pronounced with [ɛ] as in General American and Contemporary RP and *not* is pronounced with [ɒ] as in General American, *marriage* truly sounds novel, rendered as /'mæɹ.jɑz/ with secondary stress on the second syllable. In the same line, the preposition *of* is clipped to become /ə/ and the diphthong in *minds* is centered to /əɪ/. The second line, too, features very characteristic phonological features: *love* is pronounced with /ɜ/, *which* and *when* have aspirated /ʌ/, *alteration* is pronounced /altəɪɜːsɪən/, and there is again a centering diphthong /əɪ/ in *finds*. The fourth line features /ɜ/ phonemes in *remover* and *remove*. In

the sixth line, there is monophthongal /ɛ:/ in *shaken*, in place of the /eɪ/ diphthong found in the General American and Contemporary RP versions and the /eɪ/ diphthong of the Conservative RP rendition. Moreover, there is a long monophthong /o:/ in *although*, contrasting with the /əʊ/ vowel of Received Pronunciation and the /o/ vowel of General American. Line nine features again /ʌ/ in *love's* and a centering diphthong /əɪ/ in *time's*, in addition to monophthong /o:/ in *though* and *rosy* and a long /e:/ vowel in *cheeks*. Line ten features repeated /ʌ/ sounds in the words *compass come* as well as consonant elision of /ŋg/ in *bending*. In line eleven, the long /e:/ vowel appears again in *brief* and *weeks*, as does monophthong /o:/ in place of modern-day /aʊ/. In line twelve, there is another centering diphthong, /əʊ/ in *out* instead of present-day /aʊ/, and once again an instance of consonant elision, this time /v/ is dropped in the word *even*. Similarly, /v/ is also dropped in *of*, whereas the short /ʌ/ in *doom* presents a marked contrast to the long /u:/ found in the same word in the three modern renditions. In the penultimate line, *upon* is pronounced /əpən/, akin to modern General American, whereas the /ʌ/ sound in *proved* restores the rhyme with the analogously pronounced *love* in the final line. Finally, the last line features a centering diphthong /əɪ/ in *I* and a long monophthong /o:/ in *no* as distinguishing markers of this Original Pronunciation recording.

A few sounds that are characteristic of Shakespearean Phonology with regard to Crystal's (2016) reconstruction are notable for their absence; for instance, there is no final diphthong /əɪ/ in *every* and *rosy* and /h/ is not dropped at the beginning of *whose* and *height*. This is, however, entirely in agreement with the principles of the Original Pronunciation "movement" laid out, among others, by Crystal (2005) himself. Barrett (2013: 10), for instance, argues that Original Pronunciation "is best defined by a set of parameters, within which the accent will function," allowing for a variety of pronunciation styles, rather than "a definitive set of rigid pronunciations." In the following section, I will expound upon the questionnaire design in more detail, as well as on the considerations that informed this process.

4.4. The Questionnaire

4.4.1. *Piloting phase*

Once the decision was made in favor of the matched guise technique as the principal research method for this study, the choice of employing semantic differential scales for the evaluation of the speakers followed naturally. Its ability to measure the connotative meanings of attributes and its propensity for eliciting “gut reactions” from respondents without giving them too much time to think were the most compelling arguments in favor of this evaluative tool. Nevertheless, the choice of attributes and their corresponding bipolar adjective pairs proved to be a much more arduous process than expected, because of the peculiar subject of this study: a historical language variety no longer spoken by anyone. As Garrett, Coupland, and Williams (2003: 56) note, “in deciding which bipolar adjective scales to use on the questionnaire, researchers have generally made their own decisions about what to include, usually on the basis of those commonly used in earlier studies.” After carefully surveying and consulting previous language attitude studies, however, I concluded that the adjective scales used in previous research were, for the most part, not particularly well suited to the purpose of my own peculiar research topic. For similar reasons, “specific semantic-differential scales are sometimes specially constructed for studies, as adjectives that elicit reactions from particular speech communities are likely to be highly culture bound,” and “[l]anguage attitude researchers should, therefore, not suppose that the same set of traits will be salient for different populations” (McKenzie 2010: 50).

I therefore devised an extensive set of bipolar adjectives in response to the prompt “Speaker X sounds. . .”, with which to precisely capture respondents’ evaluative reactions to my audio stimuli. The number of adjectives in the initial draft was deliberately chosen high because I intended to trial these attributes with a pilot group and subsequently reduce them based on the group’s feedback. The pilot version of the semantic differential scales was thus set up as displayed below in Figure 2. As McKenzie (2010: 48) observes, in order to minimize potential test fatigue for the respondents, the ordering of bipolar adjective scales is usually reversed for roughly half the items on a questionnaire; that is to say, positively connotated and negatively connotated adjectives are scrambled to avoid any potential left-right bias, a consideration that is echoed by scholars such as Dörnyei

(2007: 106). In this initial draft, however, I preferred not to scramble the adjectives because the left-right ordering provided a clearer overview of the traits that I was seeking to define. I therefore proceeded to scramble the positive and negative poles of the attributes only once I had finalized the adjective pairs.

1) <i>SPEAKER A</i> sounds... (tick the appropriate boxes)							
←	VERY	RATHER	SOMEWHAT	SOMEWHAT	RATHER	VERY	→
humorous	☐	☐	☐	☐	☐	☐	humourless
likeable	☐	☐	☐	☐	☐	☐	unlikeable
beautiful	☐	☐	☐	☐	☐	☐	ugly
sophisticated	☐	☐	☐	☐	☐	☐	unsophisticated
cheerful	☐	☐	☐	☐	☐	☐	solemn
approachable	☐	☐	☐	☐	☐	☐	distant
not snobbish	☐	☐	☐	☐	☐	☐	snobbish
educated	☐	☐	☐	☐	☐	☐	uneducated
urban	☐	☐	☐	☐	☐	☐	rural
modern	☐	☐	☐	☐	☐	☐	old-fashioned
normal	☐	☐	☐	☐	☐	☐	strange
familiar	☐	☐	☐	☐	☐	☐	unfamiliar
friendly	☐	☐	☐	☐	☐	☐	unfriendly
serious	☐	☐	☐	☐	☐	☐	light-hearted

Figure 2. Pilot questionnaire: semantic differential scales

Another important issue in designing semantic bipolar scales, as mentioned previously, concerns the number of gradient points one wishes to include and whether the scale should be odd-numbered or even-numbered. After careful consideration, I opted for a six-point even-numbered scale because I second Oppenheim's (1992: 200) preoccupation that scores in the middle of an uneven scale are often ambiguous. My chief concern was that a neutral mid-point might have encouraged too much uncertainty in respondents and provoked too much "fence-sitting." Furthermore, given the more "trivial" nature of my inquiry (which, as noted before, dovetails with the fields of aesthetics and the arts)—in contrast with the decidedly more "serious" sociological nature of traditional language attitude studies—I held no reservations against deliberately forcing respondents to tilt toward one or the other pole of the spectrum. (It is also my personal belief that when it

comes to opinions on arts there can be no “fence-sitting.”) Nevertheless, I decided to counterbalance this potential “coercion” of my respondents by providing three gradients for each pole thus effectively splitting the neutral mid-point in two. Accordingly, I labelled the three points on each side of the scale as “somewhat,” “rather,” and “very,” to dispel any potential misapprehension on the part of the respondents. Garland (1990: 20), for instance, found that labels on a semantic differential scale are the best compromise for a researcher intending to eliminate any lingering doubt that respondents might not fully comprehend the adjective scales and their gradient points.

In addition to semantic differential scales in response to the prompt “Speaker X sounds. . .”, I also included two direct questions (Figure 3, below). The first question was intended to evaluate the intelligibility of the four speakers, and thus of the four different speech varieties. To facilitate the comprehension of the audio stimuli—particularly of the Original Pronunciation recording—I decided to include the text of the sonnet on the questionnaire. The second question was designed to elicit respondents’ opinions on how suitable the various speech varieties are for theatrical performances of Shakespeare.

2) How clear and understandable was *SPEAKER A*'s pronunciation?

←	VERY	RATHER	SOMEWHAT	SOMEWHAT	RATHER	VERY	→
clear / understandable	☐	☐	☐	☐	☐	☐	unclear / not understandable

3) How suitable do you think *SPEAKER A* would be for a role in a Shakespeare production?

←	VERY	RATHER	SOMEWHAT	SOMEWHAT	RATHER	VERY	→
not suitable	☐	☐	☐	☐	☐	☐	suitable

Figure 3. Pilot questionnaire: suitability for Shakespeare

As Garret et. al (2003: 66) rightly point out, “the nature of semantic-differential scales and the over-reliance on such scales and dimensions from previously published work may restrict the evaluative picture that emerges from language attitudes research.” According to them, there is “a strong case for supplementing them with more spontaneous and context-sensitive data,” crafting “a complex of methods and of response options that is

able to match the inherent complexity of language attitudes, as entertained by different individuals and groups” (Garret et. al 2003: 66). Heeding this advice, I incorporated an open-ended question with which to elicit respondents’ impressions regarding the provenance of the four different speakers (Figure 4, below), opting for the sentence completion format, which, as Dörnyei (2007: 107) notes, “can elicit a more meaningful answer than a simple question.” With this test item, I specifically intended to substantiate Crystal’s (2013: 7) anecdotal reports in which he describes Irish English as being the present-day English variety most commonly associated with Original Pronunciation.

4) Where do you think the *SPEAKER* comes from?

Speaker A comes from: _____

Figure 4. Pilot questionnaire: speaker provenance

In order to validate my research instrument, I then piloted the questionnaire with a peer group of eight graduate students in the English department of the University of Vienna, which also included three native speakers. As Garrett et al. (2003: 56) lament, not enough researchers “have made efforts to elicit adjectives from the respondents themselves (or what is deemed to be a comparable group),” a practice that strives to ensure that “scaled judgements are made on scales that are meaningful and salient to the respondents.” My pilot group’s feedback was thus instrumental in shaping and contributing to the final version of the questionnaire. Right away, several adjective pairs were eliminated to their perceived redundancy. For instance, the attributes *friendly* / *unfriendly* and *familiar* / *unfamiliar* were regarded as already semantically contained in *likeable*, while the pair *educated* / *uneducated* was considered to be encompassed by *sophisticated* / *unsophisticated*—the latter deemed a more preferable option by the peer group. Because of similar issues of redundancy, the attributes *cheerful* / *solemn* and *serious* / *light-hearted* were merged to *light-hearted* / *solemn*, whereas the adjective pair *approachable* / *distant* was altered to *relatable* / *unrelatable*. One last attribute was changed owing to the precious contribution of Professor Nikolaus Ritt, my supervisor for this thesis. Upon his suggestion, the bipolar couple *relatable* / *unrelatable* was refashioned as *authentic* /

inauthentic, because that formulation was regarded as much clearer and more intuitively understandable by respondents.

4.4.2. Final questionnaire

As can be seen in Figure 5 below, the semantic differential scales were thus reduced from fourteen to nine in the final version of the questionnaire (a full version of which can be found in the appendix).

1) *SPEAKER A* sounds... (tick the appropriate boxes)

←	VERY	RATHER	SOMEWHAT	SOMEWHAT	RATHER	VERY	→
pleasant	☐	☐	☐	☐	☐	☐	unpleasant
snobbish	☐	☐	☐	☐	☐	☐	not snobbish
rural	☐	☐	☐	☐	☐	☐	urban
likeable	☐	☐	☐	☐	☐	☐	unlikeable
old-fashioned	☐	☐	☐	☐	☐	☐	modern
authentic	☐	☐	☐	☐	☐	☐	inauthentic
light-hearted	☐	☐	☐	☐	☐	☐	solemn
boring	☐	☐	☐	☐	☐	☐	interesting
sophisticated	☐	☐	☐	☐	☐	☐	unsophisticated

Figure 5. Final questionnaire: revised semantic differential scales

3) On a scale from 1 (*not at all*) to 10 (*very well*), how well would you have understood *SPEAKER A* without reading the text?

→	1	2	3	4	5	6	7	8	9	10	→
not at all	☐	☐	☐	☐	☐	☐	☐	☐	☐	☐	very well

Figure 6. Final questionnaire: intelligibility of the speaker

In addition to the two direct questions that were already included in the pilot questionnaire, a further direct question was added as a result of the peer group’s feedback. The participants—most vocally the native speakers, which came as a surprise—expressed the need for an additional item that would measure the respondents’ subjective impressions of how well they would have understood the speaker without also reading the text along with listening to the recordings (Figure 6, above). The decision to include the

text of the sonnet on the questionnaire was thus not only validated but proven to be a vital part of the research instrument by the pilot group. One final item on the questionnaire asks to rank the four speakers according to participants' preference for them in relation to performances of Shakespeare.

In conclusion: how would you describe the relevance of Shakespeare's works for you personally?							
←	VERY	RATHER	SOMEWHAT	SOMEWHAT	RATHER	VERY	→
not relevant	☐	☐	☐	☐	☐	☐	relevant

Finally, please fill out some basic information about yourself.					
Age: _____					
Please indicate your level of proficiency in English:					
NATIVE SPEAKER ☐	MASTERY (C2)* ☐	ADVANCED (C1)* ☐	UPPER INTERMEDIATE (B2)* ☐		
* according to the Common European Framework of Reference for Languages					
Which variety of English do you speak (or use as a model, if you're a non-native speaker)?					
BRITISH ENGLISH ☐	AMERICAN ENGLISH ☐	SCOTTISH ENGLISH ☐	IRISH ENGLISH ☐	OTHER (please specify) ☐	NO MODEL (non-native speaker) ☐
Other variety: _____					
Are you pursuing or have completed an English degree? (please indicate the highest level reached)					
No		☐	Yes, I am pursuing a master's degree		☐
Yes, I am pursuing a bachelor's degree		☐	Yes, I have obtained a master's degree		☐
Yes, I have obtained a bachelor's degree		☐	Yes, I am pursuing or have obtained a PhD		☐

Figure 7. Final questionnaire: background variables

Lastly, a section covering personal information was included on the final questionnaire (Figure 7, above) with the specific intent of providing background variables for the data analysis. Altogether, it was decided to establish a primary background variable and three secondary background variables. The primary background variable sought to establish two main subsets of participants based on their language proficiency—effectively separating native speaker respondents from non-native speaker respondents in the process.

Accordingly, this background variable was instrumental in shaping the sampling procedures for this study, as I will explain in the following section. An additional important function of the primary variable was to determine the (self-ascribed) proficiency level of non-native speaker respondents, according to the Common European Framework of Reference (CEFR). (In order for non-native speaker respondents to assess their proficiency level as accurately as possible, I also included a link to a global description of the CEFR proficiency levels in the online version of the questionnaire.) Because proficiency in English lower than C1 level according to the CEFR was not deemed enough to obtain valid responses to the questionnaire from non-native speakers, I decided to discard data by participants who had indicated proficiency in English at B2 level. The first secondary background variable was intended to determine whether there exists a correlation between respondents' attitudes and the variety of English that they speak natively—or that they use as model variety or reference point, if they are non-native speaker respondents. An additional secondary background variable was provided by the question, "How would you describe the relevance of Shakespeare's work for you personally?", which was intended to establish a correlation between respondents' personal interest in Shakespeare's works and their attitudes toward different speech varieties used to perform these. A final secondary variable was introduced in order to seek out the correlation between respondents' attitudes and their academic background in the subject of English; that is, whether the respondents had obtained, or were in the process of pursuing, an English degree at tertiary level.

4.5. Questionnaire Administration and Data Collection

In order to obtain the most representative and wide-spread sample possible for the study, I decided to collect the data on the internet with the aid of the Google Forms platform. For that purpose, the questionnaire was disguised as a survey intended to evaluate four potential actors who had recently auditioned for an upcoming performance of Shakespeare (as can be viewed in the final version of the questionnaire in the appendix). The most significant advantage of an online platform with respect to a pen-and-paper version of a matched guise test, is the seamless integration of the audio stimuli into the questionnaire. Respondents were thus able to play the recordings of the different speakers at their leisure

and where not limited to the two listening passages that are customary with pen-and-paper questionnaire administration. One potential pitfall of online questionnaires is that respondents may skip answers, either willingly or inadvertently. Nevertheless, this can be prevented by the forced-response option, which many online survey platforms, such as Google Forms, offer. This option impedes respondents to progress or complete the survey without having filled out all the questions, flagging up incomplete items in the process.

As previously mentioned, the decision of establishing two subsets of the study population based on their language proficiency profoundly affected the process of data collection, leading to the adoption of so-called “quota sampling” (Dörnyei 2007: 98). Hereby, I determined that fifty percent of the participants be native-speakers and fifty percent be non-native speakers. In order to obtain enough responses to meet my quota requirements, an invitation to participate in the online questionnaire was sent out over various social media platforms as well as via e-mail, thereby combining elements of random sampling and snowball sampling (Dörnyei 2007: 97–98). Participation in the study was extended, in particular, to students and members of staff at the English department of the University of Vienna, as well as to English departments at universities in the United Kingdom. Even though the invitations were sent out at possibly the least propitious moment—coinciding with the end of the academic year and the summer vacations—a total of hundred answers was obtained. Because of my sampling requirements, this meant that fifty responses each were gained from native speakers and non-native speakers. In the next chapter I will undertake to analyze and discuss the results of the data thus collected.

5. Results and Discussion

5.1. Description of Participants

5.1.1. Entire sample

The mean age of participants in this study was 36.8 years. As can be seen in Figure 8, below, the age groups most represented were those of people aged 25–34 years (36%) and of people aged between 35–44 (25%). Least represented were the age groups of 55–64 year-olds (8%), as well as of seniors over the age of 65 (5%).

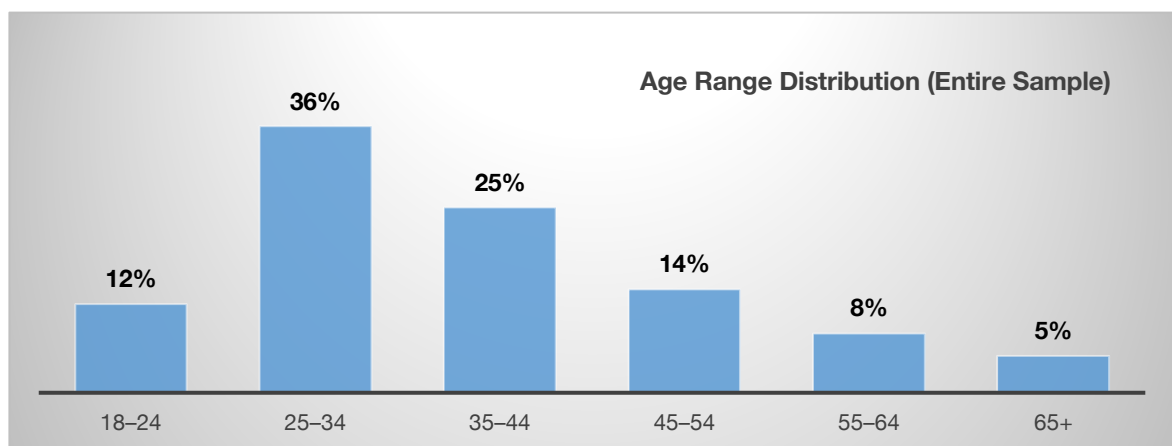


Figure 8. Age range distribution of the study population

Roughly half of the native speaker participants in this study had not obtained—or were not in the process of pursuing—an English degree at tertiary level, can is displayed in Figure 9, below. While 17.6% of the native speaker sample reported an English degree at undergraduate level, 20% of the sample indicated an English degree at graduate level and 10% specified an English degree at doctoral level. The number of non-native speaker participants who reportedly were not pursuing or had not obtained an English degree was comparatively fewer, at 42.5%, whereas 12.5% of the sample reported an English degree at undergraduate level and 35% of the participants indicated an English degree at graduate level. The share of non-native speaker participants who indicated to have obtained, or being in the process of pursuing, an English degree at doctoral level was the same as that of native speaker participants, with 10% each.

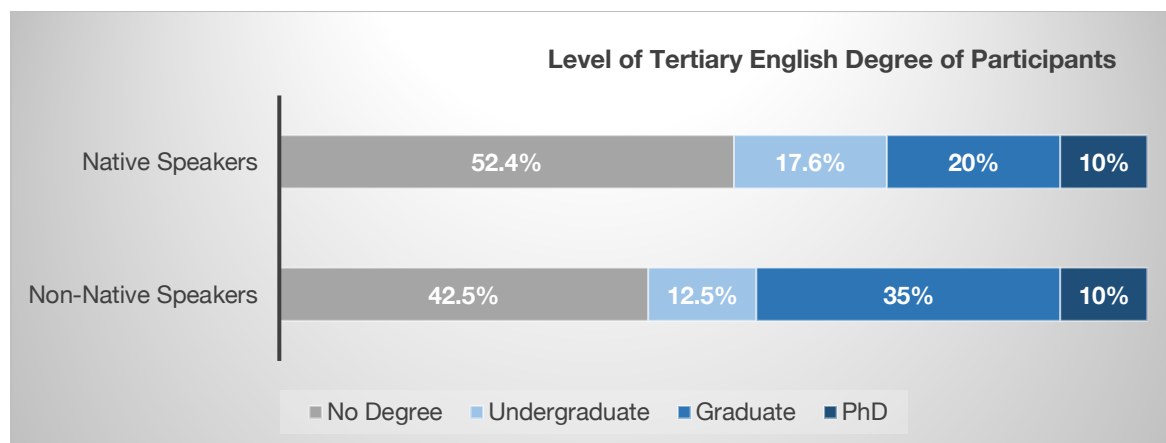


Figure 9. Level of tertiary English degree of participants

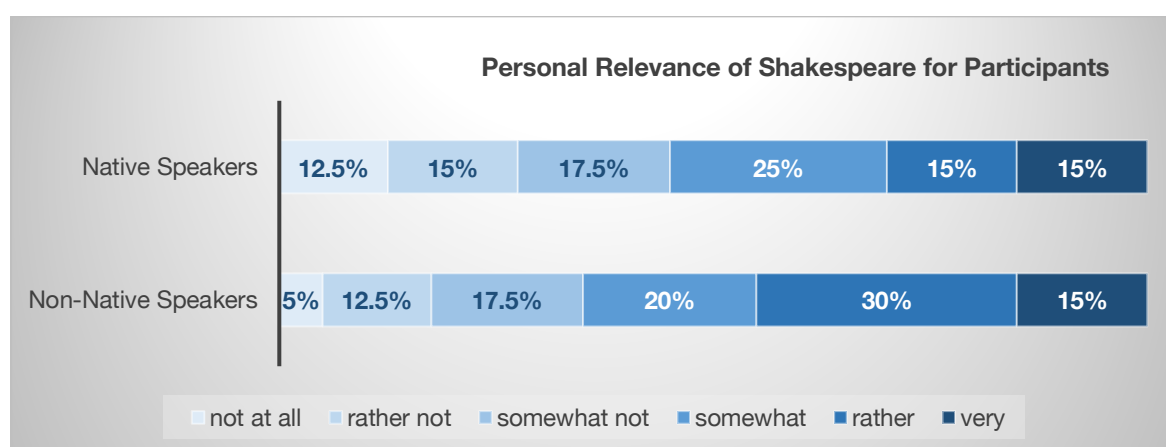


Figure 10. Personal relevance of Shakespeare for participants

Conversely, the degree of personal relevance that Shakespeare holds for the participants in this study (Figure 10, above) presented a rather marked split between native speakers and non-native speakers. Shakespeare is generally more personally relevant to non-native participants in this study, of which 65% indicated various degrees of personal relevance, whereas 35% of the same sample reported that Shakespeare is not relevant to them, to various degrees. As for native speaker participants, 55% of them indicated that Shakespeare is to various degrees relevant to them personally, while 45% percent of them reported the opposite, to various extents.

5.1.2. Native speaker respondents

The mean age of native speaker respondents (NSR) is 38.3, slightly above the mean age of the entire study population. As can be seen in Figure 11 below, roughly half of the participants in this subsample belong to the age groups of 25–34 year-olds (26%) and of

35–44 year-olds (30%), the latter cohort being the most represented among native speakers. Least represented in this subsample were the age groups of 55–64 year-olds (10%) and of seniors over the age of 65 (6%).

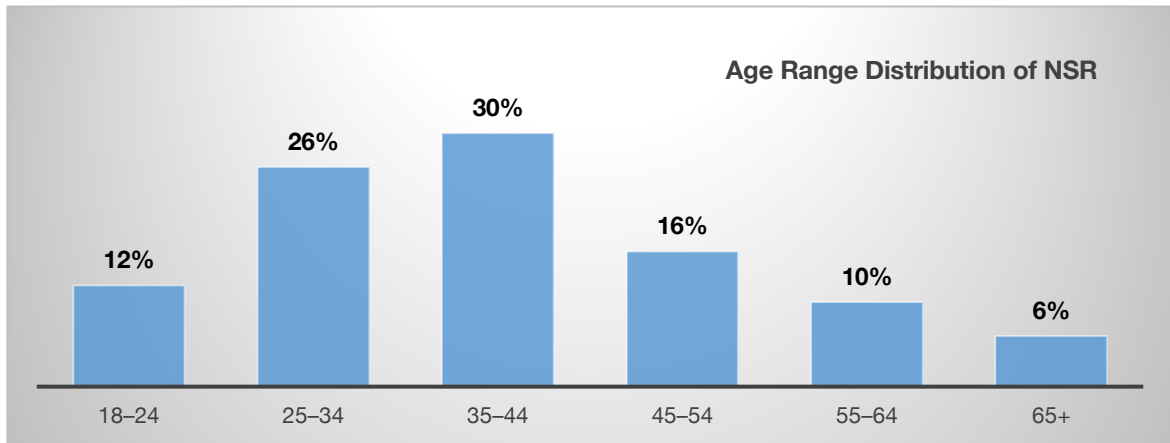


Figure 11. Age range distribution of NSR

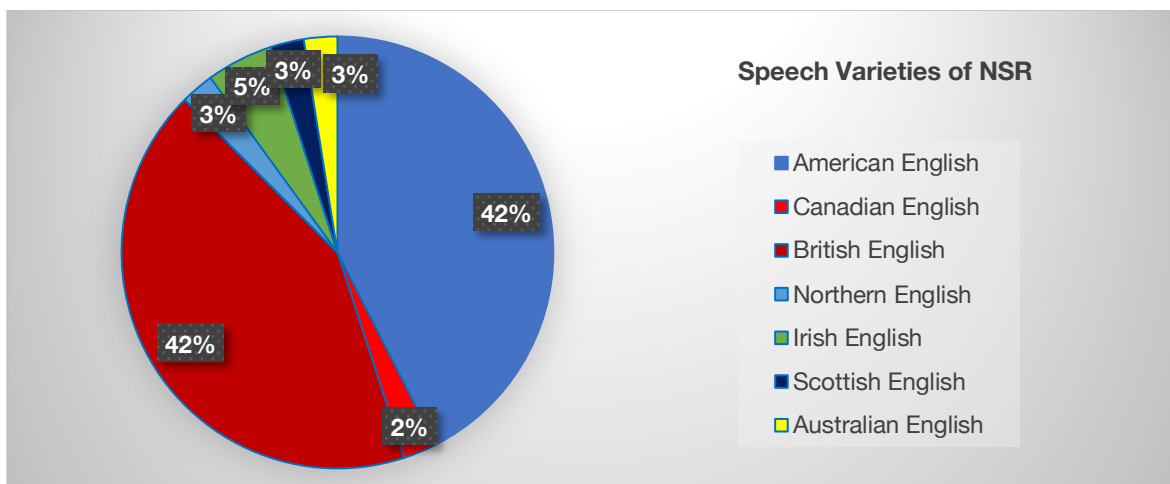


Figure 12. Speech varieties of NSR

Among the native speech varieties reported by participants in this subsample, there was an exact equivalence between British English and American English, with 42% respondents each (Figure 12, above). The native speech varieties of the remaining participants were more or less equally distributed among Canadian English, Australian English, Scottish English, Northern (England) English (3% share each), and Irish English, the latter with a marginally larger share of 5% of participants in this group.

5.1.2. Non-native speaker respondents

The mean age of non-native speaker respondents (Non-NSR) was 34.6, slightly lower than the mean age of the entire sample. As evidenced by Figure 13, below, roughly two thirds of participants in this category belong to the age groups of 25–34 year-olds (38%) and of 35–44 year-olds (26%), the former being clearly the most represented. The rest of this subsample was distributed analogously to the native speaker subsample, with the age groups of 55–64 year-olds (6%) and seniors over the age of 65 (4%) being the least represented.

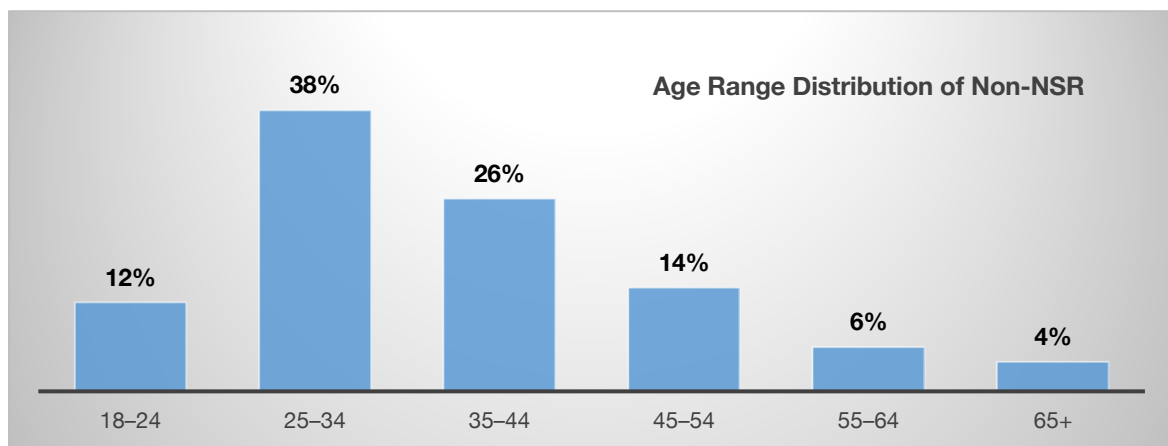


Figure 13. Age range distribution of Non-NSR

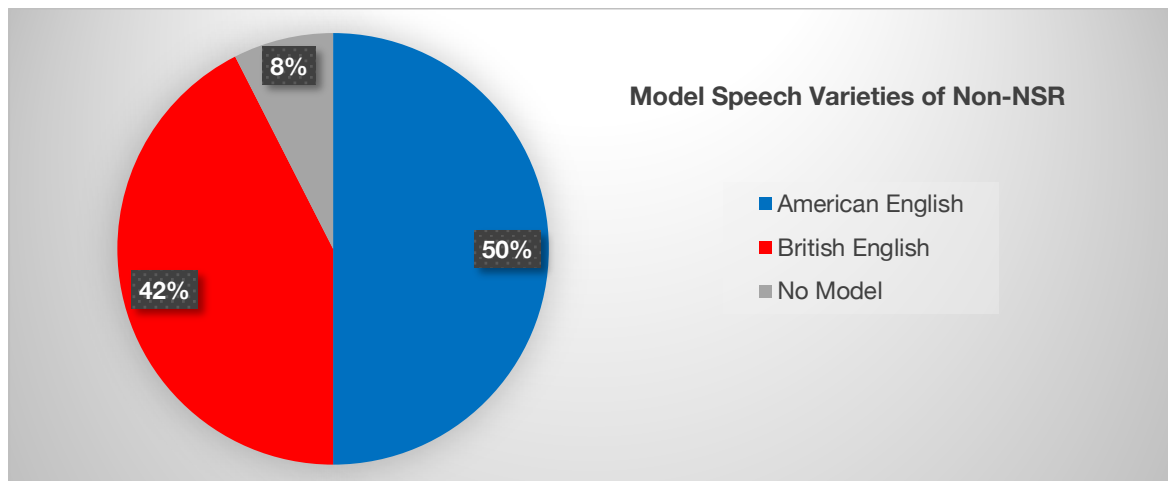


Figure 14. Model speech varieties of Non-NSR

Among non-native speaker participants there was not quite an equivalence among the varieties of English on which respondents reportedly model their speech. But with half of the subsample indicating American English and 42% reporting British English as their model variety, the gulf was quite narrow (Figure 14, above). No other English variety was

specified as speech model, while another 8% of participants in this subsample indicated that they did not acknowledge any variety serving as their reference.

5.2. Analysis of the Data

As McKenzie observes (2010: 47), the first step in the interpretation of the data is usually to conduct a form of factor analysis (which most often takes the shape of principal components analysis) in order to decrease the number of variables in the study as well as to pinpoint the salient evaluative dimensions among the traits that the respondents have judged to be important. For a variety of reasons, it was not feasible to apply this procedure in the context of this study. Instead, analysis of the data obtained through the questionnaire was conducted relying on the results of previous language attitude studies as a reference point. To that effect, the adjective pairs on the semantic differential scales were first condensed into the traits that they represent. For instance, the dichotomy *urban* / *rural* was reduced to the attribute *urbane*, while attributes corresponding to traits such as *pleasant* / *unpleasant* were renamed according to their positive poles.

Table 1. Individual Components of Evaluative Dimensions

Prestige	Accessibility	Intelligibility	Suitability for Shakespeare
<i>urbane</i>	<i>pleasant</i>	<i>clear and understandable</i>	<i>suitable for Shakespeare</i>
<i>solemn</i>	<i>not snobbish</i>	<i>understandable without text</i>	
<i>sophisticated</i>	<i>likeable</i>		
	<i>modern</i>		
	<i>authentic</i>		
	<i>interesting</i>		

These attributes were then grouped into the evaluative dimensions displayed in Table 1 above. The dimension commonly referred to as *social status* in language attitudes research (McKenzie 2010: 47)—here composed of the individual attributes *urbane*, *solemn*, and

sophisticated—has been relabeled *prestige* in the context of this study. This is because the term “prestige” was regarded to better reflect the intellectual, cultural, and aesthetic cachet that the different speech varieties possess in relation to performances of Shakespeare. Correspondingly, the term *accessibility* was considered to be a better expression of the evaluative dimension conventionally labeled as *social attractiveness* (McKenzie (2010: 47)—here composed of the individual attributes *pleasant*, *not snobbish*, *likeable*, *modern*, *authentic*, and *interesting*. As mentioned before, this evaluative dimension is partly determined by a phenomenon, known as “in-group loyalty,” that frequently contrives to lend so-called “covert prestige” to nonstandard varieties that otherwise rate low on *social status* (Dragojevic 2017). Accordingly, in the context of this study, this evaluative dimension seeks to establish how accessible—in other words, “relatable” (Lodewyck 2013: 55)—a passage of Shakespeare is when performed in different speech varieties and, furthermore, to determine if a variety that rates low on overt prestige, may indeed enjoy *covert* prestige in the context of performances of Shakespeare.

In addition to these more conventional evaluative dimensions, two additional dimensions have been identified that are highly specific to this study. The first one comprises the attributes *clear and understandable* and *understandable without (reading) the text*, and was accordingly labeled *intelligibility*. This dimension is a gauge of the respondents’ subjective impressions of how well they understand a passage from Shakespeare spoken in different speech varieties, both with and without simultaneously reading the text. The other additional dimension comprises just one attribute, *suitable for Shakespeare*, and was thus aptly named *suitability for Shakespeare*. It exactly corresponds to the direct question in the survey, “How suitable do you think Speaker X would be for a role in a Shakespeare production?”.

To elaborate the responses to the questionnaire, the four speakers were coded as follows: Speaker A = GenAm (General American), Speaker B = ContempRP (Contemporary RP), Speaker C = ConservRP (Conservative RP), Speaker D = OP (Original Pronunciation). To calculate the mean scores allocated to the individual attributes on the semantic differential scales, I first assigned numerical values to the gradient points. The label *very* on the left side of the scale was thus allocated the numerical value of 0, the label *rather* was given the numerical value of 1, and the label

somewhat was assigned the numerical value of 2. On the right side of the scale, accordingly, *somewhat* was assigned the value of 3, *rather* was allocated the value of 4, and *very* was given the value of 5. Because the numerical intervals between these values naturally mirror those on a decimal scale, the mean score of each item was converted into its decimal equivalent in order to provide better visibility of the results. In this process, the lowest possible mean score is expressed by 0 and the highest possible score is expressed by 1. For instance, if an attribute received a mean score of 3.45 on the six-point semantic differential scale, the same score is displayed as .69 in decimals. (For those attributes that had a right-left ordering, i.e. where the negative pole was located on the left side of the scale, the calculation had to be inverted.) The mean scores of the individual traits were then combined into the four dimensions outlined and discussed in the previous section: *prestige*, *accessibility*, *intelligibility*, and *suitability for Shakespeare*. In the following sections, I will present the results of the data analysis.

5.3. Main Results of the Analysis

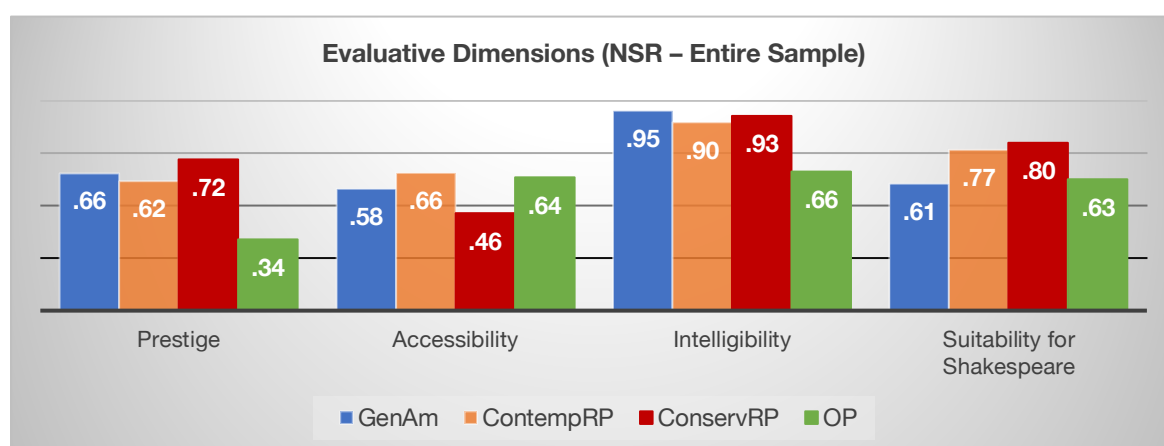
5.3.1. Native speaker respondents

The mean scores of the individual traits on the semantic differential scales expressed by native speaker respondents are displayed in Table 2 below. The GenAm speaker scored highest on the *intelligibility* traits, as well as on the *prestige* trait *urbane* and the *accessibility* trait *modern*. The ContempRP speaker scored highest on the *accessibility* traits *pleasant*, *likeable*, and *authentic*, whereas the ConservRP speaker scored highest on the prestige traits *solemn* and *sophisticated*, as well as on *suitability for Shakespeare*. The OP speaker, then, scored highest on the accessibility traits *not snobbish* and *interesting*. The resulting evaluative dimensions can be viewed in Figure 15, below. The ConservRP speaker was evaluated highest on *prestige* and was deemed most suitable for Shakespeare performances, but was scored lowest on *accessibility* by quite a wide margin. The ContempRP speaker was rated highest on *accessibility*, while the GenAm speaker was rated highest on *intelligibility*. The OP speaker was rated lowest on *prestige* by a very significant margin, with a score of less than half that of the ConservRP speaker, but was rated second highest on *accessibility*, only slightly lower than the ContempRP speaker.

Table 2. Mean Scores of Individual Traits Expressed by NSR

Trait	GenAm	ContempRP	ConservRP	OP
<i>pleasant</i>	.74	.80	.68	.64
<i>not snobbish</i>	.59	.69	.29	.76
<i>urbane</i>	.77	.64	.61	.24
<i>likeable</i>	.63	.79	.60	.73
<i>modern</i>	.59	.51	.18	.38
<i>authentic</i>	.51	.72	.71	.67
<i>solemn</i>	.51	.55	.76	.38
<i>interesting</i>	.54	.65	.57	.69
<i>sophisticated</i>	.69	.67	.80	.40
<i>clear and understandable</i>	.97	.89	.93	.61
<i>understandable without text</i>	.94	.90	.93	.72
<i>suitable for Shakespeare</i>	.61	.77	.80	.63

Note Fields marked green scored highest

**Figure 15.** Evaluative dimensions expressed by NSR (entire sample)

The GenAm speaker was rated the most *intelligible* speaker, while the OP speaker was evaluated lowest in that regard, again by a significant margin. On the dimension *suitability for Shakespeare*, the ConservRP speaker was evaluated highest, slightly above the ContempRP speaker, while the OP speaker was deemed slightly more suitable for Shakespeare than the GenAm speaker, who rated lowest on that aspect.

The rank ordering of the four speakers by native speaker respondents (Figure 16, below) appears to reflect the picture painted by the results of the evaluative dimensions. For instance, the ConservRP speaker received most first-rank mentions (42%), while also

receiving a majority of third-rank mentions (31.7%) and only a tiny minority of second-rank mentions (7.5%). The ContempRP speaker, on the other hand, received fewer first-rank rank mentions (30%) but completely eclipsed the ConservRP speaker on second-rank mentions (40%). A similar situation is displayed by the lower rungs of the preference order. For instance, although the OP speaker received the majority of fourth-rank mentions (35.9%), he actually received more first-rank mentions (15%) in this subsample than the GenAm speaker (12.5%).

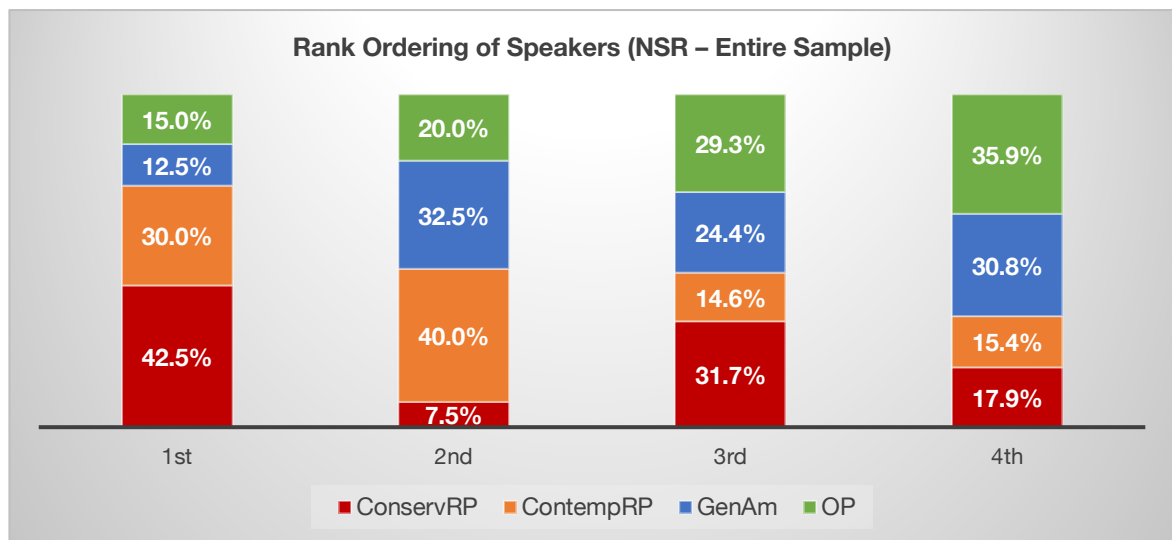


Figure 16. Rank ordering of speakers expressed by NSR (entire sample)

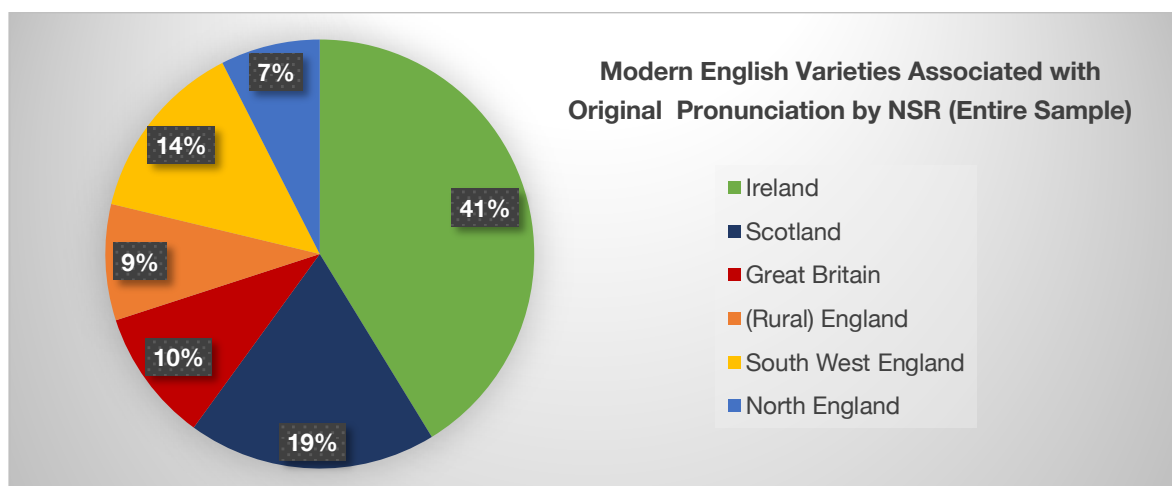


Figure 17. Modern English varieties associated with Original Pronunciation by NSR (entire sample)

The modern English varieties that respondents in the native speaker subsample associated with Original Pronunciation are displayed in Figure 17, above. A majority of native speaker respondents (41%) indicated Irish as the most likely provenance of the OP

speaker, while the next commonest association was Scottish (19%). 14% of the respondents located the OP speaker in various geographical regions of South-West England—with mentions such as Devon, Somerset, and even “darkest Bristol”—while the rest of the subsample more generically indicated the speaker’s provenance as “British” (10%), “English” or “rural English” (9%), and “Northern English” (7%).

5.3.2. *Non-native speaker respondents*

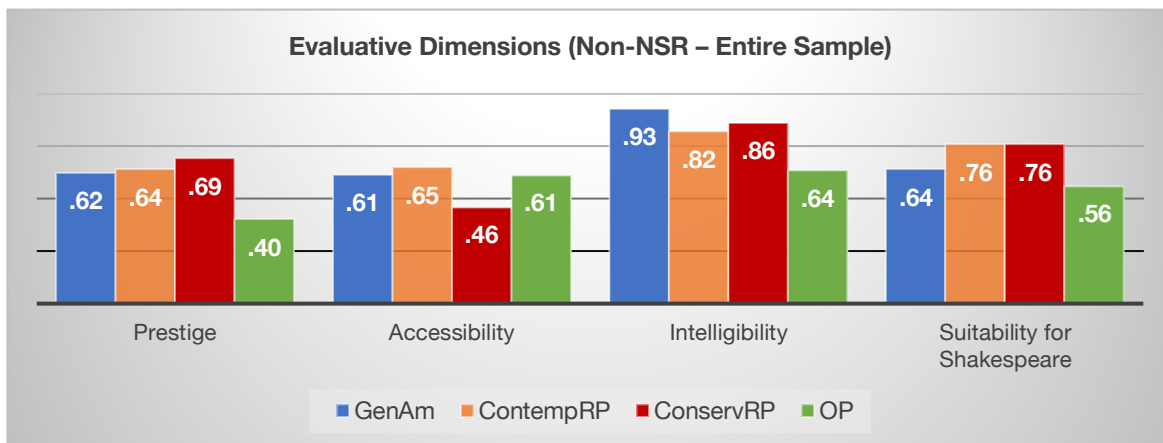
The mean scores for the individual traits of the four speakers as expressed by the non-native speaker subsample can be viewed in Table 3, below. The evaluations of the GenAm speaker mirror those expressed by native speaker respondents: highest on the *intelligibility* traits, as well as on the traits *urbane* and *modern*. The ContempRP speaker, for his part, scored highest on the *accessibility* traits *pleasant*, *likeable*, *authentic*, and *not snobbish*, and was additionally deemed the most suitable for Shakespeare performances in a tie for first place with the ConservRP speaker. The latter, in turn, scored highest on the *prestige* traits *solemn* and *sophisticated*, analogously to the evaluations expressed by native speaker respondents. The OP speaker, then, was regarded as the most *interesting* of the four speakers.

The evaluative dimensions resulting from these scores follow a similar pattern to those expressed by the native speaker subsample (Figure 18, below). The ConservRP speaker rated highest on *prestige*, followed by the ContempRP speaker and GenAm speaker, whereas the OP speaker was evaluated lowest. On *accessibility*, non-native speaker respondents rated the ContempRP speaker highest score and evaluated the OP speaker only slightly lower, on a par with the GenAm speaker, while the ConservRP speaker obtained the same lowest evaluation as expressed by the native-speaker subsample. *Intelligibility* was deemed to be highest for the GenAm speaker while the other varieties displayed an evaluative hierarchy similar to that expressed by native speaker respondents, albeit with lower ratings for both the ContempRP and the ConservRP speaker. The evaluative dimension *suitability for Shakespeare* showed the most marked difference in evaluation between the two main subsamples: while the ContempRP and ConservRP speaker tied for best evaluation in this category, the OP speaker was deemed least suitable for Shakespeare performances.

Table 3. Mean Scores of Individual Traits Expressed by Non-NSR

Trait	GenAm	ContempRP	ConservRP	OP
<i>pleasant</i>	.76	.86	.66	.67
<i>not snobbish</i>	.65	.69	.37	.68
<i>urbane</i>	.77	.70	.64	.31
<i>likeable</i>	.70	.79	.59	.70
<i>modern</i>	.60	.40	.11	.35
<i>authentic</i>	.50	.74	.72	.66
<i>solemn</i>	.50	.58	.71	.45
<i>interesting</i>	.61	.66	.49	.68
<i>sophisticated</i>	.61	.65	.73	.46
<i>clear and understandable</i>	.96	.84	.88	.60
<i>understandable without text</i>	.90	.81	.85	.67
<i>suitable for Shakespeare</i>	.64	.76	.76	.56

Note: Fields marked green scored highest

**Figure 18.** Evaluative dimensions expressed by Non-NSR (entire sample)

Analogously, the ranking order distribution expressed by non-native speaker respondents painted a markedly different picture compared to that expressed by the native speaker subset. In fact, almost half of non-native speakers (47.5%) ranked the ContempRP speaker first, while the ConservRP speaker received roughly a quarter of the top preferences, as can be viewed in Figure 19, below. In a similar vein, though the OP speaker was ranked first only by a tiny minority of the subsample (10%), he surpassed the GenAm speaker in second-rank mentions by an eight percent margin, while also receiving fewer bottom-rank mentions (35.1% vs 43.2%).

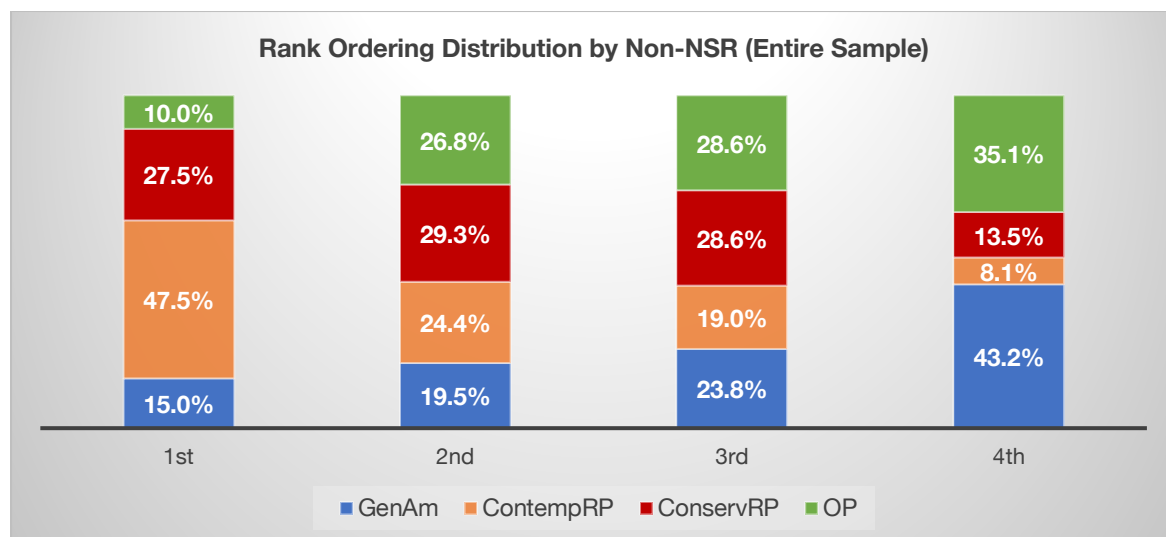


Figure 19. Rank ordering of speakers expressed by Non-NSR (entire sample)

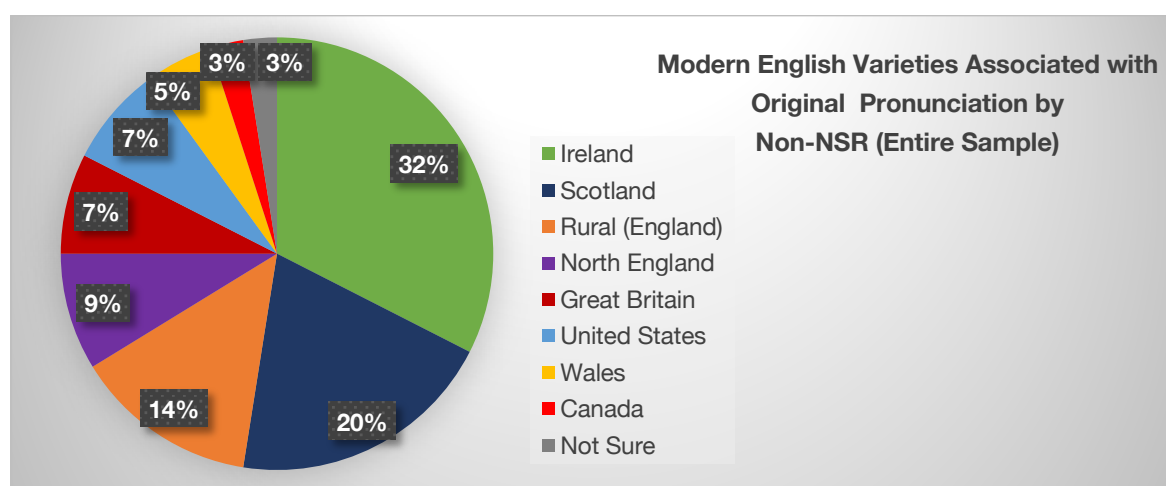


Figure 20. Modern English varieties associated with Original Pronunciation by Non-NSR (entire sample)

A majority of non-native speaker respondents also indicated Irish as the most likely provenance of the OP speaker, albeit with a lesser percentage (32%) than their native-speaker counterparts (Figure 20, above). Scottish, on the other hand, was indicated by the same share of respondents in both subsets (20%) as the next most common association with Original Pronunciation. Overall, responses by non-native speakers were more varied than those by native speakers—indicating also Wales (5%), the United States (7%), and Canada (3%) as possible provenance of the OP speaker—but did not include mentions of specific regional British areas, such as Somerset or Devonshire.

In the following sections, I will present the results of the data analysis in relation to the three secondary background variables mentioned above: the English varieties spoken by respondents, the presence or absence of an English degree at tertiary level, and

the degree of personal relevance that Shakespeare holds for the respondents. For these secondary background variables, however, I will only take into consideration the results of the evaluative dimensions as well as the rank ordering of speakers expressed by the respondents, since these two sets of data are by and large the most significant with respect to the research questions formulated in this study.

5.4. The Speech Variety Variable

5.4.1. Native speaker respondents

Because the overwhelming majority of native speaker respondents (84%) reported to speak one of just two English varieties, American English and British English, only these two sets of data were taken into account and results by the few respondents who indicated other native English varieties not considered.

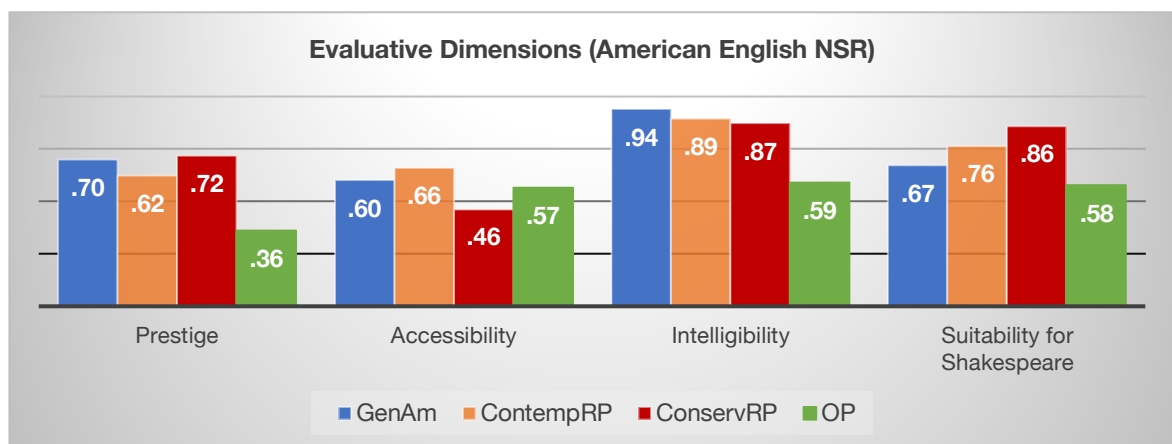


Figure 21. Evaluative dimensions expressed by American English NSR

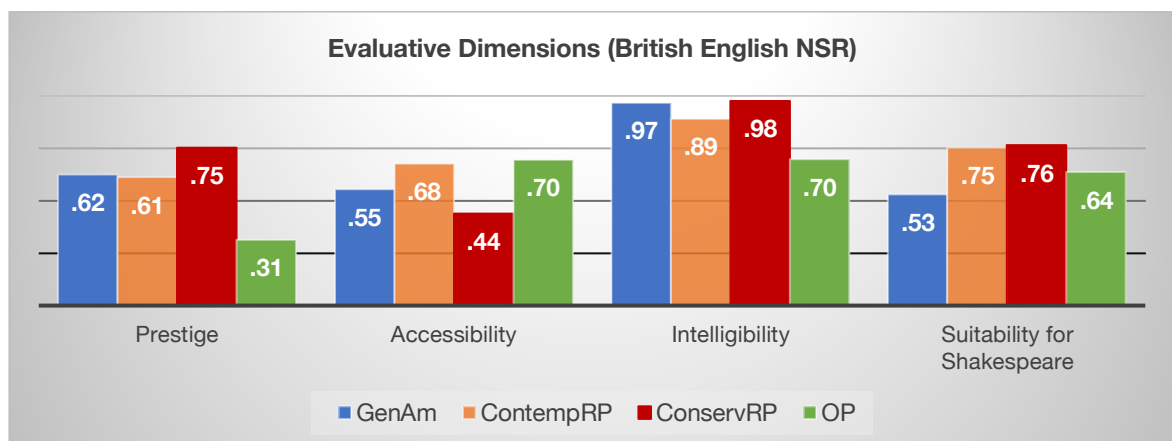


Figure 22. Evaluative dimensions expressed by British English NSR

There are both discrepancies and similarities between the evaluative dimensions expressed by American English native speakers (Figure 21, above) and those expressed by British English native speakers (Figure 22, above). Both subgroups evaluated the ConservRP speaker highest on *prestige* while evaluating the OP speaker lowest on the same dimension. American English respondents rated the ContempRP speaker highest on *accessibility*, whereas British English respondents evaluated the OP speaker highest along the same dimension. Both subgroups rated the ConservRP lowest on *accessibility* by a significant margin compared to the other speakers. American English respondents, then, regarded the GenAm speaker as the most comprehensible, while rating the OP speaker lowest on *intelligibility*, also by a substantive margin.

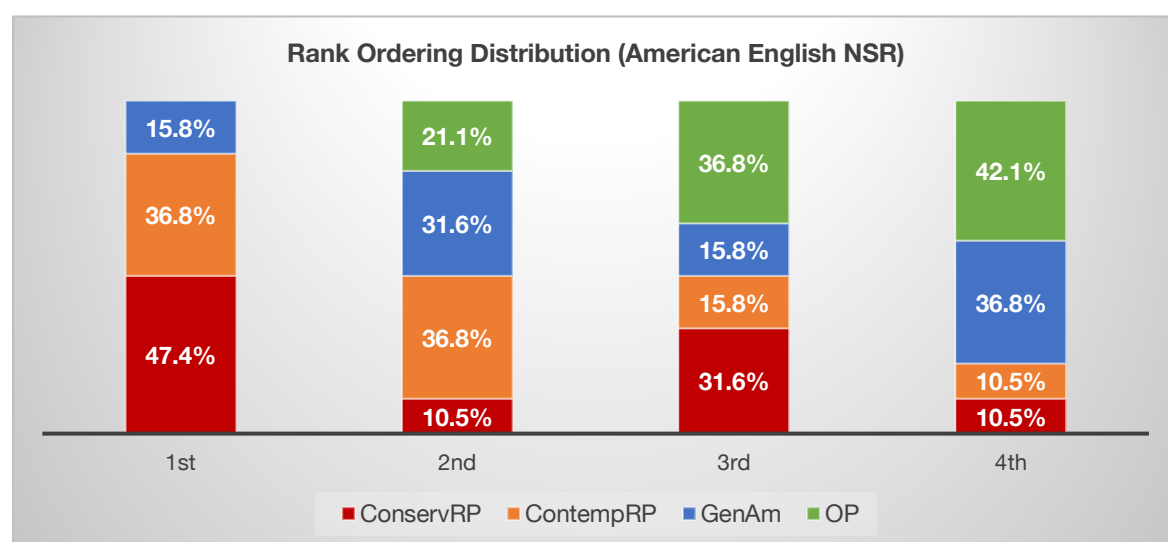


Figure 23. Rank ordering distribution expressed by American English Non-NSR

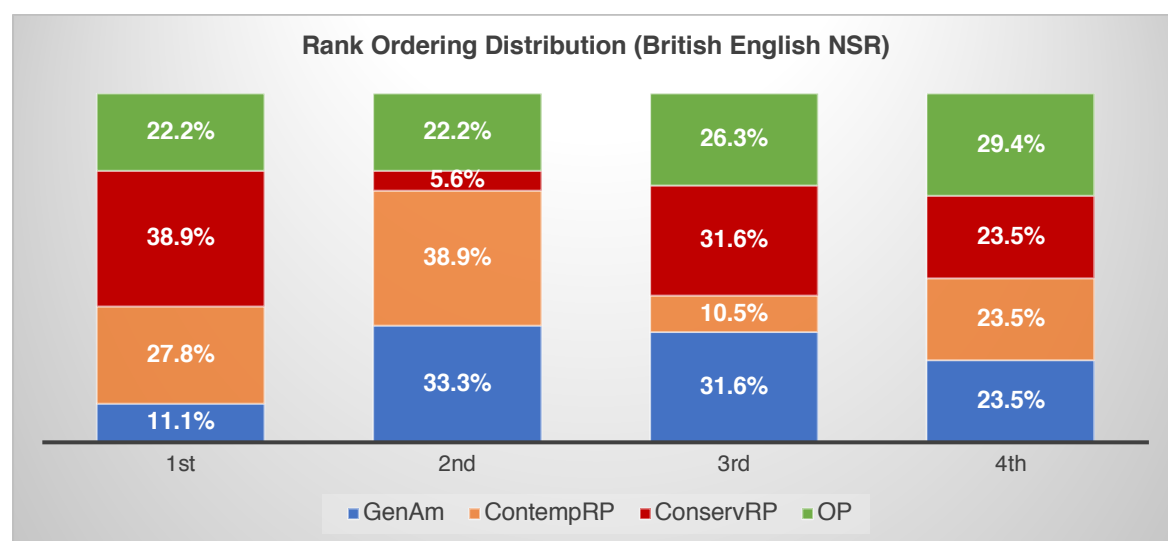


Figure 24. Rank ordering distribution expressed by British English Non-NSR

Conversely, British English respondents evaluated the ConservRP highest on the same dimension and the GenAm slightly lower, while rating the OP speaker, though also lowest, significantly higher than their American English counterparts. Finally, American English respondents deemed the ConservRP speaker as the most suitable for Shakespeare performances, while preferring the ContempRP speaker over the GenAm speaker and rating the OP speaker lowest by a significant margin. British English respondents, on the other hand, only gave a slight edge to the ConservRP speaker over the ContempRP speaker on *suitability for Shakespeare*, while deeming the OP speaker more suitable for Shakespeare than the GenAm speaker.

The rank ordering of speakers by American English native speakers (Figure 23, above) and British English native speakers (Figure 24, above) also appear to reflect the respective evaluative scores. The ConservRP speaker obtained a resounding majority of first-rank mentions by American English respondents (47.4%), whereas the OP speaker received no first-rank mention. The latter, in turn, received a clear majority of third-rank and fourth-rank mentions. British English respondents, on the other hand, though also assigning most of the first-rank mentions to the ConservRP speaker (38.9%), were far more favorable toward the OP speaker, who received roughly a quarter of mentions for each rank by respondents in this group.

5.4.2. Non-native speaker respondents

Analogously to the native speaker subsample, respondents in the non-native speaker subsample indicated only two English varieties as their chosen reference model, roughly split in half between American English (50%) and British English (42%). Results by non-native speaker respondents who indicated that they do not model their speech on any variety (8%) were not considered in the analysis of this variable. The evaluative dimensions expressed by the two subgroups also display some discrepancies. American English non-native speakers (Figure 25, below), for instance, rated the ContempRP speaker higher on *prestige* than the GenAm speaker, compared to British English non-native speakers (Figure 26, below), who gave a clear preference to the GenAm speaker. Both, however, clearly evaluated the ConservRP speaker highest and the OP speaker lowest along the same dimension.

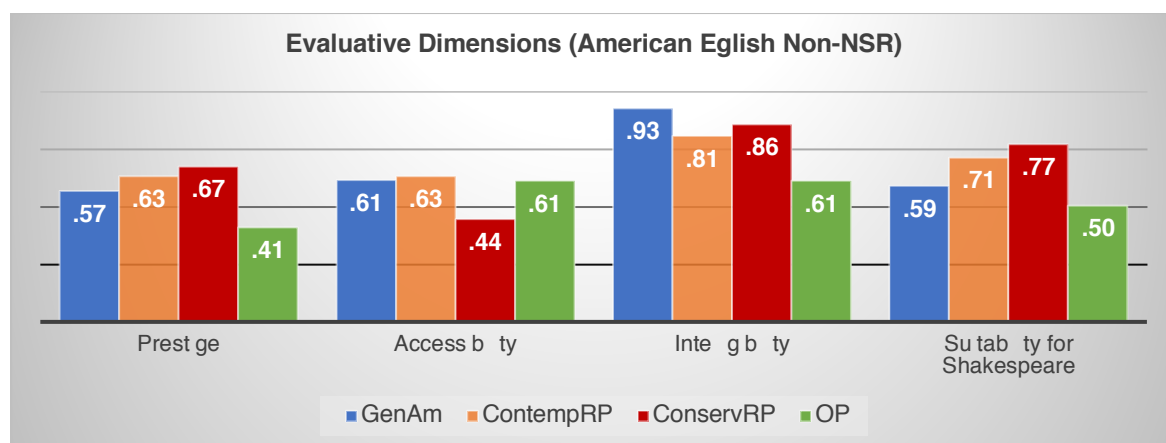


Figure 25. Evaluative dimensions expressed by American English Non-NSR

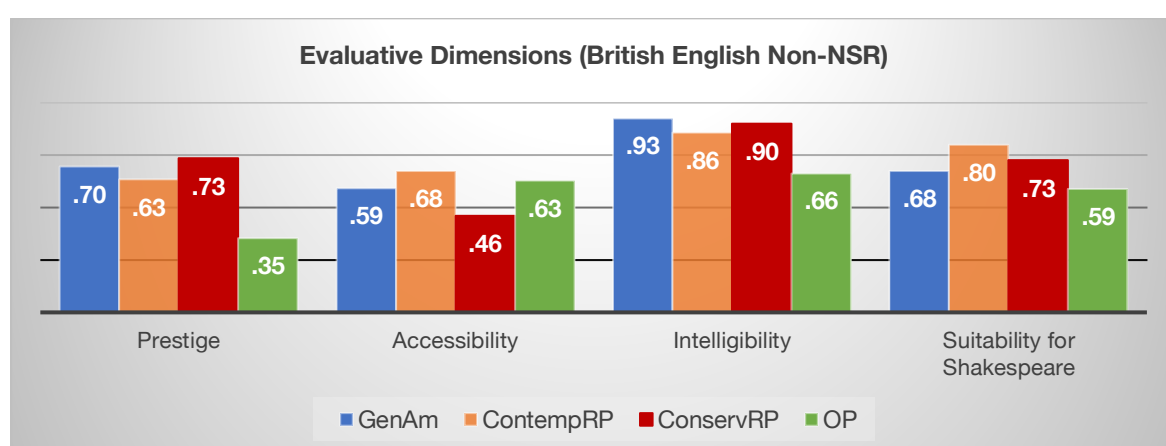


Figure 26. Evaluative dimensions expressed by British English Non-NSR

Both subgroups evaluated the ContempRP speaker highest on *accessibility*; American English non-native speakers gave him only a slight edge over the GenAm speaker, whereas that margin was significantly wider for British English non-native speakers. Both subgroups rated the OP speaker very high on the same dimension and the ConservRP very low, by a substantive margin. Evaluations of *intelligibility*, then, were found to be fairly similar between the two subgroups, exhibiting only minor differences: while American English non-native speaker found it slightly more difficult to comprehend both the ConservRP speaker and the ContempRP speaker, they expressed slightly less difficulty in understanding the OP speaker. *Suitability for Shakespeare* displayed more marked differences between the evaluations of these two subgroups. American English non-native speakers deemed the ConservRP speaker most suitable for Shakespeare, while British English non-native speakers deemed the ContempRP speaker most suitable. Both subsets,

however, preferred the GenAm speaker over the OP speaker in this regard, albeit with overall higher evaluations expressed by the British English subgroup.

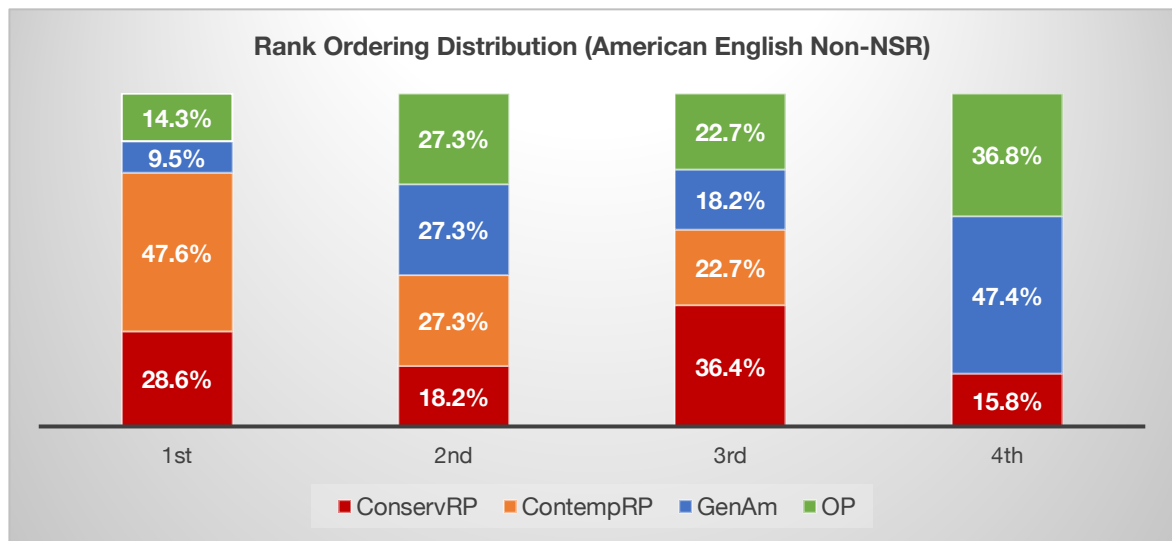


Figure 27. Rank ordering distribution expressed by American English Non-NSR

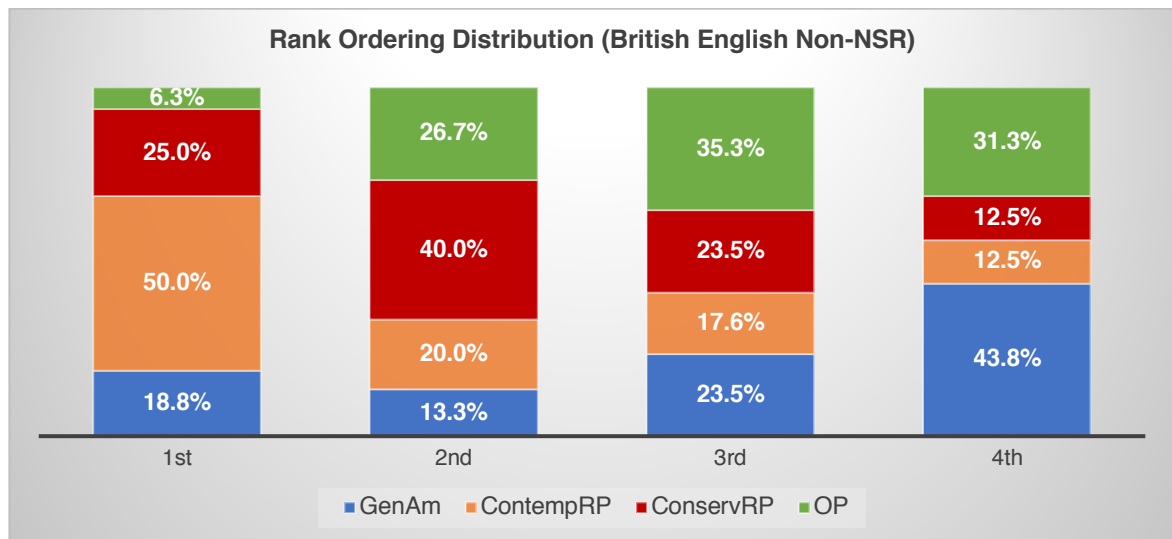


Figure 28. Rank ordering distribution expressed by British English Non-NSR

The rank ordering distribution articulated by American English non-native speakers (Figure 27, above) and British English non-native speakers (Figure 28, above) produced a far less accurate reflection of the evaluative dimensions expressed by these groups than those of the other subsets. Indeed, American English non-native speakers ranked the ContempRP speaker first with about half of the mentions, while the ConservRP speaker received roughly a quarter of their first-rank mentions (28.6%) and also obtained the fewest second-rank (18.2%) and the most third-rank mentions (36.4%). On the other hand,

the ConservRP speaker received the highest share of second-rank mentions by British English non-native speakers (40%). British English non-native speakers also ranked the ContempRP speaker first with a similar share of the mentions (50%), but ranked the ConservRP second with the highest share of mentions (40%). The OP speaker, then, received more first-rank mentions by the American English subgroup (14.3%) than by the British English subgroup (6.3%), while the GenAm speaker obtained more first-rank mentions by British English respondents (18.8%) but was ranked fourth by a majority of both groups, with a similar share.

5.5. The Tertiary Education Variable

As displayed in Table 4, below, this variable produced a notable difference in evaluations of the four speakers' *suitability for Shakespeare* among the native speaker subsample, while evaluations along the other dimensions were more or less comparable between the two subgroups. Native speakers with a tertiary English degree rated the ContempRP speaker as the most suitable for Shakespeare, while the ConservRP speaker and OP speaker were rated slightly lower in that regard and the GenAm speaker was evaluated lowest. Native speakers without a tertiary English degree, on the other hand, deemed the ConservRP by far the most suitable for Shakespeare, with a considerable margin over the ContempRP speaker, while the OP speaker was deemed least suitable by this subset of the study population, albeit only slightly so.

Table 4. Evaluative Dimensions by Tertiary English Degree of NSR

Dimension	NSR WITH Tertiary English Degree				NSR WITHOUT Tertiary English Degree			
	GenAm	Contemp RP	Conserv RP	OP	GenAm	Contemp RP	Conserv RP	OP
<i>prestige</i>	.66	.62	.75	.39	.65	.62	.70	.30
<i>accessibility</i>	.58	.66	.45	.63	.59	.65	.47	.65
<i>intelligibility</i>	.93	.88	.90	.69	.98	.91	.95	.64
<i>suitability for Shakespeare</i>	.63	.79	.73	.68	.58	.74	.87	.57

Note Fields marked green scored highest

Applying this variable to the non-native speaker subsample resulted in even more marked differences in evaluations between the two subsets based on tertiary education, as can be seen in Table 5, below. Non-native speaker respondents with a tertiary degree in English clearly rated the ConservRP speaker highest on *prestige*, by a significant margin over the GenAm speaker and the ContempRP speaker. The latter, in turn, was evaluated highest on *accessibility* by the same subgroup. The GenAm speaker was found to be the most intelligible of the four speakers, while the ContempRP speaker clearly emerged as the most suitable for Shakespeare performances in the evaluations of this subset.

Table 5. Evaluative Dimensions by Tertiary English Degree of Non-NSR

Dimension	Non-NSR with Tertiary English Degree				Non-NSR WITHOUT Tertiary English Degree			
	GenAm	Contemp RP	Conserv RP	OP	GenAm	Contemp RP	Conserv RP	OP
<i>prestige</i>	.63	.60	.71	.41	.61	.69	.67	.40
<i>accessibility</i>	.60	.68	.43	.63	.64	.61	.51	.59
<i>intelligibility</i>	.95	.87	.90	.63	.89	.76	.81	.64
<i>suitability for Shakespeare</i>	.63	.79	.71	.50	.66	.72	.82	.65

Note Fields marked green scored highest

Non-native speaker without a tertiary English degree, on the other hand, rated the ContempRP speaker highest on *prestige*, with a slight margin over the ConservRP speaker and the GenAm speaker. The latter, in turn, was evaluated highest on *accessibility* by this subset, with a slight margin over the ContempRP speaker and OP speaker. *Intelligibility* of the speakers displays an identical hierarchy of evaluations between the two subsets. Out of all the subgroups in the study population, non-native speaker respondents without a tertiary English degree overall expressed the lowest ratings along *intelligibility* of the four audio stimuli.

5.6. The Personal Relevance of Shakespeare Variable

In order to analyse participants' responses based on the degree to which Shakespeare's works holds a personal relevance for them, I separated the results of the questionnaire item, "How would you describe the relevance of Shakespeare's works for you personally?", according to whether the respondents had selected one of the gradient points (*rather, somewhat, very*) either on the "positive" pole (labeled *relevant*) or on the "negative" pole (labeled *not relevant*). Among the native speaker subsample, 52.5% of respondents indicated that Shakespeare is, to various degrees, relevant to them (*somewhat* = 42.9%, *rather* = 28.6%, *very* = 28.6%), while 47.5% indicated that Shakespeare is not relevant to them, to various extents (*somewhat* = 42.1%, *rather* = 31.6%, *very* = 26.3%).

Table 6. Evaluative Dimensions by Personal Relevance of Shakespeare for NSR

Dimension	NSR for whom Shakespeare is relevant				NSR for whom Shakespeare is NOT relevant			
	GenAm	Contemp RP	Conserv RP	OP	GenAm	Contemp RP	Conserv RP	OP
<i>prestige</i>	.65	.64	.75	.36	.67	.59	.69	.31
<i>accessibility</i>	.60	.66	.45	.65	.55	.66	.48	.64
<i>intelligibility</i>	.96	.91	.94	.67	.94	.89	.93	.66
<i>suitability for Shakespeare</i>	.62	.78	.75	.61	.58	.72	.83	.63

Note: Fields marked green scored highest

As displayed in Table 6, above, this variable produced hardly any difference in evaluations along the dimensions of *prestige*, *accessibility*, and *intelligibility* between native speakers for whom Shakespeare is relevant and those for whom Shakespeare is not relevant. There were, however, major differences in how the two subgroups judged the *suitability for Shakespeare* of the four speakers. Respondents for whom Shakespeare is relevant, in fact, deemed the ContempRP speaker most suitable for Shakespeare, with a slight margin over the ConservRP speaker. They also rated the GenAm speaker slightly higher than the OP speaker along the same dimension. The ConservRP speaker, in turn, was deemed most suitable for Shakespeare performances by those respondents in this subsample for whom Shakespeare is not relevant, with a significant margin over the

ContempRP speaker. On the other hand, the same respondents deemed the OP speaker more suitable for Shakespeare than the GenAm speaker, who was rated least on this aspect.

In the non-native speaker subsample, then, 62.5% of the respondents indicated that Shakespeare was, to various extents, relevant to them (*somewhat* = 28%, *rather* = 48%, *very* = 24%), while 37.5% of them indicated that Shakespeare was not relevant to them, to various extents (*somewhat* = 50%, *rather* = 35.7%, *very* = 14.3%). In contrast to the native speaker subsample, there were more marked differences in evaluations of the four speakers between respondents in this subgroup for whom Shakespeare is relevant and those for whom Shakespeare is not relevant.

Table 7. Evaluative Dimensions by Personal Relevance of Shakespeare for Non-NSR

Dimension	Non-NSR for whom Shakespeare is relevant				Non-NSR for whom Shakespeare is NOT relevant			
	GenAm	Contemp RP	Conserv RP	OP	GenAm	Contemp RP	Conserv RP	OP
<i>prestige</i>	.63	.60	.71	.37	.62	.68	.66	.48
<i>accessibility</i>	.59	.68	.44	.65	.66	.61	.50	.51
<i>intelligibility</i>	.94	.84	.90	.69	.90	.78	.78	.54
<i>suitability for Shakespeare</i>	.69	.78	.73	.56	.56	.71	.80	.59

Note Fields marked green scored highest

As can be seen in Table 7, above, both subgroups evaluated the GenAm speaker highest and the OP speaker lowest along the dimension of *prestige*. Nevertheless, responses by those for whom Shakespeare is not relevant established a clear hierarchy of the four speakers in this regard, whereas those for whom Shakespeare is relevant rated the ConservRP and ContempRP speaker almost at the same level. The two subsets of respondents diverged most in their evaluations of *accessibility* of the speakers. Indeed, though both rated the ContempRP speaker highest in that regard, respondents for whom Shakespeare is relevant evaluated the OP speaker only slightly lower than the ContempRP speaker on *accessibility*; whereas those for whom Shakespeare is not relevant clearly rated

the OP speaker far lower along the same dimension, by a substantial margin and only slightly higher than the ConservRP speaker, who received the lowest evaluation.

Perceived *intelligibility* of the four speakers was also evaluated with pronounced differences by the two subgroups. Overall, non-native speaker respondents for whom Shakespeare is not relevant expressed much lower ratings than their counterparts for whom Shakespeare is relevant. And while both deemed the GenAm speaker as being most comprehensible, those for whom Shakespeare is relevant expressed less difficulty in understanding the ConservRP speaker compared to the ContempRP speaker, whereas respondents for whom Shakespeare is not relevant rated both speakers equally. Both subgroups, then, judged the OP speaker to be the least intelligible, however those for whom Shakespeare is relevant did so with a much higher evaluation than respondents for whom Shakespeare is not relevant. The *suitability for Shakespeare* dimension further displayed marked differences between the two subgroups: those respondents for whom Shakespeare is relevant deemed the ContempRP speaker most suitable for Shakespeare performances, followed by ConservRP speaker, GenAm speaker, and OP speaker. Non-native speakers for whom Shakespeare is not relevant, on the other hand, deemed the ConservRP speaker most suitable for Shakespeare performances, while preferring the OP speaker to the ContempRP speaker and the GenAm speaker, who was rated lowest along this dimension. The following sections of this thesis are dedicated to a discussion of the results of the data analysis; after that, I will lay out the conclusions that can be drawn from this study regarding peoples' evaluative reactions toward Original Pronunciation as it compares with the modern British and American standard English varieties in relation to performances of Shakespeare.

5.7. Discussion of Results

5.7.1. Prestige

The hypothesis was formulated that Original Pronunciation will rate low on *social status*—which in the context of this study has been renamed *prestige*—whereas the modern British and American standard varieties of English will rate high on the same evaluative dimension. This hypothesis was substantially confirmed by the results of the main speaker subsamples, who both evaluated the ConservRP speaker highest on *prestige*

while rating the OP speaker lowest on the same dimension by a significant margin. Analysis of the various subsets of the study population, produced by the application of different background variables, shows only minor deviations from the main subsamples in two subgroups: non-native speakers without a tertiary English degree and respondents for whom Shakespeare is not relevant both rated the Contemporary RP speaker highest on *prestige*.

As has been predicted, the very low evaluation of Original Pronunciation along the evaluative dimension of *prestige*, compared to those of the modern American and British standard varieties, is in line with previous language attitude research, which has evidenced that nonstandard varieties are consistently rated lower on *prestige* than standard varieties—a phenomenon that has been observed world-wide and cross-culturally and that according to some scholars, such as Millroy (2016), reflects a so-called “standard language ideology” prevalent in many societies. In that regard, people’s attitudes toward the language used for performing Shakespeare in the realm of the arts appear indeed to be a mirror reflection of their attitudes toward languages in general: Conservative RP is assigned the most prestige by virtue of it being the “standard,” or “prototypical” English variety for Shakespeare performances. The perhaps most intriguing result from the analysis of this evaluative dimension is that General American, very surprisingly, was evaluated higher on *prestige* than Contemporary RP by all subsets of the native speaker subsample as well as by a few subsets of the non-native speaker sample—most prominently by those who indicated British English as their model speech variety. Evaluations of *prestige*, however, do not necessarily correlate with *suitability for Shakespeare*, as will yet be discussed in more detail.

5.7.2. Accessibility

The hypothesis was formulated that Original Pronunciation will rate high on *social attractiveness*—which in the context of this study has been renamed *accessibility*—whereas the modern British and American standard varieties of English will rate low on the same evaluative dimension. This hypothesis was partly rebuked by the results of the quantitative data analysis. Though Original Pronunciation was indeed evaluated high on *accessibility*, particularly by native speakers, the modern American and British standard varieties were not all rated lower than Original Pronunciation. In fact, Contemporary RP

was rated highest, albeit with only a slight margin over Original Pronunciation, by almost all subsets of each main subsample—with a few interesting exceptions: British English native speakers, for instance, evaluated Original Pronunciation highest on *accessibility* by a significant margin. Conservative RP, on the other hand, was clearly rated lowest along this dimension by all subgroups of the main subsamples, bar native speakers for whom Shakespeare is not relevant. (Respondents in this category however, proved to be a veritable outlier from the other subsets, because they rated Original Pronunciation very low on *accessibility*, by more than half, on average, compared to the other subgroups.) The very low evaluations of Conservative RP on *accessibility* are also in line with previous language attitudes research, which has shown that standard varieties frequently rate low on *social attractiveness* even though they rate high on *social status*. Conversely, the high evaluations of Original Pronunciation along the *accessibility* dimension can be interpreted as an indicator of a “covert prestige” that, as has been revealed by previous language attitudes studies, nonstandard varieties frequently exhibit. This covert prestige assigned to Original Pronunciation is evidently expressed highest by British native speaker respondents, who, as I am yet to discuss, associate reconstructed Elizabethan speech with several different British regional varieties.

The argument in favor of Original Pronunciation possessing covert prestige for Shakespeare performances is further strengthened by the results of three individual components of the evaluative dimension *accessibility*. For instance, the OP speaker was considered the least *snobbish* of the four speakers, whereas the ConservRP speaker was considered the most *snobbish*, both with a wide margin to the other speakers. Most significant, perhaps, is that the OP speaker was evaluated as the most *authentic* by the participants in this study, both native speakers and non-native speakers. In comparison, the GenAm speaker rated very low on that account. This is a clear sign that people regard General American as more *inauthentic* for Shakespeare performances, and thus contrived, than British varieties—including Original Pronunciation—who are regarded as more “genuine” modes of speech. It also provides further evidence pointing to the conclusion that the overt prestige that a speech variety has for respondents does not necessarily correlate with its perceived *suitability for Shakespeare*, as expressed by those same respondents. Crucially, in that regard, the OP speaker was also rated highest on the

individual component *interesting*. This can be viewed as a solid indicator of the considerable potential that reconstructed Elizabethan speech holds to keep theater audiences engaged in performances of Shakespeare. By comparison, Contemporary RP was also rated highly in this regard, whereas Conservative RP was rated lowest by a significant margin, indicative of the low potential this English variety has for effectively keeping audiences engaged during a Shakespeare performance. The results of this evaluative dimension appear to provide empirical evidence for Crystal's (2005) claim that Original Pronunciation is more relatable than Received Pronunciation for audiences experiencing Shakespeare. This, however, only applies to the conservative variety, whereas Contemporary RP is still seen as slightly more relatable than Original Pronunciation, according to the results of this analysis.

5.7.3. *Intelligibility*

The hypothesis was formulated that Original Pronunciation will not prove more difficult to understand than the modern British and American standard varieties of English. This hypothesis was comprehensively rebuked by the results of the data analysis. Indeed, Original Pronunciation emerged as the least comprehensible of the four varieties in relation to performances of Shakespeare—by a substantial margin and as indicated by both native speaker and non-native speaker respondents. A few additional interesting indications emerge from further analysis of the variables applied to the two main subsamples. The sub-groups who expressed least difficulty in understanding Original Pronunciation were British-English native speakers in general, and, in particular, those who have obtained, or were pursuing, a tertiary English degree. This seems hardly surprising, for, as we will see in the next section, Original Pronunciation has echoes of several English rural dialects for British native speakers. The very specific indications of regional English provenance that British English native speaker respondents associated with Original Pronunciation are a clear indication of that. Interestingly, the results expressed by non-native speakers along this evaluative dimension suggest that they did not have substantially more difficulty understanding Original Pronunciation than native speakers. (Only one sub-group stood out in that regard: non-native speaker respondents for whom Shakespeare is not relevant expressed significantly more difficulty in understanding the OP speaker.) This further adds empirical evidence to the anecdotal

claims by Crystal (2005) and Barrett (2013) that foreign spectators in attendance at Shakespeare's Globe's Original Pronunciation productions in the 2010s gave overwhelmingly positive feedback of these performances.

5.7.4. Suitability for Shakespeare

The hypothesis was formulated that Original Pronunciation will not be deemed less suitable for performances of Shakespeare than the modern British and American standard varieties of English. This hypothesis was mostly rebuked by the results of the analysis. Indeed, though Original Pronunciation was rated slightly higher than General American in this regard by native speaker respondents, it was deemed considerably less suitable for Shakespeare than both Conservative RP, which was evaluated highest, and Contemporary RP. Non-native speakers, for their part, clearly deemed Original Pronunciation the least suitable variety for Shakespeare performances by a significant margin compared to the modern American and British standard varieties. Analysis of the non-native speaker subsample produced no deviation between evaluations by respondents who indicated British English and those who indicated American English as their speech model. Analysis of the native speaker subsample, however, showed a deviation that, though not dramatically altering the overall picture, is perhaps of some significance. Indeed, American English native speakers overall rated Conservative RP highest by a very significant margin over Contemporary RP and, in turn, rated Original Pronunciation lowest on *suitability for Shakespeare* by a significant margin to General American. British English native speakers, on the other hand, evaluated Conservative RP only very slightly more suitable for Shakespeare than Contemporary RP, while they also clearly deemed Original Pronunciation more suitable than General American, which was rated lowest.

As has been suggested previously, evaluations of *suitability for Shakespeare* do not correlate exactly with the *prestige* dimension—and to an even lesser degree with the *accessibility* dimension. One could have expected *prestige* to be a good predictor of *suitability for Shakespeare*, and to a certain extent this appears true. Conservative RP, for instance, was both evaluated highest on *prestige* and deemed the most suitable variety for Shakespeare performances, by native speaker and non-native speaker respondents alike, but was rated lowest by far on *accessibility*. General American, on the other hand, was rated similarly to Contemporary RP on *prestige*—unsurprisingly, a spike in evaluations

came from American English native speakers—but was deemed substantially less suitable for Shakespeare than Contemporary RP by all categories of respondents. The very high *suitability for Shakespeare* evaluations of Contemporary RP, in turn, do not mirror the evaluations that this variety received along the *prestige* dimension, nor apparently those along the *accessibility* dimension. Undoubtedly, the evaluations of Original Pronunciation on the *suitability for Shakespeare* dimension present the most intriguing case, by displaying the most jarring difference in results compared with its evaluations along the *prestige* dimension. Indeed, the comparatively high evaluations of Original Pronunciation on *suitability for Shakespeare* seem to correlate with the high evaluations it received on the *accessibility* dimension—certainly for native speaker respondents.

As we will see in the next sections, many British English respondents had precise associations of Original Pronunciation with regional nonstandard varieties of English in mind when they evaluated the OP speaker. Upon this consideration, the comparatively high evaluations of Original Pronunciation regarding its *suitability for Shakespeare* can be explained as a direct result of Original Pronunciation having high covert prestige. As has been previously discussed, nonstandard language varieties—which Original Pronunciation can be regarded as—are frequently evaluated high on *social attractiveness*—here *accessibility*—which confers these varieties so-called covert prestige. By extension, Conservative RP emerges as the most “standard” English variety for performances of Shakespeare in the perception of respondents. The very low evaluations of Conservative RP along the *accessibility* dimension—in other terms, *social attractiveness*—further strengthen this assumption, because previous language research has shown that standard English varieties are often rated low on *social attractiveness*.

5.7.5. Ranking order of preferences

The hypothesis was formulated that Original Pronunciation will not rank lower in preference as a medium for performances of Shakespeare than modern British and American standard varieties of English. This hypothesis was also mostly rebuked. In fact, Original Pronunciation was overall clearly ranked lower than Conservative RP and Contemporary RP, and marginally lower than General American. The only exception—as had for evaluations of *suitability for Shakespeare*—were the subgroup of British English native speakers, who ranked Original Pronunciation higher than General American.

Indeed, it appears that the order of preference of the four English varieties for performances of Shakespeare as expressed by the respondents more or less directly correlates with their perceptions of the *suitability for Shakespeare* of the same varieties (with the exception of the American English and British English non-native speaker subsets).

5.7.6. Modern-day speech varieties associated with Original Pronunciation

The hypothesis was formulated that Irish English will be the modern-day speech variety most closely associated with Original Pronunciation. This hypothesis was confirmed by the analysis of responses to the question, “Where do you think Speaker X comes from?”. Irish English was the most frequent mention in both main subsamples, with 41% of native speaker respondents and 32% of non-native speaker respondents indicating “Ireland” as the most likely provenance of the OP speaker. This is an empirical validation of Crystal’s (2005) anecdotal reports that Irish English is the commonest association that people have with Original Pronunciation, even though, by his analysis, “only a few of the features of OP have a direct correspondence with modern Irish accents” (2013: 7). Additionally, a fifth of the respondents in each subsample indicated that the OP speaker hailed from Scotland. It is noteworthy, perhaps, that native speaker respondents—even those who indicated speaking American English—did not associate Original Pronunciation with any other English variety from outside the British Isles, whereas several responses by non-native speakers provided mentions such as “United States” and “Canada.” This can most likely be attributed to an unfamiliarity with different English speech varieties on the part of many non-native speaker participants in this study. That such a significant share of respondents associate Original Pronunciation with one particular English variety may, however, also prove to be a limitation of the research instrument—more specifically of the matched guise technique—as an effective means of accomplishing the goal of objectively measuring people’s attitudes toward Original Pronunciation in relation to performances of Shakespeare. I will address this issue, as well as other limitations of the study, in the following section.

5.8. Limitations of the Study

The study conducted for this thesis presents a series of limitations that were mostly imposed by the unconventional subject of the inquiry and by constraints of time and available resources. To begin with, the sample size of respondents can certainly be described as problematic. On the one hand, it can be considered as statistically valid, because all age groups were represented both in the main sample and in the two subsamples determined by the language proficiency of the respondents. On the other hand, a comparatively small sample size of hundred respondents—effectively reduced to fifty in the two subsamples—does not allow generalization to the population because of the overrepresentation of some age groups, specifically of the cohorts aged 25–44 years. The limited size of the study population also made it impossible to perform a principal component analysis on the results of the data analysis, as has been mentioned. Instead, it was necessary to rely on the statistically validated results of previous language attitude studies to group the attributes into their relevant evaluative dimensions.

Another limitation of this study concerns the research instrument, more precisely the stimulus material used for the questionnaire as part of the matched guise technique. Owing to the lack of available resources, it was not possible to create them specifically for this study. Although the recording used as stimulus for the evaluation of Original pronunciation had been produced by Ben Crystal—one of the foremost exponents of the OP movement, thus ensuring a high quality of recording and accuracy of speech—it also severely limited the choice of passages from Shakespeare that could be selected. A more valid approach would undoubtedly have been to record a passage from a play by Shakespeare containing dialogue between two characters—one male and one female—being as “neutral” as possible with regard to its tragic or comedic connotations; ideally, a dialogue that does not contain famous or memorable lines that would be too loaded with associations of “prototypical” interpretations for respondents. It would also be crucial to produce the different versions as guises by the same actors, in order to benefit from the full potential of the matched guise technique.

The research instrument itself might ultimately have been a limitation to this study, as mentioned in the previous sections. Indeed, evaluations of Original Pronunciation may have been distorted by the fact that the respondents were not aware of the existence of

Original Pronunciation. Therefore, the argument could be made that an indirect approach may not, after all, be the most suitable methodology for this research question; or, at any rate, that a study on Original Pronunciation involving the matched guise technique should only be conducted among respondents who are aware of the use of reconstructed Elizabethan speech for theatrical performances. For the above reasons, the empirical study conducted for this thesis should only be regarded as providing intriguing preliminary data that could serve as a starting point for more thorough scholarly research into the subject.

6. Summary and Conclusion

Interest in performances of Shakespeare using reconstructed Elizabethan speech, referred to as “Original Pronunciation,” began in the nineteenth century. From the beginning, it took the form of a combined effort between linguists such as Daniel Jones and theater practitioners such as Bernard Miles. The present-day Original Pronunciation revival has been similarly spearheaded by a close collaboration between Shakespeare’s Globe and David Crystal, which resulted in the Original Pronunciation productions of *Romeo and Juliet* (2004) and of *Troilus and Cressida* (2005). These performances sparked a minor “OP” movement leading to several productions over the following years, mostly in North American academic settings.

The phonology of Original Pronunciation is quite distinct from modern American and British English standard varieties, most notably in its rhoticity and pronunciation of long vowels and diphthongs. As such, contemporary audiences tend to predominantly associate Original Pronunciation with Irish English, Scottish English, as well as with some Southwestern English dialects. On the contemporary British stage, Conservative RP is still the most widely used English variety for performances of Shakespeare, although it has recently ceded ground to its more modern incarnation, Contemporary RP, particularly in the performances by a younger generation of British actors such as Benedict Cumberbatch. Against this backdrop, practitioners who have employed Original Pronunciation for Shakespeare performances in recent times have reported overwhelmingly positive reactions from spectators in attendance. Crystal (2005) judges these anecdotal reports to be proof that modern audiences can relate to and engage with Shakespeare’s works more if they are performed in Original Pronunciation.

People’s deep-seated convictions about languages are commonly referred to as “language attitudes.” The concept of “attitude” is one of the most crucial notions in social psychology. Garret et. al. (2003: 3), for instance, define it as “an evaluative orientation to a social object of some sort” that, being a “disposition,” is “at least potentially an evaluative stance that is sufficiently stable to allow it to be identified and in some sense measured.” Language attitudes can thus be described as evaluative reactions to language varieties, which engender processes known as social categorization and social stereotyping of speakers. Past language attitude research has shown that, world-wide and

cross-culturally, people consistently evaluate language varieties, and thus their speakers, primarily along two dimensions: *social status* and *social attractiveness*. Speakers of standard varieties are usually attributed higher *social status* than speakers of nonstandard varieties. The latter however, are frequently attributed so-called “covert prestige” due to them being evaluated higher along the *social attractiveness* dimension compared to speakers of standard varieties.

In order to validate the claims by theater practitioners and scholars that using reconstructed Elizabethan English speech brings modern audiences “closer” to Shakespeare, an empirical study of people’s attitudes toward Original Pronunciation was conducted for this thesis. Based on Crystal’s (2005) anecdotal reports as well as on the results of previous language attitudes research, it was hypothesized that Original Pronunciation would be evaluated lower on *social status* but higher on *social attractiveness* than the modern British and American English standard varieties. Additionally, it was postulated that Original Pronunciation would not be deemed less suitable and less intelligible for performances of Shakespeare than the modern British and American English standard varieties and that it would rank high in preference compared with these varieties. Finally, it was suggested that Irish English would be the modern variety most closely associated with Original Pronunciation.

A direct approach to language attitudes measurement was chosen to investigate these research questions, with an online questionnaire containing both semantic differential rating scales and direct questions. Respondents were asked to rate four different recordings of Shakespeare’s *Sonnet 116*, recited in General American, Contemporary RP, Conservative RP, and Original Pronunciation. Responses to the questionnaire were obtained from hundred participants, fifty of which native speakers of English and fifty of which non-native speakers of English with a proficiency of at least C1 level according to the Common European Framework of Reference for Languages. These two groups effectively constituted two primary background variables in this study. Additional secondary background variables were introduced by asking respondents to indicate which English variety they natively speak (or use as a reference model), the degree of relevance that Shakespeare’s works hold for them, and whether they had obtained, or were in the process of obtaining, a tertiary degree in the subject of English.

The results of the study both confirm and rebuke the hypotheses set out for Original Pronunciation at the beginning. Reconstructed Elizabethan speech was indeed evaluated lower on *social status*—renamed *prestige* in the context of this study—and fairly high, though not highest, on *social attractiveness*—here renamed *accessibility*. It was also most closely associated with Irish English by a majority of the respondents. Original Pronunciation, however, was deemed less suitable for Shakespeare performances than the modern British standard varieties and, accordingly, ranked lower in preference. Most notably, it was very clearly evaluated as the least intelligible English variety for performances of Shakespeare, in a particularly stark rebuke of the claims brought forth by proponents of Original Pronunciation. The background variables that most significantly affected the results were the respondents' language proficiency and the English variety that they speak. As such, British English native speakers were the only subcategory of respondents to evaluate Original Pronunciation highest on *accessibility* and comparatively high on *suitability for Shakespeare*—substantially more so than General American. British English non-native speakers, for their part, clearly deemed Contemporary RP as the most suitable variety for Shakespeare performances. Additionally, a lack of personal relevance of Shakespeare for respondents and the absence of a tertiary academic background in the subject of English were found to have deviated responses among non-native speakers in favor of Contemporary RP on *prestige* and of General American on *accessibility*. The results of this study further indicate that Conservative RP is still regarded as the “standard” pronunciation for Shakespeare performances, as evidenced by the fact that this variety was evaluated highest on *prestige* and on *suitability for Shakespeare*. Nevertheless, evaluations of Contemporary RP were found to be overall not significantly lower than Conservative RP and, crucially, substantially higher along the evaluative dimension of *accessibility*. These results suggest that the slow erosion of Conservative RP observable on the contemporary British stage is provoked by changing mindsets among the theatergoing public regarding what constitutes “proper” pronunciation of Shakespeare.

Nevertheless, the conclusions that can be drawn from this empirical study are encouraging for proponents of staging Shakespeare's plays in reconstructed Elizabethan speech. Indeed, the high evaluations of Original Pronunciation along the *accessibility* dimension point to a considerable covert prestige that this speech variety enjoys among

British English native speakers as a medium for performances of Shakespeare. This consideration is strengthened by the result that General American—the English variety dominating popular media to a wide extent also in the United Kingdom—was evaluated significantly lower than Original Pronunciation along the *accessibility* dimension by British English native speaker. These results also partly validate the claims by proponents of Original Pronunciation regarding the viability of using reconstructed Elizabethan speech for contemporary Shakespeare performances as part of so-called “original practices” in the theater. Although the very low evaluations along the *prestige* dimension suggest that a widespread adoption of Original Pronunciation for Shakespeare performances may not take hold, the covert prestige attributed to it by participants in this study appears to validate the proselytizing efforts of the “OP movement,” whose desire it is to bring Shakespeare “closer” to the public.

In this regard, it is crucial to highlight once again that the participants in this study were not aware of the existence of performances of Shakespeare in reconstructed Elizabethan speech. As a result, most of them associated Original Pronunciation with either Irish English or Scottish English, while the rest of them indicated various areas of rural England as the likeliest provenance of the speech. As such, the social categorization and the stereotypes customarily applied to these language varieties were presumably also transferred upon Original Pronunciation, which is perhaps a severe limitation of this study. Further research on this subject should therefore take these aspects into consideration and both expand and improve on the research methodology employed for this thesis. Studies on attitudes toward Original Pronunciation using the matched guise technique could, for instance, be principled on a previous knowledge or awareness of the respondents regarding the use of reconstructed Elizabethan speech in performances of Shakespeare. In that sense, a mixed methodological approach that integrates qualitative data drawn from interviews with quantitative data obtained from questionnaires would arguably represent the best approach to this peculiar research topic. Finally, for comparative purposes, it would be particularly interesting to conduct separate investigations of attitudes toward Original Pronunciation performances of Shakespeare in different anglophone areas such as Britain, Ireland, and North America.

References

- Agheyisi, Rebecca; Fishman, Joshua A. 1970. "Language Attitude Studies: A Brief Survey of Methodological Approaches." *Anthropological Linguistics* 12(5), 137–157.
- Ajzen, Icek; Fishbein, Martin. 1980. *Understanding Attitudes and Predicting Social Behavior*. Englewood Cliffs: Prentice-Hall.
- Allport, Gordon W. 1935. "Attitudes." In Murchison, Carl (ed.). *A Handbook of Social Psychology* (Vol. 2). Worcester, MA: Clark University Press.
- Allport, Gordon W. 1954. "The Historical Background of Modern Social Psychology." In Lindzey, Gardner (ed.). *Handbook of Social Psychology* (Vol. 1: *Theory and Method*). Cambridge, MA: Addison-Wesley, 3–56.
- Barber, Charles. 1976. *Early Modern English*. London: Deutsch.
- Barrett, David. 2013. "Performing Shakespeare in the Original Pronunciation." PhD thesis, School of Music and Performance, University of South Wales.
- Brewster-Geisz, Zachary. 2006. *Sonnet 116 By William Shakespeare (1564–1616)*. [MP3 file] Retrieved from <https://librivox.org/sonnet-116-by-william-shakespeare> (12 May 2018).
- Cercignani, Fausto. 1981. *Shakespeare's Works and Elizabethan Pronunciation*. Oxford: Clarendon Press.
- Crystal, Ben; et. al. 2012. *Shakespeare's Original Pronunciation: Speeches and Scenes Performed as Shakespeare Would Have Heard Them*. [CD]. London: British Library Publishing.
- Crystal, David. 2005. *Pronouncing Shakespeare: The Globe Experiment*. Cambridge: Cambridge University Press.
- Crystal, David. 2013. "Early Interest in Shakespearean Original Pronunciation." *Language & History* 56(1), 5–17.
- Crystal, David. 2016. *The Oxford Dictionary of Original Shakespearean Pronunciation*. Oxford: Oxford University Press.
- Dalton-Puffer, Christiane; Kaltenböck, Gunther; Smit, Ute. 1997. "Learner Attitudes and Pronunciation in Austria." *World Englishes* 16(1), 115–128.
- Dobson, Eric J. 1957. *English Pronunciation: 1500–1700* (Vols. 1–3). Oxford: Clarendon Press.
- Dörnyei, Zoltan. 2007. *Research methods in applied linguistics: Quantitative, Qualitative, and Mixed Methodologies*. Oxford: Oxford University Press.

- Dragojevic, Marko. 2017. "Language Attitudes." *Oxford Research Encyclopedia of Communication*. Retrieved from <http://communication.oxfordre.com/view/10.1093/acrefore/9780190228613.001.0001/acrefore-9780190228613-e-437> (27 Aug. 2018).
- Edwards, John. 1982. "Language attitudes and their implication among English speakers." In Ryan, Ellen B.; Giles, Howard (eds.). *Attitudes towards Language Variation*. London: Edward Arnold.
- Edwards, John. 1994. *Multilingualism*. Abingdon: Routledge.
- Ellis, Alexander J. 1869–1874. *On Early English Pronunciation, With Especial Reference to Shakespeare and Chaucer. Part 3 [1871]: Illustrations of the Pronunciation of the XIVth and XVIth Centuries*. London: Philological Society.
- Garland, Ron. 1990. "A Comparison of Three Forms of the Semantic Differential." *Marketing Bulletin* 1(1), 19-24.
- Garrett, Peter. 2010. *Attitudes to Language*. Cambridge: Cambridge University Press.
- Garrett, Peter; Coupland, Nikolas; Williams, Angie. 2003. *Investigating Language Attitudes: Social Meanings of Dialect, Ethnicity and Performance*. Cardiff: University of Wales Press.
- Gibbons, Brian. 2011. "'He shifteth his speech': Accents and Dialects in Plays by Shakespeare and His Contemporaries." In: Jansohn, Christa; Orlin, Lena C.; Wells, Stanley. (eds.). *Shakespeare Without Boundaries: Essays in Honor of Dieter Mehl*. Newark: University of Delaware Press.
- Giles, Howard; Coupland, Nikolas; Henwood, Karen; Harriman, Jim; Coupland, Justine. 1990. "The Social Meaning of RP: an Intergenerational Perspective." In Ramsaran, Susan (ed.). *Studies in the Pronunciation of English: A Commemorative Volume in Honour of A. C. Gimson*. London: Routledge.
- Giles, Howard; Coupland, Nikolas. 1991. *Language: Context and consequences*. Milton Keynes: Oxford University Press.
- Giles, Howard; Billings, Andrew C. 2005. "Assessing Language Attitudes: Speaker Evaluation Studies. In Davies, Alan; Elder, Catherine (eds.). *The Handbook of Applied Linguistics*. Oxford: Blackwell.
- Gimson, A. C. 1962. *An Introduction to the Pronunciation of English*. London: Arnold.
- Henerson, Marlene E.; Morris, Lynn L.; Fitz-Gibbon, Carol Taylor. 1987. *How to Measure Attitudes*. Newbury Park: Sage.
- Hudson, Richard A. 1996. *Sociolinguistics* (2nd ed.). Cambridge: Cambridge University Press.

- International Phonetic Association. 2015. "Full IPA Chart." Retrieved from <http://www.internationalphoneticassociation.org/content/ipa-chart> (29 Jul 2018).
- Kökeritz, Helge. 1953. *Shakespeare's Pronunciation*. New Haven: Yale University Press.
- Knops, Uus; van Hout, Roeland. 1988. "Language Attitudes in the Dutch Language Area: an Introduction." In van Hout, Uus; Knops, Roeland (eds.). *Language Attitudes in the Dutch Language Area*. Dordrecht: Foris.
- Labov, William. 1966. *The Social Significance of Speech in New York City*. Washington, DC: Centre for Applied Linguistics.
- Lahr, John (2005). "Talking the Talk: The Globe goes Elizabethan." *The New Yorker*, 19 Sep., 98–99.
- Lambert, Wallace E.; Hodgson, Richard; Gardner, Robert; Fillenbaum, Samuel. 1960. "Evaluational Reactions to Spoken Languages." *Journal of Abnormal and Social Psychology* 60, 44–51.
- Lanier, Douglas. "Popular Culture." In Dobson, Michael; Wells, Stanley; Sharpe, Will; Sullivan, Erin (eds.). 2015. *The Oxford Companion to Shakespeare* (2nd ed.). Oxford: Oxford University Press, 439–440.
- Likert, Rensis. (1932). *A Technique for the Measurement of Attitudes*. New York: Columbia University Press.
- Lindemann, Stephanie. 2003. "Koreans, Chinese or Indians? Attitudes and ideologies about non-native English speakers in the United States." *Journal of Sociolinguistics* 7(3), 348–364.
- Lodewyck, Laura A. 2013. "Look with Thine Ears: Puns, Wordplay, and Original Pronunciation in Performance." *Shakespeare Bulletin* 31(1), 41–61.
- McKenzie, Robert M. 2008. "Social Factors and Non-Native Attitudes towards Varieties of Spoken English: a Japanese Case Study." *International Journal of Applied Linguistics* 18(1), 63–88.
- McKenzie, Robert M. 2010. *The Social Psychology of English as a Global Language: Attitudes, Awareness and Identity in the Japanese Context*. Dordrecht: Springer.
- Meier, Paul. 2011. "A Midsummer Night's Dream: An Original Pronunciation Production." *Voice and Speech Review* 7(1), 209–223.
- Miola, Robert S. 2000. *Shakespeare's Reading*. Oxford: Oxford University Press.
- Nevalainen, Terttu. 2006. *An Introduction to Early Modern English*. Edinburgh: Edinburgh University Press.

- Oppenheim, Bram. 1982. "An exercise in attitude measurement." In Breakwell, Glynis M.; Foot, Hugh C.; Gilmour Robin (eds.). *Social Psychology: a Practical Manual*. Basingstoke: Macmillan.
- Oppenheim, Abraham N. 1992. *Questionnaire Design, Interviewing, and Attitude Measurement*, London, Pinter.
- Osgood, Charles E.; Suci, George J.; Tannenbaum, Percy H. 1957. *The Measurement of Meaning*. Urbana: University of Illinois Press.
- Pear, Tom H. 1931. *Voice and personality*. London: Wiley.
- Peirce, Charles S.; Noyes, J. B. 1864. "Shakespearian Pronunciation." *North American Review* 98, 342–69.
- Pensalfini, Rob. 2009. "Not in Our Own Voices: Accent and Identity in Contemporary Australian Shakespeare Performance." *Australasian Drama Studies* 54, 142–158.
- Perloff, Richard M. 1993. *The Dynamics of Persuasion*. Hillsdale: Lawrence Erlbaum Associates.
- Ryan, Ellen B.; Giles, Howard; Hewstone, Miles. 2005. "The Measurement of Language Attitudes." In Ammon, Ulrich; Dittmar, Norbert; Mattheier, Klaus J.; Trudgill, Peter (eds.) *Sociolinguistics: an International Handbook of the Science of Language and Society* (Vol. 2). Berlin: de Gruyter.
- Sarnoff, Irving. 1970. "Social Attitudes and the Resolution of Motivational Conflict." In Jahoda, Marie; Warren, Neil (eds.). *Attitudes*. Harmondsworth: Penguin.
- Selwyn, Bertram. 2009. *Sonnet no 116: By William Shakespeare*. [Video file]. Retrieved from <https://www.youtube.com/watch?v=hdlZVZW9wGQ> (12 May 2018).
- Smith, Bruce. R. 1999. *The Acoustic World of Early Modern England: Attending to the O-Factor*. Chicago: University of Chicago Press.
- Tynan, Kenneth. 1961. *Curtains: Selections from the Drama Criticism and Related Writings*. New York: Atheneum.
- Weber, Ann L. 1992. *Social Psychology*. New York: Harper Collins.
- Wells, John C. 1982. *Accents of English* (Vols. 1–3). Cambridge University Press.
- Shakespeare, William. 2005. *All's Well That Ends Well*. In Wells, Stanley; Taylor, Gary; Jowett, John; Montgomery, William (eds.). *The Oxford Shakespeare: The Complete Works* (2nd ed.). Oxford: Oxford University Press, 1031–1058.
- Shakespeare, William. 2005. *As You Like It*. In Wells, Stanley; Taylor, Gary; Jowett, John; Montgomery, William (eds.). *The Oxford Shakespeare: The Complete Works* (2nd ed.). Oxford: Oxford University Press, 655–680.

- Shakespeare, William. 2005. *Hamlet*. In Wells, Stanley; Taylor, Gary; Jowett, John; Montgomery, William (eds.). *The Oxford Shakespeare: The Complete Works* (2nd ed.). Oxford: Oxford University Press, 681–718.
- Shakespeare, William. 2005. *King Lear: The Folio Text*. In Wells, Stanley; Taylor, Gary; Jowett, John; Montgomery, William (eds.). *The Oxford Shakespeare: The Complete Works* (2nd ed.). Oxford: Oxford University Press, 1153–1184.
- Shakespeare, William. 2005. *Macbeth*. In Wells, Stanley; Taylor, Gary; Jowett, John; Montgomery, William (eds.). *The Oxford Shakespeare: The Complete Works* (2nd ed.). Oxford: Oxford University Press, 969–994.
- Shakespeare, William. 2005. *A Midsummer Night's Dream*. In Wells, Stanley; Taylor, Gary; Jowett, John; Montgomery, William (eds.). *The Oxford Shakespeare: The Complete Works* (2nd ed.). Oxford: Oxford University Press, 401–424.
- Shakespeare, William. 2005. *Romeo and Juliet*. In Wells, Stanley; Taylor, Gary; Jowett, John; Montgomery, William (eds.). *The Oxford Shakespeare: The Complete Works* (2nd ed.). Oxford: Oxford University Press, 369–400.
- Shakespeare, William. 2005. *Sonnets and A Lover's Complaint*. In Wells, Stanley; Taylor, Gary; Jowett, John; Montgomery, William (eds.). *The Oxford Shakespeare: The Complete Works* (2nd ed.). Oxford: Oxford University Press, 777–804.
- Shakespeare, William. 2005. *The Two Gentlemen of Verona*. In Wells, Stanley; Taylor, Gary; Jowett, John; Montgomery, William (eds.). *The Oxford Shakespeare: The Complete Works* (2nd ed.). Oxford: Oxford University Press, 1–24.
- White, Richard G. 1865. "Memorandum on English Pronunciation in the Elizabethan Era." Appendix to Vol. 12 of *The Works of William Shakespeare*. Boston: Little Brown.
- Zahn, Christopher; Hopper, Robert. 1985. "Measuring Language Attitudes: The Speech Evaluation Instrument." *Journal of Language and Social Psychology* 4, 113–123.

Appendix

EVALUATION OF THEATRICAL AUDITIONS

In the following survey, you are going to hear four people who auditioned for a part in an upcoming theatre production. To assess their suitability for the role, we test the impression they make on potential audiences based on their performance of a sonnet by William Shakespeare. The text of the sonnet is provided below for you to read along with the recordings of the speakers.

You are given six options to answer the items in the questions. There are no right or wrong answers. Please tick the box which, in your opinion, best describes the person. In the following example, for instance, the listener thinks the speaker would be a ***rather*** good actor:

Do you think the **SPEAKER** would be a good actor?

←	VERY	RATHER	SOMEWHAT	SOMEWHAT	RATHER	VERY	→
good actor	☐	☒	☐	☐	☐	☐	bad actor

The survey is entirely **anonymous** and will be treated **confidentially**. Please answer the questions truthfully. We greatly value your participation!

TEXT OF THE RECORDING

Let me not to the marriage of true minds
Admit impediments. Love is not love
Which alters when it alteration finds,
Or bends with the remover to remove.
O no, it is an ever-fixèd mark
That looks on tempests and is never shaken;
It is the star to every wand'ring barque,
Whose worth's unknown, although his height be taken.
Love's not time's fool, though rosy lips and cheeks
Within his bending sickle's compass come;
Love alters not with his brief hours and weeks,
But bears it out even to the edge of doom.
If this be error and upon me proved,
I never writ, nor no man ever loved.

SPEAKER A

1) *SPEAKER A* sounds... (tick the appropriate boxes)

←	VERY	RATHER	SOMEWHAT	SOMEWHAT	RATHER	VERY	→
pleasant	☐	☐	☐	☐	☐	☐	unpleasant
snobbish	☐	☐	☐	☐	☐	☐	not snobbish
rural	☐	☐	☐	☐	☐	☐	urban
likeable	☐	☐	☐	☐	☐	☐	unlikeable
old-fashioned	☐	☐	☐	☐	☐	☐	modern
authentic	☐	☐	☐	☐	☐	☐	inauthentic
light-hearted	☐	☐	☐	☐	☐	☐	solemn
boring	☐	☐	☐	☐	☐	☐	interesting
sophisticated	☐	☐	☐	☐	☐	☐	unsophisticated

2) How clear and understandable was *SPEAKER A*'s pronunciation?

←	VERY	RATHER	SOMEWHAT	SOMEWHAT	RATHER	VERY	→
unclear / not understandable	☐	☐	☐	☐	☐	☐	clear / understandable

3) On a scale from 1 (*not at all*) to 10 (*very well*), how well would you have understood *SPEAKER A* without reading the text?

→	1	2	3	4	5	6	7	8	9	10	→
not at all	☐	☐	☐	☐	☐	☐	☐	☐	☐	☐	very well

4) How suitable do you think *SPEAKER A* would be for a role in a Shakespeare production?

←	VERY	RATHER	SOMEWHAT	SOMEWHAT	RATHER	VERY	→
not suitable	☐	☐	☐	☐	☐	☐	suitable

5) Where do you think the *SPEAKER* comes from?

SPEAKER A comes from:

SPEAKER B

1) *SPEAKER B* sounds... (tick the appropriate boxes)

←	VERY	RATHER	SOMEWHAT	SOMEWHAT	RATHER	VERY	→
pleasant	☐	☐	☐	☐	☐	☐	unpleasant
snobbish	☐	☐	☐	☐	☐	☐	not snobbish
rural	☐	☐	☐	☐	☐	☐	urban
likeable	☐	☐	☐	☐	☐	☐	unlikeable
old-fashioned	☐	☐	☐	☐	☐	☐	modern
authentic	☐	☐	☐	☐	☐	☐	inauthentic
light-hearted	☐	☐	☐	☐	☐	☐	solemn
boring	☐	☐	☐	☐	☐	☐	interesting
sophisticated	☐	☐	☐	☐	☐	☐	unsophisticated

2) How clear and understandable was *SPEAKER B*'s pronunciation?

←	VERY	RATHER	SOMEWHAT	SOMEWHAT	RATHER	VERY	→
unclear / not understandable	☐	☐	☐	☐	☐	☐	clear / understandable

3) On a scale from 1 (*not at all*) to 10 (*very well*), how well would you have understood *SPEAKER B* without reading the text?

→	1	2	3	4	5	6	7	8	9	10	→
not at all	☐	☐	☐	☐	☐	☐	☐	☐	☐	☐	very well

4) How suitable do you think *SPEAKER B* would be for a role in a Shakespeare production?

←	VERY	RATHER	SOMEWHAT	SOMEWHAT	RATHER	VERY	→
not suitable	☐	☐	☐	☐	☐	☐	suitable

5) Where do you think the *SPEAKER* comes from?

SPEAKER B comes from:

SPEAKER C

1) *SPEAKER C* sounds... (tick the appropriate boxes)

←	VERY	RATHER	SOMEWHAT	SOMEWHAT	RATHER	VERY	→
pleasant	☐	☐	☐	☐	☐	☐	unpleasant
snobbish	☐	☐	☐	☐	☐	☐	not snobbish
rural	☐	☐	☐	☐	☐	☐	urban
likeable	☐	☐	☐	☐	☐	☐	unlikeable
old-fashioned	☐	☐	☐	☐	☐	☐	modern
authentic	☐	☐	☐	☐	☐	☐	inauthentic
light-hearted	☐	☐	☐	☐	☐	☐	solemn
boring	☐	☐	☐	☐	☐	☐	interesting
sophisticated	☐	☐	☐	☐	☐	☐	unsophisticated

2) How clear and understandable was *SPEAKER C*'s pronunciation?

←	VERY	RATHER	SOMEWHAT	SOMEWHAT	RATHER	VERY	→
unclear / not understandable	☐	☐	☐	☐	☐	☐	clear / understandable

3) On a scale from 1 (*not at all*) to 10 (*very well*), how well would you have understood *SPEAKER C* without reading the text?

→	1	2	3	4	5	6	7	8	9	10	→
not at all	☐	☐	☐	☐	☐	☐	☐	☐	☐	☐	very well

4) How suitable do you think *SPEAKER C* would be for a role in a Shakespeare production?

←	VERY	RATHER	SOMEWHAT	SOMEWHAT	RATHER	VERY	→
not suitable	☐	☐	☐	☐	☐	☐	suitable

5) Where do you think the *SPEAKER* comes from?

SPEAKER C comes from:

SPEAKER D

1) *SPEAKER D* sounds... (tick the appropriate boxes)

←	VERY	RATHER	SOMEWHAT	SOMEWHAT	RATHER	VERY	→
pleasant	☐	☐	☐	☐	☐	☐	unpleasant
snobbish	☐	☐	☐	☐	☐	☐	not snobbish
rural	☐	☐	☐	☐	☐	☐	urban
likeable	☐	☐	☐	☐	☐	☐	unlikeable
old-fashioned	☐	☐	☐	☐	☐	☐	modern
authentic	☐	☐	☐	☐	☐	☐	inauthentic
light-hearted	☐	☐	☐	☐	☐	☐	solemn
boring	☐	☐	☐	☐	☐	☐	interesting
sophisticated	☐	☐	☐	☐	☐	☐	unsophisticated

2) How clear and understandable was *SPEAKER D*'s pronunciation?

←	VERY	RATHER	SOMEWHAT	SOMEWHAT	RATHER	VERY	→
unclear / not understandable	☐	☐	☐	☐	☐	☐	clear / understandable

3) On a scale from 1 (*not at all*) to 10 (*very well*), how well would you have understood *SPEAKER D* without reading the text?

→	1	2	3	4	5	6	7	8	9	10	→
not at all	☐	☐	☐	☐	☐	☐	☐	☐	☐	☐	very well

4) How suitable do you think *SPEAKER D* would be for a role in a Shakespeare production?

←	VERY	RATHER	SOMEWHAT	SOMEWHAT	RATHER	VERY	→
not suitable	☐	☐	☐	☐	☐	☐	suitable

5) Where do you think the *SPEAKER* comes from?

SPEAKER D comes from:

We're almost done! Please recall your overall impressions of the four speakers one more time and try to rank them according to your preference. Give the highest rank (1st) to the speaker you liked best and the lowest rank (4th) to the one you liked least.

1. _____
2. _____
3. _____
4. _____

In conclusion: how would you describe the relevance of Shakespeare's works for you personally?

←	VERY	RATHER	SOMEWHAT	SOMEWHAT	RATHER	VERY	→
not relevant	☐	☐	☐	☐	☐	☐	relevant

Finally, please fill out some basic information about yourself.

Age: _____

Please indicate your level of proficiency in English:

**NATIVE
SPEAKER**
☐

**MASTERY
(C2)***
☐

**ADVANCED
(C1)***
☐

**UPPER INTERMEDIATE
(B2)***
☐

* according to the Common European Framework of Reference for Languages

Which variety of English do you speak (or use as a model, if you're a non-native speaker)?

**BRITISH
ENGLISH**
☐

**AMERICAN
ENGLISH**
☐

**SCOTTISH
ENGLISH**
☐

**IRISH
ENGLISH**
☐

**OTHER
(please specify)**
☐

**NO MODEL
(non-native speaker)**
☐

Other variety: _____

Are you pursuing or have completed an English degree? (please indicate the highest level reached)

No ☐	Yes, I am pursuing a master's degree ☐
Yes, I am pursuing a bachelor's degree ☐	Yes, I have obtained a master's degree ☐
Yes, I have obtained a bachelor's degree ☐	Yes, I am pursuing or have obtained a PhD ☐

Thank you for your time and participation!

Zusammenfassung in deutscher Sprache

Die vorliegende Masterarbeit untersucht Spracheinstellungen gegenüber Darbietungen von Shakespeares Werken in rekonstruierter frühneuenglischer Aussprache. Befürworter/innen dieses als „*Original Pronunciation*“ bekannten Bühnenakzents sind der Meinung, dass Shakespeare einem breiten Publikum dadurch „näher“ gebracht werden könne als dies durch die Standardaussprache der britischen Elite – der sogenannten *Received Pronunciation* (RP) – möglich sei. Die Studie wurde an einer Population von hundert Teilnehmer/innen durchgeführt, die zur Hälfte aus Personen englischer Muttersprache besteht und alle Altersgruppen umfasst. Zur quantitativen Datenerhebung wurde die von Lambert entwickelte *Matched Guise*-Methode in Form eines Online-Fragebogens eingesetzt. Die Respondent/innen wurden gebeten, vier verschiedene Aufnahmen von Shakespeares *Sonnet 116* einzustufen, in denen der Text in jeweils unterschiedlicher Aussprache vorgetragen wird (Standardamerikanisch, gegenwärtige RP, konservative RP, sowie *Original Pronunciation*). Die Beurteilung erfolgte anhand bipolarer Adjektivpaare des semantischen Differentials und in Beantwortung gezielter Fragen. Als Hauptvariable fungierte die Muttersprache der Respondent/innen; die gesprochene Englischvarietät und der jeweilige Bildungshintergrund bildeten weitere Variablen. Die Ergebnisse der Datenauswertung lassen den Optimismus der *Original Pronunciation* Befürworter/innen als nicht unbegründet erscheinen: Obwohl rekonstruiertes Frühneuenglisch deutlich niedrigere Evaluierungswerte in der charakterbezogenen Prestigedimension erhielt und insgesamt als für Shakespeare-Darbietungen am wenigsten geeignet eingestuft wurde, weisen die hohen Bewertungen dieser Aussprache in der charakterbezogenen Solidaritätsdimension auf ein „verborgenes“ Prestige hin, das Nonstandardvarietäten laut Spracheinstellungsstudien häufig genießen. Dies ist besonders daran ersichtlich, dass britische Respondent/innen *Original Pronunciation* dezidiert höher einstufen als andere Untergruppen der Studienpopulation und gleichzeitig den rekonstruierten Bühnenakzent mit verschiedensten regionalen Englischvarietäten assoziierten. Das „verborgene“ Prestige der *Original Pronunciation* kann somit, zumindest aus Sicht eines muttersprachlich britischen Publikums, als sicheres Indiz für das Bühnenpotential frühneuenglischer Aussprache auch in der Gegenwart interpretiert werden.