

GENDER DISGUISE AND LINGUISTIC IDENTITY PERFORMANCE IN ONLINE
WRITINGS:

Production, Perception, and Forensic Applications

JULIANE FORD

Doctor of Philosophy

ASTON UNIVERSITY

August 2021

© Juliane Ford, 2021

Juliane Ford asserts her moral right to be identified as the author of this thesis

This copy of the thesis has been supplied on condition that anyone who consults is understood to recognize that its copyright belongs to its author and that no quotation from this thesis and no information derived from it may be published without appropriate permission or acknowledgement.

Table of Contents

Table of Contents.....	2
Table of Figures	5
Chapter 1 – Introduction.....	12
Part 1 – Theories of Language and Identity	13
Part 2 – Identity Disguise	15
A. Types of Identity Guising.....	16
B. Examples of Identity Play	18
C. Overview	22
Part 3 – Overview of Chapters	22
Chapter 2 – Literature Review	25
Part 1 – Introduction	25
Part 2 – Approaches to the Study of Gendered Language	25
A. Earlier Approaches.....	25
B. Contemporary Approaches	28
C. Overview	33
Part 3 – Qualitative Approaches	34
A. Proposed Features of Gendered Language	34
B. Overview	46
Part 4 – Computational Approaches	47
A. Proposed features of Gendered language.....	48
B. Overview	53
Part 5 – Discussion/Overview	54
Chapter 3 – Methodology	57
Part 1 – Introduction	57
Part 2 – Variables	58
A. The edited Blog Authorship Corpus (eBAC)	59
B. The Features.....	60
C. Overview	79
Part 3 – Linguistic Theories	81
A. Audience Design	82
B. Accommodation Theory	86
C. Community of Practice	89
D. Speaker Contamination.....	94
Part 4 – Discussion/Conclusion	98
Chapter 4 – Non-Forensic Analyses.....	100
Part 1 – Introduction	100
Part 2 – The AGGID Blog	101

A. The Data	101
B. The Analysis	102
C. Keywords	106
D. Overview	108
Part 3 – Tinder	110
A. The Data	110
B. The Analysis	113
C. Downward Trending 'Female Coded' Features	113
D. Upward-trending 'Male Coded' Features	117
E. Mixed or Unclear Features	121
F. Overview	128
Part 4 – Tinder 2	129
A. Conversational Types	130
B. Conversational Trajectories.....	153
C. Conversational Transitions	157
D. Overview	176
Part 5 – Other Instances	177
A. Suspicion on Tinder	177
B. Suspicion Elsewhere	182
Part 6 – Discussion/Overview	183
Chapter 5 – Forensic Analyses	186
Part 1 – Introduction	186
Part 2 – Hemmert	186
A. The Data	187
B. Analysis of Non-Q Features	188
C. Overview of Non-Q Features.....	190
D. Analysis of Q Features	190
E. Overview of Q Features	194
F. Overview	195
Part 3 – Hemmert Feature Weights	195
Part 4 – Hemmert 2	202
A. The Data	203
B. The Analysis	203
C. Keyword Analysis.....	207
D. Overview	208
Part 5 – Potter	208
A. The Data	210
B. The Analysis	211
C. Overview	221
Part 6 – Overall Feature Weights.....	224
A. A comparison to the eBAC	225

B. Comparison to the Known Authors' eBAC Findings.....	227
C. Relational comparison to the known authors	230
D. Determining the overall validity of individual features & the application of context	233
E. Features indicative of gender guising and the problem of audience	235
Part 7 – Discussion/Overview	237
Chapter 6 – Discussion/Conclusion	240
Appendices	247
References.....	248

Table of Figures

Table 1.1 Ashley’s “Angels” chat bot messages (Newitz, 2015).....	19
Table 2.1 Gendered features as suggested by Lakoff, 1975.....	35
Table 2.2 Overview of Lakoff 1975 proposed features and additional research	47
Table 2.3 Gendered features as suggested by Bamman et al. (2014) & Schler et al. (2006).....	49
Table 2.4 Overview of 20 proposed features and additional research	54
Table 3.1 Overall distribution of the Original and Edited BAC	59
Table 3.2 Overall distribution of the eBAC by male and female subcorpora	60
Table 3.3 Overall pronoun distribution by person	60
Table 3.4 Overall distribution of pronouns and their alternative spellings.....	61
Table 3.5 Distribution of pronoun alternative spellings.....	61
Table 3.6 Overall distribution of emotion terms.....	62
Figure 3.1 Unique emotion terms	62
Table 3.7 Examples of non-standard, innovative emoticon use as found in the eBAC.....	62
Table 3.8 Overall distribution of female-, male-, and other-coded kinship terms	63
Table 3.9 Overall distribution of female- and male-coded kinship terms	64
Table 3.10 Distribution of female- and male-coded friendship terms	64
Table 3.11 Overall distribution of female-, male-, and other-coded friendship terms.....	64
Table 3.12 Overall distribution of abbreviations	65
Figure 3.2 Non-lol unique abbreviation variations	65
Table 3.13 Overall distribution of punctuation.....	66
Table 3.14 Overall distribution of punctuation per unit	66
Table 3.15 Distribution of standard/unlengthened and lengthened/complex punctuation	66
Table 3.16 Expressive lengthening overall	67
Table 3.17 Lengthened and unlengthened backchannel sounds	67
Table 3.18 Backchannel sound variations	68
Figure 3.3 Overall distribution and differences of normed backchannel sound frequencies ..	68
Figure 3.4 Use of nyaa (blue) and squee (red) over time via Google Trends	69
Table 3.19 Overall distribution of backchannel sounds	69
Table 3.20 Unique expressively lengthened backchannel sounds	69
Table 3.21 Overall hesitation words	69
Table 3.22 Overall distribution of assent term varieties	70
Table 3.23 Overall distribution of negation terms.....	70
Table 3.24 Distribution of female- and male-coded negation terms	71
Table 3.25 Distribution of swears and anti-swears.....	71
Table 3.26 Overall swear and anti-swear word variations.....	72
Table 3.27 f and non-f abbreviations	72
Table 3.28 Overall distribution of preposition varieties and alternative spellings.....	73

Table 3.29 Overall distribution of preposition varieties and alternative spellings.....	74
Table 3.30 Alternative spellings.....	74
Table 3.31 Distribution of conjunctions and their alternative spellings	74
Table 3.32 Articles and determiners	75
Table 3.33 Overall distribution of the eBAC by male and female subcorpora.....	75
Table 3.34 urlLink uses	75
Table 3.35 Total keywords and type/token ratio	76
Table 3.36 Top 10 overall keywords.....	76
Table 3.37 Top 10 function keywords.....	77
Table 3.38 Conjunction keyness in the top 100 keywords	77
Table 3.39 Top 10 content keywords.....	77
Table 3.40 Top 10 proper noun content words	78
Table 3.41 Top 100 keywords that are function words.....	78
Table 3.42 Content words from the top 100 keywords.....	79
Table 3.43 eBAC features in agreement with Bamman et al. (2014) & Schler et al. (2006) ..	80
Table 3.44 eBAC features not in agreement w/ Bamman et al.(2014) & Schler et al.(2006) .	81
Figure 3.5 Tumblr posts between anonymous	85
Figure 3.6 “Funny Texts from Mom”	89
Table 4.1 Dataset distribution of Amina and Tom in AGGID	101
Table 4.2 Distribution of features in which Amina better matches the Male eBAC	103
Table 4.3 Distribution of features in which Amina better matches the Female eBAC	104
Table 4.4 Distribution of features in which Amina doesn't clearly match either subcorpus ..	106
Table 4.5 Total words with a keyness above 6.53	106
Table 4.6 Top 10 keywords between Amina and the eBAC datasets.....	107
Table 4.7 Distribution of political and non-political keywords between Amina & the eBAC..	107
Table 4.8 Non-political keywords from the top 100 for Amina and the eBAC datasets	107
Table 4.9 Total negative keywords between Amina and the eBAC datasets	108
Table 4.10 Top 10 negative keywords between Amina and the eBAC datasets	108
Table 4.11 Feature findings overview	109
Figure 4.1 Catfi.sh demonstration of third-party catfishing	110
Figure 4.2 Example excerpt of the Before phase	111
Figure 4.3 Example excerpt transitioning from During to After phases.....	112
Figure 4.4 Example excerpt of the After phase	112
Table 4.12 Distribution of word counts in the Tinder data phases.....	112
Table 4.13 Overall raw and normed uses of 2nd person pronouns and variants.....	113
Table 4.14 Overall distribution of pronoun use throughout the phases	113
Figure 4.5 Distributions of emotion terms by phase	114
Figure 4.6 Examples of emotion terms as used across phases	115
Table 4.15 Total instances and variety of emotion terms used in the data	115

Table 4.16 Overall kinship term counts.....	116
Figure 4.7 Distribution of expressive lengthening variety by phase.....	116
Table 4.17 Overall expressive lengthening counts	116
Figure 4.8 Distribution of backchannel sound variety by phase	116
Table 4.18 Raw and normed backchannel sound counts by phase	117
Figure 4.9 Distribution of hesitation word variety by phase	117
Table 4.19 Raw and normed hesitation word counts by phase.....	117
Figure 4.10 Friendship terms.....	118
Table 4.20 Overall friendship/“bro” term counts by phase.....	118
Figure 4.11 Examples of friendship terms in the data	119
Table 4.21 Overall negation term counts, female and male coded	120
Table 4.22 Overall swear and anti-swear counts (raw)	120
Figure 4.12 Examples of “fuck” words in the data.....	121
Table 4.23 Overall use of determiners and articles throughout the phases.....	121
Table 4.24 Ratio of use of determiners and articles throughout the phases.....	121
Table 4.25 Most common emojis on Tinder and in the data (overall % of frequency)	122
Table 4.26 Overall counts of words per emoji, and types of emoji, by phase	123
Table 4.27 Most frequent emoji by phase	123
Figure 4.13 Examples of the most common emoji uses in the data	124
Figure 4.14 Distribution of abbreviation variety by phase.....	124
Table 4.28 Raw and normed abbreviation counts by phase	124
Table 4.29 Overall distribution of punctuation.....	125
Table 4.30 Overall distribution of punctuation per unit	126
Table 4.31 Distribution of standard/unlengthened and lengthened/complex punctuation....	126
Table 4.32 Overall assent term counts	126
Figure 4.15 Assent terms	127
Table 4.33 Overall preposition use throughout the phases	127
Table 4.34 Normed and raw frequencies of “and” and “&” throughout the phases	128
Table 4.35 Average post length throughout the three phases.....	128
Figure 4.16 Conversational trajectory types.....	130
Figure 4.17 Type A conversations	131
Figure 4.18 James and Tom (39)	132
Figure 4.19 Daniel and Joe (4)	133
Figure 4.20 Type B conversations	133
Figure 4.21 Iknoor and Bradley (53)	134
Figure 4.22 Will and Ata (54)	135
Figure 4.23 Type C conversations.....	135
Figure 4.24 Jonny and Joshua (13)	137
Figure 4.25 Tam and Sol (7) excerpt.....	138

Figure 4.26 Type D conversations.....	138
Figure 4.27 Tom and James (38)	139
Figure 4.28 Nat and Sixten (17) excerpt.....	139
Figure 4.29 Type E conversations	139
Figure 4.30 Jesse and Zäimon (48).....	141
Figure 4.31 Gerard and Sam (27) excerpt	142
Figure 4.32 Type F conversations	143
Figure 4.33 Alassane and Jamie (30).....	143
Figure 4.34 Will(A) and Will(B) (24)	144
Figure 4.35 Type G conversations.....	145
Figure 4.36 James and Andrew (26)	146
Figure 4.37 Deji and Joe (20) excerpt	147
Figure 4.38 Type H conversations.....	147
Figure 4.39 Steven and Corey (15)	148
Figure 4.40 David and Adam (37) excerpt.....	149
Figure 4.41 Type I conversations	149
Figure 4.42 Adam and Mark (52).....	149
Figure 4.43 Type J conversations.....	150
Figure 4.44 Armand and Kasim (21).....	152
Figure 4.45 Type K conversations	152
Figure 4.46 Kane and Jitesh (23) excerpt.....	153
Figure 4.47 Conversational trajectories	153
Figure 4.48 Conversational types without During phases	154
Figure 4.49 David and Adam (37) excerpt.....	154
Figure 4.50 Steven and Corey (15) excerpt.....	155
Figure 4.51 Conversational types without After phases	155
Figure 4.52 Nat and Sixten (17) excerpt.....	155
Figure 4.53 Morgan and Callum (26) excerpt	156
Figure 4.54 Tom and James (38)	157
Table 4.36 Conversational flag types	157
Figure 4.55 Type C conversations.....	158
Figure 4.56 Martyn and William (28) flags trajectory 1	158
Figure 4.57 Martyn and William (28) flags excerpt 1.....	159
Figure 4.58 Martyn and William (28) flags excerpt 2.....	159
Figure 4.59 Martyn and William (28) flags trajectory 2.....	160
Figure 4.60 Martyn and William (28) flags excerpt 2.....	160
Figure 4.61 Martyn and William (28) flags excerpt 3.....	161
Figure 4.62 James and Tom (39) flags trajectory	161
Figure 4.63 James and Tom (39) flags.....	162

Figure 4.64 David and Adam (37) topical progression excerpt	163
Figure 4.65 David and Adam (37) flags trajectory.....	163
Figure 4.66 David and Adam (37) flags excerpt.....	163
Figure 4.67 Type C conversation.....	164
Figure 4.68 Chris and Conor (5) flags trajectory	164
Figure 4.69 Chris and Conor (5) flags.....	166
Figure 4.70 Will and Dan (1) flags trajectory.....	166
Figure 4.71 Will and Dan (1) flags excerpt.....	166
Figure 4.72 Olly and Jack (14) flags trajectory.....	167
Figure 4.73 Olly and Jack (14) flags excerpt	167
Figure 4.74 Daniel and Joe (4) flags trajectory	167
Figure 4.75 Daniel and Joe (4) flags excerpt	167
Figure 4.76 Deji and Joe (20) excerpt	169
Figure 4.77 Martyn and William (28) excerpt	169
Figure 4.78 Deji and Joe (20) excerpt	170
Figure 4.79 Gerard and Sam (27) excerpt	171
Figure 4.80 Armand and Kasim (21) excerpt	172
Figure 4.81 Tom and James (38) excerpt.....	173
Figure 4.82 John and Jacob (2) excerpt	176
Figure 4.83 Adam and Mark (52).....	178
Figure 4.84 Tinder bot (3) Figure 4.85 Tinder bot (4)	179
Figure 4.86 Tinder bot (5) Figure 4.87 Tinder bot (6)	180
Figure 4.88 Tinder bot (1) Figure 4.89 Tinder bot (2)	181
Table 5.1 Distribution of word counts in Hemmert	187
Table 5.2 Distribution of female- and male-coded friendship terms	188
Table 5.3 Overall abbreviations distribution Figure 5.1 Unique abbreviation variations...	188
Table 5.4 Overall distribution of punctuation.....	188
Table 5.5 Overall backchannel sound distribution Figure 5.2 Unique backchannel sounds	189
Table 5.6 Overall hesitation word distribution Figure 5.3 Unique hesitation words	189
Table 5.7 Overall swear and anti-swear word variations.....	189
Table 5.8 Overall alternative spellings distribution	190
Table 5.9 Overall feature coding compared to Bamman et al.(2014) & Schler et al.(2006).	190
Figure 5.4 Distribution of features found in the Q	191
Table 5.10 Overall Pronoun use by person	191
Table 5.11 Overall distribution of pronouns and their alternative spellings.....	192
Table 5.12 Overall distribution of emotion terms.....	192
Table 5.13 Overall distribution of female- and male-coded kinship terms	192
Table 5.14 Overall distribution of assent term varieties	192
Table 5.15 Overall distribution of negation terms.....	193

Table 5.16 Overall distribution of preposition varieties and alternative spellings	193
Table 5.17 Distribution of conjunctions and their alternative spellings	193
Table 5.18 Articles and determiners	194
Table 5.19 Distribution of word counts in Hemmert	194
Table 5.20 Overall coding of features as compared to the eBAC.....	194
Table 5.21 Binary distinction of only the 8 features present in Q	196
Table 5.22 Binary distinction of all 16 features, including those absent in all of Q	197
Table 5.23 Distribution of all 16 binary features in Q text 16, including those absent.....	197
Table 5.24 Distribution of present features in each text only, and % of total gendering	198
Table 5.25 Overall non-skewed findings (binary in Q and % individual present w/ context)	199
Table 5.26 Q texts with mixed results across non-skewed tests	199
Table 5.27 Q texts with assent and negation terms	200
Table 5.28 Relative distribution of K Amber & K Lael compared to binary distribution in Q.	201
Table 5.29 Q Texts found to match K Amber's patterns	202
Table 5.30 Distribution of K Lael data into new data subsets.....	203
Table 5.31 Overall feature findings in "To" datasets	204
Table 5.32 Comparison of "To" features to K Lael and each other.....	205
Table 5.33 Binary distribution as compared to the eBAC.....	206
Table 5.34 Comparison of "To" datasets to Tinder Before/After phases	206
Table 5.35 Data distribution.....	211
Table 5.36 Pronouns overview in Potter	211
Table 5.37 Distribution of gendered pronouns in Potter	212
Table 5.38 Overview of emotion terms in Potter	212
Figure 5.5 Shared and unique emotion terms in Potter.....	213
Table 5.39 Distribution of kinship terms in Potter.....	213
Table 5.40 Overall kinship term use in Potter	214
Table 5.41 Kinship term comparison in Potter and the eBAC	214
Table 5.42 Overall friendship terms in Potter.....	214
Table 5.43 Overall abbreviations in Potter	215
Figure 5.6 Unique and shared abbreviations in Potter	215
Table 5.44 Overall punctuation in Potter.....	216
Table 5.45 Words per punctuation in Potter	216
Table 5.46 Standard and lengthened punctuation in Potter	216
Table 5.47 Overall expressive lengthening in Potter.....	217
Table 5.48 Overall backchannel sounds in Potter.....	217
Figure 5.7 Shared and unique backchannel sounds in Potter.....	218
Table 5.49 Overall hesitation words in Potter	218
Table 5.50 Overall assent terms in Potter.....	219
Table 5.51 Overall negation terms in Potter	219

Table 5.52 Swears and taboo words in Potter	220
Table 5.53 Overall prepositions in Potter	220
Table 5.54 Overall alternative spellings in Potter	221
Table 5.55 Overall conjunctions in Potter	221
Table 5.56 Overall articles and determiners in Potter	221
Table 5.57 Overall features in Potter as compared to the eBAC and Jenelle	222
Table 5.58 Overall comparison of Hemmert and Potter to the eBAC	226
Table 5.59 eBAC comparison features in Hemmert that did not match Known Lael	227
Table 5.60 eBAC comparison features in Potter that did not match Known Jenelle	228
Table 5.61 Distribution of friend in Potter and the eBAC	229
Figure 5.8 Distribution of features found in the Q	231
Table 5.62 Overall feature distribution in Potter	232
Table 5.63 Total preposition use in Potter	232
Table 5.64 eBAC feature distribution in the four analyses	234
Table 5.65 Overall findings of all four analyses	236

Chapter 1 – Introduction

When we interact online, we lose a lot of the usual information most other communicative interactions provide, such as the ability to see the person with whom we are speaking or hear their voice. Interlocutors receive a potentially curated conversation from other interactants and miss a variety of identity cues we as communicators tend to take for granted during in-person and even audio-only interactions. While a lack of identifying information can lead to miscommunications within the realm of the conversation itself and the information being conveyed, this can also lead to incorrect assumptions about just who it is we are speaking—or, as is most often the case online, typing—to.

This paper examines multiple such situations, in which one or more interactant's identity is disguised, intentionally or not, specifically along the lines of gender. This paper aims to answer research questions that address both the performance and perception of interlocutors in such gender-guised interactions, with particular attention paid to any variables that can provide useful authorship information in analyses of such datasets. Thus this paper asks three major research questions, discussed here.

First, what strategies do people use when performing gender, disguised or otherwise, intentionally or otherwise? This paper considers not only disguised gender performance, but also overt gender performance that may not be disguised, as both of these performance types intersect in that they are intentional. To answer this question, this paper further considers not only individual linguistic variables useful for authorship analyses, but also situates these variables in their appropriate contexts given a set of data within a variety of linguistic theories that can variably affect any performance and its perception.

Second, what are the 'tells' that give away gender as overtly performative and possibly disguised, and how perceptive are interactants to such tells? To answer this question, this paper considers not only what of the analyzed variables show particular gain for indicating gender performance, but also the observations of those interacting with such performances in real time, and how the two do or do not align. Each outcome is considered alongside such interactional frameworks as the difference between cooperation and suspicion as conversational commodities for a given exchange type in order to explain why the same tell may have more or less consequence to a given interaction.

Third, given the first two questions, can a given set of variables applied for such analyses which are derived from experimental research be successfully applied to real (specifically forensic) data in order to determine gender or gender guising? To answer this question, this paper relies on a framework of variables previously attest to have demonstrated statistical gain for gender identification on large sets of data, and applies them in a more quantitative approach to smaller and eventually forensic datasets. The outcome of the application of these variables is then paired with a more qualitative approach that considers the output of the variables as appropriately situated in the specific context of a given dataset.

In seeking to answer these questions, this paper seeks to demonstrate a model whereby quantitative variables can be responsibly paired with qualitative analyses on smaller datasets with the aim of still providing useful answers to forensic questions. In doing so, this paper will demonstrate that such analyses can be successful in indicating gender performance when appropriate baselines can be established when they are appropriately situated in the proper context of the dataset. Further, this paper will demonstrate that interactional gender confusion is just as perceptive as it is performative, and performance can be amplified in both conscious and unconscious ways.

Part 1 – Theories of Language and Identity

Sociolinguistic research has historically been categorized as being comprised of three overarching waves of analytical practice, which shifted from ideas of macro-level categories of identity characterizing language to more localized and eventually contextually based approaches. Within these three waves, the focus of sociolinguistic research has similarly shifted from more specific, essentialist identity categories to a more constructionist approach, building upon ideas of the ‘authenticity’ of these facets and those who access them.

The first wave, based on the foundational works of Labov in the 1960s and 70s, “introduced more quantitative empiricism into linguistics, with supportive theoretical underpinnings. [...] subsequent studies came to focus on filling cells defined by macrosociological categories. In this way, speakers emerged as human tokens—bundles of demographic characteristics,” (Eckert, 2012). Along with other fields of social science at this time, these first two waves relied on the idea of ‘authenticity’ in constructing their categories. That is:

The idea of authenticity gains its force from essentialism, for the possibility of ‘real’ or ‘genuine’ group members relies on the belief that what differentiates ‘real’ members from those who only pretend to authentic membership is that the former, by virtue of biology or culture or both, possess inherent and perhaps even inalienable characteristics criterial of membership. (Bucholtz, 2003).

Beyond this, the second wave of sociolinguistic analytical practice focused on more localized expressions of these identity categories, and, “began with the attribution of social agency to the use of vernacular as well as standard features and a focus on the vernacular as an expression of local or class identity,” (Eckert, 2012). Research conducted throughout the following decades served to “make it clear that linguistic variables do not index categories, but characteristics, giving an entirely new theoretical underpinning and methodological thrust to the variation enterprise in the third wave,” (Eckert, 2012). The second wave approach solidified the idea of authenticity by building on the first wave’s assumption of, “vernacular as having local value” (Eckert, 2012).

While both the first and second waves focused on more essentialist and “apparently static categories of speakers and equated identity with category affiliation,” the third wave “was from a view of variation as a reflection of social identities and categories to the linguistic

practice in which speakers place themselves in the social landscape through stylistic practice,” (Eckert, 2012). As such, the third wave has moved away from essentialist categories which relied on the idea that, “the attributes and behavior of socially defined groups can be determined and explained by reference to cultural and/or biological characteristics believed to be inherent to the group. As an ideology, essentialism rests on two assumptions: (1) that groups can be clearly delimited; and (2) that group members are more or less alike,” (Bucholtz, 2003) In moving instead toward a constructionist approach:

Rather than regard social categories, such as social class, ethnicity, and gender as primary, with linguistic behavior reflecting membership, a social constructionist approach shifts the emphasis to language as a dynamic resource used to construct particular aspects of social identity at different points in an interaction. (Holmes, 2001)

Within this third wave approach in more recent decades, which focused more on style, “Social categories are not fixed but are subject to constant change; talk itself actively creates different styles and constructs different social contexts and social identities as it proceeds,” (Holmes, 2001).

More modern computational approaches to largescale language study must generally rely on essentialist or at least macro-level categories, but then situate them in their appropriate interactional or at least local and micro-level contexts. This may occur by focusing on a gender binary, for example, when analyzing a compilation of hundreds of users’ tweets, but also considering the gender makeup of individuals’ social networks in their adherence or nonconformity to the generalizations provided by the macro-level categories to which they belong, as we will see with Bamman Eisenstein, & Schnobelen (2014) as is discussed throughout this paper. Essentially, while, “More essentialist approaches to differences tend to highlight differences between groups and treat groups as relatively internally homogenous. Constructionist approaches tend to focus more on variations within groups: not all women are alike, not all ethnic minorities are alike,” (Hearn & Louvrier, 2015).

As Stokoe (2005) points out, “Accompanying this shift were methodological moves away from experiments, surveys and statistical analyses” in the first and second waves, “to the qualitative study of talk and text, as well as a shift from ‘essentialist’ to ‘performative’ theorizations of gender,” in the third. Although general theories of language and identity have largely shifted from essentialist to constructionist approaches, that is not to say that the foundation of essentialist categories upon which the waves of analytical practice were built can or should be pushed aside and ignored. As Sidnell notes:

There is an underlying tension here in so far as many researchers advance anti-essentialist, theoretical conceptions of gender (suggesting that gender emerges through the practices of talk) but at the same time employ the very same categories in their analysis. The theoretical notion of ‘performativity’ offered as an anti-essentialist antidote, is problematic in so far as it presupposes some

‘real’ set of actors who inhabit the roles of the *dramatis personae*. (2003)

That is, constructionist approaches can over-rely on the ‘authenticity’ of identity facets both of the speaker and performer, and of the facets themselves, and as such build upon the foundation of categories as established by essentialist approaches in order to then consider how the individual relationally adheres or diverges from these categories, and in what contexts, in the style of their performances.

The idea of ‘authenticity’ underlying various theories of language and identity is of particular interest in this paper, which analyzes the language not only of individuals who can be said to ‘authentically’ belong to the various performative aspects of their identity, but those who do not (and in many cases can be considered to instead ‘authentically’ belong to mutually exclusive categories). As we will see in the following section, and throughout the paper as a whole, not only are the identity-disguised performances of these individuals interactional in their construction, but they are also largely—especially in those cases in which the identity disguise is unintentional—perceptive as well, with the value of their ‘authenticity’ being frequently measured by the commodities of cooperation and suspicion inherent to the contexts in which they occur. The following section, then, builds specifically upon interactionist views of identity, and their application to identity disguise.

Part 2 – Identity Disguise

As was overviewed in the previous section, contemporary theories of language and identity have tended to move away from essentialist positions, which view aspects of identity such as gender, age, and other characteristics as features of an individual. Instead, there is a view of identity as something that an individual performs and engages in as an activity. This is best summarized by Bucholtz and Hall (2005) who set out the characteristics of identity as being **emergent** from interaction where an individual’s conversational turns will **index** identity **positions** either overtly or through implicature, in **relational** comparisons and contrasts. Such identity positioning can be partially deliberate and partially habitual and so less than fully conscious. This view and related views can for convenience be considered an interactionist view of identity.

Such interactionist views on identity likely represent a new consensus across the social sciences, including sociolinguistic study, and are somewhat at odds with most contemporary computational linguistic work and text analytics, where more essentialist identity categories are largely taken for granted, seen either as static independent factors of analysis or as target categories for predictive machine learning studies (e.g., Koppel, Argamon, & Shimoni, 2002). Gender-play, and indeed identity-play in general, is not a recent phenomenon on the internet, long predating the most common current social media sites and applications such as Facebook, Twitter, and, as we will discuss in Chapter 4, Tinder (e.g., Danet, 1998).

When it comes to interlocutors (and not analysts), as Herring and Martinson point out, “that people can be fooled by deceptive gender performances in text-based CMC supports Walther’s (1996) claim that online communicators tend to over-attribute characteristics of their

interlocutors from minimal cues,” (2004). Such minimal cues are discussed below.

Herring and Martinson further state that, “it is important to know an interlocutor’s identity in order to understand and evaluate the interaction; this is especially true for gender, which is conventionally associated with different norms, roles, and communication styles in most human cultures,” (2004). According to their work, gender identification through language cannot be cleanly done based on a binary assumption of gender, as “it is a combination of weighted features [...] rather than any single feature alone, that allows for accurate gender attribution, lending empirical support to the notion of gender ‘styles’” (2004).

A. Types of Identity Guising

It is first worth pointing out, of course, that not all attempts at guising one’s identity or features thereof are undertaken in an attempt to commit any crimes, or necessarily trick interlocutors into any altered perceptions intended to be at the detriment of other interactants, though much effort has been paid to analyzing intentional guising for criminal purposes (e.g., Chiang, 2019; Chiang & Grant, 2019; Grant & MacLeod, 2020; Rashid, Baron, Rayson, May-Chahal, Greenwood, & Walkerdine, 2013); in some cases, research delves into identity performance along gender and other lines as a tool to counter criminal activity online (e.g., Grant & MacLeod 2020; MacLeod & Grant, 2017 & 2021). Ainsworth and Juola provide several non-criminal examples in which authors may have disguised their language, or at least obscured their authorship, which beg questions of interested scholars (2018):

Literary scholars want answers to questions such as: Did Shakespeare write all the plays and sonnets attributed to him? Or, in a more modern context, did the *Harry Potter* author J.K. Rowling also write the crime novel, *The Cuckoo’s Calling*, under the pen name Robert Galbraith? Journalists want answers to questions like: Who is the inventor of Bitcoin? Historians are interested in questions such as: [...] Which portions of the Federalist Papers were written by James Madison, which by Alexander Hamilton, and which by someone else?

For similarly benign reasons, people online may guise all or some aspects of their identity in order to retain anonymity online (for gender perception reasons online, e.g., Eklund, 2011; Postmes & Spears, 2002; Riedl, Hubert, & Kenning, 2010; Song, Restivo, van de Riit, Scarlatos, Tonjes, & Orloy, 2015; Youn & Hall 2008), in order to engage in identity-play that explores alternate identity facets not otherwise attainable by an individual in face-to-face interactions (for gender play online, e.g., Bessière, Seay, & Kiesler, 2007; Huh & Williams, 2010; Leavitt, 2015; Williams, Consalvo, Caplan, & Yee, 2009), or due to reasons beyond their control or outside their overt interest, at least when it comes to the perception of their interactants (for the effects of outside-gender on performance and perception: e.g., Banakou & Chorianopoulos, 2019; Palomares & Lee, 2010; Wang & Wang, 2009).

1. Intentional guising

Intentional identity-guising, whereby an author performs an identity or facets of an identity dissimilar to their own, can occur on a spectrum of identity-play that may or may not

involve real-world identities, in whole or in part, either of the performer themselves, or of some other person. Erard (2017) refers to this as “disappearing in style”, though he notes that “people generally prove to be enduring amateurs at identifying the right changes to make.”

At one end of the spectrum is an author’s attempt to guise their identity as the real-world identity of some other, specific individual, such as in the cases of undercover officers assuming the identity of an underaged child with whom an online child predator and groomer is suspected of having interacted (MacLeod & Grant, 2017). An amalgam of this and the second type of identity-play discussed below can be seen in the Tennessee Facebook Murders, which are discussed further in Chapter 5. In this case, the female author based her guised persona on a high school acquaintance, using his name, image, and other real-world features in her guising as a part of the guised identity creation of “Chris” the CIA agent.

Another example is an author’s attempt to guise their identity as a specific, non-real individual—who shares some or no identity facets with the author’s real-world identity. This can be seen in various datasets discussed within this thesis, in complex examples such as the creation of the Syrian-American lesbian by a white, straight, American man in the A Gay Girl in Damascus case, or as simple as the police decoy’s attempt to portray himself more generally as an underaged female in the “Erin Princess Baby” case, both of which are discussed in subsection B below.

Finally, at the far opposite end of the spectrum, an author may attempt to guise their own identity simply by being “neutral” and attempting to provide no information, or conflicting information, either through their language features or the content of their language itself, that would not give an interlocutor any of the useful cues discussed by, for example, Walther (2006). Such performances might occur frequently, for example, on online fora that allow for the anonymity of posters (such as with “throwaway” accounts, e.g. Leavitt, 2015).

Beyond linguistic features, intentional identity guising may simply come down to aesthetics and no actual attempt by any author to keep their language coded by their purported identity, or keep the information conveyed by their language consistent with the aspects of this identity. This can be seen in subsection B below, for example, with the Ashley Madison hacks.

2. Unintentional guising

Unintentional identity-guising relies on the perception of interlocutors more so than the performance of an author, and may come down to simple aesthetics, environment, or the interference of some third party. As we will see in the sections below, although such instances of false perception guising an author’s identity may be unintentional, there is some evidence to suggest that authors may, equally unintentionally, change their language in part to cooperatively conform to their interlocutors’ perceptions.

Aesthetic identity guising is most common in situations in which an author’s language is paired with visual cues, whether these be a profile image, a character or avatar such as in a game, or the formatting of their text by color or other cues, which may alter an interlocutor’s perception through no actual linguistic or informational cues performed by the author

themselves. The reverse of this can be seen in sentiments such as the “G.I.R.L.” (“guy in real life”), in which even users on some sites who present themselves as female are perceived as male, as “there are no girls on the internet”—though sentiments such as this have fallen out of vogue in much of the internet now.

Third-party identity guising, as will be discussed in depth in the case of the Catfi.sh Tinder data, is identity guising wherein non-linguistic cues beyond an author’s control present them as some other identity thanks to the interference of a third party. In the Tinder data, for example, male Tinder users had their messages routed through female accounts, one they had paired with and one that had paired with another male user, leaving both male users to assume, based on the profiles relaying the other male’s message and no attempt by their interlocutor at guising, that they were speaking to a female.

B. Examples of Identity Play

The following section provides three real-world examples of intentional identity play, including an exploration of the language (or lack thereof) of the Ashley Madison hacks, an explanation of the “Erin Princess Baby” case, and a discussion of the A Gay Girl in Damascus blog. This discussion aims to consider both performance and perception in why these examples of identity play were and were not successful, how these performances do or do not affect the language use of the guising author, and how identity facets other than gender (specifically sexuality and age) may interplay with gender in a way that skews both performance and perception.

1. The Ashley Madison hacks

An analysis of some of the breached data from the hack of the dating site Ashley Madison provided “evidence that Ashley Madison created more than 70,000 female bots to send male users millions of fake messages, hoping to create the illusion of a vast playland of available women,” (Newitz, 2015).

To understand why Ashley Madison would bother to create quite so many bots (or any bots at all), it is important to consider how the service functioned. To initiate conversations on Ashley Madison, users had to pay a fee—meaning, in effect, that the bot accounts were meant to elicit conversations to be started by the substantially larger male user base. According to Newitz¹, these bots, or at least their profiles, were referred to as “Angels”, and she provided “a verbatim list, taken directly from the code, of the random messages the chat bot was programmed to spew,” shown in Table 1.1 below (2015).

‘hi’	‘hi’	‘hi’
‘hi (s)’	‘hi there’	‘how are you?’
‘hey’	‘Hey’	‘hey there’
‘hey there’	‘Hey there’	‘u busy?’
‘you there?’	‘any body home?’	‘Hi’
‘Hi’	‘Hi’	‘hows it going?’
‘chat?’	‘how r u?’	‘anybody home? lol’

¹ <https://gizmodo.com/ashley-madison-code-shows-more-women-and-more-bots-1727613924>

'hello'	'hello'	'Hello'
'hello?'	'whats up?'	'so what brings you here?'
'oh hello'	'free to chat??'	

Table 1.1 Ashley's "Angels" chat bot messages (Newitz, 2015)

Many of the inter-company emails leaked during the Ashley Madison hack discuss these bots, appearing to pay particular attention not to their language (as is clear in Table 1.1 above) but to their profiles. Newitz (2015²) provides the following email as an example:

There needs to be some personal content written in the details along with the preselected choices. It doesn't have to be a lot, but there should be something personal in most of them.

[In the profiles you provided] the Ethnicity were all set as Rather Not Say, they must be selected as multiple profiles all looking the same is an issue.

Newitz also references another email in which someone asks, "whether it's OK to reuse photos if they are in different states, and [the response is] no—she notes that many members travel and they might spot the duplicates," (Newitz, 2015).

It is, then, clear from these email exchanges, and indeed the seeming lack of attention put into the bots' actual messages themselves, that there was at least some impression by the Ashley Madison team that the important performance issue for the "Angels" lay primarily in the believability of their profiles. That they sought to avoid elements like the repetition of images and were aware that all bots (which most such sites are suspected to have to some extent) having their ethnicity marked as "Rather Not Say" would be a marker to other users that the profiles were potential bots due to such similarities, are evidence of this fact. Indeed, as we will see in Chapter 5 below, much of the success of the Tinder data discussed therein being perceived by interlocutors as genuine came down to a user's profile, and not their language, or even always the content of their language, which is discussed in the next example.

2. The "Erin Princess Baby" case

In Australia in 2006, the *R v. Plumridge* case, referred to here as the "Erin Princess Baby" (EPB) case, was brought to bear against Mr. Plumridge with an allegation "based on the proposition that he believed the person with whom he was engaged in conversation on the Internet was under the age of 16 years. In fact, his interlocutor was a middle-aged male police officer," (Lincoln & Coyle, 2012). As with the other two cases discussed here, one facet of the identity-guising of the male officer was gender, but another, arguably intertwined factor was age. As in many such decoy cases where the undercover officer (UCO) believes that their conversee perceived their identity-guised performance as genuine, Mr. Plumridge was charged.

In this case, however, Mr. Plumridge outlined several reasons why he, as he argued, did not believe that the UCO was in fact the 13-year-old female "Erin Princess Baby" (EPB), but the adult male that the UCO indeed was. As Lincoln and Coyle report, "He [Mr. Plumridge]

² <https://gizmodo.com/the-fembots-of-ashley-madison-1726670394>

also claimed that certain linguistic cues such as the use of the terms ‘spaz’ and ‘veg’ and the sign-off of the interlocutor ‘see ya later alligator’ were not terms that 13-year-old females would use,” (2012). Though these examples were linguistic, and to do with what Schler, Koppel, Argamon, & Pennebaker (2006) might include under the category of blog language, Mr. Plumridge also gave non-linguistic examples that altered his perception of the UCO’s performance as genuine, including (bullets added):

- the capacity of EPB to ‘work out’ how to send a file in 1 minute 11 seconds after initially saying ‘I have one (referring to a picture) but I am not sure how to send it’;
- EBP stating ‘she’ ‘only knew (sic) to chat really’ yet displaying significant familiarity with protocol/procedure as evidenced by ‘her’ retort ‘wats (sic) with the CAPS’ to the use of capitals (which are considered a form of shouting/aggression in online communication);
- and EPB’s observation that ‘she’ was in her office when ‘she’ was supposed to be home from school then correcting this glaring error” (Lincoln & Coyle, 2012).

Though part of Mr. Plumridge’s purported conclusion that the UCO’s performance had not been successful was indeed directly linguistic, much of it was not.

The three bulleted reasons above, though conveyed through language, were inconsistent not with the linguistic patterns of a 13-year-old female online, but with the content and context expected to be conveyed by a 13-year-old female online. The third point, that a purported 13-year-old female was at “the office”, when, generally speaking, 13-year-olds don’t spend time at “the office” after school, is the sort of information that is consistent with the identity performance being attempted.

In the end, with these and other arguments in hand, Mr. Plumridge was found not guilty of the charges brought against him. Such cases are indicative of why, most especially in the case of UCOs attempting to find child predators online, it is not only the language of the performer, but the content of that language, which is important to the interlocutor’s perception of the genuineness of that performance. As Lincoln and Coyle conclude, “glaring content errors, such as claiming to be off school but then stating the ‘she’ was in the office, as happened in the Plumridge case, will give the game away much more readily than syntactical and other stylistic errors,” (2012).

As we will see with the exploration of the Tinder data in Chapter 4 below, such content differences can, indeed, be much more salient factors when it comes to perception than any of the linguistic features discussed in Chapter 3.

3. The A Gay Girl in Damascus blog

As will be further discussed in Chapter 4 below, the A Gay Girl in Damascus blog (AGGID) was purportedly authored by Amina Arraf, “a lesbian Syrian blogger”, when in fact it was written by “Tom MacMaster, a 40-year-old American,” (Bell & Flock, 2011)³. Much of the

³ https://www.washingtonpost.com/lifestyle/style/a-gay-girl-in-damascus-comes-clean/2011/06/12/AGkyH0RH_story.html

language of AGGID was specific to Amina, discussing the politics of Syria, the issues of her homosexuality in a country where it was disparaged, and her related activism—none of which Tom was himself actually engaged with.

As we saw with the “EPB” case, it was not necessarily the features of ‘her’ language that caused the readers of the blog to begin to question Amina’s authenticity, but the contradictions of information the blog presented. This issue came to a head when, according to a blog posted by Amina’s cousin “Rania” (also Tom) on June 6th, 2011:

Earlier today, at approximately 6:00 pm Damascus time, [...] Amina was seized by three men in their early 20’s. According to the witness (who does not want her identity known), the men were armed.

And in a later update:

[...] all that we can say right now is that she is missing.

Bell and Flock go on to state, “News of her disappearance became an Internet and media sensation. The U.S. State Department started an investigation. But almost immediately skeptics began asking: Had anyone ever actually met Amina? On Wednesday, pictures of her on the blog were revealed to have been taken from a London woman’s Facebook page,” (2011). In subsequent posts on the 12th and 13th of June, Tom came clean, apologized, and explained (or at least attempted to).

It was not the linguistic features of Amina that led to the eventual suspicion and exposure of the true author behind the AGGID blog, but the contents of her post, and other, real-world failures of performance. Specifically, in this case, the suspicion came down to the extreme scenario of Amina’s purported abduction, and the fact that the blog reported on Amina’s supposed attendance at several key public events during the Syrian spring, when it turned out that none of the other activists in her network had ever actually met or spoken to her outside of text exchanges.

More to the point, what readers might perceive as acceptable for a lesbian female but not a straight female, which may include language that is otherwise considered to be more male coded, is additionally important to consider, regardless of how accurate such features would be to lesbian females. This last point was evidenced in the Ashley Madison hack, discussed above, wherein the language of the bots, and any gender-coding of that language, was not nearly so important (if it was indeed important at all) as the believable perception of the bots’ profiles to men using the service.

Finally, as we will see in the Catfi.sh Tinder data in Chapter 5 below, there are many online arenas in which a reader has no reason to immediately question the identity of an author as it is being presented. Indeed, Amina had accounts elsewhere on the internet, including Facebook, an email address, and accounts on a variety of other forums and sites, with which she was frequently engaged. Moreover, in contexts such as the Ashley Madison and Tinder conversations, readers have every reason to ignore the kinds of cues that might have otherwise altered their perception of their interactant’s performance.

C. Overview

In forensic spheres, an investigator or an analyst may often be posed with the question of identifying an author, whether by providing a profile of their likely identity, or by determining between a set of probable suspects who the most likely author is. This may come in instances in which no identity is given as a baseline, or in which there is reason to question whether the proffered identity is genuine. This thesis, then, seeks to explore the interplay of both identity performance and audience perception in a variety of cases, both forensic and non-forensic, both when guising was intentional and incidental, both when audience and perception was and was not a factor, and how audience, genre, register, and context might influence performance. The ultimate goal of this thesis, then, is to ask first whether there are any features that provide useful indicators of an author's likely gender, and second whether any particular pattern of those features is more indicative of overt performance rather than natural gender distribution.

As outlined by sections below, this thesis first establishes a set of features to be applied to the four major datasets analyzed herein, and in doing so also establishes the need for a baseline or reference corpus, and a candidate used throughout this thesis (an edited portion of the Blog Authorship Corpus). This thesis then goes on to consider not only the standard application of these features (which can be seen, alone, as a largely algorithmic and non-interpretive step), but also how these features must be responsibly situated by the analyst. This includes not only by considering appropriate linguistic theories and their meta-applications to computer-mediated communication as a mode of discourse types including audience, discourse community, and commodity, but also the micro-applications of genre, register, and context to each individual case as it is situated within the real world.

Notably, what this thesis does *not* seek to do is to attempt to conduct or establish either a statistical approach or a one-fits-all model of analysis. Although the features proposed, tested, and applied are based on prior research that had the luxury of a statistical approach to big-data distributions, the datasets in this thesis—and indeed datasets analysts are likely to be tasked with analyzing in a forensic sphere—are often neither large nor robust enough for such an approach to be either appropriate or responsible. Further to that point, real-world forensic data does not exist in a vacuum, and (because, again, it does not often exist at such big-data scales as to make the task untenable) as such analysts have available to them not the luxury of lab-restricted variables, but the luxury of context, by which this thesis seeks to argue an analyst can reasonably and responsibly apply such statistics-based but non-statistically-applied feature sets, and interpret them as situated in their appropriate, real-world contexts.

Part 3 – Overview of Chapters

This paper is comprised of six chapters. This chapter, Chapter 1, provided an introduction and overview of theories of language and identity through the three waves of sociolinguistic research on identity from essentialist to constructionist perspectives and how they have shaped not only the modern landscape of considering identity as interactive and

performative, but also how the approaches will be applied throughout this paper. Chapter 1 goes on to provide an overview of identity disguise alongside examples of identity guising, and a brief overview of what is to follow.

Chapter 2 provides a literature review, covering both earlier, more qualitative (or at least not computational) approaches to the study of gendered language, and more contemporary and quantitative (or at least more computational, big-data) approaches to the same categories and variables. The discussion of earlier qualitative approaches outlines historical approaches to analyses of language and gender, and various features of gendered language as proposed by these approaches. The discussion of later qualitative approaches overviews more modern genres of computer mediated communication (CMC) that lends itself to more computational, big-data approaches to studies of language and gender and provides an overview of various features of gendered language as proposed by these approaches. The aim of this chapter is not to argue in favor of any specific features as accurate or blanket indicators of gender identity or at least identity performance, but rather to establish a baseline from which the 'inauthentic' performance of identity-guised, non-group members (as discussed in Chapter 1) can be discussed in relation to such pre-proposed variables.

Chapter 3 builds upon the features proposed by the literature review and establishes a methodological approach to analyzing them in the data considered in the later chapters. This is based first on an analysis of the edited Blog Authorship Corpus (eBAC), and an establishment of the twenty features and their gendered distributions considered in the subsequent analyses, and of the eBAC as a reference corpus. Chapter 3 then goes on to overview the application of various linguistic theories to computer-mediated communication, and an overview of performance and perception in identity play online, providing real-world examples. This chapter aims to demonstrate not only the quantitative approaches that consider gender as a binary category along an outline of statistically proven variables, but also the subsequent use of theoretical and qualitative approaches that can better situate the binary findings, or any more essentialist identity categories, in their interactive performative and perceptual contexts.

Chapters 4 and 5 contain applications of these approaches upon various datasets of CMC in which some form of identity disguise took place. These chapters are concerned with identity and perception, and the interplay with performance, and provide further analytical steps to better place the twenty features established in Chapter 2 in their appropriate context by considering theoretical, qualitative approaches such as those introduced in Chapter 3. Chapter 4 focuses on non-forensic data, applies this framework of analysis to the A Gay Girl in Damascus blog discussed earlier in this chapter, and establishes the use of keywords in subsequent analyses. The latter part of Chapter 4 applies this framework to an analysis of the catfi.sh Tinder data, briefly introduced earlier in this chapter, before going on to consider the features in the context of the types of conversational progressions that occur within the data, and the types of conversational strategies that contribute to these progressions. Finally,

Chapter 4 discusses both cooperation and suspicion as commodities in online exchanges, and how these dueling facets impacted many of the datasets considered so far; and briefly explores the overarching application of some of the linguistic theories proposed earlier in Chapter 3.

Chapter 5 applies the findings of the various approaches and analyses outlined in the previous chapters to two forensic cases with highly contrasting features. The first concerns the Hemmert case, in which a husband was alleged to have kidnapped and assaulted his estranged wife, and to have taken possession of her phone, sending texts as her so no one would know she was missing. This chapter first applies the twenty features established in Chapter 2 to the Hemmert data as is, and then goes on to consider the appropriate weighting of these features in the Hemmert data, and the context and perception in an analysis of the Hemmert data as introduced in Chapter 3. The next concerns the Potter case, in which a woman was alleged to have coerced her friend and family members into committing a double homicide on her behalf, via the guised identity of a CIA agent using her Facebook page and email address. This analysis combines the appropriate analytical tactics from the previous Hemmert case and those in Chapter 4 to the Potter data at once.

Finally, Chapter 6 provides an overview of the analytical chapters and approaches in this thesis, comparing the findings across analyses in order to determine whether any may be more useful indicators of shared authorship or overt gender guising, and what context might be appropriate to consider when applying these features. This chapter considers which of the twenty features, if any, worked not only in accurately indicating gender performance, but in accurately indicating gender disguise across the various datasets, and in which contexts. Finally, Chapter 6 provides the overall takeaways of this thesis.

Chapter 2 – Literature Review

Part 1 – Introduction

This chapter seeks to outline approaches to the study of language and gender, both in how theoretical baselines for analyses have shifted, and in what features have been proposed as distinct indicators of an author's gender. Chapter 1 began with a walkthrough of the three waves of sociolinguistic research; Part 2 below continues this exploration with five other approaches within the first two waves, and the contemporary and more recently computational approaches that fall within the third. Parts 3 and 4 then provide an outline of some features these various approaches have suggested, either through qualitative observation or quantitative, often statistical analyses, both historical and contemporary. The features covered in Part 3 are often more observational than proven, and were often considered along the lines of spoken more so than written language (the latter of which comprises all of the case analyses in this paper, though they each have the more communicative function of the former). However, these features are, where possible, aligned with those outlined in Part 4, which are themselves carried through the rest of this paper. Finally, this chapter considers the overall implications and applicability of these approaches to the later analyses contained within this paper.

Part 2 – Approaches to the Study of Gendered Language

Building on the three waves of sociolinguistic research as outlined in Chapter 1 above, this section seeks to further outline both historical and contemporary approaches to the theories and analysis of language and gender. This section first situates five historical approaches as suggested by Coates (2004) and Cameron (2010) within the first two waves of sociolinguistic research, their historical trajectory, and criticism. This section then shifts to more contemporary approaches and ideologies, as well as providing an introduction to more computational approaches.

A. Earlier Approaches

Coates (2004) outlines four different approaches to studies of gendered language, in basic, chronological order of their popular use: the deficit approach, the dominance approach, the difference approach, and the dynamic or social constructionist approach. Cameron (2010) proposes a fifth approach: new biologism. These approaches are briefly discussed below, to better classify the further studies discussed on gendered language as based on their appropriate theoretical foundations. These approaches fall largely within the first and second waves of sociolinguistic research on language and gender as discussed in Chapter 1 and as such tend to rely on essentialist categories of identity.

1. The deficit approach

The deficit approach “claims to establish something called ‘women’s language’ (WL), which is characterized by linguistic forms such as hedges, ‘empty’ adjectives like *charming*, *divine*, *nice* and ‘talking in italics’ (exaggerated intonation contours). WL is described as weak

and unassertive, in other words, as deficient. Implicitly, WL is deficient by comparison with the norm of male language,” (Coates, 2004). Works utilizing the deficit approach include Jespersen’s 1922 book, *Language: Its Natural Development*, which similarly defines adult male language as the standard or norm of language, leaving the language of women, children, and ‘others’ to be considered to have something ‘inherently’ wrong with it. Though Jespersen did note that this “way of viewing language is fraught with some great dangers,” (1922), his established approach nevertheless positioned women’s language as deficient with respect to the language of men.

In opposition to the deficit approach, Coates cites examples such as Lakoff’s 1975 *Language and Woman’s Place*, which posited that women’s language, or at least the language expected of women (including features such as tag questions, question intonation, and “weak” directives), in large part thanks to Jespersen’s own assertions, served to maintain women’s (inferior) role in society. Indeed, Jespersen suggested this likelihood as well, stating that the deficiency of women’s language “is not one of sex really, but of rank”. Lakoff, “was challenged because of the implication that there was something intrinsically wrong with women’s language, and that women should learn to speak like men if they wanted to be taken seriously,” (Coates, 2004), giving way to the dominance approach alluded to by Jespersen’s suggestion of ‘rank’.

2. The dominance approach

The dominance approach, derived from works such as Lakoff’s, “sees women as an oppressed group and interprets linguistic differences in women’s and men’s speech in terms of men’s dominance and women’s subordination [...] Moreover, all participants in discourse, women as well as men, collude in sustaining and perpetuating male dominance and female oppression” (Coates, 2004). This approach positions the language of women and men at a similar distance as the deficit approach, with the largest difference relying not on the deficiency of women’s language or capacity for sufficient language, but on a concerted social effort to keep women’s language as deficient to maintain the authority of men’s language. Coates cites researchers such as West and Zimmerman (1983), who describe “doing gender” as a way of “doing power”, which has been considered by numerous scholars (e.g., Carli, 1990; O’Barr & Atkins, 1980).

Other works include Spender’s 1980 *Man Made Language*. In discussing the origin of the idea that women talk a lot, for example, Tannen (1991) observes, “Dale Spender suggests that most people feel instinctively (if not consciously) that women, like children, should be seen and not heard, so any amount of talk from them seems like too much. Studies have shown that if women and men talk equally in a group, people think the women talked more. So there is truth to Spender’s view”. Thus, the approach of male language as dominant stems in part from societally induced perceptions not of their language itself, but of the roles of men and women in society.

3. The difference approach

The difference approach “emphasises the idea that women and men belong to different subcultures. [...] The advantage of the difference model is that it allows women’s talk to be examined outside a framework of oppression and powerlessness,” (Coates, 2004). The difference approach attempts to contrast from the deficit and dominance approaches as one of equality, differentiating men and women as belonging to the different ‘sub-cultures’ they have been socialized to since childhood, which in turn results in the varying communicative styles of men and women.

Cameron (2010) aligns Tannen’s 1991 *You just Don’t Understand* with John Gray’s 1992 *Men are from Mars, Women are from Venus*, in which she claims, “the emphasis was on solving relationship problems: differences between sexes were treated as simply a fact of life, and most writers had no interest in debating their underlying causes,” considering each work to fall under the joint categories of popular science and self-help rather than within true scientific research. Coates cites Tannen’s applying the difference approach to “*mixed talk*” as eliciting criticism from other authors such as Troemel-Ploetz (1991), Cameron (1992), and Freed (1992), who “argue that the analysis of mixed talk cannot ignore the issue of power,” (Coates, 2004). Some more modern approaches (e.g., Eckert & McConnell-Ginet, 1992) consider dominance and difference in conjunction, and “the processes through which each feeds the other to produce the concrete complexities of language as used by real people engaged in social practices.”

4. The dynamic approach

The dynamic approach, or social constructionist perspective, has “an emphasis on dynamic aspects of interaction. [...] Gender identity is seen as social construct rather than as a ‘given’ social category,” so while various social constructs may be affiliated to particular groups, they can be utilized by any speakers as they see fit (Coates, 2004). Coates cites Crawford (1995), who claimed “that gender should be conceptualised as a verb, not a noun!”. Other scholars include Zimmerman and West, whose 1996 study of turn-taking in conversation concluded, for example, that while men generally use minimal responses such as “uh-huh” and “yeah” less frequently than women unless it is to show only agreement, both men and women can employ minimal responses for interactional functions, rather than for gender-specific functions.

Although these approaches differ, and “the deficit approach is now seen as out-dated by researchers (but not by the general public, whose acceptance of, for example, assertiveness training for women suggests a world view where women should learn to be more like men)”, the four approaches “do not have rigid boundaries: researchers may be influenced by more than one theoretical perspectives” (Coates, 2006).

5. The new biologism approach

The new biologism approach considers not simply the social roles of males and females and their relative position to one another, but also the functional biological differences that

have influenced these social roles. According to Cameron (2010), much of this approach:

is related to the thesis that early human males maximized reproductive success by competing with one another for access to females and for resources those females valued, whereas females maximized their success through forging relationships with others—the children they cared for, the mates they depend on to protect and provide for them, the other community members, especially other women, whose assistance they might need.

New biologism, then, appears to argue that there may be some biological innateness to the difference in language between biological men and women, in that “[o]n that basis, it is logical to suppose that selection pressures would have favored competitive men and co-operative empathetic women,” (Cameron, 2010). As the deficit approach characterizes women’s language, and thus their role, as lesser, the dominance approach characterizes men’s language, and thus their role, as oppressive, and the difference approach characterizes their language (and roles) as different but equal, new biologism would consider biology as an explanation for why these differences might tend to persist even if societal roles, norms, and expectations shift.

However, as Cameron (2010) asks: “How far, though, does the linguistic evidence really accord with that supposition,” that biological differences can be understood as the cause of linguistic differences between gender? Especially when, as stated above, societal roles, norms, and expectations shift. This shift, then, is what the dynamic approach appears to seek to highlight: that in the context of biological trends, men and women still have the capacity to employ *any* otherwise biologically gendered language features for interactional purposes. However, as Freed (2014) argues:

[...] the lines being drawn to different people along gender and sexual lines are even more disturbing than in past years because of the inclusion in the discussions of so-called scientific evidence from controversial brain science research, what Deborah Cameron (2009) calls the new ‘biologism.’ Questionable and limited findings about the human brain, used in conjunction with sweeping theoretical assertions that include claims about the way women and men use language, are increasingly incorporated into public debates about male–female difference.

Other research considers whether the linguistic differences that may be inherent to sex and gender actually matter to, for example, cross-sex interactions (e.g. Mulac, 2006).

B. Contemporary Approaches

While the above approaches fall largely into the first and second wave of sociolinguistic study, more contemporary approaches tend to fall neatly into the third, with more of a focus on interactional, intersectional, and constructionist facets of identity. Finally, third wave approaches are benefitted not only by new ideologies, but new methodologies born from technological advancements, which can further enable research to consider multiple and variable identity facets at once.

Contemporary approaches to the study of language, gender, and identity are also often concerned with other intersecting identity facets. As Lazar (2017) points out, “According to

intersectionality theory, gender identity is not homogenous or singular; rather, femininities and masculinities are heterogenous and plural.” *The Handbook of Language and Gender* (HLG), originally published in 2003, has since been updated to *The Handbook of Language, Gender, and Sexuality* (HLGS). According to the introduction of this second edition, “we have attempted to highlight the ongoing importance of sexuality to the field and the close connections between gender and sexuality,” (Ehrlich & Meyerhoff, 2014). Aside from sexuality, the *Handbook* also considers “research on language, gender, and sexuality in non-anglophone locations around the world [...] from a wide range of languages and cultures,” (Ehrlich & Meyerhoff, 2014).

In many cases, contemporary approaches specifically seek to undo the ideologies instilled by the first two waves, such as, “the significant discrepancy that existed between public representations and popular perceptions of how women and men speak (and how they are expected to speak) and the empirically verifiable character of the language that people use,” which, “underscore the vitality of deeply engrained stereotypes about sex and gender and the weight and influence of societal efforts to maintain the impression that, simply put, men and women are different,” (Freed, 2014). Gender identity can, instead, be considered to be as emergent as other identity facets, able to be contextualized in how it deviates from the stereotypes established by earlier linguistic and popular science research.

Finally, while earlier waves did include quantitative methods, third wave approaches have the benefit of highly computational approaches, as well as access to vast amounts of written, communicative language across many genres and mediums in the form of computer mediated communication (CMC). As Ehrlich and Meyerhoff (2014) point out, “there are clear advantages in bringing together quantitative and qualitative approaches to the extent that ‘macrolevel quantitative research identifies the gendered norms on which speakers are drawing, the ground against which individual choices must be interpreted.’”

1. Intersectional Approaches

Contemporary approaches to the study of language and gender have largely eschewed relying whole cloth on the essentialist categories established by the earlier waves of sociolinguistic research, instead shifting to more constructionist approaches. As McElhinny (2003) outlines:

To argue that differences found in people’s behavior, including their speech behavior, can simply be explained by invoking gender is to fail to question how gender is constructed. Instead, one needs to ask how and why gender differences are constructed in particular ways and what political interests are served by such constructions. This question is often linked to an intersectional approach to the construction of identity, where gender is understood as imbricated with sexuality, race/ethnicity, class, nation, and so on.

As mentioned in the section above, the refocus of the *Handbook* from the HLG to the HLGS is itself indicative of how embedded sexuality is in much contemporary research on language and gender.

Much previous research can be categorized as studying deviation: women’s language as the deviant form of men’s language; later, men and women as deviating from one another;

later still, as other identity facets such as sexuality, race, class, and education explaining how women and men deviate internally from the norms of their own gender identity. Hall (2008) offers examples of deviant or “exceptional speakers” historically being considered to include, among others, “effeminates and feminists”, “the woman”, “hippies, historians, and homos”, “sissies and tomboys”, and “queers and the rest of us.” In much the same way that research on language and gender has moved away from considering men’s language as the norm or even baseline, contemporary approaches argue for feminism (e.g., Chen, 2021; Lazar, 2017), queer theory and trans linguistics (e.g., Chen, 2012; Milani, 2021, 2017; Zimman, 2021), and race (e.g., Bucholtz & miles-hercules, 2021; Ehrlich, 2021; Foster, 2021) as being of equal or more importance to often cisgendered, heterosexual, and Anglophonic perspectives often taken for granted as benchmarks for analysis.

Although this paper only briefly touches upon the interaction of gender with other identity facets such as age (Erin Princess Baby) and, more specifically, sexuality (the AGGID blog), current approaches to language and gender consider sexuality to be more inextricably linked to gendered linguistic performance. While, as with previous approaches to the studies of language and gender providing a foundation upon which non-normative gendered linguistic expressions can be better analyzed, “binary views of gender and sexuality remain an important ontological basis of the social order,” and continue to characterize much research into language, gender, and sexuality, it is increasingly important to consider the intersectional nature of multiple identity facets in such approaches rather than relying too heavily on pre-established norms as absolutes (Meyerhoff & Ehrlich, 2019).

In the same and only case discussed in this paper in which an author’s sexuality was part of their disguised identity (the AGGID blog), so too was race only at issue once. Though the concern either did not come up in the other non-forensic datasets analyzed—because it was either never explicitly provided (the Blog Authorship Corpus) or was anonymized (the catfi.sh Tinder data), or in either of the forensic cases analyzed (Hemmert and Potter) because all of the participants were white—race, ethnicity, nationality, remain as important and underexplored as many other intersectional facets of gender identity, and identity more broadly. As Bucholtz and -miles-hercules (2021) point out, “In short, language and gender studies had – and still has – a race problem.”

2. Emergent Approaches

Contemporary approaches to the study of language and identity consider identity facets such as gender to be emergent through performance and interacting with other identity facets, as discussed above, as broad as age, race, and social status, and as specific as individual communities of practice and other local and social linguistic communities with their own performative norms. This “third wave” approach to language and gender has moved away from the more essentialist thinking of previous approaches to more formative thinking. Gender identity is a performance for everyone, and how you perform that identity draws on the resources you have available to you. Or, as Besnier and Philips (2014) state, “gender is not a

predetermined category, but rather the emergent production of social practices, including interactional practices.”

As Ehrlich and Meyerhoff (2014) point out, “The theoretical claim of Butler’s—that identities do not exist beyond their expression—is probably most transparent when an individual’s ‘expressions’ of identity depart from what we take to be their ‘true’ identity.” While the performance of a specific woman may draw on knowledge of how she used language in the past, as situated in the identity facets which she has the access to draw from, performing simply as a woman might lead the performer to draw on their understanding of the macro-level ideological categories of ‘women’, or at least what is expected of women.

That is, as discussed elsewhere throughout this paper, identity is just as perceptive as it is performative in its interactional expression, in that the association of a performance to various identity facets such as gender is limited by the perceiver’s understanding of both the performer’s available identity resources and available identity resources as a whole, and in the performer’s intended expression. As Ehrlich and Meyerhoff (2014) note, “the issues of how best to study gender as an interactionally emergent phenomenon and how best to warrant claims about gender relevance will no doubt continue to raise questions about research design and methods.”

Worth noting in the context of this paper, in particular, is that even emergent categories “have tended to problematize a conceptualization of masculinity and femininity as a priori *analyst’s* categories” (Benwell, 2014). The analyses throughout this paper consider the categories not of masculinity and femininity, but of male and female, as they are derived from the same binary macro-level categorizations provided by the Blog Authorship Corpus (as detailed in Chapter 3), though this paper does not seek to argue any such gender binary as absolute, nor is it assumed that a given author will necessarily comport with the established patterns of men and women’s language as established in Parts 3 and 4 below. Instead, as is further discussed below, while the analytical approach is indeed structured along more essentialist baseline categories, the actual production—again, what could be considered as deviation or exception—is similarly graded as more masculine or feminine than either baseline, either by other intersectional identity facets as discussed in the previous section, or by the interactional contexts as discussed in this section.

3. Computational Approaches

As will be further overviewed in Part 4 below, many contemporary approaches to the study of language and identity are at least partially computational, with a variety of theories, methodologies, and technologies being studied and innovated upon in order to conduct large-scale and reliable quantitative (and sometimes automatic) analyses of language data along

facets of identity categories, identification, and attribution. Organizations such as PAN⁴ seek to encourage refinement of these methodologies in the fields of plagiarism analysis, authorship identification, authorship obfuscation, author profiling, and multi-author analysis. As Rosso et al. point out, “The practical importance of such technologies is obvious for law enforcement, cyber-security, and marketing,” (2019).

PAN has conducted a number of shared tasks since 2009, and shared tasks specific to profiling since 2013, the latter of which considers gender as a distinct variable for analysis. These shared tasks have included the categories of general authorship profiling (2013-2018), profiling iron and stereotype spreaders on Twitter (2022), profiling hate speech spreaders on Twitter (2021), profiling fake news spreaders on Twitter (2020), bots and gender profiling (2019), and celebrity profiling (2019, 2020). In their overviews of the submitted approaches to various relevant tasks, Rangel et al. (2013, 2016, 2017, 2018, 2019, 2020) note that they generally fall into the categories of content or style-based.

Content features include, for example “Latent Semantic Analysis, bag of words, TF-IDF, dictionary-based words, topic-based words, entropy-based words” including “named entities, sentiment words, emotion words, and slang, contractions and words with character flooding,” (Rangel et al., 2013). Style-based features have included, for example, “frequency of punctuation marks, capital letters, quotations [...] POS tags or HTML-based features as image urls or links,” (Rangel et al., 2013), “the frequency of use of function words, words that are not in a predefined dictionary, slang, [...] unique words [...] use of specific sentences per gender [...] and age [...] sentiment words,” (Rangel et al., 2016). Other approaches have included what Rangel et al. (2017) refer to as “deep learning techniques”.

In many of the submitted PAN approaches, analysts make frequent use of both character and word n -grams, either in isolation at varying lengths, or in conjunction with other features such as “word embeddings, and with beginning and ending character 2-grams” (Rangel et al., 2017), combining “POS tags n -grams with syntactic dependencies to model the use of amplifiers, verbal constructions, pronouns, subjects and objects, types of adverbials, as well as the use of interjections and profanity,” (Rangel et al., 2018), and in combination with “Support Vector Machines,” (Rangel et al., 2020), to list only a few.

The goal of individual PAN tasks and many such computational studies is accuracy, demonstrated by the ordering of aggregated participant results for each task by accuracy in each category. The models used to achieve these accuracy rates are not pre-existing models simply applied to each new task, but rather are developed and trained on relevant baselines provided by each task. These development datasets share many of the same features of the eventual evaluative datasets upon which accuracy is tested, such as belonging to the same genre, topic, and source. Thus research like that conducted for and by PAN would appear to move away from the idea that any one model can maintain predictive accuracy on any given

⁴ <http://pan.webis.de>

dataset, and move toward the idea that given the particularities of a dataset or analytical tasks, different variables that have previously demonstrated useful accuracy or lack thereof may be more or less relevant for any given next task. Similarly, as is discussed further throughout this paper, the analyses applied here do not take a given variable as an absolute demonstrator of gender identification in any and all circumstance, but rather situates each variable within the individual datasets and analytical questions at hand. Such approaches are situated within the cooperation between an algorithm and an analyst, as suggested by Swofford and Champod (2021), and discussed further in Chapter 3 below.

While many of the features considered throughout PAN and other contemporary computational studies are not necessarily specific to gender and were not necessarily derived from specific language and gender research even if they did show statistical gain for gender as a category, they nevertheless comport with many of the features outlined in Part 4 which are tested both on gender specifically, and report any statistical gain on gender explicitly.

C. Overview

Contemporary qualitative and quantitative approaches differ largely in that current qualitative approaches have trended toward constructionist categories while quantitative approaches must rely on some framework of essentialist categories. As we have seen in the sections above, both approaches are continually evolving and refining their methodologies as a means of overcoming any analytical shortcomings. Where qualitative approaches have the benefit of analyzing variation in context along the intersection of numerous variables, computational approaches have the benefit of applying a multitude of features to large-scale amounts of data with often statistically significant results. As Ehrlich and Meyerhoff (2014) point out, “innovations in technology and fresh forms of social media will allow researchers to access new kinds of data and to explore potentially new ways of constructing gendered and sexual identities”.

While many of these areas of research, briefly outlined above, that are at the forefront of gender and identity theory in contemporary approaches are not heavily considered beyond this chapter of the paper, the intersectionality of identity and identity performance is still considered throughout the analyses herein. As is discussed in the section below, variations within a given author’s gendered language features are often performative and emergent, and considered in the context of other identity facets they have access to (such as the authors at issue in the Hemmert case being not just male- and female-identifying, but also a husband and a wife, a husband and a father, and so on), and in the audiences to whom they are speaking (such as in romantically-suggestive and intended-to-be heterosexual interactions in the catfi.sh Tinder data where each female-interested male believes their male conversee to be female, and in platonically-oriented same-sex interactions between the male Jamie and his *guy friend* Chris who was the performed identity of Jamie’s female romantic partner).

Although this paper only briefly touches upon the interaction of gender with other identity facets such as age and, more specifically, sexuality, current approaches to language

and gender consider sexuality to be more inextricably linked to gendered linguistic performance. While, as with previous approaches to the studies of language and gender providing a foundation upon which non-normative gendered linguistic expressions can be better analyzed, again, “binary views of gender and sexuality remain an important ontological basis of the social order,” and continue to characterize much research into language, gender, and sexuality, it is increasingly important to consider the intersectional nature of multiple identity facets in such approaches rather than relying too heavily on pre-established norms as absolutes (Meyerhoff & Erlich, 2019).

As discussed in the sections above, and elsewhere throughout this paper, as Ehrlich and Meyerhoff (2014) point out, “much work claiming to adopt an anti-essentialist, constructivist understanding of gender tends to precategorize groups of people as women and men and then investigates how women ‘do femininity’ and how men ‘do masculinity’.” That is, much research, even that that claims to value performative emergence and move away from essentialist categories established in earlier waves of sociolinguistic analysis, nevertheless structures itself around the same baseline. As was discussed in Chapter 1 as “authenticity”, and in Chapter 4 as “commodity”, adherence to or deviation from the expected norms of these categories—or as “deviant” or “exceptional” speakers as discussed throughout this chapter—remain of importance to researchers. Indeed, in the instances of identity disguise analyzed throughout this paper, it is these same deviations or exceptional performances of a perceived identity that are of interest not only to a performer’s interlocutors, but also to this paper.

The Parts below outline numerous features as proposed by historical qualitative approaches and contemporary and often online approaches to gendered language, and their intersection with one another.

Part 3 – Qualitative Approaches

In order to establish what features will be considered in the analyses below, it is first important to establish a general background of the literature on language and gender. Section A below seeks to situate these approaches in their historical context, as they tend to build off one another, regardless of whether they are in agreement or opposition. Numerous proposed features that may indicate, to varying degrees, the gender of a given author based on their language have arisen from this literature. These features are outlined in Section A, within the context of when and why they were proposed. Although this thesis seeks to determine useful features of gender-guised language for analysis, this thesis *does not* seek to argue the validity or invalidity of any of these features as accurate indicators of gender, and much research exists that covers the theories and ideologies of language and gender research (e.g., Cameron, 2010, 2014; Newman, Groom, Handelman, & Pennebaker, 2008; Speer, 2005; Speer & Stokoe, 2011), as are touched on below.

A. Proposed Features of Gendered Language

According to Lakoff (1975), “[w]omen’s language’ shows up in all levels of the grammar

of English. We find differences in the choice and frequency of lexical items; in the situations in which certain syntactic rules are performed; in intonational and other suprasegmental patterns.” Differing features include those which occur within minimal responses, questions, turn-taking, changing the topic of conversation, self-disclosure, verbal aggression, listening and attentiveness, dominance versus subjection, and politeness. These features, outlined in the table below, rely on earlier, more essentialist approaches to gender identity categories.

	female	male
Lexical differences	more specific, weaker expletives, neutral or female-coded, empty adjectives, intensive “so”	more general, stronger expletives, neutral only (no weaker or female-coded)
Situations in which certain syntactic rules are performed	questions and tag questions as rhetorical means of engagement, to verify information, or to hedge	information-seeking questions
Tag questions	tag questions (as hedges or soft declarative statements)	regular questions (no hedges)
Intonational and other suprasegmental patterns	uptalk	
Minimal response	more frequent, cooperative use to show attentiveness; seldom pause for minimal responses, but give them often	only to show agreement; often pause for minimal responses, but seldom give them
Turn-taking	take cooperative turns, elicit turns by hedges, take hedge-elicited turns	center toward their own point, remain silent in both using and replying to hedges
Changing the topic of conversation	listen attentively, encourage and elaborate upon others’ topics cooperatively	provide topics, change the subject
Self-disclosure	tend toward self-disclosure to offer sympathy/empathy	tend toward non-self-disclosure; provide advice/solution instead
Verbal aggression	indirectly, relationally, and socially aggressive; nonverbally aggressive	overtly and explicitly aggressive; physically aggressive
Listening and attentiveness	provide feedback to show attentiveness; expect feedback in the form of agreement	remain silent or provide only minimal responses; interrupt and overlap more; expect feedback in the form of advice/solution
Dominance versus subjection	use humor as self-protection; use humor to create solidarity; concerned with relationships	use humor to exert dominance; use humor to create solidarity; concerned with power
Politeness	deferential politeness	camaraderie politeness

Table 2.1 Gendered features as suggested by Lakoff, 1975

1. Lexical differences

According to Lakoff (1975), one difference in gendered language is “in the choice and frequency of lexical items.” Women’s lexical choices tend to be more specific (at least within the realms deemed unimportant to men), “weaker” as in expletives, and neutral or female coded, while men’s lexical choices tend to be more general (again, within the realms deemed unimportant to their attention), “stronger” as in expletives, and neutral only, with female-coded

options deemed unavailable to them. Much research has been paid to lexical differences in different contexts, such as in conversations (e.g., Singh, 2001), over the phone (e.g., Boulis & Ostendorf, 2005), and on social media (e.g., Bamman Eisenstein, & Schnobelen, 2014).

Lakoff offers the example “(2) The wall is mauve,” as more indicative of a woman’s lexical choice, as women, “make far more precise discriminations in naming colors than do men.” Lakoff posits that the attention women pay to such lexical choices as color specificity, and the lack of attention men pay to the very same terms, is because such terms are deemed outside the realm of importance for men, thus relegating such distinctions to the less important female realm of experience.

According to Lakoff, “aside from specific lexical items like color names, we find differences between the speech of women and that of men in the use of particles that grammarians often describe as ‘meaningless’.” Lakoff offers the following examples:

- (3) (a) Oh dear, you’ve put the peanut butter in the refrigerator again.
- (b) Shit, you’ve put the peanut butter in the refrigerator again.

In this example, the difference between the use of “Oh dear,” and “Shit,” is the difference between women’s and men’s lexical choices, respectively, as men use “stronger” expletives, and women “weaker”. These differences, says Lakoff, are imposed upon us when we are but boys and girls—while a temper may be thought of as typical for a boy, and, eventually, man, girls and women alike are expected to be *ladies*, and thus must be tempered enough to use only the “weaker” expletives. However, as Lakoff points out, while it is becoming increasingly acceptable (or at least common) for women to avail themselves of “stronger” expletives in everyday conversation, men are still largely exempt from availing themselves of the “weaker” expletives in any genuine context.

Similarly, Lakoff offers adjectival categories, both neutral and women-only, that typify the speech of either gender. While either gender may use neutral lexical items such as “great, terrific, cool, neat”, only women can use female-coded, or “empty” lexical items, such as “adorable, charming, sweet, lovely, divine.” For a man to use female-coded words is acceptable only if he meets other criteria, such as being upper-class, or an academic. Women, however, must sometimes revert to the neutral terms as well, as in the example Lakoff offers: “However feminine an advertising executive is, she is much more likely to express her approval with (5) (a) [What a terrific idea!] than with (b) [What a divine idea!], which might cause raised eyebrows, and the reaction ‘That’s what we get for putting a woman in charge of this company.’”

Lastly, Lakoff tentatively posits the use of the intensive “so”, as in the example “(a) I feel so unhappy!”. While Lakoff states that using “so” in contexts other than as a superlative such as “very, really, utterly,” is more indicative of women’s speech, men “seem to have the least difficulty using the construction when the sentence is unemotional, or nonsubjective—without reference to the speaker himself,” as in the example “(d) Fred is so dumb!”. This would seem to fly in the face of female-coded word choices as described by Lakoff, as rather than “weaker” or “specific”, “so” is a vague intensifier; however, as we will see below with tag

questions, “so” can function as a hedge. As Lakoff puts it, “to hedge in this situation is to seek to avoid making any strong statement: a characteristic [...] of women’s speech.”

2. Situations in which certain syntactic rules are performed

According to Lakoff (1975), “when we leave the lexicon and venture into syntax, we find that syntactically, too, women’s speech is peculiar.” There exists the stereotype, for example, that women are the gender to use questions more frequently. This stems both from the idea that while men generally employ questions as requests for information, women more often use them as a rhetorical means of engaging with another’s conversational contribution or of acquiring attention from others conversationally involved. Although, according to Lakoff, there are no syntactic structures conventionally relegated to either gender in the same way as lexical choices, women are also said to employ more tag questions to verify or confirm information, or otherwise avoid making “strong” statements. Similarly, women’s use of questions promotes the same idea as their alleged higher use of minimal responses, in that questions are used in a more cooperative capacity, to collaborate with other interlocutors.

Other research on syntactic differences has found similar variations in preference or semantic use between men and women for the use of adverbial clauses as hedges (e.g., Aijmer, 1986) and their positioning (e.g., Mondorf, 2002), periods (e.g., Jespersen, 1922), gender-specific syntactic variations across languages (e.g., Johannsen, Hovy, & Søggaard, 2015), and so on.

3. Tag questions

Much of the research on syntax includes tag questions as a specifically female feature, or at least a more robust feature for females (e.g., Calnan & Davidson, 1998; Hepburn & Potter, 2011; Mondorf, 2011). Defining a tag question as something that is “used when the speaker is stating a claim, but lacks full confidence in the truth of that claim,” Lakoff (1975) offers the following examples:

- (7) Is John here?
- (8) John is here, isn’t he?

(7) fulfills the role of the “masculine” question, in that it is employed purely as a means of requesting information, and both yes and no are equally valid responses depending upon the answer and the responder’s knowledge of or devotion to the truth. While (8), on the other hand, might similarly be a request for information, the “feminine” tag “isn’t he?” initiates the preferred response of yes, again, as an avoidance of making any “strong” statement. According to Lakoff, “a tag question, then, might be thought of as a declarative statement without the assumption that the statement is to be believed by the addressee: one has an out, as with a question. A tag gives the addressee leeway, not forcing him to go along with the views of the speaker.”

However, a study by Freed and Greenwood (1996) showed that there was no significant difference in the use of questions between genders, and in writing, both genders use rhetorical questions equally as literary devices.

4. Intonational and other suprasegmental patterns

According to Lakoff (1975), “related to this special use of a syntactic rule is a widespread difference perceptible in women’s intonational patterns,” popularly referred to as “uptalk” (e.g. Tomlinson & Tree, 2011), “which has the form of a declarative answer to a question, and is used, as such, but has the rising inflection typical of a yes-no question, as well as being especially hesitant.” Lakoff offers the following exchange as an example:

- (13) (a) When will dinner be ready?
(b) Oh...around six o'clock...?

Such an inflected response, Lakoff argues, is typical of the more “polite” speech sounds expected of women over men—rather than positing a firm declarative when asked for information, women can again, in a sense, hedge their responses, “leaving a decision open, not imposing your mind, or views, or claims on anyone else.” Research has focused both on the gendered differences in production (e.g., Jiang, 2011), and perception (e.g., Hancock, Colton, & Douglas, 2014). As we will see in the section on minimal responses below, the intonation of such question-statements can be employed by interlocutors to elicit feedback markers from their listeners to ensure continuing cooperation.

5. Minimal response

According to Fellegly (1995), “minimal responses, in American English, are forms such as *mmhmm*, *yeah*, *uh-huh*, and *right* which are uttered by a listener during a speech event to signal a certain level of engagement with the speaker.” In other words, minimal responses are paralinguistic features associated with cooperative language, and the term is predominantly used in research about language and gender, usually synonymous, according to Fellegly, with assent terms, back channels, listener responses, feedbacks, accompaniment signals, and hearer signals.

According to Maltz and Borker (1982), women are more likely to use minimal responses as a sign of active listening, and not necessarily as a sign of agreement, and as such are very likely to use them more. Men on the other hand, according to Zimmerman and West (1996), generally use minimal responses for one of two reasons: only to show agreement, and not simply cooperative listening; or as a means of not sounding inattentive when they actually are. However, both men and women can employ minimal responses for their interactional function, rather than only for their gender-specific functions.

Fellegly (1995) found that while women “produce minimal responses strung out across each of the categories,” of continuous speech—end of turn, end of sentence within a turn, and elsewhere—men “predominantly produce minimal responses at the end of sentences while the speakers continue their turns.” Further, “statistically, men and women are significantly different from each other in their production of minimal responses at ends of turns; their production becomes similar at the end of sentences within turns; they diverge again at the production of minimal responses elsewhere.” Essentially, although women are far more likely to insert minimal responses at the end of a turn than men, and, to a lesser extent, far more likely to insert minimal responses elsewhere, both men and women are most likely to insert minimal

responses at the end of a sentence within a turn.

Fellegy points out a potential issue with these results:

is that men appear to use minimal responses less frequently and less attentively. Instead of appearing engaged throughout a conversation, men load up their responses at a specific syntactic unit. This can be particularly important for women who, with other women, more frequently receive a response from listeners at the end of turns. If in mixed-gender conversation these different patterns are upheld, it is understandable that an absence of end-of-turn responses could be interpreted by women as inattentiveness or the male listener's desire to extinguish the female speaker's topic of conversation.

The very fact that men and women use minimal responses in different syntactic structures dictates that women have more opportunity to employ them, "thus women appear to use minimal responses more frequently," especially in mixed-gender conversations.

Fellegy also offers up strategies by which speakers can elicit minimal responses from the listener—"the use of question statements (question intonation when making a statement) and the use of pauses; both are instances where the speaker is not yielding the floor but is, perhaps, checking in with the listener." Fellegy found that minimal responses, once again, occur exclusively at phrase boundaries, "coinciding with the end of a question statement or with a brief pause." However, Fellegy also found that men and women tend to respond to both pauses and question statements differently:

Male speakers in these groups use significantly more pauses and question statements as women, but male listeners do not respond to them frequently. Female speakers in these groups do not often use pauses or question statements, but female listeners more frequently respond when such devices are used. Once again, women and men appear to be operating in converse manners.

Essentially, "if these patterns are maintained in mixed gender conversations, women will be responding more to male speakers because males use more elicitation devices and because women respond more frequently to such elicitations," and vice versa. Again, it is not plainly that women use minimal responses *more*, but that they employ them in a larger variety of linguistic structures, and so have more opportunities to offer minimal responses to conversees. Research analyzing minimal responses and gender (e.g., Reid, 1995) has expanded to considering such avenues as how they frame scripted television shows (e.g., He, 2010) and live television (e.g., Pasfield-Neofitou, 2007), and even the language production of machines (e.g., Nass, Moon, & Green, 1997).

6. Turn-taking

According to DeFrancisco (1991), female linguistic behavior characteristically encompasses a desire to take turns in conversation with others, which is opposed to men's tendency toward centering on their own point or remaining silent when presented with such implicit offers of conversational turn-taking as are provided by hedges such as "y'know" and "isn't it". This, again, follows along with "the belief that women are particularly accomplished in the verbal arts of cooperation, empathy and rapport building, while men are more direct,

decisive, and authoritative communicators,” (Cameron, 2010).

As with many aspects of gendered language, however, the differences between male- and female-expected turn-taking strategies are often tied not only to gender alone, but also to ideas of power. Kendall and Tannen (1997) cite an example from Edelsky and Adams (1990), to illustrate the point that “women and men who do not conform to expectations for their gender may not be liked.” In Edelsky and Adams’ example:

A woman in the Arizona gubernatorial debate who was a long-time party insider, spoke in ways similar to how the men spoke in this and other debates: she took full turns that were out of turn, inserting some of these turns into otherwise orderly episodes, and she was the only woman to make a demeaning move and to engage in friendly repartee with the moderators. The authors note that by speaking in these ways, she was able to make the debate an equal forum; however, she was later lampooned in political cartoons and on local call-in talk shows for being ‘mannish’. Because she spoke in ways that were common in political debates but more typically associated with men, she was evaluated negatively in a way that the men were not.

As a result, Kendall and Tannen (1997) argue, “researchers suggest that language strategies that women use to downplay their authority are drawing on the resources available to them.” Thus, a woman employing male turn-taking strategies should expect to be perceived as overbearing or pushy, and a man employing female turn-taking strategies should expect to be perceived as weak or indecisive, and both should expect to be disliked. Research has focused on turn-taking in a variety of settings such as professional (e.g., Baker, 1991) and academic (e.g., Swann & Graddol, 1988), and on a variety of facets such as interruptions (e.g., Robinson & Reis, 1989) and overlapping talk (e.g., Schegloff, 2000).

7. Changing the topic of conversation

Dorval (1990) offers the following role differences between males and females at the dinner table: “The males provide entertaining stories emphasizing a sex-typed form of self-display. The females listen attentively and otherwise encourage the males. Thus, drastically different conversational participation is enacted by the males and females.” Similarly, a study of young American couples and their interactions revealed that while women raise twice as many topics as men, the topics that are more often picked up and elaborated upon are the men’s (Fishman, 1978), though Dorval found that, in same-sex interactions, males tended to change the subject more frequently than did females. As Goodwin (1990) observes, girls and women link their utterances to the utterances of previous interlocutors and develop one another’s topics rather than introduce new topics. This research, once again, shows indications of females as the more cooperative gender of interlocutors, who may appear to talk more based on how much they elaborate upon each topic, but whose topics are dropped in deference to any male interlocutors.

Tannen (1996) studied same-sex discourse at four age levels—grade two, grade six, grade ten, and twenty-five-year-olds—finding more consistencies within each gender than differences across the age groups. As Tannen (1996) states:

Whereas all the pairs displayed discomfort with the experimental situation and

the assigned task, at every age level, the female friends quickly established topics for talk and produced extended talk related to a small number of topics. In contrast, boys at the two younger ages produced small amounts of talk about many different topics. At the two older ages, the boys and men, like their female counterparts, produced a lot of talk about a few topics, but the level at which they discussed the topics was more abstract, less personal.

Rather than concluding these findings as evidence of males as not “engaged” or “involved” in conversations, Tannen (1990) subscribes “to a cross-cultural approach to cross-gender conversation by which women and men, boys and girls, can be seen to accomplish and display coherence in different but equally valid ways.”

Goodwin’s (1990) study of the same-sex and cross-sex discourse of young black children found similar differences in the conversational styles of same-sex interlocutors based on their gender pairs. While boys were more likely to use many commands as they played in groups (such as “give me your hanger”), girls playing in groups were more likely to use directives such as “let’s” and “we gotta” (such as “we gotta find some more bottles”). However, Goodwin also found that in certain contexts, such as playing house, the girls shifted to the type of hierarchical group organizations the boys more commonly used, rather than the more typically expected relational and cooperative strategies. And when boys and girls interacted together, the girls proved just as capable of holding their own against the boys in arguments or verbal contests. Research has focused on topic change and gender overall (e.g., Okamoto & Smith-Lovin, 2001), and in specific settings (e.g., Ainsworth-Vaughn, 1992) and its relation to power. As with many of the features covered, although there are certain styles deemed more typical of either gender, either gender is equally capable of employing whichever strategy they choose should the situation arise.

8. Self-disclosure

Scholars define self-disclosure as the sharing of information with others that they would not normally know or necessarily discover throughout the course of natural conversation. Self-disclosure, such as sharing problems and experiences with other interlocutors, often as a way of offering sympathy or empathy, involves risk and vulnerability on the part of the person sharing the information. According to popular theory, female tendencies toward self-disclosure contrast with male tendencies toward non-self-disclosure and professing advice or offering a solution when confronted with another’s problems. This can also vary depending on the gender of an audience (e.g., Leaper, Carson, Baker, Holliday, & Meyers, 1995), and more recently much research has been conducted on gender and self-disclosure online (e.g., Cho, 2007; Jaidka, Guntuku, & Ungar, 2018; Kim & Dindia, 2011).

Tannen (1991) contrasts the difference between women’s tendency toward self-disclosure, and men’s tendency toward non-self-disclosure as, instead, in lieu of giving advice or offering a solution. Women, argues Tannen, engage in self-disclosure as a form of rapport-building between other interlocutors—that is, in a same-sex conversation, if the first woman discloses personal information, the cooperative, feminine thing to do would be for the other

woman or women present to disclose something equally personal to show conversational alignment. So “women are frustrated when they not only don’t get this reinforcement but, quite the opposite, feel distanced by the advice, which seems to send the metamessage ‘We are not the same. You have the problems; I have the solutions.’”

However, Dindia and Allen (1992), in conducting a meta-analysis of self-disclosure studies, concluded that “sex differences in self-disclosure are not as large as self-disclosure theorists and researchers have suggested,” as, while there are some differences, these differences are small, or at least certainly smaller than many researchers would suggest. In example of this, Dindia and Allen found that “using the average effect size found in this meta-analysis, if approximately 45% of men would disclose a particular item, approximately 55% of women would disclose the same information.” As with minimal responses, it is not that women necessarily self-disclose more, but that they self-disclose in a larger variety of contexts, once again making it seem that women over-share, while men under-share.

9. Verbal aggression

Though the basis of much past research on aggression was primarily built upon the idea that females were non-confrontational, newer research has shown that while “boys tend to be more overtly and physically aggressive, girls are more indirectly, socially, and relationally aggressive,” (Cupach & Brian, 2011). According to Blake, Eun Sook, & Lease (2011), “collectively, *indirect, relational, and social aggression* refers to types of aggression that inflict harm by damaging the interpersonal relationships, self-esteem, and social status of victims through exclusionary and socially manipulative tactics.” Blake et al. describe indirect aggression as “attacks delivered covertly [...] the use of circuitous means to disguise the aggressor’s intention to cause harm to victims,” relational aggression as “the extent to which an aggressor is able to disrupt the interpersonal relationships of their victims in order to cause physical injury,” and social aggression as “nonconfrontational forms of aggression, although it has assumed different meanings when employed by independent research teams.”

Indirect aggression—such as gossip, exclusion, or the ignoring of the victim—relational aggression—such as a threat to terminate a friendship or the spreading of false rumors—and verbal aggression in general are all thought to be more feminine forms of aggression. This is because, according to Blake et al., “for example, girls are hypothesized to be more socially aggressive than boys because of the cohesiveness and greater disclosure and intimacy involved in female friendships.” Masculine forms of aggression, on the other hand, are thought to be more overt and physical in nature. Social aggression, which “is directed toward damaging another’s self-esteem, social status, or both, and may take direct forms such as verbal rejection, negative facial expressions or body movements, or more indirect forms such as slanderous rumors or social exclusion,” on the other hand, is a type of aggression more common in both male and female adolescent behavior (Balter, 1999).

Blake et al. (2001) found that there is “preliminary evidence that gender differences exist in children’s use of nonverbal social aggression, but additional research is needed to

further support these findings.” They cite a 2004 study by Underwood and colleagues, which found “both boys and girls engaged in verbal and nonverbal exclusionary behavior [...] but girls were observed to be more nonverbally socially aggressive than were boys.” Blake et al. explain that “if girls are more likely to monitor and interpret the nonverbal communication of their peers than are boys, it is plausible that girls may rely on nonverbal exclusionary patterns more so than boys do when girls choose to aggress against their same-gendered peers.” Studies in verbal aggression and gender have often focused on its conjunction with turn-taking (e.g., Bresnahan & Cai, 1996).

10. Listening and attentiveness

Listening and attentiveness, turn-taking, changing the topic of conversation, and minimal response are often seen to go hand in hand with regards to how many scholars expect males and females to behave, in that the listening and attentiveness of interlocutors is often gauged by how often and who interrupts to change a topic out of turn, or provides only silence or minimal feedback in response. Zimmerman and West (1996) take the position that “interruptions are a violation of the current speaker’s right to complete a turn, or more precisely, to reach a possible transition place in a unit-type’s progression.” Although a possible response to an interruption is a complaint—“i.e. a formulation of a speaker’s previous utterance as a certain kind of act. Such a complaint could be: ‘You just interrupted me’ or, in the case of a series of such acts, ‘You keep interrupting me’”—Zimmerman and West found the most common response to an interruption, and occasionally to overlaps, especially by females in cross-sex interactions, was a period of silence.

In two-person, same-sex interactions, Zimmerman and West found that speakers interrupt and overlap with one another roughly equally, with around three times as many overlaps occurring as interruptions. Contrastingly, in two-person, cross-sex interactions, interruptions occurred over five times more than overlaps, and 96% of them were in the form of males interrupting the female interlocutors. However, Zimmerman and West also found that only 18% of the cross-sex conversations contained 66% of the interruptions, “thus if the distribution of overlaps across the segments is construed as evidence of clustering, we would have to conclude that the pattern is essentially identical for both cross-sex and same-sex pairs.”

DeFrancisco (1991) found that “there are two general findings which lead to the conclusion that men were relatively silent and that their behaviors silenced the women.” The first is that “the no-response was the most common turn-taking violation, particularly for the men.” As Zimmerman and West (1996) found, however, even when men do offer responses in cross-sex conversations, when these responses are minimal responses, oftentimes “these are retarded beyond the end of the utterance,” serving as just as clear if not a clearer indication that male conversants are not being attentive listeners. Second, “results from the components of conversation combined with the ethnographic information strongly suggest that women in this project worked harder to maintain interaction than men, but were less successful in their

attempts.” According to DeFrancisco (1991), “together these findings reveal the multiple ways in which these women have been silenced.”

DeFrancisco’s study found several violations of turn-taking that indicated inattentiveness: “no-response, 32 percent women, 68 percent men; interruption, 46 percent women, 54 percent men; delayed response, 30 percent women, 70 percent men; and minimal response, 40 percent women, 60 percent men [...] Among the total violations, women were responsible for 36 percent, and men were responsible for 64 percent.” Aside from interruptions, men significantly out-violate women in every other category. Although DeFrancisco concludes that “men generally silenced the women,” she points out a second reason for this: women talk more, “139 minutes in total, 63 percent; men spoke 83 minutes, 37 percent,” with more topics raised, and less turn-taking violations. As discussed above, though women may offer more topics, it is the men’s topics—which are more likely to be interruptive—which are more often taken up.

Again, women can be seen to place more weight on the importance of cooperation in conversations, and men can be seen to be the more dominant interlocutors. This comes down to the structures in which men and women are expected to employ various communicative strategies. If, in cross-sex interactions, women are talking more, it may be because men expect they should be offering less agreement in the form of minimal responses or otherwise. If, in cross-sex interactions, men are being less attentive and cooperative, it may be because women expect feedback in the form of agreement, and not advice or topic changes. Simultaneously, when women offer agreement in the form of minimal responses or otherwise to their male counterparts, men may be more likely to take this as actual, literal agreement encouraging them to continue or develop their own topics, and not the cooperative strategies women may have intended them as. DeFrancisco chalks these differences up to “different preferences for how to avoid conflict [which] may themselves come into conflict.”

11. Dominance versus subjection

Coates (2013) offers the example of humor, which “has three main functions: first, humour can emphasise power differences; second, humour can provide self-protection; and third, it can be used to create or maintain solidarity within the group.” However, Coates goes on to point out that these usages are not equally distributed across the genders, as “male speakers use humour as a way of exerting dominance, female speakers use humour as a form of self-protection, and both male and female speakers use humour to create solidarity.” In this instance, although males and females can employ the same strategies—humor, for example—they do so for different reasons, or at least to differing results.

Coates goes on to draw on the presence and lack of collaborative conversation strategies in various contexts, differentiating their implications based on “conversations involving a one-at-a-time floor and those involving a collaborative floor.” As seen in the section above, violating certain expectations of listening and attentiveness—such as interrupting,

disrupting proper turn-taking, and so on—can have different implications for male versus female interlocutors. So, too, can context, or floor, dictate the implications. Coates offers the example of overlaps, in which “where the floor is jointly owned by all participants, then overlapping speech is an inevitable consequence [...] but where a one-at-a-time floor is operating, any overlap is potentially a violation of the current speaker’s turn at talk, specifically of their right to speak.” So if men interrupt more or expect to engage in a collaborative floor more often, and if women are more likely to eschew collaboration for politeness (as we will see below), or expect to engage in a one-at-a-time floor more often, in cross-sex interactions, men will naturally dominate.

According to Coates, a variety of research concluded with “men’s greater usage of certain strategies being associated with male dominance in conversation,” including such domains as “in the classroom [...] at the doctor’s [...] in internet chat rooms [and] in the home.” An explanation for these differing strategies, as Coates claims, is that “key components of contemporary hegemonic masculinity are hardness, toughness, coolness, competitiveness, dominance and control,” and even if most men do not fit within or subscribe to this ideal, “the concept of hegemonic masculinity ‘captures the power of the masculine ideal for many boys and men’ (Frosh et al. 2002)”. Essentially, because women are typically less concerned with power and more concerned with maintaining relationships than are men, women’s communication—such as in the way women employ humor, the conversations in which women do or do not expect overlap or interruption, and the relationship- and emotion- rather than self-centered stories they tell—is more focused on building and maintaining relationships.

12. Politeness

Much research has been conducted on politeness as a gendered language feature (e.g., Mills, 2002, 2003) and its conjunction with power (e.g., Eliasoph, 1987; Mullany, 2004). Politeness is often concerned with ‘face’, or “the public self-image that every member wants to claim for himself,” which can be preserved or threatened (Brown & Levinson, 1987). Lakoff (1975) outlines three forms of politeness: formality, or to “keep aloof”; deference, or to “give options”; and camaraderie, or to “show sympathy”. As examples of formal politeness, Lakoff offers medical and legal jargon that use technical terminology to remain aloof (“*carcinoma* rather than *cancer*”), passive voice, “the academic-authorial *we*”, hypercorrect forms, and the impersonal pronoun “*one*”. As examples of deferential politeness, which “may be used alone or in combination with either of the other two rules, Lakoff offers hesitancy in speech, question intonation and tag questioning, hedging, and euphemism. And as examples of camaraderie politeness, Lakoff offers the nonlinguistic device of back-slapping, colloquial language, “the use of four-letter words” rather than the technical term or euphemism, and the use of nicknames.

While these types of politeness can be combined with varying degrees of success depending upon the context and intent of the speaker, Lakoff points out, “how you categorize a particular act may determine whether it is to be considered polite according to one rule, or

rude according to another.” Similarly, how speakers categorize politeness types across gender lines can color how polite or impolite an interlocutor appears, depending upon their own gender and which politeness rules they are following. Women’s language, according to Lakoff (1974), is primarily comprised of formal and deferential politeness, “establishing and reinforcing distance: deferential mannerisms coupled with euphemism and hypercorrect and superpolite usage,” whereas men are exempt from worrying about such politeness rules other than, perhaps, camaraderie, as “in a man it’s ‘just like a man,’ and indulgently overlooked unless his behavior is really boorish.”

Sung’s 2012 study of gender, discourse, and impoliteness in the workplace found much the same result, stating, “it has been shown that gendered assumptions play a role in the assessment of (im)polite behaviours.” Even in the workplace, where the power-differential is purportedly equal across genders:

The negative evaluation of [a woman’s] use of relatively masculine behaviour could be explained by gender stereotyping that comes into play when assessments of (im)politeness are being made. As a result of gender stereotyping [her] relatively masculine behaviour, albeit politic in this masculine context, is assessed as impolite, inappropriate, and ‘negatively marked (Locher and Watts 2005). Also, the case study seems to raise the issue of a double bind (Lakoff 1975, 1990, Tannen 1994) that women confront on a regular basis at work, and this double bind may be explained by the coexistence of two sets of conflicting norms that women need to adhere to: the norms of the ‘masculine’ work-place as well as the stereotypical expectations that women should be polite.

While Sung admits that this study, conducted on politeness assessments made on episodes of the television show *The Apprentice*, may not be wholly representative of real-world workplaces, these findings seem to hold true across scholars. Although both men and women are equally capable of availing themselves of any type of politeness—and increasingly *do*—we are still societally conditioned to assess these strategies based not only on appropriateness of context, but on adherence to stereotypical gender expectations.

B. Overview

The table below demonstrates the overall distribution of features in the above section.

	Feature	Source(s)
Lexical Difference	choice and frequency	Lakoff, 1975
	in conversations	Singh, 2001
	over the phone	Boulis & Ostendorf, 2005
	on social media	Bamman, Eisenstein, & Schnobelen, 2014
Syntactic Rules	use of adverbial clauses as hedges	Aijmer, 1986
	position of adverbial clauses	Mondorf, 2002
	periods	Jespersen, 1922
	gendered syntactic variations across languages	Johannsen, Hovy, & Søgaard, 2015
Tags	as a female-coded feature	Canlan & Davidson, 1998; Hepburn & Potter, 2011; Mondorf, 2011
	no significant gender difference	Freed & Greenwood, 1996
Intonation	uptalk (as a female-coded feature)	Tomlinson & Tree, 2011
	gendered differences in production	Jiang, 2011
	gendered differences in perception	Hancock, Colton, & Douglass, 2014
Minimal	as features of cooperative language	Felleggy, 1995
	as a sign of active listening for women	Maltz &orker, 1982

	minimal responses and gender	Reid, 1995
	on scripted television shows	He, 2010
	on live television	Pasfield-Neofitou, 2007
	in the language production of machines	Nass, Moon, & Green, 1998
Turn-Taking	women are cooperative, men less-so	DeFrancisco, 1991; Cameron, 2010
	genders must conform to precepts to be liked	Kendall & Tannen, 1997; Edelsky & Adams, 1990
	in professional settings	Baker, 1991
	in academic settings	Swan & Graddol, 1988
	as interruptions	Robinson & Reis, 1989
	overlapping talk	Schegloff, 2000
Topic	males provide topics, females listen attentively	Dorval, 1990; Fishman, 1978
	females link utterances to previous utterances	Goodwin, 1990
	genders accomplish coherence differently	Tannen, 1990
Self-Disclosure	varies on an audience's gender	Leaper, Carson, Baker, Holliday, & Meyers, 1995
	gender and self-disclosure online	Cho, 2007; Jaidka, Guntuku, & Ungar, 2018; Kim & Dindia, 2011
	style differences in cross-gender talk	Tannen, 1991
	women self-disclose in more contexts	Dindia and Allen, 1992
Verbal Aggression	males are physically aggressive, females socially	Cupach & Brian, 2011
	indirect, relational, and social aggression	Blake, Eun Sook, & Lease, 2011
	social aggressions more common for adolescents	Balter, 1999
	in conjunction with turn-taking	Bresnahan & Cai, 1996
Listening	interruptions as violations	Zimmerman & West, 1996
	men silence more, women talk more	DeFrancisco 1991
Dominance		Coates, 2013
	hegemonic masculinity	Coates, 2013; Frosh et al., 2002
Politeness	as a gendered language feature	Mills, 2002; Mills, 2003
	in conjunction with power	Eliasoph, 1987; Mullany, 2004
	as concerned with 'face'	Brown & Levinson, 1987
	formality, deference, and camaraderie politeness	Lakoff, 1974; Lakoff, 1975
	in the workplace	Sung, 2012

Table 2.2 Overview of Lakoff 1975 proposed features and additional research

The next section, which goes on to outline additional proposed features of gendered, written, and often interactional language, does so on the basis not of historical sociolinguistic research on largely spoken language as outlined in this section, but rather on the basis of large-data, computational, and often statistically significant findings. Although often distinct, many of the features in Part 4 can nevertheless be seen to have derived from earlier research such as that outlined here in Part 3, at least in their conception if not also their findings and application. As such, where possible, the features outlined here are aligned with the features outlined in the section below and expressed in how they comport or deviate from the below findings of computational approaches on largely online data.

Part 4 – Computational Approaches

The field of internet linguistics is a vast one (e.g., Androutsopoulos, 2006; Crystal, 2001, 2005, 2011), with much research applied to online genres (e.g., Herring, Scheidt, Bonus, & Wright, 2004; Thurlow & Poff, 2013; Wright 2013, 2013, 2014), author identification online (e.g., Calix, Connors, Levy, Manzar, McCabe, & Westcott, 2008; Juola, 2008; Koppel, Schler, J. Ford, PhD, Thesis, Aston University, 2022

& Argamon, 2009; Narayanan, Paskov, Gong, Bethencourt, Stefanov, Shin, & Song, 2012; Stamataatos, 2008, 2009; Vorobeve, 2016) and gender identification more specifically (e.g., Bamman, Eisenstein, & Schnobelen, 2014; Cheng, Chandramouli, & Subbalakshmi, 2011; Schler et al., 2006). While the previous section relied largely on historical observations and findings as to categories of gendered and often spoken language, the features in this section are entirely computational, and are usually conducted on large quantities of CMC, such as the PAN studies overviewed in Part 2.

A. Proposed features of Gendered language

In their analysis of a collection of tweets, Bamman et al. (2014) who use the standard repertoire of computational techniques to explore both essentialist and interactionist aspects of identity performance, found that the features summarized in their analysis could indeed predict categorical gender as indicated by language-independent details on tweeters' profiles. Bamman et al. (2014) wish to recognize the interactionist complexity of such analyses, pointing out that in conventional studies:

there is often an implicit assumption that linguistic choices are associated with immutable and essential categories of people. Indeed, strong aggregate correlations between language and such categories enable predictive models that are disarmingly accurate. But this gives an oversimplified and misleading picture of how language conveys personal identity. (2014)

Bamman et al. (2014) found that linguistic performance of gender was influenced by the gender makeup of an individual tweeter's social group. Thus, one finding was that males with more female online contacts than male tended to exhibit more female-coded features, and conversely identified female tweeters with more male interactants used a more 'male language style'. Such a finding is unsurprising for the interactionist model of identity which would recognize the effects of linguistic accommodation and the reciprocal drawing on our fellow interactants as discursive resources in interaction (e.g., Grant & MacLeod, 2018). While this thesis does not seek to verify the features as found by Bamman et al. (2014), their findings are applied here as a measure by which to test any potential change in linguistic performance when an interactant's perception is altered.

The following analysis examines gendered language differences as found by both Schler et al. (2006) in the original BAC analysis, and as found by Bamman et al. (2014), in order to determine which and to what extent their suggested features hold up as potential gender differentiators when using the edited BAC (henceforth eBAC).

Feature	female	male	Schler
Pronouns	Including alternative spellings such as <i>u</i> , <i>ur</i> , <i>yr</i>	—	female
Emotion terms	All, including <i>sad</i> , <i>love</i> , <i>glad</i> , etc.	—	female
Emoticons	All, including <i>:)</i> , <i>:D</i> , and <i>;) </i>	—	female
Kinship terms	Most, including <i>mom</i> , <i>mommy</i> , <i>sister</i> , <i>daughter</i> , <i>aunt</i> , <i>auntie</i> , <i>grandma</i> , <i>kids</i> , <i>child</i> , <i>dad</i> , <i>husband</i> , <i>hubs</i> , etc.	<i>wife</i> , <i>wife's</i>	female
Friendship Terms	<i>bestie</i> , <i>bff</i> , <i>bffs</i>	<i>bro</i> , <i>bruh</i> , <i>bros</i> , <i>brutha</i>	female
Abbreviations	<i>lol</i> , <i>omg</i>	—	female

Punctuation	..., !, ?	—	female
Expressive lengthening	e.g., <i>coooooooooo!</i>	—	—
Backchannel sounds	<i>ah, hmm, ugh, and grr</i>	—	—
Hesitation words	<i>Um and umm</i>	—	—
Assent terms	<i>okay, yes, yess, yesss, yessss</i>	—	female
Negation terms	<i>noo, noooo, and cannot</i>	<i>nah, nobody, and ain't</i>	female
Swears and taboo words	Anti-swear <i>darn</i>	Most, excluding anti-swears	—
Prepositions	—	—	male
Alternative spellings	<i>w/a, w/the, w/my for with</i>	2 for to	female
Conjunctions	&	—	—
Articles and determiners	—	—	male

Table 2.3 Gendered features as suggested by Bamman et al. (2014) & Schler et al. (2006)

Schler et al. (2006), in their original analysis of the BAC, consider several features overlapping with Bamman et al.'s (2014), with similar results. These include pronouns, assent terms, negation terms, determiners, and prepositions, while the category provided by Schler et al. (2006) of “blog words” encompasses abbreviations and some alternative spellings. Schler et al. (2006) also consider other word categories provided by LIWC word classes, including money, job, sports, and tv, which they find to be male coded, and sleep, eating, sex, family, friends, and emotions, which they find to be female coded, some of which overlap with remaining Bamman et al. (2014) categories.

Though Schler et al. (2006) identified few individual terms most indicative of either the female- or male-coded feature variations, they did identify overall frequency per 10,000 words, indicated by the rightmost column below. Additional gender-distinguishing features considered by Schler et al. (2006) include the use of hyperlinks (male-prevalent) and overall post length (with longer posts being female-prevalent). The following walks through the collective features of Bamman et al. (2014) and Schler et al. (2006), and takes into consideration how previously considered features as outlined in Part 3 might best be considered to align with these features, as well as providing examples other contemporary computational analyses which considered these same features, and how their findings did or did not comport and why.

1. Pronouns

Bamman et al. note that “**Pronouns** are generally associated with female authors, including alternative spellings *u, ur, yr.*” (2014). According to Newman et al., “One striking result reported by Mehl and Pennebaker (2003) was that women were more likely to use first-person singular,” yet they also note that “The result is also at odds with a review by Mulac et al. (2001), which cited findings that men used first-person singular more often. However, their conclusion was based on only two studies,” but finally that, “if the entire category of personal pronouns is considered, women frequently are the higher users (e.g., Gleiser et al., 1959; Mulac & Lundell, 1986),” (2008). Contrastingly, Cheng et al. found that “men’s conversational patterns usually express ‘independence’ and assertions of vertically hierarchical power, so they use more first-person singular pronouns like I” (2011). Multiple recent PAN studies (e.g.,

Rashkin et al., 2017) found the number of pronouns to show gain for profiling.

2. Emotion terms

Bamman et al. note that “**emotion terms** (*sad, love, glad, etc.*) [...] that appear as gender markers are associated with female authors” (2014), and Schler et al. (2006) found that this female-coded trend held true for both positive and negative emotion terms. Other research has found that, “for example, women make frequent use of emotionally intensive adverbs and affective adjectives such as *really, very, quite and adorable, charming, lovely,*” though these distinctions are not made here (Cheng et al., 2011). Multiple recent PAN studies (e.g. Ghanem et al., 2020; Giachanou et al., 2020; Guo et al. 2019,) considered emotions as useful for classification.

3. Emoticons

Bamman et al. note that “**emoticons** that appear as gender markers are associated with female authors, including some that the prior literature found to be neutral or male: :) :D and ;)” (2014). Schler et al. (2006) include emojis as part of their “blog words” category, finding that blog words, generally, are more prevalent in the original BAC’s female subcorpus. Since the standardization of emojis (and then subsequent addition of non-standard varieties thanks to services like Discord and Twitch allowing for custom emojis), much research has been paid to the place of emojis in CMC, including their use as, for example, non-standard, emotive punctuation (e.g., Provine, Spencer, & Mandel, 2007). Multiple recent PAN studies have found both emoticons (e.g. Fahim et al, 2019) and emojis (e.g. Manna et al., 2020; Spezzano et al., 2020), classified into various categories, to be a useful feature in profiling.

4. Kinship terms

Bamman et al. note that “Of the **kinship terms** that are gender markers, most are associated with female authors (*mom, mommy, sister, daughter, aunt, auntie, grandma, kids, child, dad, husband, hubs, etc.*). Only a few kinship-related terms are associated with male authors—*wife, wife’s, bro, bruh, bros and brotha*” (2014). Similarly, Schler et al. (2006) consider the LIWCs category of “family” words and find that some have the “greatest information gain for gender”, including the kinship terms *mom, mommy, boyfriend, husband, and hubby* for their female data. The closest reference in PAN studies (e.g. Gencheva et al., 2016) falls to the “specific sentences per gender (e.g. ‘my wife’, ‘my man’, ‘my girlfriend’...)” (Rangel et al., 2016).

5. Friendship terms

Bamman et al. note that “Only a few kinship-related terms are associated with male authors—*wife, wife’s, bro, bruh, bros and brotha*—though many of these may be better described as friendship terms, with corresponding female markers *bestie, bff, and bffs (best friends forever)*” (2014). Schler et al. (2006) also find that “friends” as a class of words taken

from LIWC classes is a significantly female-coded category. Both friendship and kinship terms above likely fall into the general PAN category of dictionary-based words, as opposed to, for example, slang words or named entities, or topic-based words (e.g. Rangel et al., 2013).

6. Abbreviations

Bamman et al. note that, “Several **abbreviations** like *lol* and *omg* appear as female markers” (2014). Although Schler et al. (2006) do not consider abbreviations as a particular feature, they do consider a category of blog words which they define as “neologisms - such as *lol*, *haha* and *ur* - that appear with high frequency in the blog corpus”, which they additionally find to be female coded. Multiple recent PAN studies have included abbreviations (e.g. Raiyani et al., 2018; Stout et al., 2018).

7. Punctuation

Bamman et al. (2014) based their inclusion of punctuation as a category in their research on Rao, Yarowsky, Shreevats, & Gupta, (2010) who “found that women used more [...] ellipses (...), [...] complex punctuation (!! and ?!),” than do men. Lakoff (1975) and others (e.g., Mulac, Erlandson, Farrar, Hallet, Molloy, & Prescott, 1998) suggest that questions are a more female-frequent feature, perhaps owing to attempts at deferential politeness. Contrastingly, Freed and Greenwood (1996) found no difference in question use, which may be indicative that Lakoff’s observations do not hold true for online posts and communications, either because they are written, or because they are not as conversational as other mediums. PAN studies (e.g. Cardaioli et al., 2020; Russo, 2020; Spezzano et al., 2020) regularly include punctuation as a style feature for profiling tasks.

8. Expressive lengthening

Bamman et al. note that “expressive lengthening (e.g., *coooooooooo!*), [...] appear as female markers” (2014). Such terms would likely fall into the category of non-dictionary-based words or “correctness” (e.g., Pimas et al., 2016) for PAN studies, along with backchannel sounds and hesitation words below.

9. Backchannel sounds

Bamman et al. note that “backchannel sounds like *ah*, *hmmm*, *ugh*, and *grr* [...] appear as female markers” (2014). These would fall into the category of minimal responses as described by Fellegly (1995), which Maltz and Borker (1982) found to be more female coded.

10. Hesitation words

Bamman et al. note that “Hesitation words *um* and *umm* are also associated with female authors” (2014) and may be features of female speech which Lakoff (1975) refers to as deferential politeness, as they can indicate active listening by an interlocutor.

11. Assent terms

Bamman et al. note that “The assent terms *okay*, *yes*, *yess*, *yesss*, *yessss* are all female markers, though *yessir* is a male marker” (2014), while Schler et al. (2006) also found

that assent terms as a general category are more frequently used by females than males. As with the above two features, assent terms can be considered minimal responses (Felleggy, 1995; Maltz and Broker, 1982), and part of deferential politeness when they are indicators of active listening (Lakoff, 1975), both of which would leave them as female coded.

12. Negation terms

Bamman et al. note that “Negation terms *nooo*, *noooo*, and *cannot* are female markers, while *nah*, *nobody*, and *ain’t* are male markers” (2014), while Schler et al. (2006) found that negation terms overall were more frequent in their female dataset. As with assent terms above, negation terms likely fall under the same topic-based categories as kinship and friendship terms in PAN studies.

13. Swears and taboo words

Bamman et al. note that “Swears and taboo words are more often associated with male authors; the anti-swear *darn* is a female marker. This gendered distinction between mild and strong swear words was previously reported by McEnery (2005)” (2014). The closest Schler et al. (2006) come to considering taboo words (though not necessarily swears) is in their inclusion of the LIWC word class of “sex”, which they find to be a female-coded category. Lakoff (1975) expects “weaker” expletives and anti-swears from females as a means of deferential politeness, and “stronger” expletives from males as a form of camaraderie politeness. PAN studies regularly include swear words (e.g. Rashkin et al., 2017) and slang words more generally (e.g., Bouazizi & Ohtsuki, 2017; Patra et al., 2018).

14. Prepositions

Bamman et al. note that their “analysis did not show strong gender associations for standard prepositions, but a few alternative spellings had strong gender associations: 2 (a male marker) is often used as a homophone for *to*; an abbreviated form of *with* appears in the female markers *w/a*, *w/the*, *w/my*,” (2014). Schler et al. (2006), on the other hand, found that prepositions as a stylistic as opposed to content category were more prevalent in their male than female subcorpus. According to Newman et al.’s findings, “Men exceeded women on a number of linguistic dimensions including [...] prepositions,” (2008).

15. Alternative spellings

Bamman et al. note that “The word classes defined in prior work failed to capture some of the most salient phenomena in our data, such as the tendency for [...] non-standard spelling to be used more frequently by women (*vacay*, *yaay*, *lol*),” (2014). This feature in particular is recurring throughout Bamman et al. (2014), who elsewhere include alternative spellings for both pronouns (*u*, *ur*, *yr*) and prepositions (2, *w/*). Multiple other alternative spellings are considered by Bamman and not called as such. The more male-coded alternative spellings are the friendship terms *bro*, *bruh*, *bros*, *brutha*, etc., and the negation terms *nah* and *ain’t*. More female-coded alternative spellings include various kinship terms such as *hubs*, friendship

terms such as *bestie*, and expressive lengthening in general (specifically with regards to backchannel sounds, hesitation words, assent terms, and negation terms).

16. Conjunctions

Bamman et al. note that “The only **conjunction** that displays significant gender association is &, associated with female authors,” (2014). Along with prepositions above, and articles and determiners below, conjunctions are often considered in PAN studies that make general use of POS tagging (e.g. Hortenhuemer & Zangerle, 2020; Karlgren et al., 2018).

17. Articles and determiners

While Bamman et al. note that “No **articles or determiners** are found to be significant markers,” (2014), Schler et al. (2006) find that “male bloggers use more articles” (2006), and, as mentioned above, that they found such style-based features to be more reliable than other content-based features. Other studies such as Newman et al. (2008), however, cite that, “Men have been found to [...] use more articles (e.g., Gleiser, Gottschalk, & John, 1959; Mehl & Pennebaker, 2003; Mulac & Lundell, 1986).”

18. Post length

Although not considered by Bamman et al. (2014) (likely because tweets have length restrictions, which are additionally much shorter than a potential blog length, for example), Schler et al. (2006) found post length to be a distinguishing feature, with longer posts for females (213.0) than for males (201.0). Other research (e.g., Cheng et al., 2011) has found that other structural distributions, such as paragraph length, sentence length, sentences per paragraph, blank line size and distribution, and so on can be useful authorial markers. A variety of PAN studies also include features of tweet length (e.g. Buda & Bolonyai, 2020; Johansson, 2019), word length (e.g. Hortenhuemer & Zangerle, 2020; Labadie et al., 2020) and rely on previous studies of the average length of sentences (e.g. Goswami et al., 2009).

19. Hyperlinks

Schler et al. (2006) found that “male bloggers use more hyperlinks than do female bloggers” (2006). This feature was also not considered by Bamman et al. (2014), but again this is likely due to the limitations on tweets, meaning that hyperlinks in general take up valuable tweet real-estate, and as such likely conform to different frequencies and restrictions on Twitter than in other modes of CMC. Many PAN studies rely on the use of links as a feature (e.g., Alrifai et al., 201; Giachanou & Ghanem, 2019) as well as the number of URLs (e.g. Manna et al., 2020), along with other HTML-based features that are specific to CMC.

B. Overview

The table below provides an overview of research that considered each of the features discussed in this section. These 19 features, along with the broader application of keywords, are applied throughout this paper, and their exploration and application is further detailed in Chapter 3 below.

	Feature	Source(s)
Pronouns	women might use more first person singular	Newman et al., 2008
	men use more first person singular	Cheng et al., 2011
	number of uses in profiling	Rashkin et al., 2017
Emotion terms	women use intensive adverbs and affective adjectives	Cheng et al., 2011; Lakoff, 1975
	LSTM-memory network-based emotion classifications in profiling	Ghanem et al., 2020; Giachanou et al., 2020; Guo et al., 2019
Emoticons	emoticons as punctuation	Provine et al., 2007
	emoticon and emoji use	Fahim et al., 2019; Manna et al., 2020; Spezzano et al., 2020
Kinship terms	specific sentences per gender	Gencheva et al., 2016
	as dictionary- and topic-based words	Rangel et al., 2013
Friendship terms	as dictionary- and topic-based words	Rangel et al., 2013
Abbreviations	as a profiling feature	Raiyani et al., 2018; Stout et al., 2018
Punctuation	women use more questions/tags	Lakoff, 1975; Mulac et al., 1998
	no difference in question use	Freed and Greenwood, 1996
	as a profiling feature	Cardaioli 2020; Russo, 2020; Spezzano et al., 2020
Expressive lengthening	non-dictionary-based words and “correctness”	Pimas et al., 2016
Backchannel sounds	as minimal responses, female coded	Felleggy, 1995; Maltz & Borker, 1982
	non-dictionary-based words and “correctness”	Pimas et al., 2016
Hesitation words	women use as deferential politeness	Lakoff, 1975
	non-dictionary-based words and “correctness”	Pimas et al., 2016
Assent terms	women use as minimal responses	Felleggy, 1995; Maltz & Broker, 1982
	women use as deferential politeness and active listening	Lakoff, 1975
	as dictionary- and topic-based words	Rangel et al., 2013
Negation terms	as dictionary- and topic-based words	Rangel et al., 2013
Swears and taboo words	women use “weaker”, men use “stronger”	Lakoff, 1975
	as swear words	Rashkin et al., 2017
	as slang words	Bouazizi & Ohtsuki, 2017; Patra et al., 2018
Prepositions	in POS tagging	Hortenuemer & Zangerle, 2020; Karlgren et al., 2018
Alternative spellings	non-dictionary-based words and “correctness”	Pimas et al., 2016
Conjunctions	in POS tagging	Hortenuemer & Zangerle, 2020; Karlgren et al., 2018
Articles and determiners	men use more articles	Newman et al., 2008
	in POS tagging	Hortenuemer & Zangerle, 2020; Karlgren et al., 2018
Post length	paragraph length, sentence length, sentences per paragraph, blank line size and distribution, etc., can be used as authorial markers	Cheng et al., 2011; Goswami et al., 2009
	tweet length	Buda & Bolonyai, 2020; Johansson, 2019
	word length	Hortenuemer & Zangerle, 2020; Labadie et al., 2020
Hyperlinks	as a feature of use	Alrifai et al., 201; Giachanou & Ghanem, 2019
	as a count of use	Manna et al., 2020

Table 2.4 Overview of 20 proposed features and additional research

Part 5 – Discussion/Overview

As outlined here in Part 2, and above in Chapter 3, approaches to the study of gendered

language have shifted ideologically over the course of sociolinguistic research. Contemporary approaches seldom stop at considering gender as a binary, or as a facet-internal deviation, instead considering gender identity as performative and interactional, relational and intersectional, and in conjunction with other identity facets such as race and age, but most specifically sexuality and gender identity. Although the analyses within this paper focus primarily and almost exclusively on gender, with no real exploration into how other identity facets might impact gendered language performance overall, this paper does consider contextually relevant facets in each analysis. What this paper does not seek to do is argue for the absence of such contemporary considerations as outlined in this chapter, but rather, as is further explored in Chapter 3 below, apply the theoretical approaches to identity along with various theoretical approaches to interaction.

As we saw throughout Part 3 of this chapter, both men and women are equally capable of employing any conversational strategy they so choose in any discourse context, and to any audience. As Cameron (2010) puts it: “Like the idea that women talk more than men, the idea that women’s talk is typically cooperative and empathetic whereas men’s is more adversarial is a persistent folk-belief, and is often presented as a self-evident truth in both self-help writing and scientific sources (Expert as well as popular.” It is not so much what men and women can say, but how men and women are likely to be assessed based upon what they *do* say. Women are seen as more cooperative speakers, and men more competitive, and whenever either sex ventures too far into the realm of the other, however appropriate such sojourns may be to the discourse, situation, or interlocutors present, we as societally conditioned listeners of language enter an uncanny valley of sorts, left to perceive perfectly viable, rule-abiding language as somehow wrong based not upon natural language so much as chromosomal pairs.

Each new model of analyzing the gendered features of language tends to criticize its predecessors for either the too biased or not biased enough nature of their approaches, and such a pattern is likely to continue to repeat itself for as long as scholars attempt to adduce a model of gendered language that accounts for the real world. For the purposes of this study, however, such stereotypes are important not only to consider, but to understand within the binary categories within which they are often purported to exist, as language guising online is seldom conducted by learned, degreed, peer-reviewed linguists. Rather than indicative of the discourse of the genders to which they are often stereotypically argued to belong, it is hoped that any overt or excessive use of such stereotypes is indicative of attempts by laypersons to naively apply them as effective methods of gender guising their own language.

Finally, Part 4 sought to overview computationally derived features of statistical significance between gender-coded CMC, which is the focus of the data analyses throughout this paper, while situating these features, where applicable, within the historical context of their observational, often qualitative counterparts. These features are continually applied throughout this paper, which, as is more specifically outlined below in Chapter 3, seeks to situate the application and findings of such (and not necessarily explicitly these) feature sets

in the appropriate contexts and relevant theories of a given set of data.

Chapter 3 – Methodology

Part 1 – Introduction

Several approaches exist which combine quantitative and qualitative efforts in analysis, both linguistic and otherwise. Ehrlich and Meyerhoff (2014), for example, discuss, “A slightly different model for mixing methods [which] can be discerned in work that uses corpus linguistics in combination with discourse analysis,” in which the corpus analysis is largely qualitative and computational, and the discourse analysis is largely qualitative and analogue. The analyses in this paper seek to follow a similar model, combining the quantitative features outlined in the latter part of Chapter 2 with various linguistic theories as outlined in Part 3 of this chapter below.

Swofford and Champod (2021) suggest a model for such an approach, which, when it is possible to establish a baseline that would preclude various features at the outset, would fall under the category of algorithm assistance:

The human is responsible for forming an expert opinion based on subjective observations. The algorithm *may* be used *after* an initial opinion has been formed. The algorithm serves as an optional assistance tool supplemental to the expert opinion that may be used at the discretion of the examiner.

Such preclusions might include, as we will see below, the removal of any of the twenty features proposed in Chapter 2 and further tested in Part 2 here, based on the genre, context, or circumstances of a dataset.

Features such as hyperlinks are, for example, which were proposed on an analysis of blog data, are irrelevant to the genre direct messages in the catfi.sh Tinder data in Chapter 4. Features such as kinship terms did not prove useful in the context of the Hemmert case in Chapter 5 below, as the authors were an estranged husband and wife almost exclusively discussing the circumstances and fallout of their impending divorce. Features such as emojis and emoticons, while interesting in the catfi.sh Tinder data, were not consistently preserved in the Hemmert data given to the court and used for this analysis. For these and other reasons, an analyst with a given set of features may decide to exclude some at the outset as irrelevant, unhelpful, or unreliable.

It may also be the case that, when the output of included features results in unexpected outliers (either in the absence of features for comparison, or in features that appear wildly inflated or deflated on the surface) the approach may then fall into the category of what Swofford and Champod (2021) refer to as an algorithm informed evaluation:

The human is responsible for forming an expert opinion based on the output of the algorithm. The algorithm *shall* be used *before* the opinion has been formed. The algorithm serves as an integrated factor informing the opinion.

It is then up to the expert analyst to apply approaches such as a keyword analysis, or consider

exigent factors such as those discussed in Part 3 below, to determine whether these outputs are indicative of a useful pattern for determining likely common/non-common authorship or identity (in this case gender identity), or whether these outputs are indicative of some other factor that would shift their face value use. While it is not necessarily the case that the process will remain exactly the same for every analysis, both in that the algorithm (in this case the features applied) should be based on a reasonable framework of reliability, and in that the decisions made by the analyst surrounding the algorithm should be based on reasonable linguistic theories relevant to a given analysis.

This chapter seeks to outline such a methodological approach, first by providing a baseline for the algorithm—in this case the features applied—in Part 2 using an edited version of the Blog Authorship Corpus, and then by providing examples, in Part 3, of the types of linguistic theories that an analyst might consider in appropriately situating the algorithm's output in the context of the data itself. Because real-world data, forensic and otherwise, is so variable, and seldom meets the requirements for the types of statistical analyses outlined above in Chapter 2, the analyses in these paper instead draw on features already suggested by prior, big data, statistical research to demonstrate useful distinctions in gendered language, and does not further seek to apply them statistically (but still does so quantitatively).

Finally, although this paper considers around twenty features of gendered language (with various permutations within), and suggests only a handful of linguistic theories, this paper seeks neither to argue that these features and these theories are the only or best features to use in a given analysis, nor to claim to be exhaustive of what features and theories might be relevant. Rather, this paper seeks to test the twenty features proposed in Part 2 below in order to determine whether any reliably indicate gender (or gender-guised) performance and argue for the responsibility of an analyst to situate any such findings in their appropriate contexts, which may include but are not limited to those linguistic theories discussed in Chapter 2 and Part 3 below.

Part 2 – Variables

As discussed in Part 1 above, this section aims to apply the 20 features discussed in the latter part of Chapter 2 to an edited version of the Blog Authorship Corpus (eBAC) in order to reestablish the features suggested by Schler et al. (2006) and others as an appropriate baseline for a non-statistical and more qualitatively approachable analysis. The baselines established in this chapter are used in one of two ways throughout the analysis contained in this paper. First, the size of the eBAC provides more feature saturation than much of the remaining real-world data upon which to compare normed baseline frequencies of the male and female subcorpora. Second, the trends in the eBAC suggest the more male- or female-coding of a given feature expression, which can be applied when the normed use of a single author's total performance does not answer the research question of the data (as we will see in the catfi.sh Tinder data in Chapter 4, which splits a single author's relative uses along three different phases of performance).

A. The edited Blog Authorship Corpus (eBAC)

The original Blog Authorship Corpus (BAC) is described by Schler, Koppel, Argamon, and Pennebaker (2006) as follows:

The Blog Authorship Corpus consists of the collected posts of 19,320 bloggers gathered from blogger.com in August 2004. The corpus incorporates a total of 681,288 posts and over 140 million words - or approximately 35 posts and 7250 words per person.

Each blog is presented as a separate file, the name of which indicates a blogger id# and the blogger's self-provided gender, age, industry and astrological sign. (All are labeled for gender and age but for many, industry and/or sign is marked as unknown.)

All bloggers included in the corpus fall into one of three age groups:

- 8240 "10s" blogs (ages 13-17),
- 8086 "20s" blogs (ages 23-27)
- 2994 "30s" blogs (ages 33-47).

For each age group there are an equal number of male and female bloggers.

Each blog in the corpus includes at least 200 occurrences of common English words. All formatting has been stripped with two exceptions. Individual posts within a single blogger are separated by the date of the following post and links within a post are denoted by the label `urlLink`.

The below analysis considers an edited version of the BAC, with the removal of 1,855 total blogs and 106,572 total posts, totaling around 22 million removed words, with the differences shown in Table 3.1 below.

	Original BAC	eBAC	Difference
Total Blogs	19,320	17,465	1,855
Total Posts	681,288	574,716	106,572
Total Words	~140 million	~118 million ⁵	~22 million
Posts per Blog	~35	~33	~2
Words per Blog	~7,250	~6,787	~463

Table 3.1 Overall distribution of the Original and Edited BAC

First and foremost, blogger.com offers and has always offered bloggers the option to share up to a total of 100 specific authors and administrators to an individual blog—meaning, essentially, that blogs can be co-authored. This is evidenced in the original BAC by the repetition of whole blogs two or more times, with some duplications occurring in both the male and female dataset, indicative that the original blog was shared by multiple authors of both genders. Where duplicate blogs occurred within a single gender, only one version was left in the eBAC; where duplicate blogs occurred in both genders, all iterations of the blog were removed from the eBAC, as it is impossible to differentiate the gender of each individual post within the data itself.

In addition, any blog entries that did not contain English, or otherwise contained enough non-English as to be largely unintelligible to English-only readers were removed. Where code-switching occurred, language was retained if the primary language was English, as opposed

⁵ The total here includes counts of the link denotation `urlLink` for direct comparison to the original BAC counts, which total 1,184,362 hits not otherwise considered in the total word counts throughout this analysis as they are not original to the blogs themselves.

to the code-switched variety.

Otherwise, any common repetition of entire posts was removed as they were often found to be copied and pasted from other sources, including such originals as: news articles, book entries, song lyrics, horoscopes, Bible verses, quiz results, and other such commonly repeated language either repeated multiple times within the original BAC (as by individual posts), or found to have some other commonly used source. Also included in this removal were surveys, as the language of the questions did not originate with the blog authors, and the language of the responses was most often much shorter, frequently single-word answers (e.g., “How old are you?” versus “20”).

The removal of redundant, unreliable, and copied language resulted in the following distribution between the male and female halves of the eBAC, notably with around 1 million fewer words and 1,000 fewer blogs in the female half (though the words per post and posts per blog remain relatively close to one another).

	Male	Female
Total Words	59,233,203	58,108,641
Total Posts	287,381	287,335
Total Blogs	9,318	8,147
Words Per Post	206	202
Words Per Blog	6,356	7,132
Posts Per Blog	31	35

Table 3.2 Overall distribution of the eBAC by male and female subcorpora

The data for the eBAC is provided in Appendix A to this thesis.

B. The Features

The counts below consider the eBAC and are normed to 1,000,000 words. Elsewhere in this thesis, these counts are converted, and normed to 1,000 words. [Where regular expressions or lemma lists were used to obtain word counts in this and any other chapter analyzing these features, see Appendix B for the specific searches.]

1. Pronouns

As demonstrated by Table 3.3 below, total pronoun use does indeed occur more frequently in the female data than in the male. Indeed, these numbers are in line with those found by Schler et al. (2006), who found pronoun uses of around 111,380 for males and 133,410 for females per million words.

	Male			Female		
	#	norm	%	#	norm	%
first person	3,570,203	60,274.6	54.9%	4,583,007	78,869.6	59.1%
second person	715,838	12,085.1	11.0%	762,181	13,116.5	9.8%
third personal	2,214,167	37,380.5	34.1%	2,415,726	41,572.6	31.1%
TOTAL pronoun use	6,500,262	109,740.2	---	7,760,914	133,558.7	---
Schler pronoun use	--	111,380	--	--	133,410	--

Table 3.3 Overall pronoun distribution by person

In addition to pronouns overall, Bamman et al. (2014) found alternative pronoun spellings (u, ur, yr) to be female coded. These overall counts of alternative pronouns are demonstrated in Table 3.4 below, and indeed the female data demonstrates the alternative spellings at a higher rate than the male.

	Male		Female	
	#	%	#	%
you, your, youre, you're	660,767	92.3%	696,671	91.4%
u, ur, yr, yur, yer, ure	55,071	7.7%	65,510	8.6%
I	1,744,247	82.2%	2,123,459	76.7%
i	378,522	17.8%	644,555	23.3%

Table 3.4 Overall distribution of pronouns and their alternative spellings

However, when the individual pronoun counts are broken down, as in Table 3.5 below, we do find that some alternative spellings not considered by Bamman et al. (2014) occur more frequently in the male data than they do in the female data. These include the second person pronoun alternative spellings *ye*, *yer*, *yaself*, and *yerselves*, and much less difference between the male and the female data in such alternative spellings as *yur*, *ureselves*, *yurself*, *yerself*, *yoself*, and *youer*. I note, however, that even of those alternative spellings that appear to be more male coded, the male data makes significantly more use of the female-coded alternative spellings, as well. (This would appear to buck the trend suggested by Lakoff (1975) of women having more access to both male- and neutral-coded features than men do to female-coded features.)

	Male		Female	
	#	norm	#	norm
Female prevalent alternative spellings				
u	37,954	640.8	41,685	717.4
yr	1,138	19.2	1,590	27.4
ure	65	1.1	265	4.6
ya	11,024	186.1	14,773	254.2
yall	678	11.4	1,667	28.7
urself	259	4.4	339	5.8
thee	663	11.2	1,101	18.9
thy	637	10.8	1,097	18.9
thine	84	1.4	115	2.0
thou	939	15.9	1,564	26.9
Male prevalent alternative spellings				
ye	1,022	17.3	820	14.1
yer	485	8.2	375	6.5
yoselves	1	0.0	0	0.0
yaself	9	0.2	5	0.1
yerselves	3	0.1	1	0.0
Non-prevalent alternative spellings				
urselves	12	0.2	9	0.2
yurself	7	0.1	7	0.1
yerself	15	0.3	16	0.3
yoself	4	0.1	4	0.1
youer	1	0.0	1	0.0
yur	71	1.2	76	1.3
2nd Person Informal	55,071	930.0	65,510	1,127.4

Table 3.5 Distribution of pronoun alternative spellings

Overall, the eBAC contains relatively the same distribution found in both Bamman et al. (2014) and Schler et al. (2006), with more prevalent pronoun use, including that of alternatively spelled pronouns, by female authors. This data suggests that only a small number of alternative spellings (notably *ye* and *yer*) may be more male than female coded.

2. Emotion terms

Table 3.6 below demonstrates the overall distribution of the use of emotion terms,

which both Bamman et al. (2014) and Schler et al. (2006) found to be more female-coded, as defined by Clore and Ortony (1988), and considering those lexemes provided in the British National Corpus's lemma list.

	Male		Female	
	#	normed	#	normed
Total Instances	1,713,565	28,929.1	1,972,485	33,944.8
	Male		Female	
	All	Unique	All	Unique
Total Variety	451	0	453	2

Table 3.6 Overall distribution of emotion terms

Indeed, these findings demonstrate that emotion terms in general are more commonly used in the female dataset, with all but two emotion terms found in both the male and female datasets. Only *heartsore* and *warmhearted* were unique to the female dataset, though neither word was used prevalently in either, with only 2 and 4 hits, respectively, and notably the pair of unique words is both a negative and positive emotion word.

	Male	Female
unique		heartsore, warmhearted

Figure 3.1 Unique emotion terms

3. Emoticons

The eBAC itself exists in a time before the emojis (icon-based) that Schler et al. (2006) considered (as largely female-coded) were more widely used or indeed even available. As such, the instances found in the eBAC are largely the emoticons (text-based) Bamman et al. (2014) considered (as largely female-coded). However, Bamman et al. (2014) seemed to only consider standardized emoticons, such as those containing colons or semi-colons for eyes, and not any of the other variety of text-based emoticons found in the eBAC. Just as the eBAC was somewhat too early in the history of CMC for emojis to be widely used, the eBAC was also somewhat earlier than the later formalization of emoticons, and as such contains several of what can now be considered “non-standard”—or at the very least, less common—innovations in emoticon use.

T_T	^_^	@~@	0AO	-_-	^0~	>_>	o3o	0^0	X_x
-----	-----	-----	-----	-----	-----	-----	-----	-----	-----

Table 3.7 Examples of non-standard, innovative emoticon use as found in the eBAC

Many of these instances, as exemplified in Table 3.7 above, are derived from Japanese conventions of emoticon construction, and as such might involve things like an asterisk (*), apostrophe ('), or semi-colon (;) at the end of a face for a “sweat drop” (“^_^*”, “-_-’”, or “u_u;”), slashes (/) for blushing (“^///^”), parentheses (()) for face boundaries (“(^_^)”), or other additions (“^_^d” or “<(o_o <”). Within non-borrowed conventions, some of the emoticons Bamman et al. (2014) suggested may have the addition of a nose (“:-)” and not “:.)”), may have different eyes (“=”) different or duplicated mouths (“:]”, or even “^.^” or “^____^”), or other additions (“:’)”, “0:)”, or “>:)”). That these varieties are “non-standard” is evident in the fact that they do not translate into specific emojis in the way that those suggested by Bamman et al. (2014) now do (“:”) for 😊, “;”) 😏, “:(” for ☹️, and so on).

Because of this potentially unending variety in a dataset like the eBAC, which brings with it the possibility of missing a significant portion of such non-standard emoticons, the likelihood that these uses are not necessarily equivalent to contemporary uses of either emoticons or emojis, and due to the fact that none of the latter datasets (save the Tinder dataset, which does consider emojis specifically, in Chapter 4 below) analyzed in this thesis either have or preserved emoticons or emojis, this feature is not further analyzed here.

4. Kinship terms

		Male			Female		
		#	norm	%	#	norm	%
Female Coded	mom	16,437	277.5	10.6%	34,317	590.6	15.0%
	mommy	621	10.5	0.4%	2,144	36.9	0.9%
	mother	10,797	182.3	7.0%	16,178	278.4	7.1%
	sister	9,395	158.6	6.1%	15,680	269.8	6.9%
	daughter	3,126	52.8	2.0%	5,526	95.1	2.4%
	aunt	2,359	39.8	1.5%	4,492	77.3	2.0%
	auntie	558	9.4	0.4%	955	16.4	0.4%
	grandma	1,734	29.3	1.1%	3,664	63.1	1.6%
	grandmother	1,572	26.5	1.0%	2,280	39.2	1.0%
	kids	15,436	260.6	10.0%	19,814	341.0	8.7%
	child	17,342	292.8	11.2%	20,635	355.1	9.0%
	dad	15,262	257.7	9.9%	22,249	382.9	9.7%
	daddy	1,564	26.4	1.0%	2,972	51.1	1.3%
	father	10,478	176.9	6.8%	9,670	166.4	4.2%
	husband	2,904	49.0	1.9%	8,506	146.4	3.7%
	hubs	29	0.5	0.0%	28	0.5	0.0%
	hubby	112	1.9	0.1%	2,174	37.4	1.0%
	brother	12,804	216.2	8.3%	15,156	260.8	6.6%
	boyfriend	3,297	55.7	2.1%	8,598	148.0	3.8%
	bf	675	11.4	0.4%	1,955	33.6	0.9%
	uncle	3,577	60.4	2.3%	3,692	63.5	1.6%
	grandpa	671	11.3	0.4%	1,385	23.8	0.6%
	granddad	145	2.4	0.1%	147	2.5	0.1%
	grandfather	1,197	20.2	0.8%	1,096	18.9	0.5%
	cousin	4,409	74.4	2.8%	6,517	112.2	2.9%
Female Coded Subtotal		136,501	2,304.5	88.2%	209,830	3,611.0	91.9%
Male Coded	wife	9,027	152.4	5.8%	6,124	105.4	2.7%
	wifey	43	0.7	0.0%	47	0.8	0.0%
	girlfriend	5,579	94.2	3.6%	5,586	96.1	2.4%
	gf	819	13.8	0.5%	857	14.7	0.4%
Male Coded Subtotal		15,468	261.1	10.0%	12,614	217.1	5.5%
Other	sis	1,286	21.7	0.8%	2,346	40.4	1.0%
	sista	65	1.1	0.0%	340	5.9	0.1%
	mama	980	16.5	0.6%	2,290	39.4	1.0%
	papa	480	8.1	0.3%	720	12.4	0.3%
	gramps	26	0.4	0.0%	63	1.1	0.0%
	daddio	5	0.1	0.0%	12	0.2	0.0%
Other coded Subtotal		2,842	48.0	1.8%	5,771	99.3	2.5%
TOTAL		154,811	2,613.6	---	228,215	3,927.4	---

Table 3.8 Overall distribution of female-, male-, and other-coded kinship terms

The categories of kinship terms as suggested by Bamman et al. (2014) are indicated in the tables below in their distribution between female and male coded. As demonstrated in Table 3.8 above, although Schler et al. (2006) found similar LIWC categories of “family” words to have the “greatest information gain for gender”, the only male-coded kinship term used more

in the male half of the dataset is *wife*, while the two formal, female-coded kinship terms *father* and *grandfather* were also more common in the male subcorpus.

As shown in Table 3.9 below, while kinship terms are more prevalent in the female data overall, and both male and female terms are used in both datasets, the male subcorpus does indeed use Bamman et al.'s (2014) male-coded kinship terms with a higher distribution than does the female subcorpus.

	Male		Female	
	#	%	#	%
Female	136,501	89.8%	209,830	94.3%
Male	15,468	10.2%	12,614	5.7%

Table 3.9 Overall distribution of female- and male-coded kinship terms

5. Friendship terms

	Male		Female	
	normed	%	normed	%
Female Coded	54.7	22.1%	120.7	39.9%
Male Coded	192.8	77.9%	182.8	60.1%

Table 3.10 Distribution of female- and male-coded friendship terms

The distribution of friendship terms noted by Bamman et al. (2014) is shown for the eBAC in Table 3.10 above. While both the male and female subcorpora use male-coded terms more so than they do female-coded friendship terms, their distribution of use is significantly higher in the male subcorpus. Rather, the rate of use of female-coded friendship terms is significantly lower in the male data, skewing the use of only 10 more per million words much higher in the overall comparison. (This feature, then, does comport with Lakoff's (1995) suggestion that women have more access to male-coded features than males do female-coded features.)

		Male			Female		
		#	norm	%	#	norm	%
Female Coded	bestie	6	0.1	0.0%	34	0.6	0.0%
	bff	11	0.2	0.0%	141	2.4	0.1%
	best friend	3,221	54.4	4.2%	6,837	117.7	7.2%
Female Coded Subtotal		3,238	54.7	4.2%	7,012	120.7	7.4%
Male Coded	bro	2,121	35.8	2.7%	2,237	38.5	2.4%
	bruh	5	0.1	0.0%	4	0.1	0.0%
	brah	11	0.2	0.0%	3	0.1	0.0%
	brotha	82	1.4	0.1%	88	1.5	0.1%
	pal	942	15.9	1.2%	1,015	17.5	1.1%
	dude	5,407	91.3	7.0%	5,084	87.5	5.4%
	mate	2,851	48.1	3.7%	2,151	37.0	2.3%
Male Coded Subtotal		11,419	192.8	14.8%	10,582	182.1	11.2%
Other	friend	62,510	1,055.3	81.0%	76,992	1,325.0	81.4%
Total		77,167	1,302.8	---	94,586	1,627.7	---

Table 3.11 Overall distribution of female-, male-, and other-coded friendship terms

Schler et al.'s (2006) finding that the LIWC category of "friends" was female-coded holds true in the eBAC, where friendship terms in general are used at a higher rate in the female data.

Of the 11 types of kinship words considered in Table 3.11 above, only the male-coded terms *dude* and *mate* occur more frequently in the male data, while the various male-coded

bro terms proposed by Bamman et al. (2014) have a relatively similar distribution between the subcorpora. Notably, the neutral-coded word *friend* accounts for a similar percentage of overall friendship term distributions in both the male and female subcorpora, with Bamman et al.'s (2014) female- and male-coded terms accounting for the largest difference.

6. Abbreviations

Abbreviations occur in the data in varied ways, including interjections (*wtf*, *omg*, *lol*, *lmao*, *lmfao*), questions (*hbu*, *wbu*, *wby*), assent markers (*np*, *ofc*), and abbreviated phrases (*tbh*, *atm*, *btw*). Although a number of individual words are shortened (*yh* for *yeah*, *pls* for *please*, *soz* for *sorry*, etc.), along with compound words (*init* for *isn't it*, *dya* for *do you*, etc.), the abbreviations considered in the counts in the tables below largely include initialisms—where the first letter of every word, or compound component, are used for the abbreviation (the latter of which would include *gf* for *girlfriend*, *fb* for *Facebook*, etc.), with the infrequent exception of those abbreviations that include extra letters beyond initialisms for clarity (as in the case of *ofc* for *of course*, where the actual initialism *oc* has other meanings as in *original comment*).

As *lol* and other laughing-related abbreviations account for a significant portion of both datasets' total abbreviation counts, they are considered separately in Table 3.12 below from all other abbreviations.

	Male			Female		
	#	norm	%	#	norm	%
Total (lol/lmao)	15,239	252.4	61.3%	32,008	550.9	69.8%
Total (other)	9,602	162.1	38.7%	13,840	238.1	30.2%
Total (all)	24,841	414.5	---	45,848	789.0	---
Unlengthened	19,786	334.0	98.4%	38,109	655.8	99.0%
Lengthened	314	5.3	1.6%	382	6.6	1.0%
Total	20,100	339.3	---	38,491	662.4	---

Table 3.12 Overall distribution of abbreviations

As shown in Table 3.12 above, abbreviations in the eBAC are indeed more common in the female than the male subcorpus.

Overall, however, the male subcorpus has a higher distribution of non-*lol* abbreviations, and a somewhat higher distribution of lengthened abbreviations (lengthening as a feature is discussed further below). Lengthened abbreviations include two types--both those in which a character or characters is lengthened (such as *lol* versus *loool*, *lolll*, *looolll*), and those in which an abbreviation is fully repeated (such as *lol* versus *lolol*, *lololol*).

	Male	Female
unique	oomfg, wby, ihu	wtfwtf, withh, dtf, idfk, hbu

Figure 3.2 Non-lol unique abbreviation variations

Unique abbreviation variations, shown in Figure 3.2 above, were found between the male and female datasets. Though the difference between *wby* (what about you) and *hbu* (how

about you) is likely incidental as they occur 1 and 2 times, respectively, it is interesting to note that of the remaining unique varieties, only one in each the male and female datasets does not contain *f* for *fuck* (swearing as a feature is discussed further below).

7. Punctuation

As shown in Table 3.13 below, ellipses account for a higher percentage of female punctuation use, as do punctuation strings classified as questions (any combined string of . ! and ? containing ? defaults to a question in the count below) and exclamations (any remaining combined string of . and ! defaults to an exclamation in the count below), as predicted by Bamman et al. (2014).

	Male		Female	
	#	%	#	%
Total Periods	3,336,049	73.1%	3,286,722	67.4%
Total Ellipses	664,819	14.6%	823,896	16.9%
Total Exclamations	291,992	6.4%	468,700	9.6%
Total Questions	272,337	6.0%	298,791	6.1%
Total All	4,565,197	---	4,878,109	---

Table 3.13 Overall distribution of punctuation

Though this does not account for null-punctuation sentences, that the female subcorpus has a smaller word count but a larger punctuation count would appear to indicate that females use more punctuation (or have shorter sentences) than males, though the increased use of ellipses, which may break up rather than end sentences, may account for this to some extent. This can be seen as reflected in the distribution in Table 3.14 below.

	Male	Female
Words per Punctuation	13.0	11.9
Punctuation Per Post	15.9	17.0
Punctuation Per Blogs	489.9	684.0

Table 3.14 Overall distribution of punctuation per unit

Finally, Table 3.15 below demonstrates the distribution of lengthened and non-lengthened punctuation, where standard includes single units of punctuation (. ... ? !) and lengthened includes all complex punctuation (both ?? and ?!, for example). Again, the female subcorpus demonstrates a higher preference for the non-standard, complex punctuation as expected by Bamman et al. (2014).

	Male		Female	
	#	%	#	%
standard punctuation	4,228,952	92.6%	4,399,590	90.2%
lengthened punctuation	336,245	7.4%	478,519	9.8%

Table 3.15 Distribution of standard/unlengthened and lengthened/complex punctuation

8. Expressive lengthening

Table 3.16 below considers expressive lengthening across multiple of Bamman et al.'s (2014) categories, including backchannel sounds, abbreviations, assent terms, negation terms, hesitation words, and even punctuation. (Percentages provided in Table 3.16 below are as compared to the totality of that individual feature—as in lengthened backchannel sounds account for 27.5% of total male and 16.0% of total female backchannel sounds.)

	Male			Female		
	#	norm	%	#	norm	%
Lengthened Backchannel	83,411	1,408.2	27.5%	24,318	418.5	16.0%
Lengthened Abbreviations	314	5.3	1.6%	382	6.6	1.0%
Lengthened Assent	780	13.2	0.8%	1,374	23.6	1.0%
Lengthened Negation	732	12.4	0.5%	1,363	23.5	1.7%
Lengthened Words TOTAL	84,527	1,427.0	---	27,437	472.2	---
Lengthened Hesitation	6,229	105.2	24.7%	9,232	158.9	27.4%
Lengthened Punctuation	336,245	---	7.4%	478,519	---	9.8%

Table 3.16 Expressive lengthening overall

Other than backchannel sounds, the remaining categories are more prevalent in the female subcorpus than the male. When the outliers of *o*, *oh*, and *ooh* are removed from backchannel sounds, however, we find that, as with the rest of the categories, backchannel sounds are overall more frequently lengthened in the female subcorpus, as shown in Table 3.17 below.

	Male		Female	
	#	%	#	%
unlengthened	220,192	72.5%	127,387	84.0%
lengthened	83,411	27.5%	24,318	16.0%
TOTAL	303,603	---	151,705	---
unlengthened <i>o</i>, <i>oh</i>, <i>ooh</i>	183,963	71.9%	77,278	94.2%
lengthened <i>o</i>, <i>oh</i>, <i>ooh</i>	71,911	28.1%	4,790	5.8%
TOTAL <i>o</i>, <i>oh</i>, <i>ooh</i>	255,874	---	82,068	---
unlengthened other	36,229	75.9%	50,109	72.0%
lengthened other	11,500	24.1%	19,528	28.0%
TOTAL other	47,729	---	69,637	---

Table 3.17 Lengthened and unlengthened backchannel sounds

9. Backchannel sounds

Of the 33 backchannel variations demonstrated in Table 3.18 below, only five occur more frequently in the male dataset (*o*, *oi*, *ai*, *unf*, and *er*), while the remainder appear more frequently in the female datasets.

Although the backchannel sound *o* in the male subcorpus would appear at first glance to be strongly indicative of a potentially male-coded feature (used over 200,000 times when the next highest individual use, *oh* in the female subcorpus, is under 60,000), it is worth noting that this likely outlier may be artificially conflated by the regex variations, which include zeroes. In addition, the use of individual backchannel sounds may fluctuate in popularity over time, especially those that would not be considered standard backchannel sounds due to their spelling (such as *o*), confusability with other non-backchannel terms (such as the vocative *O* or *0* as a number), or their influence.

	Male		Female	
	#	normed	#	normed
oh variations	41,461	700.0	58,860	1012.9
ooh variations	1,450	24.5	2,867	49.3
o variations	212,963	3,595.3	20,341	350.1
aa variations	665	11.2	1,348	23.2
ah variations	10,398	175.5	14,136	243.3
ag(h) variations	252	4.3	365	6.3
arg(h) variations	3,000	50.6	4,120	70.9
ug(h) variations	1,986	33.5	4,805	82.7
aw variations	1,158	19.5	2,715	46.7

	Male		Female	
	#	normed	#	normed
awe/h variations	757	12.8	787	13.5
pf(ft) variations	238	4.0	254	4.4
ouch variations	746	12.6	851	14.6
ow variations	383	6.5	499	8.6
ew variations	531	9.0	1,847	31.8
yo variations	2,356	39.8	2,359	40.6
oi variations	462	7.8	397	6.8
oy variations	192	3.2	493	8.5
gr variations	1,826	30.8	3,677	63.3
rawr variations	37	0.6	85	1.5
ra(g)h variations	140	2.4	240	4.1
rr variations	122	2.1	130	2.2
ee variations	639	10.8	660	11.4
squee variations	4	0.1	87	1.5
nya variations	140	2.4	362	6.2
ai variations	787	13.3	657	11.3
oof variations	49	0.8	58	1.0
unf variations	10	0.2	8	0.1
um variations	4,461	75.3	7,953	136.9
uh variations	3,380	57.1	4,554	78.4
er variations	3,528	59.6	3428.0	59.0
erm(h) variations	1,194	20.2	1,205	20.7
hm variations	11,976	202.2	15,707	270.3
hum variations	689	11.6	868	14.9

Table 3.18 Backchannel sound variations

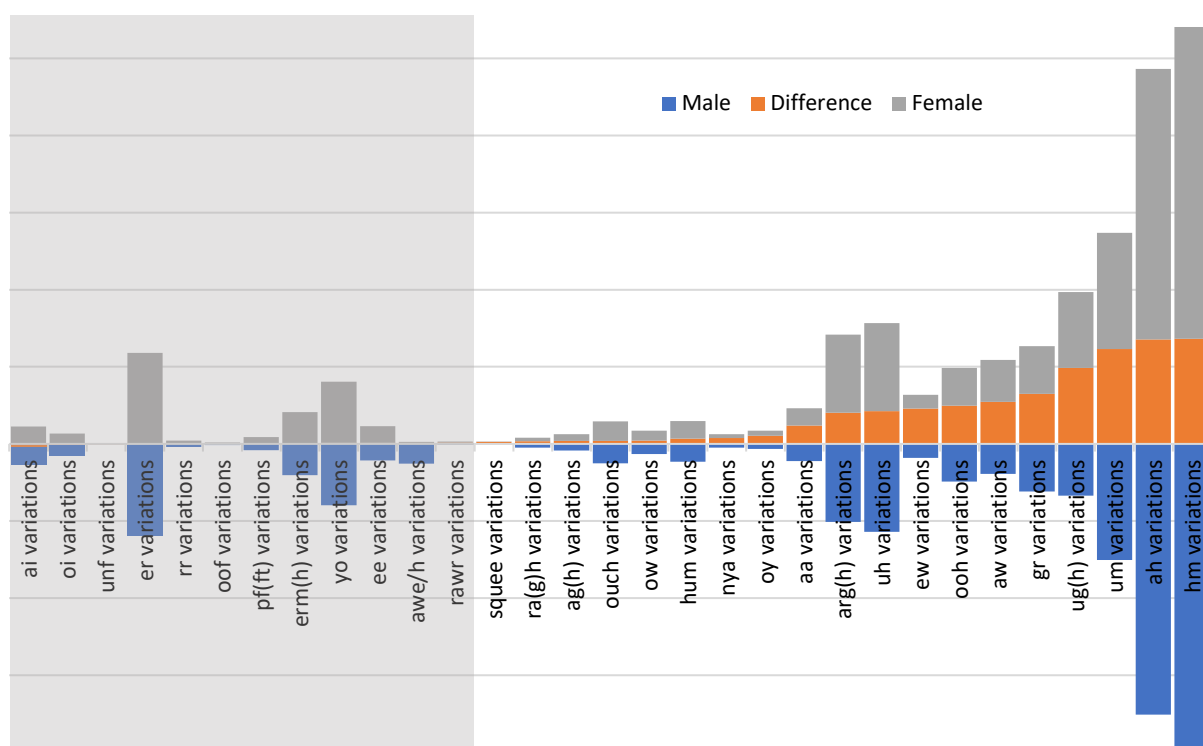


Figure 3.3 Overall distribution and differences of normed backchannel sound frequencies

Twelve of these backchannel sounds occur more frequently in one dataset by only a small margin (around one word or less per million when normed), indicated in Figure 3.3 below within the gray box. The latter can be demonstrated, for example, with the included backchannel sounds *nyaa* (blue) and *squee* (red) in Figure 3.4 from Google Trends above, both of which have peaked in their popularity of use since the 2004 conception of the BAC.

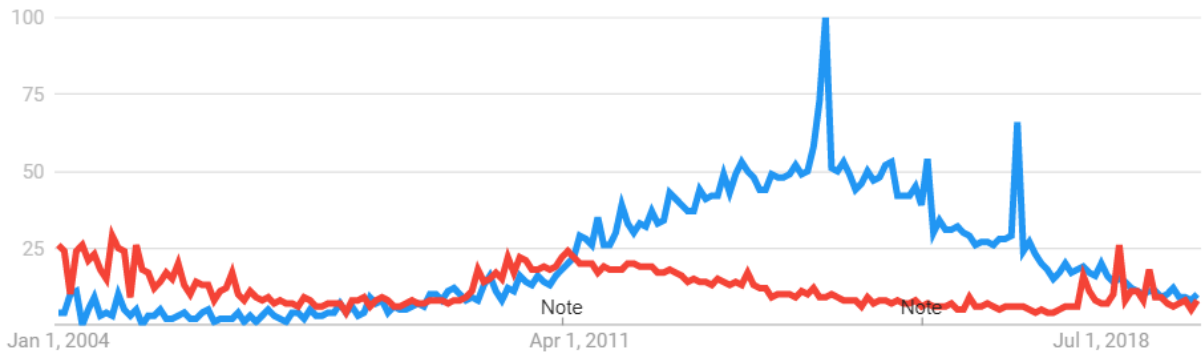


Figure 3.4 Use of nyaa (blue) and squee (red) over time via Google Trends

As demonstrated in Table 3.19 below, *o*, *oh*, and *ooh* variations make up the majority of backchannel sounds as used in both the male and female data. With these backchannel sounds removed, however, (as *o* seems to be a strong outlier for the male subcorpus) backchannel sounds go from being almost twice as frequently used in the male dataset to one and a half times more frequently used by the female dataset.

	Male			Female		
	#	normed	%	#	normed	%
<i>o</i>, <i>oh</i>, <i>ooh</i> variations	255,874	4,319.5	83.3%	82,068	1,412.3	52.8%
other variations	51,441	868.8	16.7%	73,307	1,261.6	47.2%
Total	307,315	5,188.3	---	155,375	2,673.9	---

Table 3.19 Overall distribution of backchannel sounds

The distribution of other backchannel variations is more in line with what Bamman et al. (2014) might expect. Finally, though lengthened backchannel sounds are discussed elsewhere in the section on expressive lengthening, only one lengthened variation was unique to either subcorpus, as shown in Table 3.20 below.

	Male	Female
unique		ra+wr

Table 3.20 Unique expressively lengthened backchannel sounds

10. Hesitation words

	Male			Female		
	#	normed	%	#	normed	%
um variations	4461	75.3	17.7%	7953.0	136.9	23.6%
uh variations	3380	57.1	13.4%	4554.0	78.4	13.5%
er variations	3528	59.6	14.0%	3428.0	59.0	10.2%
erm(h) variations	1194	20.2	4.7%	1205.0	20.7	3.6%
hm variations	11976	202.2	47.5%	15707.0	270.3	46.6%
hum variations	689	11.6	2.7%	868.0	14.9	2.6%
TOTAL	25228	425.9	---	33715.0	580.2	---

Table 3.21 Overall hesitation words

Hesitation words are also discussed in the expressive lengthening and backchannel sound sections above, as they are also often used as backchannel sounds and subject to expressive lengthening. As demonstrated in Table 3.21 below, hesitation words, overall, are more common in the female subcorpus, as Bamman et al. (2014) would expect, with all variations but *er* being more common in the female subcorpus individually.

As demonstrated in the section above in Figure 3.3, *er* and *erm(h)* variations appear with a distribution of less than a single word per million difference between the male and female subcorpora, while the hesitation words *um* and *hm* appear to be particularly female coded.

11. Assent terms

	Male			Female		
	#	normed	%	#	normed	%
yes variations	24,262	409.6	24.6%	30,194	519.7	22.7%
yeah variations	28,159	475.4	28.6%	34,342	591.0	25.9%
yas variations	44	0.7	0.0%	167	2.9	0.1%
yis variations	5	0.1	0.0%	0	0.0	0.0%
okay variations	12,254	206.8	12.4%	19,819	341.1	14.9%
ok variations	25,278	426.7	25.7%	35,266	606.9	26.6%
yeh variations	549	9.3	0.6%	1,420	24.6	1.1%
yep variations	1,878	31.7	1.9%	2,252	38.8	1.7%
yup variations	1,633	27.5	1.7%	2,368	40.8	1.8%
yea variations	4,478	75.6	4.5%	6,888	118.5	5.2%
TOTAL	98,540	1,663.6	---	132,726	2,284.1	---

Table 3.22 Overall distribution of assent term varieties

As mentioned above, assent terms are also subject to expressive lengthening. Table 3.22 above demonstrates the overall distribution of assent term variations. Overall, these assent terms occur more frequently in the female data (with the only exception *yis* having a difference of less than 0.1 words per million, with only 5 total uses). These results are in line with what both Bamman et al. (2014) and Schler et al. (2006) might expect.

12. Negation terms

As with assent terms, negation terms are subject to expressive lengthening, discussed elsewhere.

		Male			Female		
		#	normed	%	#	normed	%
Female Coded	noo+	623	10.5	0.4%	1,235	21.3	0.8%
	cannot	10,547	178.1	6.9%	9,416	162.0	6.0%
Female Coded Subtotal		11,170	188.6	7.3%	10,651	183.3	6.7%
Male Coded	nah	1,181	19.9	0.8%	1,315	22.6	0.8%
	n+a+h+	109	1.8	0.1%	128	2.2	0.1%
	nobody	5,163	87.2	3.4%	4,703	80.9	3.0%
	aint	4,508	76.1	2.9%	4,244	73.0	2.7%
Male Coded Subtotal		10,961	185.0	7.1%	10,390	178.8	6.6%
Other	no	131,817	2,225.4	85.6%	137,025	2,358.1	86.7%
Total		153,948	2,599.0	---	158,066	2,720.2	---

Table 3.23 Overall distribution of negation terms

As demonstrated in Table 3.23 below, although negation terms in general are more common in the female subcorpus, both female- and male-coded negation terms appear more frequently in the male subcorpus. Although the female-coded *noo* occurs more in the female data, *cannot* occurs more frequently in the male data. Similarly, although both the male-coded *nobody* and *aint* occur more frequently in the male data, *nah* occurs more frequently in the female data. As such it is largely down to the neutral-coded (by Bamman et al. (2014)) *no*, with no expressive lengthening, that the female dataset contains a higher normed frequency of negation terms.

It is, of course, worth noting that blogs, unlike the Twitter dataset considered by J. Ford, PhD, Thesis, Aston University, 2022

Bamman et al. (2014), are less likely to be conversational, and thus less likely to have assent and negation terms as a response to, for example, yes or no questions, which may account for the differing distributions of these categories between Bamman et al. (2014) and the eBAC. As shown in Table 3.24 below, both the male and female subcorpora use female- and male-coded negation terms with almost the exact same distribution (within 0.1%), and nearly 50/50 distribution each (within 0.6%). As such, negation terms in general may not be the best potential feature when it comes to indicating an author's gender.

	Male		Female	
	#	%	#	%
Female Coded	11,170	50.5%	10,651	50.6%
Male Coded	10,961	49.5%	10,390	49.4%

Table 3.24 Distribution of female- and male-coded negation terms

This may be in part due to the less conversational medium of blogs meaning that more cooperative forms of negation terms are less necessary, as an author negating something is not the same as an interlocutor providing a dispreferred response. That dispreferred responses must sometimes be given to remain felicitous may explain why research expects to find higher instances of *noo+* among female speakers. Such non-standard negation terms may fall into the category of deferential politeness, and, as with the swears below, be considered as part of a class of “weaker” negation terms (Lakoff, 1975).

13. Swears and taboo words

As demonstrated in Table 3.25 below, neither subcorpora's use of anti-swears makes up a significant portion of their total swear word use, though the use of female-coded anti-swears is higher in the female subcorpus.

	Male		Female	
	#	%	#	%
Female Coded Anti-Swears	3,182	2.6%	5,455	4.5%
Male Coded Swears	120,510	97.4%	114,527	95.5%

Table 3.25 Distribution of swears and anti-swears

Overall, as demonstrated in Table 3.26 below, *darn* and *gosh* are the anti-swears with the most prominent distributional difference between the male and female subcorpora. Although the male subcorpus does, overall, use more swear and taboo words, the difference is not large (less than 100 words per million), and the overall category for opportunities in which either a swear or anti-swear could be used is even closer (within 25 words per million).

Indeed, while *dam(n)*, *fuck*, *shit*, *bastard*, and *arse* (and to a lesser extent *fag* and *cunt*) are more frequently used in the male subcorpus, other swears such as *dam(n)it*, *fag*, *bitch*, and *ass* (and variations of *cunt*) occur more frequently in the female subcorpus. The term *bitch* (and to a lesser extent *cunt*) may not be terribly surprising when social networks are taken into consideration, and the fact that female authors tend to discuss such relationship networks more frequently, such as those of “family” and “friend” categories as included by Schler et al. (2006). But as with negation terms, these findings would seem to indicate that swears and to a lesser extent anti-swears may not be the best feature in differentiating between male and female authors.

		Male			Female		
		#	normed	%	#	normed	%
Female Coded	darn	1288	21.7	0.5%	1812	31.2	1.5%
	darnit	31	0.5	0.0%	63	1.1	0.1%
	dang	647	10.9	0.3%	878	15.1	0.7%
	dangit	61	1	0.0%	62	1.1	0.1%
	gosh	1025	17.3	0.8%	2416	41.6	2.0%
	gosh(other)	130	2.2	0.1%	224	3.9	0.2%
Female Coded Subtotal		3,182	53.6	2.6%	5,455	94.0	4.5%
Male Coded	damn	18563	313.4	15.0%	18013	310	15.0%
	damnit	721	12.2	0.6%	773	13.3	0.6%
	dam	711	12	0.6%	347	6	0.3%
	dammit	1246	21	1.0%	1733	29.8	1.4%
	dam(other)	1752	29.2	1.4%	1721	28.3	1.4%
	fuck	14075	237.7	11.4%	12018	206.8	10.0%
	fuck(other)	23214	387.1	18.8%	21358	362.9	17.8%
	shit	20481	345.8	16.6%	19237	331	16.0%
	shit(ending)	6434	106.2	5.2%	5298	88.6	4.4%
	fag	234	4	0.2%	371	6.4	0.3%
	fag(other)	517	8.1	0.4%	185	2.4	0.2%
	bastard	3116	52.6	2.5%	2033	35	1.7%
	bastard(other)	120	1.6	0.1%	63	0.7	0.1%
	bitch	5888	99.4	4.8%	7633	131.4	6.4%
	bitch(other)	1848	30	1.5%	2319	38.8	1.9%
	ass	15220	257	12.3%	15489	266.5	12.9%
	ass(other)	4729	78.2	3.8%	4850	81.5	4.0%
	arse	853	14.4	0.7%	581	10	0.5%
	arse(other)	342	5.5	0.3%	160	2.5	0.1%
	cunt	385	6.5	0.3%	311	5.3	0.3%
	cunt (other)	61	0.4	0.0%	34	0.4	0.0%
Male Coded Subtotal		120,510	2,022.3	97.4%	114,527	1,957.6	95.5%
TOTAL		123,692	2,075.9	---	119,982	2,051.6	---

Table 3.26 Overall swear and anti-swear word variations

It is also worth considering, as shown in Table 3.27 below, the distribution of abbreviations in which *f* for *fuck* can be optionally included, such as the difference between *omg* and *omfg*, and *lmao* and *lmfao*, but not including *gtfo* (get the fuck out) or *dtf* (down to fuck), for example, as *go* and *d* are not used as alternatives.

	Male		Female	
	#	%	#	%
(rotf)lmao	319	90.1%	1,381	92.4%
(rotf)lmfao	35	9.9%	113	7.6%
omg	1,696	93.0%	4,706	96.8%
omfg	128	7.0%	155	3.2%
wth	230	12.8%	168	10.4%
wtf	1,572	87.2%	1,454	89.6%

Table 3.27 *f* and non-*f* abbreviations

In comparing all three pairs above, *lmao* and *omg* are significantly more frequent for both the male and female subcorpora over their *f* alternatives, while *wtf* is significantly more frequent in both over *wth* (although it can further be noted that *h* for *hell* or *heck* may or may not be a swear or anti-swear, and is not the same *f* variation as the other two). In both of the true *f*-optional variations, at least, the *f* variation is more frequent in the male subcorpus.

14. Prepositions

		Male		Female	
		#	normed	#	normed
Prepositions	about	209,580	3538.2	216,224	3,721.0
	above	7,829	132.2	5,700	98.1
	across	10,313	174.1	8,466	145.7
	after	89,494	1,510.9	80,159	1,379.5
	against	18,214	307.5	10,208	175.7
	among	5,645	95.3	2,903	50.0
	around	60,734	1025.3	63,137	1,086.5
	at	303,794	5,128.8	306,287	5,270.9
	before	54,736	924.1	55,748	959.4
	behind	13,099	221.1	11,876	204.4
	below	4,811	81.2	3,159	54.4
	beside	1,700	28.7	2,064	35.5
	between	20,046	338.4	15,253	262.5
	beyond	4,942	83.4	3,702	63.7
	but	389,074	6,568.5	420,078	7,229.2
	by	159,805	2,697.9	120,625	2,075.9
	down	65,502	1,105.8	64,527	1,110.5
	during	20,107	339.5	15,742	270.9
	except	10,388	175.4	11,090	190.8
	for	543,749	9,179.8	517,530	8,906.2
	from	202,282	3,415.0	172,850	2,974.6
	in	811,430	13,698.9	711,862	12,250.5
	inside	11,023	186.1	12,642	217.6
	into	83,683	1,412.8	75,287	1295.6
	like	222,089	3,749.4	264,610	4,553.7
	near	8,570	144.7	6,928	119.2
	of	1,157,040	19,533.6	934,582	16,083.4
	on	437,507	7,386.2	412,637	7101.1
	onto	5,698	96.2	5,098	87.7
	out	216,899	3,661.8	231,006	3,975.4
	outside	12,258	206.9	13,189	227.0
	over	81,035	1,368.1	84,340	1,451.4
	past	20,652	348.7	19,567	336.7
	since	43,070	727.1	44,208	760.8
	through	43,511	734.6	41,078	706.9
	throughout	3,589	60.6	2,465	42.4
	to	1,714,289	28,941.4	1,728,099	29,739.1
	toward	3,217	54.3	2,289	39.4
	under	16,248	274.3	12,686	218.3
	until	29,306	494.8	33,973	584.6
	upon	9,270	156.5	7,399	127.3
	with	375,248	6,335.1	380,057	6,540.5
	within	9,007	152.1	6,294	108.3
	without	25,850	436.4	24,665	424.5
Prepositions Subtotal		7,536,333	127,231.6	7,162,289	123,256.9

Table 3.28 Overall distribution of preposition varieties and alternative spellings

As demonstrated in Table 3.28 above and Table 3.29 below, alternative prepositions are more frequent in the female subcorpus, but prepositions overall are more frequent in the male, both findings that, overall, confirm both Bamman et al. (2014) and Schler et al.'s (2006) findings.

		Male		Female	
		#	normed	#	normed
Alternative Prepositions	2	63,606	1,073.8	63,396	1,091.0
	4	33,550	566.4	31,505	542.2
	b4	1,270	21.4	1,549	26.7
	w/(o)	7,441	125.6	10,883	187.3

	Male		Female	
	#	normed	#	normed
Alternative Prepositions Subtotal	105,867	1,787.3	107,333	1,847.1
TOTAL Overall Preposition Use	7,642,200	129,018.9	7,269,622	125,104.0

Table 3.29 Overall distribution of preposition varieties and alternative spellings

15. Alternative spellings

Table 3.30 below considers alternative spellings from several categories including pronouns, prepositions, kinship terms, negation terms, and Bamman et al.'s (2014) own suggested alternative spellings. In every instance but *ain't* (which is a male-coded negation term to begin with) and *brutha*, these spellings are more frequent to the female subcorpus, which would appear to confirm Bamman et al.'s (2014) findings, though female-coded alternative spellings are more frequent in both subcorpora overall.

		Male			Female		
		#	normed	%	#	normed	%
Female Coded	vacay	12	0.2	0.0%	42	0.7	0.0%
	yaay	4,660	78.7	6.3%	11,471	197.4	11.0%
	lol	15,002	253.3	20.1%	30,689	528.1	29.5%
	u	33,884	572.0	45.5%	35,888	617.6	34.5%
	ur	3,944	66.6	5.3%	5,570	95.9	5.4%
	yr	1,138	19.2	1.5%	1,590	27.4	1.5%
	w/	7,924	133.8	10.6%	10,883	187.3	10.5%
Female Coded Subtotal		66,564	1,123.8	89.3%	96,133	1,654.4	92.4%
Male Coded	bro	2,125	35.9	2.9%	2,239	38.5	2.2%
	bruh	5	0.1	0.0%	4	0.1	0.0%
	brutha	18	0.3	0.0%	7	0.1	0.0%
	nah	1,289	21.8	1.7%	1,442	24.8	1.4%
	ain't	4,508	76.1	6.1%	4,244	73.0	4.1%
Male Coded Subtotal		7,945	134.1	10.7%	7,936	136.6	7.6%
TOTAL		74,509	1,257.9	---	104,069	1,790.9	---

Table 3.30 Alternative spellings

16. Conjunctions

As shown in Table 3.31, in contrast to what Bamman et al. (2014) would expect, all three variations of *and* considered below, including Bamman et al.'s (2014) suggested & and the informal *n*, are more common in the female subcorpus than the male.

	Male			Female		
	#	normed	%	#	normed	%
and	1,536,135	25,676.5	98.8%	1,633,275	28,107.3	96.8%
&	1,627	27.2	0.1%	25,032	441.7	1.5%
n	17,244	288.2	1.1%	29,664	510.5	1.8%
TOTAL	1,555,006	25,991.9	---	1,687,974	29,048.6	---

Table 3.31 Distribution of conjunctions and their alternative spellings

17. Articles and determiners

As shown in Table 3.32 below, determiners are indeed more frequent in the male than female subcorpus, though not categorically. The most notable exception are possessive pronouns, which account for the biggest difference in categorical distribution (along with *the*).

	Male			Female		
	#	normed	%	#	normed	%
the	2,550,149	43,052.7	30.4%	2,057,750	35,412.1	26.2%
a, an	1,478,041	24,942.9	17.6%	1,315,844	22,644.6	16.7%

this, that, these, those	1,321,175	22,304.6	15.8%	1,267,614	21,814.9	16.1%
my, your, his, her, its, our, their	1,234,507	20,841	14.7%	1,404,013	24,161.9	17.9%
one, ten, twenty, etc.	338,253	5,710.5	4.0%	314,990	4,520.7	4.0%
all, both, half, either, neither, each, every	393,919	6,550.3	4.7%	408,991	7,038.4	5.2%
other, another	133,824	2,259.3	1.6%	112,767	2,112.7	1.6%
such, rather, quite	77,549	1,309.2	0.9%	67,765	1,166.2	0.9%
wh- determiners	721,029	12,172.7	8.6%	766,893	13,197.6	9.8%
no	131,817	2,225.4	1.6%	137,025	2,358.1	1.7%
TOTAL	8,380,263	141,479.2	---	7,863,652	134,236.7	---

Table 3.32 Articles and determiners

18. Post length

The distribution of the subcorpora is recreated again below in Table 3.33, which shows that post lengths of the eBAC have the opposite distribution, with the average male post (206) larger than the average female post (202), although this difference is smaller overall than in the original BAC, at only 4 as opposed to 12 words different. Post length itself, then, is likely not a good indicator of potential gendered authorship, and the larger lengths may have come down to what was removed from the blogs—lyrics, quotes, quizzes, and so on—which were more frequently found in the female subcorpus.

	Male	Female
Total Words	59,233,203	58,108,641
Total Posts	287,381	287,335
Total Blogs	9,318	8,147
Words Per Post	206	202
Words Per Blog	6,356	7,132
Posts Per Blog	31	35

Table 3.33 Overall distribution of the eBAC by male and female subcorpora

19. Hyperlinks

Table 3.34 demonstrates the distribution of the hyperlink tag *urlLink* in the eBAC.

	Male		Female	
	#	words per	#	words per
urlLink	593,398	99.8	590,954	98.3

Table 3.34 urlLink uses

Although the use of hyperlinks is indeed still higher in the male subcorpus, the difference in words per hyperlink in the eBAC is minimal (1.5). However, the editing of the BAC was, in part, to remove language taken from other sources, and as such the remaining uses of *urlLink* can be seen to more accurately but *only* reflect those instances in which hyperlinks are used within a sentence or the body of a blog post, as opposed to separately, or as a part of a recurring tag.

20. Keywords

As compared to one another, the female subcorpus has 15,243 (or 31.2%) fewer keywords (with a keyness of 6.53 or higher) than does the male subcorpus, as shown in Table 3.35 below. This is somewhat reflected in the type/token ratio of each subcorpus; though both have relatively low type/token ratios (likely due to the non-standard nature of blogs as compared to other genres of formal writing), the male subcorpus is slightly higher.

	Male		Female		Difference	
					#	%
Keywords	48,840		33,597		15,243	31.2%
Male			Female			
Type	Token	Ratio	Type	Token	Ratio	
373.805	59,233,203	0.0063	330.024	58,108,641	0.0057	

Table 3.35 Total keywords and type/token ratio

Table 3.36 below demonstrates the top 10 overall keywords for both the male and the female subcorpora, with only *i* and *urlLink* shared between the two, though they occupy the same position within each. Neither is surprising given the genre of blogs. Of the remaining 8 for each, the male subcorpus has two function words—an article (*the*) and a preposition (*of*)—and six proper noun content words, which fall into either the category of religion (*god*) or politics (*bush*, *iraq*, *american*, *kerry*, *john* (Kerry)); the female subcorpus has only function words, with six pronouns (*my*, *she*, *it*, *me*, *he*, *we*), one backchannel sound (*oh*), and the multifunctional word *so* (which may function as an adverb, conjunction, pronoun, adjective, or interjection). These instances of *so* may include intensive uses, expressing for example “very, really, utterly”, or uses of *so* as a hedge, which Lakoff (1975) expects to be more female coded.

		Male			Female		
		hit count	keyness	keyword	hit count	keyness	keyword
Overall Keywords	1	2,087,155	858,061.7	i	2,712,757	2100438.6	i
	2	171,612	235,826.0	urllink	115,905	162091.0	urllink
	3	2,550,159	89,971.5	the	445,771	54507.1	so
	4	1,157,046	24,153.3	of	729,170	45090.4	my
	5	21,412	23,502.7	bush	221,375	31847.0	she
	6	40,251	21,185.3	god	903,555	30659.9	it
	7	12,597	15,663.2	iraq	483,565	30481.8	me
	8	15,762	14,877.1	american	328,867	29640.4	he
	9	11,638	14,346.1	kerry	57,205	27892.5	oh
	10	14,095	12,263.1	john	372,290	24831.3	we

Table 3.36 Top 10 overall keywords

Many of these categories and their distributions are in line with both the findings of Bamman et al. (2014) and Schler et al. (2006) as discussed in the sections above, with Schler et al. (2006) finding that males use more articles and prepositions, and discuss politics more as a LIWC category, and Bamman et al. (2014) finding that females tend to use more pronouns and backchannel sounds.

Table 3.37 below demonstrates the top 10 overall function keywords for both the male and the female subcorpora, with only *i* and *it* shared between the two. Notably, again, the male subcorpus has articles (*the*, *a*, *this*) and prepositions (*of*, *as*, *in*, *by*), while the female subcorpus has more pronouns (*my*, *she*, *me*, *he*, *we*) and a backchannel sound (*oh*), with none of these function keywords overlapping. Of the remaining keywords for each, the male subcorpus has the copula *is*, and the female corpus has the conjunction *but*.

Function	Male				Female			
	#	hit count	keyness	keyword	#	hit count	keyness	keyword
Function	1	2,087,155	858,061.7	i	1	2,712,757	2,100,438.6	i
	3	2,550,159	89,971.5	the	3	445,771	54,507.1	so
	4	1,157,046	24,153.3	of	4	729,170	45,090.4	my
	12	284,092	10,987.9	as	5	221,375	31,847.0	she

13	1,328,006	10,835.5	a	6	903,555	30,659.9	it
14	811,456	10,375.7	in	7	483,565	30,481.8	me
16	398,467	9,883.9	this	8	328,867	29,640.4	he
21	159,806	7,201.7	by	9	57,205	27,892.5	oh
23	651,567	6,780.1	is	10	372,290	24,831.3	we
24	856,170	6,662.7	it	12	420,078	19,484.8	but

Table 3.37 Top 10 function keywords

Although *is* is discussed elsewhere below, the keyness of the conjunction *but* is in line with what Bamman et al. (2014) might expect. In addition, although *but* does occur within the top 100 keywords for the male subcorpus as well, as shown in Table 3.38 below, not only is this with a much lower keyness than in the female subcorpus, but in addition the conjunction *and* (as well as the likely, non-standard conjunction *n*) does not occur within the top 100 keywords for the male subcorpus.

	Male			Female		
	rank	count	keyness	rank	count	keyness
but	72	389,083	3,353.7	12	420,078	19,484.8
and	--	--	--	15	1,633,275	14,188.3
n	--	--	--	45	29,664	6,292.5

Table 3.38 Conjunction keyness in the top 100 keywords

Table 3.39 below shows the top 10 content keywords for both subcorpora, excluding proper nouns (specifically names and places) and time-related words (such as times, days, months), as these are considered as specific categories further below. Of the top 10 content keywords, only *mr* and *tv* are shared, though it is worth noting *mr* has a much higher keyness in the male than the female subcorpus.

	Male				Female			
	#	hit count	keyness	keyword	#	hit count	keyness	keyword
Content Keywords	15	13,249	10,025.9	mr	14	100,841	15,243.4	love
	17	57,939	8,807.0	blog	42	9,347	6,936.3	mr
	22	19,005	7,140.4	war	48	264,610	6,208.7	like
	34	4,280	5,645.4	tyke	50	151,111	5,852.2	go
	37	12,315	5,481.9	government	54	172,658	5,708.4	know
	39	29,550	5,152.8	game	55	73,012	5,600.1	feel
	43	13,033	4,862.7	tv	57	38,392	5,402.3	happy
	44	8,204	4,798.2	national	69	11,019	4,902.0	tv
	52	81,916	4,414.2	new	76	16,898	4,442.7	baby
	64	8,968	3,791.4	cd	77	12,970	4,412.5	cute

Table 3.39 Top 10 content keywords

Notably, the female subcorpus contains 5 content keywords that function as verbs (*love*, *like*, *go*, *know*, *feel*), while the male subcorpus does not. Considering the categories analyzed by Schler et al. (2006), these findings remain in line with the original BAC. The male subcorpus contains the categories of politics (*war*, *government*, *national*), technology (*blog*, *tv*, *cd*), and sports (*game*), while the female subcorpus contains the categories of emotions (*love*, *like*, *feel*, *happy*) and family (*baby*), as well as the individual term *cute*.

These categorical distributions are further demonstrated in Table 3.40 below with the top 10 proper noun content words. While both the male (*bush*, *iraq*, *american*, *kerry*, *john*,

america, president, united) and female (*john*⁶, *american, bush*) subcorpora contain political words, almost all of the top 10 for the male fall under the political category, with the remaining 2 (*god, jesus*) falling under religion.

Content Keywords	Male				Female			
	#	hit count	keyness	keyword	#	hit count	keyness	keyword
	5	21412	23502.659	bush	11	38487	23264.138	god
	6	40251	21185.348	god	22	33128	10595.965	mom
	7	12597	15663.237	iraq	39	8636	7210.932	john
	8	15762	14877.075	american	40	155471	7055.286	don
	9	11638	14346.097	kerry	49	7372	5926.180	christmas
	10	14095	12263.123	john	51	8168	5808.413	american
	11	11139	11846.350	america	67	6645	5166.704	lord
	18	13210	8429.519	president	68	7129	4976.568	bush
	19	7376	8093.982	united	71	5542	4654.187	jesus
	20	8102	7431.106	jesus	75	6048	4498.594	chris

Table 3.40 Top 10 proper noun content words

The female subcorpus has more religion words (*god, christmas, lord, jesus*), another family word (*mom*), and proper names (*john, don, chris*). This confirms Schler et al.'s categorical findings, and their suggestions that "female writing tends to emphasize what Biber [3] calls 'involvedness', while male writing tends to emphasize 'information'" (2006), and that the "differences further suggest a pattern of more 'personal' writing by female bloggers than male bloggers," (2006).

		Male		Female	
Function Words	pronouns	i it they their which also	6	i my she it me he we him you im her u	12
	articles	the a this an some these	6	my her what all how	5
	prepositions	of as in by after from	6	then when	2
	verbs	is has	2	m am maybe go know do	6
	conjunction	but if	2	but and n also if	5
	backchannel	oh	1	oh yay	2
	abbreviations	--	0	lol omg	2
	negation	--	0	t	1
	assent	--	0	ok yes okay yeah	4
	other	however well	--	so well just anyway really like hey there cuz because anyways	--

Table 3.41 Top 100 keywords that are function words

Table 3.41 above shows the distribution of the top 100 keywords that fall into the function words type, by their general part of speech (as some may frequently function as more than one part of speech). The female corpus contains more function words within the top 100 keywords, and the distribution of these function words again follows what would be expected by Bamman et al. (2014) and Schler et al. (2006), with more frequent articles and prepositions in the male subcorpus, and more frequent pronouns, conjunctions, backchannel sounds, abbreviations, negation, and assent terms in the female corpus.

Table 3.42 below shows the content words within the top 100 keywords, separated by thematic category. Again, the male subcorpus has more top keywords in the categories of

⁶ That "john" occurs in in both subcorpora, but "kerry" only appears in the top 100 for the male subcorpus likely indicates that many more instances of "john" in the female dataset do not refer to politician John Kerry.

politics, technology, and sports, while the female subcorpus has more in the categories of emotions and family, as found by Schler et al. (2006). Shared (and thus not particularly distinguishing) categories include proper and place names (though many of them may come from the political category, especially in the male subcorpus), religion, time, and characters (many of which are likely from emoticons, such as :P :C :D :J). Overall, the male subcorpus has more content words than does the female subcorpus, which skews more toward function words.

	category	Male	Female
Content Words	politics	bush iraq american kerry john america president united war americans george states government national michael moore saddam washington clinton iraqi reagan democratic king democratic democrats republican news state bill israel	american bush america
	religion	god jesus christ christian lord	god christmas lord jesus
	emotions	--	love like feel happy want thank
	proper names	john mr george paul al michael david dave bill mike tom james	john don mr chris sarah miss matt josh mike david michael dr harry
	place	iraq america india york china canada israel california new	york america
	family	tyke	mom baby dad
	sports	game	--
	technology	blog tv microsoft cd dvd google system web pc media software linux games	tv
	time	sunday saturday friday july monday june august thursday	friday saturday sunday monday july christmas thursday now tuesday today wednesday june
	other	urlink english	urlink cute
	characters	p ii	c p b j d

Table 3.42 Content words from the top 100 keywords

It is worth pointing out, however, that politics as a category would be likely to have a different distribution now after the 2016 race between candidates Donald Trump and Hillary Clinton than it did in 2004, as the 2016 election brought gender into the political sphere in multiple ways, more so than any previous American election. In addition, because the original BAC was collected in 2004, and contains largely data from 2003-2004, a time in which a United States presidential election was taking place, other non-presidential election periods of data collection would likely yield less prevalent discourse. This, again, would affirm Schler et al.'s position that, "despite the strong stereotypical differences in content between male and female bloggers seen above, stylistic differences remain more telling than content differences," (2006).

C. Overview

Overall, the above analysis of the eBAC confirmed many (but not all) of the findings of both Bamman et al. (2014), Schler et al. (2006), and other scholars, most especially where the two agree. Table 3.42 below shows an overview of the features in which there is an agreement of likely gender indication between Bamman et al. (2014), Schler et al. (2006), and the eBAC.

#	Feature	Schler	Bamman	eBAC
1	Overall Pronouns	female	female	female

	alternative spellings <i>u, ur, yr</i>	--	female	female
	alternative spellings <i>ure, ya, yall, urself, thee, thy, tine, thou</i>	--	--	female
	alternative spellings <i>ye, yer, yoselves, yaself, yerselves</i>	--	--	male
2	Overall Emotion Terms	female	female	female
	negative emotion terms	female	--	female
	positive emotion terms	female	--	female
4	Overall Kinship Terms	female	female	female
	mom, mommy, sister, daughter, aunt, auntie, grandma, kids, child, dad, husband, hubs, etc.	female	female	female
	wife	--	male	male
	father, grandfather	--	--	male
6	Overall Abbreviations	female	female	female
	lol, omg	--	female	female
7	Overall Punctuation	--	female	female
	... ! ?	--	female	female
9	Overall Backchannel Sounds	--	female	female
	ah, hm, ugh, gr	--	female	female
	o, oi, ai, unf, er	--	--	male
	ragh, agh, ouch, ow, hum, nya, oy, aa, argh, uh, ew, ooh, aw, gr, ugh, um, ah, hm, oh	--	--	female
	<i>o, oh, ooh</i> variations	--	--	male
	other backchannel sound variations	--	--	female
10	Overall Hesitation Words	--	female	female
	um umm	--	female	female
	er	--	--	male
	uh, ermh, hm, hum	--	--	female
11	Overall Assent Terms	female	female	female
	okay, yes, yess+	--	female	female
	yeah, yas, ok, yeh, yep, yup, yea	--	--	female
	yis	--	--	male
13	Overall Swears and Taboo Words	male	male	male
	Anti-swears	--	female	female
	swears dammit, fag, bitch, ass	--	--	female
	other swears	--	--	male
	anti-swears darn(it), dang(it), gosh	--	--	female
	<i>f</i> -abbreviations (<i>lmfao, omfg</i>)	--	--	male
16	Overall Conjunctions	--	female	female
	&	--	female	female
	n, but	--	--	female
20	Overall Keyword category "politics"	male	--	male
	"technology"	male	--	male

Table 3.43 eBAC features in agreement with Bamman et al. (2014) & Schler et al. (2006)

Of the features overviewed in the table above, 11 out of 19 match with the overall findings of either or both Bamman et al. (2014) and Schler et al. (2006) with no conflicts between the two sources or the eBAC, along with additional variations of these features from the analysis above that were found to be more frequent in either the male or female subcorpora.

Table 3.44 below demonstrates the 8 remaining features in which there is some or total disagreement with the findings of either Schler et al. (2006) or Bamman et al. (2014), with 3 of these (negation terms, prepositions, and articles and determiners) being instances where Bamman et al. (2014) and Schler did not initially agree in their findings. Differences are indicated in **bold**.

#	Feature	Schler	Bamman	eBAC
5	Overall Friendship Terms	female	female	female
	bestie, bff, best friend	--	female	female
	bro, bruh, bros, brutha	--	male	mixed

	bro, brotha, pal	--	--	female
	bruh, brah, dude, mate	--	--	male
8	Overall Expressive Lengthening	--	female	male
	lengthened backchannel sounds	--	--	male
	lengthened abbreviations	--	--	male
	lengthened assent terms	--	--	female
	lengthened negation terms	--	--	female
	lengthened hesitation words	--	--	female
	lengthened punctuation	--	--	female
12	Overall Negation Terms	female	mixed	mixed
	noo+, cannot	--	female	mixed
	nah, nobody, ain't	--	male	mixed
	noo+, nah, nah+, no, 't	--	--	female
	cannot, nobody, ain't	--	--	male
14	Overall Prepositions	male	none	male
	2 for to	--	male	female
	w/ for with	--	female	female
	4 for	--	--	male
	b4 for before	--	--	female
15	Overall Alternative Spellings	--	female	female
	vacay, yaay, lol, u, ur, yr, w/	--	female	female
	bro, bruh, brutha, nah, ain't	--	male	mixed
	bro, nah	--	male	female
	brutha, ain't	--	male	male
17	Overall Articles and Determiners	male	none	male
	Pronouns	female	female	female
	Negation pronoun no	female	female	female
18	Overall Post Length	female	--	male
19	Overall Hyperlinks	male	--	female

Table 3.44 eBAC features not in agreement w/ Bamman et al. (2014) & Schler et al. (2006)

What is not explicitly discussed by Bamman et al. (2014) and Schler et al. (2006), however, and indeed what is not derived from the analysis in this chapter, is which of these features might be the most useful determiners of an author's gender (and thus may be most helpful in authorship analyses and linguistic profiling), which of these features might be more subject to context and genre than others (and thus require a more qualitative approach in conjunction with the quantitative approach taken here), and which of these features, if any, may be the most useful in determining not likely gender, but likely gender guising.

In the chapters below, the results of this analysis of the eBAC will be used as a baseline where appropriate, and the method established in this chapter will be applied to a variety of other datasets, and in conjunction with other types of analyses, with the goal of answering these questions.

Part 3 – Linguistic Theories

While Chapter 2 above focused primarily on the concept of language and gender performance, this section focuses on the idea of identity more broadly, most especially with regards to the perception of identity online and in CMC, and the interplay of other facets of identity with the facet of gender. This section covers multiple linguistic theories (audience design, accommodation theory, community of practice, and speaker contamination) and considers how such theories can be applied to identity performance and perception in CMC. As stated in the introduction to this chapter, while the theories suggested here and in Chapter

2 are applied throughout the analyses in this paper, they are not intended to be either singular or exhaustive.

Bucholtz and Hall (2004) conceptualize identity within the constraints of four semiotic processes: practice, indexicality, ideology, and performance, which are applied in varying degrees to other linguistic theories and the CMC examples below. They term this model of combined semiotic processes “tactics of intersubjectivity”, which “provides a more systematic and precise method for investigating how identity is constructed through a variety of symbolic resources, and especially language,” (Bucholtz & Hall, 2004). Bucholtz and Hall (2005) further outline five principles key to discussing the relation of language and identity: the emergence principle, the positionality principle, the indexicality principle, the relationality principle, and the partialness principle, as are discussed and applied specifically to CMC examples and the theories of audience design, communities of practice, and speaker contamination below. These principles are posited in reaction to ideas of essentialism, which “has been vulnerable to charges of operating within overgeneralized notions of similarity and difference,” (Bucholtz & Hall, 2004), and which “maintains that those who occupy an identity category (such as women, Asians, the working class) are both fundamentally similar to one another and fundamentally different from members of other groups,” (Bucholtz & Hall, 2004).

Within the essentialist approach, our identities would be considered positionally situated either in parallel or contrast to the identities of others, and derived from naturally, inevitably forming in- and out-groups (such as the suggested women, Asians, and the working class, positioned as internally similar to other women, other Asians, and other working class, and as externally dissimilar to, for example, men, non-Asians, and the upper class). Such an approach discounts multiple, sometimes conflicting identities a single individual can access—most especially online, when a pervading if erroneous sense of anonymity encourages interactants to perform in identities they might otherwise keep internalized, and sometimes with the wrong in-group—and the fact that even shared identity qualities, such as gender, race, and class, can be accessed and performed differently from individual to individual.

A. Audience Design

CMC is unique in that it exists somewhere between, and yet not quite inside the frameworks of either spoken or written language. Crystal defines speech as “time bound, dynamic, and transient; it is part of an interaction in which both participants are usually present, and the speaker has a particular addressee (or several addressees) in mind,” (Crystal, 2011). Contrastingly, Crystal defines writing as “space bound, static, and permanent; it is the result of a situation in which the writer is usually distant from the reader, and often does not know who the reader is going to be,” (Crystal, 2011). Both explanations still consider Bell’s theory of audience design, as detailed below—the addressees Crystal mentions overlap neatly with Bell’s concepts of the addressee and audience/hearers, and the ambiguous potential readers to which Crystal refers can be considered to fall within Bell’s outer categories of bystanders

and eavesdroppers.

Throughout multiple works, Bell discusses his theory of audience design, which “assumes that persons respond mainly to other persons, that speakers take the most account of hearers in designing their talk,” (1984). According to Bell, “Audience Design (Bell, 1984) treats the speaker, the first person, as the primary participant at the moment of speech, qualitatively apart from other interlocutors,” (2014). The audience, as Bell defines it, is comprised of the addressee, who is the speaker’s addressed recipient, as well as the hearer/audience, which is comprised of: unaddressed recipients, or auditors, who are not addressed; bystanders as overhearers, whose participant status in the exchange are not ratified; and bystanders as eavesdroppers whose participant status in the exchange is not known.

Although, as Bell points out, “[o]ften in an interaction, the physical distance of audience members from the speaker coincides with their role distance, with the addressee physically closest and the eavesdropper farthest away,” when it comes to online discourse, physical distance is not a factor, or at least not one of gradation (2014). However, there are other factors of speaker and audience ‘proximity,’ albeit of the non-physical kind, that can be categorized in much the same way as Bell proposes audience roles in spoken exchanges.

Bell also concedes that “[o]ne main critique made of frameworks such as audience design which attempt to systematize style is that they are reductionist. They run the risk of minimizing or discounting the complexity of speakers’ moment-by-moment, self-expressive use of language,” but further points out that, “one could argue that an approach which is as richly person-oriented as audience design is by definition not reductionist,” (2001). However, CMC offers a unique window into the study of audience design, in that it mimics the communicative and often near-simultaneous or at least quick feedback properties of spoken language, while often keeping a written record of itself ripe for study. Although generalizations are unavoidable in any method for analyzing stylistic structures, they are generalizations better made, in that style and mode can both be conventionally and logistically dictated by CMC.

In discussing Bell’s theory of audience, Holmes points out that there is, “strong support for the view that the addressee or audience is a very important influence on a speaker’s style,” (2001). Although Holmes is analyzing the phonetic differences in the pronunciations of New Zealand radio newsreaders, the comparison is still apt online. Although there is no such thing as pronunciation in the formatting of written text, a valid analogue to this can come in the form of spelling and prose choices. Things like chat speak function in much the same way when it comes to audience design, as the writers often have an audience of other chat speakers in mind.

Rogers, Fay, Mayberry, & Roulin, (2013) point out that “[w]hile ‘design’ implies a thoughtful and strategic process of linguistic adjustment, it is equally possible that such adjustment is thoughtless and non-strategic,” (2013). By studying written explanations of geometric shapes, meant first for the individual writers, and then an audience of increasing

size, Rogers et al.'s goal was to "examine the extent to which audience design is a strategic, top-down and individual-level design process or a non-strategic, bottom-up and interaction-level process," with their results showing that audience design arises as a non-strategic outcome of social interactions during group discussions (2013).

Although Rogers et al.'s experiments do not analyze online discourse, they do consider written language meant for an audience made up of many unknown variables to the speaker. Much of what Rogers et al. found, however, can be seen to apply to messages formatted online, which are done so with an audience of unknown or unverifiable numbers, as Rogers et al. found that, "audience design during multiparty communication is an adaption achieved through monitoring and adjustment during social interaction rather than being pre-planned at the individual-level," (2013). An online poster does not necessarily take into consideration the entire potential scope of their audience until feedback indicates participants outside their register, though this initially non-strategic process can have unfortunate repercussions.

1. Audience design and identity

According to Bucholtz and Hall, rather than wholly a product of our internal beliefs about ourselves, "[i]dentity is best viewed as the emergent product rather than the pre-existing source of linguistic and other semiotic practices and therefore is fundamentally a social and cultural phenomenon" (2005). That is, identity is performative—and therefore heavily tied to our language production—just as it is shaped by both the societies and cultures in which we exist. This contrasts with the traditional, scholarly concept of "identity as housed primarily within an individual mind, so that the only possible relationship between identity and language use is for language to reflect the individual's internal mental state" (Bucholtz & Hall, 2005). Rather, the emergence principle seeks to count for both an individual's "identity inside the mind [and] the social ground on which identity is built, maintained, and altered," (Bucholtz & Hall, 2005).

Similarly, we attribute our internal assignments of others' identities based on their performance, linguistic, communicative, and otherwise, and, "[s]uch interactions therefore highlight what is equally true of even the most predictable and non-innovative identities: that they are only constituted as socially real through discourse, and especially interaction," (Bucholtz & Hall, 2005). Identity, while initially internal in its conception, is shaped by interaction as much as it is by interactants—or at least perceived interactants—which constitute the audience to which discourse and identity are performed.

As it applies to audience design, Bucholtz and Hall's (2004) "tactics of intersubjectivity" demonstrate the continuum through which identity can be seen first as the internalized process of "habitual social activity, that makes up our daily lives". This is the semiotic principle of *practice*, whereby our internalized social roles in various contexts are so ingrained within ourselves that we perform them below the level of conscious thought. A speaker, then, not giving proper consideration or design to an unusual as opposed to habitually expected audience for certain behaviors may find their identity performance at odds with what their true audience expects, as we will see in the examples below.

2. Audience design in online identity performance

Figure 3.5 below shows an exchange between three Tumblr users, as responded to by **responder**, between **poster** and an **anonymous** user, and eventually reposted by **reblogger** (names changed):

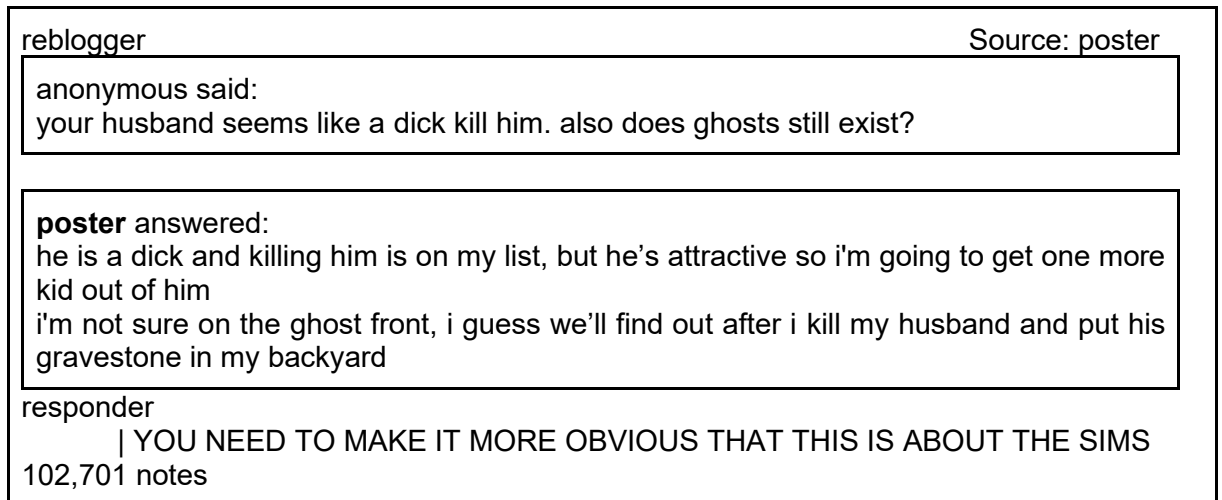


Figure 3.5 Tumblr posts between anonymous

Anonymous sent an “ask” to the **poster**, essentially a private message (PM), which can be responded to either publicly or privately. For **anonymous**, **poster** is the addressee and entire intended audience, hence the ask, and not, for example, a public response to one of **poster**’s previous posts.

Poster, on the other hand, replied to the ask publicly, meaning any visitors to the **poster**’s Tumblr could see the exchange—which is exactly what happened, as a **responder** replied to an ask that was not tagged for them. For **poster**, **anonymous** is the addressee, while **responder** would be considered an inconsequential overhearer or eavesdropper, as they are another Tumblr user, and likely follows (or subscribes to) the **poster**. Thus while the **poster**’s response to the ask is tailored for **anonymous** as the addressee, it is with the expectation that other bystanders, such as the **responder**, may make up the wider audience, and will understand the context.

It is, however, clear from the **responder**’s unsolicited reaction, that the context—that of the Sims, a life-simulation computer game—was not immediately apparent, or at the very least would not be clear to all other unintended overhearers or eavesdroppers. Neither **anonymous** nor the **poster** mention the game the Sims, although it is likely that **anonymous** knows the **poster** plays the game, either based on previous conversation, or earlier blog posts. Both have *practice*, as evidenced by their posting history, in accessing the shared identity of Sims players, an identity which includes in-group knowledge of the game that does not necessitate a contextual explanation by either **anonymous** or the **poster** for their discourse to be understood by either intended interlocutor.

Without the shared context that comes with habitual *practice* both in playing the Sims and discussing it with other players, however, and without the entirety of the exchange as shown in Figure 3.5 above, the conversation seems to be a grave if odd one about mariticide

and ghosts to any outside observer, rather than an innocent one about gaming. Although this exchange did not result in any sort of judicial action, as the threatened entity was never real, the fact that it is so easily and strikingly misunderstood (so much, in fact, that an unrelated **reblogger** decided to disseminate the post further—certainly, a bystander outside the **poster's** initially expected audience) is evident in the fact that the post received 102,701 notes as of early 2015, a large amount for a simple ask response.

The **poster** in this example obviously identified themselves as an avid Sims player at some point in their online activities, predicating **anonymous's** ask without any particular or necessary context. The circumstances of the ask led to the perception of a shared identity between the **poster** and **anonymous**, one which was not necessarily shared by all of the bystanders who inevitably stumbled across the public response. Given the potential implications of the exchange out of context, had the initial **poster** considered the larger audience of Tumblr and the internet in general, they may have been more careful with their method of response or the discourse used therein, rather than relying on a perceived shared identity of an unquantifiable audience to provide some very necessary context. Such shared identities and contexts often fall into the realm of community of practice, as we will see below.

B. Accommodation Theory

Like audience design, accommodation theory, as defined by Giles, Coupland, & Coupland (1991), is highly performative, in that it has to do with the audience for which a given discourse is (or isn't) constructed. According to Giles et al., "[t]here are many ways of performing acts we could deem to be accommodative, many reasons for doing or not doing so, and a wide range of specifiable outcomes," (1991). Accommodation theory can account for tactics employed in Audience Design both when constructed discourse is meant to highlight a shared identity—such as via 'convergence', which is "a strategy of identification with the communication patterns of an individual internal to the interaction,"—and when constructed discourse is meant to highlight differing identities—such as via 'divergence', which is "a strategy of identification with linguistic communicative norms of some reference group external to the immediate situation," (Giles et al., 1991). Giles et al. (1991) outline a number of strategies by which speakers may or may not accommodate, many of which are briefly touched upon in the paragraphs below.

'Convergence' is "a strategy whereby individuals adapt to each other's communicative behaviors in terms of a wide range of linguistic-prosodic-nonverbal features including speech rate, pausal phenomena and utterance length, phonological variants, smiling, gaze, and so on," (Giles et al., 1991), while 'divergence' "is the way in which speakers accentuate speech and nonverbal differences between themselves and others," (Giles et al., 1991). A speaker may, for example, choose to 'converge' *upward* or *downward* to match their interlocutors as a way of being cooperative, and accessing a shared or mutually acceptable identity, whereas another speaker may choose to 'diverge' in order to be confrontational, or highlight some in-group identity status to which the other interlocutors do not have access. Although these

nonverbal cues have been applied almost strictly to spoken language alone, because it is communicative, analogues do exist in CMC, as we will see below.

Speakers may 'overaccommodate' when they "overshoot" even at full convergence and 'hyperconverge'," such as by hypercorrecting to attempt to match a higher educational, social, or hierarchical level (Giles et al. 1991). Conversely, speakers may 'underaccommodate' when they "undershoot the level of implementation desired for a successful interaction (Coupland et al., 1998)" (Giles, 2016), such as speakers failing to speak slowly enough or avoid slang and jargon when conversing with a non-native interlocutor still learning the language. Finally, 'disaccommodation' is "when people switch registers in repeating something uttered by their partners [...] as a tactic to maintain integrity, distance, or identity when misunderstanding is not even conceivably an issue," such as an adult "properly" rephrasing a child's utterance, or a child correcting an adult's use of slang (Giles et al., 1991).

Speakers may employ a number of tactics of disaccommodation for various reasons, including: 'speech maintenance', 'noble selves', 'perceptual convergence', and 'perceived vitality'. 'Speech maintenance' "is a valued (and possibly conscious and even effortful) act of maintaining one's group identity," (Giles et al., 1991), such as a student refusing to use "proper" English when conversing with an authority figure in a formal setting, despite what would be considered proper and polite, maintaining, instead, the in-group identity of their age group and whatever social cliques they may subscribe to. 'Speech maintenance', then, at its most extreme, would categorize such interlocutors as 'noble selves', who "would be predicted to maintain their idiosyncratic speech and nonverbal characteristics across many situations [and] are those straightforward, spontaneous persons who see deviation from their assumed 'real' selves as being against their principles and, thus, intolerable," (Giles et al., 1991).

While speakers may not be able to employ many of the same non-verbal tactics online as they do in verbal and/or face-to-face communication, speakers can and do 'converge' and 'diverge' over CMC platforms. This may come down to the formality or informality of their writing (such as very formal written language in, for example, an email to one's boss, but very informal "chat speak" in texts to one's children), the frequency, speed, and length of responses in a chat platform, the formatting of text, the use of emoticons, even the communicative platform used, and so on. How accommodation theory applies to CMC, specifically, is discussed with regards to Figure 3.6 below.

1. Accommodation theory and identity

Taking it a step further in communication, the positionality principle asserts that, "[i]dentities encompass (a) macro-level demographic categories; (b) local, ethnographically specific cultural positions; and (c) temporary and interactionally specific stances and participant roles" (Bucholtz & Hall, 2005). That is, identity is *not* "simply a collection of broad social categories. This perspective is found most often in the quantitative social science, which correlate social behavior with macro identity categories such as age, gender and social class," (Bucholtz & Hall, 2005).

While “the traditional forms of these approaches have been valuable for documenting large-scale sociolinguistic trends; they are often less effective in capturing the more nuanced and flexible kinds of identity relations that arise in local contexts,” (Bucholtz & Hall, 2005). That is, speakers often modify their discourse and identity performance depending on the speech event, their interlocutors, and other extralinguistic factors, and do not always perform the same identity across every single set of variables, and often multiple identities are performed at once. Instead, a specific identity—most especially when the idea of shared identities, or the need for differentiating identities, arise—may be performed given not only the context, but at least the perceived audience, and the need to associate (‘converge’) or dissociate (‘diverge’) with them.

Multiple factors influence whether, how, and to what degree a speaker may ‘converge’ or ‘diverge’, such as how interested they are in gaining the approval or cooperation of their interlocutors, or how interested they are in maintaining distance and authority, and so on. According to Giles et al., “[f]actors that influence the intensity of this particular need include the probability of future interactions with an unfamiliar other, an addressee’s high social status, and interpersonal variability in the need for social approval itself,” (1991). Major influencing factors may come down to what Giles et al. term ‘social vitality’, which “refers to the extent to which group members of a social group consider sociostructural factors to be operating in their favor or not,” (1991).

So a non-standard English speaker in, for example, a court of law where highly standardized and specialized English is used may be perceived to have an identity and related language with a lower ‘vitality’ relative to the in-groups of the court and will thus be forced to attempt to ‘converge’ to accommodate the speech event, as they cannot expect the higher ‘vitality’ court to ‘converge’ to accommodate them. Such a scenario could easily lead to hypercorrection on the part of the interlocutor with a lower perceived ‘vitality’, and ‘underaccommodation’ by the interlocutors with a higher perceived vitality should “correct” and “proper” accommodation fail.

2. Accommodation theory in online identity performance

Demonstrated in Figure 3.6 below are two separate SMS conversations with two separate mother and child pairs, both demonstrating a misuse by “Mom” of the common chatspeak feature “WTF”⁷.

⁷ As demonstrated in the example definition from UrbanDictionary.com below, the general consensus of the definition “WTF” with regards to CMC is “what the fuck”:

wtf

Generally stands for 'What the fuck'. Most people use a question mark afterwards to get the point through. Rather than using the same term for the other 'w's, who, when, where, and why, it makes more sense to actually state the word and follow it with 'tf'

Capitalization doesn't really matter.

This term can also be likened to 'What the shit?' which is more comical and has a tantamount meaning.

WTF?

Mom ⁸ Nov 20, 2012 5:46 PM 1 When are you coming 2 home? 3 4 Hello?? WTF?? 5 6 Mom! Do you know 7 what that means? 8 9 Yeah, are you coming 10 home Wednesday, 11 Thursday, or Friday? 12 13 Nope. Not what that 14 means.	Mom ⁹ 15 I got an A in Chem! 16 17 WTF , well done! 18 19 Mom, what do you 20 think WTF means? 21 22 Well That's Fantastic
--	---

Figure 3.6 “Funny Texts from Mom”

In both cases, “Mom” attempted to ‘converge’ *downward*, if not with the child in question, then with the general language variety of text or chat speak, by adopting the lower register, as it has a higher ‘vitality’ in the context of texting. “Mom”, with the audience of her child in mind, a child who likely employs chatspeak such as “wtf” on a regular basis (or who at the very least shares an identity with an age group stereotypically thought to use such language), attempted to access this same identity by using the same language. Although both “Moms” had differing understandings of what “WTF” means, even in relation to one another, both were wrong, and both failed to successfully converge and access a shared identity with their child in exactly the same way.

We find similar ‘convergence’ (or at least ‘non-divergence’) in Figure 3.5 in the previous section, where both **anonymous** and **poster** employed non-standard written language to varying degrees, eschewing, for example, any standard capitalization, and much standard punctuation. Despite such writing having a lower ‘vitality’ relative to standard written English, because the non-standardness is shared, it has a vitality worth maintaining. While the **responder** also used non-standard capitalization, their use differed and ‘diverged’ significantly from the original interlocutors, the all-caps emphasizing the likely perceptions of out-group members (an identity which the **responder** is considering when making their point) as distinct from the in-group of **anonymous** and **poster**.

C. Community of Practice

Given that CMC genres tend to overlap and employ multiple or alternating registers at any given time, it is important to consider theories such as communities of practice (CoPs) (here broadly including the more specific contexts of speech communities and social networks)

WhoTH?
WhenTF?
WTS?

⁸ <http://static.boredpanda.com/blog/wp-content/uploads/2014/04/funny-texts-from-parents-9.jpg>

⁹ <http://justsomething.co/wp-content/uploads/2014/04/funniest-parents-texts-28.jpg>

when analyzing CMC genres. Johnstone (2009) covers discourse analysis as related to power and community, indexicality, stance and style, social roles and participant structure, audience, politeness, accommodation, social identity and identification, personal identity: discourse and self, and the linguistic individual in discourse. Participant analysis becomes an important point in the analyses done below, as it is precisely the differing communities, stances, social and personal identities, and ideas of audience, politeness, and accommodation that are at issue in many of the forensic cases explored. Johnstone's work illustrates the importance of individual participants in a given exchange in order to better interpret the context of the language analyzed.

According to Shuy, "Speech acts are, quite simply, the ways that people use language to get things done. Since resolution [...] hinges on whether or not the accused persona actually offers, agrees, threatens, admits, lies, promises or requests, accurate identification of these speech act is, to say the least, crucial," (1993). Although the primary focus of this thesis is not to consider the pragmatics of a given exchange, this is nevertheless pertinent, in that CoPs often dictate the framework of speech acts and play a heavy hand in infelicitous illocutionary statements being taken as felicitous locutionary ones by outsiders. Facilitators of CMC, in their attempt to accommodate specific needs, fields, or interests, are rife with CoPs, forum sites, for example, often housing as many registers as they host subforums.

Because of this, as Shuy points out, "[i]t is extremely dangerous to isolate anything from context, especially words," (1993). This holds true for all CMC just as it does for any exchanges, as, without context, as we saw with Figure 3.5 above, we as unattuned listeners are left primed to misunderstand. This is especially true when online users must balance multiple, sometimes conflicting or overlapping identities online. These identities may, as we will see in the example below, be heavily stratified, and largely compartmentalized.

That is, in your public, offline life, you may identify, for example, as a student, at a particular university, in a particular program, and you may access this identity online, by claiming your status on Facebook, joining or following accounts and groups associated with your school and your schoolwork, and interacting with classmates online. Such an identity and the associated CoPs would be widely acceptable to display not only for your classmates, teachers, and other school contacts, but for your friends, family, coworkers, and other acquaintances otherwise unaffiliated with your studies. Other identities and CoPs in which you participate online, however, may be less fit for public consumption, such as, in the example we will see below, not only identifying as a gamer—an unprofessional perceived pastime if you want future employers to think you a responsible adult—but as part of a gaming community known for hyperbolic language and interactions that are just as inappropriate to engage in with your classmates as they are to perform for the general public.

1. Community of practice and identity

Bucholtz and Hall describe the indexicality principle, from the idea of indexicality in linguistic anthropology, as:

Identity relations emerge in interaction through several related indexical processes, including: (a) overt mention of identity categories and labels; (b) implicatures and presuppositions regarding one's own or others' identity position; (c) displayed evaluative and epistemic orientations to ongoing talk, as well as interactional footings and participant roles; and (d) the use of linguistic structures and systems that are ideologically associated with specific personas and groups. (2005)

That is, "while the first two principles [...] characterize the ontological status of identity, the third principle is concerned with the mechanism whereby identity is constituted," (Bucholtz & Hall, 2005). Indexicality of identity and shared identities shape the appropriate discourse of CoPs, which may, online, be as superficial as chatspeak and jargon, or as misleading to an outside audience as strongly hyperbolic language such as infelicitous threats.

In their "tactics of intersubjectivity," *indexicality* "is the semiotic operation of juxtaposition, whereby one entity or event points to another," (Bucholtz & Hall, 2004). This can be related to speech acts in much the way speech acts can be related to CoPs. If "where there's smoke, there's fire," smoke as a the locutionary act and fire as the illocutionary act will predicate different perlocutionary acts depending upon the context in which smoke, which indexes a likely fire, is understood. The perlocutionary result could, for example, be fear, if the smoke and thus the fire is coming from inside someone's house and indicating an unexpected hazard, or relief for a person lost in the wilderness in search of warmth and aid and indicating potential help. While an extreme example, as we will see in the CMC example below, what the locution and even the illocution of a speech act indexes to an interlocutor can vary wildly depending upon what CoPs that interlocutor has access to.

Bucholz and Hall take this one step further in the performance of shared (or purposefully distinct) identities with the relationality principle, which, "emphasizes identity as a relational phenomenon," which they outline "first, to underscore the point that identities are never autonomous or independent but always acquire social meaning in relation to other available identity positions and other social actors; and second, to call into question the widespread but oversimplified view of identity relations as revolving around a single axis: sameness and difference," (2005).

Bucholtz and Hall outline "several, often overlapping, contemporary relations, including similarity/difference, genuineness/artifice, and authority/delegitimacy," when developing relational identities (2005). All three can be seen to relate to the in- and out-group status of CoPs. In-group status may result from adequation, in which, "differences irrelevant or damaging to ongoing efforts to adequate two people or groups will be downplayed, and similarities viewed as salient to and supportive of the immediate project of identity work will be foregrounded." (Bucholtz & Hall, 2005), while outgroup status may result from distinction, which "focuses on the identity relation of differentiation," (Bucholtz & Hall, 2005).

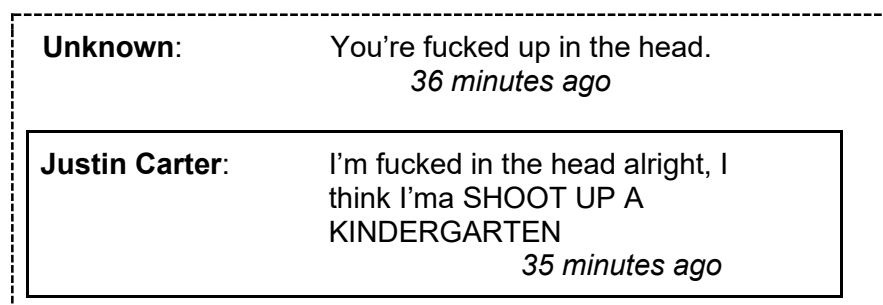
Within the concepts of adequation and distinction are the relational aspects of authentication, "by which speakers make claims to realness and artifice, respectively. While both relations have to do with authenticity, the first focuses on the ways in which identities are

discursively verified and the second on how assumptions regarding the seamlessness of identity can be disrupted.” (Bucholtz & Hall, 2005), and denaturalization, in which “[w]hat is called attention to instead is the ways in which identity is crafted, fragmented, problematic, or false. Such aspects often emerge most clearly in parodic performance and in some displays of hybrid identity, but they may also appear whenever an identity violates ideological expectations,” (Bucholtz & Hall, 2005). Relational identities, once established, can be authorized, which “involves the affirmation or imposition of an identity through structures of institutionalized power and ideology, whether local or translocal,” or illegitimized, which, “addresses the ways in which identities are dismissed, censored, or simply ignored by these same structures” (Bucholtz & Hall, 2005).

How these contemporary relations are both performed and perceived falls into the “tactics of intersubjectivity model” with the concept of *ideology*, which “organizes and enables all cultural and beliefs and practices as well as the power relations that result from these,” (Bucholtz & Hall, 2004). Like *indexes*, *ideologies* relate to CoPs in much the same way, determining, often through shared *practice*, what shared knowledge and thus performative tactics interlocutors have when accessing a shared or opposing identity. Bucholtz and Hall point out that *ideology* can eventually lead to, “the creation of a naturalized link between the linguistic and the social that comes to be viewed as even more inevitable than the association generated through indexicality,” (2005). Many of these elements come into play in the example below.

2. Community of practice in online identity performance

On February 14, 2013, just months after the Sandy Hook Elementary school shooting, 17-year-old Justin Carter was arrested for posts he made in a forum on Facebook, which were construed by the authorities as credible threats against a kindergarten. Carter was charged and convicted on the basis of a screenshot of the inculpatory exchange, which was sent to the police by another Facebook user, despite the fact that Carter alleged the threats, while in poor taste, were meant as a joke, and the fact that no other incriminating evidence was found to suggest he had any intent or ability to follow through with his statements. Although Facebook never released the entirety of the thread, Carter was convicted on the basis of the screenshot, the contents of which can be seen in the inner, solid-lined box in Figure 3.7, while the outer, dotted-line box includes parts of the conversation not shown in the screenshot, but attested to by a number of sources who were privy to the exchange before it was removed.



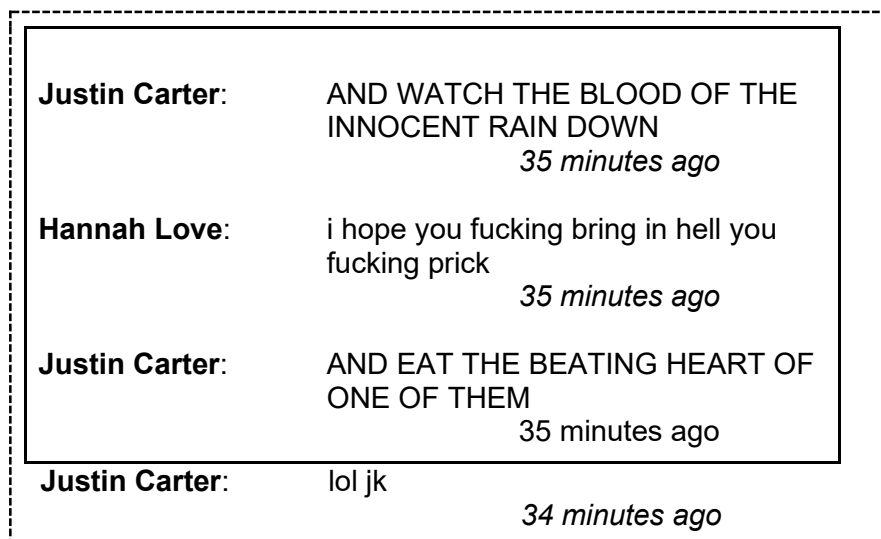


Figure 3.7 Justin Carter's Facebook forum posts, February 14, 2013

There are a number of aspects that go into an understanding of this particular case. Most broadly is the fact that, Facebook, a social media site that employs forums as one of its many applications, has heavily tied accounts. That is, users provide ample personal information, and in Carter's case, police not only knew who he was, but were able to retrieve his permanent address, which was just miles from a kindergarten, though Carter was living in another city at the time. The thread Carter was participating in, however, was topical to the "massive multiplayer online role-playing game" (MMORPG) League of Legends (LoL), which is only a remotely tied account, requiring only a valid email address and no payment information, though all games have methods of tracking a user's IP.

As a CoP, LoL is somewhat notorious in the online gaming community, known to have a very competitive atmosphere, that often leads to somewhat severe and hyperbolic discourse, both in-game, and on the game's official forum. Although the Facebook thread in which Carter was posting was viewable to other, non-in-group members, Carter's—and indeed the other in-group thread participants'—intended audience were other LoLers, who would have expected the 'convergence' of such graphic, hyperbolic discourse as par for the course within their discourse community, locutionary speech acts *indexing* not genuine, realis indirect threats. Facebook, on the other hand, would *not* expect such explicit threats, jokingly or not, as the utterances are tied to a person's reality, thus giving them real-world implications, and *indexing* them as having the illocutionary purpose of indeed making realis threats. Carter's discourse obviously was not modeled with the appropriate audience or shared *ideologies* in mind, as considering outsiders and authority figures as bystanders would likely have produced less violent discourse.

It is clear from Love's response to Carter's discourse that they do not share the same in-group status to the LoL CoP, and that because of this the perception would be of Carter as overtly 'diverging' from the vitality norm Love expects. While Love herself uses highly hyperbolic language in reaction to Carter's, their discourse is distinct in that LoL's common discourse expects threats, either direct or indirect, to be performed as infelicitous speech

acts.¹⁰ In Carter's view of his own identity, his threats authenticate and authorize his identity within the LoL CoP. But because he is accessing this CoP on Facebook where the same CoP features are neither the norm nor an acceptable alternative, Love's view of Carter's identity performance is as one that is denaturalized, and does not fit into what is expected to be performed on Facebook, and indeed it does not fit into what is acceptable according to US law, as Carter is reported, arrested, and indicted based upon his language and identity performance to the wrong audience and in the wrong community.

D. Speaker Contamination

Shuy discusses a number of conversational strategies used to create language crimes, or otherwise contaminate language, when none have actually occurred, and misconceptions involving language in law cases, as we saw in the case of Figure 3.7 above. Shuy offers the following three misconceptions about defendants in court cases where language evidence is considered:

- 1 If they are on the tape at all, they must be guilty of something; otherwise the police would not have been after them.
- 2 If they are guilty of one of the charges, they are probably guilty of the other charges as well.
- 3 The defendants hear, understand, and remember everything said by the agent or other persons in the taped conversation. (1993)

Although CMC does not immediately qualify it as a language crime, as we saw in Figure 3.5 above, aspects of the principle relate. There is, in a given thread, chat, or other CMC chain type, the misconception that each post follows another in much the same way as dialogue in a novel or script does—that if an utterance is followed by another utterance, the second utterance is a response to the first, and that each participant fully heard (or read) and understood each preceding utterance when constructing their responses. Because internet forums and chats in particular are laid out in chronological order and are essentially a written record of the conversations therein, this assumption is exacerbated.

So the third misconception, as applied to forums, can be seen to be as follows:

- 3 The defendants hear, understand, and remember everything written by the other persons in the thread.

This misconception follows along with Shuy's idea of speaker contamination, which Shuy

¹⁰ Such identity performance is so highly reified that the internet has even come up with a memetic term for those who default to such discourse, as in the example definition from UrbanDictionary.com below (emphasis added):

edgelord

A poster on an **Internet forum**, (particularly **4chan**) who expresses opinions which are either strongly nihilistic, ("life has no meaning," or Tyler Durden's special snowflake speech from the film Fight Club being probably the two main examples) or contain references to **Hitler**, Nazism, fascism, **or other taboo topics which are deliberately intended to shock or offend readers.**

The term "edgelord," is a noun, which came from the previous adjective, "**edgy**," which described the above behaviour.

Nietzsche was an edgelord before it was cool.

[#internet](#) [#4chan](#) [#forum](#) [#hitler](#) [#edgy](#)

refers to as an “‘all or none’ logical fallacy,” (1993). Shuy offers the following example to explain how speaker contamination works:

Suppose somebody recorded us at a party when we were cornered by a person we do not particularly like. That person insists on telling us an off-color joke. What we might *like* to do is to tell that person that the joke is inappropriate or insulting, but social constraints, personal weakness or sheer indifference cause us to go along with it, laugh politely and try to excuse ourselves from that person as quickly as possible. Then, later on in the evening, the person who secretly taped the incident plays it back and observes that we were *all* telling dirty jokes. (1993)

Shuy offers many other applicable scenarios of speaker contamination aside from not refuting a statement, including not being around when the statement was actually made, or missing it because of overlaps or volume, and so on. As a record of language, threads can be misconstrued in much the same way—ignoring another user’s post, joining a thread discussion late or leaving it early, or simply missing out on a post are all lost in the available data. It’s impossible to say who read what when, despite what the structured formatting may suggest, an issue that must be especially considered in forensic cases related to CMC.

1. Speaker contamination and identity

Unlike audience design and communities of practice, speaker contamination often has more to do with how an identity is perceived than how an identity is performed. Bucholtz & Hall explain this with the partialness principle, in which:

Any given construction of identity may be in part deliberate and intentional, in part habitual and hence often less than fully conscious, in part an outcome of interactional negotiation and contestation, in part an outcome of others’ perceptions and representations, and in part an effect of larger ideological processes and material structures that may become relevant to interaction. It is therefore constantly shifting both as interaction unfolds and across discourse contexts. (2005)

Of importance as the partialness principle relates to speaker contamination is “others’ perceptions” of our *performance*, the final principle of the “tactics of intersubjectivity” model which is “a highly deliberate and self-aware social display,” (Bucholtz & Hall, 2004).

Identity is thus performed within an outside-imposed concept of structure and agency for the performer, which are “intertwined as components of micro as well as macro articulations of identity,” (Bucholtz & Hall, 2005). Bucholtz & Hall define agency as something that “becomes problematic only when it is conceptualized as located within an individual rational subject who authors his identity without structural constraints,” (2005). Identity performance, then, may or may not be intentional, and “habitual actions accomplished below the level of conscious awareness act upon the world no less than those carried out deliberately.”

We saw this phenomenon of *practice* in both Figure 3.5 and Figure 3.7, where access to a shared but compartmentalized identity manifested itself in the larger world, to varying degrees of reaction, understanding, and consequence, just as we will see again in Figure 3.8 below. We saw also the issues of *performance* when an identity is not shared with one’s audience, most especially in the case of Figure 3.7, where, while Carter’s threat-based

performance was indeed highly deliberate, what he intended to *index* were the shared ideologies of his target CoP and audience, and not the unintended overhearers on Facebook such as Love, who do not share the same *ideological indexes*.

2. Speaker contamination in online identity performance

In August 2013, 16-year-old high school junior Tyson Leon was suspended from his school's athletics programs after tweeting Figure 3.8 below. Leon's school took the tweet as a threat against his fellow teammates, while Leon insisted his use of the inflammatory term "drill", which has a number of meanings in "drill" lyrics (including talking about gang activities and, most inculpatory, to kill), was not a threat, but rather a common sports term, meaning to run practices. Much as in the case on Tumblr discussed above, it was context—in this case that of sports rather than a video game—that was missing in determining whether Leon was attempting to 'converge' or 'diverge' from his intended audience, leading to the erroneous interpretation of the utterance.

In Leon's case, school officials were doubtfully his intended audience, making them eavesdroppers who wouldn't necessarily be privy to the context he assumed was shared knowledge between himself and his intended audience—in this case, specific jargon related to the CoP of sports and other athletes, for whom such jargon would constitute a proper 'convergence'. But as Leon's Twitter is a tied account—tied to his real-world identity, and accessible by school officials—his audience was wider than he netted for, leading to perceptual 'divergence' from the expected standard.

Im boutta drill my 'teammates' on
Monday.

Figure 3.8 Tyson Leon's tweet, August 2013

In this case, there is also a degree of speaker contamination, applied to Leon himself based on his past behavior and discourse. Leon was contaminated not by the other participants in an exchange, but by his own reputation—Leon had received multiple warnings previously for other behavioral misconduct, and thus his unintended audience was primed, through schema, to expect he would continue with this same pattern, leading to a contaminated interpretation of this particular tweet as *indexing* not only a genuine threat, but a threat at all. Such an issue as this case demonstrates the permanence of a person's posting history in coloring their subsequent language and identity performance, leading to the possibility not only of speaker contamination, but self-contamination. In Leon's case, his identity performance online was compounded by his school's perception of his regular behavior, and thus the eventual disciplinary actions taken against Leon were, as Bucholtz & Hall point out, partially due to the language itself, and partially due to a pre-existing perception not only of Leon's behavior, but of the appropriate behavior a student should demonstrate in school and online. The school may have additionally assumed that Leon was performing as part of a CoP of gang-affiliated youths, while Leon instead claimed to be performing as part of a CoP of team-based athletics.

Unlike Leon, perhaps the **poster** of Figure 3.5 may have benefitted from the partialness principle of identity performance and perception. Rather than contaminate their discourse as a legitimate threat, if the **poster** had suffered any legal or disciplinary action for their words, past posts concerning the Sims may have provided the appropriate, non-threatening context for the post, just as Leon's past disciplinary issues provided the inappropriate context for his tweet as a legitimate threat.

In this case in Figure 3.7, perhaps one of the most important extralinguistic features of the genre is the timestamps on the posts. The *minutes ago* markers provided in Figure 3.7 are the times as given by the screenshot sent to the police, which show that the exchange took place in a relatively condensed amount of time. Because of the nature of Facebook's threads, more rapid-fire discourse is expected, emulative of chatlogs and IMs, in part because the posts are required to be shorter than on a more traditional forum site. This leads into the idea of speaker contamination, or, at the very least, conversation contamination.

It is the nature of written communications that we as observers assume that if something was posted in a certain order, it was considered by each and every interlocutor in the same order, and that a perfect exchange of turn-taking took place. In court, it was alleged that Carter's final post, "AND EAT THE BEATING HEART OUT OF ONE OF THEM," was meant to further antagonize Love, who stepped up to point out that what he was posting was messed up, in an equally colorful way. But the nature of the utterance, a clause linked by and, certainly makes it a continuation of his previous post, regardless of whether this was meant to escalate his discourse for Love as a newly intended addressee. Although the timestamps do not provide counts of seconds, what is *more* likely in this case is that Carter continued on with his own discourse *regardless* of what Love had interjected with, and not because of it, possibly because her post was ignored, or Carter missed it at the time of posting his last screen-captured remark, though his discourse surrounding Love's comment was contaminated nevertheless, as continuing in spite of an adverse reaction on her part.

Although, especially given the lack of second counts, timestamps are not absolute when it comes to understanding the turn-taking—if indeed any intentional turn-taking took place—in an exchange, when considering forum-based CMC, it is irresponsible to dismiss them as non-data. Further irresponsible is to prosecute a case based solely on an excerpt of an exchange. Based on the fact that the posts were so rapid-fire, that the screencap was not taken to be sent to the police until 35 minutes after the inculpatory utterances had been posted, and that at least "lol jk"—which would have appeared *far* before the following 35 minutes had passed—was attested by multiple conversants and viewers, it is highly likely that the exchange continued beyond what the screencap shows. Without the entire context of the exchange—any reference to the CoP discourse being employed, the intended audience, Carter's own assertion that he was just kidding, the fact that his utterances were in response to a prompt, what occurred after the out-group Love jumped in, and so much more—the best that can be said about this data is that it is woefully incomplete, and poor evidence, linguistic or otherwise,

without anything to corroborate it.

Although gender in particular did not play much of a part in the examples provided in this chapter, as we will see below in Chapter 4 further analyzing the Tinder data, audience design, accommodation theory, community of practice, and speaker contamination all play a role in how gender is both perceived and performed.

Part 4 – Discussion/Conclusion

This chapter sought to outline twenty overall features with various permutations to be applied as the quantitative or algorithmic portion of the analyses (e.g. Swofford & Champod, 2021), and four overall linguistic theories (along with those proposed in Chapter 2 above such as intersectionality) as the qualitative or theoretical portion of the analyses. As discussed in the introduction to this chapter, this paper seeks not only to demonstrate how any such set of features should be responsibly applied and situated within the context of the data, but also how combined quantitative and qualitative approaches can responsibly do so on real-world data that is not appropriate for a purely statistical approach, but need not rely on a purely non-algorithmic approach.

Chapter 4 contains two non-forensic examples of real-world data: the A Gay Girl in Damascus (AGGID) blog, and the catfi.sh Tinder data. In the AGGID analysis, fifteen of the twenty features are applied excluding, for example, expressive lengthening because of the formal register of the blog. Because there is ample guised data for Amina, but little un-guised data for Tom, Amina's actual author, the AGGID analysis primarily compares Amina's language to the eBAC distributions, in an attempt to accurately predict Amina's performance. The analysis then considers such linguistic theories as intersectionality (in that Amina was not just a female performance by Tom, a straight white male American, but specifically a Syrian lesbian female activist) and others to situate much of the otherwise male-coded language use (as suggested by the eBAC) indicated by the keywords as contextually topical rather than gendered.

In the catfi.sh Tinder data, eighteen of the features are applied, excluding hyperlinks because of the genre, and keywords for illuminating only the specific context of men seeking women for romantic or sexual encounters. Because this analysis does not contain an authorship question—but does contain a performance question in that the data shifts from heterosexual male interlocutors thinking their conversational partner is an available and interested female into them each understanding the same conversational partner has been another heterosexual male the whole time—the baseline comparison is not some other author or even the eBAC, but the collective group of authors before, during, and after their perceptual understanding of their conversational partners has shifted. These trajectories are then compared to the eBAC to determine their performative gender skewing. A further qualitative analysis is applied in order to appropriately situate these trends within the context of the data.

Chapter 5 contains two forensic examples of real-world data: the Hemmert case, and the Potter case. In the first Hemmert analysis, seventeen overall features are applied, although

only nine of these are applied to the text messages of questioned authorship, as the remaining eight did not appear in the limited dataset. Those eight, however, are still analyzed in the known language of the husband and wife, in order to test their applicability based on the eBAC findings in this chapter. The second Hemmert analysis reapplies these features to a more refined subcorpus of the husband's language with an interest toward appropriately situating the features in the proper context of power and audience between his often-hostile communications with his ex-wife, and his more pleading communications with his current but estranged wife.

In Potter, seventeen of the twenty features are applied in a single analysis which similarly attempts to account for the same distribution of power and audience as considered in the second Hemmert analysis. This was done by splitting the questioned data for "Chris" the CIA agent between that sent to his male friend Jamie, and that posted to a wider and more hostile audience online. This analysis excluded, for example, all but a brief mention of emotion terms, as they were not reliably memorialized in the data provided, and post length to avoid conflating genres, as it was unclear from the emails analyzed which were emails and which copied from Facebook posts. Because, unlike in the Hemmert case where both the husband and wife argued that the other had authored the questioned text messages, there was no second, real-world "Chris" to compare his language to, the eBAC was used as a baseline corpus against known Jenelle's female-identifying language.

Finally, while this paper as a whole does not seek to argue that these twenty features or considered linguistic theories are the best and only features to apply to any gender-based authorship, profiling, or other analyses, the latter part of Chapter 5 seeks to compare the approaches undertaken on these four cases, in the hopes of identifying any patterns of success or failure for any given feature, and weighting those features appropriately. Although the focus of this paper is on gendered language features and gendered identity performance (sometimes but not always disguised, sometimes but not always only perceptually so), the larger intents of the analyses herein are to test and establish a method for applying any such feature set along any category or categories of identity performance, and appropriately situate such applications in the individual and often singular context of any given set of real-world data.

Chapter 4 – Non-Forensic Analyses

Part 1 – Introduction

The following chapter contains two analyses of non-forensic data in which gender performance and disguise played a role in the formation or interpretation of the language. Both cases are CMC, with the first being a set of blog posts, and the second a set of in-app Tinder messages. Both datasets are run through the relevant twenty features proposed above in Chapter 3, and various theories and approaches as discussed throughout Chapters 1 through 3 are applied to the output of these feature distributions.

The first analysis covers the A Gay Girl in Damascus blog, in which a straight, white, American male performs over the blog posts as a lesbian Syrian activist. Although the primary focus of this analysis remains the possible gender performance differences between Tom, the author's real identity, and Amina, the author's performed identity, the latter part of this analysis also considers the intersectionality of Tom and Amina's gender preferences, and takes into consideration both Audience Design and Community of Practice of a highly formal register (as compared to, for example, many of the blogs in the Blog Authorship Corpus) written within a community and for an audience of other queer Syrian activists.

The second and third analyses cover the catfi.sh Tinder data, in fifty-four sets of straight male conversees are subject to third-party catfishing by a hacker, with each initially believing their interlocutor to be a straight (or at least male-interested) female. The first analysis applies the relevant twenty features from Chapter 3 as a trajectory over the collective authors' progression from 'knowing' their conversational partner is female, suspecting all or part of their conversational partner's identity is untrue, and then 'knowing' that their conversational partner is male (or at least not who their profile purports them to be).

The second catfi.sh Tinder data analysis takes a more discourse-based approach as suggested in Chapter 3 by Ehrlich & Meyerhoff (2014), with the aim of analyzing why the shift of the obvious gender-based output of the feature performances does not prompt immediate red flags. The catfi.sh Tinder data overall considers Audience Design in the tailoring of each author's language from more female- or neutral-coded language with their female partners to more male-coded language when those partners turn out to be male, Speaker Contamination in the form of the authors' unknowingly fabricated profiles as overwriting many of the otherwise obviously glaring gender-based language cues, and finally Accommodation Theory in heightening both the performative and perceptive aspects of the prior two theories for authors with a vested interest in appeasing a potential romantic or sexual partner, which is the entire onus behind their interaction.

Finally, this chapter considers other instances of identity-disguised CMC in which gender played some part in the performance, both on and off Tinder, while introducing the difference between cooperation and suspicion as commodities in different Communities of Practice.

Part 2 – The AGGID Blog

The blog A Gay Girl in Damascus (AGGID) on blogspot.com ran from February to June of 2011, and chronicled the struggles of a Syrian-American, Amina Abdallah Arraf al Omari, living in Syria as an openly gay woman, with the blog's tagline at the time touting it as "An out Syrian lesbian's thoughts on life, the universe and so on ...". On June 6, 2011, posts on the blog by Amina's cousin Rania claimed that she had been abducted, and after the public backlash brought Amina's abduction into the media spotlight, it came out on June 12, 2011 that the actual author of the posts was not a 35-year-old Syrian-American lesbian, but a 40-year-old, white, straight, American man living in Edinburgh, Scotland, by the name of Tom MacMaster.

A. The Data

The blog contains a total of 146 posts, two of which are reportedly by Amina's cousin Rania, two of which are signed by the actual author of the blog, Tom, and the remainder of which were reportedly authored by Amina.

A number of these posts were excluded from this analysis based on a variety of criteria. The posts were gathered via the Wayback Machine, with any posts not archived by this site unable to be retrieved. Non-communicative posts such as poems, and posts that were primarily multimedia such as videos, links, or images were not included, as many of these were not retained by the Wayback Machine either. Other posts that Amina attributed to being reposted from other sources, or posts in which it was unclear whether she was taking the language from elsewhere, as any language that occurred within block quotes or citations, or any posts that were not authored in English, were additionally excluded.

Finally, chapters from Amina's purported book were excluded, as they, like the poetry, were not CMC in that they were not intended to be communicative or receive conversational responses one would generally expect from other mediums such as emails, text messages, and so on. As such, other blog posts that resembled this writing style—telling a story from the distant past as opposed to reporting on current events as most of Amina's posts do, containing numerous conversational quotations in the way a book might, and generally following along with the style of a novel rather than a blog post meant to be interacted with in some kind of timely manner—were excluded as well.

The exclusion of these posts and excerpts above left a remaining 112 posts attributed to Amina, and the final two posts Tom attributed to himself.

	Amina	Tom
Total Words	86,938	1,512
Total Posts	112	2
Words per Post	776.2	756
Total Types	8,082	540
Type/Token Ratio	0.09	0.36

Table 4.1 Dataset distribution of Amina and Tom in AGGID

The gender-guising aspect of this dataset differs from the Tinder data discussed below in a number of key ways. First, the gender-guising performed as part of this eventually self-

professed hoax was done knowingly and intentionally, while none of the Tinder authors had any knowledge that their profiles were presenting them as anyone other than themselves, and thus there was no intent to deceive on either side of the conversation. Second, although the blog posts included based on the criteria discussed above are not the same type of CMC as Tinder conversations, they are still communicative in that the primary goal of a blog is often to receive timely responses to the information given. Finally, the gender-guising performed by Tom as Amina has the added component of sexuality, which is itself intentionally guised. In the end, both datasets (AGGID and Tinder) are straight males frequently being understood as females interested in other females, whether intentional or not.

The data for the collected posts for the A Gay Girl in Damascus blog is provided in Appendix C.

B. The Analysis

As CMC, the AGGID blogs also differ from the Tinder data in that they contain no emojis or emoticons, or abbreviated terms such as *lol* or *omg* and contain very little other expressive textual language or orthography. As such, many of the eBAC features do not strictly apply to this dataset, and so are not included in the following analysis (**3. Emoticons, 6. Abbreviations, 8. Expressive lengthening, 12. Alternative spellings, and 19. Hyperlinks** are excluded below).

It is also worth noting that there is little data for Tom writing as himself, and so while all of the counts considered in the tables below are normed to 1,000 words to accommodate his relatively small word count of 1,512, many of his counts are likely to be far less probative than Amina or the eBAC's. This is largely because the focus of his two posts is identical (to explain himself and apologize, both of which call for more formal language than many of the eBAC blogs, for example). Thus, Tom's counts are considered below only as an anecdotal comparator where appropriate.

The following analysis considers 15 total features of the 20 discussed in Chapter 3 above from the analysis of the eBAC, with the remaining 5 features, as mentioned above, not applicable to the AGGID dataset. For Amina (but not necessarily also Tom):

- a. 6 of these features trended male coded (**1. Pronouns, 2. Emotion terms, 5. Friendship terms, 11. Assent terms, 13. Swears and anti-swears, and 17. Articles and determiners**)
- b. 4 of these features trended female coded (**4. Kinship terms, 14. Prepositions, 15. Negation terms, and 16. Conjunctions**)
- c. 3 of these features were ambiguously coded (**7. Punctuation, 9. Backchannel sounds, 10. Hesitation words, and 18. Post length**)

The final feature, **20. Keywords**, is analyzed separately in section 4. (Although **8. Expressive lengthening** is not considered as an individual feature, as the AGGID does not expressively lengthen any words, it is considered for individual features such as **9. Backchannel sounds** and **10. Hesitation words** within the analyses of those individual features.)

1. Male-coded features (for Amina)

Table 4.2 demonstrates the distribution of the 6 out of 14 analyzed features in which Amina was found to exhibit distributions more indicative of the Male eBAC than the Female eBAC. In two of these instances (**11. Assent terms**, and **13. Swears and anti-swears**), no hits of the feature occur in Tom, and thus they cannot be compared. In two of these instances (**2. Emotion terms**, and **5. Friendship terms**), Tom better matches the Female eBAC. And in the final two instances (**1. Pronouns**, and **17. Articles and determiners**), though Tom closer matches the Male eBAC along with Amina, the difference between Amina and Tom trends in the same direction as the difference between the Male and Female eBAC subcorpora, respectively.

#	Feature(s)	Amina	Tom	Male	Female
1	First Person Pronouns Total	55.3	89.9	60.3	78.9
	Second Person Pronouns Total	5.8	0.0	12.1	13.1
	Third Person Pronouns Total	36.8	38.4	37.4	41.6
	Total Pronouns	97.9	128.3	111.4	133.4
2	Total Emotion Terms	24.1	38.4	28.9	33.9
5	Female Coded	0.00	0.0	0.1	0.1
	Male Coded	0.02	0.0	0.2	0.2
	Other coded	0.75	1.98	1.1	1.3
	Total Friendship Terms	0.77	1.98	1.3	1.6
11	yes variations	0.4	0.0	0.4	0.5
	yeah variations	0.3	0.0	0.5	0.6
	ok variations	0.2	0.0	0.4	0.6
	yeh	0.02	0.0	0.01	0.02
	yup variations	0.1	0.0	0.03	0.04
	Total Assent Terms	1.0	0.0	1.7	2.3
13	Female coded (anti-swears)	0.0	0.0	0.1	0.1
	Male coded (swears)	0.4	0.0	2.0	2.0
17	Total Articles and Determiners	143.6	124.3	141.5	134.2

Table 4.2 Distribution of features in which Amina better matches the Male eBAC

As shown in Table 4.2 above, Amina better matches the Male eBAC for **1. Pronouns**, not only in the overall normed frequency of pronouns, but also in their distribution between first, second, and third person. Interestingly, Tom better matches the Female eBAC in the same way, with more pronouns overall than Amina.

The overall normed frequency of **2. Emotion Terms** in Amina better matches the Male eBAC, while Tom better matches the female eBAC.

Although 0.77 **5. Friendship terms** in Amina is lower than both the Male and Female eBAC counts, the eBAC finds that, in general, Males use fewer friendship terms than females. Thus, this normed count for Amina is still what the eBAC might expect for males, while Tom's count of 1.98 is closer to what might be expected for females.

The same reasoning also applies to **11. Assent terms**, as the eBAC finds that in general assent terms are more frequently used by females than males, and Amina contains a lower normed frequency of them. Tom contains no assent terms, but this is likely due to the context of the two posts, as discussed above, rather than a gender-based stylistic indicator.

Amina contains 0 **13. Anti-swears** (female coded), and a higher distribution of **13. Swears** (male coded). As compared to the other 6 features here, swears are more of a stylistic

choice than a functional one (as they are content and not function words), the implications of which are discussed further below. As with **11. Assent terms**, that Tom contains no swears or anti-swears likely has more to do with the posts' context as apologies than anything else.

Finally, the eBAC finds that males tend to use more **17. Articles and determiners** than do females, and Amina contains a closer frequency to the male subcorpus, while Tom contains a closer frequency to the female.

That these 6 features, which the eBAC found to be male coded, or are otherwise found to occur at similar relative normed frequencies between Amina and the Male eBAC subcorpus, indicates that these features may indeed exist above the level of conscious manipulation in instances of gender guising, or otherwise may not occur to an author as a gender-indicative feature.

2. Female-coded features (for Amina)

Table 4.3 demonstrates the distribution of the 4 out of 14 analyzed features in which Amina was found to exhibit distributions more indicative of the Female eBAC than the Male eBAC. In three of these instances (**4. Kinship terms**, **15. Alternative prepositions**, and **16. Conjunctions**), however, Tom also exhibits frequencies closer to the Female eBAC, and in one (**15. Negation terms**) no instances of the feature are present in Tom for comparison.

#	Feature(s)	Amina	Tom	Male	Female
4	Total Kinship Terms	3.5	3.3	2.6	3.9
12	Female coded	0.3	0.0	0.2	0.2
	Male coded	0.05	0.0	0.2	0.2
	Other coded	3.4	0.0	2.2	2.4
	Total Negation Terms	3.7	0.0	2.6	2.7
14	Alternative Prepositions	0.0	0.0	1.8	1.8
	Standard Prepositions	118.1	121.0	127.2	123.3
	Total Prepositions	118.1	121.0	129.0	125.1
16	and	35.9	44.3	25.7	28.1
	&	0.0	0.0	0.03	0.4
	n	0.03	0.0	0.3	0.5
	Total Conjunctions	36.2	44.3	26.0	29.0

Table 4.3 Distribution of features in which Amina better matches the Female eBAC

In the case of **4. Kinship terms**, both Amina and Tom have normed frequencies of kinship terms between the Male and Female eBAC subcorpora. Though Tom is closer than Amina to the Male frequency (Tom is 0.7 off while Amina is 0.9), Tom and Amina are closer to one another (0.2) than either is to the Female subcorpus (0.6 for Tom and 0.4 for Amina). However, for this feature, both Tom and Amina exist within the frequency range between the Male and Female subcorpus, thus making this particular feature a poor indicator of gender coding for this dataset. It is also worth noting that for Amina, who identifies as a lesbian, the male-coded *girlfriend* and *wife* and variations thereof would not necessarily follow the findings of Bamman et al. (2014), who did not take sexuality into consideration with such terms.

The distribution of **12. Negation terms** in Amina match relatively with both the Male and Female eBAC, with other-coded *no* being far more frequently used than either the female- (*noo+*, *cannot*) or male-coded (*nah*, *ain't*, *nobody*) negation terms. The overall normed frequency of negation terms in Amina is, however, on the higher end, as the eBAC would

expect for Female authors. It is worth noting, however, that not only are there no examples in Tom for comparison, but the difference between the Male and Female eBAC subcorpora (0.1) is significantly lower than the difference between Amina and the Female eBAC (1.0).

The frequency of **15. Prepositions** in both Amina and Tom are better matches for the Female eBAC, which predicts that Prepositions in general tend to be more frequently used by males than by females. This is the only one of the four features in which the trend between Tom and Amina (2.9) is relatively the same and in the same direction as the trend between the Male and Female (3.9) eBAC subcorpora, respectively. It is worth noting, however, that neither Amina nor Tom exhibits the non-standard alternative prepositions that were suggested as female coded (2, 4, w/, w/o).

Finally, both Amina and Tom exhibit a higher distribution of **16. Conjunctions**, better matching the Female eBAC subcorpus. Additionally, this feature does trend lower from Tom to Amina (and thus more male coded), though this may be due to the differing lexical densities of the two sets.

That these features are more female coded for *both* Amina and Tom (aside from **12. Negation terms**, which do not occur in Tom for comparison) suggests that, for these features in particular, the shared author Tom may simply have more female- than male-coded trends for **4. Kinship terms**, **15. Prepositions**, and **16. Conjunctions**. This may also suggest that the single content-word feature (as opposed to the remaining three function-word features), **4. Kinship terms**, may be within the realm of conscious manipulation for an author guising their gender. More than likely, however, this feature is influenced by the content of the AGGID blog, which discusses the interpersonal activities of Amina with other people and groups in Syria going through the ‘same’ struggle as the purported author.

3. Ambiguously-coded features (for Amina)

Table 4.4 below demonstrates the distribution of the 4 out of 14 analyzed features in which Amina did not closely or clearly resemble one eBAC subcorpus over another. Two of these features (**7. Punctuation**, and **18. Post length**) are similarly ambiguous in Tom, while the final two features (**9. Backchannel sounds**, and **10. Hesitation words**) do not occur in Tom for comparison.

Both Amina and Tom follow the pattern of the eBAC of using periods the most frequently, and ellipses the second most frequently of all their **7. Punctuation** uses. The major difference between the two is that Tom does not use either exclamations or questions—likely due to the context of the two posts—and Amina more closely matches the Female eBAC distributions of periods and ellipses. When it comes to questions and exclamations, however, Amina does not follow the same preference of use as the eBAC (which prefers exclamations over questions in both subcorpora). Additionally, both Amina and Tom exhibit closer words per punctuation to the Male eBAC subcorpus, but much higher punctuation per post (though this is likely due to the much higher average word count per post, and much lower use of expressively lengthened punctuation than either eBAC subcorpus).

#	Feature	Amina	Tom	Male	Female
7	Total Periods	67.7%	90.5%	73.1%	67.4%
	Total Ellipses	18.7%	9.5%	14.6%	16.9%
	Total Exclamations	4.7%	0.0%	6.4%	9.6%
	Total Questions	8.9%	0.0%	6.0%	6.1%
	Words per Punctuation	16.4	14.4	13.0	11.9
	Punctuation per Post	47.4	52.5	15.9	17.0
	Standard punctuation	98.9%	96.2%	92.6%	90.2%
9	Lengthened punctuation	1.1%	3.8%	7.4%	9.8%
	o, oh, ooh variations	0.1	0.0	4.3	1.4
	other variations	0.2	0.0	0.9	1.3
	Total Backchannel Sounds	0.3	0.0	5.2	2.7
	Unlengthened Backchannel Sounds	96.6%	--	72.5%	84.0%
10	Unlengthened Backchannel Sounds	3.4%	--	27.5%	16.0%
	um variations	0.03	0.0	0.08	0.14
	uh variations	0.08	0.0	0.06	0.08
	hmm variations	0.01	0.0	0.20	0.23
	Total Hesitation Words	0.13	0.0	0.43	0.58
	Unlengthened Hesitation Words	91%	0.0	75.3%	72.6%
18	Lengthened Hesitation Words	9%	0.0	24.7%	27.4%
	Total words	86,938	1,512	59,233,203	58,108,641
	Total posts	112	2	287,381	287,335
	Total blogs	1	1	9,318	8,147
	Words per post	776.2	756	206	202
	Words per blog	86,938	1,512	6,356	7,132
	Posts per blog	112	2	31	35

Table 4.4 Distribution of features in which Amina doesn't clearly match either subcorpus

For both **9. Backchannel sounds** and **10. Hesitation words** (which are themselves included in backchannel sounds), Amina exhibits far lower counts than the eBAC, and Tom exhibits none. This is likely due to the more formal nature of Amina's posts than many of the eBAC blogs (and the same is, of course, true for Tom), along with the absence of other features such as emoticons, abbreviations, and alternative spellings. In addition, though expressive lengthening does occur for these features, the 3.4% and 9% totals are both comprised of a single *hmmm*. Thus, these features are not helpful indicators of gender or gender guising for this dataset.

Finally, because the AGGID overall has much higher **18. Post length** than the eBAC, this feature remains a poor comparison as well.

C. Keywords

The following section considers keywords between Amina, the total eBAC corpus, and the Male and Female subcorpora. (Tom is not included in this analysis, as the context of the two Tom posts as an apology means that any keywords are likely more specific to that context than to any potential indicators of gender.)

	Total eBAC	Male eBAC	Female eBAC
Total Keywords	2,022	1,898	2,287

Table 4.5 Total words with a keyness above 6.53

Table 4.5 above demonstrates the total number of keywords (with a keyness of 6.53 or higher) between Amina and any of the 3 datasets. As shown, Amina has hundreds more words key to the Female than the Male subcorpus—indicating, perhaps, that Amina has more in common with the male half of the eBAC.

Total eBAC			Male eBAC		Female eBAC	
#	word	keyness	word	keyness	word	keyness
1	syria	2,377.721	syria	2,132.508	syria	2,591.367
2	regime	1,963.2	regime	1,735.090	regime	2,303.081
3	damascus	1,553.0	damascus	1,396.408	damascus	1,573.926
4	syrian	1,279.1	syrian	1,138.042	syrian	1,386.367
5	assad	1,216.3	assad	1,094.876	assad	1,197.508
6	arab	853.1	arab	746.488	arab	1,024.744
7	homs	738.7	homs	668.895	syrians	670.500
8	syrians	649.7	syrians	580.619	homs	667.334
9	alawi	575.0	we	538.509	muslim	649.103
10	muslim	527.0	alawi	523.494	al	580.485

Table 4.6 Top 10 keywords between Amina and the eBAC datasets

Table 4.6 above shows the top 10 keywords for Amina and each of the three datasets. As shown, for all three datasets but the Male eBAC subcorpus which has 'we', all of the top 10 keywords fall into the category of political language specific to the region around Syria (including individual terms, place names, and proper names).

	Total eBAC	Male eBAC	Female eBAC
Political Keywords	86	84	85
Non-Political Keywords	14	16	15

Table 4.7 Distribution of political and non-political keywords between Amina & the eBAC

As shown in Table 4.7 above, of the top 100 keywords between Amina and each eBAC dataset, around 85 are within the political category (including individual terms, place names, proper names, and religious terms). The remaining 14 and 16 keywords per dataset are demonstrated in Table 4.8 below.

Total eBAC			Male eBAC			Female eBAC		
#	word	keyness	#	word	keyness	#	word	keyness
13	we	458.5	9	we	538.5	19	they	421.4
16	katy	416.7	14	katy	415.0	21	katy	412.1
18	they	405.5	15	they	389.5	27	we	386.4
36	and	284.4	18	and	353.3	31	of	361.1
41	amina	255.2	29	amina	286.0	34	as	315.6
43	as	239.6	39	were	237.3	36	the	310.6
45	were	219.0	45	women	193.4	48	amina	227.2
46	of	218.4	48	as	178.8	49	and	223.2
49	rania	190.7	49	rania	171.6	53	were	201.2
53	us	177.6	50	us	169.4	57	us	186.0
62	youtube	159.3	59	youtube	144.3	59	rania	183.2
75	the	132.0	66	want	129.3	62	women	181.3
80	are	127.0	71	d	125.8	74	are	153.0
90	d	109.1	75	of	118.4	80	youtube	143.9
			89	are	104.4	99	in	118.1
			94	father	96.8			

Table 4.8 Non-political keywords from the top 100 for Amina and the eBAC datasets

Although politics as a feature was found by Schler et al. (2006) to be male-prevalent, that the AGGID blog contains ample words of these categories, especially as compared to the eBAC dataset of largely American and Western European varieties of English, these keywords are more telling of the content of the AGGID blog and not any gender-coding. Any keywords that appear in only two of the three top 100 lists are indicated in **bold**; these include **the** in only the total and Female, **d** (would) in only the total and Male, and **women** in only the Male and Female. Any keywords that appear in only one of the top 100 lists are indicated in underlined

bold; these include **want** and **father** in the Male, and **in** in the Female.

Because the top keywords shown below are overwhelmingly indicative of the context of the AGGID blog, the following instead considers the negative keywords, the distribution of which is shown in Table 4.9 below.

	Total eBAC	Male eBAC	Female eBAC
Total Negative Keywords	1,879	1,965	1,761
Negative Keywords above 6.53	384	331	356

Table 4.9 Total negative keywords between Amina and the eBAC datasets

Table 4.10 below shows the top 10 negative keywords between Amina and the eBAC datasets. (Note that the top word *post* comes from the <post> and </post> tags in the BAC, and so is not probative here.) Again, words that are shared by only two top 10 lists are indicated in **bold**, while words that are unique to a single top 10 list are indicated in **underlined bold**.

#	Total eBAC		Male eBAC		Female eBAC	
	word	keyness	word	keyness	word	keyness
1	post	1549.5	1549.2	post	1548.6	post
2	i	472.0	191.7	i	842.7	i
3	my	212.5	172.6	you	340.9	my
4	you	198.9	163.1	it	226.8	it
5	it	193.9	119.2	this	226.6	you
6	m	143.5	112.1	s	205.7	m
7	s	120.5	107.5	my	202.8	me
8	got	117.1	103.1	got	148.6	t
9	get	113.7	94.8	get	142.6	love
10	me	108.9	92.4	to	135.1	just

Table 4.10 Top 10 negative keywords between Amina and the eBAC datasets

The unique words for the Male eBAC subcorpus are the determiner *this* and the preposition *to*, while the unique words for the Female eBAC subcorpus are the pronoun *me*, the emotion term *love*, and the multi-use function word *just*.

Overall, because of the strongly political nature of the AGGID blogs, keywords are likely not as probative in this dataset for finding key categories Schler et al. (2006) found to be gender coded, though, as we will see in further analyses, keywords can instead prove a useful tool for setting aside features that are less probative based on the context of any given dataset. In this case, it would appear that Audience Design within the Community of Practice of other queer Syrian activists (a serious topic to the point of being deadly) suggested a more formal register for the AGGID blog as opposed to the wider array of registers within, for example, the eBAC.

D. Overview

Table 4.11 demonstrates the findings above: whether either AGGID dataset feature was found to be more indicative of **male-coded** or **female-coded** language. In addition, the color of the left columns indicates whether these features would **correctly indicate a male author**, or incorrectly indicate a female author, either because **the features were successfully guised**, or the features in the non-guised dataset would have indicated a female author to begin with. (Any features that were found to be unclear are also indicated.)

#	Feature	Amina	Tom
1	Pronouns	male	female
2	Emotion terms	male	female
4	Kinship terms	female	male
5	Friendship terms	male	female
7	Punctuation	unclear	unclear
9	Backchannel sounds	unclear	unclear
10	Hesitation words	unclear	unclear
11	Assent terms	male	--
12	Negation terms	female	--
13	Swears and anti-swears	male	--
14	Prepositions	female	female
16	Conjunctions	female	female
17	Articles and determiners	male	male
18	Post length	unclear	unclear
20	Keywords	unclear	unclear

Table 4.11 Feature findings overview

Of the 10 indicative features above, 6 would indicate a male author, and 4 would indicate a female author. Of the latter 4, 2 do not appear to have been successful guising, because Tom exhibits the female variation of these features in his non-guised writing. This may be down to these two features being less probative in differentiating gendered language, simple author preference that cannot be accounted for with such features, or the formality or some other feature of Tom's two posts that skewed these features away from his usual writing style.

Worth noting is the fact that, of the 7 features which had enough information in Tom to show findings in Table 4.11 above, 5 of them indicated his non-guised language was more female than male coded. As such, though these features might more accurately predict that his guised language in Amina is, in fact, male, the same features would have inaccurately predicted his non-guised language was more likely authored by a female. However, these findings may simply be down to the small 1,512-word dataset for Tom.

Thus, the above analysis may indicate any number of things:

1. That of the 15 features considered, pronouns, emotion terms, friendship terms, swears and anti-swears, and articles and conjunctions are the most probative for indicating gender-guised language,
2. That the dataset analyzed here is too much of an outlier to be accurately accounted for by the eBAC features, either because the context (of politics) is too specific, the gender-guising is paired with sexuality-guising, or some other underlying issue,
3. That the features found to be helpful by Bamman et al. (2014) and Schler et al. (2006) cannot be usefully applied to all datasets, either because they are too short, too dissimilar (the AGGID blog is much more formal than the average eBAC blog), or some other underlying issue, OR
4. That there are other dataset features to consider when applying features potentially indicative of gendered language, such as audience, that may skew the distribution of gender-coded features in (un)predictable ways.

This latter possibility is considered further in the following section's analysis of Tinder data.

Part 3 – Tinder

This section examines the issues of language and gender perception and performance through a series of text-based conversations associated with the dating app Tinder. Such apps are generally set up so that, for example, heterosexual males will spend most of the time conversing with heterosexual females. For such conversations we might start with two conflicting intuitions. One prediction from Bamman et al.'s (2014) study might be that heterosexual men on Tinder (who will spend their time talking with those identifying as heterosexual women) will thus adopt a more female language style. There may also be the counterintuition that because of the speech activity of chatting on a dating app the heterosexual males may adopt a hypermasculine stance and perform their maleness more overtly (e.g. Johnstone, 2009).



Figure 4.1 Catfi.sh demonstration of third-party catfishing

The data in the current study, however, is a perversion of Tinder norms; as is described below, the data is taken from a “catfish” man-in-the-middle attack. In this situation, two heterosexual males are matched to talk with one another, both seeing the profile of a heterosexual woman and thus believing that they are interacting with a woman through the app. This situation creates a natural experiment, wherein men start interacting while believing that they are talking with a woman, and then part way through the conversation realize that this is not the case.

It may be the case that on dating apps such as Tinder, males adopt more female-coded or gender-neutral features of language for this particular speech event and the community of practice on Tinder, which is likely to be made up of 100% female interlocutors. It is also very likely, as we will see below, that interlocutors on Tinder tend to attribute perhaps overmuch to the meta-information provided by Tinder profiles, and the way the application itself is supposed to work, taking it to be the ground truth for the most part, and often overlooking otherwise notable linguistic cues that may indicate inaccuracies in their assumptions.

A. The Data

The following datasets are sourced from the website <http://catfi.sh>, which describes itself as follows:

Catfish is a fun demonstration of an active *man-in-the-middle attack* against

users of the Tinder dating app. Using fake female Tinder profiles, Catfish establishes a Tinder “match” with two victim men from the same city. Messages are then relayed between these two men, creating the illusion that they are communicating over a private channel. The two victims are chatting (and likely flirting) with one another, whilst each believes he is talking to a girl.

Such conversations enable a look into not only how the conversees negotiate their own identity (most especially when the truth of their interaction is revealed), but also how they negotiate their perceptions of their interlocutor’s identity. This often sees users ignoring what the outside observer would likely categorize as red flags, overridden by the information provided by the profile displayed and the application itself (which did not make same-sex matches at the time).

The data was compiled, first, into a corpus, separated not only by conversations and individual conversees, of which there are 54 and 108 respectively, but also into three possible phases of conversation. The first phase of the conversation, which we term here the Before phase, occurs while an author takes the profile they are presented with to be the relative truth of their conversational partner’s identity, as exemplified in Figure 4.2 below. This phase makes up just over 80% of the total 35,000-word corpus, as conversees either do not reach one or both of the latter two phases or reach them outside of the available Tinder conversations, and because, when the latter phases do occur, they tend to be both awkward and brief.

Good morning cutie	
	Good morning beautiful! Hope your morning is going well.
Thank you baby How are you	
	Good stuck at work :(how about you? Everything’s good?
Just living the dream lol	
Work? We could be going to breakfast lol	
	Omg don’t tease me

Figure 4.2 Example excerpt of the Before phase

The second phase, termed here the During phase, is the initial point of an author’s half of the conversation at which they begin to question the consistency of the details they are receiving throughout the conversation with the details provided by the profile itself, as exemplified in Figure 4.3 below. Though authors do sometimes alternate between skeptical (as in the During phase) and appeased (as in the Before phase), once the During phase begins, conversational data is categorized into this subcorpus.

	I mean, you stumped me with your comment saying you do the pole planting. To me that sounded like you have male anatomy! Lol
Well I do?! 😏 😏 😏	
You’re weirding me out now 😏 😏 😏 😏 😏 😏 😏	
You’re a dude?	

Uh yes.

Figure 4.3 Example excerpt transitioning from During to After phases

The third and final phase, termed here the After phase, occurs when an author finally realizes that their conversational partner is, in fact, male, or at least is not at all who they were led to believe they were speaking with by the profile provided (as in some instances it is unclear how much a participant understands before a conversation is terminated by either party), as exemplified in Figure 4.4 below. The context of the After phase can vary in multiple ways, such as an author attempting to find out more information about the deception, whether it was intentional or incidental, or simply ending the conversation before their partner has resolved to the After phase; either taking the situation in stride, or demonstrating aggression toward their conversational partner (often determined by who an author blames for the deception); either believing their conversational partner to be the catfish himself, or reasoning out the third-party catfishing or a glitch by Tinder; and so on. Both latter phases make up just under 10% of the total corpus, each.

I'm a 21-year-old guy called Chris btw if that's not what you're seeing?

I'm so confused

Hahaha oh god that is brilliant. Well what I'm seeing is a 22 year old Asian lady called Sybil

This has made my day.

No. Fucking. Way. This is hilarious dude!

Figure 4.4 Example excerpt of the After phase

Not every conversation or individual participant reaches every phase of the conversation, frequently failing to reach the latter two phases, or skipping the During phase altogether. This is demonstrated by Table 4.12 below by the drop-off of authors from 107 of the 108 in the Before phase (only one author is immediately skeptical of the entire situation based on the user's profile alone), to the 63 and 66 authors who reach the During and After phases, respectively (and who do not always overlap). I note here that because of the disparate sizes of the subcorpora, all normed counts are normed to 3,000 to avoid inflation of the During and After counts.

	Word Count	Corpus %	Authors	Author %
Total	32,875	--	108	--
Before	27,272	83.0%	107	99.1%
During	2,717	8.3%	63	58.3%
After	2,886	8.8%	66	61.1%

Table 4.12 Distribution of word counts in the Tinder data phases

In reading the excerpts from the Tinder conversations below, note that any string of XXX (not to be confused with the tag marker x(x)(x)(etc.)) indicates identifying information redacted by Catfi.sh before making these exchanges publicly available. Additionally, while many excerpts show conversation-medial exchanges, any excerpt that concludes the exchange will culminate in a black line beneath the final turn (as this distinction is sometimes relevant to the analysis). These conventions are maintained elsewhere where the Tinder data

is excerpted.

The data for the Catfi.sh Tinder conversations is provided in Appendix D.

B. The Analysis

Below, this thesis discusses an analysis of 10 of the 20 features outlined by the eBAC in the analysis in Chapter 3, as both female- and male-coded features that follow trends through the conversational phases the eBAC analysis would expect (falling and rising, respectively, and female-coded features following a similar but notably different trend). The remaining 8 features showed more mixed results, either in that they provided inconclusive findings, or contradictory feature-internal findings. Not discussed here are Schler et al.'s (2006) two additional features, which are not relevant to the dataset (**19. Hyperlinks**, and **20. Keywords**).

C. Downward Trending 'Female Coded' Features

Of the features found in the eBAC to be female coded, six are discussed here with overall downward trending trajectories. The first two features, **1. Pronouns** and **2. Emotion terms** trend downward, with an upward spike in the During phase (which is, as explained below, likely related to the context of the During phase). The next three occur at a consistently downward trajectory; these are **4. Kinship terms**, **8. Expressive lengthening**, and **9. Backchannel sounds**. The final feature, **10. Hesitation words**, has the same consistently downward trend, with an additional male variation inversely trending upward.

1. Pronouns

	Before			During			After		
	#	norm	%	#	norm	%	#	norm	%
you, your, youre, you're	1,331	138.7	82.1%	225	229.2	83.6%	148	140.7	85.5%
u, ur, yr, yur, ure	290	30.2	17.9%	44	44.8	16.4%	25	23.8	14.5%
overall use	1,621	168.9	--	269	274.0	--	173	164.5	--

Table 4.13 Overall raw and normed uses of 2nd person pronouns and variants

As shown in Table 4.13 above, the overall use of second person pronouns rises in the During phase, while it is roughly the same in the Before and After phases, indicating that this rise is likely due to the context of the During phase.

As demonstrated in Table 4.14 below, the distribution of *you* variants and alternative spellings stays roughly the same throughout the phases, with only a 2.3% difference in total. This trend exhibits a rise in the During phase of the more female-coded alternative spellings, but a drop in the After phase.

	Before			During			After		
	#	3k	%	#	3k	%	#	3k	%
first person	2,146	223.6	50%	185	188.5	37.2%	185	175.9	43.7%
second person	1,621	168.9	37.8%	269	274.0	54.1%	173	164.5	40.9%
third person	524	54.6	12.2%	43	43.8	8.7%	65	61.8	15.4%
TOTAL	4,291	447.0	--	497	506.3	--	423	402.2	--

Table 4.14 Overall distribution of pronoun use throughout the phases

As shown in Table 4.14 above, overall pronoun use would appear to follow what the eBAC analysis would expect, as they are used most in the Before and least in the After phases.

However, pronoun use in the data in particular is likely contextually motivated, as with many of the other potentially gendered language features, as both first and second person singular pronouns in particular are key to the exploration of identity, both of an author and their addressee, in the During phase in particular.

2. Emotion terms

The emotion terms found to collocate with either *I* or *you* in the data are exhibited in Figure 4.5 below (“I like/love...”, “I am/feel happy/sad...”, etc.). Few of these words are unique to (and introduced in) the During and After phases. The terms unique to the latter two phases, *stumped* in the During phase and *skeptical*, *baffled*, and *annoyed* in the After phase, are indicative of the context of the conversation, expressing confusion (and indeed *confused* is one of 4 terms shared by all three phases) and annoyance. These uses, as well as the differing uses of shared words by these three phases, are exemplified below.

Before	During	After
wish, hope, bored, excited, liked, angry, afraid	want, sorry, need, confused	wish, hope, bored, excited, liked, angry, afraid
love, good, glad, happy, unfortunately, enjoy, hate, curious, confident, struggle, miss, faith, trusted, scared, regret, intrigued, surprised, surprise, mesmerized, lost, keen, judge, hoping, excite, devastated, awful, admire	stumped	skeptical, baffled, annoyed

Figure 4.5 Distributions of emotion terms by phase

The contextual uses of *need* are, for example, good demonstrations of how even the shared emotion terms differ in their uses. “I need a dance partner” in the Before phase is an author expressing romantic interest in their conversee; “I need to call Nev for u” in the During phase is a reference to the MTV show *Catfish*, in which the host, Nev, investigates whether an online correspondent is indeed who they say they are; “I need to delete this account” in the After phase is clearly indicative of understanding the reality of the interaction that just occurred. Other shared emotion terms, such as *excited*—“you are getting me excited” in the During phase, and “I was excited” in the After—are used in differing tenses that indicate much the same trajectory, with further examples provided in Figure 4.6 below.

Before	During	After
I can be pretty naughty when I want to be		Obv u don't want to be coming up as a girl on guys apps an I don't want same either
Sorry I get naughty at night xxxxxxxxxxxx	U might want to update ur profile photo	Sorry to disappoint bro
No I need a dance partner, it's ok I'll lead!	Im sorry but if you do actually have a cock I aint about that life	But I need to delete this account before I turn gay!
I'm so confused		I'm so confused

I wish you could babe ;)	I need to call Nev for u	Thanks. I wish more girls thought that
I hope you like anal	I'm so confused	Just unmatched and continue to hope this never happens again hah
Single, bored , looking to see who else is on and maybe go out on the weekends. You?	I mean, you stumped me with your comment saying you do the pole planting. To me that sounded like you have male anatomy! Lol	Or that you were bored and made a profile with girl pics...
Would be nice, you are getting me excited just thinking about it		I was excited about sybil the chinese security guard from XXX

Figure 4.6 Examples of emotion terms as used across phases

Notably, Table 4.15 below demonstrates two different methods for determining what qualifies as an emotion term, and their subsequent counts. The first, discussed alongside Figure 4.6 above, determines emotion terms based on their position following “I (am/feel)...”, and considers all relevant instances and lemmas thereof, a method easier to employ on smaller datasets such as this, but somewhat more prone to subjective interpretation. The second relies on the method discussed in Chapter 3, which uses the emotion terms as defined by Clore and Ortony (1988), and the British National Corpus’s lemma list, and is more practical for larger datasets such as the eBAC, but somewhat more subject to erroneous inclusions.

After “I...”	Before		During		After		Total	
	#	3k	#	3k	#	3k	#	3k
Total Instances	328	36.1	28	30.9	32	32.3	388	35.4
After “I...”	Before		During		After		Total	
	All	Unique	All	Unique	All	Unique	All	Unique
Total Variety	38	27	5	1	15	3	42	--
BNC Lemmas	Before		During		After		Total	
	#	3k	#	3k	#	3k	#	3k
Total Instances	1,154	120.2	89	90.7	108	102.7	1,351	116.1

Table 4.15 Total instances and variety of emotion terms used in the data

Although the total numbers differ significantly, not only do their trajectories remain the same, but their distribution relative to one another remains similar (ranging from 3.2 to 3.5 times higher). (As both methods appear to give relative distributions, the eBAC method is used elsewhere in this thesis for consistency.)

As the eBAC might predict, the overall frequency of use, as well as the variety of emotion terms used, are the greatest in the Before phase. The variety is notable in the During phase, which similarly has the lowest instances of emotion terms used; this is possibly indicative of the During phase having largely the same task in every conversation: determining whether the conversee is who their profile purports them to be. Similarly, the Before phase has the largest variety of potential tasks as it has the largest range of possible conversational topics given the context. In the After phase, while it is similar to the During phase in that there is mostly one task—in this case dealing with the outcome of the During phase—authors handle this in a larger variety of ways: politely, humorously, angrily, dismissively; and by either believing the other person is equally innocent or that they perpetrated the ruse themselves.

4. Kinship terms

Kinship terms occur in the data, as shown in Table 4.16 below, exclusively in the Before phase. They are, however, not exceptionally common in the data overall, in total accounting for 0.06% of the total words used, and some of them refer to the author themselves (such as *dad*, *daddy*, and *boyfriend*).

	Before		During		After		Total	
	#	%	#	%	#	%	#	%
Total	21	2.3	0	0.0	0	0.0	21	2.0

Table 4.16 Overall kinship term counts

This trajectory, again, follows along with what the eBAC would predict. Notably, where non-standard or non-female-coded variations are used, they are more frequent—such as *mum* over both *mom* and *mother*, and *bro* over *brother*.

8. Expressive lengthening

Before	During	After
ooo, lol, hmm, erm		
what, nah, err		what, nah, err
yeah, no, oh, ahh, um, yes, xxx, wtf, you, aww, mm		

Figure 4.7 Distribution of expressive lengthening variety by phase

The Tinder data exhibits a variety of expressive lengthening, with the most given to interjections (*ohh*, *yes*, *no*, *umm*, *lol*, etc.) and *what* (which often, as in the case of *wtf*, function as interjections more than wh-pronouns/questions). This lengthening notably includes another marker of accommodation specifically related to affection, *xxx*, where *x* has the literal *kiss*, as in the *x* from *xoxo*, but where the marker itself serves no grammatical function.

	Before		During		After		Total	
	#	norm	#	norm	#	norm	#	norm
Total (xxx)	51	5.3	0	0.0	0	0.0	51	4.4
Total (other)	86	9.0	8	8.1	4	3.8	98	8.4
Total (all)	137	14.3	8	8.1	4	3.8	149	12.8

Table 4.17 Overall expressive lengthening counts

In all categories demonstrated in Table 4.17 above, the use of expressive lengthening trends downward throughout the three conversational phases, which follows along with what the eBAC would predict for this female-coded feature. This holds true for the robustness of variety of words lengthened as the three phases progress, shown in Figure 4.7 above (which does not include words lengthened a single time by only a single author for brevity). The only words lengthened are introduced in the Before phase, which shares only *oo*, *lol*, *hm*, and *erm* with During, and only *what*, *nah*, and *erm* with After.

9. Backchannel sounds

Before	During	After
	oh	
hmm, ooo, erm,		
er		er
ah, ugh, ooh, aw, awh, um, pfft, ouch, ow	uh, eww, yo	oi

Figure 4.8 Distribution of backchannel sound variety by phase

Many of the backchannel sounds that occur in the data are subject to expressive lengthening, as demonstrated above in bold. Other backchannel sounds occurring in the data include *pfth*, *eww*, *ouch*, *ugh*, and *ha*, as well as some more male-coded backchannel markers such as *oi* and *yo*, but these occur less frequently than do others. This same trend follows with the distribution of varieties used, as shown in Figure 4.8 above.

	Before		During		After		Total	
	#	norm	#	norm	#	norm	#	norm
Total	188	19.6	13	13.2	7	6.7	208	17.9

Table 4.18 Raw and normed backchannel sound counts by phase

The decreasing use of backchannel sounds as each phase progresses follows the trajectory the eBAC would expect. Notably, backchannel sounds that may tacitly indicate agreement or understanding (*ah*, *ooh*, etc.) and thus more likely to be female coded (see assent markers below) occur only in the earlier phases, while the more male-coded backchannel sounds (*yo*, *oi*) occur exclusively in the latter two phases. As such, the decrease of female-coded backchannel sounds can be seen to occur inversely to the increase in male-coded backchannel sounds.

10. Hesitation words

While eBAC found hesitation words to be more female coded overall, some hesitation words, as are found in the data here, were found in the eBAC to be more male coded: *erm*, *hm*, and particularly *er*. Again, as indicated in bold, all but *uh* were additionally subjected to expressive lengthening.

Before	During	After
<i>hmm, erm,</i>		
<i>er</i>		<i>er</i>
<i>um</i>	<i>uh</i>	

Figure 4.9 Distribution of hesitation word variety by phase

That this feature—which expresses hesitation—occurs most frequently in the During phase is not unexpected due to the context of the phase, as this is when interlocutors are most negotiating the information they are receiving, which disagrees with their profile-based preconceptions. Context of the dataset aside, overall, the trajectory of this feature, which occurs less than half as much in the After than the Before phase, follows the trajectory the eBAC would expect (and though only one occurs in the After phase, as a matter of proportion, male-coded hesitation words can be seen to increase inversely with the overall decrease in their use).

	Before		During		After		Total	
	#	%	#	%	#	%	#	%
Total	26	2.7	6	6.1	1	1.0	33	2.8

Table 4.19 Raw and normed hesitation word counts by phase

D. Upward-trending ‘Male Coded’ Features

Of the features found by the eBAC analysis to be male coded, four are discussed here with overall downward trending trajectories. The first and last features, **5. Friendship** terms and **17. Articles and Determiners**, demonstrate consistent upward-trending use. The middle

two features contain downward trajectories with—**12. Negation Terms**—and without—**13. Swears and taboo words**—a spike in the During phase (discussed below), with the addition of female-coded variations of those features demonstrating the inverse pattern.

5. Friendship terms

The friendship terms that occur in the data analyzed here are *friend*, *bro*, *pal*, *dude*, and *mate*. These terms are generally used in one of two ways in the conversations—either to refer to the author or conversee’s friends, or to address the author or conversee themselves. (This also excludes *bro* as a shortening for a fraternal *brother*, as opposed to a friend.)

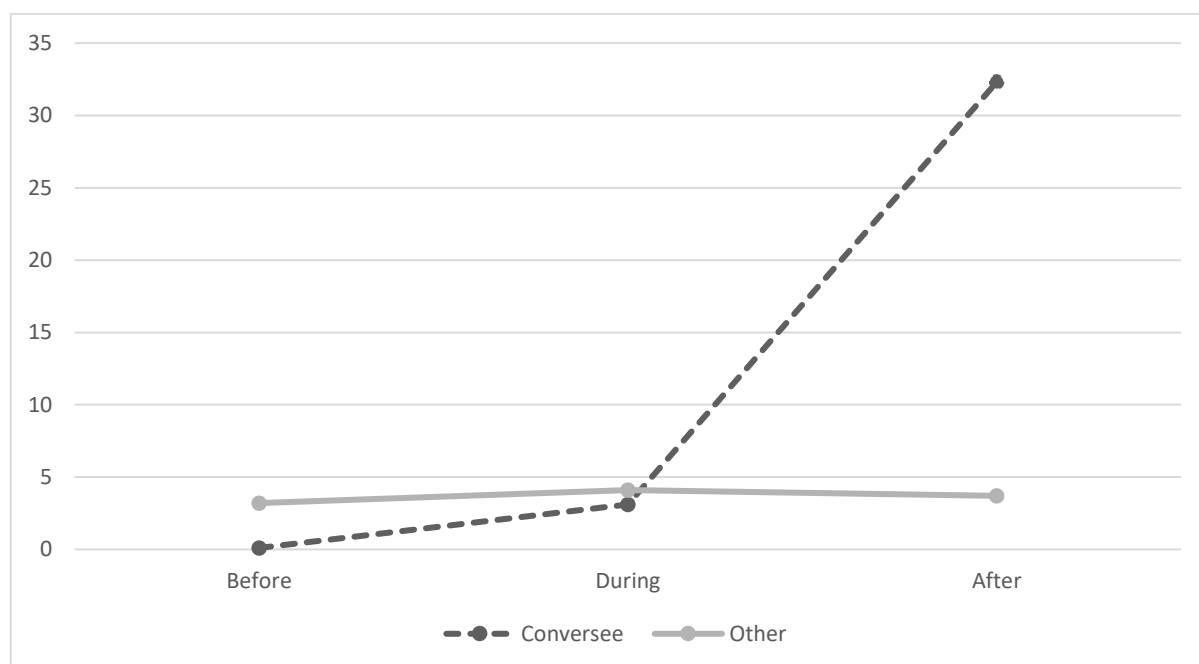


Figure 4.10 Friendship terms

Notably, as shown in Figure 4.10 above and Table 4.20 below, the term *friend* is used only to refer to other entities and not as a term of address, while the terms *pal*, *bro*, and *dude* occur only after the author realizes they are conversing with another male. Only the term *mate* shifts from referring to friends only Before, and then additionally and primarily to the conversee After.

	About Conversee						About Others					
	Before		During		After		Before		During		After	
	#	3k	#	3k	#	3k	#	3k	#	3k	#	3k
friend(s)	0	0.0	0	0.0	0	0.0	14	1.5	2	2.0	1	0.9
bro(s)	0	0.0	0	0.0	5	4.8	0	0.0	0	0.0	0	0.0
pal(s)	0	0.0	0	0.0	2	1.9	0	0.0	0	0.0	0	0.0
dude(s)	1	0.1	3	3.1	10	9.5	0	0.0	2	2.0	1	0.9
mate(s)	0	0.0	0	0.0	17	17.7	17	1.8	0	0.0	2	1.9
Total	1	0.1	4	3.1	34	32.3	31	3.2	4	4.1	4	3.7

Table 4.20 Overall friendship/“bro” term counts by phase

Examples of these uses, where they occur, are demonstrated in Figure 4.11 below.

Before	During	After
Im at my friends house extremely hung over	In your pic you are	I can't wait to tell my friends this shit!!
		Funny as fuck this BTW

Sorry niggas sayin u a dude but im like nah thats bae	wearing a green dress with two friends	bro , no idea how this happened
		Let's never speak of this again pal . Safe!
Because my mate from uni is near there and there is a really good gig on!	So you're a dude ?	Tinder's got you fucked then dude !
		Mate I've heard that tinder got hacked and guys have been chatting people up.

Figure 4.11 Examples of friendship terms in the data

Both *friend* and *mate* are used in the Before phase exclusively to refer to other people. *Friend* maintains this pattern even in the After phase, while *mate* shifts to the conversee. In the case of *dude*, in particular, the single instance Before is the author reporting someone *else* stating the conversee sounds like “a dude”, and they are not addressing them as *dude* themselves; and in the single instance During, the author is referring to themselves, stating, “I am a dude”.

Notably, only *friend* and *dude* cross into the During phase. In the *friend* example, the conversees are attempting to negotiate their identities by describing and attempting to reconcile their individual profile pictures, with *friend* used to describe other people in said pictures. In the *dude* example, along with most of the *dude* examples in During, *dude* is used to inquire as to the conversee's real-world gender.

Only *pal* is specific to the After phase, only ever used to refer to the other conversee. Although the terms used are overall more male coded in nature, the use of the terms does trend in line with what the eBAC might expect. Their use to refer to other parties (such as an author's friends) or to the conversee is clearly related to the phases of conversation and the understanding of the conversee's gender. The use of neutral and male-coded friendship terms drops as the conversation progresses, inversely with the rise in use of male-coded terms of address as an author becomes more aware that their conversee is in fact another male.

12. Negation terms

As with assent terms, negation terms are subject to expressive lengthening, and occur with relative stability throughout the three phases of the conversation. Table 4.21 considers the only attested variations as suggested by Bamman et al. (2014), with expressively lengthened *noo+* as a female-coded feature, *nah* as a male-coded feature, and *no* as the non-coded variation.

		Before		During		After		Total	
		#	norm	#	norm	#	norm	#	norm
female coded	noo+	4	0.4	1	1.0	0	0.0	5	0.4
	cannot	3	0.3	0	0.0	0	0.0	3	0.3
female coded subtotal		7	0.7	1	1.0	0	0.0	8	0.7
male coded	nah	18	1.9	3	3.1	6	6.2	27	2.3
	nobody	0	0.0	0	0.0	0	0.0	0	0.0
	aint	9	0.9	0	0.0	0	0.0	17	1.5

male coded subtotal		27	2.8	3	3.1	6	6.2	44	3.8
other	no	107	11.1	24	24.4	25	23.9	156	13.4
Total		141	14.7	34	33	31.4	32.2	208	17.9

Table 4.21 Overall negation term counts, female and male coded

Like assent terms, which may also be indicative of the cooperation inherent to Tinder conversations, the frequency of the negation terms in the Before phase likely owes to the nature of Tinder conversations, and the yes/no questions often asked by conversees attempting to get to know one another. While the female-coded negation term shows little variation and little use overall, male-coded *nah* follows the trajectory the eBAC might expect—as do the use of negation terms overall.

13. Swears and taboo words

Though anti-swears like *gosh* and *dang* occur infrequently and only in the Before phase of the conversations, swears occur throughout all three stages of the conversations. Of course, *fuck* can be used both as a swear word, and as a taboo term.

		Before		During		After		Total	
		#	3k	#	3k	#	3k	#	3k
male coded	damn	10	1.0	0	0.0	0	0.0	10	0.9
	fuck	24	2.5	7	7.1	48	45.6	79	6.8
	shit	12	1.2	0	0.0	6	5.7	18	1.5
	faggot	0	0.0	0	0.0	1	1.1	1	0.1
	bastard	2	0.2	0	0.0	0	0.0	2	0.2
	ass	10	1.0	1	1.0	0	0.0	11	1.0
	cunt	0	0.0	0	0.0	1	1.0	1	0.1
Subtotal		56	5.8	8	8.1m	54	51.3	118	10.8
female coded	gosh	1	0.1	0	0.0	0	0.0	1	0.1
Subtotal		1	0.1	0	0.0	0	0.0	1	0.1
Total		57	5.9	8	8.1	54	51.3	119	17.1

Table 4.22 Overall swear and anti-swear counts (raw)

These variable uses follow similar inverse trends as the friendship terms discussed above, in that the uses in the Before phase are frequently taboo terms, and the uses in the After phase are almost entirely swears, as shown below.

Before	During	After
I'd be very happy to give you a hard fucking	What the fuck	The actual fuck is going on tinder?!
dat would make me fuck u so bad		
	ur fucking my brain	FUCK OFF
The ocean said ' fuck off beach I'm tired of you pushing me around'		
If you're not fake it's a fucking miracle	What the fuck 😏😏😏	I'm gonna find you and fuck you up
Ah well tinder done fucked up	I fucking am!	You fucking creepy fuck
Well youre also pretty as fuck	What the fuck ?	And dolly Is fit as fuck

I'm an ugly fucker called Luke :)		I thought was strange a 19 old girl coming out saying id fuck u 😏 😏
I'm not fucking with you	This is fucked up	

Figure 4.12 Examples of “fuck” words in the data

The frequency of swears throughout the three phases, however, shows a clear trajectory toward their use. Like kinship terms, swears can be a relatively conscious, hypermasculine feature, and the fact that they can be employed whether an author is angry or understanding of the situation likely accounts for their spike in frequency.

17. Articles and determiners

Table 4.23 below demonstrates a difference in use of overall use of articles and determiners. The latter two phases, During and After, have only a difference of 2.2, while the difference between them and the Before phase is over 100.

	Before		During		After	
	#	norm	#	norm	#	norm
the	428	44.6	36	36.7	35	33.3
a, an	750	78.1	119	121.2	135	128.4
this, that, these, those	458	47.7	59	60.1	81	77.0
my, your, his, her, its, our, their	526	54.8	106	108.0	78	74.2
one, ten, twenty, etc.	87	9.1	11	11.2	6	5.7
all, both, half, either, neither, each, every	90	9.4	11	11.2	20	19.0
other, another	22	2.3	1	1.0	3	2.9
such, rather, quite	39	4.1	0	0.0	1	1.0
wh- determiners	689	71.8	70	71.3	63	59.9
no	107	11.1	24	24.4	25	23.8
TOTAL	3,196	332.9	437	445.2	447	425.0

Table 4.23 Overall use of determiners and articles throughout the phases

Despite this upward trend, which can largely be accounted for by the rise in use of both indefinite articles (*a, an*) and pronouns and possessive determiners (*my, your, his, her, its, our, their*), the definite article *the*, the use of numbers as determiners, and the use of wh-determiners are highest in the Before phase. This is reflected in the overall ratio of uses per each phase, where the definite article, number determiners, and wh- determiners are the highest in the Before phase. The use of pronouns and possessive determiners are the highest in the During phase, and the remaining determiner classes are highest in the After phase.

	Before	During	After
the	13.4%	8.2%	7.8%
a, an	23.5%	27.2%	30.2%
this, that, these, those	14.3%	13.5%	18.1%
my, your, his, her, its, our, their	16.5%	24.3%	17.5%
one, ten, twenty, etc.	2.7%	2.5%	1.3%
all, both, half, either, neither, each, every	2.8%	2.5%	4.5%
other, another	0.7%	0.2%	0.7%
such, rather, quite	1.2%	0.0%	0.2%
wh- determiners	21.6%	16.0%	14.1%
no	3.3%	5.5%	5.6%

Table 4.24 Ratio of use of determiners and articles throughout the phases

This trajectory would be predicted by the eBAC.

E. Mixed or Unclear Features

The eight remaining features in this section are mixed, in that they include one or more feature variations that follow the expected trends, but one that does not, or their results are otherwise unclear (discussed below). these include **3. Emoticons**, **6. Abbreviations**, **7. Punctuation**, **11. Assent terms**, **14. Prepositions**, **15. Alternative spellings**, **16. Conjunctions**, and **18. Post length**. In many cases, Bamman et al. (2014), the eBAC, or other sources found inconclusive or contradictory conclusions regarding the below features, though many of their aspects are still noteworthy for discussion here.

3. Emoticons

Emojis (such as 😊, 😊, and 🙌) and emoticons (such as :), ;), and 😊) occur 930 times, or roughly once every 32 words, in 72 varieties total. For brevity, the following only considers emojis, though emoticons follow similar trends.

Most common (Tinder)				Most common (data)					
		String	Total	Collective Strings			Individual Totals		
1	😊	3.3%	1.1%	😊	53	44%	😊	128	34%
2	🙌			🙌 😊	19	16%	🙌	34	9%
3	👉			😊	12	10%	😊	23	6%
4	🍷			😊 👍	9	8%	😊	14	4%
5	🌍			😊 🙌	8	7%	😊 🙌	12	3%
6	😍	4.1%	2.2%	😊 🙌	7	6%	🙌	11	3%
7	🍷			😊	6	4%	😊	10	3%
8	☕			😊 😊	5	4%	👍	9	2%
9	🍷			😊 😊 😊 😊	4	3%	😊	8	2%
10	🙌			😊 😊 😊 🙌 🙌	3	3%	🙌 😊 😊 😊 😊	7	2%
Other					36%		Other	20	20%

Table 4.25 Most common emojis on Tinder and in the data (overall % of frequency)

Although emoticons and emojis are thought to be female-coded features, it is difficult to tell from the distribution found here, as compared to Bamman et al. (2014) and Schler et al.'s (2006) findings, how marked the emoticon usage is within this data. Table 4.25 demonstrates the top 10 used emojis on Tinder according to the company itself in 2016, as compared to the top 10 emojis found in the data. Of the top 10 emojis listed by Tinder, only two occur in the data (1 and 6), and while they do occur within the top 10 most frequent in the data, both account for less than 2% of the overall emoji usage. More than 50% of the emoji usage in the data is distributed among the top 5 emojis, none of which make Tinder's list.

This comparison alone might suggest that females on Tinder *do* use much more emojis than males do, and that the emojis they most frequently use differ greatly from those most frequently used by males. It is worth noting, however, that many of the top emojis proposed by Tinder—such as the beer, pizza, music notes, coffee, and wine—are often used when arranging a date between conversants, an act that does not occur with overmuch frequency in the data analyzed here. (So is dancing, though this emoji is of a woman dancing, and not the male equivalent.)

Another possible explanation for the use and variety of emojis found in the data is the fact that the authors suspect they are talking to females. As shown in Table 4.26, the total

words per emoji reaches a peak in the During phase where emoji use is the most concentrated, and drops to the least frequent use in the After phase. The range of emoji types used, on the other hand, gets consistently and significantly smaller as the conversations progress through the three phases.

	Total Individual Emojis	Total Instances (cumulative)	Total Types of Emojis	Words Per Emoji	Words Per Instance
Before	237	169	54	115.1	161.4
During	86	47	18	31.6	57.8
After	44	19	9	65.6	151.9
Total	372	241	54	88.4	136.4

Table 4.26 Overall counts of words per emoji, and types of emoji, by phase

That the range (number of types) and variety (use related to meaning) of emojis used follow the trajectory expected by Bamman et al. (2014), but the frequency of use, either by string or individual use, does not, may itself be explained by the context of the After phase, or the conversations in their entirety. The much higher use of emojis in the During phase is perhaps explained by politeness, accommodation, and negotiating.

	Before	During	After	Total
1				
2				
3			--	

Table 4.27 Most frequent emoji by phase

In most of the instances, however suspicious an author might be, they attempt to negotiate the discrepancies in the information they are being given with the information provided in the profiles, and the use of emojis is a possible way of reducing the tension. But as with the After phase, the range of emojis required to express either politeness, acceptance, or anger, is much lesser than in the less specific conversation trajectories in the Before phase.

Before	During	After
Switch the food for alcohol and it's like the best date ever 😊	Please tell me I haven't been catfished 😊 😊 😊	I am not Sybil no 😊
Hahahha you seem really funny 😊	😊 😊 😊 Is this a joke	No here now your pic is a Chinese women 😊
Haha why so shy 😊	Haha I give up who ever this really is, nice trolling 😊 😊	Lmfao, tinders got you fucked then dude! 4 pictures of some girl and the name dolly 😊 😊 😊
I am a mind reader yes I won't bore you with 😊	Don't think thats me 😊	This has actually made my night tbh 😊 hahahaha
Noooooo you have to tell me 🙄	Wtf is going on here 😊	Nah man me neither haha 😊
Good morning 🙄	It's someone else as your profile pic 🙄 🙄	This is so fucked up man 😊
I'm not really making a	I give up. 🙄	

good impression here 🤔	U got woman pics and say u got beard 🤔🤔	yeah it's a lass 🤔
------------------------	--	--------------------

Figure 4.13 Examples of the most common emoji uses in the data

In all three phases, the top emoji remains the same as the overall top emoji used—known as the tears of joy emoji, that depicts laughing and crying. Although this emoji itself is likely used differently in all three phases, the next most frequent emojis are perhaps more demonstrative of the shifts.

In the Before phase, the second most common emoji is a winking face—something expected in flirtation.

Both the Before and During phases use the “see no evil” monkey, commonly used “as a way to express embarrassment in an amusing way,”¹¹ and “shifty eyes”, which, “is good for drawing attention to concerning behavior, albeit in a somewhat judgmental way. In this same vein, it can be used to express disbelief or disapproval at a situation.”¹²

The After phase uses various smiling faces *not* considered flirtatious once each and hence has no third most frequent, but also uses the “grinning face with sweat” second most frequently—an emoji not frequent in either other phase. This emoji is “intended to depict nerves or discomfort but commonly used as a means of expressing ‘whew!’ or ‘close call!’ that would be implied when a person wipes sweat from their brow in an exaggerated manner.”¹³

6. Abbreviations

Before	During	After
	wtf, lol	
smh, btw, omg, tbh		smh, btw, omg, tbh
ru, fb, ig, bbm, bj, gf, dtf, wyd, idk, wby, ofc	ffs	np

Figure 4.14 Distribution of abbreviation variety by phase

Because the use of such initialisms as *lol* (including *lmao*, *lmfao*, *loo+l*, and *lol+*) vastly outnumbers even the total count of every other initialism, the two counts are separated.

	Before		During		After		Total	
	#	3k	#	3k	#	3k	#	3k
TOTAL (lol/lmao)	168	18.5	16	16.3	15	15.6	199	18.1
TOTAL (other)	86	9.5	4	4.1	13	13.5	103	9.4
TOTAL (all)	254	27.9	20	20.4	28	26.6	302	25.9

Table 4.28 Raw and normed abbreviation counts by phase

Overall, the use of abbreviations trends downward, a trend is matched by the *lol* abbreviations, which follow what would be expected by the eBAC—that in the After phase, when authors know they are conversing with another male, their overall use of the female-coded feature of abbreviations should drop. Non-*lol* abbreviations, however, have the opposite trend.

There is a drop in the trend for all three in the During phase, where abbreviations are

¹¹ <https://www.lifewire.com/less-obvious-emoji-meanings-3485884>

¹² <http://www.dictionary.com/meaning/eyes-emoji>

¹³ <https://emojipedia.org/smiling-face-with-open-mouth-and-cold-sweat/>

used notably less than in either other phase. This can possibly be attributed to a need for clarity, as in the During phase, authors are attempting to negotiate the identity of the interlocutor and navigate between the information their profile presents (of a female looking for a male partner), with the information presented within the conversation itself (of a male looking for a female partner). That the use of non-*lol* variations does not follow the expected trajectory may come down to the necessity of informality in the After phase, as authors often attempt to brush off the misunderstanding; if this is indeed the case, overt informality may itself trump some otherwise female-coded features.

The lack of abbreviations in the During phase for the sake of clarity is perhaps bolstered by the trend of the *lol* variations, which are themselves a better match for the trends the eBAC might expect. The *lol* variations maintain a consistent, downward trajectory throughout the three phases, though the difference of 2.9 (normed) words from the Before and After phases is not large. This might suggest an overall interest in using accommodative *lol* variations that would not be inherent in other abbreviations.

Finally, although the overall frequency of abbreviations does not follow the pattern the eBAC would expect, the variety of variations used does, as shown in Figure 4.14 above. Only variations of *wtf* and *lol* occur throughout all three phases, and only *ffs* (*for fuck's sake*) is unique to During, while only *np* (*no problem*) is unique to After. Before and After share 4 more abbreviations, while a further 11 are unique to Before. Thus the variety of use both follows the eBAC findings in that the variety goes from more to less robust, and the idea of accommodation for clarity in the During phase which has the least robust variety, containing only three variations total.

7. Punctuation

The eBAC found that males use more periods as a percentage of total sentence-final punctuation, and that females contain more of every other variety. This is starkly different in the Tinder data, in which questions are the most frequent across every phase. This may in part be due to the genre of such conversations, in which one of the goals is to learn information about the person with whom you are speaking. However, that the question marks spike drastically in the During phase as compared to the other two, and are almost even with periods in the After phase, is clearly demonstrative of the context of this dataset. That is, the During phase can be seen as a particularly investigative one. As such, this feature may not be the most probative in a context this specific, or even in the more general context of conversations more broadly.

	Before		During		After	
	#	%	#	%	#	%
Total Periods	734	31.5%	57	20.9%	65	35.5%
Total Ellipses	59	2.5%	4	1.5%	17	9.3%
Total Exclamations	261	11.2%	14	5.1%	29	15.8%
Total Questions	1,279	54.8%	198	72.5%	72	39.3%
Total All	2,333	--	273	--	183	--

Table 4.29 Overall distribution of punctuation

sentences. This trend, then, would follow what the eBAC would expect of the After phase being one in which authors are overtly performing their maleness.

	Before	During	After
Words per Punctuation	11.6	10.0	15.8

Table 4.30 Overall distribution of punctuation per unit

Finally, the eBAC found that women tended to have a higher distribution of lengthened punctuation. This feature does not appear to follow that trend in the Tinder data, but this may be due to a similar context as the first component of this feature. That is, because the Tinder data occurs in a context rife with opportunities for confusion and disbelief, that the lengthened punctuation is much higher in all three phases than in either the male (7.4%) or female (9.8%) eBAC subcorpus, and particularly higher in the During phase, seems to be indicative of this likelihood.

	Before		During		After	
	#	%	#	%	#	%
standard punctuation	1,806	77.4%	173	63.4%	130	71.0%
lengthened punctuation	527	22.6%	100	36.6%	53	29.0%

Table 4.31 Distribution of standard/unlengthened and lengthened/complex punctuation

11. Assent terms

As mentioned above, assent terms are also subject to expressive lengthening. Unlike many of the other features, assent terms overall remain frequent in use throughout all three phases of the conversation. This is likely because question-answer pairs, including yes/no questions, are a staple of Tinder conversations, as the point is, first, to get to know one another, and in the data, second, to determine the veracity of the profiles conversees initially rely upon.

		Before		During		After		Total	
		#	norm	#	norm	#	norm	#	norm
female coded	yes	42	4.4	7	7.1	1	1.0	50	4.3
	yeah	161	17.1	11	11.2	19	18.1	194	16.7
	yass	1	0.1	0	0.0	0	0.0	1	0.1
	okay	20	2.1	1	1.0	1	1.0	22	1.9
	ok	51	5.3	1	1.0	5	4.8	57	4.9
Subtotal		275	29.0	20	20.4	26	24.7	324	27.9
male coded	yeh	14	1.5	1	1.0	0	0.0	15	1.3
	yep	6	0.6	0	0.0	1	1.0	7	0.6
	yup	7	0.7	1	1.0	1	1.0	9	0.8
	yea	11	1.1	2	2.0	0	0.0	13	1.1
Subtotal		38	4.0	4	4.1	2	1.9	29	2.8
Total		316	32.9	24	24.4	28	29.1	368	31.6

Table 4.32 Overall assent term counts

That the assent terms dip in use overall in the During phase may simply be because an author in the During phase (who is likely conversing with someone else still in the Before or already in the After phase) will likely be more interested in asking than answering questions as a means of clarifying any perceived disparities between the identity presented in their conversee's profile, and the identity presented by their conversee's language.

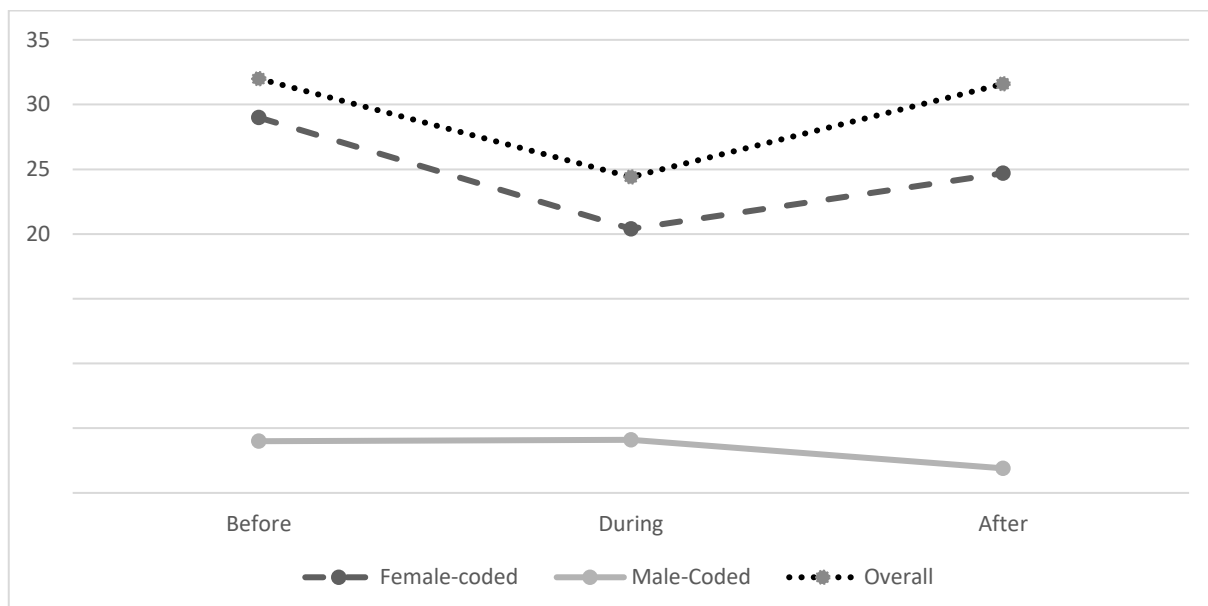


Figure 4.15 Assent terms

The eBAC finds that assent terms are overall female-coded markers. Notably, although less frequently, the data contains more male-coded assent terms, such as *yeh*, *yep*, and *yup*, and only once, the more female-coded assent term *yas*. That the male-coded varieties drop in use along with the overall trend of female-coded assent terms may simply be down to the context of the exchanges—there is less to agree to as the conversations progress, as authors are confronted with something other than that which they expected.

14. Prepositions

As shown in Table 4.33 below, the data would appear to follow the opposite trajectory that the eBAC would expect as male coded, in that prepositions are most used in the Before phase, and least used in the After phase, with a difference of 90.1. It is worth noting, however, that some sources found only alternative spellings, such as *2* for *to* (male coded) and *w/* for *with* (female coded). No such alternative spellings were found in the data, which, following the eBAC findings, would indicate male authors throughout.

	Before		During		After	
	#	norm	#	norm	#	norm
overall preposition use	2,672	278.4	213	217.0	198	188.3

Table 4.33 Overall preposition use throughout the phases

15. Alternative spellings

While most specific alternative spellings considered in this analysis are female-coded, it appears that the use of alternative spellings overall is not purely male- or female-coded. However, as shown throughout the sections above, alternative spellings found in the Tinder data do largely appear to follow along with what Bamman et al. (2014), the eBAC, and others would expect throughout the three phases—that is, more female-coded in the Before phase, and more male-coded in the After phase. That some female-coded alternative spellings, such as negation terms and pronouns, occur most frequently in the During phase, may be explained by the suggestion that female-coded language tends to be more cooperative and collaborative, a useful tool in the During phase when at least one speaker is attempting to negotiate their

conversee's identity (e.g., Tannen, 1992; Ersoy, 2008).

Finally, as this feature is relevant throughout the features above, and few alternative spellings are exhibited in the data that fall outside the realm of the other features considered, and are not clearly stylistic rather than typos, this feature is not considered further here.

16. Conjunctions

	Before			During			After		
	#	norm	%	#	norm	%	#	norm	%
and	456	50.2	97.6%	37	40.9	100.0%	29	30.1	93.5%
&	11	1.2	2.4%	0	0.0	0.0%	2	2.1	6.5%

Table 4.34 Normed and raw frequencies of “and” and “&” throughout the phases

The normed frequency of overall *and* conjunction use drops as the phases progress, as shown in Table 4.34 below. The use of the more female-coded & stops completely in the During phase and rises in the After phase. This pattern is shared by the ratio of use between *and* and &, where the more female-coded variant & accounts for a higher ratio of *and* variants in the After phase. The overall use of conjunctions would follow what the eBAC would expect, while the use of & would not, and may simply be down to the difference between ease of use in texting versus typing.

18. Post length

Length achieved mixed results in the eBAC analysis, which just barely contradicted Schler et al.'s (2006) findings that women wrote more words per blog post than men, but also found that women averaged more posts and thus more words per blog, even with slightly fewer words per post. Obviously, it is not possible to directly compare all averages between the three phases as the Before phase is around ten times larger than the latter two phases, with just under twice the author participants. But the average words per message would appear to follow the trajectory Schler et al. (2006) would expect of males writing less per post/message.

	Before	During	After	Total
Total Words	27,272	2,717	2,886	32,875
Total Messages	2,889	465	404	3,758
Total Users	107	63	66	108
Words Per User	254.9	43.1	43.7	304.4
Messages per user	27	7.4	6.1	34.8
Words Per Message	9.4	5.8	7.1	8.7
% of total words	83%	8.3%	8.8%	--
% of total messages	76.9%	12.4%	10.8%	--
% of total users	99.1%	58.3%	61.1%	--

Table 4.35 Average post length throughout the three phases

Because of the differences between blogs and Tinder conversations, however, it is unclear whether this feature is particularly probative in this case.

F. Overview

As demonstrated by the above features, the findings of the eBAC appear to largely hold even in the subverted context of hacked Tinder conversations. In the case of at least ten of the eighteen features considered above, the all-male interlocutors tend to begin interactions with more female-coded language features and end them (as in the After phase) with more male-

coded language features. Of the remaining eight, five had at least some feature-internal variations (such as the range and variety of emoticons but not also the frequency of their use following a male-coded trajectory) that followed their expected trajectory. Although there are exceptions to this, and some inconsistent variations, they may, rather than contrary to the eBAC's findings, be largely contextually motivated by the particularities of such interactions, as explained above.

Such deviations would include any of the female- or male-coded features that spike or dip, seemingly unexpectedly, in the During phase (despite ending with an overall matching trajectory in the After phase). The spike in female-coded hesitation words in the During phase may be due to the confusion inherent in the transition. The spike in male-coded negation terms in the During phase may be contextually motivated by an author disagreeing with the information they are being provided with (or deflecting accusations themselves, in some cases), which does not match their previously profile-based understanding of their interlocutor's identity.

What we may draw from this is confirmation of Bucholtz and Hall's partialness principle. In this data, identity performance seems to be both partially unconscious, and partially deliberate. This is evident in a couple of ways. The features that tend to exhibit hypermasculinity in the After phase of the conversations—especially kinship terms and swear words—are those that likely exist at least partially *above* the level of conscious thought. Other features that trend more female coded in the Before parts of the conversation—such as emoji use and expressive lengthening—likely exist at least partially *below* the level of conscious thought. Along those same lines, in the instances where male and female features inversely rise and drop, respectively—hesitation words, negation terms, and swears and anti-swears—would appear to indicate that these gender-subverting variations of otherwise oppositely gender-coded features can be particularly useful markers.

These gendered features can be seen to be **emergent** as interlocutors' identity performances shift through the contexts of the three phases, and **indexed**, for example, by using not only overtly male coded, but also overtly hypermasculine features as are found in the After phase. Conversely, aspects of the remaining six features that did not follow the expected trajectories may have done so for a variety of reasons, which are discussed individually in subsection 3 above. It may be, then, that these features are too context-driven to be useful in every such analysis or are otherwise poor markers of gender.

Part 4 – Tinder 2

As discussed in the corpus analysis of the Catfi.sh Tinder data in Part 3 above, an individual author's half of a conversation progresses through a possible three phases, termed here Before, During, and After, explained below:

- Before occurs at the outset of a conversation, when an author takes the information provided in their conversee's profile to be the ground truth upon which their interaction is based. In this phase, all participants believe they have been matched with an

attractive female interested in men. Negotiation of information that appears to contradict a conversee's may begin in the Before phase, resulting not in suspicion as to misinformation by either the conversee or their profile, but in the subversion of their expectations of the profile. All authors but one begin in the Before phase, and it is often the longest phase by far, making up roughly 80% of the total data.

- During occurs once an author begins to express suspicion as to the identity of their conversational partner. In this phase, authors finally have enough reason to take note of inconsistencies between the information they are being provided by their conversee and what details their profiles entail. Negotiation is often prevalent in this phase, in which an author begins to question whether there is some trickery afoot, either by their conversee, some glitch on Tinder, or elsewhere. Only one author begins with the During phase, but does so after their conversee has begun the conversation with information that does not match their Tinder profile.
- After occurs once an author has received confirmation of their suspicions that the information they have been provided by their conversee indicated an identity other than the one presented by the Tinder profile. This occasionally occurs without suspicion expressed in either previous phase, or without any During phase, as their conversee or some other evidence flat-out indicates that they are male, and this information is immediately accepted as true. The After phase does not usually contain negotiation of their conversee's identity, but may contain negotiation as to the nature of the misunderstanding. This phase is often short, and any given author can reconcile this phase in a variety of ways, such as with anger or humor. While many conversations reach the After phase, none of the 108 authors begin in the After phase.

Although almost all conversations begin in the Before phase, not all make it to both or either the During and After phases, for a number of reasons.

A. Conversational Types

These types of transitions are shown in Figure 4.16 below, and account for 11 different Types of conversational progression, termed here A-K, showing the transitions of pairs of conversees throughout all three phases.

	Before	During	After	Before	During	After	Total
A	X	X	X	X	X	X	12
B	X	X	X	X	—	—	10
C	X	X	X	X	—	X	8
D	X	—	—	X	—	—	5
E	X	X	X	X	X	—	5
F	X	X	—	X	—	—	4
G	X	—	X	X	—	—	4
H	X	—	X	X	—	X	3
I	X	—	—	—	X	X	1
J	X	X	—	X	X	—	1
K	X	X	—	X	—	X	1

Figure 4.16 Conversational trajectory types

An example of each of these conversational Types is demonstrated below, along with short analyses of the transitions, and examples of how these conversations are either typical

of their conversational Type, or differ from other possible norms or exceptions. Although mentioned here, strategies of negotiation are discussed in Section C below.

1. Type A – James and Tom (39)

	Before	During	After	Before	During	After	Total
A	X	X	X	X	X	X	12

Figure 4.17 Type A conversations

The most common type of conversational progression, as shown in Figure 4.17 above, is one in which both participants make it through all three stages of conversational progression—Before, During, and After. This type, Type A, accounts for 12 of the 54 total conversations, or roughly 22%, and an example of such conversations is the one between James and Tom, shown below.

	James	Tom
1	Hi	
2		Hey
3		How's it going?
4	All good thanks, caning through coffee	
5		Needed on a Monday
6		Just had a red bull
7		00wide awake!
8	Haha	
9	So what do you work as?	
10		I'm a carpenter/foreman on a site in XXX at the moment
11		How about you?
12	Joining XXX XXX in a month but at the minute labouring from site to site	
13	You're really a foreman?	
14	I just find it hard to believe a pretty girl like you works on a building site	
15		Your labouring!?
16		Pretty girl haha?
17		What's going on here
18	Yeah I don't mind it	
19	I think honesty, haven't tried it in a while	
20		I'm confused
21	Why?	
22		So you're a pretty girl working on a site.
23		What do you mean about honesty?
24	No your a pretty girl working on site	
25		Haha!
26		No you are
27	Well this is weird	
28		Haha
29		Your the girl working on site!
30		Unless your a ladyboy
31	I'm pretty sure I'm a man, I can prove it	
32		Your a man?
33		Albertina?
34		This is strange
35		I'm pretty sure I'm a man also
36		My pictures are a lot more manly then your haha
37	James actually	
38	My name is	
39		Well that's not what it says!

40	Getting strange now
41	Well...
42	Well...
43	Ever made love to a man?
44	Haha
45	Why are your pics of women?
46	And no

Figure 4.18 James and Tom (39)

Both conversants start out in the Before phase of the conversation. At 10, Tom states that he works as a foreman, and at 13 and 14, James responds, beginning his During phase by questioning the likelihood that “a pretty girl like [Tom] works on a building site.” Similarly, James states that he is “laboring from site to site” in 12, which prompts the same skepticism in 15 from Tom—“Your labouring!?”—and begins his During phase.

Both conversants negotiate this phase in varying ways. Tom takes the mismatch of information as a joke, as in 25-26 with “Haha! No you are”. Both During phases continue until 31, with the statement from James “I’m pretty sure I’m a man, I can prove it”, which prompts Tom’s After phase. Similarly, in 35, Tom’s statement of “I’m pretty sure I am a man also” prompts the After phase for James.

Both sides of the conversation transition for the same reason—because a conversee is presented with information that does not appear to match their conversational partner’s Tinder profile—which they have taken for the ground truth of that person’s identity. In this case, both James and Tom have physically demanding jobs which, while not impossible, are unlikely professions for young females.

	Daniel	Joe
90		Just text you x
91	Dolly & joe? 😊	
92		Haha
93	I'm confused haha	
94		Joe and dolly?
95	Explain 😊?	
96		No dolly and joe sounds better
97		Confused about what?
98	Yeah?	
99	“Guy you been talking to on tinder?” 😊	
100		Doesn't matter
101	Explain 😊	
102		Haha well you said to text you then asked who it was
103		What else was I going to put
104	You said guy? 😊	
105		And explain what the joe and dolly comment?
106		Haha this is getting confusing
107	You text me saying “hey dolly, it's joe”	
108		What?
109		Yeah...
110	Who's dolly and who's joe 😊 😊	
111		What? Your display name is dolly and my name is Joe? What is going on?
112	What? My name is Dan 😊 and your name says Sybil	

113	Am I missing something?
114	I'm really confused now!
115	Am I?
116	My name says dolly?
117	I genuine don't know what is going on. I think tinder has fucked us over. My name is Joe not Sybil. And your name is dolly apparently. I take it you're a guy?
118	It's starting to look that way?! What the fuck?! Yeah my name is not dolly
119	You even have pictures of this "Sybil" girl on your profile! The actual fuck is going on tinder?!
120	Haha sorry mate. I honestly don't know what has gone on here

Figure 4.19 Daniel and Joe (4)

Other such Type A conversations transition in similar ways, such as the conversation between Daniel and Joe (4), excerpted above.

Unlike James and Tom, who reach all three phases of progression within a single, relatively swift conversation, Daniel and Joe's interaction lasts days, with the first 89 messages occurring before the During phases shown in the excerpt above. Daniel and James also discuss their professions, but a personal trainer and a bartender don't set off any red flags for either party. The rest of their Before phases are spent discussing relatively gender-neutral topics: the areas they live in, food and cooking, what they are doing for the upcoming holiday, and so on.

It is only between 90 and 91, when they attempt to transition their conversation from Tinder to texting, that Joe's text of "Dolly and Joe" begins to confuse Daniel—whose name is not, in fact, "Dolly", and who hasn't been talking to a Joe, but "Sybil". Again, this information contradicts the ground truth of their partner's Tinder profile, this time because of the names. This is confirmed for Daniel when Joe texts him "guy you've been talking to on Tinder" off of Tinder to clarify any misgivings about "Dolly and Joe" or "Joe and Dolly".

In the case of Type A conversations, both participants are able to negotiate any misinformation and come to the same conclusion—that they have each been talking with a male the whole time—regardless of their reasoning behind it. As we will see in other types below, however, this is not always the case.

2. Type B – Iknoor and Bradley (53)

	Before	During	After	Before	During	After	Total
B	X	X	X	X	—	—	10

Figure 4.20 Type B conversations

As with Type A above, Type B conversations are those in which at least one participant reaches all three phases of the conversation. However, as shown in the conversation between Iknoor and Bradley below, in Type B conversations, their conversee never advances past the Before phase. This is the third most frequent type of conversation, accounting for 19 of the 54 interactions, or roughly 19%.

	Iknoor	Bradley
1	Sadie. You are too fucking cute lol	
2		Ur not real

	Iknoor	Bradley
3	Lmfao why would you say that	
4		So many people are fake on this thing
5		If ur real send me a pic on snapchat
6	Lmfao i mean the fake ones usually send like a "webcam link"	
7	Ok wb instagram?	
8	I dont have snapchat	
9		I don't have IG
10		So you think I'm too fucking cute?
11		What would u do to something this cute?
12	Ugh why not	
13	But yes you have pretty eyes	
14	Lmao i dont think you want to know right now 😏	
15		Well I'm kinda only looking for a hook up on here
16		Oh really? Send me some naughty pics then
17	Where are you from?	
18	XXX im guessing ?	
19	WhT?	
20	How about you first	
21		XXX. U?
22		Ok but will u send one in return?
23	XXX.. But im out rn. Can send u one later lol	
24		What's RN?
25	Right now lol	
26		Ohhhh hahaha. So are u looking to fuck?
27	I would love to but i cant at the moment 😏	
28		Hahahahahaha Maybe this week?
29	No like i cant for a while	
30	I got cut a few weeks back. Doc said no lmfao	
31		Cut in ur vagina?
32	No i got circumcised	
33	🙄	
34		What does that mean
35	Lmfao wait you dont know what that is?	
36		Men get that not women
37	Im a male...	
38		Oh cool peace then
39	😏 how did i look like a girl wtf	
40	But ya take care	
41		Huh
42	Did you think i was a girl	
43		Ur pics are
44	I have a beard	

Figure 4.21 Iknoor and Bradley (53)

In the case of the above conversation, although it is brief, Bradley's During phase occurs between 34 and 36, in which he expresses confusion as to why a female would need to get a circumcision. When Iknoor states in 37 that he is male, Bradley immediately transitions to the After phase, and the conversation subsequently ends.

On Iknoor's end of the conversation, however, he has been given no indication that Bradley is not who his female Tinder profile states he is, and he is never given any reason to transition to the During or After phases before the conversation ends, only expressing

confusion as to why Bradley could have thought he was female.

This pattern remains relatively the same throughout all Type B conversations. The conversation that continues for the longest time after one conversee reaches the After phase is between Will and Ata (54), excerpted below.

	Will	Ata
14		0XXX add me on what's app. 😏
15	Why is your profile picture of a guy?	
16		That guy is me...? lol
17	Your a guy?	
18		I've got a wide 8" that makes me pretty sure of that fact little lady.
19	What do you get by trolling tinder? I was really excited to have matched with a pretty girl and now it turns out you're a dude with an 8" cock 😏	
20		Haha, so you dtf? 😏
21		I'm cool with 3 ways...
22		We could team up and try find another girl who's game. 😏
23	Do you know any allot of girls here in XXX? I'm new here and if you could help me find a girl that would be amazing	
24		Depends... Are you down for a threesome? X
25	Why do you want me?	
26		Let's put it this way... I'd jump in my car right now and drive to you in heartbeat if you but gave me the word. 😏 X
27	Why? I'm not good looking	
28		I think you're beautiful 😏
29	Thanks. I wish more girls thought that 😏	
30		Do you not find men attractive at all?
31	No, sorry	

Figure 4.22 Will and Ata (54)

As with the conversation above, Will transitions quickly between all three phases, the During phase lasting only in 15 after they have transitioned to WhatsApp. Unlike the conversation above, Will and Ata continue to speak, during which Ata could transition to the latter phases as well. However, both instead negotiate their new understanding of the interaction. Will understands Ata to be a male, looking for another male to have a threesome with a female; Ata understands Will to be a female looking to meet other females, which is why he suggests a threesome between the two of them and another female.

As we see, although the most negotiations between identity and information occur in the During phase, they can occur in the After and Before phases as well, leading to one or both participants not fully transitioning between the three phases.

3. Type C – Jonny and Joshua (13)

	Before	During	After	Before	During	After	Total
C	X	X	X	X	—	X	8

Figure 4.23 Type C conversations

The third most common conversational type, shown in Figure 4.23 above, is Type C, in which one participant reaches all three phases, and the other skips the During phase. This

conversational type occurs in 8 out of the 54 conversations, or roughly 15%, and is exemplified below with the conversation between Jonny and Joshua (13).

	Jonny	Joshua
1	hey!	
2		Heyy
3		You alright? x
4	I'm not too bad thankyou, how're you x	
5		I'm alright thanks x
6		Where you from? x
7	good and I'm from XXX, XXX XXX. yourself? x	
8		Ohh cool, I'm from London but I'm in XXX for Uni x
9		Do you go to Uni? 😊
10	aw that's cool I'm doing my degree but at a college because it's cheaper :) what dya study x	
11		Ohh that's smart 😊, mechanical engineering x
12		You?
13	Jesus! I struggle to make a pop up tent:(hospitality and business management x	
14		😂😂😂 that's funny
15		Oh cool, do you like it? x
16	nope I hate it :(x	
17		Ahh why? :(x
18	it's just shit x	
19		Lool that's peak
20		Have you got whatsapp ?
21	yeah I do, 07XXX	
22		Is that your number ? x
23		It's someone else as your profile pic 🤖🤖
24	yeah it's my number?	
25		Pkay
26		Didn't wanna whatsapp the wrong number 😊
27	what's your nhmver	
28	number	
29		07XXX
30		Just whatsapp'd you
31	you're a guy?	
32		Yh lol
33		Aren't you a girl? :/
34		Can't you see my pic ? 😊
35	no:L	
36	yeah it's a lass 🤖	
37		:///
38	it's a mass called dolly	
39	lass	
40		You're acting as a girl and saying I'm a guy :///
41		I'm not into guys
42		Sorry
43	my profiles a guys profile... you've got the girls profile	
44		But you swipe right on guys? :/
45		Makes no sense
46		Don't think so lol
47	you're not a guy on mine!	
48		Sure
49		But I have guy pictures etc
50		And you're into girls..

Figure 4.24 Jonny and Joshua (13)

In the conversation above, Joshua transitions between all three phases similarly to the transitions we saw in the Type A examples above. It is only when both conversants attempt to transition off Tinder and onto WhatsApp—a messaging application on which both conversees have genuine profiles—that Joshua transitions to the During phase. In 22 he questions whether the number he was given was correct, because, as he states in 23, “It’s someone else as your profile pic”—that someone else being, of course, Jonny himself. Rather than taking this information in immediately, Joshua continues to negotiate throughout his During phase to reconcile this new information with the ground truth provided by the Tinder profile.

Although Jonny gets the exact same information through WhatsApp—a male name and picture—he has no During phase of negotiating the information, simply stating, “you’re a guy?” in 31 after seeing Joshua and not “Dolly” on WhatsApp.

	Tam	Sol
42	Haha looks like I win you must be a bright young lady to be studying math	
43	So shall I give myself away or do you want to carry on guessing	
44		Tam
45		‘Bright young lady’...
46		I quite sure that’s definitely something I am not haha
47	Haha really what makes you say that ?	
48		You called me a bright young lady haha
49	Yeah I know what I meant was what makes you feel that your not bright	
50		Because I’m not a lady tam
51		Haha I’ve just double checked
52		And still definitely not
53	Haha okay “girl” is that better do you feel much younger when I say that :p	
54		I’m a guy hahahaha
55	And the devastating truth is revealed haha	
56		Tam 🤔
57		This is hilarious, which girl did you think i was? Want me to hook you up hahah
58	I’m getting trolled - _ - aren’t I	
59		Can you see my pictures?
60	Well yeah of course	
61		Do I look extra feminine in them or something haha
62	Haha extra feminine nah you don’t	
63	You look pretty normal and cute	
64		Why the gender confusion then tam ??
65	Oh no when I said the truth has been revealed I was joking lol	
66	I know your a woman	
67		Is that right
68		Like a reverse lady boy
69	Haha yeah like a “reverse lady boy”	
70		Haha thnk you might be slightly too confusing for me tam

	Tam	Sol
71		With the wrong name calling me a woman...
72	Yeah I'm so confused as well lol you know I'm a guy right	
73	In you pictures all I see is girls so I'm assuming your not a guy but I'm so tatty confused	
74		And the devastating truth is revealed haha

Figure 4.25 Tam and Sol (7) excerpt

Other Type C conversations have similar transitions for one conversee, where the new information they have been given is so strikingly dissimilar from the Tinder profile, or is otherwise backed up by convincing enough evidence (in this case their genuine WhatsApp profiles), that there is no room for negotiation. In other cases, such as in the conversation between Tam and Sol (7), excerpted above, this transition from Before to After occurs not only because the information they are given doesn't match the profile, but because of the conversee who transitions to the During phase.

Although Tam's statement in 42 of Sol being a "bright young lady" is what sets off Tam's eventual During phase, and confusion for Sol, Sol is not confused about Tam's gender, but about what Tam thinks of his own gender. Although Sol states in 50 "I'm not a lady", and in 54 "I'm a guy", Tam remains in the During phase, confused as to whether he's being trolled, but still believing Sol is female. It is within this During phase for Tam, and cooccurring Before phase for Sol that, in 72, Tam asks "you know I'm a guy right".

Because Tam has begun to question whether Sol is female, he decides preemptively to clear up any confusion as to his own gender. Thus, Sol skips from the Before to the After phase. It is only later in the conversation that Tam, too, at last transitions to the after phase himself.

4. Type D – Tom and James (38)

	Before	During	After	Before	During	After	Total
D	X	—	—	X	—	—	5

Figure 4.26 Type D conversations

In the next most common conversation type, Type D, shown above, neither participant transitions past the Before phase, as with the conversation between Tom and James below. This conversational type accounts for 5 of the 54 conversations, or roughly 9%.

	Tom	James
1	Are you real?	
2		No I don't actually exist. I'm just a 60 year old man pretending to be a 21 year old called James in the hope of picking up girls.
3	Hahah well that was what i was hoping for to be honest. I mean who wants to talk to attractive girl when you can talk to old men 😊	
4		Exactly my thought process as well haha. Take it you're real then? 🙄
5	Unfortunately yes, so I'm guessing this isn't James then	
6		Well I was kind of joking about the 60 year old man thing, so unfortunately, yes I am 21 year old James 😊
7	Well since you're another guy I guess we should talk	

	Tom	James
	about manly like sports and cars then 😊	
8	Or we could talk about how good th body of the girl in the picture is 😊	
9		Haha sports and cars are good, but the body is better 😊
10	Definitely! Yeah she looks like shed be a great fuck 😊	
11		Haha she does, but I wouldn't know, you'd have to ask someone else about that 🙄
12	It's a real she I Can't talk to her, I'd love to tell her the things we could do together	
13	Because that's what men talk about 😊	
14	I wonder if she's flexible omg	

Figure 4.27 Tom and James (38)

Because both Tom and James are attempting to negotiate their partner's messages with their Tinder identity, both take any errant comments as jokes, a tactic common to such negotiations. Because of this aspect of joking, neither Tom nor James has any concrete reason to question the other's identity. Even though James states in 2 "No I don't actually exist. I'm just a 60 year old man pretending to be a 21 year old called James," Tom has no reason to believe the James identity is any more or less of a joke than the 60-year-old man identity.

Eventually, the conversation ends in 14, though James stops responding in 11, the deflection and paired emoji (🙄) an indication of the oddness of their conversation. James never responds again, and Tom does not press. Although this is a common progression for these types of conversations (in that any oddities in the conversation are ascribed to joking as a negotiation tactic), it may be the case, as we see in the conversation excerpt between Nat and Sixten (17) below, that the latter phases eventually do occur off Tinder.

	Nat	Sixten
50	Dolly i think you should give me your number so you can take me XXX sometimes :-)	
51		I don't have an English phone number though, and it is quite expensive to write back and forth to a XXXish number. Do you use WhatsApp or something like that?
52	Oh fairs yh i got whatsapp 07XXX x	

Figure 4.28 Nat and Sixten (17) excerpt

After a lengthy conversation, the interaction finally ends in 52 after Nat and Sixten exchange WhatsApp information. Based on this, it is unclear whether they kept talking on WhatsApp and eventually transitioned to either or both latter phases, or if, as we have seen in many other conversations, their interaction ended as soon as they saw one another's WhatsApp profiles.

5. Type E – Jesse and Zäimon (48)

	Before	During	After	Before	During	After	Total
E	X	X	X	X	X	—	5

Figure 4.29 Type E conversations

Shown in Figure 4.29 above, Type E conversations are relatively similar to Type B conversations, in that while one conversant makes it through all three phases of the

conversation, their conversee does not. In this case, however, the conversee does transition into the During phase, but not the After phase, as shown below in the conversation between Jesse and Zäimon. Like Type D conversations, Type E accounts for 5 out of the 54 conversations, or roughly 9%.

	Jesse	Zäimon
1	You look like you bite	
2		Ooh yeah! Im dangerus! ;)
3		Or maybe the opposite :P
4		Hi
5	With lios like those I don't want you using too much teeth	
6		Haha promise ;)
7		Are u from XXX?
8	Maybe a little bit ☹️	
9	Yeah	
10		Nice it's a sick city ;) fking love it :P
11	Well maybe you could stay the night	
12		Ofc ;)
13		Wouldnt complain if i did :P
14	How kinky are you	
15	Like what you're into 😊	
16		Hahaha not sure if im worse than u now actually ;)
17	I hope you like anal	
18		Hahahaha
19		Ur dirty as fk ;)
20	I'm very dirty	
21	You wanna text me	
22		Allways ;)
23	XXX-XXX-XXX	
24	I'll be there	
25		On the phone? ;)
26	Just text me	
27		Did, didnt get it?
28	I have a surprise when you do	
29		;)
30		U being a guy or what? Hahah
31	No what's your number	
32		07XXX
33	Lmao just wait you'll see	
34		Key ;)
35	Get it?	
36	Hun	
37		Naa :/ but try *46XXX
38		+46*
39	Where are you from	
40		Sweden
41	What about now	
42		Naa but got fb? The messenger app?
43	Yeah what's your name	
44		XXX XXX
45		Says zäimon XXX XXX when u serch it
46		Well im gonna go out and eat, but ill send u a text when i find some internet ;) just add me on fb
47	?	

	Jesse	Zäimon
48	Hun	
49	You freak you're a guy	
50		No shit
51		Thats why my profile says so....
52		And ur says that ur a girl but im not rly sure bout' that

Figure 4.30 Jesse and Zäimon (48)

In the conversation above, both participants attempt to transition off Tinder, as many conversations do, but run into trouble because of Zäimon's international phone number. Because they decide to try Facebook instead, Zäimon gives his full name in 45, which tips Jesse off to Zäimon's gender, either because of his name, or because of his Facebook page. Because of this, Jesse transitions—again quickly for the During phase—between all three phases. At the end, in 52, Zäimon himself expresses suspicion that Jesse may also be male, though whether this is because of the mix-up with his own profile, or some other, unexpressed reason (such as Jesse's language or their interaction thus far), is unclear. As with the Type C conversations, the conversation ends before Zäimon can fully transition, and fully realize his suspicions as true.

In other cases, such as in the conversation between Gerard and Sam (27), excerpted below, this final transition is complicated in one or both participants' During phases.

	Gerard	Sam
211		Please confirm that you are a woman !
212	No I'm definitely male	
213	Sorry to disappoint	
214	How old are you?	
215		You're lying to me! Please stop it!
216		And I'm 20 you?
217	I'm also 20	
218		My timer says you're 23
219	07XXX, check my whatsapp if you have it, that's me 👍	
220		And I'm a female by the way of you're a male
221		Just to clarify
222	Well at least I haven't been talking to a bloke	
223	Silver lining and all that	
224		Ok and yeah that is good
225		You don't have a profile picture
226		And how long are you going to lie to me and say that you're male
227	You're joking right 😏	
228		No
229	Definitely being catfished or something here 🙄😏	
230		I fucking am!
231		If you're make this conversation ends now
232		Male*
233	Gerard XXX, Facebook me 👍	
234		So you'd better own up or I'm deleting the chat 😏
235	What's your second name? 😏	
236		No this is over I'm not gay
237	😏😏	
238		👍

	Gerard	Sam
239	Enjoy the rest of your day 😊	
240		I don't know whether you're lying to me or that you are an actual bloke?
241	Do you have snapchat?	
242		Yes
243		XXX
244		Prove to me you're not a bloke
245	Well I am a bloke 😊	
246		Well we can't be talking
247	Not sure now I don't want some randomer having my SC	
248	Link me your Facebook 😊	
249		Sam XXX
250	No I mean actual link 😊	
251	I'll be there all day	
252		I've only got the app you will have to search for me
253	Right what's your twitter 😊	
254		@XXX
255	And you're female yeah? 😊	
256		No I'm not I'm having you on!
257	😊	
258		I'm waiting for you to say that you're not male
259		So you're make
260		Male *
261		if you say yes I'm unmatching you
262		Cya

Figure 4.31 Gerard and Sam (27) excerpt

Prior to 211, both Gerard and Sam have been going back and forth with regards to their true genders, with each instance of either participant expressing that they are male being taken as a joke, as we see again in 211-215. Eventually, rather than negotiating this discrepancy with simple joking, Sam transitions to trolling in 220 by stating, “And I’m a female by the way if you’re male”.

Such a tactic may be considered reverse-trolling by Sam, as he perceives Gerard to be continually messing with him by claiming he is male but having a female profile—thus leading to the suspicion that Gerard is trolling Sam, either because he’s a female joking about being male, or a male using a female profile. Sam appears to want to give Gerard a chance to confess by stating he is female, or simply troll Gerard back. This trolling, however, further convinces Gerard that Sam *has* been joking this entire time about being male, and is actually female, as his Tinder profile states.

Gerard finally transitions in 257 after Sam confesses “No I’m not I’m having you on!” in 256, and then does not respond again (indicating this is indeed an After phase for Gerard). Sam, on the other hand, keeps giving Gerard opportunities to admit that he has actually been female this whole time—but Gerard has already stopped responding. As with the conversation between Jesse and Zäimon, the conversation ends before the participant in the During phase can confirm their suspicions, but this happens not because of negotiation via joking, but because of trolling.

6. Type F – Alassane and Jamie (30)

	Before	During	After	Before	During	After	Total
F	X	X	—	X	—	—	4

Figure 4.32 Type F conversations

Type F conversations, as shown in Figure 4.32 above, are similar to Type D conversations, in which neither participant reaches the After phase, but differ in that one does eventually reach the During phase. This is shown in the conversation below between Alassane and Jamie. Type F conversations account for 4 out of the 54 conversations, or roughly 7%.

	Alassane	Jamie
1	Oh hello 😊	
2	How u doing	
3		You Alrite trouble :) , yeah good thanks you ?
4	U look like u cause more trouble then me 😊 Yeh im good what u up to ?	
5		Na never that halo above my head :) , good good and just went had me beared trimmed at the barbers now jus popped to see my mate at work . What you doin ?
6	Wdf?!! 😊😊	
7	You man or woman	
8		I beg your pardon haha
9	Are u a man or a woman talking about beard 😊😊😊😊	
10		What you mean haha my beard had me looking like
11		A alley tramp
12		It's still there jus trimmed and shaped I can't stand my beard to long
13	Lmfao	
14	😊😊😊	
15	I wonder what kinda beard thats	
16		What's ya number il txt ya a picture
17	How can i trust u 😊😊😊	
18	U got woman pics and say u got beard 🤔🤔	
19		Hahaha you could sell out the O2 with jokes like that
20	I should tell u that	
21		:)
22	So u telling me u aint the person on the pix ?	
23		What you mean corse I am
24		You ain't a catfish are ya lol
25	Lol do i look like one to u ?	
26		Na lool I'm only messin
27	I need to call Nev for u	
28	Til i see u live i wont believe its u lol	
29		Lool you joker , see ya when I see ya then :) , I ain't watched that programme in ages last one I watched this girl thought she was chatting to lil bow wow
30	Lool yeh remember that still	
31	Whats ure facebook name	
32		Jamie XXX , you ?
33		What's your last name sheen ?
34		Where you from anyways ?

Figure 4.33 Alassane and Jamie (30)

In the case of the conversation above, the transition to During begins for Alassane

because of Jamie's mention in 12 of having a beard. While this is initially taken as a joke by Alassane, once Jamie offers to send a picture of it in 16, Jamie transitions to the During phase out of suspicion. Alassane even states in 27, "I need to call Nev for u", a reference to the MTV show Catfish, in which people who want to believe the identity of an online romantic partner seek out the help of the hosts (including Nev) to investigate if they are who they say they are.

As with the Type D conversation between Nat and Sixten, it is unclear whether Alassane transitions to the After phase off of Tinder, as he asks for and is given Jamie's full name to look up on Facebook. For Jamie, on the other hand, the conversation simply ends before he has reason to be suspicious of Alassane (other than jokingly in 26 and 29). Although this is the same kind of end and potential off-Tinder transition as found in the Type D conversations, possible off-Tinder information is not the only reasoning behind Type F conversations ending in the way they do, as we see below in the conversation between Will(A) and Will(B) (24).

	Will(A)	Will(B)
1	If I offered you a pizza with pineapple would you: a:devour it B:pick off the pineapple and carry on C:flip it off the table and ask what sort of heathen puts pineapple on a pizza	
2		Can I not give you d as my answer?
3	Well if you want the d...	
4		Lubricate yourself. Be ready in 15 minutes.
5	I need to be lubricated?!? No deal	
6		Maybe I'll excite you enough tat you won't need it ;) are you free tonight? It's pretty hard for me to fit stuff in ;)
7	If only I was... I'm in XXX till fairly late	
8		I can wait up.
9	🐱are you a bloke called Steve that's going to drag me away in his van	
10		No van, but something else you can ride.
11	A lot of your references sound like I'm getting penetrated	
12	Which....isn't ideal	
13		You make it sound like you've already decided ;)
14	I've decided nothing's going in me	
15	That's for sure if	
16		Well, have you got a strap on?
17	What on earth would I need one for	
18		Well, if you don't want to be penetrated, I might as well take the opportunity?
19		Oh well, I have my fleshlight.
20		Sorry that was my Mum...
21	That was your mum	
22	Seriously	
23	On reading this back, if it genuinely was your mum then she's a bit of a legend	
24		Sorry I got a bit excited last night, got carried away by the unbelievable chat. Had a good day today?
25	Well now I don't know who I'm talking to	
26		Billy big-cock right here
27		Sorry was a friend. Want to meet up tonight? Xxx
28		Spoons?

Figure 4.34 Will(A) and Will(B) (24)

The conversation above between the two Wills quickly becomes sexual, lending itself to numerous innuendos, and eventual confusion. While Will(B) does not appear to have any cause to question the identity of Will(A), Will(A) expresses both confusion and suspicion multiple times, before terminating his side of the conversation at 25 with the line, “Well now I don’t know who I’m talking to”. Again, the conversation ends before Will(B) can transition past the Before phase, but not because of any outside information.

7. Type G – James and Andrew (26)

	Before	During	After	Before	During	After	Total
G	X	—	X	X	—	—	4

Figure 4.35 Type G conversations

Similarly to Type C above, one conversee in Type H conversations makes the jump from Before to After with no During phase, as is shown in Figure 4.35 above. However, as demonstrated in the conversation below between James and Andrew, the other conversee never leaves the Before phase. As with Type F conversations, Type H conversations account for 4 out of the 54 total conversations, or roughly 7%.

	James	Andrew
1	Good evening how are you today	
2		Good.....
3		U?
4	Im good thank you what you looking for on here	
5		Where ru from?
6	Im from XXX XXX where are you from	
7		XXX. U know it?
8	Yeah sort of	
9		So ru studying here or something?
10	I work in XXX XXX in XXX what about you	
11		Barrister
12	Wow im an in house porter	
13		What's that
14	I move desks build desks and move cabinets and stuff like that	
15		Like handyman?
16	Yeah like that	
17		U born in uk?
18	Yeah was you	
19		Yeah. Where's ur pics taken then. Doesn't look like uk
20	One was in my room and the other was in the living room	
21		Have u even looked at ur pics. They are outside
22	They aint	
23		The first one is u walking down a road...?!
24	What you looking for on here	
25		Fun. U
26	The same and see what happens	
27	The first one is of me in a white shirt in the kitchen	
28		The other two?
29	One is with my dog in my dads and the other one is me one the sofa	
30	What kind of fun you after	
31		What do u thin
32	Sex lol	

	James	Andrew
33		Course
34	Sounds good to me ;)	
35	Your very beautiful	
36		Oh yeah? What do u like
37	I like going out for a dew drinks something to eat walks what about you	
38		U have Facebook
39	No I have what's app I closed my Facebook account down	
40		What's ur number
41	My number is 07XXX	
42		So when do u try and scam me
43	I dont thats my phone number	
44	Im real not fake	
45		U might want to updated your profile photo
46		It's of a man
47	I am a man	
48		U gay
49	No I aint im straight as anything whats your number I'll whats app you	
50		Sorry mate u don't make any sense
51		Why are u pretending to be a girl
52	I ain't pretending to be a girl im a man im not gay im straight	
53	How am I pretending to be a girl	
54	I can prove I am a man and 100% real	
55	Are you a man	
56	What positions do you like doing	
57	What you up to	
58	Hi	
59	Hi	
60	You still up for fun	

Figure 4.36 James and Andrew (26)

In the conversation above, as with many others, Andrew transitions to the After phase immediately upon obtaining James' WhatsApp information, and stops responding at 51. (Again, though not explicitly expressed, the foray into WhatsApp may itself count as James' During phase.) Although James takes 9 more conversational turns, attempting to continue with the conversation, he does not appear to ever really question whether Andrew is female, and never gets any further information.

The conversation in 20, which follows the same progression, does so for a very different reason than any other conversation in the set of 54. The interaction below, between Deji and Joe (20), is the only conversation in which the participants make plans to meet in person and appear to attempt to follow through.

	Deji	Joe
97	Hi dolly.	
98	Is tonight still on?	
99		Hello
100		Yeah all good
101	How's it going?	
102	Nice. So, what's the crack for tonight?	
103	XXX for 8?	

	Deji	Joe
104		Good thanks
105		Been doing cw all day
106		Yeah XXX sounds good
107		Maybe half 8 but depends how long it takes to have food once im back I will let you know
108	No worries. Drop me a message whatever the crack is.	
109		Hey hey
110		Sorry is 9 ok
111	9's good. Don't be sorry.	
112	I live 2 seconds from XXX, so it's cool	
113		Oh cool. Im on my way now will message you when im 5 mins away x
114	Cool	
115		Just waiting for a bus.
116	Cool 👍	
117		Be there in 10
118	See you soon	
119		Im outside
120	Outside XXX?	
121	What's your number? I'll text you when I'm outside. Heading down now.	
122	I've got no data so when I leave the building, my wifi goes off and I can't contact you on here	
123		Im inside now downstairs
124	Ok	
125		07XXX
126		07XXX
127		Hilarious

Figure 4.37 Deji and Joe (20) excerpt

As with many conversations in which the Before phases last for a long time for both participants, Deji and Joe manage to inadvertently avoid overly gender-coded topics, or otherwise negotiate them. Such negotiations include the name “Dolly”, which Deji uses in 97 above. Although an odd pet name for a male, by this point, Joe has apparently negotiated it as acceptable, or at least unavoidable as is used by Deji. Because neither has had any real reason to suspect the other’s Tinder identity is not genuine, they make plans to meet, which they begin to discuss on the day of in 97. Because of Joe’s connectivity issues, they have not been able to transition to text messaging prior.

Both remain in the Before phase, and Joe finally provides his phone number in 125-126. It is unclear exactly what transpires after the two find themselves physically in the same location, but Joe, at least, makes one final comment in 127, “Hilarious”, an indication that he has transitioned to the After phase upon their eventual meeting, or at least based on their WhatsApp profiles. Although not verbally indicated on Tinder, it is likely Deji reaches the After phase himself off Tinder, as he offers no more confused responses attempting to negotiate “Dolly” not showing up.

8. Type H – Steven and Corey (15)

	Before	During	After	Before	During	After	Total
H	X	—	X	X	—	X	3

Figure 4.38 Type H conversations

Unlike the above conversational Types, Type H conversations, as shown above, skip the During phase entirely, and both participants transition from Before immediately to After, as shown in the conversation between Steven and Corey below. Type H conversations occur in 3 out of the 54 total conversations, accounting for roughly 6%.

	Steven	Corey
1	Evening	
2		Hey there!
3	How are u sybil	
4	Don't mind me saying you look quite nice in your photo 😊	
5		Thanks! And the same to you 😊
6	Your welcome ;) so anyway what u up to	
7		Just in bed, peeving on your pictures 😊
8		You?
9	Me just laying in bed and really what you think 😊	
10		That I'd very much like to see you naked and wet 😊
11	Really now well maybe I am thinking the same thing about u 😊	
12		We should do something about this then 😊
13		You want my number?
14	Okay teh	
15	So where this number at 🐼🐼	
16		07XXX XXX
17		😊
18		I want to see your wet naked body!
19	I show u my naked body but not wet 😊	
20		I'll make you wet then ;)
21	Okay 😊	
22		So you're a guy?
23	Yes and why your pic on this a women	
24		Your pic is a woman,
25		Mine's a guy!? Haha
26	No here now your pic is a Chinese women 😊	
27	So tinder is fucked uo	
28	Up	

Figure 4.39 Steven and Corey (15)

As we see above, as with many other conversations, the immediate transition from Before to After occurs when the participants exchange off-Tinder numbers, both realizing the other is male through WhatsApp, in 22 for Corey, and in 23 for Steven. Although Steven does negotiate this discrepancy, it is in the After phase, where he decides that the reason their profiles were inaccurate must be because of a glitch through Tinder. (Arguably, both engaging via WhatsApp was their During phase, though it is not explicitly expressed here.)

In the case of David and Adam (37), excerpted below, the conversation similarly transitions at once, but because of on-Tinder information provided by David.

	David	Adam
44	Can I tell you something without you un matching me straight away?	
45		Sure lol
46	Please don't be scared off! I'm a Daddy	

47	So you're a guy...
48	Er yeah! I thought you were a girl
49	Better luck next time then

Figure 4.40 David and Adam (37) excerpt

Because the conversation has been going well prior to 44, David decides to make it known that he has a child—but as the word “Daddy”, as is used in 46, is a male-coded term, Adam immediately understands in 47 that David is male. Based on Adam’s response, David’s own transition is immediate in 48 as well, and both participants transition from Before straight to After.

9. Type I – Adam and Mark (52)

	Before	During	After	Before	During	After	Total
I	X	—	—	—	X	X	1

Figure 4.41 Type I conversations

Type I conversations perhaps differ the most from any other type, in that the participants experience inverse phases. As shown in the conversation between Adam and Mark below, one participant remains only in the Before phase, while the other participant begins at the During phase, skipping the Before phase entirely, before transitioning to the After phase. This occurs only once out of the 54 conversations, accounting for roughly 2%.

	Adam	Mark
1	Hey Sadie! How goes it? Here's my number...XXX XXX XXX...would love to get to know you... :-)	
2		Slut
3	Umm not looking for just a hookup so don't rush to judgement	
4	Thnx	
5		I know but its still some sort of scam
6	Not really.. I am Just a guy looking to chat and see where things go from here... 100% not a scam	
7		Ur a dude posing as a woman?
8		Nuff said
9	Um no.....	
10	Sounds like you have been catfished quite a bit... I am actually a pretty genuine guy just looking for something real and not just for a flnig or to try and scam someone like you are assuming.	
11		But ur a dude using women's pics as a way to meet dudes.

Figure 4.42 Adam and Mark (52)

Because Adam begins the conversation with a very forward offer of transitioning to text messaging (and possibly because of the use of the name “Sadie” while addressing Mark), Mark begins the conversation immediately suspicious, and thus in the During phase. Upon Adam’s insistence that he is “Just a guy looking to chat” in 6, Mark’s suspicions are confirmed, and he transitions from During to After. The conversation ends before Adam can transition to any stage past the Before stage. It may be that, as we will see in Part 5 below, Adam’s introduction follows too closely with the sort of introduction Tinder bots are notorious for, thus prompting Mark’s immediate suspicion.

10. Type J – Armand and Kasim (21)

	Before	During	After	Before	During	After	Total
K	X	X	—	X	X	—	1

Figure 4.43 Type J conversations

Type J conversations are those in which both conversants transition only from Before to During, and not After. This is shown below in the conversation between Armand and Kasim, and as with Type I conversations, occurs only once in the 54 total interactions, accounting for roughly 2%.

	Armand	Kasim
1	Dolly!! X	
2		Dolly? 🤖
3	Just trying to get ur attention 🤖🤖🤖 x	
4	Where u from 🤖🤖 x	
5		From XXX. How about you?
6	Not far. From XXX xx	
7	Youre pretty hott btw 😊😊 x	
8	How tall are you ? X	
9		Haha thanks.. XXX? It says you're only five miles from me.
10	6ft 3, you?	
11	It says 45 kilometres on mine x	
12	U like the tallest person i know 🤖	
13	Im a bit shorter	
14	🤖🤖🤖🤖🤖🤖	
15	🤖🤖🤖🤖🤖	
16		Most people say that. Haha.
17		A bit? 😊
18	Haahaha i bet u the tallest in the bunch ? X	
19	Yea not tellinf 😊😊😊	
20		Oh go on, tell me.
21		Yeah I am. Giraffe here.
22		What do you do?
23	Hahaha if i tell we still be talking right 😊😊	
24	I like giraffes got like stuffed toys of them when i was little 😊😊😊	
25		Yes, we'll still be talking. Get on with it. 😊
26	I go to XXX doing sports studies xx	
27	Yourself ? Xx	
28		Where was the picture with the binoculars taken?
29		I work in XXX, in technology.
30		Studied History
31	5inches short 😊😊😊	
32	What binoculars. I dont think that's my profile 😊😊	
33	Oh so u finished ? Thats impressive	
34	What do u work as xx	
35		What? In your second picture! With the water in the background.
36	Dont think thats me 😊😊😊	
37	My second pic is just my face😊😊	
38		Send me a flipping pic of yourself. I'm confused.
39	Hahaha u cant do that in here x	
40		Oh. One of your pics is you with some giraffe thing?

	Armand	Kasim
41	No i dont think so	
42	I just got 2 pics a selfie and a selfie zoomed in 😂😂😂	
43		Wtf is going on here. 😂😂😂😂
44		This'll help. What race are you?
45	Wait lemme chexk urs	
46	U got a pic in a beach	
47	With ur hand on ur head	
48	Definintely not white	
49		Hahaha. No. You're winding me up.
50	U got 4 pics ? X	
51	U got a bit curly hair ? X brunette or black	
52	And one of ur pics u got headband on ?	
53	This is craazyy 😂😂😂	
54		Uh. I have a hat on in one pic.
55		Black hair. Pretty straight hair.
56		Hahaha
57		You look far eastern Asian....
58	One of ur pic is brown hair ?	
59	With a hat ?	
60	And like pink lipstixk ? X	
61		No 😂😂😂😂😂😂😂😂😂😂
62		Definitely! That's me.
63	Which one xx	
64	Yea east asia filipino to be exact	
65		Well yea. Then your pictures are right.
66		I definitely am not wearing lipstick. Stop winding me up. 🙄🙄🙄🙄
67	I dont have pic with binoculars and giraffe tho 😂😂	
68	Im not gahahahah its just the pic that turns up z	
69		Are you wearing big glasses in your selfie?
70	Pink lipstick and a brown hair with like a black thing on your head ?	
71		😂😂😂😂😂
72	No theyre not me 😂😂😂	
73		What are you smoking?
74	Are u winding me up 😂😂😂	
75	Seriously this never happened before 😂😂😂	
76		No😂
77	Wait is this dolly ?	
78		Who the hell is dolly
79		Lol
80	Hahahaha	
81	Youre funny 😂😂😂	
82		My name is Kasim
83	Stop winding me up 😂😂😂	
84		I give up. 🙄
85	No its not 😂😂😂	
86	Its dolly	
87	My first message was dolly coz of ur name 😂😂	
88	U wanna do something crazy 😂	
89	Shall we meet up to find out 😂😂	
90		Erm. Nah.

	Armand	Kasim
91	Why not	
92		You mean why would I.
93	So u would ? X	
94		No 🙄
95		You're odd.
96	I meant meet up 🙄🙄🙄	
97		I know what you meant
98	🙄🙄🙄	
99	🙄🙄🙄	
100	Dolly!	

Figure 4.44 Armand and Kasim (21)

As with many previous conversation types, Armand and Kasim are able to continue on within the Before phase for a number of turns, first by discussing the relatively neutral topic of their location. This eventually leads into discussing issues with their locations according to Tinder, and issues with the pictures they are seeing of one another. This takes some time to transition to the During phase because of negotiation via joking, such as in the exchange between 58 and 62, where Armand asks if one of Kasim's pictures is of him wearing pink lipstick. Kasim's joking response of "Definitely, that's me!" in 62 perpetuates the negotiation. By matter of sheer coincidence, Armand's race, "east asia filipino" as he states in 64, matches the fake Tinder profile, further confusing the negotiations.

Once both reconcile that their pictures appear to be wrong, and the other is not joking to that effect, Armand asks in 77 "Wait is this dolly?"—which it, of course, is not. Kasim "gives up" in 84 trying to figure out who Armand actually is if not the provided Tinder profile, and finally rejects the suggestion they meet up to find out who the other person actually is in 90. Kasim does not respond past 97, despite Armand's apparent continued interest until 100, and the conversation ends before either one gets any definitive answer as to who they have been chatting with.

11. Type K – Kane and Jitesh (23)

	Before	During	After	Before	During	After	Total
K	X	X	—	X	—	X	1

Figure 4.45 Type K conversations

Finally, Type K conversations are those in which one participant does not reach the During phase, and the other participant does not reach the After phase. As with Types I and J, Type K occurs only once in 54 interactions, accounting for roughly 2%. However, the context of the conversation between Kane and Jitesh may indicate that this Type K instance is better defined as a permutation of Type C, a much more common conversational Type. This is because, as we can see in the excerpt below, Kane *does* reach the conclusion that Jitesh's biological gender is not that of a female, even if "her" gender presentation is.

	Kane	Jitesh
11		Do you live on a chicken farm?
12	No why? 🤔	
13		Because you sure know how to raise a cock 😏😏
14	Hahaha	

	Kane	Jitesh
15	Are you implying you have a cock? 🤔🤔	
16		😏 i do have a cock
17	Ohhh 🤔🤔	
18		Are you an archaeologist?
19	No why?	
20		Because i got a bone for you to examine ;) 😏
21	Haha	
22	Im sorry but if you do actually have a cock	
23	I aint about that life	
24		Haha banter !
25	Ffs 🤔🤔	
26		I do actually have a cock but tinder is purely banter
27	Oh fair enough	
28		Haha why so shy 😏
29	Like i said	
30	I aint about that life	
31		About the cock life?
32		What life are you about :p
33	Pussssy	
34		Pussy is the best life
35	Yeah it is	

Figure 4.46 Kane and Jitesh (23) excerpt

Kane *does* have an After phase of sorts, but it is one in which he still comes to the incorrect conclusion about Jitesh's identity. This conversation continues for 29 more turns, during the entirety of which Kane continues to accept Jitesh's Tinder profile as accurate, but for that of a trans female. This conversation in particular is emblematic of conversees continuing to give primacy to the information presented to them in their interlocutor's Tinder profile, and misunderstanding their interlocutor's language and the information it provides in the context of the profile, enabling them to bypass the red flags discussed later in Section C.

B. Conversational Trajectories

Although the overall conversational Types involve 10 different overall makeups when considered by speaker pairs, for individual speakers, there are 5 total trajectories found per conversant (with two possibilities found *not* to occur).

Before	During	After	Before	During	After	Total
—	X	—	X	X	X	47
—	—	X	X	—	X	19
Trajectories not found			X	—	—	29
			X	X	—	12
			—	X	X	1
Total who reach this phase			107	61	66	108
Total who do not reach this phase			1	47	42	

Figure 4.47 Conversational trajectories

Although the two possible trajectories that do not occur would be unlikely to occur in this dataset, as we will see in Part 5 below, there are indeed conversational types in which trajectories without a Before type phase are more common (such as in online child porn communities). That there is only one conversee in the Tinder dataset who skips the Before phase is likely indicative of the necessity for low suspicion on the part of Tinder users in order

for their conversations to stand a chance of becoming real-world encounters. Exceptions to this are discussed further in Part 5 below. Here we consider why, in general, conversees may miss either or both of the During and After phases.

1. No During phase

As shown in Figure 4.48 below, of the 10 Types of conversations, 7 occur wherein at least one participant misses out on the During phase of the conversation, totaling 36 out of a total 54 conversations. Of these, in only 3 of the types do *both* participants miss out on the During phase (Types D, G, and H), totaling 12 out of 54 conversations.

	Before	During	After	Before	During	After	Total
B	X	X	X	X	—	X	11
C	X	X	X	X	—	—	9
D	X	—	—	X	—	—	6
F	X	X	—	X	—	—	3
G	X	—	X	X	—	X	3
H	X	—	X	X	—	—	3
I	X	—	—	—	X	X	1

Figure 4.48 Conversational types without During phases

Although precise instances of the flags that most often result in transitions through conversational phases are discussed in Section C below, of interest here are the trajectories in which the During phase is skipped for direct transition to the After phase, and trajectories in which neither latter phase is reached, although Type D conversations are discussed further in subsection 2 below.

In 3 of the conversational Types (B, G, and H), at least one conversee skips the During phase but does reach the After phase, for a total of 17 conversations and 20 conversees. The most exemplifying Type of conversation is, then, Type G, wherein both participants make the skip. As seen in two of these conversations, shown below, this transition occurs because of unavoidably gender-specific real-world information.

	David	Adam
43	Can I tell you something without you un matching me straight away?	
44		Sure lol
45	Please don't be scared off! I'm a Daddy	
46		So you're a guy...
47	Er yeah! I thought you were a girl	
48		Better luck next time then

Figure 4.49 David and Adam (37) excerpt

	Steven	Corey
15	So where this number at 🐼🐼	
16		07XXX XXX
17		🐼
18		I want to see your wet naked body!
19	I show u my naked body but not wet 😊	
20		I'll make you wet then ;)
21	Okay 😊	
22		So your a guy?
23	Yes and why your pic on this a women	
24		Your pic is a woman,
25		Mines a guy!? Haha

	Steven	Corey
26	No here now your pic is a Chinese women 😊	
27	So tinder is fucked uo	
28	Up	

Figure 4.50 Steven and Corey (15) excerpt

No real negotiation (which is typical of the During phase) occurs in these instances, as the admission of being a “Daddy” by David and exposure to Corey’s real-world WhatsApp profile provide inarguable evidence that both are male. That David immediately understands Adam to be a guy (Corey likely also sees Steven’s real-world profile when they connect via WhatsApp) may be down to the cooperation typical of these types of conversations.

Although not all such red flags are so immediately ratified by participants, as is discussed further in Section C below, the lack of a During phase seems to be frequently triggered by them. As such, as we will see in subsection 2 below, the lack of the latter two phases by either participant is not necessarily indicative of successful identity play (however inadvertently) by their interlocutor.

2. No After phase

As shown in Figure 4.51 below, of the 10 Types of conversations, 7 occur wherein at least one participant misses out on the After phase of the conversation, totaling 31 out of a total 54 conversations. Of these, in only three of them (Types D, F, and J) do both participants miss out on the After phase, totaling 10 out of 54 conversations.

	Before	During	After	Before	During	After	Total
C	X	X	X	X	—	—	9
D	X	—	—	X	—	—	6
E	X	X	X	X	X	—	5
F	X	X	—	X	—	—	3
H	X	—	X	X	—	—	3
I	X	—	—	—	X	X	1
J	X	X	—	X	X	—	1

Figure 4.51 Conversational types without After phases

It is perhaps notable that conversations in which neither participant reach the After phase are the least likely. Conversations such as Type D, where neither participant gets past the Before phase, may be due to non-linguistic reasons, such as a general disinterest in the app or match, an unexpressed sense by either participant that something is off (thus discontinuing the conversation), or for reasons not exhibited in the data.

The latter of these possibilities is demonstrated in the Type D conversation between Nat and Sixten (18), below.

	Nat	Sixten
51	Dolly i think you should give me your number so you can take me XXX sometimes :-)	
52		I don't have an English phone number though, and it is quite expensive to write back and forth to a XXXish number. Do you use WhatsApp or something like that?
53	Oh fairs yh i got whatsapp 07XXX x	

Figure 4.52 Nat and Sixten (17) excerpt

It is unclear from the data what occurs after 53, as the conversation ends there. Based

on other conversations in which WhatsApp or similar information is shared, however, it is likely that either or both Nat and Sixten reached the latter two phases off Tinder, either hashing it out over WhatsApp, or not at all.

Likely disinterest is found in a Type D conversation between Morgan and Callum (26), below.

	Morgan	Callum
44	Care to expand?	
45		Not particularly
46	Not the sharing type eh, I'm curious as to what you imply by insightful	
47		Nope, I prefer not to.. Well guess away

Figure 4.53 Morgan and Callum (26) excerpt

The conversation peters out at 47, where Morgan does not take up the challenge from Callum to “guess away”. As Morgan points out in 46, Callum is “Not the sharing type”, frequently responding to other attempts at communication by Morgan with curt responses that do not move the conversation along. Although not expressed here, it is likely that Morgan finds Callum less interesting of a partner than the effort required to keep engaging in conversation, and is not turned off by any particular cues as to Callum’s real-world identity.

In most cases, however, it appears that one or both conversees stop responding because of a general, unspecified feeling that something is off. This is frequently from overtly sexual or suggestive language, such as the following conversation between Tom and James (38).

	Tom	James
1	Are you real?	
2		No I don't actually exist. I'm just a 60 year old man pretending to be a 21 year old called James in the hope of picking up girls.
3	Hahah well that what i was hoping for to be honest. I mean who wants to talk to attractive girl when you can talk to old men 😊	
4		Exactly my thought process as well haha. Take it you're real then? 🙄
5	Unfortunately yes, so I'm guessing this isn't James then 😊	
6		Well I was kind of joking about the 60 year old man thing, so unfortunately, yes I am 21 year old James 😊
7	Well since you're another guy I guess we should talk about manly like sports and cars then 😊	
8	Or we could talk about how good the body of the girl in the picture is 😊	
9		Haha sport and cars are good, but the body is better 😊
10	Definitely! Yeah she looks like shed be a great fuck 😊	
11		Haha she does, but I wouldn't know, you'd have to ask someone else about that 🙄
12	Its a real she I Can't talk to her, I'd love to tell her the things we could do together	
13	Because that's what men talk about 😊	

Figure 4.54 Tom and James (38)

Although they both refer to a “her”, Tom is talking about James’ profile, and James is talking about Tom’s. That James is likely under the impression that Tom is overtly sexually objectifying “herself” likely has something to do with his lack of response after 14, stating in 11 “you’d have to ask someone else about that” in regard to how good of a fuck “she” is. Such misapplications of information are discussed further in Section C below.

Although many of the conversational flags discussed in Section C that lead to phase transitions are indications of a conversant’s real-world identity, as we have seen with these examples, it is not necessarily a successful identity performance that keeps conversees from reaching latter conversational phases.

C. Conversational Transitions

In order for these conversations to progress from phase to phase, or even linger in one phase, conversees must negotiate a variety of strategies, consciously or not, and either ratify informational discrepancies as acceptable variations within the perceived identity of their conversee, or redefine their understanding (or lack thereof) of their conversee’s identity. At the crux of these strategies are informational cues, or A. Conversational Flags, in which a conversant provides identifying information about themselves or their conversee. The presence of these flags may be mitigated by the 4. Accommodation of Gender-coded features by either or both parties, 2. Turn Taking Strategies inherent to CMC, or other 3. Strategies of Negotiation that may be largely restricted to such specific types of interactions as these.

1. Conversational Flags

As discussed above, the conversations progress through a possible three phases, termed here Before, During, and After. Although not all conversations or conversees progress through all three phases, progression tends to be at the introduction of informational flags, which may or may not be negotiated by conversational participants. These flags are of four general kinds:

◆	Gendered References to Conversees	referring to a conversee as a girl, woman, lady, etc.
		referring to a conversee by the female names in their profile
		referring to a conversee by usually female-coded terms
▲	Non-Identifying Real-World Information	hobbies, occupations, orientation, or other real-world information that is neutral-coded as far as gender, or can be negotiated as female-possible if male coded
▲	Identifying Real-World Information	an author providing their (non-neutral-coded) real name
		an author providing their gender
		links to non-Tinder accounts (Instagram, Facebook, WhatsApp, etc.) that likely contain real-world, gender-identifying information

Table 4.36 Conversational flag types

The first two flags are made by speaker A about speaker B, and may prompt speaker B to question speaker A’s impression of their identity. The second two flags are made by speaker A about speaker A, and may prompt speaker B to question their own impressions of speaker B’s identity.

In many cases, as we will see with the examples below, these flags can be either

catalysts for conversational phase progression, or can otherwise be negotiated by conversees as potentially questionable but still explainable discrepancies about their assumptions as to their conversational partner’s identity as provided by their Tinder profile.

2. Turn Taking Strategies

	Before	During	After	Before	During	After	Total
A	X	X	X	X	X	X	12

Figure 4.55 Type C conversations

While Chris and Conor take a Type C trajectory (Conor skips the During phase), Martyn and William engage in a Type A conversation, the most common type, and very similar to Type B, differing only in that both participants transition through all three phases. For both Martyn and William, each of these phases is lengthy in comparison to many other conversations.



Figure 4.56 Martyn and William (28) flags trajectory 1

Similarly to Chris and Conor, the conversation here begins with each thinking their audience is a female—this time by the names Shawnee and Skye, respectively—and accommodating their speech accordingly until the truth is eventually revealed around 105.

	Martyn	William
102	Well im gonna use my real name	
103		So where's Skye come from? You've got me a little confused here
104	Cos i like it more than shawnee	
105	Martyn	
106	Marty	
107	Thats your name on here	
108		Well this is weird
109	Seriously....	
110		Yep
111	Uh...you have blonde hair and super white teeth in your photo	
112	Is this u?	
113		Sure why not
114	Have we been hacked	
115	Or is tinder playing games	
116		One of the two
117		If it makes you feel better reading the messages from a <u>lads</u> perspective you make a mint girl. Kudos 👍
118	What?	
119		I'm a <u>guy mate</u> .
120		You've come up as a very attractive blonde girl named shawnee
121	Wtf	
122		I'm guessing I've come up as the same picture called skye

	Martyn	William
123		No idea
124		It's kindof funny though
125	Ok....	
126	Guess we'll leave it at that	
127		Do you feel abit weird like you need a shower now?
128	Sth like that	
129		This is going to end up on the internet isn't it . . .
130	Not on my end	
131		Well good chatting martyn. Let's never speak of this again pal. Safe!

Figure 4.57 Martyn and William (28) flags excerpt 1

The shift, for them, is not all that drastic—with only a single “wtf” from Martyn and both a “pal” and a “mate” from William—but their confusion takes much longer to sort out, in large part thanks to accommodation and speaker contamination, as discussed in Chapter 3, but also, as discussed here, due to the turn-taking strategies found here and indeed in many such device-mediated online conversations.

The conversation continues, relatively ungendered as the discussion turns to rocks, until 32, when Martyn says that he began his rock collection while traveling “Since I had no luck with women I settled for rocks”. This is a joke, by which Martyn alludes that because he is bad at picking up women, he picked up rocks while out (something a straight female is unlikely to say). William neither responds nor comments, to either this joke, or the one made by Martyn in 37, but this is likely not because he understood or questioned the joke, but because of a pervasive element of speaker contamination caused by the rapid flow of their conversation.

The relative speed with which they are having their exchange is evident in two ways. First, unlike Chris and Conor, who tended to have much longer messages, and only one each per turn, both Martyn and William tend to send multiple, shorter messages before the other responds.

	Martyn	William
41	Btw i wasnt ignoring you I'm a teacher	
42		Ignoring me? I only said hello 15 minutes ago but sure.

Figure 4.58 Martyn and William (28) flags excerpt 2

Second, the excerpt above implies that, not only have they sent 40 messages within the past 15 minutes or so, but also that William does not perceive the conversation as going particularly slowly. This rapid messaging likely contributes to the missed opportunities for either Martyn or William to pick up on various flags throughout the conversation. Both features of the conversation lead to clear signs of a particular aspect of speaker contamination—the “drive-by”.

In the context of a typed, rather than spoken, but still rapid conversation, “drive-by” speaker contamination occurs in ways it may not in a spoken conversation, as discussed in Chapter 3. Because of their nature as written conversations, logged in Tinder and thus easily accessible past the immediate point of conversing, there is often the assumption that everything written was read, comprehended, and implemented in further conversational construction. What we find instead, as demonstrated in the examples below, is that online

conversees may pick up only certain strains of conversation, often multiple at once, and continue on with them to varying degrees of fruition. In the case of Martyn and William, and indeed other conversations discussed here, such rapid, uneven, and sometimes undone topic fluctuation can contribute to flag information that does not conform to a profile's proposed identity *not* being picked up on.

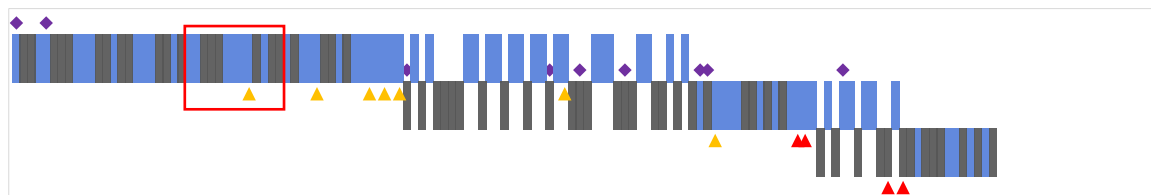


Figure 4.59 Martyn and William (28) flags trajectory 2

The first major flag occurs in 32, shown in the excerpt below, within the exchange about rocks and rock collections between 24 and 36:

	Martyn	William
24	I have a collection	
25	No really i do	
26		I have a pet rock
27		I called it dwayne
28		I really hope you have a rock collection. Can I ask why though?
29	This rock i got is pretty awesome. I picked it up in Ecuador. its volcanic ash that solidified !	
30	Lol I hope he has a good eyebrow	
31	I went travelling in south america so i just had a habit of picking up something in each country	
32	Since i had no luck with women I settled for rocks	
33		See now that would be a rock worthy of a present. An otter brings that as a gift he is getting some
34	Much easier to carry around	
35		Yeah I do this on nights out aswell
36		Lot of rocks in my room

Figure 4.60 Martyn and William (28) flags excerpt 2

Martyn's statement in 24 is not directly responded to by William until 28, though Martyn in turn responds immediately in 29, taking four turns in a row. His second, 30, is his response to messages 26 and 27 by William. He then goes back to answering the posed question in 28 in message 31. Finally, 32 is the joke about picking up women. When William responds in 33, it is to 28, and his responses in 35-36 are to 31. No mention is made, or recognition given by William that he considered 32 at length.

Their conversation is not an even series of statements to responses, but rather a continual, uneven flow, in which some information is not taken up, and possibly barely read, ignored, or not fully considered. Such mismatched responses occur frequently throughout their conversation—and, indeed, occur frequently in many such conversations where a quick back and forth of text is sent between two or more parties; just because a message was sent does not necessarily mean it was caught and accepted into the full flow of both participants' negotiations of the conversation at face value.

Just like Chris and Conor, Martyn and William eventually come to the realization that

something is off about their interaction, or, more specifically, their interactant. William becomes suspicious first after Martyn begins to explain his job and some of his background at 41, as the timeline does not match up with what is likely for the 20-year-old “Shawnee”. William states his skepticism in 69 with, “Just seems a little to good to be real is all,” and in 72 with “Beautiful, smart, well travelled starting a masters and only 20. Gotta admit it’s unusual”. Throughout, Martyn is distressed at the insinuation that he’s “a psychopath using this to lure in victims”, and eventually they negotiate the misunderstanding, in 74-75, as the background Martyn provides is much more realistic for a 27-year-old:

- 74 Im gonna have to break your heart a little though, I’m not
20.I was 20 2,848 days ago
- 75 27 then fair enough, that’s fine with me. I’ll check your
credentials then.

Figure 4.61 Martyn and William (28) flags excerpt 3

It is only at 92 that *both* clearly suspect something is up, and at 119 where William finally flat-out states, “I’m a guy mate,” having realized from Martyn’s reveal of his real name in 105 what is going on. Martyn goes from doing most of the talking, or at least sending more frequent messages, to shorter, less frequent responses, until the conversation finally peters out at 131.

Another conversation, between James and Tom, has a similar mismatch in conversational strains. Although this progression does not keep James and Tom from discovering the truth of the other’s identity for as long as it does Martyn and William, at only 46 turns in length as opposed to 105, the mismatched turn-taking does exacerbate the confusion, requiring the conversees to backtrack through prior, not fully ratified statements in order to consolidate some shared truth between them.

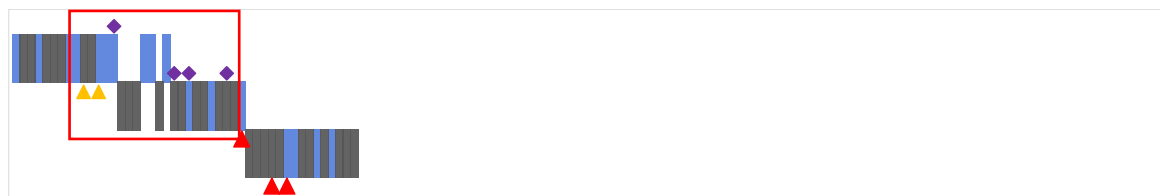


Figure 4.62 James and Tom (39) flags trajectory

As shown in the color-coded excerpt from 9 to 27 below, both James and Tom give their profession after having been asked by the other. Both happen to work in a physically-demanding profession which is generally more common for males than the 20-something females presented in the fake Tinder profiles. Both express disbelief for this exact reason, and yet it is only via some backtracking that they finally sort out the confusion.

- | | James | Tom |
|----|--|--|
| 9 | So what do you work as? | |
| 10 | | I’m a carpenter/foreman on a site in XXX at the moment |
| 11 | | How about you? |
| 12 | Joining XXX XXX in a month but at the minute labouring from site to site | |
| 13 | You’re really a foreman? | |
| 14 | I just find it hard to believe a pretty girl like you works on a building site | |
| 15 | | Your labouring!? |
| 16 | | Pretty girl haha? |

	James	Tom
17		What's going on here
18	Yeah I don't mind it	
19	I think honesty, haven't tried it in a while	
20		I'm confused
21	Why?	
22		So you're a pretty girl working on a site.
23		What do you mean about honesty?
24	No your a pretty girl working on site	
25		Haha!
26		No you are
27	Well this is weird	
28		Haha
29		Your the girl working on site!
30		Unless your a ladyboy
31	I'm pretty sure I'm a man, I can prove it	
32		Your a man?
33		Albertina?
34		This is strange
35		I'm pretty sure I'm a man also
36		My pictures are a lot more manly then your haha
37	James actually	
38	My name is	
39		Well that's not what it says!
40		Getting strange now
41	Well...	
42		Well...
43	Ever made love to a man?	
44		Haha
45		Why are your pics of women?
46		And no

Figure 4.63 James and Tom (39) flags

After beginning in 9 by **asking for Tom's profession**, Tom reciprocates by **asking for James' profession** in 11. This begins the back and forth between the parallel topics, as each conversee expresses the very same disbelief at the other's profession, and which of them is "a pretty girl working on a site". They resolve the issue relatively quickly, but the back and forth due to the uneven conversational flow adds in a bit of otherwise unnecessary confusion, as pointed out in line 20.

Such potentially confusing conversational trajectories are not necessarily endemic to these types of conversations. As shown in the conversation between David and Adam below, many participants address multiple topics in longer messages, or otherwise complete topics before either conversant introduces a new one, aiding in conversational flow, and minimizing confusion.

	David	Adam
1	Hi, how are you?	
2		Hi. I'm good and yourself?
3	Not bad thanks. What brings you to the wonderful world of tinder?!	
4		The usual i guess. Single, bored, looking to see who else is on and maybe go out on the weekends. You?
5		Thats not the Santa Monica pier in your photo is it? I

	David	Adam
		love that place
6	Pretty much the same. Meet new people and all that. No it's not. I've never been there. Only part of America I've been to is New York. Would love to go to California one day though	
7		It looks just like the one in your photo then again most piers do look alike =). You from England then I take it?
8	I am from England yes. Where are you from?	
9		XXX
10		The central part though unfortunately not the coast
11	Oh right! What's made you cross the pond?	
12		Work. Ive been here almost 4 years now and I like it a lot for most part.

Figure 4.64 David and Adam (37) topical progression excerpt

Despite some topical overlap, both David and Adam are able to keep up with the conversational flow along every topical strain due to their conversational turn strategies.

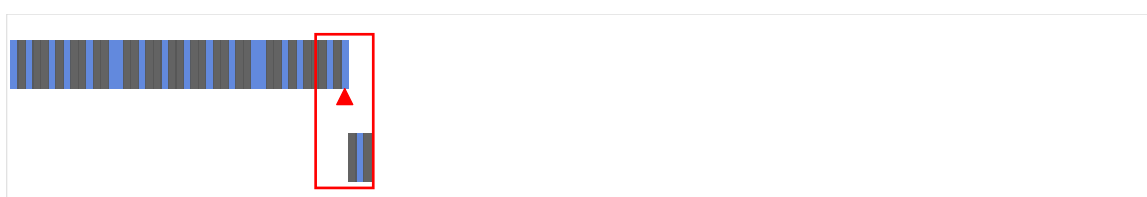


Figure 4.65 David and Adam (37) flags trajectory

This continues on throughout the conversation, to the point that, when a red flag is introduced at 45, both David and Adam are immediately on the exact same page, and the identity confusion is resolved swiftly, with the conversation ending only three turns later in 48.

	David	Adam
43	Can I tell you something without you un matching me straight away?	
44		Sure lol
45	Please don't be scared off! I'm a Daddy	
46		So you're a guy...
47	Er yeah! I thought you were a girl	
48		Better luck next time then

Figure 4.66 David and Adam (37) flags excerpt

While turn-taking strategies and conversational flow can contribute to the identity confusion inherent to these conversations, as we will see in the section below, such mismatched, unratified topical exchanges are but an exacerbating factor of other strategies of negotiation that facilitate misunderstanding beyond the informational red flags that the outside reader may believe should lead to a reveal or at least suspicion much sooner than they sometimes do.

3. Strategies of Negotiation

As discussed above, conversations, in general, tend to be cooperative. Tinder conversations in particular are very likely to fall into the category of cooperative communication, as both participants are motivated by a shared end goal to maintain their own face value, as well as the value of their conversational participant and hopeful romantic or sexual partner. It is precisely because of the overt cooperation and accommodation motivated

by the genre of Tinder and similar chat types that interlocutors are able to favorably negotiate information that may not immediately conform to their initial preconceptions regarding their conversational partner—in the case of the data here, specifically, their gender. (This, as well as the opposite, is discussed further in Part 5 below.)

What are termed strategies here are not necessarily intentionally implemented by speakers, and are certainly not intentionally meant to hinder their understanding of the real-world facts presented to them throughout the conversation, but rather are employed or accepted due to the often-shared desire to cooperate. Often, the strategies discussed and demonstrated below exist in some combination with one another, and are exacerbated by the non-linear conversational flow, as discussed in the section above.

4. Accommodation via Gender-coded features

As discussed above, Chris and Conor's conversation follows a Type C trajectory, wherein Chris transitions through all three conversational phases, and Conor skips the During phase, but does reach the After phase.

	Before	During	After	Before	During	After	Total
C	X	X	X	X	—	X	8

Figure 4.67 Type C conversation

The trajectory of their conversation is demonstrated in Figure 4.68 below, with the leftmost speaker, Chris, in blue, and the rightmost speaker, Conor, in grey.

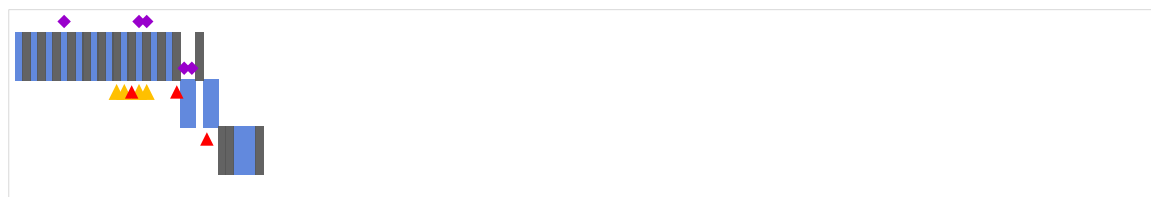


Figure 4.68 Chris and Conor (5) flags trajectory

A variety of flags occur throughout their conversation as shown below.

	Chris	Conor
11	Hey, I'm really bad at starting conversations so do you wanna hear a terrible joke?	
12		I'll beat you to the punchline and say my love life
13	Hahaha, it's that bad?	
14		Well I can say tinder isn't really my first resort of finding new people to talk to, more like the last
15	I wouldn't have thought you'd have trouble finding people interested in you	
16		What makes you think that haha
17	Well you're rather gorgeous haha	
18		Thanks haha >.< as much as I don't believe that, its nice to receive such compliments from a stranger :)
19	You're very welcome, you should believe it though, trust me :)	
20		I've always been one to return a compliment, and have to say you're quite attractive yourself haha
21	Aww thank you :) so I have to ask, what are you looking for on here, just people to talk to?	
22		You're quite welcome :) Mostly yeah, just got the 'whatever happens, happens' kinda mentality towards it though

	Chris	Conor
23	Same as me then really :) so what are your interests?	
24		I don't have high hopes though, with this app, to me, seeming more like a 'smash or pass' sorta thing haha. Mostly cars, it sounds generic, but I work at a garage and enjoy working on them, we get some modified stuff in most days and its fun learning new things :) what about you, what catches your attention
25	Haha, well as much as I would 'smash' you, that's not all I'm looking for, don't worry :p and that's not generic at all, pretty damn cool actually! I'm a musician so that's what I'm doing most of the time, just finishing up an album at the moment	
26		Ha! Well I can honestly say Id have no objections if it came to that haha. I suppose its more generic for 15 year old boys that 21 year old guys haha, not a bad work choice though, it isn't often people do things they enjoy, so Ive heard. That sounds pretty neat, what instrument and genre do you play? :)
27	Maybe we'll have to make that happen sometime then ;) and yeah I guess haha, it's also way less generic for a 20-year-old girl to be into. I play lots of stuff! Mainly guitar, I sing, bit of piano, drums etc. And this album is singer-songwriter but I play in a rock band too :) what music are you into?	
28		Maybe, if Id ever get that lucky :p that is true, but there are often some women that are in to it, just a dime a dozen haha. A bit of a multi-talented it seems then haha, you'll have to let me listen to some of your work some time though :) Im mostly in to the heavier sorta music, but when it comes to chilling out I enjoy a bit of deep house, makes for a nice break from the heavier stuff :)
29	Pfft, YOU get lucky? I would most definitely be the lucky one! And you are obviously that dime then haha. Umm I could send you a link to a couple of tracks from my new album if you'd like?	
30		Oh you, haha, I think we'll have to see about that if or when the time comes ;) if they're youtube links then yeah sure, I can't see much else working on iphone, they're ridiculously limited which is the only thing I hate about them haha
31	They're not on YouTube yet unfortunately. Do you have facebook? I could send you the links on there and you can check them out next time you get to a computer? (also it means I get your Facebook ;))	
32		Fair enough haha, good god that was smooth though, I love it ;) search for Conor XXX, Id honestly be surprised if anyone else has the same name as I do haha
33	Wait, so you're a guy..?	
34	Has tinder glitched out on us or are you pretending to be a 20-year-old girl called dolly?	
35		Well I mean I do believe my profile covers that aspect pretty well. Haha
36	I'm a 21-year-old guy called Chris btw if that's not what you're seeing?	
37	I'm so confused	
38		Hahaha oh god that is brilliant. Well what I'm seeing is a

	Chris	Conor
39		22 year old asian lady called Sybil
40	No. Fucking. Way. This is hilarious dude!	This has made my day.
41	I'm gonna add you on Facebook hold on	
42	Don't delete me on here, I'll send you screenshots hahaha	
43		Someone at tinder hq has some issues to fix it seems hahaha

Figure 4.69 Chris and Conor (5) flags

Clearly, both Chris and Conor at least initially believe their audience is an available female with whom they have matched on a dating application. Both believe this to be true until in 22, Conor provides his full real name, which is *not* “Dolly”, and from 24 to the conclusion of the conversation at 33, the truth is discovered by both of them. Although this is a single conversation with only two participants, it is worth noting various features that stay exclusive to the conversation prior to 22, and after.

Both Chris and Conor use emojis, usually :) and ;), *only* while they still believe the other to be a female. They also tend to have much longer messages, especially once the conversation picks up at 11. Interestingly, these are both features suspected to be more indicative of feminine writing. Although neither male is actively pretending to be female, both believe they are interacting with a female, and this likely colors their methods of discourse to, if not something more overtly feminine, then something more perceptually gender neutral. Each attempt to accommodate the other’s perceived identity, as is further discussed below.

After 22, the emojis stop, and the posts get notably shorter. In addition to this, we find the only swear throughout the conversation in 30: “No. Fucking. Way.” In that same turn, we also find a hypermasculine term in “This is hilarious **dude!**” In this conversation, and in similar conversations on Catfi.sh, the exchange remains cordial after the reveal, and both Chris and Conor laugh at the prank. In similar situations, where both participants become aware that they have *both* been trolled by a third party and not by the other person, such hypermasculine terms pop up very commonly, as in the excerpted examples below:

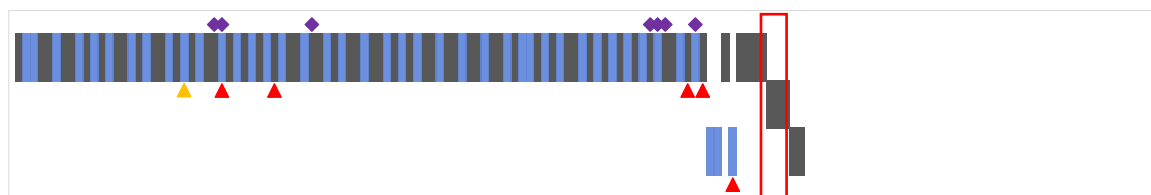


Figure 4.70 Will and Dan (1) flags trajectory

	Will	Dan
94		So your not a girl called dolly? Lmfao
95	Definelty not dolly haha	
96		Lmfao, tinders got you <u>fucked</u> then dude! 4 pictures of some girl and the name dolly 🤔🤔🤔🤔

Figure 4.71 Will and Dan (1) flags excerpt

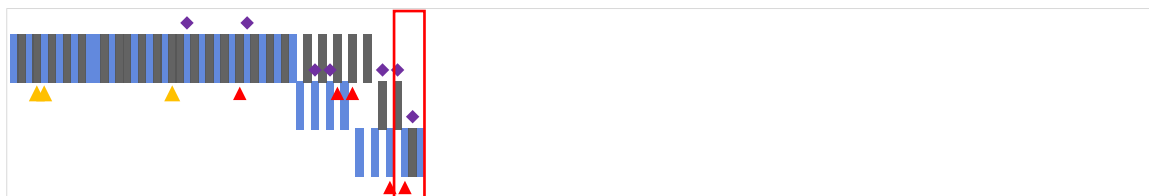


Figure 4.72 Olly and Jack (14) flags trajectory

	Olly	Jack
52		Haha take your a lad
53	Yeah that's true, well <u>fuck</u> . Tinder seems to have <u>ducked</u> me over	
54		Hahaha thats class, take it easy pal .
55	<u>You too</u>	

Figure 4.163 Olly and Jack (14) flags excerpt

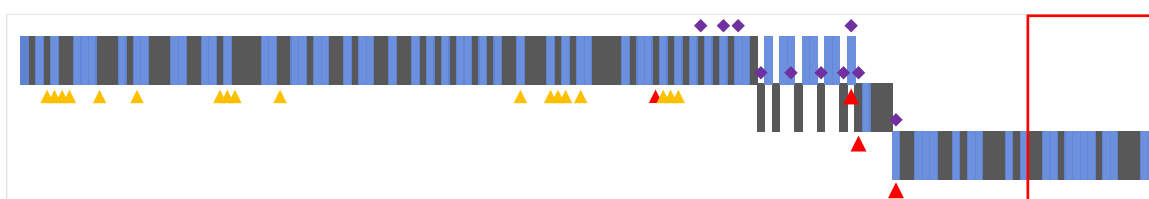


Figure 4.74 Daniel and Joe (4) flags trajectory

	Daniel	Joe
136	Have you not seen any of my moments today then 😊 they're dead give aways Imma guy hahahaha	
137	Can't see yours either	
138		No mate nothing
139		Must be a hack
140	This is so <u>fucked</u> up man 😊	
141		Haha
142		Right mate
143		Sorry for wasting your time
144		Made my <u>fucking</u> night
145	Or a glitch, who knows what the fuck has gone on	
146		Need I screenshot this shit
147		Mv mates aren't going to believe this
148	Haha nah <u>man</u> it's cool. I think this is the best tinder match yet haha	
149	Hahahaha 😊😊😊 no shit they won't!	
150	Hide my number tho when you do. 😊 thanks <u>man</u>	
151		Yeah will <u>mate</u> don't worry
152		See you <u>mate</u>
153		All the best

Figure 4.75 Daniel and Joe (4) flags excerpt

When the third-party ruse is revealed, both participants tend to revert to a shared discourse of masculinity. As shown in the final excerpt above between Daniel and Joe, 5 out of 7 of Daniel's messages contain at least one "man", "shit", or "fuck(ed)", and 7 out of 11 of Joe's messages contain "mate", "shit", or "fucking". Although not excerpted or fully analyzed here, in conversations where one or both parties believe the other person is the one trolling them (and not a third party), very similar masculine discourse is often defaulted to, albeit

sometimes far more aggressively than the relatively friendly and cordial exchanges shown here.

Both Chris and Conor miss multiple cues that the wary observer may see as a clear red flag. In 7, for example, Chris comments that Conor is “gorgeous”, a term usually—but not exclusively—expected to be used to refer to feminine features (and indeed this is exactly what Chris believes he is doing).

Beginning at 14, the two begin to discuss how usual it is for Conor to be into cars and working at a garage. For Chris, who believes he is talking to a 20-year-old female named “Dolly”, the job hardly seems typical. In 16, when Conor comments “I suppose its more generic for 15 year old boys [than] 21 year old guys haha,” the 21-year-old guy he is referring to is himself, for whom he expects the interest seems typical and generic. Obviously, Chris misses this comparison, and the blanks filled in between the 15-year-old boy and 21-year-old guy take on a different meaning—he’s probably the 21-year-old guy “Dolly” means. When, in 17, Chris says “it’s also way less generic for a 20-year-old girl to be into,” that girl is “Dolly”, but, clearly, Conor assumes the girl is “Sybil”, or Chris himself.

Both accommodate any potential miscommunication or misunderstanding by giving primary consideration to the identity they have been presented with—those of the 20-something females Sybil and Dolly—in the other’s profile, and negotiating the real-time information they are given within the context of those profiles. Not only are both likely accommodating their language to be more similar to or neutral against the other’s perceived gender, as discussed above, but both are also accommodating of any small blip that seems to contrast with or confuse what they believe the other’s identity to be.

These miscommunications, as we will see below with Martyn and William, can be exacerbated in both directions, leading to much more overcompensating accommodation, and eventual, compounded confusion.

5. Misapplication of information

The most common types of informational misapplication have to do with names and gendered pronouns. It is unclear from the data itself, and the information provided by the hacker of the hack, why the names chosen were chosen. Although the names and profiles themselves are not provided, the names that appear to be most commonly used based on the conversations themselves are Sybil and Dolly, though other evidenced names include Anette, Quiana, Shawnee, Albertina, and many others.

Some specific names, such as Dolly in particular, appear to cause more confusion than others given their status as apparent pet-names—odd for a female referring to a male, but not entirely unprecedented. In the below excerpt, the use of “dolly” (likely compounded with the lack of capitalization further solidifying the use as that of a pet and not a proper name) by Deji draws no suspicion from Joe. The two plan to meet up, and in fact do, despite the use of the name.

Deji

Joe

98	Hi dolly.	
99	Is tonight still on?	
100		Hello
101		Yeah all good

Figure 4.76 Deji and Joe (20) excerpt

The same can be said to happen with other references to a speaker's conversee, such as in the earlier mentioned conversation between Martyn and William.

	Martyn	William
1	Sup cutie pie! :)	
2		Hey 😊
3		How's it going?
4	Pretty damn good	
5	Anyone tell u how pretty u look today?	
6		Good to hear, it's pretty late you abit of a night owl?
7		Constantly all day, it got very tiring to be honest.
8		It's a curse really
9	Yeah I try to sleep early but never seems to happen	
10	I get distracted	
11	Haha...ok ill stop	

Figure 4.77 Martyn and William (28) excerpt

The accommodation begins almost immediately, with Martyn using “cutie pie” in 1, and asking “Anyone tell u how pretty u look today?” in 5, both statements that sound, generally, more likely to come from a male and be directed at a female. Because William isn't a female, and as far as her profile leads him to suspect, “Shawnee” isn't a male, William takes this as an obvious joke, stating in 7 that he was called pretty “Constantly all day,” and that “it got very tiring to be honest,” and, in 8, “It's a curse really”. Because Martyn meant this as a flirtatious compliment, and not a joke, based on “Skye's” not entirely favorable response, he says, in 11, “Haha...ok ill stop” talking about her looks.

6. Deflection via humor

As discussed in Chapter 2, humor is a common tactic for negotiating politeness and cooperation (e.g., Coates, 2013), especially in exchanges like those in the Tinder data which are rife for misunderstanding. In many cases, the misapplication of information by a ‘listener’ is paired with their perception of humor being employed on the part of the ‘speaker’. Perhaps the most common strategy of negotiation is deflection via humor, generally by the ‘speaker’, wherein either participant either takes a statement that is inconsistent with either of their understood identities as a joke, or otherwise jokes about their own identity, usually not understanding their conversee actually believes the joke. This occurs with the aforementioned name “Dolly”, discussed in the section above, with Deji and Joe.

	Deji	Joe
1	I want to paint you green and spark you like a naughty avocado.	
2	Spank	
3		Wow!
4		Thanks ha
5		Sure that can be arranged if thats really what you want
6	Haha. I like you already.	

	Deji	Joe
7	I'm Dej. You're dolly. When do you fancy going for a drink?	
8		Haha im afraid I don't get your dej and dolly reference
9		And im slightly suspicious of your forwardness
10		But im free right now
11	My names Dej. And you're dolly. And i would like to take you out for a drink. As for my forwardness, i could talk to you for a couple of days about the same mindless things or I could cut the chase and take you for a drink and get to know you better.	
12	Not forward. I'd just rather skip to the interesting part	
13		I like your style
14		So you're not called sybil
15		Haha was just making sure you're not taking the piss!
16	No dolly. My name is not Sybil. Haha	
17	I said it at the start	
18	I'm dej	
19		Where you from dej?
20		Well nice to meet you dej my names dolly

Figure 4.78 Deji and Joe (20) excerpt

While Joe first takes Deji's use as some "dej and dolly reference" in 8, he takes Deji's insistence on him being "dolly" in 11 as a joke, evident by his eventual statement in 20.

Alternately, something that is not meant humorously by the author can be assumed to be meant as humor by the reader. This occurs in the conversation between Gerard and Sam, excerpted in the subsection below, wherein Gerard repeatedly flat-out states that he is "definitely male", and Sam refuses to believe these claims multiple times. This issue is exacerbated by Sam himself trolling, as we will see below.

7. Trolling and reverse trolling

In some cases, the joking extends into trolling—or, depending on the perception of the troll, reverse-trolling—as in the conversation below between Gerard and Sam. Sam first takes every single attempt by Gerard to insist that he is, in fact, male as a joke (212, 245). Even in the face of real-world identifying information, such as Gerard's actual facebook account (233), Sam continues to insist, all the way until the very last turn of their conversation, that Gerard is joking with him, and is, in fact, the female his profile presents him as.

	Gerard	Sam
211		Please confirm that you are a woman !
212	No I'm definitely male	
213	Sorry to disappoint	
214	How old are you?	
215		You're lying to me! Please stop it!
216		And I'm 20 you?
217	I'm also 20	
218		My timer says you're 23
219	07XXX, check my whatsapp if you have it, that's me ^	
220		And I'm a female by the way of you're a male
221		Just to clarify
222	Well at least I haven't been talking to a bloke	
223	Silver lining and all that	
224		Ok and yeah that is good

	Gerard	Sam
225		You don't have a profile picture
226		And how long are you going to lie to me and say that you're male
227	You're joking right - -	
228		No
229	Definitely being catfished or something here 00 :D	
230		I fucking am!
231		If you're make this conversation ends now
232		Male*
233	Gerard XXX, Facebook me ^	
234		So you'd better own up or I'm deleting the chat ;)
235	What's your second name? ☺	
236		No this is over I'm not gay
237	:D :D	
238		^
239	Enjoy the rest of your day ☺	
240		I don't know whether you're lying to me or that you are an actual bloke?
241	Do you have snapchat?	
242		Yes
243		XXX
244		Prove to me you're not a bloke
245	Well I am a bloke :D	
246		Well we can't be talking
247	Not sure now I don't want some randomer having my SC	
248	Link me your Facebook :D	
249		Sam XXX
250	No I mean actual link :D	
251	I'll be there all day	
252		I've only got the app you will have to search for me
253	Right what's your twitter :D	
254		@XXX
255	And you're female yeah? :D	
256		No I'm not I'm having you on!
257	:D	
258		I'm waiting for you to say that you're not male
259		So you're make
260		Male *
261		if you say yes I'm unmatching you
262		Cya

Figure 4.79 Gerard and Sam (27) excerpt

Where this differs from other situations, such as the example of Deji and Joe above, is the fact that Sam decides to (reverse) troll, stating in 220, “And I’m a female by the way if you’re a male”. It’s unclear whether this is an attempt to get Gerard to admit that he is, in fact, a male, as Sam continues to go along with his own joke in 224 after Gerard in 222 states, “Well at least I haven’t been talking to a bloke”.

Although at this point in the conversation Gerard is aware that Sam has somehow gotten the wrong profile information about him, because of Sam’s continued trolling, Gerard continues to believe until at least 225 that Sam is still a female, and not subject to the same kind of mismatched profile as he was.

8. Incidental real-world correctness

There are, however, a few overt examples of negotiation via incidental real-world correctness, such as the example below between Armand and Kasim. Both of them are in the During phase of their conversation and trying to figure out why it does not appear that the profile pictures the other person can see match the profile pictures they know they posted. In 57, Kasim states that Armand looks “far eastern Asian”, and in 64 Armand confirms “east asia filipino to be exact”.

Figure 4.80 Armand and Kasim (21) excerpt

9. Assumptions as to meaning

Tom James

1 Are you real?

2 No I don't actually exist. I'm just a 60 year old man pretending to be a 21 year old called James in the hope of picking up girls.

3 Hahah well that what i was hoping for to be honest. I mean who wants to talk to attractive girl when you can talk to old men 😎

4 Exactly my thought process as well haha. Take it you're real then? 🙄

	Tom	James
5	Unfortunately yes, so I'm guessing this isn't James then 😊	
6		Well I was kind of joking about the 60 year old man thing, so unfortunately, yes I am 21 year old James 🤔
7	Well since you're another guy I guess we should talk about manly like sports and cars then 😊	
8	Or we could talk about how good the body of the girl in the picture is 😊	
9		Haha sport and cars are good, but the body is better 😊
10	Definitely! Yeah she looks like shed be a great fuck 😊	
11		Haha she does, but I wouldn't know, you'd have to ask someone else about that 🤔
12	Its a real she I Can't talk to her, I'd love to tell her the things we could do together	
13	Because that's what men talk about 😊	
14	I wonder if she's flexible omg	

Figure 4.81 Tom and James (38) excerpt

Both apply every instance of discussing “her” to the other’s profile (and not to themselves) from the beginning to the end of the conversation, exacerbated by James’ attempt at humor in 2.

Perhaps the most striking conversation of this type is that between John and Jacob, the entirety of which is provided below, for the number of conversational terms in which both overtly flirt with the other, but neither ever transition past the Before phase.

	John	Jacob
1	Thou art very beautiful	
2		Shall I compare thee to a summers day?
3		🤔
4	Thou canst	
5	You must be a thespian	
6		Thou art more lovely and temperate
7		I am a lover of the arts but no thespian and yourself
8	Wunderschön, I love the arts too	
9	People say I look like a Greek statute	
10	Thou art a model? Ballerina?	
11		More beautiful than a Greek statue
12		and I am not a model so I must be a ballerina
13	Moi? Haha eres demasiada simpática!	
14	Yup I knew it! I can tell from your arms and graceful posture	
15		Hahaha yup just call me twinkle toes
16	Can do! I'm quite the dancer myself	
17	We should enter the county championships	
18		I wish I was that good but am always up for making a show of mesel if you need a student 😊
19		*meself
20	No I need a dance partner, it's ok I'll lead!	
21		Hahahaha you actually need a dance partner random child
22		Is it time to polish the dancing shoes then?
23	It sure is! Of course I need a partner to throw up into the air and spin around on my shoulders	
24		Yeh I can see that happening
25		Don't get me wrong I'll give it a go if you think you can

	John	Jacob
		handle it 😓
26	Of course I can handle the scandal	
27	What style shall we do?	
28		The dad dance?
29	That's a style I struggle with unfortunately	
30	But I'll give it a go	
31	Will that be tonight?	
32		Hey I can lead Wer abouts you from?
33	Sounds like a plan! I am from XXX	
34	And your from Venezuela? What's the country that has produced the most Miss Worlds? That's where you're from!	
35		Red or blue?
36	Red	
37		So does that mean your from Greece the country with the most goddesses
38		Haha would love to meet up but am working in XXX tomorrow
39		Need to get up 😓
40		U getting out tonight are you?
41	Haha of course! Have you got a shoot tomorrow or is it a catwalk?	
42	No I think I will rest this evening! I've been doing some gymnastics and am quite tired!	
43	What are you doing tonight?	
44		No I am the dancing queen?
45	That is true, I thought was Monday to Friday, modelling on Saturday, resting on Sunday	
46		Hahahaha Ye if you want to call wearing a chicken suit and holding a sign point to kfc modelling
47	That's how all the greatest models start, Kate moss worked at McDonald	
48		Ye I worked in maccies for two weeks then quit was to good for them ovs
49	Of course you are, ovs you don't eat maccies or kentuckies too	
50		Nope 3 grains of rice 1 cup of tea and two fingers after each course starving here
51	Jeese, you should be dining like a queen, lobster and caviar every night!	
52		Hahaha am not treated aswell as the goddesses you still being spoon fed the good stuff
53		I hear you took a Angels wings off her for not peeling your grapes before feeding them to you
54	Haha I wouldn't do that, I love grape skins, that's just propaganda	
55	Where did your wings go?	
56		Ay don't dis propaganda it's got us this far
57		My wings they were taken I was told I would find a goddess and win they back
58		Any ideas?
59	Very true! I will get your wings back, even if I have to go to the pits of hell and wrestle Satan himself	
60		Fought you where my girl
61		I definitely wanna see that fight
62	You will, and much more	

	John	Jacob
63		Now am regretting my work ethics 🙄
64		Tell me, did you fall for a shooting star– One without a permanent scar? And did you miss me while you were looking for yourself out there
65	Wow you are very poetic, that Should be a song!	
66	And yes I did fall from shooting star	
67		Ye I think this could be a one hit wonder if I play me cards right 🙄
68	Nah you're more the triple platinum type	
69		Bloody heck you put me any higher on this pedestal am gona get vertigo and fall
70		You don't want the worlds greatest superstar singer/model/ ballerinas blood on your hands
71	I'll get you your wings back before I make the pedestal any higher, of not then I will catch you in my arms	
72		Well you are going to lead when we take over the dancing scene I think I'd like to be caught up in your arms
73	Ah well they are very strong	
74		I think it's time to test the theory
75	How shall I prove my strength unto thee?	
76		i hast many things in mind but a meeting of mind corse and soul is the only true way to defineth this strength thou speaketh of
77	Indeed, thou beauty is only rived by your reason my ladyship	
78		Damn girl your none stop
79	That's how I roll	
80		Hahaha legend when do we roll together
81	Very soon I hope	
82	When the hurley burley is done, when the battle is lost and won	
83		It's time to suit up for battle then
84	My chi is fully channelled and body honed to perfection, both for war and love, whatever my lady desires	
85		My chi is like a Tibetan monks
86		Do you meditate?
87	Of course, I do plenty of chi kung, tan tien and iron shirt	
88	Hence my incredibly chiseled abs	
89		Standard practice with goddesses
90	Sure is, I'm in a constant spiritual meditative state these days, I've got through sitting my sitting on mountain tops contemplating the Tao stage	
91	Ignore the first sitting	
92		Heavy salad
93	Yeah too many olives	
94		Moving swiftly onto my olives
95	Mmm I bet your olives are amazing	
96		Haha am sure u will love them me house mate wrote that
97	Your housemate has good taste, is she a bodacious babe too?	
98		If your into 50 year old women then Ye haha
99	Women attain perfection at 50	

	John	Jacob
100	So you'll gain in beauty and elegance for 28 more years	

Figure 4.82 John and Jacob (2) excerpt

Both refer to one another very frequently by several female-coded terms, such as “beautiful”, “goddess”, “bodacious babe”, “lovely”, “ballerina”, and as an “angel”. Both also make overt reverences to their and their conversee’s gender, such as referring to the other as “ladyship” and “my girl”, and Jacob referring to himself as doing “the dad dance”.

Their conversation combines several instances of deflection via humor, assumptions as to meaning, and misapplication of information, and still neither participant ever transitions. The conversation does come to an end at 100, though with no clear reason other than, perhaps, either or both participants had picked up on the fact that something was at least a little off about their exchange. That neither remarked on it, however, and that it continued for so long is again an example of the type of cooperative conversations inherent to Tinder-like exchanges.

D. Overview

As we have seen in the varying conversational Types and trajectories demonstrated above, although not all authors transition through the same phases, and not all pairs of conversees follow the same progression, what does occur tends to follow similar patterns throughout. That is, many conversees are often able to maintain the Before phase as long as they cover gender-neutral topics, or the information they are provided is otherwise able to be reconciled with the Tinder profile they take to be the ground truth; many authors transition when this information is not able to be reconciled, and strategies of negotiation discussed in Part 4 begin to fail, or when they are presented with overwhelming evidence to the contrary, such as off-Tinder information or their conversational partners’ own language.

Additionally, as we saw above in Part 4, what authors negotiate tends to remain relatively similar, such as names, ages, and other relatively inconsequential information that does not match the provided Tinder profile; information that appears unlikely but not impossible for a female such as their occupation and language; and sometimes even flat-out statements of their conversee’s real gender. Similarly, how authors negotiate tends to fall into many of the same categories, such as via the misapplication of information, deflection via humor and joking, deflection via trolling and reverse-trolling, incidental real-world correctness (such as a Tinder profile coincidentally having the same minority race as the conversee), and other assumptions as to meaning.

Bolstering these strategies of negotiation are Turn Taking Strategies, also discussed in Part 4 above, by which conversees confuse responses due to the non-linearity of computer-mediated communication such as the messaging type found on Tinder. Oftentimes, these uneven turns result in multiple overlapping and parallel conversations, by which conversees may further negotiate any otherwise contradictory information as belonging to another, less or non-contradictory conversational string.

As we have seen throughout this chapter, several factors can come into play at once in these conversations, either causing or deflecting suspicion, all coloring a conversant's perception of their conversee's presumed and performed identities as either genuine or ingenuine. A number of these factors often intersect, such as turn-taking strategies leading to the misapplication of information that might otherwise be considered a red flag in the flow of a non-computer-mediated conversation.

These strategies often come down to the influences of the linguistic theories discussed earlier in Chapter 3—audience design, accommodation theory, community of practice, and speaker contamination. Even in conversations where both participants amicably reach the same conclusion (regardless of the truth of that conclusion) these theories can be seen at work with the frequent shift to more male-coded linguistic features, such as the use of swear words and kinship terms discussed in Parts 2 and 3 above.

As we will see in Part 5 below, however, the general transition of these conversations, both through the three possible phases, and in how the red flags discussed in section A above are handled, have much to do with the genre of these particular conversation types. Speakers are more willing to accommodate in instances where it is to their benefit to do so—in the case of Tinder, a lack of perceived accommodation runs the risk of missing out on a potential romantic or sexual partner. In other settings, as will be discussed below, conversees may not be so primed to be accommodating, and in some cases, may be particularly primed to be overly suspicious.

Part 5 – Other Instances

As discussed in the previous Parts, most of the conversations considered in the Tinder data transition past the Before phase, but only a single participant in a single conversation begins in the During phase all together. Cooperation is, as mentioned above, a commodity for such conversations, where the goal is frequently a romantic or sexual pairing between conversees.

Such cooperation is not always typical—and indeed may itself be considered suspicious—in many other conversational settings, however. As we will see below, this occurs in instances where it is not cooperation but suspicion that is the prime currency of an exchange (such as in online child predator groups), where language and identity performance is seldom taken at face value regardless of how convincing it may be. Even on platforms such as Tinder, however, it is indeed identity performance that sparks immediate suspicion. This occurs *not* when a conversee fails to maintain a convincing performance of their presented profile, as in the data discussed in previous Parts, but rather when a conversee's language fits particular archetypes endemic to services like Tinder.

A. Suspicion on Tinder

Only once in the Catfi.sh Tinder data does immediate suspicion occur, in the Type F conversation between Adam and Mark below, where Mark begins the conversation

immediately in the During phase, stating in 5 that he already knows the conversation is “some sort of scam”.

	Adam	Mark
1	Hey Sadie! How goes it? Here's my number...XXX XXX XXX...would love to get to know you... :-)	
2		Slut
3	Umm not looking for just a hookup so don't rush to judgement	
4	Thnx	
5		I know but its still some sort of scam
6	Not really.. I am Just a guy looking to chat and see where things go from here... 100% not a scam	
7		Ur a dude posing as a woman?
8		Nuff said
9	Um no.....	
10	Sounds like you have been catfished quite a bit... I am actually a pretty genuine guy just looking for something real and not just for a fling or to try and scam someone like you are assuming.	
11		But ur a dude using women's pics as a way to meet dudes.

Figure 4.83 Adam and Mark (52)

Although Mark does not accuse Adam of being a bot, it is likely a combination of the wrong name “Sadie” and Adam’s phone number that leads to Mark’s response in 2. Scams on Tinder are not uncommon, sometimes coupled with the tell-tale language of a bot.

That such interactions are frequent is demonstrated by the examples below, posted online due to the overtness of the bot’s failure at convincing their target that they are a genuine Tinder user. The human participants go out of their way to subvert the bot-preferred trajectories of these exchanges (in which the human will be successfully scammed), which is itself demonstrative of the fact that such exchanges are hardly unprecedented.

Certain giveaways are expected, such as in the following two examples in Figure 4.84 and Figure 4.85. In the first example, the user immediately predicts the bot by parroting their messages, not expecting to get any pushback from pre-programmed responses. This is possibly due to the user’s profile, as the image used is of a famous 90s-era Tejano performer, Selena Quintanilla, who was murdered in 1995, 17 years before Tinder was launched.

In the second example, the user explains exactly what sort of profile makeup is typical of a bot: “There’s a rule of thumb on tinder (for guys) about bots. 1-2 photos, no profile info, and messages first. You hit all of those haha.”



Figure 4.84 Tinder bot (3)¹⁴

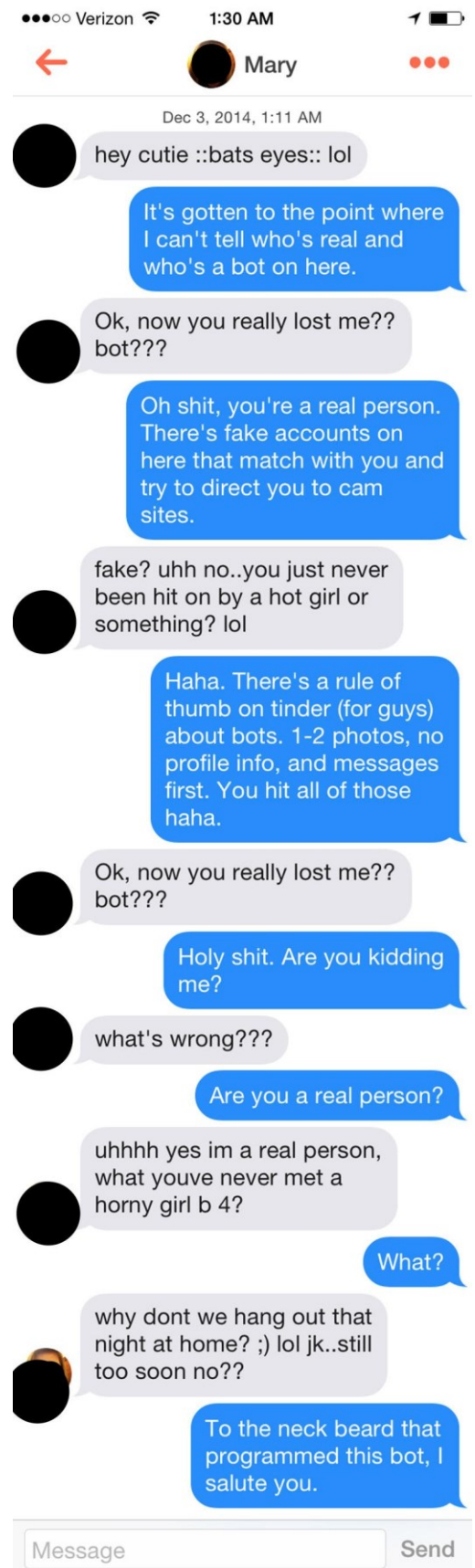


Figure 4.85 Tinder bot (4)¹⁵

It appears that in some cases, simply a profile alone (and not the additional “messages first”) is enough to arouse suspicion, such as in both of the below instances. The non-bot

¹⁴ <https://i.imgur.com/2ExrSoZ.png>

¹⁵ <https://img.17qq.com/images/diopfpogijz.jpeg>

participant messages first, outright asking if one is a bot, and saying “Hello world!” (a common programming example) to the other, without any information aside from the profiles and match to go off of.

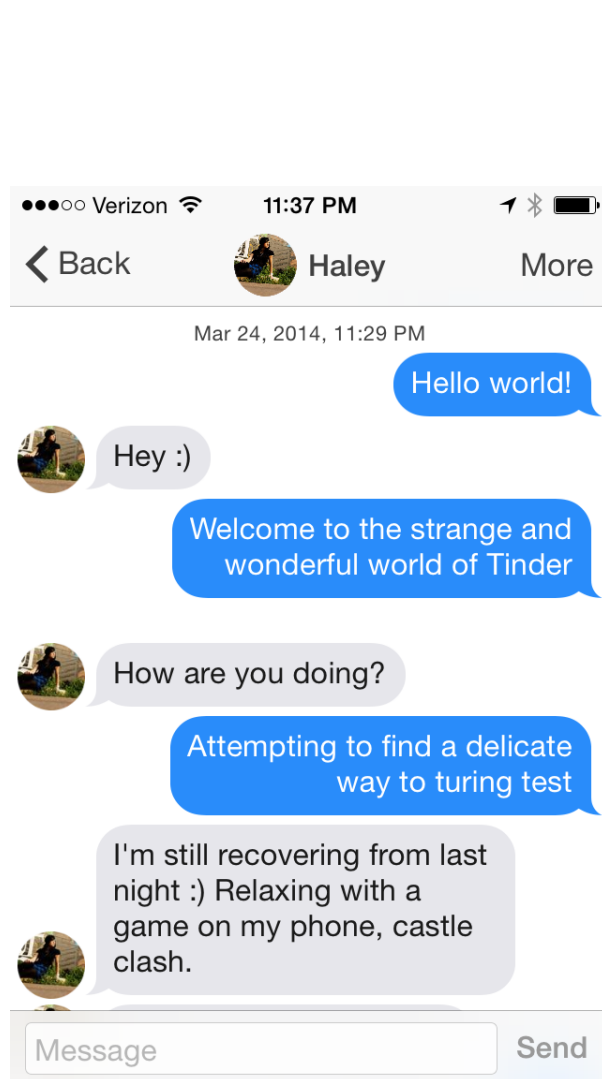


Figure 4.86 Tinder bot (5)¹⁶



Figure 4.87 Tinder bot (6)¹⁷

That profiles are enough to give a Tinder user away as a potential bot, or at least potentially not who they claim to be, suggests that while cooperation is a commodity, it is not the only approach users need to take when performance indicates otherwise. This further suggests that the profiles of the Catfi.sh data were crafted in such a way as to not raise such suspicion—and the fact that the users were not, indeed, bots, but real people with genuine, non-pre-programmed responses, likely helped.

Similarly, that in both Figure 4.86 and Figure 4.87 the users can predict the next turn or the conversational progression of the bots is indicative of bots as an expectation.

¹⁶ <https://www.dailydot.com/wp-content/uploads/5f0/71/tinderbot1.png>

¹⁷ <https://i.imgur.com/RON9Ntg.png>

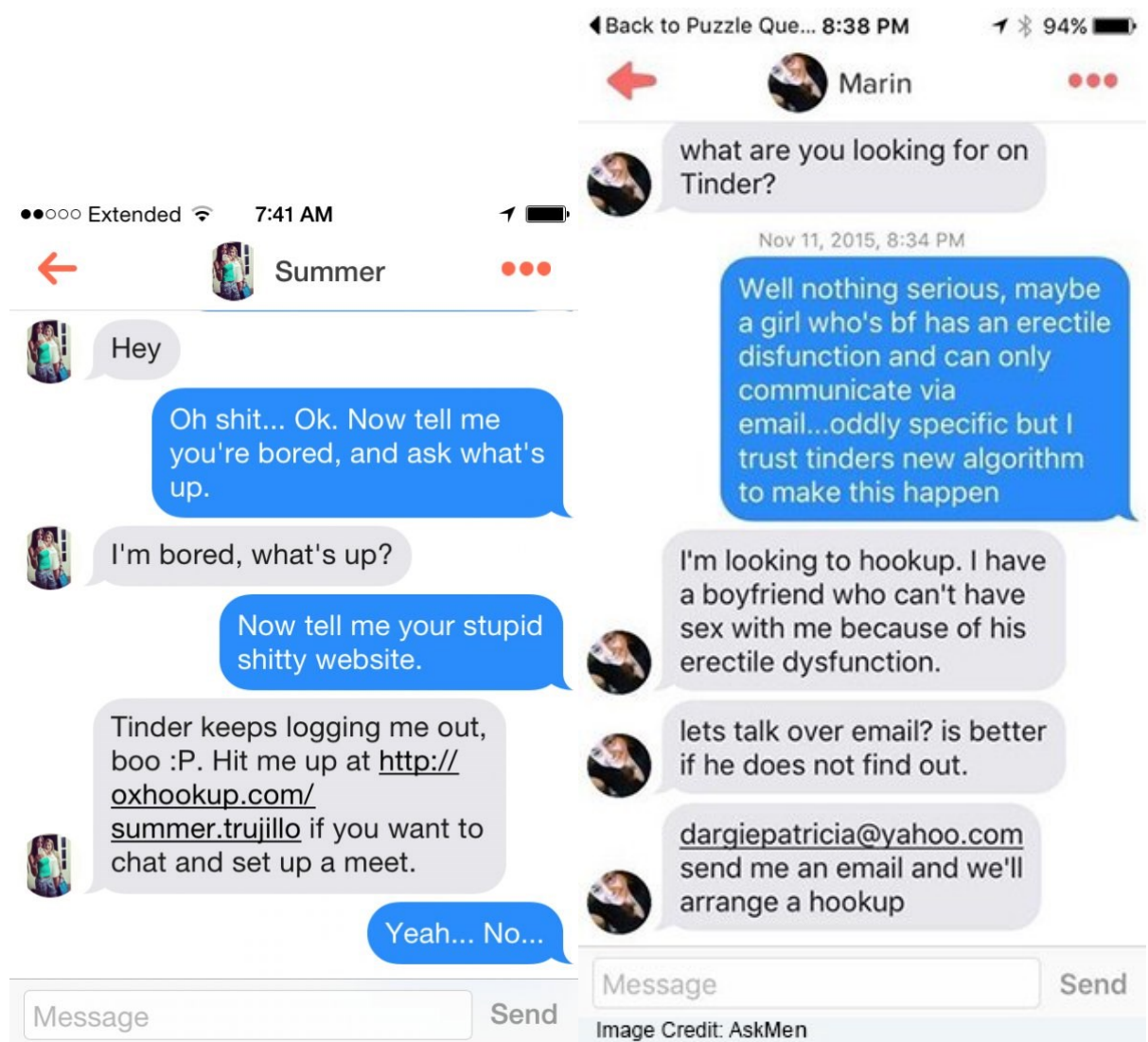


Figure 4.88 Tinder bot (1)¹⁸

Figure 4.89 Tinder bot (2)¹⁹

As with the Ashley Madison bots, Tinder bots do not appear to put much effort into the actual language of their exchanges, frequently relying on the bare minimum for both profiles and communication to hook enough users into whatever likely scam they are peddling. That the Catfi.sh approach appeared largely successful (with 107 out of 108 users not being immediately suspicious of their interlocutor) would suggest that the profile itself, the naturalness of a user's language, and perceived consistency between the two, are more important in identity perception than the potentially gender-coded or gender-indicative performance by a user.

It is worth noting, however, that in the case of the Catfi.sh data, no users were actively attempting to participate in catfishing (save for the times they temporarily reverse trolled after assuming they were being trolled themselves), and thus there was no overt guising to their language in the way that there is likely to be with bots or other scammers in such exchanges. As we will see in section B below, where performers take an active role in portraying a false identity, for a suspicious interlocutor, the types of flags discussed in Part 4 hold much more

¹⁸ <https://breakbrunch.com/wp-content/uploads/2015/09/online-dating-suck-091315-3.png>

¹⁹ <https://www.techjunkie.com/wp-content/uploads/2017/04/amtinderbot22111.jpg>

weight in their perception of the proposed identity of their conversee.

B. Suspicion Elsewhere

MacLeod and Grant, 2017, provide the following example of an instant message exchange between users in a darkweb child porn community (p 3)”

Extract 1: Suspicion of identity disguise on the Darkweb

11/26/15 9:58 PM	Username 1	im a 14 f
11/26/15 9:59 PM	Username 2	whatever you say Officer Username 1 [...]
12/10/15 3:52 PM	Username 3	I just want to pm pics for a good vid site
12/10/15 3:53 PM	Username 2	if they enter room wanting to pm odds are they are LEA [...]
12/15/15 5:08 AM	Username 4	let me rephrase...ONLY PM me if you are into private trade
12/15/15 5:10 AM	Username 2	no thanks officer Username 4

As with the formula provided in Figure 9.5 above for Tinder users to recognize bots, MacLeod and Grant provide similar behavioral flags that are likely to tip interactants off to the possibility that something is amiss:

Extract 1 shows the individual ‘Username 2’ flagging particular behaviors – such as claiming to be a teenage girl or requesting PM (private message) chats – as indicative of an individual being ‘LEA’ (a Law Enforcement Agent). For obvious reasons, members of these online communities are constantly on their guard, regarding all their interlocutors with suspicion.

Such a community is one that would exist in stark contrast to the data analyzed in Parts 3 and 4 above. While some of the catfished Tinder users eventually ratified red flags, suspicion was not the same commodity as it is in the darkweb communities of Extract 1. Indeed, even Figure 4.84 demonstrates that the Tinder user involved is willing to be convinced that his initial suspicions are false, hence prompting his explanation of the general formula for suspicion of other users as bots.

It is also equally possible in the catfished Tinder data that users either do not express their suspicion overtly, or are otherwise suspicious enough that the culmination of their suspicion is to terminate the conversation without investigation (and therefore are not seen to transition to the During and/or After phases). Indeed, MacLeod and Grant go on to say, “Although we have no specific examples across our datasets of UCOs’ cover being blown, it seems safe to assume that this is because such a discovery would lead to an immediate end to interactions, and thus a lack of data.” This is echoed in Chiang 2021 who suggests another potential formula for suspicion in a similar study between UCOs and suspected child groomers:

While the UC’s move use is generally close to that of suspected offenders, it is possible that the small differences observed (in particular the UC’s comparatively limited tendency to describe sexual and abusive interests, experiences and events, and increased tendency to request media) mark a notable departure from the linguistic behaviours of genuine offenders in these sorts of interactions, raising a red flag for offenders ever-suspicious of covert online police activity.

Chiang similarly observed that in cases where there is no suggestion of suspicion on the part of the suspected offenders, “it is possible that suspicions went unvoiced.”

It is worth noting, however, that despite the examples above from child porn and child grooming communities and members, not every community in which a user is primed to be more suspicious than cooperative is criminal in nature. Indeed, we saw this in the Chapter 1 example of a “G.I.R.L.”—or “guy in real life”. While a somewhat outdated perception by (male) internet users that “there are no girls on the internet”, the concept, even if jokingly, is reified enough to have a TV Tropes²⁰ entry, multiple Urban Dictionary²¹ entries, and a young adult fiction book²² named for and exploring the subversion of the concept.

Part 6 – Discussion/Overview

The issues of perception in the Tinder data can be seen as issues of commodity—in this case, cooperation—overwriting various conversational flags in favor of accommodative approaches to negotiating language. Chapter 3 detailed four instances in which the theories of cooperative language discussed therein are what provide the context for similar misunderstandings of device-mediated language.

As with the example of a Tumblr user discussing the murder of their Sims husband with the assumption that the audience shared their schema of being Sims players, the Tinder authors constructed their halves of the conversation on the baseline schema that they were the male, their partner and audience was the female, and they both understood the same gender distribution. This was clear again and again when authors not only gave gendered information about themselves, but also used gendered language to describe their interlocutor.

In the examples of “Mom” attempting to converge her language downward into the format of chatspeak, attempts at accommodation (itself indicative of cooperation as the likely commodity) failed. Similarly, the Tinder authors appeared to attempt to accommodate via their own language, frequently converging upward in the Before phase, and downward in the After phase. This is evidenced by the trajectory of, for example, swear words and friendship (specifically “bro”) terms, as demonstrated above. Although these features are not exclusive to males, they are both preferred by males and perceived as male coded. So it is perhaps not odd, then, that these features hardly showed up in the Before phases if they showed up at all for individual authors, only to spike drastically in the After phase. While not conscious, this trend in the Before phase was likely accommodative to a female audience, and then in the After phase, accommodative of a male one.

In the single Chapter 3 example that resulted in real-world legal consequences (or at least actual jail time), Justin Carter made the mistake of performing his in-group identity as part of the community of practice of League of Legends players to a wider audience, and not accommodating his language to be less hyperbolic and charged with infelicitous threats for an audience that would not expect such language as a baseline or at all commonplace. While nothing quite so drastic occurred in the Tinder data, assumptions about the gendered

²⁰ <https://tvtropes.org/pmwiki/pmwiki.php/Main/GIRL>

²¹ <https://www.urbandictionary.com/define.php?term=GirL>

²² <https://www.goodreads.com/book/show/18599748-guy-in-real-life>

distributions of various communities of practice were frequently negotiated. As an example, James and Tom (39) argued over which of them was the “pretty girl working on a site” in Parts 3 and 4. It was both of their access to this shared community of practice—that of manual laborers—that clued them in to the fact that something was likely amiss between the Tinder profiles they had been presented with, and the person with whom they were actually communicating.

Finally, the example of Tyson Leon considered the application of speaker contamination, not only conversation-internally, but as it relates to an author’s entire, online, public linguistic history. While the latter aspect had no real place in the Tinder data (as, hopefully, none of the unwilling participants were conned into participating more than once), the former was a large point in Part 4’s turn-taking strategies. Just as it is incredibly important for analysts and investigators to take into account the genre of CMC before deciding that every single word of a conversation was read, ratified, understood, considered, remembered, and replied to in exact turn, it likely would have been useful for the Tinder interactants to pay much more explicit attention, and even *call* much more explicit attention, to the various turns they appeared to gloss over in favor of cooperation as the commodity.

Chapter 1 introduced three examples of instances in which an identity or identities were disguised primarily (but not exclusively) on the basis of gender. In all of these cases, both cooperation and suspicion were variably commodified.

Ashley Madison, a site that caters to users not simply interested in starting a romantic or physical relationship (as on Tinder), but specifically interested in an extramarital affair, was one in which users had to walk the line between cooperation and suspicion. Like with the Tinder data, users, especially men, are benefitted by being cooperative when their end goal is to transition on-site communication into the real world. Although both Tinder and Ashley Madison users, or at least adept ones, couch this cooperation within the baseline suspicion of possible bots, catfishes, or other scams, Ashley Madison users are doubly motivated in their suspicion as they are at risk of getting “caught” by their spouses. The team behind Ashley Madison, however, could only devote resources to the former—not to help users avoid bots, but to avoid users from detecting their own, company-sanctioned bots. Knowing suspicion was a commodity for the exact types of reasons mentioned on Tinder above, almost all of the resources expended upon “Ashley’s Angels” was in their profiles being flagged by users as potential bots, and nothing further. Indeed, it seems that the efforts the company took to avoid suspicion of their bots was a useful business model.

In the “Erin Princess Baby” case, while the UCO expected suspicion as the commodity, Mr. Plumridge performed cooperation as the commodity. One might expect a child predator to be primed for suspicion, as in the examples from Grant and MacLeod and Chiang in the section above. In such a case, the slip of the UCO masquerading as a 13-year-old girl being “at the office” would have aroused suspicion similar to the transitional phases of the Tinder catfishing data, or the termination of conversations expected in the same above examples. Instead, Mr.

Plumridge apparently did not remark on the slip by the UCO, or other such indicators of his true identity, because in Mr. Plumridge's schema of them both being adult men engaging in age play online, the cooperative thing to do was to let lapses in an overt and mutually understood to be infelicitous performance slide. Indeed, although not in so many words, the court agreed with Mr. Plumridge's schema of events.

In the case of the A Gay Girl in Damascus blog, cooperation apparently was the commodity. MacMaster had purportedly been performing Amina's identity, or some permutation thereof, for a number of years before the blog grew to popularity, and alongside this performative experience, situated Amina's identity well within his own, making it easier to keep many of the details of her background, at least, straight, as they were genuinely his own. At the time of the Damascus Spring, activists such as Amina were part of a community built upon social justice and inclusivity. Community members and allies, unlike users of dating sites or people engaging in age play online, are not necessarily primed to be suspicious of purported allies, especially when being an ally was such a dangerous move for many of them. Suspicion only shifted to the forefront when Amina went missing, and the details of her abduction were dramatic and unrealistic enough (as they were, indeed, unreal) to start raising red flags. It was only retrospectively that people within Amina's community began to pick her story and identity apart.

It is interesting, then, that the catfishing Tinder data was so apparently successful in stringing along some of the 108 users, who, as with Amina, only seemed to pick apart their interlocutor's identity retrospectively, if indeed they did so at all. It is worth noting that these data differ from the other Tinder and dating site examples provided in this chapter, but as with MacMaster's performance as Amina, their performed identities were situated in their own real-world identities, because as far as they knew, they were not performing anyone but themselves. And even though they existed in a community at least somewhat primed to expect bots, they differ in that there was not a program but a person behind the keyboard.

In Part III below, this latter point is particularly relevant, as we will see two contrasting instances of identity performance. In the first, an estranged husband performs as his own wife during her alleged kidnapping and assault, maintaining a real-world identity over text messages. In the second, a woman performs the identity of a non-existent CIA agent named "Chris"—whose identity is an amalgam of a former classmate and pure fantasy—over Facebook and email, and allegedly successfully convinces her boyfriend and father to commit a double homicide on her behalf. In both cases, there was apparently no (or at least not enough) cause for suspicion until it was too late, and the crimes had already been committed.

Chapter 5 – Forensic Analyses

Part 1 – Introduction

The following chapter contains two analyses of forensic data in which gender performance and disguise played a role in the formation or interpretation of the language. As with the previous chapter, both cases are CMC, with the first being a set of text messages, and the second a set of emails and Facebook posts. Both datasets are run through the relevant twenty features proposed above in Chapter 3, and various theories and approaches as discussed throughout Chapters 1 through 3 are applied to the output of these feature distributions.

The first two analyses cover the Hemmert case and apply the relevant twenty features to two different subsets of the data. The first analysis takes all of the husband's language as a whole, as compared to the language of the wife and the questioned text messages sent from her phone, taking into consideration the limited Community of Practice of both authors as husband and wife in contextualizing some of the features. The weighting of these feature findings is then considered, as are other issues as indicated by the feature distributions of the dataset.

The second Hemmert analysis seeks to rectify these issues by reapplying the relevant twenty features, with particular consideration to Audience Design and power, by splitting up the language of the husband between those to his ex-wife, and those to his current estranged wife. A final keyword analysis again considers a shared Community of Practice that differentiates both the husband and wife from the baseline of the more general eBAC, but not necessarily from one another (as they share much of their demographic background).

The third analysis covers the Potter case, by applying the relevant twenty features and already accounting for differences in Audience Design and Accommodation Theory as demonstrated in both Hemmert and the catfi.sh Tinder data, this time between questioned data. This data, which is purportedly from a male CIA agent "Chris", is split between "Chris's" correspondences with Jamie, another male who, in his deposition claims "Chris" performed many of the same male-to-male, hypermasculine features demonstrated in Chapter 4's catfi.sh Tinder After phases, and "Chris's" more hostile language to a wider audience of his and Jenelle's perceived enemies.

Finally, this chapter considers the overall feature weights of all analyses undertaken throughout Chapters 4 and 5 in order to determine whether any particular feature patterns or theoretical applications produced more or less reliable results in predicting gender performance.

Part 2 – Hemmert

The following data come from a case adjudicated in the early 2010s, in which an estranged husband was accused of kidnapping and assaulting his estranged wife. During the

time of the alleged kidnapping, multiple texts were sent from the wife's phone to other family members and friends concerned over her whereabouts, in which the sender insisted everything was fine. In court, the wife claimed her estranged husband wrote the texts while in possession of her phone, and the husband claimed his wife maintained possession of her phone the entire time and wrote the texts herself, but was attempting to frame him for a crime he did not commit. Thus, one of the issues in the case became who likely authored the texts in question.

A. The Data

The data consist of text messages from both the wife and husband (which are referred to here as the known or K datasets), as well as the messages sent from the wife's phone alleged to have been sent by the husband (which is referred to here as the questioned or Q dataset). All of the texts provided in court were sent in the weeks just prior to the questioned texts being sent. As such, many of the texts were sent during the time period leading up to and during which the wife left her husband and attempted to file for divorce. Many of his texts to her were attempts to get her to call him, and were frequently repeated with slight if any variation. Because such texts were not particularly probative, and run the risk of inflating his feature counts due to mindless repetition, many texts were excluded from the analysis.

Such exclusions included any text containing phrases that were repeated 3 or more times without any sufficiently different context, such as "(I) love you" and "call me (back) (please)", and any text with 3 or fewer words, as these are smaller than the smallest texts included in the questioned texts. This largely excluded texts from the husband, and some texts from the wife, as shown in Table 5.1 below. Because the husband's texts still significantly outnumbered the wife's, they were randomly split into four separate datasets of relatively equal word counts, which are closer in size to the wife's dataset.

	Q	Amber	Lael	Lael1	Lael2	Lael3	Lael4
Original Word Count	188	2,962	22,481	--	--	--	--
Word Count	188	2,694	11,418	2,855	2,837	2,863	2,863
Message Count	18	171	764	192	197	193	182
Words per Message	10.4	15.8	14.9	14.9	14.4	14.8	15.7

Table 5.1 Distribution of word counts in Hemmert

As a better comparison to the earlier analyses, the findings below, when normed, were normed to 1,000, which likely conflates the numbers for the much smaller Q dataset, though anything smaller would likely deflate the numbers for many of the other datasets. Thus in many instances the relative frequencies provided by percentages may be a more accurate representation of any two datasets' similarities or dissimilarities, or it may well be that datasets as small as the one in this case are not conducive to such analyses.

Of the twenty features considered in previous analyses, 9 are found for comparison in the Q, 8 are not found in the Q but are still compared below between the K, and 3 are not considered in this analysis. The latter 3 features are emoticons, which were not preserved in the text logs provided to the court, hyperlinks, which are not found in any dataset as they are not as relevant in text messages as in CMC posted online, and keywords, as the Q dataset is so small in comparison to K Lael, and the context of the Q is extremely limited (and likely very

similar due to the timeline).

The following sections consider these individual features and the variations within them, which may prove more probative, as well as how their distributions comport with the findings in earlier analyses.

The data for the Hemmert case (in the form of word lists) is provided in Appendix E.

B. Analysis of Non-Q Features

The following features were not found for comparison in the Q, but can still be considered in their relation between the K documents as markers of their gender. These features include **5. Friendship terms**, **6. Abbreviations**, **7. Punctuation**, **8. Expressive lengthening**, **9. Backchannel sounds**, **10. Hesitation words**, **13. Swears and taboo words**, and **15. Alternative spellings**. It is worth noting that many of these features are commonly associated with CMC, and their absence may be due to the formality of the Q dataset, the age or technological illiteracy of the possible authors, or other such limiting factors.

5. Friendship terms

As demonstrated in Table 5.2 below, neither K dataset had either female- or male-coded friendship terms. That Amber has a higher rate of such words than Lael thus comports with the eBAC findings.

	Amber			Lael		
	#	norm	%	#	norm	%
other (friend)	6	2.2	100%	4	0.4	100%

Table 5.2 Distribution of female- and male-coded friendship terms

6. Abbreviations

The eBAC analysis found that abbreviations in general are female coded, which include the particularly female-coded *lol* and *omg*, both of which are unique to Amber.

	Amber			Lael		
	#	norm	%	#	norm	%
lol	7	2.6	63.6%	0	0.0	0.0%
omg	2	0.7	18.2%	0	0.0	0.0%
other	2	0.7	18.2%	1	0.1	100%
Total	11	4.1	100.0%	1	0.1	--

	Amber	Lael
	#	#
unique	omg, lol	--
shared	fb	

Table 5.3 Overall abbreviations distribution Figure 5.1 Unique abbreviation variations

7. Punctuation

While both Amber and Lael have minimal punctuation (and both only have periods and no other sentence-final punctuation), Amber's higher percentage of punctuated texts relative to Lael's comports with the eBAC findings.

	Amber		Lael	
	#	% of texts	#	% of texts
periods	4	2.3%	10	1.3%

Table 5.4 Overall distribution of punctuation

8. Expressive lengthening

Expressive lengthening is considered by Bamman et al. (2014) to be a female-coded

feature (though as discussed below it appears to be a more male-coded feature in the eBAC). As only K Lael and not also K Amber contain any expressive lengthening, counts are not included here. K Lael most frequently contains instances of *sexy* with two or three x's, and no other expressive lengthening.

9. Backchannel sounds

The eBAC found most backchannel sounds to be female-coded features, which comports with the distribution in Table 5.5 below. Of the possible backchannel sounds, the eBAC found *o* to be a particularly male-coded feature, which also follows with that term being found only in K Lael but not also K Amber, though the remaining variations found in K Lael are more female coded (*oh*, *um*, *yo*).

	Amber			Lael		
	#	norm	%	#	norm	%
oh	4	1.5	100.0%	1	0.1	7.1%
o+	0	0.0	0.0%	10	0.9	71.4%
yo	0	0.0	0.0%	1	0.1	7.1%
um	0	0.0	0.0%	2	0.2	14.3%
Total	4	1.5	--	14	1.2	--

Table 5.5 Overall backchannel sound distribution

	Amber	Lael
unique	--	o, yo, um
shared	oh	

Figure 5.2 Unique backchannel sounds

10. Hesitation words

The eBAC analysis found the overall use of hesitation words to be female coded, though none occur in K Amber as compared to 2 total in K Lael, and those that occur in K Lael are not the more male-coded variety (*er*).

	Amber			Lael		
	#	norm	%	#	norm	%
um	0	0.0	0.0%	2	0.2	100%
Total	0	0.0	0.0%	2	0.2	--

Table 5.6 Overall hesitation word distribution

	Amber	Lael
unique	---	um

Figure 5.3 Unique hesitation words

13. Swears and taboo words

The eBAC found swears in general to be male coded, which comports with the distribution of this feature in Table 5.7 below.

		Amber			Lael		
		#	norm	%	#	norm	%
Female Coded Subtotal		0	0.0	0.0%	0	0.0	0.0%
Male Coded	damn	0	0.0	0.0%	4	0.4	12.5
	fuck	0	0.0	0.0%	6	0.5	18.8%
	shit	5	1.9	83.3%	8	0.7	25%
	bitch	0	0.0	0.0%	3	0.3	7.9%
	ass	1	0.4	16.7%	11	1.0	9.4%
Male Coded Subtotal		6	2.2	100.0%	32	2.8	100%
TOTAL		6	2.2	--	32	2.8	--

Table 5.7 Overall swear and anti-swear word variations

Neither dataset has any instance of anti-swears. Some swears exhibited by the K were found to be more female coded in the eBAC, including *ass* in both datasets, and *bitch* in K Lael (though the latter may be attributable to context, as K Lael is either talking to females or about females, and K Amber is not).

15. Alternative spellings

Alternative spellings in general are found by the eBAC to be female coded, which is consistent with the distribution in Table 5.8 below. However, it is worth noting that neither K dataset has much variety of alternative spellings as compared to that found in, for example, the eBAC.

		Amber			Lael		
		#	norm	%	#	norm	%
Female Coded	lol	7	2.6	63.6%	0	0.0	0.0%
	u	0	0.0	0.0%	2	0.2	28.6%
	w/	0	0.0	0.0%	1	0.1	14.3%
Female Coded Subtotal		0	2.6	63.6%	3	0.3	42.9%
Male Coded	ain't	4	1.5	36.4%	4	0.4	57.1%
TOTAL		11	4.1	--	7	0.6	--

Table 5.8 Overall alternative spellings distribution

C. Overview of Non-Q Features

Table 5.9 provides the overall findings comparing the features of the K documents to Bamman et al. (2014) and Schler et al.'s (2006) findings relative to one another, and the normed eBAC findings.

#	Feature	Bamman and Schler		eBAC	
		Amber	Lael	Amber	Lael
5	overall friendship terms	female	male	female	male
6	overall abbreviations	female	male	male	male
7	overall punctuation	female	male	--	--
8	expressive lengthening*	male	female	--	--
9	backchannel sounds	female	male	female	mixed
	o variations	female	male	female	female
	female coded variations	mixed	female	--	--
10	overall hesitation words	male	female	male	male
13	overall swear words	female	male	male	male
15	overall alternative spellings	female	male	female	male

Table 5.9 Overall feature coding compared to Bamman et al. (2014) & Schler et al. (2006)

Overall, the features appear to accurately differentiate Amber and Lael when considering only who has more of a specific gender-coded feature as compared to the other. Of the features for which the findings do not work, the eBAC results demonstrated the opposite of Bamman et al.'s (2014) findings for expressive lengthening, which follows with what is found in the data; for the other feature, hesitation words, that it does not occur at all for comparison in K Amber may skew the findings.

Table 5.9 also compares these features to normed findings from the eBAC analysis, which proves about as successful for the male-coded features of the husband, but slightly less successful for the female-coded features of the wife.

D. Analysis of Q Features

The following features were found in the Q dataset for comparison to the other two datasets, including **1. Pronouns**, **2. Emotion terms**, **4. Kinship terms**, **11. Assent terms**, **12. Negation terms**, **14. Prepositions**, **16. Conjunctions**, **17. Articles and determiners**, and **18. Message length**. It is worth noting that many of the features not included in this section (and covered in the non-Q features in the sections above) are those more commonly associated with CMC or "chat speak."

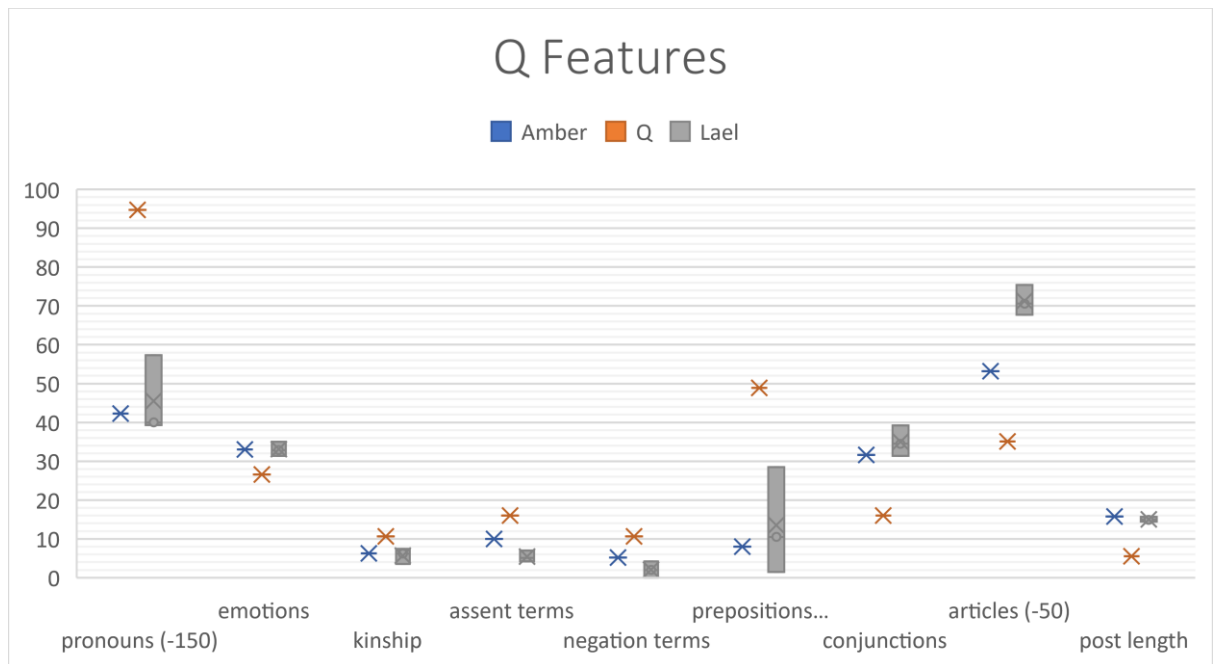


Figure 5.4 Distribution of features found in the Q

As demonstrated in Figure 5.4 above, none of these features for the K appear in range for the Q, although some are closer than others. Pronouns, emotion terms, kinship terms, prepositions, and post length all trend closer between the Q and K Lael to varying degrees. Assent terms, negation terms, conjunctions, and articles and determiners all trend closer between the Q and K Amber. However, in almost every case, the two K datasets trend closer to one another than either does to the Q dataset, which may indeed indicate such analyses of overbroad features are better conducted on larger (or at least more robust) datasets.

1. Pronouns

	Amber		Q		Lael			
					Total		Range	
	#	norm	#	norm	#	norm	#	norm
pronouns	518	192.3	46	244.7	2,161	189.3	544-588	190-207.3

Table 5.10 Overall Pronoun use by person

The eBAC found that pronouns in general, and alternative spellings in particular, are a more female-coded feature. Thus that we find a large number of pronouns in the Q, as demonstrated by Table 5.10 above, may indicate that the Q is female-authored.

However, that we find no examples of alternative spellings of *you*, and minimal examples of the lowercase *i*, as shown in Table 5.11 below, would appear to indicate more male-coded authorship. It is worth noting that Amber, the female author, does not use the apparently more female-coded alternative spellings, while Lael, a male author, does.

	Amber		Q		Lael			
					Total		Range	
	#	%	#	%	#	%	#	%
you your youre you're	84	100%	7	100%	725	99.5%	157-204	98.1-99.5%
u ur yr yur yer ure	0	0.0%	0	0.0%	4	0.5%	1-4	0.5-1.9%
I	156	98.7%	13	81.3%	559	96%	125-151	95-97.8%

i	2	1.3%	3	18.7%	24	4%	3-8	2.2-5%
---	---	------	---	-------	----	----	-----	--------

Table 5.11 Overall distribution of pronouns and their alternative spellings

2. Emotion terms

The eBAC found that the use of emotion terms in general tended to be female-coded. Again, it is worth noting that the variety and uniqueness of emotion terms found in the husband's dataset appear to be more female-coded, though this is likely due to the much larger size of his dataset compared to the other two. That the Q's overall use of emotion terms, as shown in Table 5.12 below, is relatively low as compared to both K, however, is indicative of likely male authorship.

	Amber		Q		Lael			
					Total		Range	
	#	norm	#	norm	#	norm	#	norm
Total Instances	89	33.0	5	26.6	376	32.9	90-100	31.4-35
	Total	Unique	Total	Unique	Total	Unique	Total	
Total Variety	27	8	5	0	59	41	29-36	--

Table 5.12 Overall distribution of emotion terms

4. Kinship terms

The eBAC found that kinship terms in general are female-coded features, with specific kinship terms such as *wife* being more male coded. Thus that the Q overall uses more kinship terms than either K would appear to indicate female authorship. However, because the context of many of the texts has to do with family members and their relationship as husband and wife, it is likely that this feature is not particularly probative in this instance.

	Amber			Q			Lael					
							Total			Range		
	#	1k	%	#	1k	%	#	1k	%	#	1k	%
female	17	6.3	100%	2	10.6	100%	67	5.9	94.4%	8-20	2.8-7	80-100
male	0	0.0	0.0%	0	0.0	0.0%	4	0.4	5.6%	0-2	0-0.7	0-20%
TOTAL	17	6.3	--	2	10.6	--	71	6.2	--	10-21	3.5-7.4	--

Table 5.13 Overall distribution of female- and male-coded kinship terms

11. Assent terms

The eBAC found assent terms in general to be female coded, specifically *okay* and *yes*.

	Amber			Q			Lael					
							Total			Range		
	#	1k	%	#	1k	%	#	1k	%	#	1k	%
yes	3	1.1	11.1%	2	10.6	66.7%	8	0.7	13.8	1-3	0.4-1	7.7-23.1
yeah	0	0.0	0.0%	0	0.0	0.0%	20	1.8	34.5%	4-6	2.1-4	20-46.2
okay	0	0.0	0.0%	0	0.0	0.0%	2	0.2	3.4%	0-1	0-0.4	0-8.3%
ok	12	4.5	44.4%	1	5.3	33.3%	27	2.4	46.6%	5-11	1.7-3.9	41.7-55
yep	6	2.2	22.2%	0	0.0	0.0%	1	0.1	1.7	0-1	0-0.4	0-0.4
yea	6	2.2	22.2%	0	0.0	0.0%	0	0.0	0.0%	0	0	0
TOTAL	27	10.0	--	3	16.0	--	58	5.1	--	12-20	4.2-7	--

Table 5.14 Overall distribution of assent term varieties

Thus that we find more overall use of assent terms in the Q than either K would appear to indicate female-coded authorship.

However, these assent terms may be context-specific in their use, as will be discussed in a later section, such as being prompted in response to yes or no questions. As such, it may then be more probative to consider which assent terms are being used rather than their overall

distribution. The Q contains only *yes* and *ok*, which are within the variations of what the eBAC might expect for a female author. However, Q does not include *yeah* as found in the husband's K, or *yep* and *yea* as found in the wife's K. Thus this dataset is perhaps too small to get a useful distribution of assent term variations for such analysis.

12. Negation terms

The eBAC found negation terms to be somewhat more female coded, with specific negation terms such as *nah*, *nobody*, and *ain't* being male coded. Thus the relatively higher use of negation terms in the Q as compared to either K might indicate a female author, though it is again worth noting that, as with assent terms, there could be outside contexts that prompt their use, such as yes or no questions. No dataset considered in this analysis contains any of the female-coded negation terms such as *noo+* and *cannot*, and only Q contains no instances of the male-coded assent terms found in both of the K.

		Amber		Q		Lael			
						Total		Range	
		#	1k	#	1k	#	1k	#	1k
Male coded	nobody	0	0.0	0	0.0	1	0.1	0-1	0-0.1
	aint	4	1.5	0	0.0	4	0.4	0-3	0-0.3
Subtotal		4	1.5	0	0.0	5	0.4	0-3	0-0.3
Other	no	10	3.7	2	10.6	18	1.6	0-9	0-3.1
Total		14	5.2	2	10.6	23	2.0	0-12	0-4.2

Table 5.15 Overall distribution of negation terms

14. Prepositions

The eBAC found prepositions overall and their alternative spellings to be more male coded (the latter of which are not found here). Thus the overall more frequent use of prepositions in the Q would appear to indicate a male author.

		Amber		Q		Lael			
						Total		Range	
		#	norm	#	norm	#	norm	#	norm
prepositions		291	108.0	28	148.9	1,262	110.5	288-368	101.5-128.5

Table 5.16 Overall distribution of preposition varieties and alternative spellings

16. Conjunctions

		Amber			Q			Lael					
								Total			Range		
		#	1k	%	#	1k	%	#	1k	%	#	1k	%
and		85	31.6	100	3	16.0	100	393	34.4	99.7	90-112	31.4-39.2	99-100
n		0	0.0	0.0%	0	0.0	0.0%	1	0.1	0.3%	0-1	0-0.3	0-1
TOTAL		85	31.6	100	3	16.0	100	394	34.5	--	90-112	31.4-39.2	--

Table 5.17 Distribution of conjunctions and their alternative spellings

The eBAC found all conjunctions to be more female coded; thus this feature would appear to indicate the Q has more likely male authorship, though K Lael has a higher distribution of conjunctions than both K Amber and the Q.

17. Articles and determiners

The eBAC found articles to be a male-coded feature. Thus the relatively low number of articles in the Q as compared to either K would appear to indicate likely female authorship.

	Amber		Q		Lael			
					Total		Range	
	#	1k	#	1k	#	1k	#	1k
articles	278	103.2	16	85.1	1,377	120.6	336-359	117.7-125.4

Table 5.18 Articles and determiners

18. Post length

Not considered by Bamman et al. (2014), Schler et al. (2006) found that longer posts indicated female authorship, while the eBAC indicated male authorship. As Schler et al. (2006) analyzed tweets, which are requisitely shorter than the blog posts of the eBAC and thus more emulative of text messages, the use of much shorter messages in the Q would appear to indicate male authorship.

	Amber	Q	Lael	
			Total	Range
Word Count	2,694	188	11,418	2,837-2,863
Message Count	171	34	764	182-197
Words per Message	15.8	10.4	14.9	14.4-15.7

Table 5.19 Distribution of word counts in Hemmert

E. Overview of Q Features

In cases of authorship analysis, as opposed to simple profiling, it is important to consider any features in the context of the individual potential authors and not simply overall numbers as compared to outside research, as individual authors tend to vary.

#	Feature	Amber	Q	Lael
1	pronouns	female	female	female
	alternative spellings	male	male	male
	lowercase i	male	male	male
2	emotion terms	female	male	female
4	kinship terms overall *	female	female	female
	male coded kinship terms	female	female	male
11	assent terms	female	female	female
12	negation terms overall	female	female	mixed
	male coded negation terms	male	female	male
14	prepositions overall	male	female	male
	alternative prepositions	male	male	male
16	conjunctions overall	female	male	female
	female coded conjunctions	male	male	male
17	articles and determiners	female	female	female
female coded features		8/14	8/14	6/14
male coded features		6/14	6/14	7/14
matched with Q		10/14	--	9/14
matched with Q and gender		6/14	--	4/14

Table 5.20 Overall coding of features as compared to the eBAC

Table 5.20 demonstrates these feature distributions as compared to the eBAC findings, with highlighting indicating a match with either K and the Q, and **bold** indicating a match with the K author's gender. Of the 14 categories considered in the 7 features above, the wife has 8 female- but 6 male-coded features as compared to the eBAC findings, and the husband has 7 male-, and 6 female-coded features, and one that is within range of either (mixed). While the Q has 8 female-coded features and only 6 male-coded features, only one of these features distinguishes between the husband and wife as authors, while 9 match both, and 4 match neither. In fact, the husband and wife appear to be better matches for one another than either

of them is to the Q.

Most of the Q features ending up indicated as female coded in Table 5.20 above are because the features were found to be more common to female authors, and the normed Q counts are much higher than the eBAC counts (and thus much higher than the lower male eBAC counts). Many of these results, then, may indicate that this dataset is indeed too small (either due to the total word count or the short individual message counts averaging 5 to 6 words), or at the very least too context-specific and not robust enough to get accurate normed counts without unhelpfully inflating many of the numbers.

F. Overview

Indeed, in part due to linguistic evidence, but also because of other evidence, the husband in this case was given a prison sentence on charges of attempted strangulation and domestic battery, which occurred during the time in which the wife would have had to have had possession of her phone to have been the author. Thus it would appear that while these features work fairly well in accurately categorizing the gender of the K authors, they did not accurately indicate the gender of the Q author by Bamman et al. (2014) and Schler et al.'s (2006) original findings, nor did they helpfully distinguish the gender of their author based on the eBAC findings.

This, then, would appear to indicate that the 5 Q features that did not match K Lael were successfully guised. However, of the 5, only 2 (male-coded negation terms and overall prepositions) shifted from male coded in K Lael to female coded in the Q, and neither matched the more male-coded preferences of K Amber. Certainly, then, it is hard to see these features as successful (or indeed intentional) guising, either of the generalizations of the wife's gender, or of the wife's own individual patterns. Other than the small size of the Q dataset, other issues with attempting such an analysis on data like this are discussed in the section below.

Part 3 – Hemmert Feature Weights

Because the above analysis, which considered the Q as a cohesive dataset, failed to consistently output a likely gendered authorship based on the 8 features found to occur in the Q, the following considers each text message individually. The figures below also set aside relative frequencies of the occurring features (because, as discussed above, the relatively small size of the total Q dataset is highly likely to lead to artificially inflated numbers), and instead consider the presence and absence of these features as compared to their overall gender-coding.

If a feature occurs and was determined by the eBAC findings to have a more likely gender-coding, its presence is indicated with a matching encoding below; if that feature is absent, this absence is indicated by the opposite encoding. As an example, kinship terms are considered to be female coded; thus, because text 1 below contains *kids*, it is marked as female (F), but because text 2 contains no kinship terms, it is marked as male (M). Any preponderance of gender-coding, even by one, is highlighted in green, and the outcome of that

preponderance is encoded in the rightmost “overall total” column.

Not included in this analysis are conjunctions, for which Bamman et al. (2014) found no significant gender distinction. In addition, the only feature that differs from a presence/absence distinction is message length, which instead has a cutoff of 15 words, with anything 15 or above being female coded, and anything below being male coded. This is because Schler et al. (2006) found longer posts to be indicative of female authorship, and Amber averages 15.8 words per post, while Lael averages 14.9.)

1. Considering Only the Binary Presence and Absence of Features

Q Texts		features								total		
		pronouns	emotions	kinship	assent	negation	prepositions	articles	length	FEMALE	MALE	OVERALL
1	I am going to work things out with Lael we will be over in a bit too get my stuff and the kids	F	M	F	M	M	M	M	F	3	5	M
2	We are going to work it out it over for us I figured out all your doing is using me	F	M	M	M	M	M	F	F	3	5	M
3	Yes I love him	F	F	M	F	M	F	F	M	5	3	F
4	We are working it out im sorry	F	F	M	M	M	F	F	M	4	4	-
5	Yes it is over	F	M	M	F	M	F	F	M	4	4	-
6	You will send my stuff to me and I will have someone come get the kids	F	M	F	M	M	M	M	F	3	5	M
7	no I will not be back i will send someone to get everything	F	M	M	M	F	M	M	M	2	6	M
8	You let me take it but i will call you in a bit	F	M	M	M	M	M	M	M	1	7	M
9	I will call you in a bit we are still talking	F	M	M	M	M	M	M	M	1	7	M
10	I said I will call you in a bit	F	M	M	M	M	M	M	M	1	7	M
11	Lael wants to know if he can come down to	F	F	M	M	M	M	F	M	3	5	M
12	No we are going to work it out	F	M	M	M	F	M	F	M	3	5	M
13	Just stay out of it I am Ok	F	M	M	M	M	M	F	M	2	6	M
14	Don't worry I am fine	F	F	M	M	F	F	F	M	5	3	F
15	Don't worry about it	M	M	M	M	F	M	F	M	2	6	M
16	Why	M	M	M	M	M	F	F	M	2	6	M
17	Lael did not beat me I fell out of the truck I just said that because I was upset	F	F	M	M	F	M	M	F	4	4	-
18	Get your stuff packed and shyanns we are leaving Rick's	F	M	M	M	M	F	F	M	3	5	M
	TOTAL F	16	5	2	2	5	6	11	4	51	-	2
	TOTAL M	2	13	16	16	13	12	7	14	-	96	13

Table 5.21 Binary distinction of only the 8 features present in Q

As demonstrated by the above table, of the total 18 texts, 2 are a majority female coded, 13 are a majority male coded, and 3 display no majority coding (in that they contain an even number of features for both gender-coding). Because this analysis of individual messages considers only the presence and absence of these features, the remaining 8 features above not found in the Q can similarly be considered absences. As demonstrated in the table below, it is worth noting that all of the 8 missing features were found by the eBAC to be more female coded.

	features in Q	features not in Q	total
--	---------------	-------------------	-------

	pronouns	emotions	kinship	assent	negation	prepositions	articles	length	friendship	abbreviations	punctuation	lengthening	backchannel	hesitation	swears	spellings	FEMALE	MALE	OVERALL
1	F	M	F	M	M	M	M	F	M	M	M	M	M	M	F	M	4	12	M
2	F	M	M	M	M	M	F	F	M	M	M	M	M	M	F	M	4	12	M
3	F	F	M	F	M	F	F	M	M	M	M	M	M	M	F	M	6	10	M
4	F	F	M	M	M	F	F	M	M	M	M	M	M	M	F	M	5	11	M
5	F	M	M	F	M	F	F	M	M	M	M	M	M	M	F	M	5	11	M
6	F	M	F	M	M	M	M	F	M	M	M	M	M	M	F	M	4	12	M
7	F	M	M	M	F	M	M	M	M	M	M	M	M	M	F	M	3	13	M
8	F	M	M	M	M	M	M	M	M	M	M	M	M	M	F	M	2	14	M
9	F	M	M	M	M	M	M	M	M	M	M	M	M	M	F	M	2	14	M
10	F	M	M	M	M	M	M	M	M	M	M	M	M	M	F	M	2	14	M
11	F	F	M	M	M	M	F	M	M	M	M	M	M	M	F	M	4	12	M
12	F	M	M	M	F	M	F	M	M	M	M	M	M	M	F	M	4	12	M
13	F	M	M	M	M	M	F	M	M	M	M	M	M	M	F	M	3	13	M
14	F	F	M	M	F	F	F	M	M	M	M	M	M	M	F	M	6	10	M
15	M	M	M	M	F	M	F	M	M	M	M	M	M	M	F	M	3	13	M
16	M	M	M	M	M	F	F	M	M	M	M	M	M	M	F	M	3	13	M
17	F	F	M	M	F	M	M	F	M	M	M	M	M	M	F	M	5	11	M
18	F	M	M	M	M	F	F	M	M	M	M	M	M	M	F	M	4	12	M
	16	5	2	2	5	6	11	4	0	0	0	0	0	0	18	0	69	-	2
	2	13	16	16	13	12	7	14	18	18	18	18	18	18	0	18	-	219	18

Table 5.22 Binary distinction of all 16 features, including those absent in all of Q

The outcome of these features being reintroduced into the analysis is that 18 of the 18 Q texts have a preponderance of male-coded features by a large proportion. Of these 8 additional features, it is worth considering that two, expressive lengthening and hesitation words, were found in the K to have the opposite encoding as compared to the eBAC's findings. That is, the husband had more of the female-coded expressive lengthening and hesitation words than the wife (though it is also worth pointing out that expressive lengthening was found to have the opposite distribution in the eBAC analysis). However, even in considering these features to be more female coded and thus indicative of authorship by the wife (+2F, -2M in every instance) still leaves, in every text message, a preponderance of features as male coded.

Q Texts	features in Q								features not in Q								total		
	pronouns	emotions	kinship	assent	negation	prepositions	articles	length	friendship	abbreviations	punctuation	lengthening	backchannel	hesitation	swears	spellings	FEMALE	MALE	OVERALL
16 Why	M	M	M	M	M	F	F	M	M	M	M	M	M	M	F	M	3	13	M

Table 5.23 Distribution of all 16 binary features in Q text 16, including those absent

However, there exist some issues with making such binary distinctions, one of which can best be exemplified by taking a closer look at Q text 16. Of the 16 features considered in Table 5.23 above, Q 16 can only be considered for two (at a single word, it is better indicative of the shorter lengths expected for males by Schler et al. (2006)), with the remaining 15 features simply marking their absence. At only one word, however, there is no single word environment that could account for the remaining 14 features (additionally excluding punctuation, which is the only other non-lexical feature considered).

That Bamman et al. (2014) and Schler et al. (2006) largely proposed features that are female coded (12 as opposed to 3 male coded), the simple absence of features is bound to trend as male coded, especially in instances such as Q 16 where there do not exist nearly enough environments for many if any of the features to occur.

2. Considering the Relative Distribution of Present Features

Also at issue with such binary distinctions is that each feature is given equal weight, though neither Bamman et al. (2014) nor Schler et al. (2006) described any particular feature as being the most indicative of gender, or at least more indicative of gender than others (though Schler et al. (2006) does refer to style). As such, the above gives primary weight to the absence of features, and so Table 5.24 below considers only the constellation of features present in each individual text, and the preponderance of their overall encodings (as such length being removed from the below findings) in order to mitigate the issue of each feature being equally weighted here.

	features							total counts			% of total		
	pronouns	emotions	kinship	assent	negation	prepositions	articles	FEMALE	MALE	OVERALL	FEMALE	MALE	OVERALL
1	F	-	F	-	-	M	M	2	2	-	40%	100%	M
2	F	-	-	-	-	M	-	1	1	-	20%	50%	M
3	F	F	-	F	-	-	-	3	0	F	60%	0%	F
4	F	F	-	-	-	-	-	2	0	F	40%	0%	F
5	F	-	-	F	-	-	-	2	0	F	40%	0%	F
6	F	-	F	-	-	M	M	2	2	-	40%	100%	M
7	F	-	-	-	F	M	M	2	2	-	40%	100%	M
8	F	-	-	-	-	M	M	1	2	M	20%	100%	M
9	F	-	-	-	-	M	M	1	2	M	20%	100%	M
10	F	-	-	-	-	M	M	1	2	M	20%	100%	M
11	F	F	-	-	-	M	-	2	1	F	40%	50%	M
12	F	-	-	-	F	M	-	3	1	F	40%	50%	M
13	F	-	-	-	-	M	-	1	1	-	20%	50%	M
14	F	F	-	-	F	-	-	3	0	F	60%	0%	F
15	-	-	-	-	F	M	-	1	1	-	20%	50%	M
16	-	-	-	-	-	-	-	0	0	-	0%	0%	-
17	F	F	-	-	F	M	M	3	2	F	60%	100%	M
18	F	-	-	-	-	-	-	1	0	F	20%	0%	F
totals	16	5	2	3	5	-	-	31	-	8	34.4%		5
	-	-	-	-	-	12	8	-	20	4		55.6%	11

Table 5.24 Distribution of present features in each text only, and % of total gendering

At issue with the total counts of these features is the fact that the male total cannot exceed 2, while the female total can (but notably does not go higher than 3) go up to 5. By this measure, 8 of the texts are female coded, 3 of the texts are male coded, and the remaining 7 are mixed. As such, Table 5.24 also considers the % of these features to be found, the result of which are 5 female-coded texts, 11 male-coded texts, 2 texts with mixed encodings with one of each leaning male and female (because the uneven number of female features considered cannot match the 50% of only one of the two male-coded features occurring, instances in which only 2 or 3 of the 5 female-coded features occur with only a difference of 10% are still

considered mixed), and 1 with no features.

3. Considering both binary results and relative distribution

Thus it may be useful to consider how each of these binary test variations' results vary (if they do), and why this might occur. The overall determinations are demonstrated in Table 5.25 below. This excludes the presence and absence of all features in Table 5.24, as it skewed overwhelmingly and likely unhelpfully male, and the counts of only present features which skewed overwhelmingly and likely unhelpfully female, though the percentages of the latter are maintained.

	binary features present and absent in all Q	the percentage of only present features in each individual Q	FEMALE	MALE	OVERALL
1	M	M	0%	100%	M
2	M	M	0%	100%	M
3	F	F	100%	0%	F
4	-	F	75%	25%	F
5	-	F	75%	25%	F
6	M	M	0%	100%	M
7	M	M	0%	100%	M
8	M	M	0%	100%	M
9	M	M	0%	100%	M
10	M	M	0%	100%	M
11	M	M	0%	100%	M
12	M	M	0%	100%	M
13	M	M	0%	100%	M
14	F	F	0%	100%	F
15	M	M	0%	100%	M
16	M	-	25%	75%	M
17	-	M	0%	100%	M
18	M	F	50%	50%	-

Table 5.25 Overall non-skewed findings (binary in Q and % individual present w/ context)

It is worth pointing out that, save for Q text 18, the only features that differ between the four tests are ones in which at least one other test found the outcome to be inconclusive but not contradictory. Thus it may be that both tests in conjunction offer the most robust and consistent results, with 12 of the 18 texts having the same outcome in all four of the tests considered (and of the remaining 6, 5 have a preponderance of male or female over the other). As such, it is worth considering why these 6 texts, shown below, might offer differing results.

4	We are working it out im sorry
5	Yes it is over
14	Don't worry I am fine
16	Why
17	Lael did not beat me I fell out of the truck I just said that because I was upset
18	Get your stuff packed and shyanns we are leaving Rick's

Table 5.26 Q texts with mixed results across non-skewed tests

Already discussed above is 16, which, due to its length, has minimal environments in which any of the considered features could occur. This would appear to indicate that length can be a factor in successfully applying such analyses. At the very least that, for a text to be so short, such an analysis is only applicable if more than two features can be considered as present (in this case post length and punctuation). As an example, a hypothetical single word text of *lol* would show the presence of some considered features (both abbreviations and alternative spellings) as opposed to their absence (though this also demonstrates an absence

of, for example, expressive lengthening), and a text that simply said *fuck* or *darn* would show the presence of either a swear or an anti-swear (and thus the absence of the opposite). So it may be that while length can be problematic, the real issue with a message like 16 is the lack of robustness—the absence of environments—and not a set amount of additional words.

This may also be true of 5, which along with 16, make up two of the four shortest texts in the Q dataset. As discussed above with regards to context, both of these texts contain assent terms that, when provided context, may not be the best matches for the feature considered by the eBAC analysis, as they are not naturally occurring so much as they are prompted by a yes or no question. Thus 5 is left only with pronouns, emotion terms, post length, and punctuation, a trait shared by Q text 4 even though it has almost twice the words. Again, it would appear that not word counts but a robustness of environments is the most useful in such analyses for giving clearer results.

Of course, the findings of Bamman et al. (2014) and Schler et al. (2006), as well as the findings of these features as compared to the eBAC, are not absolutes, but a trend in variations. So it may be that a preponderance of text messages, in this case, is enough to demonstrate likely male authorship of most if not all of the 18 Q texts.

4. Other Issues

While the application of the features to each individual message gave much less mixed results with regards to gender-coding than was found in the previous analysis that considered the relative distribution of these features, neither analysis considers the context of these features. As an example, the below four Q texts contain either an assent term (*yes*) or negation term (*no*), both of which are considered female-coded features by Schler et al. (2006) (though negation terms are more mixed in Bamman et al. (2014) and the eBAC).

3	Yes I love him
5	Yes it is over
7	no I will not be back i will send someone to get everything
12	No we are going to work it out

Table 5.27 Q texts with assent and negation terms

Comparatively, K Amber did indeed have a higher rate of both assent (10.0) and negation (5.2) terms than did K Lael (5.1 and 2.0, respectively), and so this feature in particular would appear to be a good match for K Amber. However, the four Q texts, in context, vary from the assent and negation terms as they are found in K Amber (and K Lael), in that they are not a response to a yes or no question, but are ‘spontaneous’ uses of these terms. K Amber, contrastingly, has sentence-initial assent and negation terms only in response to yes or no questions.

Thus, in K Amber, this feature was prompted, and if both *yes* and *no* are considered female coded, there is no response the texter could have given to a binary question that would not have been female coded. This, then, differs from the spontaneous use of unprompted assent and negation terms, though Bamman et al. (2014) and Schler et al. (2006), in looking at much larger (and also much less conversational, in that they are not forensic) datasets do not consider any such distinctions.

Of course, just as it was not realistic for Bamman et al. (2014) or Schler et al. (2006) to consider context in their research (and in the BAC, which is much less conversational than the tweets considered by Schler et al. (2006), which may themselves have only been conversational at times, such prompting as by yes or no questions is unlikely to have occurred save by the author themselves), it may not always be realistic for a forensic analyst either.

5. Overall Encoding as Compared to the K Datasets

Finally, this section considers how the above Q findings map to the actual distribution of features by K Amber and K Lael, which, as demonstrated in the previous analysis, do not always match the gender-coding that the eBAC would expect.

	Binary Distribution											K Set Matches		
	features								total					
	pronouns	emotions	kinship	assent	negation	prepositions	articles	length	FEMALE	MALE	OVERALL	K Amber	K Lael	OVERALL
K Amber	M	-	-	F	F	F	F	-	4	1	F	-	-	-
K Lael	F	-	-	M	M	M	M	-	1	4	M	-	-	-
1	F	M	F	M	M	M	M	F	3	5	M	0	5	K Lael
2	F	M	M	M	M	M	F	F	3	5	M	1	4	K Lael
3	F	F	M	F	M	F	F	M	5	3	F	3	2	K Amber
4	F	F	M	M	M	F	F	M	4	4	-	2	3	K Lael
5	F	M	M	F	M	F	F	M	4	4	-	3	2	K Amber
6	F	M	F	M	M	M	M	F	3	5	M	0	5	K Lael
7	F	M	M	M	F	M	M	M	2	6	M	1	4	K Lael
8	F	M	M	M	M	M	M	M	1	7	M	0	5	K Lael
9	F	M	M	M	M	M	M	M	1	7	M	0	5	K Lael
10	F	M	M	M	M	M	M	M	1	7	M	0	5	K Lael
11	F	F	M	M	M	M	F	M	3	5	M	1	4	K Lael
12	F	M	M	M	F	M	F	M	3	5	M	2	3	K Lael
13	F	M	M	M	M	M	F	M	2	6	M	1	4	K Lael
14	F	F	M	M	F	F	F	M	5	3	F	2	3	K Lael
15	M	M	M	M	F	M	F	M	2	6	M	3	2	K Amber
16	M	M	M	M	M	F	F	M	2	6	M	3	2	K Amber
17	F	F	M	M	F	M	M	F	4	4	-	1	4	K Lael
18	F	M	M	M	M	F	F	M	3	5	M	2	3	K Lael
Q	F	M	M	M	M	M	F	M	2	6	M	2	3	K Lael
K Amber	2	-	-	2	5	6	11	-	-	-	-	-	-	4
K Lael	16	-	-	16	13	12	7	-	-	-	-	-	-	14

Table 5.28 Relative distribution of K Amber & K Lael compared to binary distribution in Q

Table 5.28 maps K Amber and K Lael relative to one another, with 5 of the below 7 features showing a distinction from each other. (Emotion terms and kinship terms are too close to differentiate, as K Amber falls within the range of K Lael, and length is not clearly comparable here.)

Of the 18 total texts, 14 are a better match with K Lael, and the remaining 4 are a better match with K Amber, suggesting either that Amber indeed had access to her phone toward the beginning and end of the text messages, or that K Lael successfully guised his language in these 4 texts. Notably, however, these all occur in a 3/2 split in favor of K Amber, while of the remaining texts, 5 have a 4/1 split and 5 have a 5/0 split in favor of K Lael. This would suggest instead that Lael either made no attempt to guise his features or was otherwise unsuccessful.

3	Yes I love him
5	Yes it is over
15	Don't worry about it
16	Why

Table 5.29 Q Texts found to match K Amber's patterns

Further, in considering the content of the 4 Q texts that better matched K Amber's patterns, shown above, we find that they are, yet again, the four shortest texts in the Q dataset. 3 and 5 have already been discussed explicitly above for the issues of context with the assent term feature in this case, and 16 has already been explicitly discussed above for being a single word with minimal environments. As such, the distribution of features in these four messages may come down not to a failure or inconsistency in Lael's performance, but a failure of this method when applied to such sparse data.

Of the five features considered, only one matched K Amber more than half the time. Given that this feature is articles and determiners (which are unlikely to be conscious gender-coded features), and the above-discussed likelihood that Lael either made no attempt or otherwise failed to guise his language, it is unlikely that this is indicative of articles and determiners as particularly probative features for cases of gender guising.

Indeed, unlike the Tinder data discussed in Parts I and II above, there do not appear to have been any features in the Q texts that demonstrated the same dramatic shift toward the performed gender (in this case female, in the case of the Tinder data male) as found, for example, with the content word features of friendship terms and swear words. The only outliers for Q occurred in function word features, which, again, are less likely to be consciously manipulated. This, too, would lend credence to the idea that Lael made no overt attempts to guise his language on the basis of his and his wife's genders. Part 4 below, however, explores how the language of K Lael *does* vary in a performative manner based on his audience.

Part 4 – Hemmert 2

It is worth noting in this case that the 20 gendered features considered in the earlier analysis did not indicate a significantly male skewing as compared to the eBAC counts, with only 6 of the total 14 features coming out as male coded. However, as discussed throughout the Part above, this is true for K Lael as well, in which only 7 of the total features were male coded, and the remaining 7 were female or mixed. This, then, would appear to indicate that as compared to the average author in the eBAC, Lael is somewhat evenly distributed between male- and female-coded features (although of the 5 comparable features found in both K Lael and the Q, he did skew male as compared to K Amber 4 out of 5 times, even if not also as compared to the eBAC).

While it is important to consider that the eBAC analysis features are not absolutes and will not always hold true for every author given a specific gender, and that variation occurs within the eBAC and it is not necessarily a representative corpus of every author of English, other issues may be at play in this case aside from K Lael simply deviating from previous findings incidentally. The data for K Lael is comprised of texts, and largely of those sent to

Amber, his estranged wife with whom he pled to work on their marriage. Such a scenario brings to mind not gendered roles, but the distribution of power.

As discussed in the earlier review on approaches to the study of gendered language in Chapter 2, a reconciliation of the deficit approach to the study of gendered language (that women's language is somehow deficient as compared to men's language, which was considered the standard of measurement) by the dominance approach involved examining the language of women in the context of power. That is, features that are considered more female coded, or are otherwise socialized to be imposed upon female language production, are not indicative of some sort of inherent difference between the genders, but of the usual differences in position men and women have as compared to one another. Thus that men are often in more powerful positions than women because of the way English-speaking societies tend to be structured means that "male-coded" features are more likely "power-coded" features.

In the case of being a husband trying to reach out to his estranged wife who wants nothing to do with him or their marriage, Lael can thus be considered to be in a position of powerlessness with regards to fixing their marriage as compared to Amber. Though most of the texts in K Lael *are* to Amber to discuss their marriage, some of the texts are between Lael and a previous ex-wife, against whom Lael makes multiple threats and derogatory comments, and thus Lael potentially has a different power differential between her as compared to that which he has with Amber. Thus this analysis considers the differences between these two subsets of K Lael, in order to test whether the gender-coding of the eBAC features shifts between the two datasets.

A. The Data

Table 5.30 below demonstrates the distribution of K Lael between the two subsets of data, as well as the data excluded (which was texted to other parties aside from Lael's wife and ex-wife).

	Original K Lael	"To Amber"	"To Haley"	To Others
Words	22,481	17,040	2,283	3,158
Texts	1,872	1,405	147	320
Words per text	12.0	12.1	15.5	9.9

Table 5.30 Distribution of K Lael data into new data subsets

This data is not missing any of the text excised from edited K Lael, and so includes all of the repetitive language (largely in "To Amber"). Because the "To Amber" subset is so much larger than the "To Haley" subset, frequencies in the analysis below are normed to 1,000 words. Because texts to others are mixed with regards to the audience in question (some of whom Lael looks to for sympathy, others of which he threatens, all related to his estranged marriage), this third subset is not directly compared here.

B. The Analysis

As with the above analysis of both K datasets in comparison to the Q Texts, this analysis does not include emoticons as they were either not used or not preserved in the data provided to the court, or hyperlinks which are not found in any dataset as they are not relevant

in text messages as they are in CMC posted online. Despite the disparate sizes of the datasets here, keywords will be compared in this instance, along with the remaining 17 features considered in the analysis above.

Table 5.31 shows the overall distribution of each feature in the “To” datasets as compared to the Q Texts and K Lael, with **gray** indicating the closest range to the Q Texts. Where available, **bold** text indicates counts in the “To” datasets that exist outside the range of the four K Lael subsets in the analysis above.

#	Feature	Q Texts	To Amber		To Haley		K Lael edit
		norm or %	#	norm or %	#	norm or %	norm or %
1	overall pronouns	244.7	3,715	218.0	442	193.6	189.3
	standard <i>you</i> variations	100%	1,298	98.3%	142	97.3%	99.5%
	non-standard <i>you</i> variations	0.0%	22	1.7%	4	2.7%	0.5%
	uppercase <i>I</i>	81.3%	949	96.6%	96	95.0%	96%
	lowercase <i>i</i>	18.7%	33	3.4%	5	5.0%	4%
2	total emotion terms	26.6	957	56.2	70	30.7	32.9
	total variety	5	86	--	26	--	58
4	female-coded kinship terms	10.6	70	3.9	16	7.1	5.9
	male-coded kinship terms	0.0	7	0.7	1	0.4	0.4
	total kinship terms	10.6	78	4.6	17	7.5	6.2
5	non-coded friendship terms (<i>friend</i>)	2.2	4	0.2	0	0.0	0.4
6	total abbreviations	0.0	4	0.2	2	0.9	0.1
7	texts with periods	0.0%	25	1.7%	2	1.3%	1.3%
8	expressive lengthening	0.0	12	0.7	0	0.0	0.1
9	backchannel sounds	0.0	33	1.9	5	2.2	1.2
10	hesitation words	0.0	1	0.1	1	0.4	0.2
11	assent terms	16.0	128	7.5	18	7.9	5.1
12	male coded negation terms	0.0	3	0.2	2	0.9	0.4
	non-coded negation terms	10.6	22	1.3	3	1.3	1.6
	total negation terms	10.6	26	1.5	5	2.2	2.0
13	swears and taboo words	0.0	34	2.0	15	6.6	2.8
14	total prepositions	148.9	1,718	100.8	246	107.8	110.5
15	alternative spellings	0.0	5	0.3	2	0.9	0.6
16	conjunctions	16.0	545	32.0	86	37.7	34.5
17	articles	85.1	1,698	99.7	284	124.4	120.6
18	text length	15.8	--	12.1	--	15.5	14.9

Table 5.31 Overall feature findings in “To” datasets

Of the 6 overall features found to exist in “To Amber” outside the range of K Lael (overall pronouns, total emotion terms, assent terms, total prepositions, articles, and text length), all 6 trend more toward female coded than their overall use in K Lael.

Table 5.32 below demonstrates the trend of features as compared to their overall use in K Lael, with **gray** indicating a “match” to the expected power differential, and “--” indicating that the feature falls within the range of the K Lael set (and so does not overtly vary in the given “To” circumstance).

#	Feature	Compared to K Lael		Compared to each other	
		To Amber	To Haley	To Amber	To Haley
1	overall pronouns	female	--	female	male
	non-standard <i>you</i> variations	--	female	male	female
	lowercase <i>i</i>	--	--	male	female
2	total emotion terms	female	male	female	male
	total variety	--	--	female	male

4	female-coded kinship terms	--	--	male	female
	male-coded kinship terms	--	--	male	female
	total kinship terms	--	female	male	female
5	non-coded friendship terms (<i>friend</i>)	male	male	female	male
6	total abbreviations	female	female	male	female
7	texts with periods	female	--	female	male
8	expressive lengthening	female	male	female	male
9	backchannel sounds	female	female	male	female
10	hesitation words	male	female	male	female
11	assent terms	female	female	male	female
12	male coded negation terms	--	male	female	male
	non-coded negation terms	--	--	--	--
	total negation terms	--	--	male	female
13	swears and taboo words	female	male	female	male
14	total prepositions	female	--	female	male
15	alternative spellings	male	female	male	female
16	conjunctions	--	--	--	--
17	articles	female	--	female	male
18	text length	female	--	female	male
Found in the Q	Total male coded	3	5	11	11
	Total female coded	11	7	11	11
Total "correct/incorrect" to gender		3/11	7/5	--	--
Total "correct/incorrect" to power		11/3	5/7	--	--

Table 5.32 Comparison of "To" features to K Lael and each other

In both cases, more features are more female coded than in the total K Lael than they are male coded, but this difference is starkest in the "To Amber" set. In addition to the 6 features discussed above, 5 more (total abbreviations, texts with periods, expressive lengthening, backchannel sounds, and swears and taboo words) trend more female coded in "To Amber". While it would appear that powerlessness may be a factor in the larger number of female-coded features in "To Amber" as compared to K Lael (and as compared to "To Haley" as compared to K Lael), the same does not necessarily seem to hold for any overt power in "To Haley".

Table 5.33 below provides binary findings as compared to the eBAC counts, in which **gray** indicates overlap with the Q Texts, and **bold** indicates overlap with the gender of the K Author. While K Lael and "To Haley" each share 8 matching features out of the 13 found in the Q Texts, "To Amber" shares one less at 7. This would appear to indicate that, at least given the features that are found in the Q Texts, gendered features in K Lael that are motivated by power do not play a significant role in the analysis and comparison of available gendered features.

#	Feature	Q Texts	To Amber	To Haley	K Lael	K Amber
1	overall pronouns	female	female	female	female	female
	non-standard <i>you</i> variations	male	male	male	male	male
	lowercase <i>i</i>	male	male	male	male	male
2	total emotion terms	male	female	male	female	female
4	male-coded kinship terms	female	female	male	male	female
	total kinship terms	female	female	female	female	female
5	non-coded friendship terms (<i>friend</i>)	--	male	male	male	female
6	total abbreviations	--	male	female	female	female
7	texts with periods	--	--	--	--	--
8	expressive lengthening	--	mixed	male	male	male
9	backchannel sounds	--	female	female	female	female

10	hesitation words	--	male	male	male	male
11	assent terms	female	female	female	female	female
12	male coded negation terms	female	male	male	male	male
	total negation terms	female	male	male	mixed	female
13	swears and taboo words	--	male	male	male	male
14	total prepositions	male	female	female	male	male
15	alternative spellings	--	male	male	male	female
16	conjunctions	male	female	female	female	female
17	articles	female	female	female	female	female
18	text length	male	female	male	male	female
Found in the Q	Total male coded	6	5	7	6	5
	Total female coded	7	7	6	6	8
Total "correct/incorrect" to gender		--	9/10	12/8	11/8	13/7
Total "correct/incorrect" to power		--	10/9	12/8	11/8	13/7

Table 5.33 Binary distribution as compared to the eBAC

This may be due to the fact that, of the 11 features found to be female coded in "To Amber" as compared to K Lael, 5 are not found in the Q Texts for comparison. The takeaway, however, is the fact that powerlessness can have an effect on the distribution of gendered features for a given author, and as such should be taken into consideration when attempting such binary analyses for authorship purposes.

In addition, these findings appear to largely comport with the previous analysis of Tinder data. That analysis found that several features considered in the eBAC analysis to be overall female coded in use occurred more during the Before phase, and less in the After phase.

#	Feature	Before	Compared to K Lael		After
			To Amber	To Haley	
2	total emotion terms	female	female	male	male
4	total kinship terms	female	--	female	male
5	non-coded friendship terms (<i>friend</i>)	female	male	male	male
8	expressive lengthening	female	female	male	male
9	backchannel sounds	female	female	female	male
10	hesitation words	female	male	female	male
11	assent terms	female	female	female	male
12	male coded negation terms	female	--	male	male
	total negation terms	male	--	--	female
13	swears and taboo words	female	female	male	male
Total matching Tinder data		--	5 of 7	5 of 9	--

Table 5.34 Comparison of "To" datasets to Tinder Before/After phases

Similarly, that analysis found that a number of male-coded features occurred more in the After phase, and less in the During phase. These 10 features are compared to the "To" datasets below. As discussed above, context may rule out some features as less useful, and less accurate depending on the datasets. These are, for the most part, things like kinship terms, which are discussed more in "To Haley" in the logistics of an estranged marriage involving children, and friendship terms, which are discussed more in the K Lael data to people (other than Amber and Haley) who are friends. This is true to a lesser extent of assent and negation terms, which, especially in smaller datasets, may simply depend on the answer to a yes or no question, and not a gendered preference. Other features such as hesitation words, backchannel sounds, and expressive lengthening can only be so helpful in that they are rarely used in this case, and thus there is not much differentiation to be had.

terms) leaves “To Amber” to largely match the Before phase of the Tinder data, and “To Haley” to largely match the After. The findings in this section, then, would appear to comport with the previous Tinder analysis, demonstrating that outside factors—whether they be powerlessness as in K Lael, overt cooperation as in the Tinder data, or as simple as audience design in either—can affect the distribution of some features. In particular, between both analyses, emotion terms and backchannel sounds appear to be overtly female coded and swears appear to be overtly male.

C. Keyword Analysis

Excluding proper names, typos, non-English words, and hapax legomena, the top twenty keywords for K Lael as compared to the Male eBAC are *please*, *sexxxy*, *unblock*, (*sexxxy*), (*sexy*), *you*, *talk*, *call*, *love*, *beautiful*, *me*, *yeppers*, *baby*, *chance*, (*ttys*), *ok*, *very*, *give*, *much*, *straightened*, *marriage*, *bunches*, and *horse*. Many of these keywords fit into the female-coded categories of expressive lengthening (*sexy*, *sexxxy*), emotion terms (*love*), assent terms (*ok*, *yeppers*), and pronouns (*you*, *me*), all of which are found to be overall female-coded features in K Lael. Other words appear as part of repeated phrases, such as “love you [beautiful/baby/sexy]/[bunches]/[very much]” (*love*, *you*, *beautiful*, *baby*, *sexy*, *bunches*, *very*, *much*), “please [*call me*]/[give me a chance]/[unblock me]” (*please*, *call*, *me*, *give*, *chance*, *unblock*), “talk to you later/soon” (*talk*, *you*, *ttys*), and “(get everything/shit) straightened out” (*straightened*). (The only remaining words are *marriage*, which is an overarching theme of all messages in K Lael, and *horse*, which is part of a repeated threat in K Lael to sell Amber’s horse that often co-occurs with other repeated keywords, as in “**Call me** let’s **talk** about the **horse**”.)

The top 20 keywords in K Amber are (*ttys*), *rodeos*, *texting*, *rodeo*, *paperwork*, *divorce*, *number*, *horse*, *ok*, *stamps*, *pickup*, *filing*, *ranch*, *yea*, *won*, *am*, *gas*, *he*, *calling*, *you*, and *cause*. Most of these words can be seen as context-driven as they are related to specific events and locations that are more common in rural areas that the eBAC is not confined to (*rodeo*, *rodeos*, *ranch*, *pickup*), their marriage (*paperwork*, *divorce*, *filing*), phones as the mode of communication between them (*texting*, *calling*) and other such contextual terms without any gender-coding (*stamps*, *won*, *gas*, *cause*). That does leave *horse*, in the same context as it is found in K Lael, as well as assent terms (*yea*, *ok*) and pronouns (*he*, *you* and to some extent *am* collocating with *I*), both of which are generally considered to be female coded, and are also found to be female coded in K Amber.

That assent terms appear as keywords in both K Lael and K Amber as compared to the eBAC Male and Female subcorpora respectively may then be indicative of text messages, which are conversational, being more prone to responses to yes or no questions, which blogs are not.

Thus it would appear that keywords may be a helpful tool in determining which of the 20 gendered features might be outliers specific to an author’s linguistic patterns, and not necessarily indicative of an author of the opposite gender when found to trend similarly in any

Q datasets. It would also appear that keywords can provide insight into repeated words or phrases that are specific to the context of the K data, and thus not a blanket indicator of gender, but context. This can be seen to hold true for the female-coding of kinship terms in K Lael, which, while they do not appear in the top 20 keywords, still occur with a keyness above 6.8 (*wife, child, kids*, etc.) in K Lael.

D. Overview

As this case progressed in court, the wife realized that she may have herself sent some of the 18 text messages marked as Q and sent on the day of the alleged assault, though which of the 18 texts this may have applied to was not specified to the court. However, in such instances where the husband sending any amount of the texts from her phone during this time frame could be considered evidence for the prosecution, that these findings demonstrate not a back and forth, but a chunk of consistently male-coded texts (6-13) sandwiched between mixed and female-coded texts may still prove valuable.

Overall, the features included in these above analyses proved most consistent in considering overall datasets when they were applied to the larger K datasets and considered as a matter of distributions both against the other K dataset, and against the findings of the eBAC. However, in the case of a smaller (and possibly mixed) dataset like the Q, the results of these features proved most consistent when applied to individual texts, with context included, and by looking both at the binary distribution of all 16 features, as well as the percentage distribution of possible male and female-coded features as compared to one another.

Whether these findings hold true is considered in the following Part, which deals with a much larger dataset for both the Q and the K, though, as we will see below, has its own host of difficulties, as real-world forensic cases often do. Regardless, the preponderance of texts appeared to best match the patterns of K Lael, even when his patterns differed from what the eBAC findings would have expected. As discussed previously, the distribution of features would appear to make it unlikely that Lael actively attempted to disguise his language, at least on the basis of gender, though both his audience (and thus performance) and the context of the data may have played a notable role in the distribution of features therein.

Part 5 – Potter

The following analysis covers the language of a murder case. The victims, a couple, were alleged to have harassed the Defendant, Jenelle, online, to the point that her friend in the CIA, “Chris”, facilitated their murders. According to *State of Tennessee v. Jenelle Leigh Potter*:

“This case involves the murders of two victims ... For these crimes, a ... grand jury indicted the Defendant’s father ... (“Buddy”), the Defendant’s mother ... (“Barbara”), the Defendant, and ... (“Jamie”).

(*State v. Potter*, No. E2015-02261-CCA-R3-CD (Tenn. Crim. App. Feb. 5, 2019))

The Defendant was the person who first told Jamie about “Chris.” She said “Chris” was a friend of the family and “like a brother to her.” The Defendant said

that she and "Chris" were the same age and that he had lived next door to her in Pennsylvania. She said that, when she was in high school, "Chris" would take her to school and pick her up after track practice. The Defendant told Jamie that "Chris" worked for the CIA on "cases and stuff" and that "Chris" had a house in Tennessee and one in Pennsylvania, and he "would go back and forth."

(*State v. Potter*, No. E2015-02261-CCA-R3-CD (Tenn. Crim. App. Feb. 5, 2019))

"Chris" introduced himself to Jamie as a "a friend of the Defendant's[.]" and he told Jamie that he worked for the CIA. "Chris" did not have an email account, so he would contact Jamie through text or by email sent to Jamie's cell phone from the Defendant's email address. Jamie never spoke to "Chris" on the phone or met him in person. "Chris" told Jamie that he had a "phobia of phones" and did not like using them. The Defendant told Jamie that "Chris" also communicated with her through email.

(*State v. Potter*, No. E2015-02261-CCA-R3-CD (Tenn. Crim. App. Feb. 5, 2019))

When Jamie would receive an email from the Defendant's email address from "Chris," he could identify "Chris" as the author based, in part, on the words used in the emails. The emails from "Chris" addressed Jamie in a manner that the Defendant did not. For example, "Chris" would begin emails with "hey man, or hey, dude, how's it going[.]" According to Jamie, the Defendant never cursed, called people names, or spoke hatefully to people; however, "Chris" would "rant and rave about everything" in his emails. He cursed, called people names, and wished harm on others. "Chris" would also typically sign his emails as "Chris."

(*State v. Potter*, No. E2015-02261-CCA-R3-CD (Tenn. Crim. App. Feb. 5, 2019))

Jamie provides a number of linguistic and metalinguistic cues that, in his perception of their correspondences, differentiated "Chris" and Jenelle. While he does not explicitly state these differences are (at least to his understanding) inherent to their genders, the distinguishing features he notes do comport with the eBAC findings, and the findings of this thesis thus far: that male-coded kinship terms (like "bro" terms) such as *man* and *dude*, and swears and taboo words in general are both male coded, and, to Jamie, indicative of "Chris's" language and not Jenelle's.

These features are not the only finding that appears to comport with the Tinder data, which demonstrated the difficulty interlocutors had separating their conversational partner from the profile of their purported identity even in the face of "red flags" or other linguistic or informational cues to the contrary. Although "Chris" does not have a profile attached to his messages, which come from Jenelle's own Facebook account, according to Jamie, "Chris would also typically sign his emails as 'Chris.'" According to the same Opinion:

The Defendant identified "Chris" and "Matt Potter" on Facebook as her "brothers"; however, Christie Groover and the Defendant are Buddy and Barbara's only children. Law enforcement officers were unable to locate any CIA agent named "Chris" working in Johnson County but located a Chris Tjaden, as identified in the Defendant's writing, in Delaware.

(*State v. Potter*, No. E2015-02261-CCA-R3-CD (Tenn. Crim. App. Feb. 5, 2019))

And:

When shown a photograph attached to an email sent from the Defendant's email address to Barbara's email address on September 13, 2011, Mr. Tjaden identified himself in a photograph, recognizing it as one of his profile pictures

from his Facebook account. Mr. Tjaden identified the same photograph in an email sent from Barbara's email address to Barbara's email address with the subject line "Pic of Chris-Barb cropped & enlarged/plus auto adjust/contrast/brightness..5X7[.]" Mr. Tjaden identified himself in two additional photographs, once at a baseball game and the other standing next to his patrol car, in an email sent from the Defendant's email address to Jamie's email address on October 22, 2011, with the subject line, "Me Chris." Mr. Tjaden explained that he had previously used both photographs as profile pictures on Facebook.

(*State v. Potter*, No. E2015-02261-CCA-R3-CD (Tenn. Crim. App. Feb. 5, 2019))

The use of Mr. Tjaden's photographs and name on "Chris's" Facebook profile and emails purportedly from "Chris", then, served much the same function as the fake profiles on Tinder. Unlike the Tinder data, however, this offered identity was then bolstered by an author actively performing the identity of "Chris" to multiple people (and then corroborating witnesses that "Chris" existed), leading to the linguistic cues Jamie detailed above.

As detailed in the conversational analysis of the Tinder data, what often broke the "immersion" for conversees and called the purported identity of their interlocutors into question was the same sort of thing that ultimately lost the case for the police in the "Erin Princess Baby" decoy case in Australia: language containing facts or statements inconsistent with the purported identity of their author.

In the end, as summarized in the same opinion above:

Jamie explained that he asked investigators about CIA involvement because Buddy had told him "he was with the CIA." He said that he was "hoping the CIA had [his] back" and that "Chris" was real; however, he acknowledged that "Chris" did not exist and never existed. Jamie testified that he was manipulated by the Potter family. He explained, "Well, I mean, I thought Chris was real. I mean, I thought that there was a, you know, someone that I was talking to there and the Defendant the way she would talk to me . . . it was like a -- a bonding . . . a family. And it's like it's all a lie."

(*State v. Potter*, No. E2015-02261-CCA-R3-CD (Tenn. Crim. App. Feb. 5, 2019))

A. The Data

According to *State v. Potter* (in which the defendant was not Janelle, but her mother Barbara):

Investigators learned that the email account associated with Defendant was bmp9110@aol.com, the email account associated with Jenelle was BUL2DOG@aol.com, and the email account associated with Jamie was sleepiingbear@yahoo.com. Investigators obtained subpoenas for records from Yahoo and AOL in relation to the aforementioned email accounts, and Agent Lott received a DVD containing thousands of emails.

(*State v. Potter*, No. E2015-02262-CCA-R3-CD (Tenn. Crim. App. Feb. 5, 2019))

Although the data was provided as emails, the contents of the emails were largely comprised of Facebook messages copied and pasted into the emails. Almost all of these were posted by Jenelle Potter's account, and signed by either Jenelle, Barbara, or "Chris".

Table 5.35 demonstrates the distribution of both the language known to have been written by Jenelle, and the language purported to have been written by "Chris". ("Chris" also allegedly went by both Cody and Matt, but only the messages signed "Chris" are included in J. Ford, PhD, Thesis, Aston University, 2022

the data and analysis below.) Based on the difference in feature distribution of K Lael in the Hemmert case above dependent upon his recipient, questioned “Chris” is split between those correspondences sent to Jenelle, her mother, and others on Facebook, and those sent to Jamie.

Dataset Name	message approximation	Word Count
Known Jenelle	83	11,521
Questioned “Chris”	46	26,107
“To Jamie”	18	3,513

Table 5.35 Data distribution

The numbers of messages are approximated due to the fact that many messages were sent multiple times, and sometimes by both Jenelle and Barbara, but were sometimes included with multiple other messages, sent by themselves, or split across emails.

The data for the Potter case (in the form of word lists) is provided in Appendix F.

B. The Analysis

This case differs from the above analyzed Hemmert case on a number of points. The first major difference is the fact that the Q Data is not only substantially larger than the Q Texts in Hemmert, but also more than three times larger than the body of K Data. The second is that there is no second set of K Data to compare the Q Data to—there was no real world “Chris from the CIA” the way there was a real-world Amber Hemmert to compare—and as such the analysis more largely relies on the ranges provided by reference corpora such as the eBAC.

1. Pronouns

Table 5.36 demonstrates overall pronouns in Potter, and these findings as compared to the eBAC. Other than the lower distribution of non-standard lowercase *i*, both Q “Chris” and “To Jamie” match K Jenelle as compared to the eBAC. It is worth noting that K Jenelle does not seem to prefer non-standard pronouns in general, which while at a more male-coded rate than the eBAC would predict, is still consistent with both Q.

	Chris			to Jamie			Jenelle		
	#	norm	%	#	norm	%	#	norm	%
first person	1,872	71.7	39.7%	186	53.0	26.1%	1,127	97.8	51.7%
second person	721	27.6	15.3%	202	57.5	28.4%	466	40.5	21.4%
third personal	2,120	81.2	45.0%	324	92.2	45.5%	589	51.1	27.0%
TOTAL use	4,713	180.5	--	712	202.7	--	2,182	189.4	--
			Chris		to Jamie		Jenelle		
			#	%	#	%	#	%	
you, your, youre, you're			716	99.3%	196	97.0%	459	98.5%	
u, ur, yr, yur, yer, ure			5	0.7%	6	3.0%	7	0.5%	
I			993	73.3%	66	45.2%	377	55.0%	
i			361	26.7%	80	54.8%	308	45.0%	
			As compared to the eBAC						
			Chris		to Jamie		Jenelle		
1	pronouns		female		female		female		
	alternative spellings		male		male		male		
	lowercase i		male		female		female		

Table 5.36 Pronouns overview in Potter

Of particular interest is the difference, above, in the distribution between first, second, and third pronouns across the three sets, with Jenelle strongly favoring first person, “Chris”

slightly favoring third person over first, and “To Jamie” more clearly favoring third person because of the more even distribution between first and second persons. This would appear to comport with the agenda as ascribed by the prosecution in this case, of “Chris” building interpersonal rapport with Jenelle’s family and defending Jenelle from perceived online bullying (“I” and “you”), and Chris “To Jamie” also encouraging his relationship with Jenelle (“she”), and serving the purpose of attempting to convince Jamie (“you”) to be complicit in the eventual murders.

	Chris			to Jamie			Jenelle		
	#	norm	%	#	norm	%	#	norm	%
he/him/his/himself	454	17.4	31.5%	10	2.8	4.3%	109	9.5	35.0%
she/her/hers/herself	666	25.5	59.5%	204	58.1	95.7%	202	17.5	65.0%

Table 5.37 Distribution of gendered pronouns in Potter

This is especially evident in the distribution of gendered pronouns, where female-gendered pronouns go from a relatively similar majority in K Jenelle and Q “Chris” to an almost exclusive use in “To Jamie”. Thus while these features appear to be overall matches across all three datasets, this demonstrates that it is also useful to look at the context of the data to check whether any apparent anomalies in the nuance of the features are in fact anomalous.

2. Emotion terms

	Chris		to Jamie		Jenelle	
	#	norm	#	norm	#	norm
Total Instances	1,045	40.0	223	63.5	474	41.1
	Chris		to Jamie		Jenelle	
	All	Unique	All	Unique	All	Unique
Total Variety	74	22	40	4	61	13
As compared to the eBAC						
			Chris	to Jamie	Jenelle	
2	emotion terms		female	female	female	

Table 5.38 Overview of emotion terms in Potter

Both K Jenelle and Q “Chris” have almost the exact same normed distribution of emotion terms, with the strongest spike in “To Jamie”, all of which are female-coded distributions as compared to the eBAC.

	Chris	to Jamie	Jenelle
shared	alone, bad, best, better, crazy, fine, good, happy, hate, hope, hurt(s), like, love(d/s), mad, moved, sad, scared, sick, sure, unhappy, upset, want(ed/ing), well, worried		
	hated, low, safe, scares, wondering		
		welcome	
	careful, funny, guilty, hate(s/ful), high, kind, liked, lost, number, odd, sorry, strong, tired, trust, upsetting, wants, worse	--	careful, funny, guilty, hate(s/ful), high, kind, liked, lost, number, odd, sorry, strong, tired, trust, upsetting, wants, worse
unique	angry, awful, hated, higher, hoping, hungry, hurting, longer, losing, loving, lucky, moving, pain, proud, resting, shame, sicker, sleepy, touched, trusted, warm, worries	blue, merry, restless, weak	annoying, aware, cheating, concerned, confused, depressed, forgive, happier, lazy, peaceful, serious, thankful, threatening

Figure 5.5 Shared and unique emotion terms in Potter

3. Emoticons

Although it is unclear from the data whether emojis (icons) were preserved or simply not used, the same, singular emoticon (text) was found in both K Jenelle and Q “Chris”, ‘:’). This was found only 3 times in Q “Chris”, and 5 times in K Jenelle, making it more abundant per word in K Jenelle (2,304) than in Q “Chris” (8,702). This trend would follow what Bamman et al. (2014) and Schler et al. (2006) would suggest for emoticon use with a higher frequency in the female (K Jenelle) than the male (Q “Chris” and “To Jamie”) datasets, but the paucity of the feature is hardly enough to go on. Rather, it is more interesting to note, in this case, the overall lack of emoticons, and the fact that only one variety is found and shared between the two datasets with any emoticons at all.

4. Kinship terms

As shown in Table 5.39 below, all three datasets strongly prefer female-coded kinship terms over male-coded. This can be largely attributed to the context of this case, in which Jenelle’s parents actively played a part, and many of the communications were between “Chris” and Barbara, who “Chris” referred to as “mom”.

	Chris		to Jamie		Jenelle	
	#	%	#	%	#	%
Female	212	95.1%	41	100.0%	94	96.9%
Male	11	4.9%	0	0.0%	3	3.1%

Table 5.39 Distribution of kinship terms in Potter

Worth noting, then, may be the non-standard variants of the female-coded terms that can be argued to be more stereotypical of female language users. *Mommy*, *daddy*, *hubby*, and *bf* are all found in K Jenelle, and both *daddy* and *bf* are found in Q “Chris”. Both occur with a frequency in their respective datasets with only a 0.1 difference, with Jenelle the more frequent user.

		Chris			to Jamie			Jenelle		
		#	norm	%	#	norm	%	#	norm	%
Female Coded	mom	70	2.7	31.4%	21	6.0	51.2%	18	1.6	18.6%
	mommy	0	0.0	0.0%	0	0.0	0.0%	2	0.2	2.1%
	mother	15	0.6	6.7%	1	0.3	2.4%	3	0.3	3.1%
	sister	16	0.6	7.2%	4	1.1	9.8%	6	0.5	6.2%
	daughter	4	0.2	1.8%	1	0.3	2.4%	1	0.1	1.0%
	aunt	12	0.5	5.4%	0	0.0	0.0%	1	0.1	1.0%
	grandmother	5	0.2	2.3%	0	0.0	0.0%	2	0.2	2.1%
	kids	14	0.5	6.3%	0	0.0	0.0%	4	0.4	4.1%
	child	1	0.1	0.5%	0	0.0	0.0%	2	0.2	2.1%
	dad	37	1.4	16.6%	10	2.9	24.4%	34	3.0	35.1%
	daddy	4	0.2	1.8%	0	0.0	0.0%	3	0.3	3.1%
	father	0	0.0	0.0%	0	0.0	0.0%	3	0.3	3.1%
	husband	6	0.2	2.7%	0	0.0	0.0%	3	0.3	3.1%
	hubby	0	0.0	0.0%	0	0.0	0.0%	1	0.1	1.0%
	brother	16	0.6	7.2%	4	1.1	9.8%	3	0.3	3.1%
	boyfriend	6	0.2	2.7%	0	0.0	0.0%	6	0.5	6.2%
	bf	2	0.1	0.9%	0	0.0	0.0%	2	0.2	2.1%
	uncle	3	0.1	1.4%	0	0.0	0.0%	0	0.0	0.0%
	cousin	1	0.1	0.5%	0	0.0	0.0%	0	0.0	0.0%
Female Coded Subtotal		221	8.1	95.1%	41	11.7	100%	94	8.1	96.9%

Male Coded	wife	6	0.2	2.7%	0	0.0	0.0%	2	0.2	2.1%
	girlfriend	2	0.1	0.9%	0	0.0	0.0%	0	0	0.0%
	gf	1	0.1	0.5%	0	0.0	0.0%	1	0.1	1.0%
Male Coded Subtotal		9	0.4	4.0%	0	0.0	0.0%	3	0.3	3.1%
Other	sis	2	0.1	0.9%	0	0.0	0.0%	0	0.0	0.0%
Other coded Subtotal		2	0.1	0.9%	0	0.0	0.0%	0	0.0	0.0%
TOTAL		223	8.5	--	41	11.7	--	97	8.4	--

Table 5.40 Overall kinship term use in Potter

Thus it may be that such particular variations, in situations where kinship terms may be driven by context rather than simply gender, can be better indicators of likely shared gender, and thus better indicators of likely shared authorship.

		As compared to the eBAC			As compared to Jenelle	
		Chris	to Jamie	Jenelle	Chris	to Jamie
4	kinship terms overall *	female	female	female	female	female
	male coded kinship terms	male	female	male	male	female

Table 5.41 Kinship term comparison in Potter and the eBAC

5. Friendship terms

The friendship terms as shown in Table 5.42 below are of particular interest considering that Jamie pointed them out (along with swears) as being one of the factors that convinced him that “Chris” was not Jenelle but was male. Interestingly, “Chris” has almost the exact distribution of overall friendship terms as found for males in the eBAC (“Chris’s” 1.2 versus the eBAC’s 1.3), while Jenelle and “To Jamie” are both significantly higher than even the eBAC’s female distribution (1.6). While friendship terms overall are considered a female-coded feature, other than the general term “friend”, all three datasets use exclusively male and no female coded friendship terms. (The distribution of what previous research suggested to be male-coded friendship terms was much closer, with males using it 0.19 times to females’ 0.18, though female-coded features were a larger difference at 0.05 to 0.12. Both the male and female eBAC made a higher frequency use of male than female-coded kinship terms, which comports with the findings of “To Jamie” and Jenelle below.)

		Chris			to Jamie			Jenelle		
		#	norm	%	#	norm	%	#	norm	%
Female Coded	bestie	0	0.0	0.0%	0	0.0	0.0%	0	0.0	0.0%
	bff	0	0.0	0.0%	0	0.0	0.0%	0	0.0	0.0%
	best friend	0	0.0	0.0%	0	0.0	0.0%	0	0.0	0.0%
Female Coded Subtotal		0	0.0	0.0%	0	0.0	0.0%	0	0.0	0.0%
Male Coded	bro	0	0.0	0.0%	4	1.1	16.0%	0	0.0	0.0%
	bruh	0	0.0	0.0%	0	0	0.0%	0	0.0	0.0%
	brah	0	0.0	0.0%	0	0	0.0%	0	0.0	0.0%
	brotha	0	0.0	0.0%	0	0	0.0%	0	0.0	0.0%
	pal	0	0.0	0.0%	0	0	0.0%	0	0.0	0.0%
	dude	0	0.0	0.0%	4	1.1	16.0%	2	0.2	3.6%
	mate	0	0.0	0.0%	0	0.0	0.0%	0	0.0	0.0%
	man	4	0.2	13.3%	16	4.6	64.0%	3	0.3	5.5%
Male Coded Subtotal		4	0.2	13.3%	24	6.8	96.0%	5	0.5	9.1%
Other	friend	26	1.0	86.7%	1	0.3	4.0%	50	4.3	90.9%
Total		30	1.2	--	25	7.1	--	55	4.8	--

Table 5.42 Overall friendship terms in Potter

Thus the comparatively high use in “To Jamie” would appear to be the most strongly male coded of the three datasets for kinship terms, which would comport with Jamie’s explanation

of events; but this feature may also be indicative of an over-exaggeration of a feature typically understood as highly male coded, as was found in the Tinder data (even if the eBAC findings would dispute this lay understanding).

6. Abbreviations

The eBAC findings demonstrate that abbreviations are overall more frequent in the female- than in male-coded data (0.7 as compared to 0.3), with *lol/Imao* variations accounting for roughly two-thirds of all abbreviations in both subcorpora of the eBAC.

	Chris			to Jamie			Jenelle		
	#	norm	%	#	norm	%	#	norm	%
Total (lol/Imao)	79	3.0	90.1%	6	1.7	100	22	1.9	92.6%
Total (other)	8	0.3	9.9%	0	0.0	0.0%	5	0.4	7.4%
Total (all)	87	3.3	--	6	1.7	--	27	2.3	--
Unlengthened	78	3.0	98.7%	6	1.7	100	22	1.9	100.0%
Lengthened	1	0.0	1.3%	0	0.0	0.0%	0	0.0	0.0%
Total	79	2.6	--	6	1.7	--	22	1.9	--

Table 5.43 Overall abbreviations in Potter

As demonstrated in Table 5.43 above, the frequency of use of abbreviations is much higher in all three datasets (and thus more resembling female-coded uses), and with a much higher frequency of *lol/Imao* variations than found in the eBAC (though these variations did account for slightly less in the male and slightly more in the female subcorpora, as reflects “Chris” and Jenelle above). Both the distribution of the *lol/Imao* and other variations, as well as their frequency of use, are much higher than those found in the eBAC, and thus would follow the findings of them as more highly female coded. This is increasingly true in “Chris” as compared to both Jenelle and “To Jamie”.

Also worth noting are what non-*lol/Imao* variations are used. As demonstrated in Table 5.43 above, there is only a single use of expressively lengthened abbreviations in any dataset; notably, this female-coded feature variation occurs in “Chris”, the alleged male author. Additionally, Figure 5.6 below demonstrates the limited range of abbreviations used between the datasets (to Jamie used none), all of which are shared by both Chris and Jenelle, other than a single instance of fb for Facebook (the unabbreviated form of which is found once in Chris and twice in Jenelle).

	Chris	Jenelle
unique shared	gf, bf, idk	
unique	--	fb

Figure 5.6 Unique and shared abbreviations in Potter

Thus it may be that not only are the overall and relative distributions of such features useful as indicators of likely gender, but that the precise variations used may help in a determination of likely common authorship.

7. Punctuation

All three datasets follow the same preference of punctuation use, preferring periods,

question marks, ellipses, and exclamations from most to least. This differs from what is found in the eBAC, in which both datasets prefer periods, ellipses, exclamations, and then question marks from most to least, with a much more varied distribution (73.1% periods in the male and 67.4% in the female subcorpus as compared to the 88.7-96% in these datasets). Thus this preference of use, shared by all three datasets, may be indicative of a single author's pattern of use as compared to the general use in the eBAC.

	Chris		To Jamie		Jenelle	
	#	%	#	%	#	%
Total Periods	1,779	95.1%	237	96.0%	631	88.7%
Total Ellipses	29	1.6%	2	0.8%	23	3.2%
Total Exclamations	4	0.2%	0	0.0%	9	1.3%
Total Questions	59	3.3%	8	3.2%	48	6.8%
Total All	1,871	--	247	--	711	--

Table 5.44 Overall punctuation in Potter

The words per punctuation below do not account for null punctuation (which is also true of the eBAC), and as such the relatively equal distribution between "Chris" and "To Jamie" as compared to Jenelle may be due to the genre of email in the former as compared to the genre of direct messages in the latter, which tend to be more formal and more informal, respectively. All three datasets have more words per punctuation (and thus presumably longer sentences) than in the eBAC, existing above the higher end of male coded (13.0, while the female is 11.9). Whether this difference is indicative of authorial preferences or genre is unclear.

	Chris	To Jamie	Jenelle
Words per Punctuation	14.0	14.2	16.2

Table 5.45 Words per punctuation in Potter

Finally, the use of lengthened punctuation, as found in Table 5.46 below, is found by the eBAC to be more female coded. While lengthened punctuation is most frequently used in Jenelle by a large percentage, all three datasets, even "To Jamie", use lengthened punctuation more frequently than either the male (7.4%) or female (9.8%) eBAC subcorpora. This, again, may thus be indicative of an authorial preference, as well as register (with Jenelle, again, being the more informal of the datasets), rather than a strictly gendered distribution.

	Chris		To Jamie		Jenelle	
	#	%	#	%	#	%
standard punctuation	53	68.8%	8	88.9%	22	36.7%
lengthened punctuation	24	32.3%	1	11.1%	38	62.3%

Table 5.46 Standard and lengthened punctuation in Potter

8. Expressive lengthening

The eBAC find expressive lengthening to be an overall female-coded feature, which appeared to hold true for abbreviations, assent terms, negation terms, hesitation words, and punctuation (while backchannel sounds were significantly higher in the male subcorpus). Of these categories, only backchannel sounds and punctuation are found to be lengthened across all three datasets (with "Chris" containing an additional lengthened *lol* abbreviation), as demonstrated in Table 5.47 below. That the datasets appear to make both much less frequent use of expressive lengthening overall would appear to make them more male than female (because of the exceptionally high use of backchannel sounds in the male eBAC, which

accounts for 1.4 of the 1.42 used); this also holds for the absence of non-backchannel expressive lengthening (which accounts for 0.05 of the female subcorpus and 0.02 of the male, and none or essentially none of the three datasets below).

	Chris			to Jamie			Jenelle		
	#	norm	%	#	norm	%	#	norm	%
Lengthened Backchannel	4	0.2	11.8%	1	0.3	50.0%	2	0.2	9.1%
Lengthened Abbreviations	1	0.0	1.3%	0	0.0	0.0%	0	0.0	0.0%
Lengthened Assent	0	0.0	0.0%	0	0.0	0.0%	0	0.0	0.0%
Lengthened Negation	0	0.0	0.0%	0	0.0	0.0%	0	0.0	0.0%
Lengthened Words TOTAL	4	0.3	0.02%	1	0.3	0.02%	2	0.2	0.02%
Lengthened Hesitation	0	0.0	0.0%	0	0.0	0.0%	0	0.0	0.0%
Lengthened Punctuation	24	--	32.3%	1	--	11.1%	38	--	62.3%

Table 5.47 Overall expressive lengthening in Potter

However, in considering which categories are lengthened versus which categories are not, this feature's distribution among the three datasets may again be more indicative of an authorial preference over a gendered indicator. This also holds true for punctuation, which, as discussed above, is much more frequently lengthened in both "Chris" and Jenelle (and to a lesser extent "To Jamie") than in the eBAC.

9. Backchannel sounds

All three datasets have a much lower normed distribution of backchannel sounds than the eBAC (5.2 in the male and 2.7 in the female), which indicated that overall use was actually more male coded. However, discounting the abundance of *o/oh/oo* variations, this feature would appear to be more female coded in the eBAC (at 0.9 for males and 1.3 for females), with 83.3% of male instances being *o/oh/oo* variations as compared to 52.8% of female variations. Jenelle and "Chris", at least, seem to follow more closely the female eBAC pattern of a roughly 50/50 split between non- and *o/oh/oo* variations. Additionally, the non-*o/oh/oo* variations have a very close normed frequency in all three datasets.

	Chris			to Jamie			Jenelle		
	#	norm	%	#	norm	%	#	norm	%
<i>o, oh, oo</i> variations	15	0.6	41.2%	0	0.0	0.0%	13	1.1	59.1%
other variations	19	0.7	58.8%	2	0.6	100.0%	9	0.8	40.9%
Total	34	1.3	--	2	0.6	--	22	1.9	--
Unlengthened	30	1.2	88.2%	1	0.3	50.0%	20	1.7	90.9%
Lengthened	4	0.2	11.8%	1	0.3	50.0%	2	0.2	9.1%
Total	34	1.3	--	2	0.6	--	22	1.9	--

Table 5.48 Overall backchannel sounds in Potter

This is also noteworthy in that "Chris" and Jenelle share both the only *o/oh/oo* variation of *oh* (with no expressive lengthening), and most of the non-*o/oh/oo* variations, with only *grr* expressively lengthened in any dataset (and only *grr* also shared by "To Jamie"). Thus it would appear that the exact backchannel sounds used may be helpful in determining likely common authorship.

	Chris	Jenelle
shared	oh, ugh, aww, grr +, eww	

unique	er	hmm
--------	----	-----

Figure 5.7 Shared and unique backchannel sounds in Potter

In determining the likely gender of either author, however, it is worth mentioning that the backchannel sound exclusive to “Chris”, *er*, is also found to be more male coded in the eBAC, while the backchannel sound found in Jenelle, *hmm*, is found to be more female coded. Given the otherwise nearly identical overlap of backchannel sound varieties shared by the two datasets, however, it seems unlikely that this was a calculated effort in gender guising with regards to backchannel sounds as a whole.

10. Hesitation words

As with the above backchannel sounds, though the eBAC found that hesitation words were overall a female-coded feature, *er* variations in particular were male coded, accounting for both a higher normed frequency than in the female subcorpus, and a higher proportion of overall hesitation word use. Thus the exact variations of this feature would appear more male coded in “Chris”, and ambiguous in Jenelle (*hm* variations were more frequent as a normed feature in the female subcorpus of the eBAC, but accounted for a higher distribution of the variations in the male subcorpus, and was the most frequently used variation by both, at 47.5% in the male and 46.6% in the female). The male coding also holds true for the lack of expressive lengthening, though this is consistent across all datasets, including Jenelle. The comparative frequency of the hesitation words in “Chris” and Jenelle, however, are higher than in the eBAC (0.4 in the male and 0.6 in the female), and thus appear more female coded, while their absence in “To Jamie” altogether can be seen as more male coded, though this may simply be an artefact of the size of the dataset.

	Chris			to Jamie			Jenelle		
	#	norm	%	#	norm	%	#	norm	%
um variations	0	0.0	0.0%	0	0.0	0.0%	0	0.0	0.0%
uh variations	0	0.0	0.0%	0	0.0	0.0%	0	0.0	0.0%
er variations	3	1.2	100%	0	0.0	0.0%	0	0.0	0.0%
erm(h) variations	0	0.0	0.0%	0	0.0	0.0%	0	0.0	0.0%
hm variations	0	0.0	0.0%	0	0.0	0.0%	1	0.9	100%
hum variations	0	0.0	0.0%	0	0.0	0.0%	0	0.0	0.0%
TOTAL	3	1.2	--	0	0.0	--	1	0.9	--

Table 5.49 Overall hesitation words in Potter

11. Assent terms

The eBAC found assent terms to be generally female coded. All three datasets more frequently use assent terms than the eBAC (at 1.7 for the male subcorpus and 2.3 for the female), which may in part be due to the more conversive nature of emails and direct messages than blog posts. More notable, then, may be the distribution of varieties, as all three datasets contain only *yes* and *ok* variations, with no expressive lengthening, and no others (not even the common interchange between *ok* and *okay*).

	Chris			to Jamie			Jenelle		
	#	norm	%	#	norm	%	#	norm	%
yes variations	74	2.8	74.8%	5	1.42	41.67	29	2.5	55.8%

yeah variations	3	0.1	3.0%	0	0.0	0.0%	3	0.3	5.8%
yas variations	0	0.0	0.0%	0	0.0	0.0%	0	0.0	0.0%
yis variations	0	0.0	0.0%	0	0.0	0.0%	0	0.0	0.0%
okay variations	0	0.0	0.0%	0	0.0	0.0%	0	0.0	0.0%
ok variations	22	0.8	22.2%	7	1.99	58.33	20	1.7	38.5%
yeh variations	0	0.0	0.0%	0	0.0	0.0%	0	0.0	0.0%
yep variations	0	0.0	0.0%	0	0.0	0.0%	0	0.0	0.0%
yup variations	0	0.0	0.0%	0	0.0	0.0%	0	0.0	0.0%
yea variations	0	0.0	0.0%	0	0.0	0.0%	0	0.0	0.0%
TOTAL	99	3.8	--	12	3.4	--	52	4.6	--

Table 5.50 Overall assent terms in Potter

12. Negation terms

The eBAC found that negation terms were overall female (at 2.6 for males and 2.7 for females). While Jenelle is the only dataset with a female-coded variant, Jenelle and “To Jamie” are closer to the male eBAC frequency, and “Chris” is closer to the female.

		Chris			to Jamie			Jenelle		
		#	norm	%	#	norm	%	#	norm	%
Female Coded	noo+	0	0.0	0.0%	0	0.0	0.0%	0	0.0	0.0%
	cannot	0	0.0	0.0%	0	0.0	0.0%	2	0.2	6.9
Female Coded Subtotal		0	0.0	0.0%	0	0.0	0.0%	2	0.2	6.9
Male Coded	nah	0	0.0	0.0%	0	0.0	0.0%	0	0.0	0.0%
	n+a+h+	0	0.0	0.0%	0	0.0	0.0%	0	0.0	0.0%
	nobody	0	0.0	0.0%	0	0.0	0.0%	0	0.0	0.0%
	aint	0	0.0	0.0%	0	0.0	0.0%	0	0.0	0.0%
Male Coded Subtotal		0	0.0	0.0%	0	0.0	0.0%	0	0.0	0.0%
Other	no	70	2.7	100	8	2.3	100	27	2.3	93.1
Total		70	2.7	--	8	2.3	--	29	2.5	--

Table 5.51 Overall negation terms in Potter

Perhaps more indicative, then, is the absolute lack of any variations other than no for no (noo+, nah+), and the shared lack of expressive lengthening. Their distributions are also fairly close and, as with assent terms, may be down to the genre of more conversive language (thus possibly answering more yes/no questions) than the eBAC.

13. Swears and taboo words

The eBAC showed a small distinction between the male (2.08) and female (2.05) subcorpora, and a distinction between the male (0.05) and female (0.09) subcorpora for taboo words. This appears to hold true for the distinction between Jenelle and “Chris” as well, at least for swear words.

		Chris			to Jamie			Jenelle		
		#	norm	%	#	norm	%	#	norm	%
Female Coded	darn	0	0.0	0.0%	0	0.0	0.0%	0	0.0	0.0%
	darnit	0	0.0	0.0%	0	0.0	0.0%	0	0.0	0.0%
	dang	1	0.1	0.3	0	0.0	0.0%	5	0.4	7.8
	dangit	0	0.0	0.0%	0	0.0	0.0%	0	0.0	0.0%
	gosh	1	0.1	0.3	0	0.0	0.0%	2	0.2	3.1
	gosh(other)	0	0.0	0.0%	0	0.0	0.0%	0	0.0	0.0%
Female Coded Subtotal		2	0.1	0.7	0	0.0	0.0%	7	0.6	10.9
Male Coded	damn	54	2.1	17.9	0	0.0	0.0%	13	1.1	20.3
	damnit	0	0.0	0.0%	0	0.0	0.0%	0	0.0	0.0%
	dam	0	0.0	0.0%	0	0.0	0.0%	0	0.0	0.0%
	dammit	0	0.0	0.0%	0	0.0	0.0%	0	0.0	0.0%
	dam(other)	13	0.5	4.3	0	0.0	0.0%	0	0.0	0.0%
	fuck	7	0.3	2.3	0	0.0	0.0%	2	0.2	3.1

	fuck(other)	84	3.2	27.9	4	1.1	57.1	6	0.5	9.4
	shit	53	2.0	17.6	1	0.3	14.3	1	0.1	1.6
	shit(ending)	3	0.1	1	1	0.3	14.3	0	0.0	0.0%
	fag	0	0.0	0.0%	0	0.0	0.0%	0	0.0	0.0%
	fag(other)	0	0.0	0.0%	0	0.0	0.0%	0	0.0	0.0%
	bastard	0	0.0	0.0%	0	0.0	0.0%	0	0.0	0.0%
	bastard(other)	0	0.0	0.0%	0	0.0	0.0%	0	0.0	0.0%
	bitch	24	0.9	8.0	0	0.0	0.0%	10	0.9	15.6
	bitch(other)	0	0.0	0.0%	0	0.0	0.0%	0	0.0	0.0%
	ass	39	1.5	13.0	0	0.0	0.0%	10	0.9	15.6
	ass(other)	9	0.3	3.0	0	0.0	0.0%	11	1.0	17.2
	arse	0	0.0	0.0%	0	0.0	0.0%	0	0.0	0.0%
	arse(other)	0	0.0	0.0%	0	0.0	0.0%	0	0.0	0.0%
	cunt	0	0.0	0.0%	0	0.0	0.0%	0	0.0	0.0%
	cunt (other)	0	0.0	0.0%	0	0.0	0.0%	0	0.0	0.0%
	whore	13	0.5	4.3	1	0.3	14.3	4	0.4	6.3
	whore (other)	0	0.0	0.0%	0	0.0	0.0%	0	0.0	0.0%
Male Coded Subtotal		299	11.5	99.3	7	2.0	100%	57	5.0	89.1
TOTAL		201	11.5	--	7	2.0	--	64	5.6	--

Table 5.52 Swears and taboo words in Potter

Anti-swears, on the other hand, remain very consistent between “Chris” and Jenelle. For both anti-swears and swears, “Chris” and Jenelle use all of the same variations of swears (save no *bullshit* in Jenelle, but *shit* in both), and “To Jamie” contains some of the shared swears, but no examples of any of the non-shared swears. That is, their variety of uses seems consistent, though it is worth pointing out that, contrary to Jamie’s observations about “Chris’s” language, the “To Jamie” dataset contains a smaller distribution and variety than Jenelle. (It may also be that Jenelle’s language to Jamie as herself better curates her swearing, though that is not included in the data analyzed here.)

14. Prepositions

The eBAC found prepositions to be more male coded, and alternative prepositions to be more female coded. All three datasets below contained lower uses of prepositions than both the male (129.0) and female (125.1) eBAC subcorpora. Though their distributions are relatively similar, and no dataset uses any of the more female coded alternative spellings, it is worth noting that “Chris” can actually be considered to have the most male coded distribution, while “To Jamie” has the most female coded.

	Chris		to Jamie		Jenelle	
	#	norm	#	norm	#	norm
Alternative Prepositions	0	0.0	0	0.0	0	0.0
Prepositions Subtotal	2,920	111.9	374	106.5	1,262	109.5

Table 5.53 Overall prepositions in Potter

15. Alternative spellings

The eBAC found that alternative spellings tend to be more frequent for female authors, as with additional terms suggested by other features to be male coded. The alternative spellings in all three datasets are higher than those found in the eBAC (1.3 in the male and 1.8 in the female subcorpora), but are more consistent with the female-coded distribution. This feature perhaps shows the least consistency across the three datasets, in that “Chris” has the unique alternative spelling *yaay*, “To Jamie” is the only dataset with *bro* and the only male-

coded alternative spelling, and Jenelle is the only to contain *ur* (though both “Chris” and “To Jamie” have at least an instance of the non-contracted *u*). However, as all of these varieties were found to be female coded in the eBAC, the likely gender of all three would remain consistent.

		Chris			to Jamie			Jenelle		
		#	norm	%	#	norm	%	#	norm	%
Female Coded	vacay	0	0.0	0.0%	0	0.0	0.0%	0	0.0	0.0%
	yaay	1	0.1	1.4	0	0.0	0.0%	0	0.0	0.0%
	lol	71	2.7	97.3	6	1.7	50	19	1.7	82.6
	u	1	0.1	1.4	2	0.6	16.67	3	0.3	13.0
	ur	0	0.0	0.0%	0	0.0	0.0%	1	0.1	4.4
	yr	0	0.0	0.0%	0	0.0	0.0%	0	0.0	0.0%
	w/	0	0.0	0.0%	0	0.0	0.0%	0	0.0	0.0%
Female Coded Subtotal		73	2.8	100	8	2.3	66.67	23	2	100
Male Coded	bro	0	0.0	0.0%	4	1.1	33.33	0	0.0	0.0%
	bruh	0	0.0	0.0%	0	0.0	0.0%	0	0.0	0.0%
	brutha	0	0.0	0.0%	0	0.0	0.0%	0	0.0	0.0%
	nah	0	0.0	0.0%	0	0.0	0.0%	0	0.0	0.0%
	ain't	0	0.0	0.0%	0	0.0	0.0%	0	0.0	0.0%
Male Coded Subtotal		0	0.0	0.0%	4	1.1	33.33	0	0.0	0.0%
TOTAL		73	2.8	--	12	3.4	--	23	2.0	--

Table 5.54 Overall alternative spellings in Potter

16. Conjunctions

The eBAC found all three variations of the conjunction “and” to be female coded. The varieties, *and* and *&*, and the normed rates, which are higher than both the male (26.0) and female (29.0) eBAC subcorpora, are more indicative of likely female authorship.

	Chris			To Jamie			Jenelle		
	#	norm	%	#	norm	%	#	norm	%
and	1,530	58.6	99.9%	240	68.3	100.0%	528	45.8	95.8%
&	1	0.1	0.1%	0	0.0	0.0%	23	2.0	4.2%
n	0	0.0	0.0%	0	0.0	0.0%	0	0.0	0.0%
TOTAL	1,531	58.6	--	240	68.3	--	551	47.8	--

Table 5.55 Overall conjunctions in Potter

17. Articles and determiners

The eBAC found articles and determiners to be more male coded. All three datasets use a much lower frequency of determiners than both the male (141.5) and female (134.2) eBAC subcorpora, and thus all three are more female coded. Notably “Chris” and “To Jamie” have almost the exact same distribution, with Jenelle actually being the most comparatively male coded.

	Chris			to Jamie			Jenelle		
	#	norm	%	#	norm	%	#	norm	%
TOTAL	2,723	104.3	--	367	104.5	--	1,254	108.9	--

Table 5.56 Overall articles and determiners in Potter

C. Overview

Table 5.57 demonstrates the overall distribution of features in the datasets (where possible) as compared to both the eBAC frequencies, and as compared to Jenelle as a baseline.

	As compared to the eBAC	As compared to Jenelle
--	-------------------------	------------------------

		Chris	to Jamie	Jenelle	Chris	to Jamie
1	pronouns	female	female	female	male	female
	alternative spellings	male	male	male	female	female
	lowercase i	female	female	female	male	female
2	emotion terms	female	female	female	male	female
4	kinship terms overall *	female	female	female	female	female
	male coded kinship terms	male	female	male	male	female
5	male-coded friendship terms	male	male	male	male	male
	non-coded friendship term <i>friend</i>	male	male	female	male	male
6	lol abbreviations	female	female	female	female	male
	total abbreviations	female	female	female	female	male
8	expressive lengthening	male	male	male	female	male
9	backchannel sounds	female	female	female	female	female
10	hesitation words	female	male	female	female	male
11	assent terms	female	female	female	male	male
12	negation terms overall	female	male	male	female	male
	male coded negation terms	female	female	female	female	female
13	swears and taboo words	male	female	male	male	female
14	total prepositions	female	female	female	male	female
15	alternative spellings	female	female	female	female	female
16	conjunctions overall	female	female	female	female	female
17	articles and determiners	female	female	female	female	female
Total matching purported gender		6/21	6/21	15/21	9/21	8/21
Total matching Jenelle		19/21	17/21	--	--	--
Male findings not matching Jenelle		2/7	2/6	--	--	--

Table 5.57 Overall features in Potter as compared to the eBAC and Jenelle

This shows first that while the 21 features may appear to be poor indicators of the purported male genders of “Chris” and “To Jamie”, they actually show significant matches between both datasets and Jenelle (and indeed they share 16 features in common with each other). That the features are not only much more accurate in predicting Jenelle’s purported (and true) female gender, and highly accurate in matching both Jenelle’s largely female but also sometimes male-codedness, would appear to indicate that these features work as indicators of likely common authorship and shared gender, at least in this instance.

As compared to Jenelle, the results seem to be a less clear endorsement of Jamie’s assertion that the “To Jamie” language was more male coded than Jenelle. This does hold true for some features, but not for most. As discussed in the analysis above, this holds true for the use of male-coded kinship terms, as Jamie observed, but not for swear words, which are actually less frequent in “To Jamie” than in Jenelle.

This would also appear to indicate that there was no consistent attempt, on the part of Jenelle, to guise these specific features when portraying “Chris”, as only three of the 21 (male-coded friendship terms, non-coded friendship terms, and assent terms) are more male coded in both “Chris” datasets than in Jenelle’s writings as herself.

Thus it would appear that there are a few main takeaways from this analysis. First, it may be that a case like this with a larger dataset of questioned writings as compared to the previous analysis of Hemmert is likely to produce better results with this analysis. This may be incidental to these two cases, though it is generally true that the more language one has to analyze, the more robust it is likely to be; that is, it is more likely that the environments in which these features could occur for analysis are to be present. Of the 17 overall features able to be

analyzed in Potter (still excluding post length, hyperlinks, and keywords), only 8 were able to be analyzed in the questioned data for Hemmert (with an additional 9th of post length, which was of the same genre across the questioned and known datasets, unlike in Potter).

Further analysis in that Part showed that considering the absence of features which were generally considered more female coded left them more likely male coded, and thusly inflated the likelihood of male authorship. It also showed that a binary distinction with such a small dataset can be much more problematic than in the larger dataset here, where features were largely not binarily determined by their presence or absence, but by their frequency, distribution, and variety. This appears to have led to more accurate results in this instance, while more weighted tests of the findings appeared more accurate in the smaller Potter dataset.

Second, it may be that most of the features considered in this analysis are beyond the level of conscious thought for gender-guiseurs. This is demonstrated in the fact that there appears to be less consistency between “Chris” and “To Jamie”, both of which are supposed to have been authored by the same male individual, than there is between either dataset and Jenelle.

Additionally, even though Jamie himself reported that “Chris” used more (male-coded) swears and kinship terms, the actual analysis shows that this only holds true for the male-coded kinship terms. (However, as mentioned above, it may be that, in much the same way Jenelle seemed to curate her male-coded kinship terms as “Chris” when conversing with Jamie, she also curated the swears in her own language as herself to Jamie, which is not represented in the data provided in this case.) This observation would, however, appear to comport with the earlier analyses on the catfishing Tinder data in two ways—both because swear words and male-coded friendship terms showed the same conscious spike toward more highly male trending between the Before/During and After phases, and because the perception of some features by interlocutors may be influenced less by the actual language of their user than by their purported identity (until, of course, red flags are ratified).

Finally, it may be that there is a distinction this analysis can offer between what may be helpful for determining the likely gender of an author as compared to the larger findings of the eBAC analysis, and what may be helpful in determining likely common authorship when comparing questioned and known datasets. While the binary distribution of features as more female or male coded did appear to provide an accurate assessment of all three authors as female, it also seems worth considering the distribution of exact varieties as they are found within specific, individual features. That this analysis appears to solidly show that not only are all three authors likely female, but also that the non-female features seem to comport with Jenelle’s idiosyncratic non-female feature tendencies, and that they link all three datasets with individual varieties within features, seems to be an indicator that, given an appropriate amount of data, it can be a useful predictor for analysts.

And to that end, it may be that some features are better indicators of likely gender, likely gender guising, and likely common authorship than others. Whether this is explicitly due

to the context or size of any given dataset, or it holds true across all applications of this analysis in this thesis, is discussed further in Part 6 below.

Part 6 – Overall Feature Weights

Hemmert and Potter each posed unique challenges in the application of gendered language features to their data. For Hemmert, the largest issue was the lack of features. As discussed in Part 2, it is unclear whether this is because the amount of data was too small for the Known and Questioned Texts, if the individual consideration of the Questioned Texts left each one too small to consider solo, or whether it was not the total word count, but that the various features—proposed based on blogs and tweets—were not present either due to chance, context, or happenstance of Lael and Amber’s stylistic preferences. In Potter, discussed in Part 5, the largest issue was the lack of a comparator to contrast Jenelle as a potential female author of the purportedly male-authored Questioned data.

As such, both analyses required tailoring the approach to each case’s particularities. This may have come in the exclusion of some features, such as occurred in the Hemmert analysis, which weighed, in Part 3, the possible benefits and detriments of applying which present and absent features to which cross-sections of the data. This may have come in the form of splitting up the data in various ways, such as in both Hemmert and Potter where the writings of relevant authors were split between their audiences: including considering the power differential between an ex-wife and an estranged wife in the language of Known Lael in Hemmert, and the gendered differences between Questioned “Chris” “To Jamie” a male correspondent, and Questioned “Chris” to Jenelle, her mother, and a larger mixed-gender audience on Facebook in Potter. This may also have come in the filtering of features through the exigent context of each case, such as in Hemmert, for example, where family and familial relationships were always going to be a key topic, as all of the data occurred in the context of an ongoing divorce.

Despite the various challenges posed by each case, however, neither analysis was entirely unsuccessful in producing at least a preponderance of features that favored the real-world outcome. While this may have come down in part to having a large reference corpus like the eBAC, and features already probed for statistical significance such as by Bamman et al. (2014) and Schler et al. (2006) in much larger datasets upon which statistics can be reasonably and reliably applied, it was also in part due to the aforementioned consideration of context.

Both Hemmert and Potter are emblematic of the sorts of issues curated, quantitative-only approaches can face when being applied to real-world datasets, because they *are* real-world datasets, which vary in size, suitability, and uniformity in ways experimental or big data does not. This is elaborated upon by Koppel, Schler, and Argamon (2012):

The simplest kind of authorship attribution problem—and the one that has received the most attention—is the one in which we are given a small, closed set of candidate authors and are asked to attribute an anonymous text to one of them. Usually, it is assumed that we have copious quantities of text by each candidate author, and that the anonymous text is reasonably long. A number

of recent papers amply cover the variety of methods used for solving the problem.

Unfortunately, the kinds of authorship attribution problems we typically encounter in forensic contexts are more difficult than this simple version in a number of ways. First, the number of suspected writers might be very large, possibly numbering in the many thousands. Second, there is no guarantee that the true author of an anonymous text is among the known suspects. Finally, the amount of writing we have by each candidate might be very limited and the anonymous text itself might be short.

Thus this Part explores what of this analytical approach may work uniformly across datasets, what may work only in specific contexts and how, and what and when some deeper probing into the data and additional analyses may be required to achieve not only useful but also reasonable results. This Part is concerned with the exploration of three major points:

1. The usefulness of various **baselines**, often tied to the size, distribution (as between candidate authors), and sometimes context of the dataset
2. The import of **context**, including the genre, register, speech event, relevant Communities of Practice or social networks, as well as the topics in which the datasets are situated
3. The awareness of gender as a factor potentially altering the language patterns of candidate authors; that is, when language is overtly **performative**

A. A comparison to the eBAC

This section considers first the eBAC as a **baseline**, as such gendered reference corpora can be gainfully applied to datasets in which there is a paucity of data, either because of the data per author (not necessarily due to size, but due to environments for which these features can be appropriately analyzed), or the distribution among authors (such as in Potter, where there is no secondary (male) author to contrast with Jenelle).

Table 5.58 provides the gendered distribution of features found in both Hemmert and Potter as compared to the eBAC. Cells in gray indicate those findings that match the purported gender of the author, while findings in **bold** indicate those findings that match the distribution of a given feature of the eventual Known author (Lael in Hemmert, and Jenelle in Potter, both of whom were convicted on more than just the linguistic evidence).

		Hemmert					Potter		
		Q Texts	To Amber	To Haley	K Lael	K Amber	Chris	to Jamie	Jenelle
PURPORTED GENDER:		female	male	male	male	female	male	male	female
EXPECTED POWER SKEW:		--	female	male	--	--	male	male	--
1	overall pronouns	female	female	female	female	female	female	female	female
	alternative spellings	male	male	male	male	male	male	male	male
	lowercase <i>i</i>	male	male	male	male	male	female	female	female
2	total emotion terms	male	female	male	female	female	female	female	female
4	kinship terms overall	female	female	female	female	female	female	female	female
	male coded kinship terms	female	female	male	male	female	male	female	male
5	male-coded friendship terms	--	--	--	--	--	male	male	male
	non-coded friendship term <i>friend</i>	--	male	male	male	female	male	male	female
6	lol abbreviations	--	male	female	female	female	female	female	female
	total abbreviations	--	--	--	--	--	female	female	female

7	texts with periods	--	--	--	--	--	--	--	--
8	expressive lengthening	--	mixed	male	male	male	male	male	male
9	backchannel sounds	--	female	female	female	female	female	female	female
10	hesitation words	--	male	male	male	male	female	male	female
11	assent terms	female	female	female	female	female	female	female	female
12	negation terms overall	female	male	male	mixed	female	female	male	male
	male coded negation terms	female	male	male	male	male	female	female	female
13	swears and taboo words	--	male	male	male	male	male	female	male
14	total prepositions	male	female	female	male	male	female	female	female
15	alternative spellings	--	male	male	male	female	female	female	female
16	conjunctions overall	male	female	female	female	female	female	female	female
17	articles and determiners	female	female	female	female	female	female	female	female
18	text length	male	female	male	male	female	--	--	--
Found in the Q	Total male coded	6	5	7	6	5	6	6	6
	Total female coded	7	7	6	6	8	15	15	15
Total matching purported gender		7/13	9/19	12/20	11/19	13/20	6/21	6/21	15/21
Total matching expected to power		--	10/19	12/20	11/19	13/20	--	--	--
Total matching Known		9/13	14/19	17/19	--	--	19/21	17/21	--
Female findings not matching Lael		3/7	3/7	1/6	--	--	--	--	--
Male findings not matching Jenelle		--	--	--	--	--	1/6	2/6	--

Table 5.58 Overall comparison of Hemmert and Potter to the eBAC

As this demonstrates, using only the eBAC as a reference corpus, Potter was much more successful than Hemmert in matching the gender of the Known author, with “Chris” and “To Jamie” ~67-71% female coded, but the Q Texts in Potter only 46% male coded. Using only a gendered reference corpus and these features (and their sub-features), which are partially derived from the original BAC, only Potter would have returned both a useful (in that Hemmert is roughly 50/50) and accurate (in that Potter matches the real-world findings in court) results.

To that end, the eBAC only did slightly better in predicting Known Lael (58%) and Known Amber’s (65%) genders. However, as is also demonstrated above, the features were much more successful in matching the Questioned data to the Known authors’ own eBAC findings, which were not entirely male for Lael or entirely female for Jenelle, either. That is, the Questioned Texts matched 69% of the corresponding Lael features, and “Chris” and “To Jamie” matched 81-90% of the corresponding Jenelle features. Thus while comparing only to the eBAC may be less useful overall, the tool is not entirely without merit in instances of linguistic profiling, in which there is no Known author or authors to compare the Questioned data to.

This use, however, demonstrates that none of these features can be taken as absolutes, wholesale differentiating one gender from another. While reference corpora such as the eBAC are useful in establishing **baselines**, individuals vary, both from person to person, and from a single person’s audience to audience, or because of other non-gendered facets to their identity (e.g., Eckert & McConnell-Ginet, 1992). Language and gender norms vary as well, over time for an individual who may find themselves in a mostly contrasting community of practice, and over time for society at large, where some features that enter the lexicon as highly gender-encoded (such as the “bro” terms, as discussed in Chapter 4) may over time transition to neutral or other-gendered indicators. As such, reference corpora such as the eBAC may not provide the best or most appropriate **baseline** for a given dataset, at least alone

(the latter of which is considered in sections below). So long as the analyst keeps in mind that every author is unique, that these features cannot be doled out as absolutes, and that when **context** is applied it must be grounded so as to avoid some of the issues purported of non-linguistic profiling (e.g., Dern, Dern, Horn, & Horn, 2009; Snook, Cullen, Bennell, Taylor, & Gendreau, 2008), this analysis may yet provide a useful tool in profiling in conjunction with other analyses that take into account a more robust field of variables beyond gender alone.

B. Comparison to the Known Authors' eBAC Findings

Thus it may be useful to consider what features of the Questioned datasets *did not* match their Known author in gender-coding, and whether that may be because of some exigent **context** of the Questioned data itself, or because the Known author was able to successfully guise that feature in their **performance**. In this, both the eBAC and an individual author's linguistic patterns of distribution can be considered the **baselines**.

In Hemmert, we have the luxury of comparing the Questioned dataset not only to the person found to have likely been in custody of the phone (Lael), but the person whose identity they were purported to have maintained (Amber). As such, we might expect to find that the features that failed in a binary sense to match Lael's Known writings as compared to the eBAC were features that he successfully **performed** in the guise of Amber's known writings, instead. However, as we find in Table 5.59 below, this does not entirely appear to be the case.

			Amber	Q	Lael
4	kinship terms	male-coded kinship terms	female	female	female
		kinship terms TOTAL	female	female	male
12	negation terms	negation terms overall	female	female	mixed
		male coded negation terms	male	female	male
16	conjunctions	female coded conjunctions	male	male	male
		conjunctions overall	female	male	female

Table 5.59 eBAC comparison features in Hemmert that did not match Known Lael

Of the 4 sub-features, which come from only 3 total features, the overall use of negation terms and conjunctions can perhaps be immediately explained away by the **context** of the Q Texts. On average, the Q Texts were about 10.4 words per message, while both Amber and Lael were closer to 15. That the Q texts were shorter, and that there were much fewer of them, may itself deflate the environments in which conjunctions could occur. Indeed, the Q Texts did not match *either* K Amber or K Lael for either conjunctions or negations, which were a better match for each other. Notably, the conjunctions that don't match either Amber or Lael are more *male* coded, which would not appear to indicate either success or even an attempt to consciously guise this feature (as is likely equally true with the other three).

The male-coded kinship terms, on the other hand, do appear at first glance to have been successfully guised, as they match what the eBAC would expect for females in Amber and the Q Texts, and males in Lael. However, this is simply because Lael used the word *wife* when talking to Amber, his wife, and not because of some more prevalent pattern in the Q Texts.

Moving on to Potter, a total of 5 sub-features from 5 different overall features do not

match Jenelle’s eBAC findings in either or both “Chris” and “To Jamie”. Only one of the features, the non-coded friendship term *friend*, is a match between “Chris” and “To Jamie”, indicating that “Chris” was not consistently **performed** throughout his Q; this then would beg the question of whether this is because Chris’s language was tailored to Chris’s audience. However, as we see in Table 5.60 below, where we would expect the “To Jamie” features, when they do not match either “Chris” or Jenelle, to do so because they trend more male coded for the male audience of Jamie, this is not often the case.

			Chris	To Jamie	Jenelle
4	kinship terms	kinship terms overall *	female	female	female
		male coded kinship terms	male	female	male
5	friendship terms	male-coded friendship terms	male	male	male
		non-coded friendship term <i>friend</i>	male	male	female
10	hesitation words	hesitation words	female	male	female
12	negation terms	negation terms overall	female	male	male
		male coded negation terms	female	female	female
13	swears	swears and taboo words	male	female	male

Table 5.60 eBAC comparison features in Potter that did not match Known Jenelle

The single feature that only “Chris” does not match to Jenelle is overall negation terms. As mentioned in Hemmert and in the analyses above, it may be that negation terms as they are counted as distinctive in the original BAC analysis, do not function the same way in conversational language like emails and Facebook as they did on blogs, and thus we might then consider the **context** of a dataset against a **baseline** like the eBAC. Again, this is unlikely to be demonstrative of active guising on Jenelle’s part, as this feature shifts more female than male coded in “Chris”.

The three features in “To Jamie” that do not match Jenelle (or “Chris”) are male-coded kinship terms, hesitation words, and swears. Male-coded kinship terms, however, come down to terms for female significant others, including *wife*, *girlfriend*, and *gf*; “Chris” doesn’t talk about his purported wife to Jamie, and so the fact that these words do not appear is hardly surprising. More indicative is the overall use of kinship terms, which marks not the specific **context** of *who* is mentioned by differentiating female significant others from every other familial relation, but of how often any familial relations are mentioned. Overall this is a more female-coded feature, and all three datasets have a more female-coded frequency of use as expected by the eBAC.

Hesitation words differ in “To Jamie” for the same reason as male-coded kinship terms, in that they do not exist in “To Jamie”. This is possibly for the same reason that lowercase *i* might vary: hesitation words can often be a marker of formality, varying not only in distribution but also in appearance in the differing **contexts** of various datasets.

As discussed in Part 5 above, both male-coded kinship terms and swears were remarked by Jamie himself as being confirming indicators that “Chris” was, in fact, a man using Jenelle’s accounts, and not, in fact, Jenelle using Jenelle’s accounts. However, as demonstrated in Table 5.60 above, the use of swears was actually more female coded in “To Jamie”, in contrast to their male-coding in “Chris” and Jenelle. This is not because swears did not occur—they did; this is because they occurred with a much lower frequency than in either

“Chris” or Jenelle. However, both “Chris” and Jenelle share the **context** of aggressively defending Jenelle on Facebook (or otherwise talking about their aggressive defense of Jenelle).

Finally, the use of the term *friend* is only partially more male coded in “Chris” as a matter of how it is normed, as shown below, while in “To Jamie” it exists in stark contrast to Jenelle.

	Male eBAC		Female eBAC		Chris		to Jamie		Jenelle	
	norm	%	norm	%	norm	%	norm	%	norm	%
friend	1.1	81.0%	1.3	81.4%	1.0	86.7%	0.3	4.0%	4.3	90.9%

Table 5.61 Distribution of friend in Potter and the eBAC

As a percentage of overall friendship terms in “Chris”, it is actually above female coded, and closer to Jenelle’s use than the female or male eBAC. As such, as with kinship terms, it may be more useful to look at the overall use rather than the individual varieties (though the individual varieties between “Chris” and Jenelle are still quite close, and the majority preference of all friendship terms). This occurs in stark contrast to “To Jamie”, which, out of 25 total friendship terms, has only 1 instance of “friend”, while all remaining instances are split between the male-coded *dude* and *man*. This, again, was brought up by Jamie as an indicator that “Chris” was Chris, not Jenelle.

Thus, it may be that these final two features indeed *were* successfully—but not consistently—guised by Jenelle. While Jenelle also used a higher (more male-coded) frequency of swears than the female eBAC subcorpus, “Chris” used them twice as often, with a normed rate of 11.5 as compared to Jenelle’s 5. While the same did not hold true for “Chris”, at least in “To Jamie”, Jenelle’s **performance** starkly prioritized male-coded friendship terms over all alternatives. As discussed further below, this is entirely in line with what the Tinder data found: that these more overtly **performative**, male-coded features were more overtly **performed** by male users in the After phase of their conversations once their audience was revealed to be male. This is bolstered by the fact that these same overtly male-coded features (friendship terms and swears, for example) were overtly **performed** in the Tinder data in an array of dynamics, both when an author or pair of authors took the reveal as humorous for whatever reason, and when an author or pair of authors responded with vitriol to something they assumed was intentional targeting by their interlocutor challenging their gender or sexual identity; and in situations when both authors were fully aware of each other’s true gender identities, and in situations where the knowledge was one-sided. (That is, these features are *not* simply indicative of accessing a shared identity or inducing camaraderie, but rather are overtly **performative**.)

In comparing Hemmert and Potter, kinship terms and negation terms failed, in whole or in part, to accurately predict the gender of the Questioned author. When it comes to negation terms, this is perhaps not so strange. After all, Bamman et al. (2014) and Schler et al. (2006) did not fully agree on their findings with regards to negation terms, with Bamman et al. (2014) finding various non-standardisms to be gender coded, and Schler et al. (2006) finding that not any particular variations but negation terms overall were more female coded. As expressed

multiple times throughout this thesis when **context** was brought back into the forefront, both negation and assent terms may be less helpful simply because of the type of discourse being analyzed. This, then, begs the question of whether these (and possibly some similar) features fail to hold up across genres, or at all, as will be considered further in the discussion below.

When it comes to kinship terms, however, Bamman et al. (2014) and Schler et al. (2006) were in much closer agreement, each even providing some of the same examples of specific female-coded terms (*mom*, *mommy*, *husband*, *hubs/hubby*) for what they agreed to be an overall female-coded feature. However, as mentioned multiple times throughout this thesis when **context** is considered, very few variations are considered male coded, and those that are restricted to female significant others. In the **context** of guising, the idea that a speaker will mistakenly use *wife* when the guised target is a woman or forget to refer to the significant other of a male **performed** identity as a *girlfriend* not a *boyfriend* as **performative** slips seems highly unlikely. It thus seems more useful to consider the overall use of kinship terms and discussions of family, which Schler et al. (2006) found as an overall category to have the “greatest information gain for gender”. Even so, this begs the issue of **performance**, as in Hemmert, which is mired in family dynamics as both a case and an array of data. This, then, similarly begs the question of whether these (and possibly some similar) features fail to hold up across **contexts**, or at all, as will also be discussed in further sections below.

C. Relational comparison to the known authors

While the previous two sections and much of the analysis in this thesis are able to employ the eBAC as a reference corpus of contemporary, gendered CMC, it may not always be the case that this is the best or the only **baseline** to consider. In Hemmert, in particular, the data provides the luxury of having the Known language of two suspects in the case to compare and contrast to both one another as **baselines** as well as the Questioned data. Because, as we have both seen and discussed, individuals tend to vary, the **baseline** of an individual for any of these 20 features is likely to be more probative than the more generic **baseline** of something like the eBAC. Thus, in reverse order of the previous two sections, this section compares Questioned feature distributions in relation to the **baselines** of the Known data first, and the eBAC second.

This is best exemplified in Hemmert in Part 2, reproduced here in Figure 5.8 below, which shows the comparative normed frequencies of the 9 overall features found for analysis in the Q Texts.

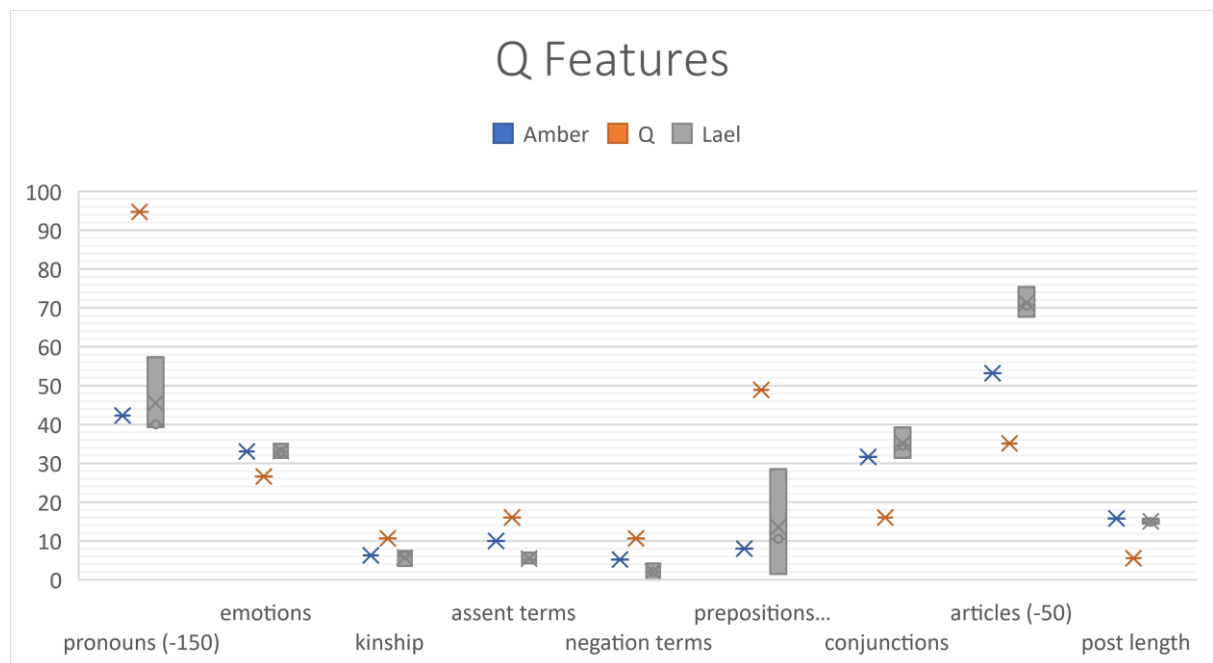


Figure 5.8 Distribution of features found in the Q

In comparing the Q Texts to Amber and Lael, we find that in every single one of the 9 features, Amber and Lael are closer in their range of use to each other than either is to the Q Texts. In many cases, Amber is even encompassed within the range of the split Lael datasets, including pronouns, emotion terms, kinship terms, and prepositions, and just barely misses in a couple of others (such as conjunctions and post length).

Thus we are left to consider not what is a *good* match, but what is the *better* match for the Q Texts between Amber and Lael. In considering their overall averages, only pronouns, emotion terms, prepositions, and post length are more similar between Lael and the Q Texts, while the extreme ranges of Lael would include kinship terms and conjunctions. In no cases do the Q Texts fall within the range of all four Lael subsets, and this still leaves assent terms, negation terms, and articles as much better matches between Amber and the Q Texts regardless of how Lael's frequencies are considered. To that end, in considering only the averages, Amber is, instead, the better match for kinship terms and conjunctions.

It is true that the three features that better match Amber may be because they were so well **performed** by Lael; similarly, that many of the Q Text features are, again, such stark outliers may be because they were not a **performance** of Amber specifically, but a **performance** of non-Lael more generally. This, however, seems unlikely. Instead, that the Q Texts are such a stark outlier in their own analysis may, as repeatedly mentioned, be due to the fact it is so small, compartmentalized, and specific in its **context**, thus inflating or deflating many of the features, at least as far as the relative distribution of normed frequencies is considered.

As such, it may be more probative to **contextualize** or to simply scrap this sort of analysis on the overall Q dataset as a whole, as was done in Part 2, or to consider other facets of the analysis as was employed in Part 4. Here, however, we move on to a discussion of Potter, which had its own set of difficulties, before reaching an ultimate conclusion.

		As compared to the eBAC			As compared to Jenelle	
		Chris	to Jamie	Jenelle	Chris	to Jamie
1	pronouns	female	female	female	male	female
	alternative spellings	male	male	male	female	female
	lowercase i	female	female	female	male	female
2	emotion terms	female	female	female	male	female
4	kinship terms overall *	female	female	female	female	female
	male coded kinship terms	male	female	male	male	female
5	male-coded friendship terms	male	male	male	male	male
	non-coded friendship term <i>friend</i>	male	male	female	male	male
6	lol abbreviations	female	female	female	female	male
	total abbreviations	female	female	female	female	male
8	expressive lengthening	male	male	male	female	male
9	backchannel sounds	female	female	female	female	female
10	hesitation words	female	male	female	female	male
11	assent terms	female	female	female	male	male
12	negation terms overall	female	male	male	female	male
	male coded negation terms	female	female	female	female	female
13	swears and taboo words	male	female	male	male	female
14	total prepositions	female	female	female	male	female
15	alternative spellings	female	female	female	female	female
16	conjunctions overall	female	female	female	female	female
17	articles and determiners	female	female	female	female	female
Total matching purported gender		6/21	6/21	15/21	9/21	8/21
Total matching Jenelle		19/21	17/21	--	--	--
Male findings not matching Jenelle		2/7	2/6	--	--	--

Table 5.62 Overall feature distribution in Potter

Unlike Hemmert, Potter had no second author to compare the Questioned data to, and as such we are left to consider what the threshold may be for determining any given feature as a match with Jenelle. As was also undertaken in Hemmert, Table 5.62 above, then, takes Jenelle as the **baseline**, and the directional variations from Jenelle as gender-indicative based on the eBAC findings.

To explain this, we might use the overall distribution of pronouns, which the eBAC considers more female coded, as an example. Jenelle's Known writings provide the **baseline** for overall pronoun use of 189.4. Thusly, anything above 189.4, such as "To Jaime", can be considered to be more female coded, and anything below, such as "Chris", can be considered more male coded, regardless of the eBAC's much lower gendered distribution of 133.6 for females and 109.7 for males (which leaves "Chris", "To Jamie", and Jenelle as highly female coded).

	Chris		To Jamie		Jenelle		Male eBAC		Female eBAC	
	#	norm	#	norm	#	norm	#	norm	#	norm
TOTAL pronoun use	4,713	180.5	712	202.7	2,182	189.4	6,500,262	109.7	7,760,914	133.6

Table 5.63 Total preposition use in Potter

Of course, this only works for confirming female-codedness for a female Known author, and does not necessarily suggest male-codedness for features on the wrong side of the Known threshold. It is worth noting that this approach *did* produce a preponderance of features in favor of female authorship for both "Chris" (12/21) and "To Jamie" (13/21), though that may not always be the case in every **context**.

As we can see when directly comparing the analytical approaches and resulting issues

between Hemmert and Potter, however, it was not the number and gender distribution of Known authors that determined the success of the analysis in identifying an author's gender. Contrastingly, Potter was much more successful than Hemmert, both in the robustness of analyzable features, and in their matches between the Questioned data and Known author. And this, as mentioned above, is not because Lael consistently or successfully guised any of his features as Amber, though Jenelle arguably did successfully **perform** at least one feature in both "Chris" and "To Jamie".

Rather, it seems much more effective to use the eBAC as the primary **baseline**, as was done in the previous two sections, but then to match the binary distribution not to the eBAC subcorpus genders, but to the gender-coding of that feature for any given Known author. That is, while the distribution of alternative spellings in both "Chris" and "To Jamie" are considered more male coded, and thus would not appear to be a match for Jenelle's gender, Jenelle's distribution of the feature is *a/so* more male coded in relation to the eBAC.

As we have discussed throughout this thesis, the very few times in which Questioned features do not match the threshold of Known features are almost always explainable by **context**. A systematic ability to apply **context** as the appropriate approach, versus determining a finding is simply indicative of a non-match, then, is the ultimate question of usefulness for such approaches.

D. Determining the overall validity of individual features and the application of context

Although every analysis conducted in this thesis was able to make use of at least some portion of the features, not every one was applicable to every dataset. Each of the four datasets were different CMC genres, including blogs (AGGID), instant or direct messages (Tinder), text messages (Hemmert), and both Facebook posts and emails (Potter). However, it was not simply the genre that precluded features from being applicable; Herring et al.'s 2004 genre analysis of blogs, for example, included many of the features of this analysis not found in the AGGID blog. Thus it appears the suitability of features may come down to formality, individual author preference, or specific **context**, as is explored via keywords in the AGGID blog.

Table 5.64 below demonstrates the distribution of features in each of the four analyses, and not what they found or whether it was probative, but simply a **yes** or no to indicate whether the feature appeared for analysis at all in the Questioned data.

	Feature	AGGID (blogs)	Tinder (instant messages)	Hemmert (text messages)	Potter (posts and emails)
1	Pronouns	yes	yes	yes	yes
2	Emotion terms	yes	yes	yes	yes
3	Emoticons	no	yes	no	yes
4	Kinship terms	yes	yes	yes	yes
5	Friendship Terms	yes	yes	no	yes
6	Abbreviations	no	yes	no	yes

7	Punctuation	yes	yes	no	yes
8	Expressive lengthening	no	yes	no	yes
9	Backchannel sounds	yes	yes	no	yes
10	Hesitation words	yes	yes	no	yes
11	Assent terms	yes	yes	yes	yes
12	Negation terms	yes	yes	yes	yes
13	Swears and taboo words	yes	yes	no	yes
14	Prepositions	yes	yes	yes	yes
15	Alternative spellings	no	yes	no	yes
16	Conjunctions	yes	yes	yes	yes
17	Articles and determiners	yes	yes	yes	yes
18	Post length	yes	no	yes	no
19	Hyperlinks	no	no	no	no

Table 5.64 eBAC feature distribution in the four analyses

This primarily (but not exclusively) came down to those features that can be seen as most indicative of chatspeak, as was at issue in both the Hemmert case as well as the A Gay Girl in Damascus analysis. Neither had any apparent use of emoticons, abbreviations, expressive lengthening, or alternative spellings. Although the AGGID blog contained other chatspeak features that Hemmert did not, namely backchannel sounds and hesitation words, these were so infrequently used in the AGGID blog that they did not end up being particularly demonstrative.

Aside from chatspeak features, Hemmert also did not include friendship terms, punctuation, or swears and taboo words. While Potter did include the chatspeak feature of a small number of text-based emoticons, it was unclear in Potter (and indeed in Hemmert) whether image-based emojis were not used at all, or were simply not preserved in the data provided to the courts. No dataset made any use of hyperlinks, including the AGGID blog, although it is unclear from the Wayback Machine preservation of the data whether this feature existed for analysis or not.

Beyond chatspeak features, the absence of a feature in any given analysis was situational. Hemmert happened to have no punctuation whatsoever—not an unusual outcome for text messages. Nor did Hemmert happen to have any swears or taboo words or friendship terms in its 188-word total. Potter, of course, had post lengths; but it was unclear from document to document whether any given Q or K file was originally a Facebook post, or an email, and as such it was not possible to clearly or accurately compare the post lengths of the individual genres involved.

Keywords were employed in both the AGGID blog and the Hemmert case analyses, but not because they were probative or demonstrated any likely gender distinction. Rather, in both instances, they were strong indicators of the **context** of the dataset. For the AGGID blog, the **context** was highly political, with 84 of the top 100 keywords relating to the Damascus spring as compared to the eBAC. For the Hemmert case, the top keywords largely revolved around familial interactions, related to family and the ongoing divorce, or were otherwise specific to the context and topics of the interactions taking place, such as the horse the husband and wife were arguing over selling. Very few of the keywords would have fallen into the LWIC categories Schler et al. (2006) found probative, and even if they did (such as family

terms), they were so motivated by **context** as to be highly inflated. Similar can be said for both the Tinder data and Potter, which both occurred in highly specific **contexts** as compared to the variety found throughout the eBAC.

Though the keyword analyses did not demonstrate any gender findings, they did demonstrate the probativeness of particular features in their two datasets. The highly political nature of the AGGID blog can be seen as a potential contributing factor to the lack of chatspeak features, which can be considered highly informal when juxtaposed against the highly formal and serious topics the blog covered. The **context** of family members in the Hemmert case makes a feature like kinship terms potentially less probative in the Q as a standalone dataset, and indeed both the husband and wife in Hemmert had an almost identical distribution of kinship terms in their K, at 6.2 and 6.3, respectively.

Aside from via keyword analysis, other features can be present but not probative, such as was the case with punctuation in the Tinder data, which a keyword analysis would not have necessarily captured. To a degree not exhibited by any other dataset considered for analysis in this thesis, the Tinder data contained a strong preference for question marks over every other sentence-final punctuation category, across all three phases. That this occurred most drastically in the During phase is itself indicative of the investigative nature of such an interaction, which would beg even more questions than the general theme of getting to know one's interlocutor even in an average Tinder-like interaction. This was reflected in the much higher distribution of lengthened punctuation as well, which is consistent with the numerous opportunities the very specific **context** of third-party catfished Tinder conversations provides for disbelief and uncertainty, emphasized through punctuation.

So it appears that the usefulness of any given feature or set of features may be driven by genre (to be absent) or **context** (to be overly present). That any given feature does not occur in any given dataset for analysis, then, is not a failure of the 20 features and their internal permutations. Although the Tinder data and Potter case showed that more features are better in demonstrating clearer gendered trends, both the AGGID blog and Hemmert case analyses demonstrated some findings with the features that they had. In the latter two, and indeed in many similar applications of this approach, the keyword analyses were helpful in demonstrating why certain features might be absent or less probative even when they were present, and both produced a preponderance of evidence in favor of the true gender of their eventual Known author.

Thus it appears that any constellation of these features as applicable to any given analysis can be useful in determining the likely gender of an author, especially when **context**, genre, and the Known author suspects' language are considered in forming a **baseline** for these features. What this section does not demonstrate, however, is whether any of these features usefully demonstrated overt, successful gender **performance** by an author, which is discussed further in the section below.

[E. Features indicative of gender guising and the problem of audience](#)

	Feature	AGGID (blogs)	Tinder (instant messages)	Hemmert (text messages)	Potter (posts and emails)
	Dataset	Amina	After	Q Texts	Chris
1	Pronouns	male	male	unclear	female
2	Emotion terms	male	male	male	female
3	Emoticons	--	female	--	--
4	Kinship terms	female	male	female	female
5	Friendship Terms	male	male	--	male
6	Abbreviations	--	male	--	female
7	Punctuation	unclear	unclear	--	--
8	Expressive lengthening	--	male	--	male
9	Backchannel sounds	unclear	male	--	female
10	Hesitation words	unclear	male	--	female
11	Assent terms	male	male	female	female
12	Negation terms	female	male	unclear	female
13	Swears and taboo words	male	male	--	male
14	Prepositions	female	female	unclear	female
15	Alternative spellings	--	unclear	--	female
16	Conjunctions	female	male	male	female
17	Articles and determiners	unclear	male	female	female
18	Post length	unclear	unclear	--	--
Total matching gender		5/8*	13/18*	2/5*	12/15
Total matching K Author		7/11	--	5/8	14/15

Table 5.65 Overall findings of all four analyses

Table 5.65 above is an overview of the findings of the four analyses in this thesis, with the overall outcome of each feature as compared to the eBAC listed as male, female, or unclear; matches to the K author's gender are indicated in **bold**, while matches to the K author's own distributions is indicated in **gray** (though obviously this is not possible for the Tinder data), and the few Q features not available for comparison in the corresponding K are indicated in *italics*. Finally, the total counts correctly predicting gender via match do not include features found to be unclear in their totals.

As this demonstrates, although all but Hemmert achieved success by simply looking at these features against a **baseline** like the eBAC, in every instance where there was a comparable K author as an established **baseline**, the findings were much more accurate. However, no single feature or set of features was found to be wrongly encoded by the eBAC across all datasets (which would, then, indicate it was prone to guising in all cases). Only one feature never matched the patterns of the K author—negation terms—though, as discussed frequently throughout this thesis, it is unclear whether that is because of the difference between largely non-conversational blogs like the eBAC, and the variety of datasets analyzed here.

What is interesting to note, however, is a comparison between Potter, which was found to be a female author guising herself in a male persona, and the other three datasets, which included the language of male authors tied to female authors in a couple of different ways. In every case where “Chris” **performed** a male feature, those features, where able to be analyzed, were also always male in the other datasets. Friendship terms and swears, in particular, were even confirmed by Jamie in the Potter case to be indicators, for him, that “Chris” was indeed male, though less can be said for expressive lengthening.

Conversely, although no other features that Jenelle failed to disguise were also always

female in the other three datasets, two came close: kinship terms, and prepositions. It is unlikely that prepositions stick out to language users as overtly male coded, prompting male authors to attempt to lessen their uses of prepositions in order to be perceived as more female-sounding. As a category of function words, as compared to kinship terms, friendship terms, and swears as content words, prepositions are much less likely to be at the level of conscious thought for any given person, though they may, instead, be indicative of other structural differences in the language being used that is not illuminated by their sheer distribution.

Although a content word category, kinship terms do not appear to be the same female-**performative** feature as friendship terms and swears, however. This is largely due to the issue of **context** for the three datasets in which it is female coded—clear in Hemmert, at least, for the fact that K Lael's use of kinship terms was also female coded—which is discussed in each analysis individually. Conversely, the Tinder data is a genre in which topics of family are likely best avoided, save for the few occurring, logistical instances of mentioning siblings and such. It is not beyond the realm of possibility that kinship terms could function in the same way, of course; just that that does not appear to be what these four analyses, at least, indicate.

Though friendship terms, swears, and expressive lengthening—and possibly other more conscious-level categories—can be seen as surface-level categories easily managed by someone disguising their gender, what is more likely is that they are less indicators of deception and more indicators of **performance**. It is also worth noting that these features shift somewhat when audience can be factored in: the male-codedness of features in the Tinder data's After phase often shifted from some slightly more female-codedness in the Before phase; K Lael's own comparable features shifted depending on the power differential between him and the female he was conversing with; and "Chris's" language shifted between Jamie and the world at large.

As such, while it appears that no feature can be seen as a blanket indicator that achieves success in illuminating likely gender guising in the same way that the overall analysis achieves success in determining likely gender, considering the overtness of any given feature in any given analysis, then, becomes another step for the analyst in considering the output.

Part 7 – Discussion/Overview

Data are limited. If nothing else holds true in a forensic **context** (or indeed in a linguistic **context**), it is this. Data may often be small, or messy in that language is being compared across genres or mediums, or may involve a host of other issues for the analyst to overcome, or walk away from. Statistical approaches to datasets are not always possible or wise, for a variety of reasons. Some analytical approaches (e.g., Fobbe, 2021) achieve more success in authorship analyses with short sets of questioned data such as text messages (albeit not explicitly on gendered lines), including those at issue in the Hemmert case, meaning they are not impervious to every linguistic toolkit. But the features proposed in this thesis, while not employed statistically in a way that would have inflated or deflated their findings on the datasets considered here, were derived from original research that previously proved their statistically

probative value in determining an author's likely gender (e.g., Bamman et al., 2014; Schler et al., 2006).

What these four analyses demonstrated and this Part reviewed is that such features can still provide useful toolkits for analysis when particular approaches are employed. It is first useful to establish a **baseline** like the eBAC, which is large enough and general enough to have had such features found statistically useful. It is secondly useful, where possible (again, forensic data are limited!) to establish a **baseline** of known author suspects, because the only thing people do not vary in is their variability.

It may also be helpful to refine this second **baseline** (and the first, if a sufficient reference corpus makes this possible) by considering the audience and community of practice. There was a difference in the Tinder phases of Before, During, and After as an author's perception of their audience's gender and trustworthiness shifted. There was a difference not in gender but in power in Hemmert between the husband's interactions with his ex-wife and his soon-to-be-ex-wife. There was a difference in Potter between how "Chris the CIA agent" conducted himself with a frequently antagonistic (and antagonized), mixed audience on Facebook he was defending his friend Jenelle from, and Jamie, his guy friend who he ultimately convinced to be complicit in a double homicide.

And **context**, as always, is key. Genre may provide a useful explanation where the previous **baselines** fail, especially if those **baselines** are themselves of a different genre. While a keyword analysis along LIWC categories may indeed be probative of gender as Schler et al. (2006) found, it can also be useful in situating both the presence and absence of other features, as well as their distributions, in **context**. But it should always be the case that an analyst determines **context** and its effect on such pre-determined lists of features before running off with their output to erroneous ends. This falls into the 'algorithm assistance' and 'algorithm informed evaluation models' proposed by Swofford and Champod (2021), as discussed in Chapter 3.

As a methodological approach, however, the above does not answer the question of whether any of these features can be probative of overt (and also sometimes disguised) gender **performance**. As Holmes & Meyerhoff (2003) point out:

The focus is on the way individuals 'do' or 'perform' their gender identity in interaction with others, and there is an emphasis on the dynamic aspects of interaction. Gender emerges over time in interaction with others. Language is a resource which can be drawn on creatively to perform different aspects of one's social identity at different points in an interaction.

This was evident in the datasets in which an author's audience affected the gendered encoding of their language as compared to the eBAC; this would include instances such as the difference between Tinder users thinking they are talking to a female in the Before phase and knowing they are talking to a male in the After phase, or the difference between the writings of Known Lael and his ex-wife as compared to with his soon-to-be-ex-wife.

This analysis did not necessarily find any particular, consistent indicators of overt attempts by an author to disguise the gender-coding of their language via any of the twenty features or their permutations considered in this thesis. Rather, these analyses, when they were not complicated by **context**, genre, or register, found that certain features do appear to be indicative of the overt **performance** of a gender identity. That these features are not features of disguise but instead features of **performance** is best exemplified by a comparison of the Tinder data and Potter, which both saw, for example, kinship terms and swears as their strongest indicators of overt maleness. While this **performance** may indeed have been motivated by an attempt at disguise by Jenelle as “Chris” in Potter, the same deliberateness of disguise cannot be said for the catfished Tinder users. Yet, when both had need for a strongly male identity, as **performed** to a male audience, both overtly employed the same features in order to shore up the perception of their gender identities.

In the end, it appears that only overtness provides any insight as to potential guising into the findings of these features in any given analysis. The features by themselves were largely successful in predicting gender, with variations often coming down to audience—as was, indeed, pointed out by Schler et al.’s (2006) initial analysis—and existing within an individual person’s linguistic variation. However, for such features to be a useful tool in predicting an author’s gender, and perhaps predicting when an author’s gender is **performed**, a variety of exigent factors, as laid out in this chapter, must be considered.

Chapter 6 – Discussion/Conclusion

This thesis sought to explore the interplay of both identity performance and audience perception, while paying specific focus to the identity of gender, and language varieties within the realm of CMC. This thesis took multiple approaches in analyzing and contextualizing linguistic data, first by applying the 20 features outlined in Chapter 3, second by filtering the usefulness of the findings through the four linguistic theories proposed in Chapter 3, and third by considering perception and performance more specifically as they relate to the two commodities at opposition in the latter part of Chapter 4. Also briefly considered here are the broader implications for concepts outlined in Chapter 2, including intersectionality and constructivist and essentialist approaches to such analyses.

Of interest first was what strategies people use when performing gender, such as the variables and linguistic theories proposed in Chapter 3, both when the performance is guised and overt. Additionally of interest was how a constellation of such features might work as predictors of gender, both when performance was natural and overt, and how these findings might then be applied to instances of authorship and profiling. Finally, of interest was how successfully such variables can be applied to a variety of datasets and analytical questions, and what analytical frameworks might be responsibly undertaken to situate such quantitatively derived variables in the appropriate qualitatively derived contexts.

A variety of cases in which gender was a major identifying factor were considered, some of which were natural examples of attempted guising (i.e., Chapter 1's Ashley's Angels, Chapter 4's A Gay Girl in Damascus blog), some of which were incidental instances of guising (i.e. Chapter 4's Catfi.sh Tinder data), and some of which were forensic cases that put both facets of identity more broadly (i.e., Chapter 1's "Erin Princess Baby" case, Chapter 5's Justin Carter case) and gender more specifically (i.e., Chapter 5's Hemmert and Potter cases) at the forefront of their investigations.

As overviewed in the previous chapter, many of the chapters in this thesis applied the relevant Chapter 3 features to four main datasets, using the eBAC as a baseline reference corpus: Chapter 4's A Gay Girl in Damascus blog and Catfi.sh Tinder Data, Chapter 5's Hemmert and Potter cases. But this thesis also sought to funnel the various analytical approaches through a variety of scenarios in order to establish an approach that appropriately couched the set of 20 features and their permutations. Chapter 3 provided appropriate linguistic theories for an analyst to consider when surveying the outcome of Chapter 3's feature-based approach or indeed any approach that seeks to analyze real-world, forensic CMC (namely audience design, accommodation theory, community of practice, and speaker contamination), and Chapters 4 and 5 explored a variety of avenues through which these theories can be appropriately applied to the genre, register, and context of any individual case.

As a result, although this thesis found the twenty Chapter 3 features to be relatively reliable markers of the likely gender of a dataset of unknown authorship (and thus unknown

gender), this thesis did not find any specific output of any feature to be indicative of gender guising per se. While not indicators of guising or even deliberate deception, however, this thesis did find features that proved to be consistently overtly performative: swear words and friendship terms. This held true across two very different datasets, with stark findings in both the Catfi.sh Tinder data—in which no unwitting participant was actively attempting to guise their language, but in which many participants did later attempt to actively and overtly perform their maleness after it was discovered they had been erroneously perceived as female prior—and in the Potter case—in which a female author was specifically attempting to disguise herself as a male—both of which had a male audience.

That these two features are both apparently male-performative and no features were found by this thesis to be consistently female-performative likely has more to do with the datasets considered than the linguistic possibilities. The other two cases analyzed, the A Gay Girl in Damascus blog—wherein a man performed the identity of a female for an extended period of time—and the Hemmert case—wherein a husband is alleged to have texted as his wife over a short period of time—had their own host of issues that made such overt features improbable.

The A Gay Girl in Damascus blog situated Amina's identity well within the male author's actual identity, differing only in gender, ethnicity (but not also nationality or indeed even local geography), and sexuality (though both shared an attraction to women). The author of the blog also had a strong, long-term understanding of the topic of the blog—politics and social justice in Damascus—and as such his performance may well have been highly advantaged by all of the other, non-gendered facets. His audience, too, likely had a more open perception, as the identity he was claiming (that of a lesbian activist in an environment that would put her in danger of arrest or worse for being so open about her identities and views) was a tenuous one, with cooperation rather than suspicion as their prime commodity. Additionally, the blog was, unlike any of the other datasets considered in this thesis, highly formal, and often written more in the style of a book than the more conversational datasets upon which the Chapter 3 features were firstly and primarily based.

The Hemmert case, on the other hand, seemed to have the opposite constraint. Where Amina's long-term performance was situated well within a familiar identity and had access to a variety of experiences shared by the overlapping identity facets, the overall dataset for Hemmert was not only short as a matter of words, but short as a matter of time (and indeed likely even smaller than was considered here, according to the wife's later admission that she may have herself authored some of the early or late tweets). Hemmert and his estranged wife obviously did share some part of their identity, which, as discussed throughout this thesis, skewed for example the usefulness of kinship terms. But as previous research and indeed instances like the Catfi.sh Tinder data demonstrate, performative identity is emergent, through either time (as the AGGID blog had) or intentional overtness. The few questioned texts the Hemmert dataset hinged on were largely logistic, and while not formal in the same way the

AGGID blog is, were nevertheless lacking in many of the chatspeak categories the 20 features rely on. In contrast to the “Erin Princess Baby” case in Chapter 1, wherein the UCO apparently attempted to focus on overt feature performance at the expense of logistics (by telling Mr. Plumridge, for example, that he was at the office), Hemmert focused on logistics rather than overt performance.

It is unlikely, then, that based only on these two examples of male-to-female guising, there are no features likely to be indicative of overt female identity performance in the same way swears and (specific) friendship terms appear to indicate overt male performance. Indeed, neither does this thesis seek to argue that these two features and these two features alone are always indicative of overt male performance, or the only features potentially indicative of overt male performance. Other features, such as preposition use, demonstrated a pattern that may indeed be indicative of the feature as a performative one, albeit one below the level of conscious thought and thus likely indicative of some socially gender-aligned structural differences not captured by this analysis, with the potential of demonstrating function categories as potentially as probative as content ones. Instead, it is hoped that the takeaway from the analyses above is that all output of any such analyses must be situated and considered in their appropriate contexts.

That is, any such content-based features that an author has access to may be prone to overt performance, indicative either of some attempt at deception, or some exigent attempt at further reifying a genuine performance. As stated in Chapter 2, it is increasingly important to consider the intersectional nature of multiple identity facets in such approaches rather than relying too heavily on pre-established norms as absolutes. That is, not only do gender identity and performances thereof exist in conjunction with other facets of identity, such as sexuality in the AGGID blog and Catfi.sh Tinder data, and parenthood in the Hemmert case, for example, but they also exist on a gradient, as we saw in the shift from more passive to more overt performances in the Catfi.sh Tinder data, and in the difference between Jenelle’s female performances to the male audience of Jamie specifically and everyone else as a whole in Potter. It is, then, perhaps not so surprising that no feature or set of features can always predict performance in any given circumstance.

As such, the takeaways from this thesis are a few. First, these features, when context, genre, register, and audience are appropriately taken into consideration, largely work in identifying likely gender. Second, when these features can be paired against a general baseline like the eBAC, or better yet a specific baseline like the known language of suspect authors, their findings are more accurate and robust. And third, it is not simply enough for an analyst to apply the features and run with their output, but rather consider (again based on the context of each individual case) whether their distribution is likely indicative of natural or overt gender performance. Methodologically, then, while this approach is couched in various limitations and must be undertaken with various other steps in mind other than a one-fits-all, purely mathematical approach, this thesis nevertheless proposes numerous findings for the

linguistic theories discussed throughout as well.

Audience design in particular came into play in all four datasets. It is perhaps obvious that this would be the case with the AGGID blog, and the two forensic cases, where deception was overt, but it also proved to be the case in the Catfi.sh Tinder data. In every respect, the attention paid to audience design was clearly performative, even in the After phases of the Tinder data, where authors tended to overtly perform their maleness to an audience they now knew was male. In the four datasets in Chapter 3, audience design caused similar issues to those in the earlier Tinder phases, though not for the same reasons. The Tumblr post about killing a husband in the Sims expected only an audience of other Sims players situated in the same context, working on the same commodity of cooperation, and discounted the possibility of suspicion as a commodity for bystanders, overhearers, and eavesdroppers who were *not* the target audience. In the Tinder data, authors similarly focused on their primary audience over the possibility of any broader audience, but the issue was instead one of perception: their audience, while composed of the same entity they were targeting and who was providing feedback, did not share the identity demographics with the author's perception.

Accommodation theory, too, seemed the most prevalent in the Tinder data, where authors accommodated their writing with more female-coded distributions in the Before phases to a female audience, and more male-coded distributions in the After phases to a male audience. This may partially explain why the Before phase accounts for such a large portion of the data (roughly 80%), and some interactions go on for hundreds of turns over multiple exchanges. Accommodation may also have played a part in more of the Potter data than was provided to the court, and motivated Jamie's claim that Jenelle had a clear female performance, while Chris had a clear male performance—that is, in matters of linguistic identity deception, an author may accommodate their performance not to their audience's identity but to what their audience would perceive as genuine performance of the author's proposed identity. This, then, would leave accommodation as a performative aspect of identity, though it is in conjunction with the perception of a social network that accommodation is crafted and performed, thus leaving it, like audience design, and least partially perceptive. In all cases, accommodation seems to work only when the commodity is cooperation, as when suspicion is at the forefront, in instances as benign as the concept of a "G.I.R.L." in Chapter 1, and as serious as UCOs on darkweb child porn rings discussed in Chapter 4, any attempt at overt accommodation can be seen as infelicitous and suspicious.

Communities of practice, speech communities, and social networks can be seen as both performative and perceptive, and came into play in both audience design and accommodation in many instances. Justin Carter tailored his Facebook posts to a community of practice of other League of Legends players who participated in similarly hyperbolic language, but in discounting his larger audience and failing to accommodate to the bystanders, overhearers, and eavesdroppers likely on a site like Facebook, suffered dire legal consequences. A speech community, however small, likely played some part in Hemmert, in

which Lael and Amber often exhibited feature distributions much, much closer to one another than they did to anything else. They were, after all, not only husband and wife, but also very well versed in texting back and forth with one another because of Lael's job and constant travel away from home, and as such could have developed similar texting styles that made a dataset as small as the Q difficult to parse. And social networks were seen at play in the AGGID case, where Tom situated Amina's identity performance in a social network that valued cooperation as a commodity for ingroup members, into which Amina neatly fell.

Finally, speaker contamination was most prevalent in the Tinder data and Potter case for fairly similar reasons, both of which had to do with cooperative commodity. In the Tinder data, the profile information falsely provided to authors greatly skewed their perception of their conversee's language; with cooperation as such a high commodity in such exchanges, many of the red flags discussed in Chapter 4 were neatly navigated as acceptable deviations by authors on both sides. In the Potter case, Jamie's preconceptions of Jenelle (who, again, may have performed an overtly female identity with Jamie in data or in-person interactions not provided to court) as a female, and of "Chris" as a male led to his statement as to *why* he was convinced that "Chris" was real. However, as the data itself showed, the two specific features Jamie found to be most convincing from "Chris" were actually not as male coded as Jamie perceived. Speaker contamination, thus, can be seen as entirely perceptive, and falls in line with the various perceptive misconceptions about gender performance discussed in Chapter 2, such as the idea that women talk more than men, when the reality is often that men can dominate a conversation while still harboring the perception that the women spoke more because of when and how women took their turns and either party elaborated upon various topics.

If essentialist approaches tend to focus on differences between opposing identity categories, such as male and female, and constructionist approaches tend to focus on the difference within a group, such as the range of variation within either gender, then the application of the above variables and theories would seek to do both given the context of the research question. That is, the model proposed in this paper handled both differences between gender identity categories (such as the difference between the Tom and Amina performances in the AGGID blog, and in the different distribution language features of the husband and wife in Hemmert) and within an individual gender identity category (female for Jenelle between Jamie and Facebook at large in Potter, and male for the Catfi.sh Tinder data in the Before and After phases, for example), as well as the various relevant intersectional categories as mentioned above. While the variables at their outset would rely on an essentialist binary between male and female performances, the conjunction with considering contextually relevant linguistic theories allows for a jointly constructionist consideration of how gradient feminine and masculine performances might indicate overtness or covertness, genuineness or disguise, or interplay with other identity facets or exigent factors.

The original contributions of this study are as follows:

1. Qualitative analysis of smaller or less robust datasets, when based on quantitative features found to have statistical gain for, in this case, gender coding, works. This was demonstrated in a detailed qualitative analysis of four major datasets, and supports non-statistical approaches for forensic problems, which, as discussed, may fall short in some ways of the ideals set by experimental data.
2. Such features work when they are specifically situated in the proper context of the dataset, including considerations for genre, register, audience, community of practice, commodity, reference and authorial baselines, and so on. The features refined through the eBAC did not work on every genre to which they were applied, in whole or in part, because of, for example, the register issues with something like the AGGID blog, or the contextual issues of family in the Hemmert case.
3. Interactional gender confusion is as much perceptive as it is performative. That is, if an author perceives interactional gender confusion (whether they perceive their own gender to be confused or their own perception of an interlocutor is shifting, and whether or not someone is actively attempting to guise their gender), their own gender performance will be amplified to the point of the overt features mentioned in, for example, the Catfi.sh Tinder data and Potter. Such performances often conform to the partialness principle, in that they are likely to be amplified in both conscious and unconscious ways.

That is, the features proposed in this thesis can be useful in qualitative applications for forensic problems of authorship analysis and linguistic profiling, both by finding a consistent constellation of gender-coded features as established by an appropriate baseline, and by finding overtly performative gender-coded features. Whether the latter overtly performed features are indicative of someone actively engaging in gender performance for reasons of disguise or otherwise—such as the difference between a woman disguising herself as a man such as in Potter, and a man hyper performing his masculinity such as in the Catfi.sh Tinder data—requires situating features and their distributions within the appropriate context.

Further research, then, may consider the context-specific usefulness of these and other features, the interplay of gender performance with other facets of identity, and how the perceptive aspect of performance may have an impact on the analysis of forensic problems. The latter, most bolstered by the secondary findings of the Catfi.sh Tinder data that authors may have a difficult time separating linguistic cues of identity with the infelicitous baseline identities established by the catfished profiles they are presented with, may have applications in, for example, police decoy cases, both in how UCOs might better establish and maintain an appropriate baseline identity when suspicion is the commodity, and in explaining how victims of poor policing procedures may become entrapped by defaulting to such profiles and not the subsequent linguistic features or revelations provided by further interaction.

As mentioned in Chapter 1's introduction, when we interact online, we lose a lot of the useful information that other interaction types can provide. We are, as such, first beholden to

what we are presented with—whether that be a convincingly verbose blog all about the author’s purported identity, a Tinder profile with the barest bones information to securely suggest a female user identity, text messages coming from the known number and informationally situated in owner’s identity, or even in emails and Facebook messages through someone else’s account who assures us that this secondary identity exists and plants ample non-linguistic evidence to that effect—and are left to navigate the identities with which we communicate on a baseline of cooperation or suspicion that may, as that person’s identity continues to further emerge through performance, gradually or abruptly begin to shift.

Appendices

All data used in this thesis is outlined in the table below. Non-forensic data is provided either in the Chapter in which it is discussed, or in the linked Appendices. Forensic data is not provided. Also included in Appendix B are the regular expressions and lemmas used for the 20 features considered in this thesis, as well as the parser used to collect them.

Dataset	Provided In	Chapter(s)	Link
Ashley Madison hacks	Chapter 1	1	none
“Erin Princess Baby”	not provided	1	none
A Gay Girl in Damascus blog	Appendix C	4	https://drive.google.com/drive/folders/1dGwYXh05ugtjiegfVlJwq_yqOjTxsQTzi
eBAC	Appendix A	3, 4, 5	https://drive.google.com/drive/folders/1d2NzkbUJL5r1LRYVwmkVdGrEo9_2aJ_i
Features	Appendix B	2, 3, 4, 5	https://drive.google.com/drive/folders/1ZvN3VGFMZaay12xDXqU6To0KC2BrEvKW
Catfi.sh Tinder conversations	Appendix D	4	https://drive.google.com/drive/folders/14vmMwnNg-aLdHhcz7lZjhk3KPa3a1YN3
Tumblr post	Chapter 3	3	none
“texts from Mom”	Chapter 3	3	none
Justin Carter Facebook post	Chapter 3	3	none
Tyson Leon tweet	Chapter 3	3	none
Darkweb UCO example	not provided	4	none
Hemmert texts	Appendix E	5	https://drive.google.com/drive/folders/1yBT_ZBgQU-xOhqzcjgp7U9cKuz3fJ28O
Potter data	Appendix F	5	https://drive.google.com/drive/folders/1nq5p-FzeTLOX87iyeW3JdM0kbZGailQa

All appendices are available at the following link:

<https://drive.google.com/drive/folders/19sny9wwNaut3WwslMAf-Z2u9dtJmTp78>

References

- Aijmer, K. (1986). Discourse variation and hedging. *Corpus Linguistics II. New studies in the analysis and exploitation of computer corpora*, 1-18.
- Ainsworth, J., & Juola, P. (2018). Who wrote this: Modern forensic authorship analysis as a model for valid forensic science. *Wash. UL Rev.*, 96, 1159.
- Ainsworth-Vaughn, N. (1992). Topic transitions in physician-patient interviews: Power, gender, and discourse change. *Language in Society*, 409-426.
- Androutsopoulos, J. (2006). Introduction: Sociolinguistics and computer-mediated communication. *Journal of sociolinguistics*, 10(4), 419-438.
- Balter, L. (1999). *Child psychology: A handbook of contemporary issues*. Philadelphia, PA: Psychology Press.
- Baker, M. A. (1991). Gender and verbal communication in professional settings: A review of research. *Management Communication Quarterly*, 5(1), 36-63.
- Bamman, D., Eisenstein, J., & Schnoebelen, T. (2014). Gender identity and lexical variation in social media. *Journal of Sociolinguistics*, 18(2), 135-160.
- Banakou, D., & Chorianopoulos, K. (2010). The effects of avatars' gender and appearance on social behavior in online 3D virtual worlds. *Journal For Virtual Worlds Research*, 2(5).
- Bell, A. (1984). Language style as audience design. *Language in society*, 13(2), 145-204.
- Bell, A. (2001). Back in style: reworking audience design. *Style and sociolinguistic variation*, eds Eckert P & Rickford JR.
- Bell, A. (2014). *The Guidebook to Sociolinguistics*. Chichester, UK: Wiley Blackwell.
- Bell, M., & Flock, E. (2011). 'A Gay Girl in Damascus' comes clean. *Washington Post*, 12.
- Bessi re, K., Seay, A. F., & Kiesler, S. (2007). The ideal elf: Identity exploration in World of Warcraft. *Cyberpsychology & behavior*, 10(4), 530-535.
- Blake, J. J., Eun Sook, K., & Lease, A. (2011). Exploring the Incremental Validity of Nonverbal Social Aggression: The Utility of Peer Nominations. *Merrill-Palmer Quarterly*, 57(3), 293-318.
- Boulis, C., & Ostendorf, M. (2005, June). A quantitative analysis of lexical differences between genders in telephone conversations. In *Proceedings of the 43rd Annual Meeting of the Association for Computational Linguistics (ACL'05)* (pp. 435-442).
- BreakBrunch. (2017, September 23). Welcome to the jungle of online dating. Retrieved April 13, 2021, from <https://breakbrunch.com/welcome-to-the-jungle-of-online-dating/>
- Bresnahan, M. I., & Cai, D. H. (1996). Gender and aggression in the recognition of interruption. *Discourse processes*, 21(2), 171-189.

- Brezenoff, S. (2015). *Guy in real life*. New York, NY: Balzer + Bray.
- Brown, P., & Levinson, S. C. (1987). *Politeness: Some universals in language usage* (Vol. 4). Cambridge university press.
- Bucholtz, M., & Hall, K. (2004). Language and identity. *A companion to linguistic anthropology*, 1, 369-394.
- Bucholtz, M., & Hall, K. (2005). Identity and interaction: A sociocultural linguistic approach. *Discourse studies*, 7(4-5), 585-614
- Bustle.com. Retrieved June 15, 2017 from: <https://www.bustle.com/articles/172726-the-top-10-emojis-used-on-tinder-in-honor-of-world-emoji-day>
- Calix, K., Connors, M., Levy, D., Manzar, H., McCabe, G., & Westcott, S. (2008). Stylometry for e-mail author identification and authentication. *Proceedings of CSIS research day, Pace University*, 1048-1054.
- Calnan, A. C., & Davidson, M. J. (1998). The impact of gender and its interaction with role and status on the use of tag questions in meetings. *Women in Management Review*.
- Cameron, D. (1992). *Feminism and linguistic theory*. Springer.
- Cameron, D. (2010). Gender, Language, and the New Biologism. *Constellations*, (4), 526. doi:10.1111/j.1467-8675.2010.00612.x
- Cameron, D. (2014). 14 Gender and Language Ideologies. *The handbook of language, gender, and sexuality*, 281.
- Carli, L. L. (1990). Gender, language, and influence. *Journal of personality and social psychology*, 59(5), 941.
- Catfi.sh (2015) *About Catfish*. Retrieved January 3, 2017 from: <https://catfi.sh/about>
- Catfi.sh (2015) *Chris and Connor*. Retrieved January 3, 2017 from: <https://catfi.sh/conversation/2750>
- Catfi.sh (2015) *Daniel and Joe*. Retrieved January 3, 2017 from: <https://catfi.sh/conversation/2770>
- Catfi.sh (2015) *Martyn and William*. Retrieved January 3, 2017 from: <https://catfi.sh/conversation/1053>
- Catfi.sh (2015) *Olly and Jack*. Retrieved January 3, 2017 from: <https://catfi.sh/conversation/2575>
- Catfi.sh (2015) *Will and Dan*. Retrieved January 3, 2017 from: <https://catfi.sh/conversation/2637>
- “chalupamonk”. (2003, March 3). Wtf. Retrieved April 13, 2021, from <http://www.urbandictionary.com/define.php?term=wtf&defid=48041>
- Cheng, N., Chandramouli, R., & Subbalakshmi, K. P. (2011). Author gender identification from text. *Digital Investigation*, 8(1), 78-88.

- Chiang, E. (2019). *Rhetorical moves and identity performance in online child sexual abuse interactions* (Doctoral dissertation, Aston University).
- Chiang, E., & Grant, T. (2019). Deceptive identity performance: offender moves and multiple identities in online child abuse conversations. *Applied Linguistics*, 40(4), 675-698.
- Chiang, E. (2021). 'Send Me Some Pics': Performing the Offender Identity in Online Undercover Child Abuse Investigations. *Policing: A Journal of Policy and Practice*, 15(2), 1173-1187.
- Cho, S. H. (2007). Effects of motivations and gender on adolescents' self-disclosure in online chatting. *CyberPsychology & Behavior*, 10(3), 339-345.
- Clore, G. L., & Ortony, A. (1988). The semantics of the affective lexicon. In *Cognitive perspectives on emotion and motivation* (pp. 367-397). Springer, Dordrecht.
- Coates, J. (2004) *Women, Men and Language: A Sociolinguistic Account of Gender Differences in Language*. 3rd ed. New York: Routledge, 2004, 5-7.
- Coates, J. (2006). Gender. In C. Llamas, L. Mullany, & P. Stockwell (Authors), *The Routledge companion to sociolinguistics*. Abingdon England: Routledge.
- Coates, J. (2013). *Women, men and everyday talk*. Springer.
- Crawford, M. (1995). *Talking difference: On gender and language*. Sage.
- Crystal, D. (2001). *Language and the Internet*. Cambridge, New York: Cambridge University Press.
- Crystal, D. (February 2005). The scope of Internet Linguistics, Paper given online to the American Association for the advancement of Science meeting. Retrieved October 7, 2013, from http://w.davidcrystal.com/DC_articles/Internet2.pdf
- Crystal, D. (2011). *Internet Linguistics*. Cambridge University Press.
- Cupach, W. R., & Spitzberg, B. H. (2011). "Girls' Social Aggression." *The Dark Side of Close Relationships II*. New York: Routledge, 2011. 297-316. Print
- Danet, B. (1998). Text as mask: gender, play and performance. *Cybersociety*, 2, 129-158.
- "Definition of edgelord". UrbanDictionary.com. Retrieved November 1, 2016, from <http://www.urbandictionary.com/define.php?term=edgelord>
- "Definition of WTF". UrbanDictionary.com. Retrieved November 1, 2016, from <http://www.urbandictionary.com/define.php?term=WTF>
- DeFrancisco, V. (1991). "The sound of silence: how men silence women in marital relationships." *Discourse and Society* 2 (4):413-24.
- Dern, H., Dern, C., Horn, A., & Horn, U. (2009). The fire behind the smoke: A reply to Snook and colleagues. *Criminal justice and behavior*, 36(10), 1085-1090.
- Dickson, E. J. (2020, March 02). Tinder is being plagued by these mobile game bots. Retrieved April 13, 2021, from <https://www.dailydot.com/debug/tinder-hack-bots/>

- Dindia, K., & Allen, M. (1992). Sex differences in self-disclosure: a meta-analysis. *Psychological bulletin*, 112(1), 106.
- Dorval, B. (1990). *Conversational Organization and its Development*. Norwood, NJ: Ablex.
- "Dovas". (2015). 22 reasons why parents shouldn't text. Retrieved April 13, 2021, from <https://www.boredpanda.com/funny-parent-texting-fails/>
- Eckert, P., & McConnell-Ginet, S. (1992). Think practically and look locally: Language and gender as community-based practice. *Annual review of anthropology*, 21(1), 461-488.
- Edelsky, C., & Adams, K. (1990). Creating inequality: Breaking the rules in debates. *Journal of language and social psychology*, 9(3), 171-190.
- Eklund, L. (2011). Doing gender in cyberspace: The performance of gender by female World of Warcraft players. *Convergence*, 17(3), 323-342.
- Eliasoph, N. (1987). Politeness, power, and women's language: Rethinking study in language and gender. *Berkeley Journal of Sociology*, 32, 79-103.
- Erard, M. (2017). Write yourself invisible. *New Scientist*. 236(3153):36-39. doi:10.1016/S0262-4079(17)32310-2
- Ersoy, S. (2009). "Men compete, woman collaborate A study on collaborative vs: competitive communication styles in mixed-sex conversation."
- Felleggy, A. M. (1995). Patterns and Functions of Minimal Response. *American Speech*, 70(2), 186–199. <http://doi.org/10.2307/455815>
- Fishman, P. M. (1978). Interaction: The work women do. *Social Problems*, 25(4), 397–406. <https://doi.org/10.1525/sp.1978.25.4.03a00050>
- Fobbe, E. (2021). Text-Linguistic Analysis in Forensic Authorship Attribution. *International Journal of Language & Law (JLL)*, 9.
- Freed, A. (1992, April). We understand perfectly: A critique of Tannen's view of cross-sex communication. In *Locating power: Proceedings of the second Berkeley women and language conference* (Vol. 1, pp. 144-152). Berkeley: Berkeley Women and Language Group.
- Freed, A., & Greenwood, A. (1996). Women, men and type of talk: What makes the difference? *Lang. Soc. Language in Society*, 25(1), 1-26.
- Giles, H. (2016). *Communication accommodation theory: Negotiating personal relationships and social identities across contexts*. Cambridge, United Kingdom: Cambridge University Press.
- Giles, H., Coupland, J., & Coupland, N. (1991). 1. Accommodation theory: Communication, context, and. *Contexts of accommodation: Developments in applied sociolinguistics*, 1.
- Goodwin, M. (1990). *He-said-she-said. Talk as Social Organisation among Black Children*. Bloomington: Indian Press.

- Grant, T., & Macleod, N. (2018). Resources and constraints in linguistic identity performance: a theory of authorship. *Language and Law= Linguagem e Direito*, 5(1), 80-96.
- Grant, T., & MacLeod, N. (2020). *Language and Online Identities: The Undercover Policing of Internet Sexual Crime*. Cambridge: Cambridge University Press.
doi:10.1017/9781108766425
- Gray, J. (1992). *Men are from Mars, women are from Venus*. HarperCollins.
- Hancock, A., Colton, L., & Douglas, F. (2014). Intonation and gender perception: Applications for transgender speakers. *Journal of Voice*, 28(2), 203-209.
- Hayes, R. (2020, October 2). How to tell if a tinder profile is fake (or a bot). Retrieved April 13, 2021, from <https://social.techjunkie.com/tinder-profile-fake-bot/>
- He, Y. (2010). An Analysis of Gender Differences in Minimal Responses in the conversations in the two TV-series *Growing Pains* and *Boy Meets World*.
- Hepburn, A., & Potter, J. (2011). Recipients designed: Tag questions and gender. *Conversation and gender*, 135-152.
- Herring, S. C., Scheidt, L. A., Bonus, S., & Wright, E. (2004, January). Bridging the gap: A genre analysis of weblogs. In *37th Annual Hawaii International Conference on System Sciences, 2004. Proceedings of the* (pp. 11-pp). IEEE.
- Herring, S. C., & Martinson, A. (2004). Assessing gender authenticity in computer-mediated language use: Evidence from an identity game. *Journal of Language and Social Psychology*, 23(4), 424-446.
- Holmes, J. (2001). *An Introduction to Sociolinguistics*. London, UK: Longman.
- Holmes, J., & Meyerhoff, M. (2003). Different voices, different views: An introduction to current research in language and gender. *The handbook of language and gender*, 1-17.
- Howard, L. (2015). The Top 10 Emojis Used On Tinder, In Honor Of World Emoji Day.
- Huh, S., & Williams, D. (2010). Dude looks like a lady: Gender swapping in an online game. In *Online worlds: Convergence of the real and the virtual* (pp. 161-174). Springer, London.
- Imgur. (2017, January 19). More fun with tinder bots! Retrieved April 13, 2021, from <https://imgur.com/gallery/biF71>
- Jaidka, K., Guntuku, S., & Ungar, L. (2018, June). Facebook versus Twitter: Differences in Self-Disclosure and Trait Prediction. In *Proceedings of the International AAAI Conference on Web and Social Media* (Vol. 12, No. 1).
- Jespersen, O. (1922). *Language: Its nature, development and origin*. London: G. Allen & Unwin ;.
- Jiang, H. (2011). Gender Difference in English Intonation. In *ICPhS* (pp. 974-977).
- Johannsen, A., Hovy, D., & Søgaard, A. (2015, July). Cross-lingual syntactic variation over

- age and gender. In *Proceedings of the nineteenth conference on computational natural language learning* (pp. 103-112).
- Johnstone, B. (2009). Chapter 4: Participants in Discourse: Relationships, Roles, and Identities. In *Discourse Analysis* (pp 76-127). Malden, MA: Blackwell
- Johnstone, B. (2009). Stance, style, and the linguistic individual. *Stance: sociolinguistic perspectives*, 29, 52.
- Juola, P. (2008). *Authorship attribution* (Vol. 3). Now Publishers Inc.
- Just Something (2017, May 18). The 36 funniest text ever sent from parents to their kids. I couldn't help laughing at #9! Retrieved April 13, 2021, from <https://justsomething.co/36-funniest-text-from-parents/>
- Kendall, S., & Tannen, D. (1997). Gender, power and practice: or, putting your money (and your research) where your mouth is. *Gender and Discourse*. London: Sage, 81-105.
- Kim, J., & Dindia, K. (2011). Online self-disclosure: A review of research. *Computer-mediated communication in personal relationships*, 156-180.
- Koppel, M., Argamon, S., & Shimon, A. R. (2002). Automatically categorizing written texts by author gender. *Literary and linguistic computing*, 17(4), 401-412.
- Koppel, M., Schler, J., & Argamon, S. (2009). Computational methods in authorship attribution. *Journal of the American Society for information Science and Technology*, 60(1), 9-26.
- Koppel, M., Schler, J., & Argamon, S. (2012). Authorship attribution: What's easy and what's hard. *JL & Pol'y*, 21, 317.
- Lakoff, R. (1975). *Language and woman's place*. New York: Harper & Row.
- League of Legends (4/27/2013). *Home > Forums > League of Legends > GeneralDiscussion > "18 year old facing 8 years prison because of LoL comment..."* Retrieved November 19, 2014, from: <http://forums.eune.leagueoflegends.com/board/showthread.php?t=595222>
- Leaper, C., Carson, M., Baker, C., Holliday, H., & Myers, S. (1995). Self-disclosure and listener verbal support in same-gender and cross-gender friends' conversations. *Sex Roles*, 33(5-6), 387-404.
- Leavitt, A. (2015, February). "This is a Throwaway Account" Temporary Technical Identities and Perceptions of Anonymity in a Massive Online Community. In *Proceedings of the 18th ACM conference on computer supported cooperative work & social computing* (pp. 317-327).
- Lincoln, R., & Coyle, I. R. (2012). No-one Knows You're a Dog on the Internet: Implications for Proactive Police Investigation of Sexual Offenders, Psychiatry, Psychology and Law, DOI:10.1080/13218719.2012.672274.
- LINE. (n/d). *Tinder clever* (page 1). Retrieved April 13, 2021, from <https://line.17qq.com/articles/suqahqurhx.html>

- Malislow, C. (Feb. 13, 2014). *The Facebook Comment that Ruined a Life*. The Dallas Observer Online. Retrieved November 19, 2014, from: <http://www.dallasobserver.com/2014-02-13/news/the-facebook-comment-that-ruined-a-life/full/>
- Maltz, D., & Borker, R. (1982). A cultural approach to male-female miscommunication. S.I.: [s.n.].
- MacLeod N., & Grant T. (2017). “Go on Cam but Dnt Be Dirty”: Linguistic Levels of Identity Assumption in Undercover Online Operations against Child Sex Abusers.’ *Language and Law/Linguagem e Direito* 4(2):157–175.
- MacLeod, N., & Grant, T. (2021). Assuming Identities Online: How Linguistics Is Helping the Policing of Online Grooming and the Distribution of Abusive Images. In *Rethinking Cybercrime* (pp. 87-104). Palgrave Macmillan, Cham.
- Mills, S. (2002). Rethinking politeness, impoliteness and gender identity. *Gender identity and discourse analysis*, 69, 90.
- Mills, S. (2003). *Gender and politeness* (No. 17). Cambridge University Press.
- Mondorf, B. (2002). Gender differences in English syntax. *Journal of English Linguistics*, 30(2), 158-180.
- Mondorf, B. (2011). *Gender differences in English syntax* (Vol. 491). Walter de Gruyter.
- Moreau, E. (2020, September 23). These 10 emoji probably don't mean what you think they mean. Retrieved April 13, 2021, from <https://www.lifewire.com/less-obvious-emoji-meanings-3485884>
- Mulac, A. (2006). *The gender-linked language effect: Do language differences really make a difference?*. Lawrence Erlbaum Associates Publishers.
- Mulac, A., Erlandson, K. T., Farrar, W. J., Hallett, J. S., Molloy, J. L., & Prescott, M. E. (1998). “Uh-huh. What's that all about?” Differing interpretations of conversational backchannels and questions as sources of miscommunication across gender boundaries. *Communication Research*, 25(6), 641-668.
- Mullany, L. (2004). Gender, politeness and institutional power roles: Humour as a tactic to gain compliance in workplace business meetings. *Multilingua*, 23(1-2), 13-37.
- Narayanan, A., Paskov, H., Gong, N. Z., Bethencourt, J., Stefanov, E., Shin, E. C. R., & Song, D. (2012, May). On the feasibility of internet-scale author identification. In 2012 IEEE Symposium on Security and Privacy (pp. 300-314). IEEE.
- Nass, C., Moon, Y., & Green, N. (1997). Are machines gender neutral? Gender-stereotypic responses to computers with voices. *Journal of applied social psychology*, 27(10), 864-876.
- Newitz, A. (2015). Ashley Madison code shows more women, and more bots. *gizmodo* [Online] <https://gizmodo.com/ashley-madison-code-shows-more-women-and-morebots-1727613924> [09.01. 2019].
- Newman, M. L., Groom, C. J., Handelman, L. D., & Pennebaker, J. W. (2008) “Gender differences in language use: An analysis of 14,000 text samples.” *Discourse Processes* 45, no 3: 211-236.

- O'Barr, W. M., & Atkins, B. K. (1980). "Women's language" or "powerless language"?.
 Okamoto, D. G., & Smith-Lovin, L. (2001). Changing the subject: Gender, status, and the dynamics of topic change. *American Sociological Review*, 852-873.
 Palomares, N. A., & Lee, E. J. (2010). Virtual gender identity: The linguistic assimilation to gendered avatars in computer-mediated communication. *Journal of Language and Social Psychology*, 29(1), 5-23.
 Pasfield-Neofitou, S. E. (2007). The gender differential use of minimal responses in daytime TV interviews: a preliminary investigation. *Monash University Linguistics Papers*, 5(2), 43-52.
 Petrus4. (2015, February 19). Edgelord. Retrieved April 13, 2021, from <http://www.urbandictionary.com/define.php?term=edgelord&defid=8113420>
 Postmes, T., & Spears, R. (2002). Behavior online: Does anonymous computer communication reduce gender inequality?. *Personality and Social Psychology Bulletin*, 28(8), 1073-1083.
 Provine, R., Spencer, R., & Mandel, D. (2007) Emotional Expression Online: Emoticons Punctuate Website Text Messages. *Journal of Language and Social Psychology*, 26(3), 299-307.
 Rao, D., Yarowsky, D., Shreevats, A., & Gupta, M. (2010, October). Classifying latent user attributes in twitter. In *Proceedings of the 2nd international workshop on Search and mining user-generated contents* (pp. 37-44).
 Rashid, A., Baron, A., Rayson, P., May-Chahal, C., Greenwood, P., & Walkerdine, J. (2013). Who am I? Analysing Digital Personas in Cybercrime Investigations. *Computer*, 46(4), 54-61. <https://doi.org/10.1109/MC.2013.68>
 Reid, J. (1995). A study of gender differences in minimal responses. *Journal of Pragmatics*, 24(5), 489-512.
 Riedl, R., Hubert, M., & Kenning, P. (2010). Are there neural gender differences in online trust? An fMRI study on the perceived trustworthiness of eBay offers. *MIS quarterly*, 397-428.
 Robinson, L. F., & Reis, H. T. (1989). The effects of interruption, gender, and status on interpersonal perceptions. *Journal of nonverbal behavior*, 13(3), 141-153.
 Rogers S., Fay N., Mayberry M., & Roulin, A. (2013). Audience Design through Social Interaction during Group Discussion. *PLoS ONE* 8(2), E57211.
 Schegloff, E. A. (2000). Overlapping talk and the organization of turn-taking for conversation. *Language in society*, 29(1), 1-63.
 Schler, J., Koppel, M., Argamon, S., & Pennebaker, J. W. (2006, March). Effects of age and gender on blogging. In *AAAI spring symposium: Computational approaches to analyzing weblogs* (Vol. 6, pp. 199-205).
 Shuy, R. (1993). *Language Crimes: The Use and Abuse of Language Evidence in the Courtroom*. Cambridge, MA: Blackwell Publishing.

- Singh, S. (2001). A pilot study on gender differences in conversational speech on lexical richness measures. *Literary and Linguistic Computing*, 16(3), 251-264.
- SirPainsalot. (2017, March 14). Girl. Retrieved April 13, 2021, from <https://www.urbandictionary.com/define.php?term=Girl>
-  smiling face with open mouth and cold Sweat Emoji. (n/d). Retrieved April 13, 2021, from <https://emojipedia.org/smiling-face-with-open-mouth-and-cold-sweat/>
- Snook, B., Cullen, R. M., Bennell, C., Taylor, P. J., & Gendreau, P. (2008). The criminal profiling illusion: What's behind the smoke and mirrors?. *Criminal Justice and Behavior*, 35(10), 1257-1276.
- Song, H. G., Restivo, M., van de Rijt, A., Scarlatos, L., Tonjes, D., & Orlov, A. (2015). The hidden gender effect in online collaboration: An experimental study of team performance under anonymity. *Computers in human Behavior*, 50, 274-282.
- Speer, S. A. (2005). The interactional organization of the gender attribution process. *Sociology*, 39(1), 67-87.
- Speer, S. A., & Stokoe, E. (Eds.). (2011). *Conversation and gender*. Cambridge University Press.
- Spender, D. (1980). *Man Made Language*. London: Routledge and Kegan Paul.
- Stamatatos, E. (2008). Author identification: Using text sampling to handle the class imbalance problem. *Information Processing & Management*, 44(2), 790-799.
- Stamatatos, E. (2009). A survey of modern authorship attribution methods. *Journal of the American Society for information Science and Technology*, 60(3), 538-556.
- Stokoe, E. H. (2005). Analysing gender and language. *Journal of Sociolinguistics*, 9(1), 118-133.
- Swofford, H., & Champod, C. (2021). Implementation of algorithms in pattern & impression evidence: A responsible and practical roadmap. *Forensic Science International: Synergy*, 100142.
- Sung, C. M. (2012). Exploring the interplay of gender, discourse, and (im)politeness. *Journal Of Gender Studies*, 21(3), 285-300. doi:10.1080/09589236.2012.681179
- Swann, J., & Graddol, D. (1988). Gender inequalities in classroom talk. *English in education*, 22(1), 48-65.
- Tannen, D. (1991). *You just don't understand: Women and men in conversation*(pp. 1990-1990). London: Virago.
- Tannen, D. (1992). *You Just Don't Understand: Women and Men in Conversation*. London: Virgo press.
- Tannen, D. (1996). *Gender and discourse*. Oxford University Press.
- Thurlow, C., & Poff, M. (2013). Text messaging. *Pragmatics of computer-mediated communication*, 9, 163-190.

- Tomlinson, J., & Tree, J. (2011). Listeners' comprehension of uptalk in spontaneous speech. *Cognition*, 119, 58-69.
- Troemel-Ploetz, S. (1991). Review essay: Selling the apolitical. *Discourse & Society*, 2(4), 489-502.
- Tudury, L. (n/d). Eyes emoji - what does the eyes emoji mean? Retrieved April 13, 2021, from <http://www.dictionary.com/meaning/eyes-emoji>
- Tumblr (n.d.). <http://www.tumblr.com>
- TV Tropes. (n/d). G.I.R.L. Retrieved April 13, 2021, from <https://tvtropes.org/pmwiki/pmwiki.php/Main/GIRL>
- Twitter (n.d.). <http://www.twitter.com>
- Underwood, M. K. (2011). "Collin College Interdisciplinary Undergraduate Student Research Conference | Keynote Speaker." Collin College. 2011.
- Vorobeva, A. A. (2016). Forensic linguistics: automatic web author identification. *Научно-технический вестник информационных технологий, механики и оптики*, 16(2).
- Walther, J. B. (1996). Computer-mediated communication: Impersonal, interpersonal, and hyperpersonal interaction. *Communication research*, 23(1), 3-43.
- Wang, H. Y., & Wang, Y. S. (2008). Gender differences in the perception and acceptance of online games. *British journal of educational technology*, 39(5), 787-806.
- West, C., & Zimmerman, D. H. (1983). "Small insults: a study of interruptions in conversations between unacquainted persons." In B. Thorne, C Kramarae, and N. Henley (eds.) *Language, Gender and Society*, 102-17. Rowley, MA: Newbury House.
- Williams, D., Consalvo, M., Caplan, S., & Yee, N. (2009). Looking for gender: Gender roles and behaviors among online gamers. *Journal of communication*, 59(4), 700-725.
- Wright, D. (2014). *Stylistics versus Statistics: A corpus linguistic approach to combining techniques in forensic authorship analysis using Enron emails* (Doctoral dissertation, University of Leeds).
- Wright, D. (2013). Using corpora in forensic authorship analysis: Investigating idiolect in Enron emails. *Corpus Linguistics* 2013, 296.
- Wright, D. (2013). Stylistic variation within genre conventions in the Enron email corpus: developing a textsensitive methodology for authorship research. *International Journal of Speech, Language & the Law*, 20(1).
- Youn, S., & Hall, K. (2008). Gender and online privacy among teens: Risk perception, privacy concerns, and protection behaviors. *Cyberpsychology & behavior*, 11(6), 763-765.
- Zimmermann, D. H., & West, C. (1996). Sex roles, interruptions and silences in conversation. *AMSTERDAM STUDIES IN THE THEORY AND HISTORY OF LINGUISTIC SCIENCE SERIES* 4, 211-236.