

UCLA

UCLA Electronic Theses and Dissertations

Title

A Corpus-based Cognitive-Functional Study of the Meaning and Use of 'Always' and 'Never', and Related Phenomena, in American English

Permalink

<https://escholarship.org/uc/item/36n1j6gp>

Author

Lindley, Jori

Publication Date

2015

Peer reviewed|Thesis/dissertation

UNIVERSITY OF CALIFORNIA

Los Angeles

A Corpus-based Cognitive-Functional Study of the Meaning and Use of
Always and *Never*, and Related Phenomena, in American English

A dissertation submitted in partial satisfaction of the
requirements for the degree Doctor of Philosophy
in Applied Linguistics

by

Jori Lindley

2015

© Copyright by

Jori Lindley

2015

ABSTRACT OF THE DISSERTATION

A Corpus-based Cognitive-Functional Study of the Meaning and Use of
Always and *Never*, and Related Phenomena, in American English

by

Jori Lindley

Doctor of Philosophy in Applied Linguistics

University of California, Los Angeles, 2015

Professor Hongyin Tao, Chair

This is a multi-faceted corpus study of the adverbs of frequency *always* and *never*, in which their meanings (exaggerated or literal), tense-aspect preferences, and functions (specifically, the function of *always* or *never* followed by the progressive) across genres are all investigated. I also briefly investigate the maximizers *utterly*, *completely*, *totally*, and *fully*. Using four corpora and quantitative and qualitative methods, I show that *always* and *never* are not as straightforward as they appear. On the contrary, their distribution, meaning, and use are highly dependent on context, both in a larger sense (i.e., genre, pragmatic concerns) and a more specific, local sense (i.e., the immediate linguistic environment, including the verbs, tense-aspect, etc.). For example, I argue that concerns about accountability explain the observed rates of exaggeration across different types of news articles and across academic journals in different fields, and that social concerns involving politeness explain the finding that the grammatical subject in complaints is more often third person than second person. Throughout, it is shown that a cognitive-functional approach is the most useful for understanding how these very common words are used.

The dissertation of Jori Lindley is approved.

Robert Kirsner

Marianne Celce-Murcia

Hongyin Tao, Committee Chair

University of California, Los Angeles

2015

TABLE OF CONTENTS

CHAPTER 1: INTRODUCTION	1
1.1 TOPIC AND BACKGROUND	1
1.2 PROBLEM STATEMENT, OBJECTIVES, THESIS, AND LIMITATIONS.....	6
1.3 TERMS AND CONCEPTS	7
1.4 SIGNIFICANCE	9
1.5 OVERVIEW	10
CHAPTER 2: LITERATURE REVIEW	12
2.1 COGNITIVE-FUNCTIONAL APPROACHES	12
2.2 ANTONYMS.....	17
2.3 INTENSIFIERS AND MAXIMIZERS	19
2.4 ADVERBS OF FREQUENCY.....	21
2.5 CONCLUSION	25
CHAPTER 3: METHOD	26
3.1 RESEARCH DESIGN	26
3.2 CORPORA.....	29
3.3 RESEARCH INSTRUMENTS AND TYPES OF DATA	34
3.4 ANALYSIS OF COLLECTED DATA.....	37
3.5 LIMITATIONS.....	38
3.6 CONCLUSION.....	39
CHAPTER 4: EXAGGERATION QUOTIENTS & GENRE.....	40
4.1 INTRODUCTION	40
4.2 LITERATURE REVIEW	41
4.3 METHOD	45
4.3.1 Corpora and Data Sets.....	45
4.3.2 Genre Characterizations	45
4.3.3 Exaggeration Quotients	48
4.4 RESULTS.....	49
4.4.1 Exaggeration Quotients.....	49
4.4.2 Negation of <i>Always</i>	53
4.5 DISCUSSION	55
4.5.1 The Nature of Informal Speech	55
4.5.2 Accountability.....	59
4.5.3 The Encyclopedic View of Meaning.....	63
4.6 CONCLUSION.....	65
CHAPTER 5: FUNCTIONS WITH THE PROGRESSIVE	67
5.1 INTRODUCTION	67
5.2 LITERATURE REVIEW	68
5.3 METHOD.....	70
5.3.1 Initial Collection of Tokens	71
5.3.2 Quality Control Check of Tokens	71
5.3.3 Coding of Function	72

5.3.4 Coding of Subject Type.....	75
5.3.5 Division into Genres	76
5.3.6 Data Analysis.....	76
5.4 RESULTS.....	77
5.4.1 Distribution of Functions.....	77
5.4.2 Grammatical Subject.....	80
5.4.3 Genre	82
5.4.4 Differences Between <i>Always</i> and <i>Never</i>	82
5.5 DISCUSSION.....	85
5.5.1 Distribution of Functions	85
5.5.2 Grammatical subject.....	91
5.5.3 Genre	97
5.5.4 Differences Between <i>Always</i> and <i>Never</i>	99
5.5.5 On Generalizations	101
5.6 CONCLUSION.....	103
CHAPTER 6: TENSE-ASPECT	105
6.1 INTRODUCTION	105
6.2 LITERATURE REVIEW	106
6.3 METHOD	107
6.3.1 Initial Collection of Tokens.....	107
6.3.2 Quality Control Check of Tokens.....	107
6.3.3 Coding of Tense-Aspect	108
6.4 RESULTS.....	110
6.4.1 Tense-Aspect.....	111
6.4.2 Verbs.....	113
6.4.3 Set Phrases.....	114
6.5 DISCUSSION	118
6.5.1 Tense-Aspect.....	118
6.5.2 Verbs.....	121
6.5.3 Set Phrases.....	122
6.6 CONCLUSION.....	128
CHAPTER 7: MAXIMIZERS.....	129
7.1 INTRODUCTION.....	129
7.2 LITERATURE REVIEW	129
7.3 METHOD.....	133
7.4 RESULTS.....	134
7.4.1 Non-Scalar Collocates.....	134
7.4.2 Semantic Prosody	136
7.5 DISCUSSION.....	136
7.5.2 Extreme Collocates and Emphasis	137
7.5.2 Non-Scalar Collocates and Construal.....	137
7.6 CONCLUSION.....	141

CHAPTER 8: CONCLUSION	143
8.1. SUMMARY OF FINDINGS AND INTERPRETATIONS	143
8.1.1 Exaggeration.....	143
8.1.2 Functions of <i>Always</i> + Progressive	143
8.1.3 Tense-Aspect.....	144
8.1.4 Maximizers.....	145
8.2 IMPLICATIONS AND SIGNIFICANCE.....	145
8.2.1 A Cognitive-Functional Perspective	145
8.2.2 Furthering Genre Studies	147
8.2.3 Characterizing Similar Words.....	148
8.3 FURTHER RESEARCH.....	149
8.4 CLOSING REMARKS	150
WORKS CITED.....	151

LIST OF FIGURES

Figure 4.1: ExagQ, written & spoken language.....	50
Figure 4.2: ExagQ, written & spoken language, normalized	50
Figure 4.3: ExagQ, academic writing.....	51
Figure 4.4: ExagQ, academic writing, normalized	51
Figure 4.5: ExagQ, written news.....	52
Figure 4.6: ExagQ, written news, normalized	52
Figure 4.7: ExagQ+, all data	53
Figure 4.8: ExagQ-, all data	53
Figure 4.9: <i>Always</i> tokens preceded by <i>not</i> , all data.....	54
Figure 5.1: Function of <i>always</i> + progressive	78
Figure 5.2: Function of <i>always</i> + progressive by subject.....	81
Figure 5.3: Function of <i>always</i> + progressive by genre	82
Figure 5.4: Function of PAF + progressive, <i>always</i> versus <i>never</i>	84
Figure 5.5: Percent of PAF + progressive tokens containing <i>come/go (back/again)</i>	84
Figure 6.1: Total and finite tokens	110
Figure 6.2: Tense-aspect (all finite tokens)	111
Figure 6.3: Aspect (all finite tokens).....	112
Figure 6.4: Tense (all finite tokens).....	112
Figure 6.5: Tense-aspect, written language.....	113
Figure 6.6: Tense-aspect, spoken language.....	113
Figure 6.7: Tense-aspect, academic writing	113
Figure 6.8: Tense-aspect, casual spoken language.....	113
Figure 6.9: Verb preferences	114
Figure 6.10: Verb preferences (by tense-aspect)	114
Figure 6.11: TV <i>as always</i>	115
Figure 6.12: TV <i>always</i>	116
Figure 6.13: Mental verbs (totals).....	120
Figure 6.14: Mental verbs (of ten strongest collocates).....	120
Figure 7.1: Intensifiers	130
Figure 7.2: Scalability of strongest collocates.....	135
Figure 7.3: Scalability of most frequent collocates	135
Figure 7.4: SP of strongest collocates	136
Figure 7.5: SP of most frequent collocates.....	136

ACKNOWLEDGEMENTS

My committee members have helped me in so many ways, academically and otherwise, that I could not begin to enumerate them here. For the dozens of recommendation letters, alone, I am extremely grateful. Thank you for the advice, feedback, and support.

I also need to thank the hundreds of students I had in discussion sections and summer courses throughout my years at UCLA. As convenient as it is to work with students who produce top quality papers with ease, I am thinking, especially, of those of you who earned your grade through long hours spent revising, and many office hour visits. Seeing you accept constructive criticism with grace and even gratitude never failed to humble and inspire me. This project put me in the same position. To complete it, I had to regain a sense of discipline I feared I had lost, work harder than I had ever worked—even through difficult personal circumstances, as I know a few of you had to do, too—and learn to love feedback. Part of how I did that was by drawing on the memories of my most resilient students. Thank you, all of you, for that.

VITA

Jori Lindley earned her bachelor's degree in Linguistics and Russian from the University of Michigan, and, after two years teaching English in Hungary, Ukraine, and the Czech Republic, her master's degree in Applied Linguistics from the University of California, Los Angeles. She also holds a graduate-level certificate in Teaching English as a Second/Foreign Language. Jori has been awarded, by her department, the Marianne Celce-Murcia Outstanding Teaching Award and, by the graduate division, the Dissertation Year Fellowship. A revision of her master's thesis will be published in the June 2015 issue of *Intercultural Pragmatics*.

CHAPTER 1: INTRODUCTION

Writing about “common,” everyday words, Sinclair (1999) explains that “the very frequent words of English form a large proportion of any text, and yet their particular qualities are not fully recognized. They are not given adequate provision in theories of language, and their role is not very clearly described in either grammars or dictionaries” (p. 158). He goes on to argue that these words “need to be studied and described in their own terms,” using corpora. Sharing this sentiment, I have conducted a multi-faceted corpus study of two common yet understudied adverbs of frequency, *always* and *never*.

Making use of four large corpora and quantitative and qualitative methods, I investigate these words’ meanings, tense-aspect preferences, and functions across genres. The findings indicate that these aspects of *always* and *never* are highly dependent on context, in both a larger sense (i.e., genre, and pragmatic concerns such as who is speaking to whom and in what setting) and a local sense (i.e., the details of the immediate linguistic context such as the verbs they modify and those verbs’ tense-aspect). Because *always* is a maximizer in addition to being an adverb of frequency, I also conduct a preliminary analysis of four other maximizers: *completely*, *totally*, *utterly*, and *fully*. Throughout the work, it is argued that a cognitive-functional approach best accounts for the empirical data which I present.

1.1 Topic and Background

Adverbs and adjectives indicating degree and, in particular, those that refer to the extreme of a scale, e.g., maximizers such as *utterly*, have attracted considerable attention among corpus linguists. Antonyms and certain adverbs of frequency, such as *often*, are also well-researched. Yet, *always* and *never* (both of which are adverbs of frequency, and the former a maximizer as well) have been left largely unstudied. This may be due to assumptions that their meanings—unlike the meanings of adverbs that indicate intermediate frequency, such as *often*—are never

ambiguous or context-dependent. This is not true, though; *always* and *never* are complicated, in ways already recognized and in ways yet to be discovered.

The first complication, regarding *always*, involves what is quantified. Sometimes this is time, as when *always* means something like “at all times” or “forever.” It does not mean this every time it is used, though. “Liam always has breakfast,” for example, means Liam has breakfast not incessantly but every morning (Kearns, 2000, p. 136). In such utterances, it is not time that is quantified but something like cases or situations in which the activity is relevant (see, e.g., Von Fintel, 1995, p. 9)—hence both the “at all times” and the “in all relevant cases” meanings appearing in dictionaries, as in the Oxford English Dictionary Online (2015, *always*, defs. 1 and 3). The interesting question is how users know exactly which “case” is relevant. Cognitive linguists would say that, just as we do for so many other words, we rely on context (e.g., the fact that we eat breakfast in the morning) to fully unpack *always*. And, despite the generative stance that semantics and pragmatics are separate systems, even many formal semanticists (e.g., von Fintel, *ibid.*) concede that pragmatics can fill in necessary contextual information. However, *always* and *never* are more complicated still.

For example, they can be used in a literal or in a non-literal/exaggerated way. Consider the following authentic utterances¹:

- (1) People always want to know why the generic [grocery product] is cheaper.
- (2) Abigail is always quizzing me about my wardrobe. It’s really none of her business what I wear.
- (3) Although it proved difficult to identify abnormal liver functions uniquely related to the HEV infection, HEV RNA detection always coincided with or was followed by an increase in alanine aminotransferase.

Is *always* used in the same way in each? In which cases should one take it literally, and in which cases not? Not only that, but in which case would it be considered dishonest if it could not be taken literally? The only example above in which the *always* can be taken literally is (3), from a medical journal; in contrast, (1) is a generalization (and thus not true of literally all people) and

¹ All from the Corpus of Contemporary American English (Davies, 2008-). Search conducted November 3, 2013.

(2) almost certainly an exaggeration. Formal semanticists typically treat adverbs of frequency as correlates of quantifying determiners (*always* corresponds to *all*, *often* to *most*, and so forth), but even if pragmatics can be called upon to resolve the question of what set of cases or situations is relevant, this still does not account for non-literal uses.

In the examples above, it is the discourse-level context, the genres in which they appear, that guides our interpretation of the *always* tokens. For example, (2), being a complaint about a personal matter, would likely appear in casual spoken language, while (3) is an instance of experts in the medical field sharing facts with interested parties, readers of a peer-reviewed medical journal. Subject matter and knowledge is important as well (for example, in the context of a controlled experiment, we can refer to things that quite literally always or never happened, whereas it is hardly possible for it to be literally true that every person in the world, or even in one city, has burning questions about generic grocery products). Note, however, that these issues tie right back in to genre and pragmatic concerns such as the norms of and expectations regarding particular genres, their purposes, their participants, and so forth.

The significance of context to our understanding of *always* and *never* also encompasses the local linguistic context. Thus, this corpus-based study focuses on the complete eco-system in which a particular form is used, encompassing any co-textual factors that seem prominent in shaping the meaning, function, and distribution of the items of interest. Overall, the main elements of this system found to be relevant are genre, tense-aspect, function (as in one's communicative goals), the semantic features of the sentences' or utterances' verbs, and the semantic features of their grammatical subjects.

To check for the effects of genre on the literal versus non-literal use of *always* and *never*, I use a measurement I call the Exaggeration Quotient (ExagQ), obtained by dividing the number of tokens of *always* or *never* by the number of tokens of adverbs indicating more intermediate rates of frequency, e.g., *often* and *rarely*. I obtain this quotient for categories such as local, national, and international news, and spoken and written language. A higher ExagQ value

suggests a propensity to use *always* or *never* less literally but, given that an *always* preceded by negation (*not always*) retains its literal meaning, I also check the rates at which *always* is negated. There appears to be a strong correlation between formality and small ExagQ values, along with higher rates of negation of *always*. One functional explanation I posit for this finding is that it is not formality per se that results in more literal usage of the words, but concerns regarding accountability and the real-world consequences of producing inaccurate statements in medical journals or newspapers. This is the subject of chapter 4.

As stated earlier, one local linguistic factor I investigate is tense-aspect, first focusing on a construction, *always/never* + progressive, involving one particular tense-aspect, and then on tense-aspect in general. Consider “You’re always asking questions.” Asked to decide, absent of context, if this is a compliment or a complaint, a native speaker would probably feel it was a complaint.² This intuition is reflected in many grammars (Carter & McCarthy, 2006; Quirk, Greenbaum, Leech, & Svartvik, 1985; Sinclair, 1990), which tell us that *always* followed by the progressive often conveys a negative feeling, especially annoyance. Asked to reformulate the above utterance into a compliment, one would likely find the present simple more suitable, e.g., “You always ask great questions.” Yet, not all instances of *always/never* + progressive involve unpleasantness. Consider “The trip genes are always being ... transcribed and then translated”³ and “Medicine is always evolving.”⁴ These involve nothing negative; the speaker is merely neutrally stating facts that hold true over time. Why *always/never* + progressive is said to be associated with negative situations, the degree to which it really is, and how or why it also serves other purposes, are some of the questions addressed in chapter 5.

Following the study of *always/never* + progressive is a more general study of the tense-

² To test my intuitions, I confirmed this with four other native English speakers, between the ages of 20 and 30—one of whom clarified “It sounds like they’re saying you’re always asking *dumb* questions.”

³ Example from the Michigan Corpus of Academic Spoken English (Simpson, Briggs, Ovens, & Swales, 2002), available online at www.hti.umich.edu/m/micase/, accessed November 3, 2013.

⁴ Example from the Corpus of Contemporary American English (Davies, 2008-), accessed November 3, 2013.

aspect preferences of *always* and *never*, the subject of chapter 6. Tense-aspect is important from a cognitive-functional perspective because certain tense-aspects naturally lend themselves to certain functions. Multiple grammars (e.g., Celce-Murcia & Larsen-Freeman, 1999; Carter & McCarthy, 2006, p. 47) state that *always* and/or preverbal adverbs of frequency in general are likely to appear with the simple aspect and unlikely to appear with the progressive. I test the accuracy of this claim, and also refine it (by differentiating between the past simple and present simple). My other goal in investigating the tense-aspect preferences of *always* and *never* is to check for previously unrecognized systematic differences between the two items.

The study of tense-aspect results in the finding that, though *always* and *never* do both appear most often with the simple aspect, they have opposing tendencies regarding the past and present simple: *Always* favors the present while *never* favors the past. I argue that this is partly due to the adverbs' verbal collocates: *Always* shows up more often than *never* with verbs that describe more enduring states of affairs, or continuous states or actions, e.g., *love* versus *eat*, and that are thus more likely to call for the present tense.

This brings us to another factor examined here, the semantic features of local linguistic elements. Such elements include verbs (e.g., in chapter 6, as was just mentioned, and in chapter 7, discussed shortly) and, in the study of *always/never* + progressive, the grammatical subjects of the modified verbs, which were coded for humanness (human versus non-human) and person (first, second, or third person). One finding was that *always/never* + progressive complaints most often contained third person human subjects, and least often contained non-human subjects. I explain this in terms of efficacy, politeness, and other social concerns, and, once again, genre. For example, the neutral, descriptive statement "The trip genes are always being ... transcribed and then translated" was uttered by an instructor in a university classroom, where the main purpose of the discourse is to convey objective facts. This is but one example of how linguistic context, genre, pragmatics, and functions are all linked.

It is hoped that the methods and approach used here for *always* and *never* will prove

useful in studies of related topics as well. To that end, I conduct a preliminary investigation of maximizers (Quirk et al., 1985, p. 590), e.g., words such as *totally* and *fully*, which “express an absolute degree” (Altenberg, 1991, p. 129), yet often modify items that are non-scalar, “i.e. items that do not normally permit grading ... or already contain a notion of extreme or absolute degree” (ibid.). On a formalist, generative view, combinations such as *totally devastated* would have to be deemed redundant at best and illogical at worst. Building on previous work, I confirm the appearance of maximizers with non-scalar items, and offer cognitive-functional explanations for this phenomenon. This is the subject of chapter 7.

In sum, although *always* and *never* are not used with any highly idiosyncratic meanings, a full explanation of their meanings and the details of how they are used requires consideration of a variety of factors (namely, genre, tense-aspect, goals/functions, and the semantic features of their verbs and grammatical subject), most of which have not been considered in previous work on adverbs of frequency. This work expands our knowledge of these and similar words, as well as demonstrates the effectiveness of a corpus-based, cognitive-functional approach.

1.2 Problem Statement, Objectives, Thesis, and Limitations

The current literature and the increasing popularity and proven usefulness of corpus methods suggests that there is much to be discovered about the distribution, functions, and even intended meanings of the adverbs *always* and *never*. Thus, the main objective of this work is to establish, through the use of corpus methods and large-scale corpora, new facts about these words. More specific research questions or objectives (frequently, the verification and/or expansion of existing claims) appear in the four body chapters. The main thesis is that explanations coherent with a cognitive-functional stance are especially appropriate and effective for making sense of the findings.

This thesis is motivated by the fact that, in investigating *always* and *never*, it becomes apparent that we cannot adequately understand their behavior without taking into consideration

context, both linguistic and non-linguistic. Generative grammar and formal semantics, at least in their original incarnations, call for a strict separation of syntax and lexicon, and semantics and pragmatics. Even as that view has shifted to allow for some cross-over, the degree to which linguists working within these frameworks incorporate context is not as great as the degree to which cognitive-functionalists do. Factors such as genre, tense-aspect, syntactic constructions, the relationship between the speaker and referent of the grammatical subject of an utterance, one's goals and/or intents, etc., all have a role to play in the meaning and use of words. It is this latter approach with which I align myself, investigating the full range of factors that may help explicate *always* and *never*.

As is the case for all studies, this work analyzes a finite set of data and phenomena. The corpora used in this study, totaling approximately 455 million words, consist of modern American English and represent multiple genres of spoken and written language. The focus is restricted mainly to two adverbs of frequency, *always* and *never*, which are studied intensively. Other adverbs of frequency, such as *often* and *rarely*, are touched on but not discussed in depth. Moreover, at times, the focus is restricted to *always* and *never* in preverbal position. This is so that claims specifically about preverbal adverbs can be tested. At the same time, the study expands into a related topic, maximizers (a lexical category which includes *always*), as a further check on the validity of the methods and approach.

1.3 Terms and Concepts

Adverbs of frequency (*always*, *often*, etc.) indicate how frequently an action occurs. *Always* and *never*, the main lexical items of interest, are also antonyms, which are words that native speakers easily and intuitively recognize as opposites (Jones, 2002, p. 1), i.e., two terms that contradict one another yet at the same time are far more similar than they are different (Jones, Murphy, Paradis, & Willners, 2012, p. 3). For example, *upward* and *downward* are opposites but both are prepositions that indicate motion in a specific, vertical direction. The study also

involves maximizers. These are words, such as *absolutely* and *fully*, that “indicate degree” and, specifically, “denote the upper extreme of the scale” (Quirk et al., 1985, p. 590).

In interpreting my findings, I take a cognitive-functional approach. Cognitive linguistics is characterized by two key commitments (Evans & Green, 2006, p. 27), the Cognitive Commitment and the Generalization Commitment (Lakoff, 1990, pp. 39-41). The first is a dedication to explaining language, whenever possible, in terms of general facts about cognition, facts supported by research in other disciplines (ibid., p. 39). The second amounts to a rejection of the separation of language into fully distinct systems such as semantics, pragmatics, and syntax, as so many phenomena do not fit into only one category. Functional linguistics, compatible and overlapping with cognitive linguistics, emphasizes the “communicative and social functions of language” (Evans & Green, 2006, p. 759), and perhaps is best summarized by the notion of “language as action”. It is these core cognitive-functional ideas, all of which are expanded on in the next chapter, that I mean to refer to by “cognitive-functional”.

Other concepts, related to such approaches, are relevant in this work, but I reserve discussion of them mostly for the literature review (chapter 2) or as they become relevant (in chapters 4-7). Here, I briefly mention only a few more. First, the concept of affect and/or emotion proves relevant in several of the studies. Ochs & Schieffelin “take affect to be a broader term than emotion, to include feelings, moods, dispositions and attitudes associated with persons and/or situations” (1989, p. 7). Thus, they argue that dispositions and attitudes are an important aspect of affect (see also Besnier, 1993, p. 163). In this way, affect blurs into stance, or affective stance, defined as “a cover term for the expression of personal feelings and assessments” (Conrad & Biber, 2000, p. 57). Linguists, of course, are specifically interested in affect and affective stance as communicated via linguistic or at least paralinguistic means (see Ochs & Schieffelin, 1989, p. 7; Besnier, 1993, p. 163; Du Bois, 2007, p. 139).

In this work, I sometimes refer to constructions, and also idioms and set phrases. “Construction” is a term used by proponents of Construction Grammar, another approach that

overlaps with cognitive-functional approaches. Goldberg (1995) explains that “Phrasal patterns are considered constructions if something about their form or meaning is not strictly predictable from the properties of their component parts or from other constructions” (p. 4). Goldberg’s work is based on that of Fillmore, Kay, & O’Connor (1988), who define idioms quite similarly, i.e., as expressions or constructions not knowable based solely on knowledge of a language’s grammar and vocabulary (p. 504). Construction-based approaches stress that constructions, far from being exceptional, account for a large portion of our language, and include things as basic as the passive construction or ditransitive. Here, I use “construction” as a more general term and reserve “idiom” for chapter 6, in which Fillmore et al. (1988) is particularly relevant. I also use, in chapter 6, the term “set phrase,” to distinguish more rigid idioms/constructions from more general and flexible ones, such as the passive construction.

Another term relevant to this work, specifically in chapter 7, is semantic prosody. Similar to connotation, semantic prosody is the characterization of a word or phrase as positive (or pleasant), negative (or unpleasant), or neutral, as determined by its collocates, with its primary function being to convey “the attitude of [a] speaker or writer towards some pragmatic situation” (Louw, 2000, p. 57; see also Sinclair, 1996, p. 87). A standard example is “break out” (intransitive), said to have negative semantic prosody because it is so often used with nouns denoting unpleasant things such as disorder, hell, and epidemics (Sinclair, 1990, p. xi).

1.4 Significance

The significance of this study of two very ordinary words lies precisely in that ordinariness: Being willing and eager to embrace the messy and irregular aspects of language, cognitivists studying items such as idioms or complex metaphors might think there is little to be learned about *always* and *never*. However, as Sinclair (1999) argues, corpus studies of everyday words result in surprising and important findings. For example, a study of *the* and *of*, two of the most common words in English, unearthed evidence that each could be considered a word class unto

itself (ibid., p. 165; Sinclair, 1989). This project endeavors to help us understand two other very everyday words, *always* and *never*.

Another general contribution this study makes is that it shows the relevance of some highly specific pragmatic factors, not normally studied, to the meaning and function of the items of interest. I show, for example, that the semantic features of a clause's grammatical subject, and characteristics of the readers of local news as opposed to the readers of national news, likely influence the way *always* and *never* are used and understood. On a related note, this study builds on what others have already done with respect to genre-analysis by continuing to divide genres into sub-genres (e.g., by comparing and contrasting academic writing in three scholarly fields, including two science fields), showing that limiting ourselves to large genre categories can prevent us from noticing significant details. Guided by the practical, communicative concerns of speakers, which are not limited to single genres or contexts, we can go beyond characterizing genres in broad strokes and interpret corpus data more insightfully.

Finally, this project demonstrates that making a priori generalizations across lexical items in the same word class results in mistakes and overlooking potentially important idiosyncrasies. *Always* and *never*, and all adverbs of frequency, have their own collocates, tense-aspect preferences, functions, etc., and must be treated accordingly. Such a view is applicable to all lexical items, in fact, and especially to antonyms.

1.5 Overview

This work consists of three studies of how *always* and *never* are used, in spoken and written language in multiple corpora, and one study of an overlapping topic, maximizers. This introduction is followed by the literature review (chapter 2) and method (chapter 3), in which I describe the corpora, corpus techniques, and tools and programs used. Additional literature and more details regarding the methods used for each study appear in the relevant chapters.

In the first study (chapter 4), I calculate Exaggeration Quotients (as defined earlier) for

different genres, and the rate at which *always* and *never* are negated. In the second study (chapter 5), in which exaggeration is again relevant, I look at the functions of the *always/never* + progressive construction. In the third study (chapter 6), I analyze the overall tense-aspect preferences of *always* and *never*. The fourth and final study (chapter 7) involves the collocates and semantic prosody of four maximizers, *totally*, *utterly*, *completely*, and *fully*. This is followed by chapter 8, which summarizes and concludes the dissertation.

CHAPTER 2: LITERATURE REVIEW

In this work I analyze original findings about *always* and *never* from a cognitive-functional perspective, the key tenets of which are outlined below. I then discuss more specific literature, starting with studies of related words (antonyms and maximizers) and ending with studies of adverbs of frequency. The body chapters contain literature reviews as well, which provide further details. As I will show, adverbs in general and adverbs of frequency in particular are understudied. Nevertheless, the successful use of corpus techniques on similar words suggest aspects of *always* and *never* which can be studied, how to go about these studies, and what sorts of things we can expect to learn, and the small set of work specifically on adverbs of frequency contains specific claims that can be confirmed and/or expanded on.

2.1 Cognitive-Functional Approaches

I use the term “cognitive-functional” to refer to the core tenets of cognitive approaches (i.e., Lakoff’s (1990) Cognitive and Generalization Commitments) and/or of functional, “language as action” approaches, which emphasize the “communicative and social functions of language” (Evans & Green, 2006, p. 759). The family of approaches that adhere to these tenets includes, in addition to those explicitly labeled as such, usage-based approaches, construction grammars, corpus linguistics, the Columbia School, and conversation analysis.⁵ I expand, below, on the key cognitive-functional stances and notions that inform my analyses.

The Cognitive Commitment is the commitment to avoid language-specific cognitive mechanisms and instead aim to explain language in terms of general cognitive principles (Lakoff, 1990, p. 40; Bybee, 2010, p. 7; Langacker, 1988a, p. 4)—and only those principles “faithful to empirical discoveries about the nature of the mind/brain” (Lakoff, 1990, p. 39), such

⁵ On the overlap between cognitivist, functionalist, and usage-based approaches, see Tomasello (1998, p. xiv; 2003, pp. 1–2) and Evans & Green (2006); on how construction grammar fits in, see Evans & Green (ibid., p. 743).

as in neuroscience or psychology (ibid., p. 41). That is, because language users are human beings with “particular abilities, biases, and limitations that affect the nature and patterning of linguistic forms” (Contini-Morava, 1995, p. 17), we should only accept explanations of linguistic phenomena that are “cognitively plausible” (Langacker, 1998a, p. 29).

Related to the Cognitive Commitment is the embodiment thesis. This holds that “language cannot be investigated in isolation from human embodiment” because embodied human experiences and “human-specific cognitive structure and organization” have significant consequences for how we perceive the world and, thus, for language (Evans & Green, 2006, p. 44). An example of this is construal (Langacker, 1987, 1991, 2008, etc.), our ability to perceive the same scene/situation in different ways, which is relevant in chapter 7.

The Generalization Commitment (Lakoff, 1990) refers to the rejection of the separation of language into distinct, autonomous components (as well as the rejection of the related notion that semantics is fully compositional (see Langacker, 1988a, pp. 17-18)). Seen by cognitive-functionalists as especially problematic is the semantics-pragmatics distinction, or that between linguistic and “extra-linguistic” knowledge (ibid., p. 13). It is not that such distinctions are meaningless, but that so many phenomena do not fit neatly into one category. Instead, there is a “subtle interplay of semantic and contextual factors” (p. 5), whereby “general knowledge, knowledge of the immediate context, communicative objectives, esthetic judgments, and so on” all play a role in the process of constructing and understanding linguistic expressions (p. 14). This consideration of all types of linguistic and extra-linguistic contextual factors is visible in the studies I present here, both in the search for what affects the use of *always* and *never*, and in the search for explanations for the patterns found.

In other words, cognitivists take a broad view of meaning (encompassing pragmatics); they also take a broad view of the forms to which meanings is linked. Going beyond morphemes and words, they also study larger syntactic units, constructions. These are “phrasal patterns” for which it is true that “something about their form or meaning is not strictly predictable from the

properties of their component parts or from other constructions” (Goldberg, 1995, p. 6). Many cognitivists work with constructions and, in turn, many who work on construction grammars hold cognitive or functional views.⁶ Because they cannot be fully broken down compositionally, constructions are problematic for a generative, algorithmic system (Langacker, 1988a, p. 4), but cognitive-functionalists treat them no differently from other form-meaning units.

I consider, in chapter 5, how multiple aspects of the *always/never* + progressive construction, such as tense-aspect, the semantic content of the adverbs, and the semantic features of the grammatical subject of the clause, all interact to make the construction suited to serve various specific functions. In chapter 6, I study another construction, which I call TV *As Always*, which takes the approximate form “[name], pleasure to [have you/be here], as always”. In general, all of the work I present here assumes that our understanding of *always/never* often comes partly from various aspects of the constructions in which they appear.

Crucially, cognitive-functionalists see the form-meaning link as loose and dynamic. They argue that semantic structure is encyclopedic rather than dictionary-like (Evans & Green, 2006; pp. 160, 173). That is, “words do not represent neatly packaged bundles of meaning” but instead “serve as ‘points of access’ to vast repositories of knowledge” (ibid., p. 160, citing Langacker, 1987). Thus, while conceding that “words do have relatively well-entrenched meanings stored in long-term memory (the coded meaning)” (Evans & Green, 2006, p. 213), they argue that a more accurate understanding of meaning is that “words (and other linguistic units) serve as ‘prompts’ for the construction of meaning” online (ibid., p. 173)—a process largely guided by the aforementioned encyclopedic knowledge, and context. On this view, word meanings are (partly) “protean,” “prone to shift depending on the exact context of use” (ibid., p. 213).

Another approach that adopts cognitive-functionalist views is Columbia School grammar (as noted by Kirsner, 2002, p. 339; Huffman, 2002, p. 311; Reid, 2002, p. xi; Diver, 1995, p. 83),

⁶ For example, Goldberg (1995) credits her version of construction grammar to Brugman (1988), Lakoff (1987), Lambrecht (1994), and Langacker (1987, 1991), as well as Fillmore’s generativist Construction Grammar (as in Fillmore & Kay, 1993; Fillmore, Kay, & O’Connor, 1988).

which stems from the work of Diver (1987, etc.) and the Prague School.⁷ Notably, its proponents argue that meaning is “sparse and imprecise,” such that “meanings which are postulated for forms are viewed not as little semantic building blocks or atoms ... but rather as mere clues, hints at messages, the details of which the hearer is held to fill in by a process of intelligent inference” (Kirsner, 2002, p. 340; see also Diver, 1995, p. 74). This view is quite similar to the cognitive-functionalist “words as prompts” view.⁸

Closely related to cognitive approaches are functional approaches, which include “any approach that places particular emphasis on the communicative and social functions of language, and attempts to explain the grammatical properties of language in terms of how it is used” (Evans & Green, 2006, p. 759). An example is Halliday’s (1985) *Systemic Functional Grammar*. As Halliday & Matthiessen (2004/1985) explain, language has two metafunctions: (1) to share ideational meaning, i.e., to “construe human experience” (p. 29), and (2) to share interpersonal meaning, i.e., to “inform or question, give an order or make an effort, and express our appraisal of and attitude towards whoever we are addressing and what we are talking about” (ibid). The latter they call “language as action” (p. 30),⁹ which can be, itself, defined as “what, and how, people *do* things with language” (Nevile & Rendle-Short, 2007, p. 30.1). Not surprisingly, this perspective is also common to studies of language in a sociological context. It can be traced to Searle’s “speech acts” (1962, 1969), Austin’s “illocutionary acts” (1962), and Wittgenstein’s “language games” (1953/2009), and runs all through the field of conversation analysis (e.g., Atkinson & Heritage, 1984; Schegloff & Sacks, 1973).

Cognitive and functionalist approaches are united, then, in recognizing the importance

⁷ The Columbia School is represented by Reid (1977, 2006), Zubin (1979), Otheguy (1980), Contini-Morava (1991), Davis (1995, 2004), Huffman (2001), Kirsner (2002, 2011), and others. The Prague School was an influential group of linguists including Roman Jakobson and Nikolai Trubetzkoy. On the latter, see Tobin (1988) for an overview.

⁸ However, they take it in a unique direction: A core tenet of the Columbia School is monosemy, the “one form, one meaning” principle, as seen in the work of Jakobson and Diver.

⁹ Halliday’s ideational metafunction appears to be an instance of “language as action” too, though: Even non-emotional, non-speech acts such as the sharing of information, ideas or meanings (e.g., when narrating a story) is by definition an interactive, social affair, and a sort of action.

of the language user: The former focus on our cognitive capacities and embodied experiences, and the latter on our needs and desires as social creatures interacting with others. In this work (especially in chapter 5), I rely on concepts relevant to sociology and evolutionary psychology, e.g., politeness theory, the negativity bias, and the social value of gossip.

Another defining feature of cognitive-functionalism is an insistence on empirical evidence derived from authentic data, analyzed using qualitative and/or quantitative methods. The connection between cognitive-functionalism and empirical evidence runs in the opposite direction as well: Functional interpretations are typically considered an “essential” part of a corpus study (Biber, Conrad, & Reppen, 1998, p. 5). Thus, in corpus work one finds the defining features of cognitive-functional work, e.g., a focus on cognition and communication, cross-over between pragmatics and semantics, and, of course, empirical data. More information on corpus methods and the specific methods and corpora used in this work appears in chapter 3.

Usage-based models (e.g., Bybee, 2010), as well, align with the cognitive-functional stance adopted in this work. In this case, *use* refers not to what we use language to do (as it does in functional accounts) but to the effect on our minds and on the language itself of our repeated exposure to specific instances of form-meaning pairs. That is, “the structural phenomenon we observe in the grammar of natural languages can be derived from domain-general cognitive processes as they operate in multiple instances of language” and “the repetitive use of these processes ... has an impact on the cognitive representation of language and thus on language as it is manifested overtly” (Bybee, 2010, p. 1). Grammar is thus “emergent,” which explains the gradient nature of so many aspects of language (ibid.). Bybee and others (e.g., Bybee 2001, 2002, 2003; Bybee, Perkins, & Pagliuca, 1994; Bybee & Thompson, 1997) focus on domain-general cognitive concepts such as exemplars and frequency effects, and on language change and grammaticalization.¹⁰ In the studies presented here, most phenomena (such as how literally we should take *always*) are seen as gradient, and I rely heavily on frequencies, ratios, and other

¹⁰ On grammaticalization see also Heine & Kuteva (2002), Hopper (1991), Givón (1973, 1979), and Traugott (1989).

calculations as indicators of how language is used and (regarding one of the idioms discussed in chapter 6) of change in progress.

In sum, cognitive-functional approaches—as heterogeneous as they are—share certain core tenets: The Cognitive Commitment (to cognitively plausible explanations involving domain-general processes and supported by work in other disciplines), the Generalization Commitment (to viewing language as involving interplay between semantics and pragmatics), the understanding of language as a form of action directed by social goals, and the reliance on authentic qualitative and quantitative data, including large-scale corpus studies. All of these tenets serve as useful guiding principles for the work I present here.

As it happens, “the domain of linguistics that has arguably been studied most from a corpus-linguistics perspective is lexical ... semantics” (Gries & Otani, 2010, p. 121), resulting in studies of a number of word categories (antonyms, maximizers, and adverbs of frequency) that include *always*, or both *always* and *never*. In addition, adverbs of frequency are characterized in a number of corpus-based and functional grammars. I summarize these works below.

2.2 Antonyms

As this is a corpus study of antonyms, I will touch upon some major studies using this methodology. Notable corpus studies of antonyms in English include a large collection of work by Jones, Murphy, Paradis, and Willners (e.g., Jones, 2002, 2006, 2007; Murphy, 2003, 2006; Willners, 2001; Jones & Murphy, 2005; Paradis & Willners, 2006; Jones, Paradis, Murphy, & Willners, 2007; Paradis, Willners, & Jones, 2009; Jones, Murphy, Paradis, & Willners, 2012), and others, e.g., Gries & Otani (2010) on antonyms of size. These studies well-demonstrate the value of corpus methods in this area. However, adverbs are not the focus of these works, and *always* and *never*, in particular, appear in few or no antonym studies.

The alternative to corpus study of antonyms is a formal semantics approach that relies on intuition to define and categorize antonyms (as in Leech, 1974; Kempson, 1977; Cruse, 1986).

Jones (2002, pp. 23–24) summarizes the problems with this approach. On the one hand, the intuitive criteria used by formal semanticists have led to a proliferation of categories and terms and, on the other, the one major distinction they agree on is not supported by evidence. This is the distinction between gradable and non-gradable antonyms. Palmer (1976, p. 87) and others show that some pairs appear to be simultaneously gradable and non-gradable. (A similar phenomenon, the topic of chapter 7, occurs with the collocates of maximizers).

Corpus studies on antonyms began with Justeson & Katz (1991), using the approximately one million word Brown corpus (Francis & Kučera, 1964). This was followed by Fellbaum (1995) (also using the Brown corpus), Mettinger (1994), Muehleisen (1997), and Murphy (1994). The key finding of Justeson & Katz was that words “tend to occur in the same sentence as their antonyms far more frequently than expected” (p. 142), and also “in parallel and often essentially identical phrases” (ibid.). This shows that antonyms are a context-based phenomenon, well-suited for study at the discourse and corpus level. Mettinger (1994) contrasts with Justeson & Katz in adopting a structural, “autonomy of syntax” approach, but nevertheless argues that empirical data is crucial, i.e., that textual evidence, not theoretical arguments based on intuition, helps us categorize antonyms in useful ways.

These important early studies had their limitations, of course. The antonyms Justeson & Katz (1991) studied were the 40 pairs identified by Deese (1964), a list Deese obtained using word-elicitation techniques now considered suspect (see Jones, 2002, pp. 27–28), and honed using overly loose criteria (ibid.) Mettinger (1994) is also problematic: His antonyms are derived from 43 British novels, mostly by Agatha Christie—an overly specific corpus—and from *Roget’s Thesaurus* (Roget, 1972), which he treats as a corpus. *Roget’s Thesaurus* was not created on the basis of empirical evidence about antonyms, nor was it kept up to date. The first version of it was published in 1852 and, as Mettinger himself tells us (p. 94), many of words in the 1972 version were, by 1994, “hardly used in contemporary English”. Finally, none of the antonyms studied by Justeson & Katz (Deese’s list contained only adjectives) or Mettinger were adverbs.

The more recent work cited above, e.g., that of Jones, Murphy, Paradis and Willners, is based on more carefully selected antonyms (e.g., see Jones, 2002, pp. 29-39). However, even when the antonyms studied are relevant and canonical, and the corpora large and appropriate, adverbs are still largely overlooked. For example, the most comprehensive recent work, Jones et al. (2012), does not include adverbs. The authors focus on adjectival antonyms, “for the simple reason that common adjectives have antonyms more often than common nouns and verbs do” (p. 4). Adjectival antonyms may indeed be the most readily available, but it is revealing that adverbs are not just excluded from the study but are, in fact, not even included in the authors’ list of what is not included. It is also worth noting that, though adjectival antonyms are more common overall, than adverbial antonyms, many of the individual antonyms studied by Jones et al. are far less common than *always* or *never*.¹¹ Adverbs are mentioned briefly in Jones (2002), but the conclusion is that “writers use antonyms to much the same ends, regardless of whether those antonyms are adjectives, nouns, adverbs or verbs” (p. 148; see also p. 153). Thus, despite themselves being a canonical antonym pair, *always* and *never* are, like most adverbs, generally excluded from studies of antonyms.

2.3 Intensifiers and Maximizers

Another category of words which includes adverbs of frequency and which has attracted a lot of attention is intensifiers and, specifically, maximizers. Intensifiers are words “broadly concerned with the semantic category of degree,” i.e., they scale quantities upward or downward (Quirk et al., 1985, p. 589; see also Bolinger, 1972). Being one of the most active and flexible classes of words, intensifiers are “an antidote to the overconfident description of language as a system” (Bolinger, 1972, pp. 18-19), and there is much to be learned about them. Work on this class of words traces back as far as Kirchner (1955), Greenbaum (1970, 1974), and Bäcklund (1973), and,

¹¹ A search of the Corpus of Contemporary American English (Davies, 2008-) resulted in 207,151 *always* and 101,117 *never* tokens, but only 58,503 *hot* and 58,540 *cold* tokens. This makes *always* 3.5 times more common and *never* 1.7 times more common than either item in the canonical pair *hot* and *cold*. Searches conducted September 1, 2013.

overall, has shown that intensifiers “are subject to a number of syntactic, semantic, lexical and stylistic restrictions affecting their use in various ways” (Altenberg, 1991, p. 142), and thus highly amenable to corpus study.

Maximizers are adjectival or adverbial intensifiers that “denote the upper extreme of the scale” (Quirk et al., 1985, p. 590). The use of corpus methods to study maximizers has resulted in surprising findings, sometimes contradictory to intuitions, including subtle distinctions between near-synonyms, and seemingly paradoxical lexical pairings. Regarding near-synonyms, Partington (1998) shows that *sheer*, *complete*, *pure*, and *absolute* vary in their collocates and uses, and Partington (2004) discusses distinctions between *absolutely*, *perfectly*, *entirely*, *completely*, *thoroughly*, *totally*, and *utterly*, e.g., *absolutely* is the only one of the set that frequently appears in hyperboles. Likewise, Tao (2007) demonstrates that *absolutely* has unique discourse functions. Another maximizer that stands out is *utterly*, which Greenbaum (1970), Partington (1993, 2004), Louw (1993), and W. Anderson (2006) note has especially negative semantic prosody (defined in §1.3). Regarding paradoxical lexical pairings, it has been found that maximizers often modify items that are non-scalar, i.e., “do not normally permit grading ... or already contain a notion of extreme or absolute degree” (Altenberg, 1991, p. 129). From a formal, generativist standpoint, pairings such as *totally devastated* are redundant or illogical. This is the subject of chapter 7.

Overall, the body of work on maximizers and intensifiers shows that there is much to be learned about these lexical items, and that corpus techniques are effective for doing so—especially when applied to modern, large-scale corpora (Altenberg, 1991, p. 142). Unfortunately, and as was the case in antonym studies, among intensifiers adverbs have been largely neglected. Although many of the maximizers that have been studied can act as adverbs (e.g., *completely* in “She completely defeated him”), they are far more likely to be adjectives.¹² And even Bolinger, so

¹² A search of COCA (Davies, 2008-) resulted in 13,348 and 2,871 tokens, respectively, of *completely* and *utterly* immediately followed by adjectives (e.g., *completely non-intuitive*, *utterly unscriptural*), and 12,827 and 1,047 tokens

keen on stressing the significance of intensifiers, in his monograph on the topic treats “nouns and verbs mainly” and, when he does discuss adverbial intensifiers, refers to them as “degree adjectives”—adding that “there is nothing of interest here that can be said about *He gave a beautiful lecture* that does not apply equally to *He lectured beautifully*” (1972, p. 15.) This exemplifies the recurring pattern whereby, in studies of both antonyms and intensifiers, adverbs are systematically excluded.

2.4 Adverbs of Frequency

Quantifiers (e.g., *all*, *some*, and *none*) are a staple of predicate calculus, the language of symbolic logic used by formal semanticists, and well-studied. Adverbs of frequency (*always*, *often*, *never*, etc.), which are a type of quantifier, are studied too, but to a lesser degree. For example, the entry in the *Encyclopedia of Language and Linguistics* (K. Brown, 2006) entitled “Quantifiers: Semantics” states that “the best understood type of quantification” is that involving quantifying determiners such as *all* in *all* poets (Keenan, 2006, p. 302); nowhere in the six-page entry are adverbs of frequency even mentioned. This echoes the same imbalances described above in studies of antonyms and maximizers. In formal semantics the likely reason for the imbalance is that (1) adverbial quantifiers are generally explained in terms of adjectival quantifiers, and thus appear to be simply a type of them, and (2) they are especially difficult to explain.

The adverbial quantifier *always* is said to be related to the adjectival quantifiers *all*, *each*, or *every*; *often* is said to be related to *most*; *never* is said to be related to *none* or *no*; etc. For example, “Liam always has breakfast” is said to mean Liam has breakfast “every morning” (Kearns, 2000, p. 136). Treatments involving paraphrasing of this sort are found in both formalist accounts and (though without the predicate calculus) in functional accounts as well (e.g., Celce-Murcia & Larsen-Freeman, 1999, p. 508). Though intuitively satisfactory, such

of the same words followed by items tagged as nouns, but the “nouns” largely turned out to be adjectives (specifically, adjectival participles) (e.g., *utterly dumbfounded*, *completely baffled*). Search conducted March 25, 2015.

treatments do not fully explain the meanings of adverbial quantifiers.

Implied but not directly addressed by paraphrases like the one above is what, exactly, is quantified over. This is typically said to be either time (referring to all time, i.e., *always* would mean “constantly/incessantly”) or “cases” (i.e., situations or opportunities in which it would be relevant for the given action to occur). Works which argue for quantification over time include the dissertations of Stump (1981, published as Stump, 1985) and Rooth (1985). Though “all the time” may be our first response if asked to define *always*, instances in which it means this are rare, though they do exist, e.g., “I will always love you” presumably refers, at least, to all present and future time. (The use of *always* with durative verbs such as *love* is relevant to §6.5.1.) This is where “cases” are helpful. For example, “Jim always orders dessert” means he does so every time he dines out, because this is the relevant opportunity for dessert-ordering.¹³ See also Lewis (1975/2002, p. 179), Von Stechow (1995, p. 9), and Heim (1990). It is usually quite easy to identify the “case” relevant to a given utterance, while arguments for quantification over time (as the only or default possibility) are often difficult to apply to specific examples.¹⁴

In sum, *always* has at least two different (but related meanings), one something like “all the time” and the other something like “in all relevant cases.” (*Never*, its antonym, may be more straightforward, since occurring in no specific cases is the same as occurring at no time.) These meanings appear in dictionaries; e.g., the Oxford English Dictionary Online (2015) gives the definitions “at all times” (def. 1), “for all time, forever; ... continually, perpetually, without any interruption” (def. 3) alongside (in fact, mixed with, suggesting our lack of conscious awareness that the two are different) “on all occasions, ... at every occasion, every time” (Def. 1).

¹³ Of course, other situations can be imagined too, e.g., perhaps we know he is a wedding planner who frequently deals with caterers. In this case, our knowledge might lead us to see “for every wedding” as the relevant case.

¹⁴ For an extended analysis of quantification over time, see Stump (1985). The arguments are a bit hard to follow, though. First, the examples (as Stump is aware; see p. 98) frequently include phrases that spell out what others would say are the “cases” quantified over. Second, Stump works with “time intervals,” but these do not feel applicable to all situations. Consider “John always went for a walk after he ate supper” (p. 183): If the time interval begins right after supper and has no preordained endpoint (it ends whenever the walk ends), it is unclear (to me) what is gained by focusing on time intervals rather than the situations/cases which serve as or precede them.

Sometimes the information necessary for a speaker to understand which “case” is intended is provided, explicitly, in the words of the utterance, e.g., “When she figured out her taxes, Jane often used a calculator” Stump (1985, p. 180), or can be retrieved from earlier in the discourse. Other times, though, the “case” is not explicitly stated anywhere. “Liam always has breakfast” (Kearns, 2000, p. 136) is understood to be about mornings purely due to our world knowledge about breakfast—an example of the encyclopedic and “words as prompts” view of meaning discussed earlier.¹⁵ This calls for cross-over between pragmatics and semantics not traditionally permissible in formal semantics. This rigid view has partly given way, actually (Portner & Partee, 2002, p. 5), and for decades now formal semantics has included work combining semantics and pragmatics (e.g., Keenan, 1971; Heim, 1983), as well as semantics and syntax (e.g., Heim, 1982; Rooth, 1985). As Von Stechow puts it, “That the domain of quantifiers is often subject to pragmatic influences is hardly shocking news” (1995, p. 11).

The situation regarding adverbs of frequency and what they quantify over is strong evidence that treating semantics and pragmatics as completely separate domains is not tenable, as the two things tend to blend together. This is true of another aspect of adverbial quantifiers as well, which is the topic of chapter 4: The number of cases in a given set that a quantifier¹⁶ is meant to refer to can be largely subjective. Studies of the subjectivity of the meanings of adverbs of intermediate frequency, such as *often*, *rarely*, and *frequently*, appear as early as the 1970s, and include Pepper & Prytulak (1974), Nakao & Axelrod (1983), Wallsten, Fillenbaum, & Cox (1986), and McGlone & Reed (1998). These studies, discussed in more detail in the next chapter, focus on the effect of context—specifically, the typical or “base” rate at which a given event occurs (derived from general/world knowledge)—on how we use and interpret these items. Unfortunately, none of these studies involve *always* or *never*, with the implication being that,

¹⁵ World knowledge must help us resolve ambiguous sentences—as many or all of the example sentences technically are. For example, “Tai always eats with chopsticks” presumably means “... whenever he eats” but, arguably, could also mean that he truly uses them incessantly (Von Stechow, 1995, p. 9).

¹⁶ I say *quantifiers* because this applies to quantifying determiners as well: What number of items or percentage of items in a set that constitutes *a few*, or *most*, etc., is variable.

because they refer to the extremes of frequency, their meanings are stable.

They are not, though. A recent corpus study on exaggeration (Claridge, 2011) shows that even *always* and *never* are “prone to taking on hyperbolic interpretation” (p. 51). Particularly in complaints, they can take on extremely exaggerated, non-literal meanings. Just like *often* or *sometimes*, then, they can be used and interpreted in a variety of ways—and this is dependent on not just typical/base rates but many other factors as well, including emotion and communicative goals. This is discussed in more detail in chapter 5.

In sum, although research has been conducted on the subjectivity of the meanings of adverbs referring to intermediate rates of frequency, little or no work has been done on those at the ends of the scale, which are also subject to the effects of context. (Claridge (2011), mentioned above, does touch on *always* and *never*, but is a study on exaggeration in general and thus leaves room for more focused study of these words.) That being said, I now turn to studies that do focus explicitly on *always* and/or *never*. This final set of literature consists of functionally-based characterizations of adverbs of frequency found in grammars of English and which mainly address tense-aspect and function. The claims made in these works— though not yet empirically proven—are a useful starting point for further studies.

Several grammars of English (Quirk et al., 1985; Celce-Murcia & Larsen-Freeman, 1999; Carter & McCarthy, 2006) comment on the tense-aspect preferences of *always* or preverbal adverbs of frequency (PAFs) in general, saying these are unlikely to appear with the progressive and most likely to appear with the simple and present perfect. The posited reason for these associations is an overlap in function: Both PAFs and the tenses they commonly appear with are conducive to making statements about habitual actions and/or general states of affairs (see Celce-Murcia & Larsen-Freeman, 1999, p. 509; Praninskas, 1975; Carter & McCarthy, 2006, p. 47). In chapter 6, I test these claims about PAFs and tense-aspect.

The present and past progressive tense-aspect, in particular, merits special comment in the context of PAFs. The three grammars mentioned above, as well as Sinclair (1990), all note

that the *always*/PAF + progressive construction often refers to negative, unpleasant situations or experiences, i.e., “regular events or states which are problematic or undesired (Carter & McCarthy, 2006, p. 47) or “a subjective feeling of disapproval” (Quirk et al., 1985, p. 199) and/or annoyance (Sinclair, 1990, p. 249). Celce-Murcia & Larsen-Freeman (1999) take a broader view, referring simply to “emotional overtones” (p. 510), including overtones of approval, not just disapproval (p. 117). Empirically testing these claims about the function of PAF + progressive constructions—as well as testing the implication that different PAFs behave roughly similarly—is the objective of the study presented in chapter 5.

2.5 Conclusion

Given that adverbs are less common than other parts of speech, and adverbs of frequency even less common, adverbs of frequency have not been studied as thoroughly as other words. Yet there is much to be learned about them, as suggested by work on similar words and by the smaller set of work on adverbs of frequency. This study verifies and expands on these previous studies with the aim of relating usage patterns to meanings and functions, as well as identifying various factors that affect the way these words are used. Corpus analyses have shown that everyday words such as pronouns and certain conjunctions, for example, are used in different ways depending on the genre in which they appear (e.g., Biber, 1988); this is likely to be true, also, for adverbs of frequency, yet no analysis of this has been conducted.

This work is comprised of multiple such analyses of *always* and *never*. I show that, for example, the semantic features of the grammatical subject (such as human versus non-human), genre (in a broad sense, e.g., academic writing versus newspaper articles, and in a very narrow sense, e.g., national versus international news articles), and so forth all interact with the meanings and discourse functions (e.g., complaints versus laments) of these words. The overall goal is to contribute to an accurate and thorough characterization of these words, as well as to show what types of factors can be relevant in lexical studies.

CHAPTER 3: METHOD

The purpose of this study is to determine and explain certain aspects of the behavior and meanings of *always* and *never*. To do this, it is necessary to select an appropriate data-set and to obtain results from that data which enable one to characterize the distribution of *always* and *never*, as well as reasonably confidently determine their functions and meanings in various contexts. These results are then discussed in light of cognitive-functional considerations.

3.1 Research Design

The corpus method consists of much more than just large-scale quantitative analyses of corpora; “it is essential,” too, “to include qualitative, functional interpretations of quantitative patterns” (Biber, Conrad & Reppen, 1998, p. 5; see also pp. 4, 9). The qualitative-quantitative combination has been proven powerful and effective and, as a result, corpus techniques have been adopted by researchers, creators of dictionaries and grammars, and language teachers, who use corpus techniques to inform their teaching or even have their students work with corpora.

Authentic data and empirical evidence is far more reliable than unverified introspection and intuition, which are problematic for a number of reasons. To begin with, the invented sentences that appear in linguistic articles typically differ greatly from what is actually found in corpora (Sampson, 1992, p. 48). More interestingly, though, even sets of examples consisting of authentic utterances can be misleading. For example, I show in chapter 5 that the examples showcased in grammars of English of the *always/never* + progressive construction, though quite representative in terms of the nature of the grammatical subject of the clause, are skewed in terms of function. The use of corpora helps control for the effect of biased samples because a large corpus will contain instances of the lexical item or construction of interest being used in a variety of ways and contexts.

In terms of quantitative measurements, a recurring problem is that people notice the

unusual more than the usual (Biber, Conrad, & Reppen, 1998, p. 3), and rarely even think about the mundane yet most pervasive aspects of language, such as function words (cf. Pennebaker, 2011, pp. 24-25). Specifically, speakers have very inaccurate intuitions or none at all about comprehensive patterns of usage and especially frequency (McEnery & Wilson, 1996, p. 15; Reppen, 2010, p. 31; Sinclair, 1991, p. 39; Stubbs, 1996, pp. 28–29). The strengths of large-scale quantitative analyses are clear here: Corpus studies not only reveal specific usage patterns and frequencies, but they do this for phenomena we were not even aware of (e.g., inconspicuous yet common syntactic patterns). And the larger the corpora, the clearer the patterns.

Like any method, of course, corpus methods can be misapplied, but there are established ways of addressing the most common concerns. One concern is that a corpus may not be representative of what it is purported to represent. Another—relevant to all quantitative work—is that the data will be decontextualized and thus oversimplified. The former is a matter of corpus design and can be addressed as such, while the latter is the impetus for the insistence that quantitative analyses be complemented with qualitative ones.

A corpus study is only as good as its corpus, which must be designed in a principled way and for a principled reason (Biber, Conrad & Reppen, 1998, p. 248), so that it can represent the language or parts of a language that it is seeking to (ibid., p. 246). For example, to create a corpus based on which one could make general comments about a particular language, one would need to include a diverse array of texts that represent the various dialects, genres, and registers of that language (ibid., p. 248). In contrast, a corpus designed to help us to study the creative writing of seventh graders attending a particular school in 1992 would only need to contain creative writing samples from that very limited set of language users.

Two other important corpus features are size and balance. Whether one aims to represent an entire language or a specific aspect of it, all other things being equal, a larger corpus will be more representative. Equally important, though, is that it be derived from more unique sources (and more unique sources per category/genre, if relevant). In addition, a good

corpus is “balanced” (Biber, Conrad, & Reppen, 1998, pp. 248-249), which means the portions of it dedicated to each sub-genre, dialect, source, etc., are roughly equal in size.

This raises a different concern, which is that, in such a heterogeneous collection of texts, any important differences will be averaged out. However, this can be avoided by looking at corpora both as a whole and in sections, in accordance with the researchers’ goals and focus. Doing this, one might discover differences between genres, or that the pattern in a given sub-category of a corpus matches that found in the corpus overall, indicating that the sub-category (with respect to the feature being analyzed) is also typical of the language in general.

Concerns about oversimplification and decontextualization are handled, as mentioned earlier, by complementing quantitative work with qualitative work. Qualitative analyses are typically flexible and multifaceted, allowing us to consider any number of factors which may be relevant. This can greatly enrich our analysis and, moreover, can feed back into what we next study quantitatively. Even qualitative analyses of the simplest nature, such as the scanning of KWIC views (see §3.3), often reveal oversights which can then be corrected. For example, if *always* is especially frequent in a particular genre, one might assume that, in this genre, people frequently speak in a hyperbolic, exaggerated manner. If, however, upon checking the KWIC views, we find (as is true for medical articles, discussed in chapter 6) that many of those *always* tokens are preceded by *not*, we will reach quite different conclusions.

The weaknesses of qualitative analyses are that they are inherently more subjective and more time-consuming. The main concern in this particular study is reliability with respect to coding schema, which can be hard to operationalize. For example, while the difference between a human and non-human noun is clear, the difference between a description and a complaint might be less obvious. In addition, because qualitative analyses are time-consuming and potentially mentally taxing, one must strike a balance between output and controlling carefully for errors. In this study, such concerns are mitigated by the use of clear and consistent coding criteria, and multiple rounds of quality control checks.

This study involves various specific qualitative and quantitative methods designed to reveal different aspects of how *always* and *never* are used in natural language; the details of these procedures appear in the method sections of the respective body chapters. Here, I discuss broader aspects of the method: the corpora used, the research instruments and types of data, and the general ways the data was analyzed.

3.2 Corpora

This study involves four corpora: the Corpus of Contemporary American English (COCA) (Davies, 2008-), the Switchboard Corpus (Godfrey, Holliman, & McDaniel, 1992), the Michigan Corpus of Academic Spoken English (MICASE) (Simpson, Briggs, Ovens, & Swales, 2002), and the Santa Barbara Corpus of Spoken English (SBC), parts 1-4 (Du Bois, Chafe, Meyer, Thompson, Englebretson, & Martey, 2000-2005). COCA, MICASE, and SBC are available for free at coca.byu.edu/coca, micase.elicorpora.info, and linguistics.ucsb.edu/research/santa-barbara-corpus, respectively, and Switchboard is available for purchase in text and audio format at catalog.ldc.upenn.edu/LDC2001S15. The constant across the corpora is that they all consist of spoken American English (COCA also contains a written component), and only very recent language (spoken, broadcast, or published between the late 1980s and 2012).

The inclusion of a large amount of spoken language is crucial. I am seeking not just facts about *always* and *never* but, ultimately, general truths about how our minds handle language. Spoken data is most relevant to this endeavor because it shows not what we can be trained to do with language but what comes naturally. That is, the ability to use spoken language is possessed by all able-bodied and neuro-typical human beings, whereas literacy is an acquired skill, and not acquired by all individuals or in all cultures. Not surprisingly, then, there are striking differences between spoken and written language (e.g., Biber, Johansson, Conrad, & Finegan, 1999, pp. 15-17; Carter & McCarthy, 2006, pp. 9-11, 164-175; Linell, 2005, pp. 17-33). Broadly speaking, “written language is structurally elaborated, complex, formal, and abstract, while spoken

language is concrete, context-dependent, and structurally simple” (Biber, 1988, p. 5). This means it is unacceptable to base claims about spoken language or about “language” in a general sense on analyses of written language.

The decision to use only very recent American English was made for important practical reasons. First, this is the native language of the investigator, which means I will be less likely to make judgment errors or overlook nuances of meaning. Second, I limit the study to a single time period and single dialect because, being a primary investigation of *always* and *never*, this cannot also be a study of diachronic variation or synchronic variation. I am most interested in the role of function and context, and limiting the data to American English from one brief time period increases the likelihood that any variation observed will be related to these things, whereas the inclusion of data from multiple time periods and dialects would introduce too many confounding variables that could not be properly treated.

At more than 450 million words, COCA is the largest of the four corpora, by far, and also one of the largest fully-functional corpora freely available.¹⁷ Widely used by linguists and other researchers, COCA is not only large, but balanced, drawing equally from spoken language, fiction books, magazines, newspapers, and academic journals. Moreover, each category is represented by a variety of authors, television shows, etc. For example, the “Academic Journal” section contains articles from over 100 different journals. The content is also categorized according to year, dating from 1990 to 2012, with 20 million words devoted to each year. All of this information (source, year, etc.) is available to users, along with the larger linguistic context (i.e., a block of approximately 180-190 words surrounding the search item).

The spoken portion of COCA consists of 95 million words of “unscripted conversation from more than 150 different television and radio programs.” While this language is not entirely natural, it is a suitable approximation of natural language (as explained on COCA’s website).

¹⁷ The information and quotations given here about COCA were obtained from its website, under “GENERAL: Introduction (general),” “TEXTS/TYPES: Overview” or “TEXTS/TYPES: Spoken Transcripts”. Accessed April 8, 2015.

This claim is supported by the findings presented here (e.g., in chapter 4), which characterize COCA's spoken language as lying somewhere between written language and casual spoken language. Thus, and especially when appropriately supplemented and contrasted with other spoken data, this data provides useful insights into natural, casual spoken language.

With the explosion in popularity of corpus work, a number of enormous new corpora have appeared, many of which are associated with COCA. These include the Corpus of Historical American English (COHA, 400 million words dating from 1810 to 2009) (Davies, 2010-), which makes use of the same interface as COCA, the Google Books corpus (the American English portion consisting of 155 billion words from the 1500s to 2000s) (Davies, 2011-), and the Corpus of Global Web-based English (GloWbE, 1.9 billion words, 387 million of which are categorized as American English) (Davies, 2013). Given that these corpora of English approach or exceed COCA in size, I must explain why they were not suitable for the present study.

First, COHA and Google Books cover too large of a time period (two centuries and half a century, respectively). As stated earlier, this study focuses on very modern English. Because languages change so rapidly, the COHA and Google Books data would introduce too many confounding variables. Second, GloWbE's American English portion likely contains other types of English, actually; dialects are deduced from the country code in the website's URL (e.g., ".ca" for Canada), but a person's URL does not indicate her native language or native speaker status—or even her actual location, for that matter. Third, the interfaces offered for Google Books and GloWbE are not as helpful as COCA's. The only contextual information offered by GloWbE is the country, not the source text (though one could click through to see), date, or other information, and Google Books cannot be searched in the same manner that COCA can; clicking on a word leads not to a convenient list of instances of the keyword in context, but to a regular search results page. This makes the data much harder to process.

Returning to the corpora selected for this study: The next largest is Switchboard, consisting of approximately three million words (more than 240 hours) of recorded telephone

conversations participated in “by over 500 speakers of both sexes from every major dialect of American English.”¹⁸ The calls were collected by Texas Instruments, and the creation of the corpus was funded by the Defense Advanced Research Projects Agency, an agency of the United States Department of Defense. It was intended to be used for a variety of purposes, including facilitating studies on topics such as “the phonetic characteristics of spontaneous speech” and the effects of different telephones on voice recognition. The corpus consists of the original audio files along with text (.txt) files, but I had access only to the text.

Switchboard was compiled very carefully. Callers agreed to have their conversations recorded, and were given pre-set topics to discuss that they were also free to stray from. Recording was done automatically, with no human intervention, to limit the effect on the callers of being knowingly recorded. In addition, each conversation’s “naturalness”, along with other features such as how easy the audio was to understand, was ranked by transcribers (mostly court reporters, following explicit instructions). The transcripts were double-checked for errors by an automated script and by humans in two rounds of quality control inspections.

The third largest corpus in the study is MICASE, a corpus designed to be used for general linguistic research and for the development of teaching and testing materials for teachers and learners of English as a Second Language, especially English for Academic Purposes. The corpus contains nearly 1.8 million transcribed words of spoken academic English used by students, faculty, staff and visitors of both genders and various ages at the University of Michigan, Ann Arbor. The content represents fifteen types of speech events (e.g., lectures, advising consultations, and student study group meetings). As such, it goes beyond “scholarly discussion” and contains “such speech acts as jokes, confessions, and personal anecdotes”.¹⁹

A strong point of MICASE is how carefully its data is labelled. Transcripts are identified

¹⁸ The information and quotations given here about Switchboard were obtained from www ldc.upenn.edu/Catalog/readme_files/switchboard.readme.html, accessed July 31, 2013.

¹⁹ The information and quotations given here about MICASE were obtained from its website, accessed May 5, 2013.

according to Speech Event Type, Academic Division, Academic Discipline, Participant Level, and Interactivity Rating, and speakers according to Gender, Age, Academic Position/Role, Native Speaker Status, and First Language. (In the studies conducted here, searches were restricted to utterances of native speakers of American English, which comprised about 1.6 million words). Another strong point of MICASE is that a portion of the recordings are available online. These could be used in studies of prosody and intonation or, for the purposes of this study, to provide additional contextual information that could clarify points of confusion.

Finally, the SBC, the smallest of the four corpora, consists of roughly 249,000 words of natural speech from face-to-face conversations (and some phone calls) during typical daily life activities such as story-telling, town hall meetings, and food preparation.²⁰ The recordings come from all over America and from speakers of various genders, ages, and backgrounds, and thus constitute a representative sample of American English. Designed as a general resource for linguistic research, the SBC is also “the main source of data for the spontaneous spoken portions of the American component of the International Corpus of English.” In addition to the text of the transcripts, the SBC website also offers all of the recordings as audio files, along with an explanation of the general context of each conversation.

I said earlier that the differences between the corpora were as crucial to my study as their similarities. First, they supplement one another. Concerns about COCA’s roughly 95 million words of spoken data not being entirely natural, for example, are alleviated by the inclusion of Switchboard and the SBC, which are much smaller than COCA yet valuable for containing casual spoken language used in natural or fairly natural settings. Second, the four corpora serve as a check on one another. For example, what is characteristic of academic spoken English is not taken to be characteristic of all speech unless observed in the other three corpora also. Third, their differences allow me to investigate how speakers use *always* and *never* across a range of situations, and seek correlations between genres and/or context.

²⁰ The information and quotations given here about SBC were obtained from its website, accessed July 30, 2013.

In sum, the four selected corpora represent extremely modern spoken American English, with the largest, COCA, containing written language as well. As always, even larger and more natural corpora would be desirable. For obvious logistical reasons, corpora of spoken language tend to be smaller than corpora of written language. However, the chosen corpora are carefully crafted—which is of the utmost importance—and the similarities in terms of time period and dialect combined with their differences in language type (i.e., spoken and written, and a variety of genres) make them an effective set of corpora for investigating *always* and *never* across contexts and genres of modern American English, including spoken English.

3.3 Research Instruments and Types of Data

This study attempts to get at the facts regarding *always* and *never* in as many ways as possible, and therefore makes use of both typical corpus measurements, like the Mutual Information (MI) scores of collocates, and also entirely new ones, like Exaggeration Quotients (in chapter 6). Speaking more generally, I use quantitative data obtained using corpus tools and also some quantitative data which could only be computed after manual coding that relies on qualitative methods. Sometimes additional analyses are conducted (and presented in the chapters' discussion sections) in order to test the validity of my interpretations.

The main tools used were, for Switchboard and SBC, the Antconc concordancer (Anthony, 2011) and, for COCA and MICASE, the tools available on their respective websites, corpus.byu.edu/coca/ and quod.lib.umich.edu/m/micase/. COCA's tools work as a sophisticated concordancer, while MICASE's function more like an internet search engine (but with options to filter the results according to type of speaker or interaction). Statistical significance tests were calculated using online tools provided at vassarstats.net/newcs.html by Vassar College. Some coding decisions (e.g., regarding the semantic characteristics of grammatical subjects, such as human or non-human) were done manually, according to pre-set criteria.

For certain procedures, corpus data (especially from large corpora) is manageable only

when its words have been tagged according to the parts of speech. This is known as parts of speech (POS) tagging, though it typically goes beyond just parts of speech. For example, items are labeled not just as verbs but by tense-aspect, and nouns may be labeled as pronouns, regular nouns, proper nouns, etc. In this study, the fact that the largest corpus, COCA, is pre-tagged enabled me to easily obtain instances of *always* + progressive appearing in it.

The tagger used for COCA is CLAWS (the Constituent Likelihood Automatic Word-tagging System), the latest version of which (CLAWS4) (Leech, Garside, & Bryant, 1994) was also used to tag approximately 100 million words of the BNC (British National Corpus) (ibid.). Its output has been determined to be 96-97% accurate (BNC2 POS-tagging Manual, available at ucrel.lancs.ac.uk/bnc2/bnc2error.htm). Moreover, all the hits I obtained using it underwent manual quality control checks (e.g., to ensure that a participle was not accidentally counted as an instance of the progressive.). The number of instances of *always* and *never* in the untagged corpora, MICASE, Switchboard, and SBC, were few enough that POS-tagging was not necessary; for the analyses conducted in chapter 5, manual searching for instances of the progressive among lists of *always* and *never* tokens in context sufficed.

Not all corpus work involves POS-tagging, but it does almost all involve concordancers, software which allows us to examine relationships between linguistic items and/or other features. These relationships, or “association patterns”, are at the heart of corpus research and fall into two types, linguistic and non-linguistic (Biber, Conrad & Reppen, 1998, p. 6). Linguistic associations are either lexical (“how [a] grammatical feature is systematically associated with particular words”) or grammatical (“how [a] linguistic feature is systematically associated with grammatical features in the immediate context”) (ibid.). Non-linguistic associations refer to associations between a linguistic item or feature and a particular genre, register, dialect, section of a text (e.g., the introductions of papers as opposed to their conclusions), and so forth.

One of the ways concordancers reveal association patterns between words (or lemmas; a lemma includes all grammatical forms of a given word) is by determining either the frequency of

particular collocations (i.e., pairings) or, alternatively, their “strength”, as determined by their Mutual Information (MI) score. The reason to use MI scores rather than or in addition to frequency counts is that frequency counts can “present a biased measure of the strength of associations” due to the fact that “more common words are more likely to appear in a collocate pair simply by chance” (Biber, Conrad, & Reppen, 1988, p. 265). The exact formulas used to calculate MI scores vary but, in all cases, the goal is to “compare the probability of observing the two words together with the probability of observing each word independently” (ibid., p. 266), thereby correcting for the problem inherent to using frequency counts alone.

Another commonly used feature of concordancers is the KWIC (“key word in context”) view. This feature displays all instances of the search item (which can be a word or, if the corpus is POS-tagged, a lemma) and the sentences/utterances in which it was found. This layout offers researchers an easy way to qualitatively assess a word’s common uses and meanings, along with gaining a sense of anything unusual or unexpected. Thus, it can be a preliminary step in the research that feeds into further analysis. Once an experiment is under way, KWIC views can also assist researchers in coding or in conducting additional qualitative analyses.

The concordancer used for Switchboard and SBC was AntConc (Anthony, 2011). Antconc calculates MI score using the equations in Stubbs (1995). For COCA, I was able to use COCA’s website tools. COCA’s built-in search mechanisms can interact with the POS-tagged corpus to allow users to search for particular words/phrases, lemmas, parts of speech, verbs in particular tense-aspects, or search phrases combining these things, as well as calculate the frequencies of search items and MI scores of those items’ collocates. It also offers users the option to limit their searches to particular genres or sub-genres, which can be combined or contrasted (e.g., newspaper articles only, or local news only, or humanities publications in comparison to science publications). COCA calculates MI score using the formula $MI = \log \left(\frac{AB * sizeCorpus}{A * B * span} \right) / \log(2)$ where A = frequency of node word, B = frequency of collocate, AB = frequency of the collocate near node word, sizeCorpus = number of words in corpus, and span = span of

words searched (e.g., 1 to the left and 1 to the right of the node word = 2).²¹

The website tools available for MICASE are not as advanced as COCA's, and MICASE is not tagged for parts of speech. Nevertheless, there are some distinct benefits to using it. Its interface allows one to restrict search results to the utterances of native speakers only, provides both KWIC views and entire transcripts, and provides information about the setting (e.g., office hours, lecture) of an interaction and each speaker (e.g., student, professor).

The final research instrument was the researcher herself. Manual coding was necessary at times because, for certain features, no automated coding schema is possible. Concerns involving the validity of such subjective/manual judgments have been addressed in the previous section, on research design: Clear coding criteria, rigorously applied and checked multiple times helped operationalize the process. Whenever possible, automated methods were favored, and in those cases manual processing served as quality control rather than the coding itself. Examples of features I coded manually are function (e.g., "He's always nagging me" would be coded as a complaint) and subject type (e.g., *lawyer* would be coded as human). The coding schemas used are described in more detail in the method sections of the respective studies.

3.4 Analysis of Collected Data

The data to be analyzed consisted, generally speaking, of association patterns, both lexical and linguistic, e.g., patterns of negation of *always* across genres, and associations between the function of a clause and the semantic characteristics of its grammatical subject. As a final step, observed correlations were tested for statistical significance using the chi-square test. These association patterns could then be interpreted functionally.

It is important to discuss, here, the validity of the chi-square test in this situation. Some (e.g., Davis, 2002; Kilgariff, 1996) have expressed the concern that, if the data points are mutually dependent on one another (as is often the case in corpus studies) rather than being

²¹ Information obtained from <http://corpus.byu.edu/MutualInformation.asp>. Accessed February 21, 2013.

fully independent, then the chi-square test is inappropriate. Yet this test “is probably the most commonly used significance test in corpus linguistics” (McEnery & Wilson, 1996, p. 84; see also Biber, Conrad & Reppen, 1998, pp. 273, 278), and offers the advantage of being more sensitive than the Student’s t-test (McEnery & Wilson, 1996, p. 84), though it cannot be used with very small numbers, or with frequencies rather than raw numbers (ibid.) Because my analysis does not critically depend on the significance test used (the patterns are generally very robust), I elected to follow common practice in corpus linguists and use the chi-square test, with raw numbers and only numbers large enough for it to work properly.

As explained in the research design section, it is important to balance quantitative techniques with qualitative techniques. This is crucial for keeping the researcher grounded in the actual utterances that lie behind the numbers and statistics under discussion. Otherwise, rather than being unable to see the forest for the trees, we will be blind to the trees and see only the forest. Thus, in coming to tentative conclusions about what *always* seems to mean in English, or what functions it may serve, I frequently returned to the corpora—not just to clarify points of confusion but also for leads and general insights that could feed into further analyses, and/or serve as examples in my arguments.

3.5 Limitations

As is the case with all studies, the generalizability of the findings presented here may be limited. Implications for cognition, in particular, must be stated with caution, given that this is not cognitive or psychological research, or even a cross-linguistic study. The addition of a biological component, e.g., EEGs or fMRIs, would be welcomed, and in keeping with the Cognitive Commitment (see §2.1). However, insofar as the purpose of the study is to describe two particular lexical items in modern American English, and given the choice of corpora (and their size), the results should be reasonably representative of how *always* and *never* work in various genres and contexts of that dialect.

This study is limited also in terms of the linguistic phenomena studied: For the most part, only *always* and *never* are investigated, and not other adverbs or even other adverbs of frequency. Moreover, a disproportionate amount of time is spent on constructions containing the progressive, an aspect actually rare with *always*. This was to investigate specific, prior claims regarding the progressive. It would be preferable, if time permitted, to have studied all tense-aspects equally thoroughly, as well as all adverbs of frequency.

3.6 Conclusion

In this chapter, I described the four corpora drawn upon in the study and the reasons for choosing them, as well as the general techniques (including which software and which measurements) I use. As was mentioned earlier, the methodologies are described in greater detail in the respective body chapters. In chapter 4, I calculate Exaggeration Quotients (as defined earlier) for certain genres and sub-genres, and the rate at which *always* and *never* are negated. In chapter 5, I determine the typical functions across genres of *always/never* followed by the progressive; in chapter 6, I study the tense-aspect preferences of *always* and *never*. Finally, chapter 7 contains a brief study of maximizers that confirms their tendency to pair with non-scalable items, and of near-synonyms to differ in collocation patterns.

CHAPTER 4: EXAGGERATION QUOTIENTS & GENRE

4.1 Introduction

This is a study of the distribution across genres of the adverbs of frequency *always* and *never* with respect to other adverbs of frequency, such as *often* and *rarely*, and of the rates at which *always* is negated. It is hardly news that the distribution, function and intended meaning of words can vary according to genre and context, but *always* and *never* are understudied in this respect—despite the fact that (as seen in chapter 2) degree words, antonyms, maximizers, and even other adverbs of frequency have all been investigated. This may be due to the assumption that these are not the sort of words that pick up unexpected meanings, and that their meanings are fairly inelastic (in contrast to the more subjective and context-dependent meanings of words like *often*). In formal semantics textbooks, e.g., Kearns (2000, p. 130), *always* and *never* are said to “express the same quantifications as the quantificational determiners,” i.e., *always* corresponds with *every*, and *never* with *no*, and little more is said.

While *always* and *never* have no highly idiosyncratic meanings, the present work shows that they are nevertheless far from simple. In this chapter, I analyze three data sets (written and spoken language, newspaper articles, and academic journal articles) which each contain multiple sub-genres. The analysis reveals a systematic pattern that suggests a propensity to exaggerate, i.e., to use *always* and *never* with non-literal meanings, in certain major genres (e.g., casual spoken language as opposed to written language) and sub-genres (e.g., humanities articles as opposed to medical articles, and local news articles as opposed to international)—a pattern whereby, roughly speaking, exaggeration is associated with informality.

I explain these differences across and within data sets in terms of (a) specific traits of casual language, such as a tendency to be more emotional and personal, and (b) accountability. That is, formality is best seen as a proxy for certain pressures, expectations, and purposes (which are associated with formality or informality) that motivate the discovered patterns.

4.2 Literature Review

The work done here echoes an earlier set of work on the subjectivity of the meanings of intermediate frequency adverbs such as *often*, *rarely*, and *frequently*. Pepper & Prytulak (1974), in *Sometimes 'frequently' means 'seldom'*, show the effect of context—namely, the typical, expected or base rate of frequency at which a given event occurs—on our understanding of quantitative expressions, and similar work has been conducted by Wallsten, Fillenbaum, & Cox (1986) and McGlone & Reed (1998). For example, the *often* in “She checks her email often,” could easily mean “several times a day,” since people spend a lot of time using the internet on computers and/or phones. In contrast, “She exercises often” is unlikely to mean she does it several times a day because doing so would be, at least for most people, highly impractical and, moreover, abnormal, perhaps a sign of a body image disorder.

Our successful interpretation of *often* in the examples above cannot come from the *often* alone, nor the basic definitions of *exercise* or *check email*; instead, it is at least partly derived from encyclopedic or world knowledge about those activities, as well as our personal experience, habits and inclinations—all of which traditional formal semantics relegates to the domain of pragmatics. Cognitive linguists, on the other hand, see such things as integral to the construction of word meanings. After all, barred from relying on context, we would have to interpret adverbs of frequency using a mental equivalent of Reid’s table, which matches quantifying phrases to explicit “absolute values” (e.g., “a few” is said to mean “3 to 5”).²² Because it strips the terms of context and massively restricts their range of possible meanings, Reid’s table cannot be a true representation of how these words are understood.

As Wallsten, Fillenbaum, & Cox explain, the flexible, context-sensitive range of possible interpretations of intermediate adverbs of frequency has important theoretical and practical

²² This table appears in an adapted form in Cazes (1983, p. 1557), and in what appears to be its original form at www.carelife.com/gordon/reids_table.html [sic; *Reid* is spelt incorrectly], accessed October 9, 2013.

implications (p. 571). This point is stressed most strongly by Nakao & Axelrod (1983), who argue in the very title of their work that *Verbal* [as opposed to numerical] *specifications of frequency have no place in medicine*, and who show how differently doctors and patients interpret words such as *often* and *infrequent*. Even Nakao & Axelrod, however, do not study the words on the far ends of the scale of frequency rates, *always* and *never*, the implication being that the meanings of these words, at least, are straightforward.

This is not true though; context affects our interpretations of *always* and *never*, as well, resulting in their meanings being much more flexible than they at first seem. Specifically, I focus on the fact that these words are frequently used in a non-literal, exaggerated sense. That we cannot always take *always* literally is clear from exaggerated and grouchy complaints (e.g., “You’re always losing things!” and “You never listen!”), as well as certain types of statements and advice (e.g., “You never can tell” and “You can always count on ...”). Used in such a way, *always* can refer to something that occurs merely frequently (or perhaps “too frequently for the speaker’s comfort”—clearly a very subjective measurement).

In a study of intensifiers (words “broadly concerned with the semantic category of degree” (Quirk et al., 1985, p. 589), a class which includes adverbs of frequency) in English, German, and Chinese, Jing-Schmidt (2007) found that, when used emotively, these words “are non-descriptive and resist a literal reading” (p. 425). Although her study did not include *always* or *never*²³ (as explained in §2.3, studies of intensifiers rarely do, and in fact rarely include adverbs of any sort), it suggests that *always* and *never* can be used non-literally as well.

More explicit evidence is found in the corpus study of Claridge (2011), who writes that “as long as no restriction of their scope is explicitly provided,” words such as *always* or *never*, or *every*, *anything*, *nothing*, etc., are “prone to taking on hyperbolic interpretation,” as in the comment “she’s allergic to everything” (file SBC1) (p. 51). The OED, as well, observes that

²³ Examples from the study include *very*, *quite*, *awful lot*, *tremendously*, *terribly*, and *damned*, as in *terribly kind* (p. 429) and *damned nice and cozy* (p. 431).

always is used to describe not just things that happen truly perpetually but also merely “for or throughout a long period.” In other dictionaries, hyperbolic meanings are even more explicitly listed as conventional sub-senses of a word (Claridge, 2011, p. 172).

The scope of Claridge (2011) encompasses much more than *always* and *never*; Claridge calculates the overall rate at which exaggeration, in general, occurs in casual spoken language. Specifically, she analyzed a portion of the spoken English section of the British National Corpus (BNC) (2001) (restricting the data to segments containing casual conversation, totaling 314,725 words; see Claridge, 2011, pp. 270-273) as well as all of Part 1 of the SBC (roughly a quarter of the SBC’s 249,000 words). In the BNC, she found “on average one overstated utterance per 1,000 words” and, for the SBC, nearly the same rate, 0.97 overstated utterances per 1,000 words (p. 72), thus verifying that word meanings in casual speech diverge from their literal interpretations, and quantifying the rate at which they do. The fact that this happens is not in itself remarkable; however, the dearth of research on *always* and *never* suggests a lack of awareness that it happens with these words, too, and of how, how often, and for what purposes. This study begins to fill that knowledge gap by delineating the patterns of how *always* and *never* work across genres, and when they can and cannot be taken literally.

Because it is not possible to determine, for every instance in a large data set, how literally *always* and *never* are being used, the propensity to exaggerate can only be extrapolated based on features of the text. In this study, this is done by dividing the number of instances of *always* by the number of instances of its less extreme counterparts *usually* and *frequently*, and by doing the same for *never*, *rarely*, and *infrequently*. I also determined the percentage of *always* tokens that were negated. The resulting measurements I refer to as the Positive Exaggeration Quotient (ExagQ+), Negative Exaggeration Quotient (ExagQ-), and negation rate, respectively.

These measurements are somewhat similar to a measurement in political psychology called “integrative complexity” (e.g., Tuckman, 1966; Tetlock, Hannum, & Micheletti, 1984; Tetlock, 1986; Suedfeld, Tetlock, & Streufert, 1992) which refers to precision, nuance, and

complexity. Integrative complexity scores are based on a number of factors—such as the author/speaker indicating awareness of “elaborate, contingent relationships” rather than only “dichotomous, good-bad categories,” and the ability to take others’ perspectives (Tetlock, 1985, p. 1570)—that go far beyond the use of quantifiers. Nevertheless, words such as *always* and *never* (and *all*, *absolutely*, *forever*, etc.) serve as “content flags” indicating that the lowest complexity rating may be appropriate, while words such as *usually* (and *probably*, etc.) suggest that the next-highest rating is called for (see Baker-Brown et al., 1990 for the precise coding schema). This suggests a connection between the use of adverbs of intermediate frequency, as opposed to *always/never*, and greater subtlety, caution, and accuracy in one’s statements.

It is imperative, in studying this phenomenon, to compare genres that are different but not too different. Although the most general distinction one can draw is probably that between written and spoken language, neither of these “is a unified phenomenon” (Chafe & Danielewicz, 1987, p. 84), leaving room for variation within and overlap between them (*ibid.*). It matters a lot, for example, if language is content-focused or not, narrative or not, formal or informal, planned or unplanned (Tannen, 1982, pp. 5-6)—characteristics that vary independently of a genre being written or spoken. Because of this, we must be cautious about exactly which samples of language we study, and take care not to “report ... findings as reflecting spoken vs. written modes” when, actually, they reflect “other facets of the genres chosen” (*ibid.*, p. 5), such as their different subject matter, different purposes, and differing degrees of formality. A pattern typical of early written versus spoken language studies, for example, is the comparison of especially casual spoken language to especially formal written language (*ibid.*), which makes it impossible to tell if discovered differences are due to mode (written versus spoken) or formality.

To control for such issues, this study includes three kinds of spoken language, three disciplines of academic journals and three kinds of news articles—all discussed in further detail in the next section. As Chafe & Danielewicz explain, as significant as the dichotomy between written and spoken language is, “the context of language use, the purpose of the speaker or

writers, [and] the subject matter” (1987, p. 84) are extremely important also. Comparison of highly similar genres should, in holding constant as many variables as possible, make it easier to observe the effects of such aspects of language on how *always* and *never* are used.

4.3 Method

4.3.1 Corpora and Data Sets

This study draws on all four corpora described in chapter 3, i.e., the Corpus of Contemporary American English (COCA), the Michigan Corpus of Academic Spoken English (MICASE), the Santa Barbara Corpus of Spoken English (SBC), and the Switchboard Corpus (henceforth Swb), out of which three major data sets were composed. The first data set, WR-SP, consists of written language (represented by the written portion of COCA) and three categories of spoken language: Highly natural spoken language (represented by Swb and SBC, combined); academic spoken language (MICASE), and unscripted TV/radio dialog (the spoken portion of COCA). The second data set, NEWS, consists of international, national, and local news articles (represented by the eponymous categories of COCA, in their entirety). The third data set, ACAD, consists of academic writing found in medical, science/technology, and humanities journal articles (again, represented by the eponymous categories of COCA, in their entirety).

The analysis of multiple data sets and very different sorts of language allows for investigation of particular genres while keeping the study grounded in “the big picture,” thereby discouraging erroneous overgeneralizations. In addition to representing both general and specific genres, the data sets represent a range of levels of formality. Brief characterizations of the genres and sub-genres are given below, followed by the details of how I obtained Exaggeration Quotients and negation rates.

4.3.2 Genre Characterizations

Regarding the first data set, WR-SP, written language is more formal than spoken language,

unscripted TV/radio dialog and academic spoken language (MICASE) are both more formal than casual conversation, and academic spoken language is likely more formal than unscripted TV/radio dialog. Regarding the ACAD data set, medical and science journals are more formal than humanities journals. Categories less easily ranked are NEWS with respect to ACAD or other genres of written language, and, within NEWS and ACAD, the different types of news articles with respect to each other and the two types of hard science articles with respect to each other. These characterizations are explained in greater detail below.

Spoken language, “arguably the most basic form of human communication” (Biber, Johansson, Conrad, & Finegan, 1999, p. 16), differs from written language in many ways, especially published written language (*ibid.*, pp. 15-17; Carter & McCarthy, 2006, pp. 9-11, 164-175; Linell, 2005, pp. 17-33), and these differences center around features which can, for the sake of simplicity, be viewed in terms of formality. We can start by considering their respective focuses and functions. Spoken language is typically interactive and focused “on the lives and interests of the interlocutors,” and its main purpose is communication (Biber et al., 1999, p. 16), while written language is not interactive, not directed at a particular interlocutor, and frequently used for disseminating information (*ibid.*), often of a technical nature. For example, a study by Biber, Conrad, & Reppen (1998)²⁴ found non-public face-to-face and telephone conversations to be much more personal than any of the types of written language analyzed (p. 164), as well as more involved (p. 152), where “involved” refers to containing, e.g., more personal pronouns, hedges, contractions, and amplifiers, and fewer agentless passives (p. 153).

As was emphasized earlier, however, there is cross-over between these two major categories. Because strategies typical of spoken language can be used in written language, and vice versa (Tannen, 1982; Chafe & Danielewicz, 1987, p. 84), we cannot treat spoken or written language as monolithic categories. To this end, I include in my study different types of spoken language, different types of academic articles, and different types of news articles.

²⁴ Based on the “five dimensions of language” established in Biber (1988).

Regarding the spoken language in this study, casual conversation like that found in the SBC and SwB is, by definition, the least formal. Beyond that, academic spoken language (MICASE) is likely more formal than unscripted TV/radio dialog. Admittedly, both are rather “unnatural”: Unscripted conversation from TV/radio shows is produced in a highly public and formal or semi-formal setting, constricted by a certain agenda, and recorded and broadcast. Academic spoken language is somewhat formal too (the university constitutes, for the instructors, a professional environment, and, for the students, at least a type of work environment), is quasi-public, adheres to an agenda, and is heavily content-focused. Moreover, academic spoken language is permeated with aspects of academic written language—which itself represents “the extremes of what writing permits” (Chafe & Danielewicz, 1987, p. 111). This is why the spoken language at universities can be as unfamiliar and “bewildering” to incoming students as the language in their textbooks (Biber, 2006, p. 1). University lectures, in particular, often emulate “some of the elegance and detachment of formal writing” (ibid). In fact, and though it happens to a lesser extent, the blurring of the conventions of academic speech and writing begins as early as elementary school (Cook-Gumperz & Gumperz, 1981).

The second data set is ACAD, academic written language (containing medical, science-technology, and humanities articles from academic journals). Descriptions of academic writing that discriminate between disciplines indicate that the sciences exhibit more features associated with formality than the humanities do. Biber, Conrad & Reppen (1998), for example, found science articles (ecology articles, specifically) to be less narrative-like (p. 160) and more impersonal (p. 164) than humanities articles (history articles, specifically). More recently, Hardy & Römer (2013) ranked the 16 disciplines represented in the Michigan Corpus of Upper-level Student Papers (2009) according to four dimensions, whereby positive rankings indicate, e.g., that texts were involved and narrative-like, and contained the present tense, active verbs, and expressions of emotion, attitudes, and opinions. The study found that, for all four dimensions, the two most positively-ranked disciplines were in the Arts/Humanities or Social Sciences, and

the two most negatively-ranked were in the Physical Sciences or the Biological and Health Sciences. Such research indicates that humanities writing is more casual (in that it is more like spoken language) than science writing. The two hard science categories I investigate, though, Sci-tech and Medical, are harder to differentiate. Thus, one sub-goal of this first study is to determine if or how these closely related categories differ.

It is also difficult to rank the formality of the third data set, NEWS, with respect to ACAD or other written genres, and to rank the different types of news articles (international, national, and local) with respect to each other. However, based on the intuition that international news articles focus on events with bigger, potentially global, consequences, I would predict that they have more formal features than national articles, and, based on the same line of reasoning, that national news articles have more formal features than local news articles.

4.3.3 Exaggeration Quotients

The data was analyzed in two ways. First, I used a ratio I call the Exaggeration Quotient, (ExagQ) which can be positive (ExagQ+) or negative (ExagQ-). ExagQ+ is the number of *always* tokens divided by the combined number of tokens of its less extreme potential alternatives *often* and the more formal *frequently*. Likewise, ExagQ- is the number of *never* tokens divided by the combined number of tokens of *rarely* and *infrequently*. The empirical basis for using ExagQ is discussed earlier, in the literature review. All data was collected on July 19, 2013. Differences in ExagQ values across categories within a data set were checked using the chi-square test for statistical significance at the $p < .001$ level.

Second, I calculated the percentage of *always* tokens preceded by *not* or contracted versions of *not*. This complements and serves as a check on ExagQ; high ExagQ+ values cannot be taken to indicate exaggeration if most instances of *always* are negated, thereby expressing something more like *frequently*. While obtaining the number of *not always* tokens was straightforward, the method of obtaining the number of negated *always* tokens containing

contractions varied by corpus: For COCA, I used COCA's website (which treats *'nt* and other contractions as independent words) to search for “*nt* always”; for Swb and SBC, I used Antconc (Anthony, 2011) to search for “**n't* always”; and, for MICASE, I used MICASE's website (which does not allow wildcards) to search, separately, for *always* preceded by *isn't*, *aren't*, *wasn't*, *weren't*, *don't*, *doesn't*, *didn't*, *hasn't*, *hadn't*, *won't*, *wouldn't*, *can't*, *couldn't*, and *shouldn't*. Differences in negation rates across categories within data sets were checked for statistical significance at the $p < 0.001$ level using the chi-square test.

4.4 Results

The ExagQ values and negation rates of *always* vary systematically across genres in ways that suggest a correlation between (to simplify the matter) exaggeration and formality. I discuss first ExagQ, then negation rates.

4.4.1 Exaggeration Quotients

The first data set, WR-SP, focuses on the most general distinction, between written and spoken language. As was stated earlier, the spoken data was further divided into academic spoken language (MICASE), unscripted TV/radio dialog (the spoken portion of COCA), and casual conversation (Swb and SBC, combined). The results (see Fig. 4.1) reveal that both ExagQ+ and ExagQ- are lower for written language than for all three spoken categories, and higher for casual spoken language than for the two other types of spoken language. Regarding the latter, ExagQ+ and ExagQ- are both lower for academic spoken language than for unscripted TV/radio dialog. It appears, then, that the propensity to exaggerate is weaker in more formal categories. Also noteworthy is how much higher the ExagQ- values are than the ExagQ+ values, such that the highest ExagQ+ value for any category (6.42) does not exceed the lowest ExagQ- value (14.26). This is characteristic of not just WR-SP but the other two data sets as well.

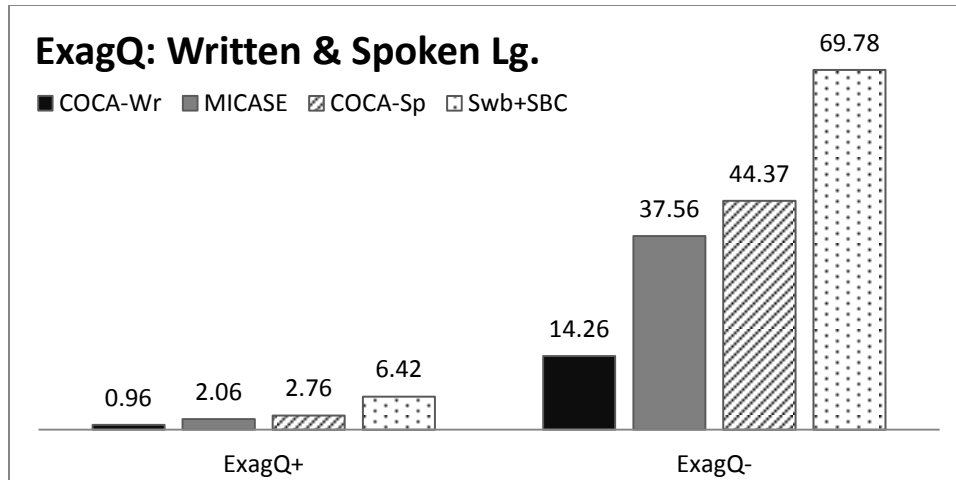


Figure 4.1: ExagQ, written & spoken language

Overall, though, the similarities in the patterns of ExagQ+ and ExagQ- values are more striking than any differences. They result in identically ordered category rankings, and the magnitudes of the differences between categories are similar too. For both types of ExagQ, the difference between casual spoken language and televised spoken language is the largest, the difference between academic spoken language and written language the second largest, and the difference between academic and televised spoken language the smallest. Normalizing the two data sets (see Fig. 4.2) accentuates these similarities.

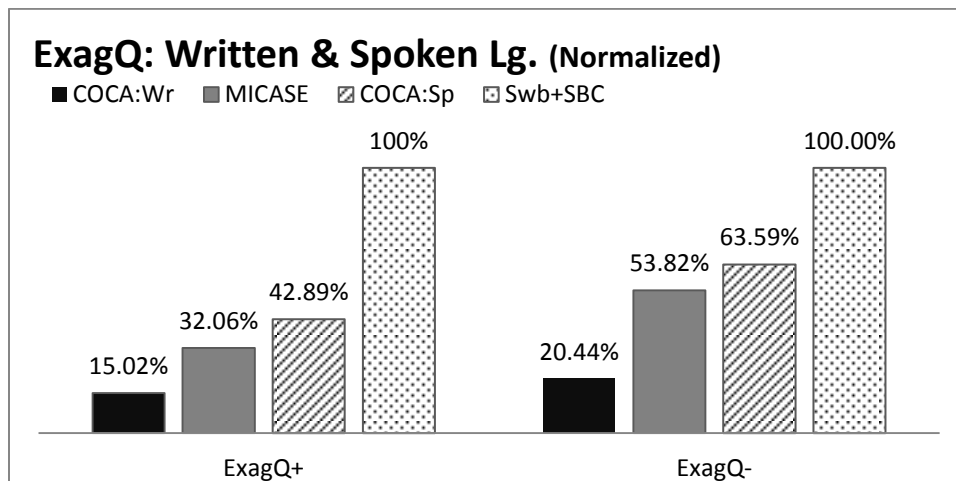


Figure 4.2: ExagQ, written & spoken language, normalized

This general pattern whereby the order of categories, in this case from lowest to highest

ExagQ values, is COCA written, academic spoken, COCA spoken, and casual language appears in the studies in chapters 5 and 6, as well. This supports the claim of Davies, creator of COCA, that the unscripted TV/radio dialog used as the spoken portion of COCA is roughly representative of casual conversation. That is, in terms of ExagQ, COCA's spoken language lies somewhere between written language and spoken language, and outranks academic spoken language.

The next data set, ACAD, contains three genres of academic writing: medical (Med), science/technology (Sci-Tech), and humanities (Hum) journal articles. The findings (see Fig. 4.3 and, for the normalized results, Fig. 4.4) show that both ExagQ+ and ExagQ- were lowest for Med, somewhat higher for Sci-Tech, and highest for Hum.

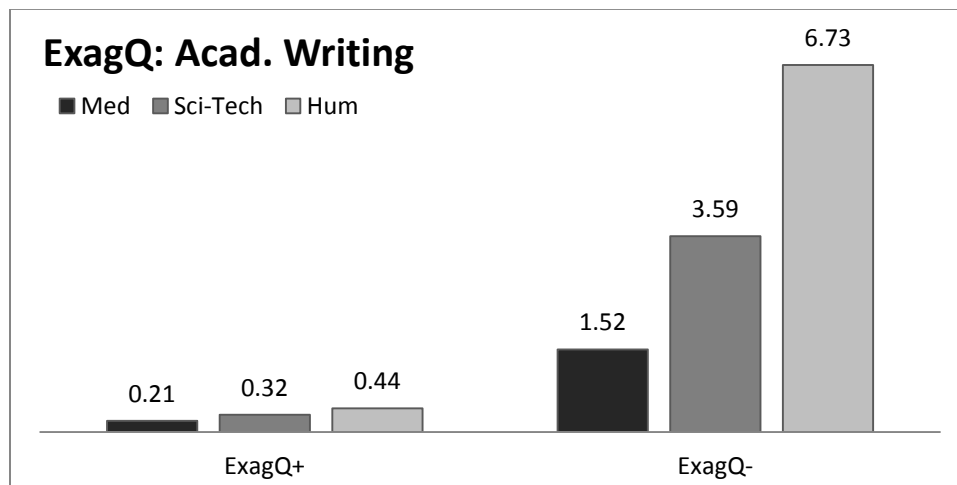


Figure 4.3: ExagQ, academic writing

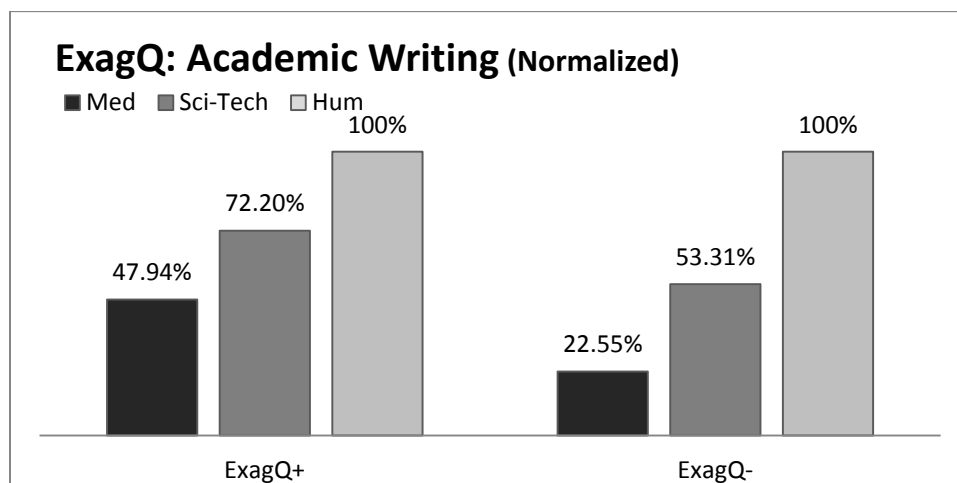


Figure 4.4: ExagQ, academic writing, normalized

The third data set, NEWS, consisted of international, national, and local news articles. It was found (see Fig. 4.5 and, for the normalized results, Fig. 4.6) that both types of ExagQ were lowest for international news, intermediate for national news, and highest for local news.

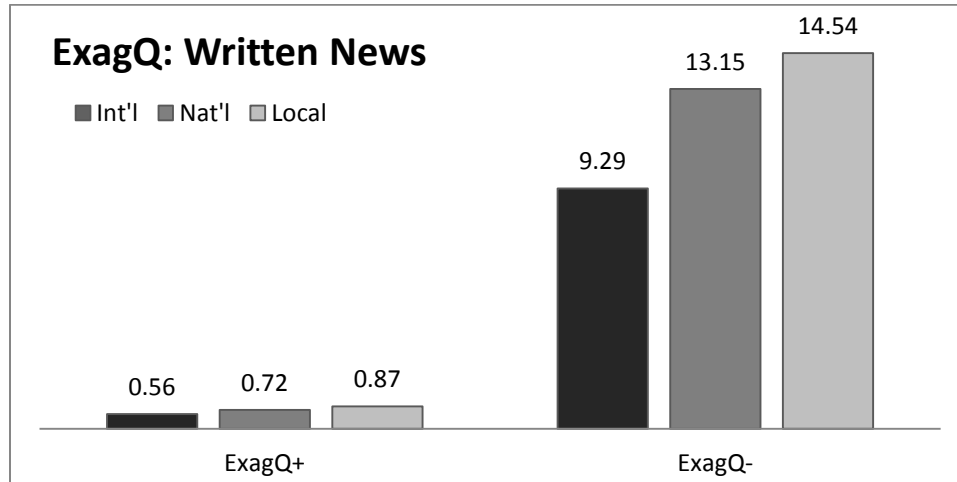


Figure 4.5: ExagQ, written news

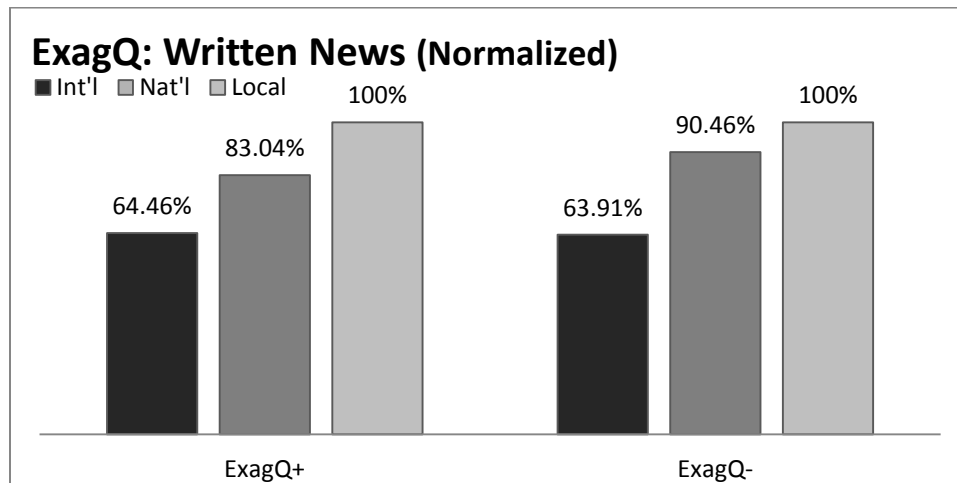


Figure 4.6: ExagQ, written news, normalized

Finally, stepping back, one sees clear patterns not just within each of the three data sets, but across sets (see Figs. 4.7 and 4.8, which show the ExagQ+ and ExagQ- values, respectively, for all ten categories). Both types of ExagQ values steadily increase as we move from ACAD to NEWS to WR-SP. There are, in fact, nearly no instances of overlap between the highest ExagQ value in one data set and the lowest in another (the one exception is that the ExagQ- value of

local news is slightly greater than that of COCA's written language overall).

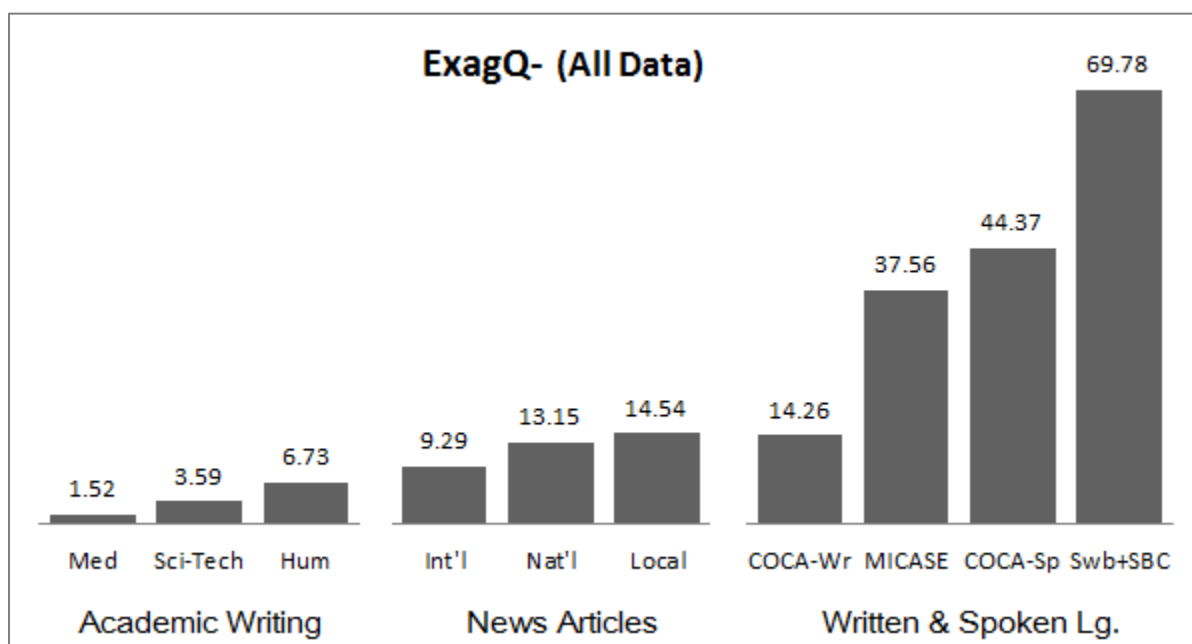


Figure 4.7: ExagQ+, all data

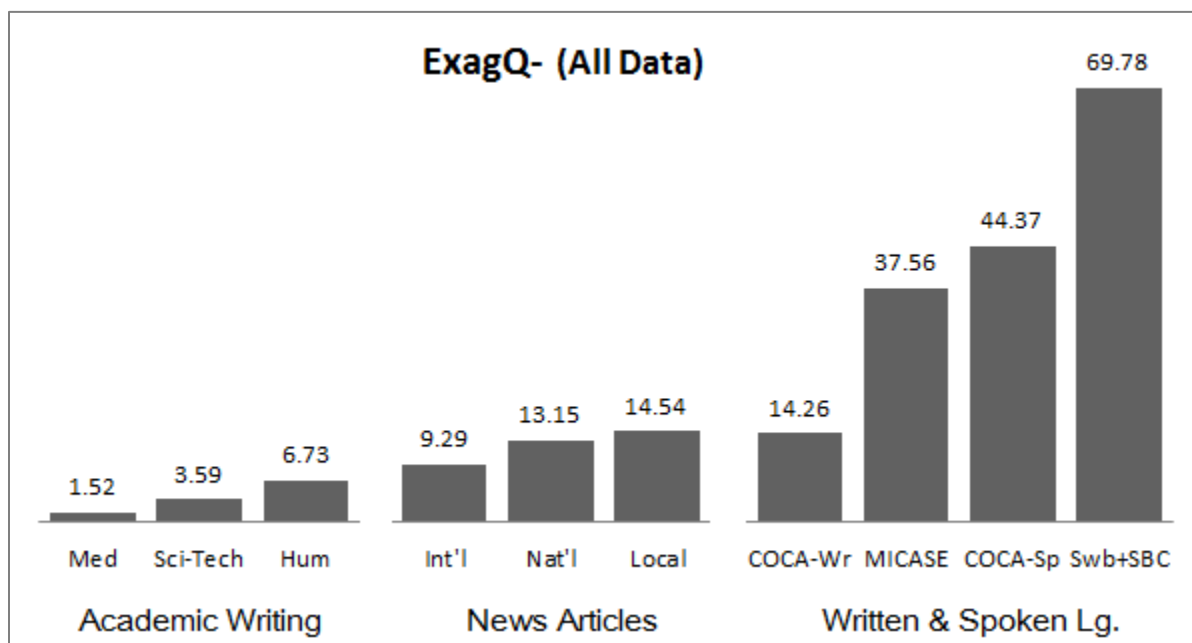


Figure 4.8: ExagQ-, all data

4.4.2 Negation of Always

For all categories, in addition to ExagQ, the rate of negation of *always* was calculated (see Fig.

4.9; lighter bars indicate results not statistically significant at the $p < .001$ level). Aside from one exception—and for the one data set, NEWS, for which the cross-category differences were not significant—the pattern of negation rates is the inverse of the pattern of ExagQ values: Within the data sets WR-SP and ACAD, those categories for which ExagQ+ and ExagQ- are lowest are the same categories in which *always* is most often negated. Likewise, those for which ExagQ+ and ExagQ- were highest are also those in which *always* is least often negated.

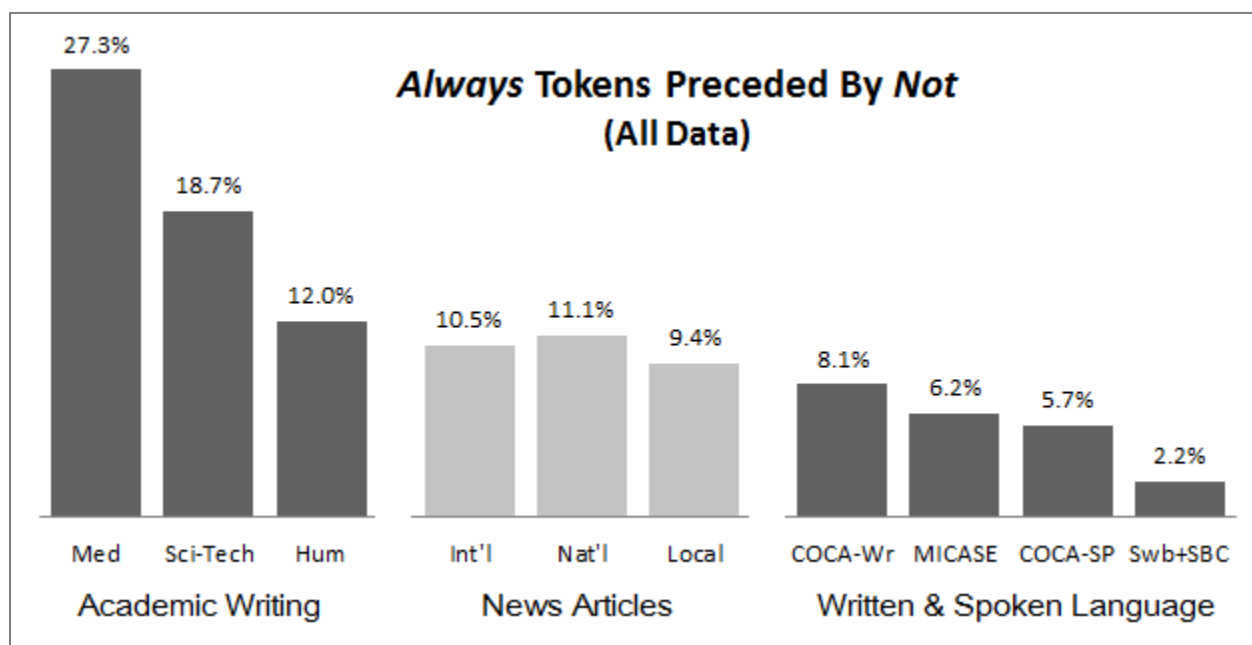


Figure 4.9: *Always* tokens preceded by *not*, all data

Specifically, within the ACAD data set, 27.4% of the *always* tokens in Med were negated, 18.7% in Sci-Tech, and only 12% in Hum. Within the NEWS data set (for which cross-category results were not statistically significant; $p = 0.2198$), *always* appeared with *not* 10.5% of the time in international news articles, slightly more often (11.1%) in national news articles, and least often in local news articles (9.4%). Finally, within the WR-SP set, *always* tokens were negated most often in written language (8.1% of tokens), less often in academic spoken language (6.4%), a little less often in unscripted TV/radio dialog (5.7%), and much less often (relatively speaking) (2.2%) in the most casual spoken language.

As was the case for the ExagQ values, there is no overlap in negation rates across data sets: *Always* was most likely (at a rate of between 12% and 27.4%) to be negated in ACAD, less likely to be negated in NEWS (rate of between 9.4% and 11.1%), even less likely to be negated in the written portion of COCA taken as a whole (rate of 8.1%), and least likely to be negated in the three categories of spoken language (rate of between 2.2% and 6.4%). Speaking more generally, *always* was more likely to be negated in the written portion of COCA as a whole, especially in each of the six sub-sets of COCA's written genres examined, than in any of the spoken categories, especially the most casual spoken language.

4.5 Discussion

The findings presented above reveal a consistent pattern whereby the same categories that exhibit lower ExagQ+ values also exhibit lower ExagQ- values and higher rates of negation of *always*. This makes intuitive sense, given that all three appear to indicate more careful, accurate language. The categories characterized by low ExagQ values and high negation rates are also those categories known or reasonably expected to be more formal, which seems to indicate that one should seek reasons for a connection between exaggeration and informality. If we are even more specific than that, however, we can explain a great deal more.

I propose two explanations, the first involving the (more personal, emotional) nature of casual conversation and the second involving the effect of accountability on published, written language. I then argue that, overall, the data presented here support the view that the meanings of lexical items are derived via not just dictionary-like definitions but also via context and the encyclopedic knowledge to which the items serve as a point of access.

4.5.1 The Nature of Informal Speech

It was mentioned earlier (§4.2) that Claridge (2011), in studying portions of the BNC and SBC, found exaggeration to occur at a rate of 1 and 0.97 instances per 1,000 words, respectively.

Because of her deliberate focus on extremely casual language, she did not study more formal genres as well. However, the work of others and even Claridge's work on casual language shows a connection between literal language and formality. Claridge discovered that the slightly lower exaggeration rate of the SBC was accounted for mainly by the fact that four of the 14 analyzed files contained no overstatements at all. And, even more interestingly, "three of those files are highly task-related interactions with a concentration of factual information" and two of them "represent 'public', business-like events, in a lawyer's office ... and in a bank" (p. 71)—i.e., they actually involved rather formal language. In addition, Claridge shares the complementary finding that "the highest amount of hyperbole is found in the *SBC* in informal conversations between friends," i.e., in files SBC6 and SBC13 (ibid.). Thus, even within data selected to be highly casual, more literal language was associated with formality.

Further support for a connection between casual settings and exaggeration is found in Link & Kreuz (2005, cited in Claridge, 2011, p. 73), who studied "non-literal language, including hyperbole" in emotional talk and found it to occur at the rate of five instances per 1,000 words, five times the rate Claridge found in the BNC and SBC; Claridge attributes this directly to the emotionality (p. 73). On a similar note, Jing-Schmidt (2007) shows that intensifiers used emotively are typically not intended to be taken literally. Thus, Claridge (2011), Link & Kreuz (2005), and Jing-Schmidt (2007) all converge in demonstrating that non-literal, exaggerated, or hyperbolic language is associated with informal conversation, which is typically non-public, non-task-related, more emotional, between close friends, etc..

What is most important, however, is not formality itself but the more specific aspects of informal language (private, personal, emotional, less task-oriented, etc.) that Claridge and Link & Kreuz highlight. Not coincidentally, these characteristics are closely associated with the factors stressed by Chafe & Danielewicz, i.e., "the context of language use, the purpose of the speaker or writers, [and] the subject matter" (1987, p. 84). It is these practical and functional matters that largely determine how literally *always* and *never* are used and interpreted.

Consider, for example, the “more emotional” characterization of casual language. The expression of emotion has particular effects and thus can help speakers accomplish certain things. For example, in her very function-oriented definition of emotive intensifiers (words like *very* or *terribly*), Jing-Schmidt (2007) explains that these items

can be used (1) to boost the speaker’s illocutionary force and especially to maximize the dramatic effect in communication and (2) to elicit attention from the hearer in conversation and, under certain circumstances, (3) to establish rapport between interlocutors. (p. 428)

These effects are due to the “metaphorical mapping between the emotional domain and the linguistic,” which “is made possible by the metonymic understanding of both domains in terms of their most noteworthy element—HIGH INTENSITY” (p. 433), i.e., we associate intense emotions with lexical items indicating extremes. Moreover, because of the “intersubjective” (p. 425) and “interpersonal” (p. 426) nature of the aforementioned functions, these words are not used to provide “evaluations of measurable degrees” that are basically “accountable,” as found in “impersonal” descriptions (*ibid.*); on the contrary, they likely involve exaggeration.

By all three measures (ExagQ+, ExagQ-, and negation of *always*) the three types of spoken language studied appeared to be less literal in their use of *always* and *never* than any of the varieties of written language. In addition, among the spoken language categories (and, again, by all three measures), casual language (Swb and SBC) was the least literal, unscripted TV/radio dialog (COCA-Sp) intermediately literal, and academic spoken language (MICASE) the most literal. This is consonant with the findings of Claridge (2011) and Link & Kreuz (2005) in that the more casual (i.e., less task-oriented, and potentially more emotional) spoken language contains more instances of hyperbole than the academic or televised language.

Moreover, it makes sense that—as formal as some radio/television shows might be—it is academic spoken language that contains the least exaggeration. The topic and purpose of discourse in classrooms and discussion sections and office hours and so forth is largely pre-set, explicitly task-oriented and focused on factual information. Students and instructors are

expected to focus on the topic and task at hand and unlikely to devote much time to personal topics such as friends and family—or, if they do devote time to these topics, they typically do so in an emotionally reserved way. That is, the goal would usually be to relate one’s personal experiences to the current academic discussion, and not so much to vent frustrations, share emotions, seek empathy, and so on. Finally, as was stated earlier (§4.3) academic spoken language shares traits with academic written language, and literate/written language in general—which, as the results for the ACAD (academic journals) and NEWS (news articles) data sets here indicate, tend toward more literal usage of *always* and *never*.

Unscripted TV/radio dialog, as well, is influenced by pre-set agendas. Personal topics do come up, perhaps even quite frequently, but one assumes that most people do not present themselves, their personal lives, and their emotions as openly on a publicly broadcasted TV/radio show as they would in completely naturally-occurring spoken language among close friends. Additionally, unlike the interactions in a university setting, many talk show conversations are far less factually-oriented. It is less necessary to be accurate and literal if the conversation is more geared toward social functions/topics such as gossip, entertainment, and so on than toward conveying objective scientific facts, or teaching certain skills.

The final connection between casual language and exaggeration to be discussed involves complaints: It turns out that “the majority of hyperbolic expressions are used for the purpose of negative evaluations” (Claridge, 2011, p. 162; see also p. 81, table 4.5), and, more specifically, for rather emotional, impolite complaints (as opposed to, e.g., an objective but critical evaluation in a literature course of a novel written by a non-present author). For example, Hartung (1996, p. 157, cited in Claridge, 2011, p. 81) shows that, in German, hyperbole is commonly used in negative evaluations and especially to express irritation. These types of complaints, in turn, elicit strong emotions in the addressee; Legitt & Gibbs (2000, cited in Claridge, 2011, p. 162) show that addressees rank the use of hyperbole in negative contexts as having a strong adverse emotional impact—which makes them quite effective (Claridge, 2011, p. 142), actually.

Because casual language tends to be more emotional and/or personal, the connection between hyperbole and (emotional, emotion-inducing) complaints further explains the high ExagQ values and low rates of negation of *always* associated with the most casual spoken language. (In the next chapter, I discuss in further detail how the words *always* or *never*, and the progressive aspect, help achieve this exaggeration effect, as well as more general reasons for the prevalence of the expression of negative affect in language.)

The observed differences within the two single-genre data sets, ACAD and NEWS (i.e., between the three types of academic articles, and between the three types of news stories), on the other hand, cannot be explained in terms of the characteristics of casual versus non-casual speech. Presumably, neither academic articles nor newspapers are appropriate arenas for exaggerated complaints. Instead, for these data sets, the more suitable explanation involves accountability. This is the topic of the next section.

4.5.2 Accountability

The two data sets, NEWS and ACAD, which consist entirely of published writing, exhibited within-set variation in ExagQ values. In NEWS, exaggeration was least common in international news articles, more common in national news articles, and most common in local news articles. In ACAD, exaggeration was less common in science articles than humanities articles and, regarding the two science categories, rarer in Medical than in Sci-Tech articles. Beginning with Med and Sci-Tech articles, to understand these findings we can—thinking in terms of very practical concerns—ask why writing in the medical field would be more literal in its use of these words. The main reason for this, I propose, is that the stakes (the risk of harm as a result of one’s words being misinterpreted) are higher in Med than in Sci-Tech.²⁵ In other words, writers of medical articles are especially constrained by concerns about accountability.

²⁵ Science and technology are extremely influential in our lives too, of course, and crucial to medicine, but the medical field always directly pertains to people’s bodies and health, while science and technology may or may not.

This issue is the focus of Nakao & Axelrod (1983), who are so concerned about fuzzy terminology in medicine leading to misunderstandings that could cause bodily harm that they argue that all statements of frequency in this context should be expressed numerically. These concerns are valid. Ideally, doctors follow the literature in their field and specialty, and the information they gain from these sources informs their actions (their interpretation of symptoms and resulting diagnoses, and what medicines or courses of actions they prescribe). Doctors also relay information obtained from the literature to colleagues and patients. Since doctors and patients operate with different sets of background knowledge, the use of non-numerical terminology by doctors has the potential to be understood by patients differently than the doctor intended, and could also be misunderstood by other doctors.

Although accountability may be especially relevant for medical journals, it is presumably relevant in all published academic writing. This explains why authors in the ACAD data set appear to use *always* more carefully than authors of newspaper articles (and, of course, more carefully than participants in casual conversation). Newspapers constitute published writing as well, and one hopes that their writing, too, is held to a high standard of accuracy. However, journalists are generally employees of the newspaper, and their work—though it goes through an editorial process—is not subject to the rigorous peer-review process that an academic manuscript goes through. There is no simply time, or we would receive breaking news months after the fact. Furthermore, newspaper journalists typically do not, as is the case in Med and a lot of Sci-Tech writing, report precise numerical results of complex experiments.

Continuing to pursue the line of thought that higher ExagQ values and lower instances of negation of *always* correspond with lesser accountability, we are led to surmise that local news stories writers are less affected by concerns about accountability than writers of national and international news stories. I propose that this is likely true, and that this is so for two reasons, which I expand on below: First, readers of international news compared to readers of local news are more highly educated and thus more discerning. Second, international news is of an

inherently higher stakes nature than national news, as is national news compared to local news, and thus likely to be associated with higher expectations regarding accuracy.

Those who are highly educated are in a position to be more discerning of the accuracy of news stories for at least two reasons: To begin with, they have undergone explicit training in critical thinking skills, including how to evaluate the quality and reliability of information and sources. In addition, they, on average, possess more background knowledge regarding the situations reported on. In a representative sample of Americans surveyed (Pew Research Center, 2007), education was “the single best predictor of knowledge” of current affairs in the news and, “holding all other factors equal, levels of knowledge r[ise] with each additional year of formal schooling” (p. 3). Specifically, holders of postgraduate degrees correctly answered an average of 17 of the 23 questions about current affairs, whereas respondents who had not finished high school answered only about eight correctly (ibid.). It is likely that this gap exists because education is positively correlated with news consumption overall (Benz, Tompson, & Rosenstiel, 2014, p. 31; see also Dutta-Bergman, 2004, p. 52, regarding online news).

For the proposed argument to hold true, we would also need to establish a link between education and readership of international versus national versus local news. Unfortunately, most studies break news down by networks or sources (which tend to report on all three types of news, even if they are a “local” news station, etc.) or medium (namely, print versus online), and not by the topics of stories being local, national, or international. Moreover, the one study (Pew Research Center, 2012) that does categorize stories in such a way and categorizes respondents by education level actually shows a higher correlation between education and the viewing of national news than the viewing of international news (p. 31) (and omits the results regarding local news and education level). Nevertheless, multiple studies, discussed below, support the proposed connection between international news consumption and education.

The measurement mentioned earlier, Integrative Complexity (partly determined based on the use of words such as *always* versus words such as *frequently*), has been found to be

correlated with author/speaker IQ (Simonton, 2006). We also know that IQ scores are related to one's level of attained education (Keage et al., 2014, Tommasi et al., 2015, etc.). To the extent that journalists with higher IQs attract readers who are similar, the observed ExagQ values suggest that the average reader of international news is more educated than the average reader of national news, and even more so than the average readers of a local news story.

In general, international news outlets do appear to be more sought out by the more highly educated. For example, the Pew Research Center (2007) found that a greater percentage of college graduates follow National Public Radio (NPR) (40%) and CNN (30%) (which focus on non-local news) than local TV news (25%) and network morning shows (25%) (which include news from all three categories, the former stressing local news more) (p. 12). Similarly, Benz, Thompson, & Rosenstiel (2014) found that, although the correlation between education and news consumption held true for all four categories tested (newspapers, radio news, magazines, and newswires such as the Associated Press (AP) and Reuters), the effect was strongest for newswires, used as a news source by 44% of Americans in the "college or more" category, 30% of those in the "high school or some college" category, and 14% in the "less than high school" category (p. 32). AP and Reuters are multi-national/international news agencies, the former based in New York and the latter in London, focusing mainly on international stories. Dutta-Bergman (2004), as well, studying online news consumption, found that "interest in international and global issues increases with an increase in education" (p. 57).

Some studies divide news stories not into international, national, and local but instead into "hard" and "soft", with the former often involving international news. Iyengar, Hahn, Bonfadelli, & Marr (2009) tested American and Swiss citizens' knowledge of hard and soft news topics. At least from the perspective of the Americans, the four hard news topics used were foreign/international news (e.g., which countries were part of the coalition to disarm Iraq), and three of the four soft news topics were national issues (e.g., the Michael Jackson molestation

charges). The finding was that “although education-related differences were pervasive in both news dimensions, they were more pronounced for the hard news subjects” (p. 353).²⁶

The second reason I offered above for international news being held to greater standards of accountability than local news is that it is of a higher stakes nature. Admittedly, news should be accurate and reliable at all levels; writers of national or local news stories are not immune to accusations of libel by the politicians, celebrities, organizations, or even local people that they write about, for example. However, inaccurate statements about other nations (and especially other nation’s leaders) could have especially detrimental repercussions for the journalist, his/her news corporation, and even the nation—particularly statements attributing acts of war or other acts of violence to nations or their leaders.

In a sampling of my own data, stories in which the journalists attribute acts of violence (most commonly war but sometimes other acts, such as violently repressing a mob, or murder) to nations or, less often, to individuals were most common in stories about international affairs and absent in stories about local affairs.²⁷ Specifically, 11 of 30 randomly sampled international news story headlines, nine of 30 national news story headlines, and no local news story headlines were of this nature. This suggests that accountability is indeed likely to be a greater concern in international news than national news, and local news.

4.5.3 *The Encyclopedic View of Meaning*

The significance of these results from a theoretical standpoint is not just that the distribution of *always* and *never* change depending on the genre in which they appear, but that their perceived

²⁶ It is arguably a problem that three of the four soft questions involved American celebrities/athletes, but the Swiss did better on these questions than the Americans, and despite less coverage of them in Switzerland, because the Swiss are (as was the study’s main finding) more attentive to news, overall (p. 348).

²⁷ I downloaded the list of COCA’s sources (available at http://corpus.byu.edu/coca/help/coca_2012_06_22.zip, accessed September 20, 2013), deleted everything besides news stories labeled as international, national, or local, randomly shuffled them, and checked the first 30 of each type (skipping articles if I could not discern the topic from the title, e.g., “Letters to the Editor”, and also could not look up the article online). Whenever possible (even if the topic was obvious based on the headline), I found the full text of the articles online and double-checked the topic.

meanings change as well, showing that these words are flexible and sensitive to context. In general, the findings support the view of meaning crucial to cognitive linguistics which holds that meaning is encyclopedic (Evans & Green, 2006, pp. 160, 173). Words are not “neatly packaged bundles of meaning” but instead “‘points of access’ to vast repositories of knowledge” (p. 160). That is, they are “‘prompts’ for the construction of meaning” (p. 173), which is itself a complex process largely guided by encyclopedic knowledge and context. It is our experience with and knowledge about breakfast, for example, that enables us to understand “I always eat breakfast” as having the “at every opportunity” meaning rather than the “constantly” meaning, and, more specifically, as meaning “every morning”. It is also what would allow us (because we know that the “base rate” for eating breakfast is once a day) to interpret “I often skip breakfast” as meaning something like “several times a week”.

The situation regarding the literal versus non-literal meaning of *always* and *never* is similar. For example, as Claridge (2011, p. 51) explains, we know not to take “she’s allergic to everything” (SBC1) literally because truly being allergic to all things in the world “would imply the impossibility of leading any sort of normal life, which is obviously not the case for the person talked about” (who spends time at stables, and rides horses, etc.). We would likewise know not to take “She’s always yelling at him” literally because in this case too the corresponding real-life situation is incredibly unlikely, or impossible.

Our interpretations of *always* and *never* varying based on encyclopedic knowledge of the concepts and situations referred to in their immediate linguistic context, and they are also likely to be understood differently based on genre or other aspects of context. In some genres, such as medical journals, more accountability and a more objective, less personal, less emotional tone can be assumed and, thus, the adverbs of frequency can safely be taken more literally. In other genres, such as casual spoken language, the opposite is true. Interpretations also vary within genres. In certain sub-sets of genres the likelihood of the adverbs being used literally or not may be stronger or weaker (for example, in a heated argument between unhappy spouses, words are

especially likely to be used non-literally, perhaps in complaints or accusations).

4.6 Conclusion

In this chapter, I analyzed three data sets, each containing three to four sub-genres of modern American English, in terms of ExagQ+ and ExagQ- values, and the rate at which *always* is negated, and found the three measurements to be highly correlated: Those categories within a set for which ExagQ+ was the highest were also those categories for which ExagQ- was highest and the rate of negation of *always* was the lowest, and vice versa. And, speaking broadly, the data sets and sub-genres within them for which ExagQ+ and ExagQ- were high and the negation rates low were also less formal.

Taking a cognitive-functional approach, I asked what practical concerns surrounding the nature or function of each of the genres might motivate the observed patterns. I argued that the differences between written and spoken language, and between the three types of spoken language, were related to the nature of casual versus formal language, e.g., casual language is more personal, more emotional, more conducive to complaints, and less content-driven and task-oriented. The differences between the three types of academic articles and between the three types of news stories, on the other hand, I argued were due largely to concerns about accountability. Finally, in general, I argued that the meanings of *always* and *never*, despite seeming relatively inflexible, are underspecified, often requiring encyclopedic knowledge, genre, or other aspects of context to help us to fully flesh them out.

These explanations are functional in that they are grounded in speakers' actual experiences using language to accomplish particular goals in their daily lives. The greater accountability associated with certain genres is a very real pressure, related to real consequences writers and speakers could face. Likewise, any socially skilled language user knows and puts to use the knowledge that strong expressions of emotion, complaints, and blatant exaggeration are appropriate only in certain contexts. In such cases, the use of non-literal *always* is an effective

way to express certain sentiments and achieve one's goal. Thus, it is not the general, abstract concept of formality that explains the ExagQ values and negation rates reported here, but the restrictions, expectations, and functions inherent to certain genres and registers.

CHAPTER 5: FUNCTIONS WITH THE PROGRESSIVE

5.1 Introduction

Adverbs of frequency in preverbal position (called preverbal adverbs of frequency, or PAFs) combined with the progressive aspect have drawn special attention. This is so despite the fact that the progressive is rare, overall, in English (as stated in §2.4 and as is verified in chapter 6). Grammars of English, including several influential ones discussed below, typically concur that when the progressive aspect is used in combination with PAFs in general, or with *always* in particular (and similar PAFs, such as *continually*), the situation is often negative in some way. To test this claim for *always* and *never*, I collected *always/never* + progressive tokens from four corpora and analyzed their functions (coding them as serving the emotionally/socially neutral function Describe, the negative functions Complain or Lament, or the positive function Praise). In doing this, I was also testing the legitimacy of extending characterizations of one PAF to another. If different PAFs behave very differently—a distinct possibility with antonyms—such generalizations are unjustified. Finally, I analyzed the *always*-containing tokens further by testing for correlations between a given sequence's function and the nature of its grammatical subject, and between function and genre.

First, and contrary to the literature, PAF + progressive sequences most often served the neutral function, Describe, and could be divided into several sub-functions. At the same time, the negative functions were more common than the positive function, and corresponded closely, in their details, to characterizations given in the literature. Second, the more rigorous analysis of *always* + progressive revealed connections between both function and grammatical subject, and between function and genre. Regarding grammatical subjects, first person humans are the least likely to be complained about, and third party humans the most likely. In addition, first person humans are the most likely to be described neutrally. Regarding genre, description was more common in academic spoken English than in other genres. Moreover, that genre contained no

instances of praise, and fewer complaints than other genres.

Third, the largest differences discovered between the two adverbs involved tense. The *be going to* future (henceforth simply the *go*-future) was much more common with *never* than with *always*, and, even exclusive of the *go*-future, three-quarters of *never* + progressive tokens carried a future meaning, while none of the *always* + progressive tokens did. On a related note, a new sub-function of Describe, “to make a vow,” appeared in the *never* data. Finally, the verb phrases *come/go (back/home)* were very common with *never* but rare with *always*.

The analysis in this chapter is driven by a functionalist, “language as action” perspective. That is, I base the analysis on the understanding that speakers use language to accomplish, in interactions with other human beings, various goals (see Evans & Green, 2006, p. 110). This focus on practical and social considerations is extremely useful for explaining the connections between function and grammatical subject, and between function and genre. I conclude by highlighting the need to be sensitive, in selecting examples, to multiple components of the linguistic utterance of interest. Particularly worth considering, due to their many social and practical ramifications, are the semantic features of grammatical subjects.

5.2 Literature Review

Several comprehensive English grammars comment on *always*/PAF + progressive, saying that it is fairly rare (Carter & McCarthy, 2006, p. 47; Celce-Murcia & Larsen-Freeman, 1999, p. 509) and that, when the progressive does occur with PAFs, the situation referred to is often negative or unpleasant (Carter & McCarthy, *ibid.*; Quirk, Greenbaum, Leech, & Svartvik, 1985; Sinclair, 1990). Specifically, Carter & McCarthy write that *always* + progressive “often refer[s] to regular events or states which are problematic or undesired” (p. 47), while Quirk et al. say that, combined with *always*, *forever*, or *continually*, the progressive “often imparts a subjective feeling of disapproval to the action described” (p. 199). Likewise, Sinclair writes that PAFs combined with the present and past progressive “emphasize” the frequency of the action and (p.

249, p. 253) and thus are “often [used] to express disapproval or annoyance” (p. 249 on the present progressive, cf. p. 253 on the past progressive.) Moreover, he notes that *always* and *forever* are the ideal adverbs of frequency for accomplishing this (p. 253)²⁸.

Celce-Murcia & Larsen-Freeman (1999) explain things slightly differently, saying PAF + progressive is used “when the speaker’s message carries emotional overtones” (p. 510), overtones which can be either disapproving or approving, actually (p. 117). And, similarly to Sinclair, though they write about PAFs in general (p. 510), they clarify that their comments are especially relevant to *always* and *forever* (p. 117). Finally, it is worth noting that even the Oxford English Dictionary Online includes, in its first definition of *always*, “sometimes with the implication of annoyance” (2015, *always*, def. 1). In sum, PAF + progressive and especially *always* + progressive is said to often involve negative situations, and/or situations that involve strong (and usually but not always negative) emotions, such as disapproval or approval.

A trend shared by the above grammars is the adoption of a functional approach. As Evans & Green explain, this “is any approach that places particular emphasis on the communicative and social functions of language, and attempts to explain the grammatical properties of language in terms of how it is used” (2006, p. 759). Sinclair is the most explicit about this, identifying his work as a “functional grammar”, i.e., one which “puts together the patterns of the language and the things you can do with them” (1990, p. v). Similarly, Carter & McCarthy state that they are interested in “the communicative acts most typically performed by particular items” (2006, p. 14). Celce-Murcia & Larsen-Freeman, likewise, clarify early on that they “view grammar with a communicative end in mind” (1999, p. 4). As such, they organize their explanations around the uses of grammatical structures, in terms of both what they do/express, and in what contexts they are typically, appropriately used.

The other trend regarding these grammars is that they are mostly corpus-based. (The

²⁸ With the past progressive; regarding the present progressive (p. 249), he is less explicit, but in both sections his examples contain only *always* and *forever*.

exception is Celce-Murcia & Larsen-Freeman (1999), but they do incorporate quantitative data in the form of acceptability judgments). Carter & McCarthy (2006) derive many of their examples from the 700 million word Cambridge International Corpus (CIC), which contains everyday written and spoken English of different national varieties, such as American and Irish (Carter & McCarthy, 2006, p. 11). CIC also includes the five million word Cambridge and Nottingham Corpus of Discourse (CANCODE), which consists of “naturally-occurring, mainly British (with some Irish), spoken English, recorded in everyday situations” (ibid.). Every example in Sinclair (1990) is from the 20 million word Birmingham corpus (Sinclair, 1990, p. xvi). Quirk et al. (1985) appear to create their examples; however, in making acceptability judgments they rely on a combination of “elicitation experiments with informants in the United States and Britain” and frequencies found in multiple corpora, mainly the Survey of English Usage (SEU) (spoken and written British English), the Brown University corpus (published American English) (Francis & Kučera, 1964), and the Lancaster-Oslo/Bergen corpus (LOB) (published British English) (Quirk et al, 1985, p. 33).

Despite their reliance on corpus or other quantitative data, the grammars mentioned above present their claims about PAF + progressive without reference to quantitative evidence. And, as useful as the authentic examples in Carter & McCarthy (2006) and Sinclair (1990) are, the example sets are not large enough or complete enough to enable us to draw any reliable conclusions. In the same functional spirit of the grammars, but endeavoring to be more precise, I use quantitative and qualitative analyses of corpus data to learn specific details about the PAF + progressive construction, and what influences the observed usage patterns.

5.3 Method

The method consisted of obtaining a sub-set of instances from COCA and all instances from Switchboard, SBC, and MICASE of *always* immediately followed by the progressive, along with the immediately-surrounding context, and coding them in terms of their discourse function. In

addition, the *always*-containing tokens (which were far more numerous than *never*-containing tokens) were also coded according to their grammatical subject type and genre.

5.3.1 Initial Collection of Tokens

The tokens obtained from COCA consisted of all instances in the 2010-2012 data of *always* or *never* immediately followed by a main verb in progressive form, and in past or present tense.²⁹ Using the online interface, I searched the 2010, 2011, and 2012 data for all instances of “always [v?g*]” and “never [v?g*]”. (The “[v?g*]” denotes a verb in progressive form.)³⁰ For all tokens, I also collected the immediate context (KWIC view) and slightly extended context.

The methods used for the other corpora were similar but simpler. To obtain the MICASE tokens, using MICASE’s online interface and limiting the search to the utterances of native speakers of American English, I searched for all instances of *always/never* and downloaded the results, including KWIC views. For certain ambiguous tokens, I collected an even larger sample of the context. To obtain the SBC and Swb tokens, I used AntConc to search the .txt files of these corpora for all instances of *always/never*, and collected the KWIC views. I then manually omitted, from the MICASE, SBC, and Swb tokens, all instances except those in past or present progressive and in which the PAF was immediately followed by the main verb.

5.3.2 Quality Control Check of Tokens

The types of utterances permitted were declaratives, questions, and passives (including the “*get* passive,” as it so frequently indicates agentivity, e.g., “You’re always getting your hair done”).

Tokens that were rejected include the following:

²⁹ The reason for analyzing only tokens containing the past or present progressive (and auxiliary *be* before, not after, the PAF) is that these are the utterances discussed in the grammars. Moreover, the other types of *always/never* + progressive tokens, i.e., involving future or perfect progressive and/or auxiliary *be* after the PAF, are very rare; in the 2010-2012 data (search conducted February 15, 2015), I found only about 30 instances, all containing *always*.

³⁰ Search conducted January 27, 2015. Note: Technically, the search string used is also compatible with a version of the future progressive, but no such instances were found.

1. **Adjective:** Tokens in which the *-ing* form was an adjective.
2. **Go-future:** Tokens in which the *-ing* form was an instance of the *go*-future rather than the progressive. For example, “You’re always going home early” would be included but “You’re always going to find that...” would be rejected.
3. **Fragments:** Tokens that consisted of fragments that could be interpreted either as a participle or as a phrase with an elided subject. For example, “In and out, everything coming and going all the time. Always getting your hopes up so someone can stomp on them. That’s how it is, life” (COCA) was rejected.
4. **Non-finite clauses:** Tokens consisting of non-finite clauses, e.g., “they find themselves always having to” rather than “they are always having to”. This includes phrases coming after a preposition, e.g., “I’m talking about always trying again,” or phrases in which the *-ing* word serves as a gerund that is either the entire noun phrase or part of one, e.g., “a personality that can withstand not always being liked” (COCA).
5. **Duplicates and instant repeats:** Duplicates are instances in which the same line is repeated verbatim for a reason other than the writer/speaker truly repeating it. This happened, for instance, in magazines if the caption under a photo repeated a line already in the article. Instant repeats are instances when a speaker utters the exact or nearly exact same phrase twice in rapid succession or starts then restarts a token, repeating and finishing it. For example, “they’re always getting different - they’re always getting different angles” (COCA), was counted as one token, not two.
6. **Hypothetical statements, modals, imperatives:** To maintain a focus on statements about actual recurring actions, tokens involving hypothetical statements, modal verbs, and/or imperatives were discarded. For example, “Make sure you’re quiet with the feet and always going toward first base on the throws” (COCA), an imperative, was discarded.

5.3.3 Coding of Function

Tokens were coded as fulfilling four functions: Complain, Lament, Describe, and Praise. I also refer to these as negative (Complaint, Lament), neutral (Describe), or positive (Praise) functions, according to the actions and/or emotions associated with them.³¹ The decision to use these four functions was based partly on the literature and partly on pilot coding sessions: The grammars do not, in their characterizations of PAF/*always* + progressive, mention any of these functions explicitly, but their descriptions (all very function-oriented) and examples suggest that complaints, laments, and praise are relevant. I explain this further below.

³¹ This is not to be confused with positive, neutral, and negative semantic prosody (a focus of chapter 7), as the categories do not perfectly overlap. For example, one can complain about an action most would consider good or at least neutral, such as singing, and one can praise a person for doing evil things.

The complaint function accounts for the claims in Quirk et al. (1985, p. 199), Sinclair (1990, p. 249), and Celce-Murcia & Larsen-Freeman (1999, p. 117) that *always*/PAF + progressive can express disapproval. Indeed, many of their examples are complaints, such as “He’s forever acting up at these affairs” (ibid., p. 117) and “You’re always looking for faults” (Sinclair, 1990, p. 249). The other negative function coded for, lament, does not involve anger, annoyance, or disapproval but merely sad or unpleasant situations. For example, if one says that a deceased loved one is “never coming back,” this refers to an emotionally negative situation, but is not an indictment of the behavior of the deceased, and not likely to be an instance of complaining to the person responsible (if there even is such a person), either.

Both laments and complaints are compatible with Carter & McCarthy’s (2006, p. 47) claim that *always*/PAF + progressive refers to “regular events or states which are problematic or undesired”, as well as with the use given by Celce-Murcia & Larsen-Freeman (1999) of making an “emotional comment on [a] present habit” (p. 117). Moreover, though it is not labeled as such, one example in Celce-Murcia & Larsen-Freeman (ibid.) may be a lament. In “Orville is always hearing noises. (i.e., he hallucinates)” (p. 509), the situation is clearly negative, but Orville is an innocent victim of circumstances, so disapproval is probably not applicable.³² A clearer example of a lament appears in MICASE, in a description of Charles Darwin’s time aboard the HMS beagle. The speaker explains that Darwin “was always seasick he was just one of those people that never got over being seasick. so you know he was always throwing up and he never could really sleep well and, you can just imagine how hard it is if you're seasick”.

When the “emotional comment” made in an *always* + progressive sequence is positive, it may be an instance of praise. The Praise function is in line with Celce-Murcia & Larson-Freeman’s specification that the construction can be used to comment on a habit in a way that is “approving” rather than only disapproving (p. 117), as demonstrated by their example “He’s always delivering in a clutch situation” (ibid.) Another example, from COCA, is “She always

³² Absent further context, we cannot be certain; poor Orville’s hallucinations might be irritating to the speaker.

gives us stuff and she's always playing with us", said by a young boy heaping praise on his sister during a TV interview about her.

In sum, the Complain, Lament, and Praise functions are all derived from the literature, which tells us that PAF + progressive typically has "emotional undertones" (Celce-Murcia & Larsen-Freeman, 1999, p. 510), and can express disapproval (Complain) or sometimes approval (Praise), or that a situation is undesirable yet does not involve disapproval or blame (Lament). These three functions accounted for only part of the data, however, necessitating the addition of a fourth, more emotionally neutral function (Describe), which is "To state facts or habitual states of affairs" (exclusive of descriptions which are also complaints, laments, or praise).

To determine a token's function, I looked at not just the KWIC view but, often, the much larger context, in order to fully understand the discourse. In a few cases, I even looked at the entire texts in which tokens appeared; these could often easily be found online. At the same time, to be consistent, I ultimately focused on the X + *always* + progressive portion of the utterance. For example, "Good opportunities were always popping up, but I never took them" is a lament but would be coded as Describe because the X + *always* + progressive portion is background information for the lament, not the lament itself.

Tokens were coded as Complain if they involved a statement about something that seemed to be unpleasant or irritating to the speaker and for which the speaker appeared to blame the subject. They were coded as Lament if they were statements about a sad, unfortunate, or otherwise negative situation but one that the speaker is not complaining about or blaming anyone for; the goal in this case is expressing (or seeking) pity/sympathy, rather than expressing anger/annoyance. An example of a lament would be "For the rest of your life, you're always waiting for something bad to happen," said by a person lamenting the effects of experiencing a traumatic event early in life. Tokens were coded as Praise if they pointed out a positive trait in a person (including oneself) and the goal appeared to be to draw attention to the grammatical subject's trait/habit in order to flatter her, him, or them.

Finally, tokens were coded as Describe, the neutral function, if they did not fall into any of the previous categories involving strong negative or positive emotions. An example from MICASE (in reference to researchers in a particular field) is “We’re always dealing with the dipoles.” Tokens could be coded as Describe even if they involved positive things. For example, even though the discovery of new things is good, “Scientists are always discovering new things, yet this remains elusive” would be coded as Describe because in this context the statement serves merely as background information for the next clause.

5.3.4 Coding of Subject Type

Based on qualitative analysis of the immediate or sometimes extended context, the *always* + progressive tokens were divided into four categories corresponding to their grammatical subjects. This process was generally very straightforward. The four categories are defined below, the first three consisting of human subjects and the fourth non-human subjects.

1. **Human, first person pronoun (1P):** This category consisted of the first person pronouns *I* and *we* in reference to human beings.
2. **Human, second person pronoun (2P):** This category consisted of the second person pronoun *you* (both singular or plural, e.g., *you all*, though the latter was very rare) in reference to human beings. The category also included instances of generic *you*,³³ which was fairly uncommon.
3. **Human, third person pronoun or other (henceforth “Third party”) (3P/O):** This category consisted of the third person pronouns *he*, *she*, and *they* (3P) when they referred to human beings, and of other (O) subjects with third party humans as their referent. This includes, for example, proper names, kinship terms, certain noun phrases, *one*, anything that begins a relative clause (such as *who*, *that*), and conjoined subjects (even if they contained 1P or 2P pronouns, e.g., “me and her” or “you and John”). Instances of non-human noun phrases (e.g., *the left* to refer to the political left, or *the media*) were not counted as human. The test used in these cases was whether or not one could answer “yes” to the question “is/are X(s) human beings?” Thus, while *left wing voters* are people, *the left* is not, nor is *the media*, etc.

³³ Another option would have been to omit instances of generic *you*. If so, however, it would also make sense to omit instances of what I would call “societal *we*” (see §5.4.1) and similar subjects, such as *people* in statements about human nature or society as a whole. Since the question of whether *we* or *people* is being used in such a way can be difficult to answer, I did not make such a distinction, either for words like *people* and *we*, or for generic *you*.

4. **Not human (NH):** All non-human subjects were included in this category, even if the non-human subject was referred to using a personal pronoun. Special situations involving human-like entities (zombies, aliens, sentient robots) were categorized as non-human, but God (four instances) and Jesus (two instances) were considered human.³⁴

5.3.5 Division into Genres

To study genre effects, I divided the *always* tokens into three broad genres, with some tokens appearing in more than one genre. These genres, in addition to All (all collected tokens), are Written (all 2010-2012 tokens from COCA's written portion), Spoken (all tokens from spoken data, i.e., all 2010-2012 tokens from COCA's spoken portion combined with all MICASE, Swb, and SBC tokens), and Academic Spoken (the MICASE tokens). Casual spoken language (Swb and SBC) was not analyzed independently because it contributed only two tokens.

5.3.6 Data Analysis

Given the multiple corpora and categories, four subject types, and four functions, numerous comparisons were possible. At the same time, due to great variability in the number of tokens from each corpus and category and for each PAF, not all analyses could reveal meaningful differences. In studying *always*, I sought correlations between function and subject type (with all genres combined) and between function and genre (with all subject types combined). However, in studying the less numerous *never* tokens, I analyzed only function (with all tokens combined). Additionally, I determined (a) the frequency of the *go*-future in discarded *always* and *never* tokens, and the percent of non-discarded *always* and *never* tokens that carried a future meaning; (b) how often a fifth function, Vow, appeared in the *never* data; and (c) how often certain verb phrases appeared in both data sets.

³⁴ The reasoning behind this is that a great many Christians (the God referred to in the four tokens was the Christian God) are officially of the view that Jesus existed, historically, and was 100% human (while also 100% a god) and is fully the same as God the Father, and that Jesus/God is as involved in their lives as any friend or family member.

5.4 Results

The number of *always* + progressive tokens obtained after the described collection methods and quality control checks was 764, with COCA contributing 709 (538 written, 171 spoken), MICASE 43, Swb ten, and SBC two. In terms of genre, the method resulted in 538 Written tokens (all from COCA), 226 Spoken tokens (171 from COCA, 43 from MICASE, ten from Swb and two from SBC), and 43 Academic Spoken tokens (all from MICASE). The number of *never* tokens was 53, with COCA contributing 51 (13 spoken, 38 written), MICASE two, and Swb and SBC zero. The total number of *always* and *never* tokens combined was 817.

This section is organized as follows: First, I discuss the overall distribution of function (neutral functions are the most common, negative functions second-most common, and the positive function least common). I also describe the observed sub-types of the neutral function. Second and third, I discuss the correlations between function and grammatical subject, and function and genre. Fourth, I discuss distinguishing features of the *never* data.

5.4.1 Distribution of Functions

The most general finding is that, contrary to expectations based on the literature, negative functions are not especially common with *always* + progressive, nor with *never* + progressive. A little over 70% of the 764 *always* tokens were coded as Describe, whereas only about 25% were coded as Complain or Lament. (See Fig. 5.1, which indicates the frequency of each of the four functions, given as a percentage of the entire set of tokens. Note that Fig. 5.1, as well as Figs. 5.2, 5.3, and 5.4, indicate the percentage of tokens with the negative functions Complain and Lament both combined and separately.) The *never* data is discussed later, in §5.4.4, but resembles the *always* data in that the Describe function is most common.

If we consider only tokens with either a positive or negative valence, however, then we can see the truth in the claims made in the literature: Complain and Lament account for six times the number of *always* tokens that Praise does (24.35% versus 4.06% of the tokens).

(Likewise, in the *never* data, Complain and Lament account for 4.5 times as many tokens as Praise does, or 16.98% versus 3.77%). Moreover, Complain, alone, accounts for a little over four times the *always* tokens that Praise does. Thus, depending on one’s focus (positive, negative, and neutral functions, or only positive versus negative functions), the results may or may not merit associating *always/never* + progressive with negative functions.

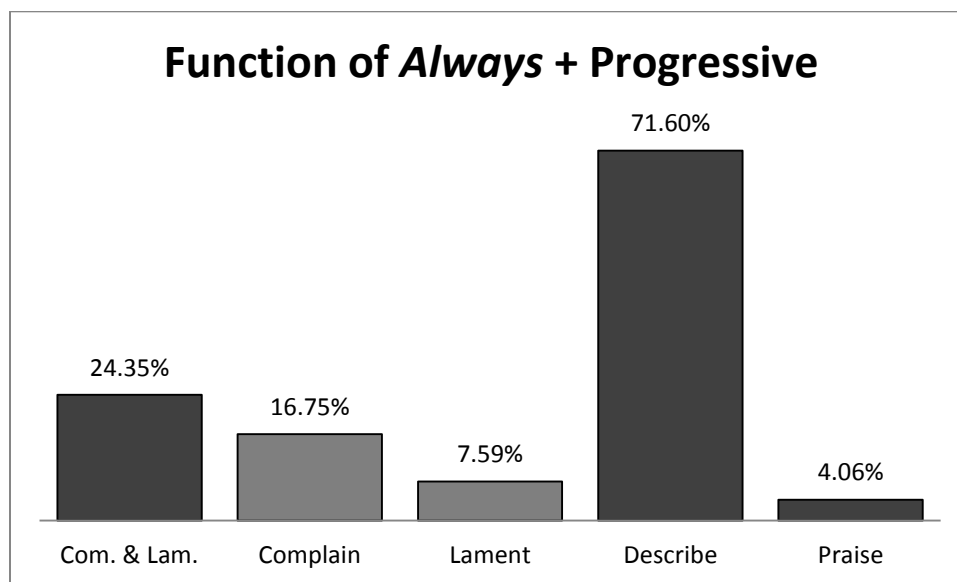


Figure 5.1: Function of *always* + progressive

The Describe function was defined as “To state facts or habitual states of affairs” (exclusive of descriptions which are also complaints, laments, or praise). Examples found in the data include “Climates are always changing” (COCA) (a fact about the planet) or “I’m always trying to catch little critters and show the kids” (COCA) (a characterization of oneself, and which served as background information to a larger story about a salamander).

Qualitative analysis revealed that Describe was not only the most common function but that—in addition to being used in the very general way demonstrated by the two examples above—it serves at least four specific sub-functions: (1) Explain how things work; (2) State facts about human beings, in general; (3) State facts about members of a particular group, especially a professional group; and (4) Describe/analyze, objectively, the states or behavior typical of

characters in novels. I expand on each of these below.

Explain how things work: Explanations of how things work were particularly common in the academic spoken language (MICASE). For example, four of the 43 MICASE tokens appear in the context of explaining how genes and evolution work (e.g., “the trip genes are always being ... transcribed”). Another example of this sub-function, also from MICASE, and part of an explanation of how to use semicolons correctly, is “this has to be lower case. it can't be [uppercase] because, that's always signalling the beginning of a new sentence”.

State truths/facts about human nature: Examples of this sub-function include “we're always looking out, to define ourselves with [a] sense of responsibility” (MICASE) and “We're always trying to find [the silver lining]” (COCA). Tokens used for this sub-function often involved what one could call “societal *we*” (i.e., *we* used to refer to human beings or society in general) or “societal *people*” (in contrast to *people* in an example such as “people are always coming up to me saying ...” (COCA), which is not about human nature but, rather, about a recurring personal experience involving various specific people).

State facts about members of a particular group, especially a professional group: An example of this is “We're always dealing with the dipoles”, where *we* refers to professionals and students in the speaker's field (MICASE).³⁵ Generally, the speaker is discussing a group she/he is a member of. Not just *we* appeared in several such tokens, but also generic *you* (or what at first looks like generic *you*). Examples include “You're always looking at markets and trying to get an edge. It's like a drug” (COCA), said about working in the world of high stakes international finance, and “You are always working on multiple projects with so many talented people, which means you're constantly challenged and inspired” (COCA), said of the fashion industry. These *you*'s are generic in that they refer to more people than just the speaker, but

³⁵ Another example including *deal with*, involving the present simple and *often*, was noted as well: “... the guys down the corridor who do deal with vascular plants ... while in your case you're often dealing with microns rather than millimeters” (MICASE). In fact, *deal with* appears twice in this example, to contrast what the in-group studies with what another group studies. In a larger data set, *deal with* might appear quite often in descriptions of what people in particular professions do on a daily basis.

specific in that the referents are limited to people in a particular field/industry.

Describe/analyze, objectively, the states or behavior typical of characters in novels:

Like the first sub-function described above, this one is particularly suited to academic contexts, and was prevalent in MICASE. Examples are “Tita occupies this, interesting middle position. between the progressive and the conservative. and she's always trying to do both” and “Newman who's always been a ... go-getter, a doer, you know, if there's a problem, you grab it and you solve it this is what he's always saying to Valentin” (both from MICASE).

Unlike Describe, the negative functions Complain and Lament did not seem divisible into sub-categories. Examples of complaints include “friends [of the teenager next door] are always honking their horns when they pull into her driveway to pick her up. I can't stand the noise! Whatever [sic] happened to parking a car and ringing a doorbell?” and “She's always asking for some money out of me” (both from COCA); the latter was about a woman who later launched “an all out assault on [the speaker's] character and even his manhood”. Examples of laments include “He's always putting things in his mouth (said of a child who likely suffered from a rare eating disorder, and who eventually died because of it) and “she was always thinking of her missing child, Elizabeth. her kidnapping had stopped the clock on life” (both from COCA). Laments typically involved sadness rather than anger.

No sub-categories of Praise were noted either. Examples of praise include “You're always taking care of others” and “He has a very big heart and he's going to try to help people who ask him to do things. I've seen him at basketball games helping an old lady down from the stands. He's always putting his hand out to help people” (both from COCA).

5.4.2 Grammatical Subject

In order to see if correlations existed between subject type and function, I coded subjects as human first person pronouns (1P), human second person pronouns (2P), human third party (3P/O, encompassing pronouns and also names, noun phrases, and other references to third

party humans), or non-human (NH). Overall, the four subjects behaved more similarly than differently (see Fig. 5.2). The 1P, 2P, and NH subjects exhibit the same pattern as found in the entire data set (with all subjects combined, as in Fig. 5.1): Describe is most common, followed by Complain, then Lament, then Praise. 3P/O is the exception, but only in that the rankings of Lament and Praise are reversed. Still, the variation between subject types is revealing.

Function of *Always* + Progressive by Subject

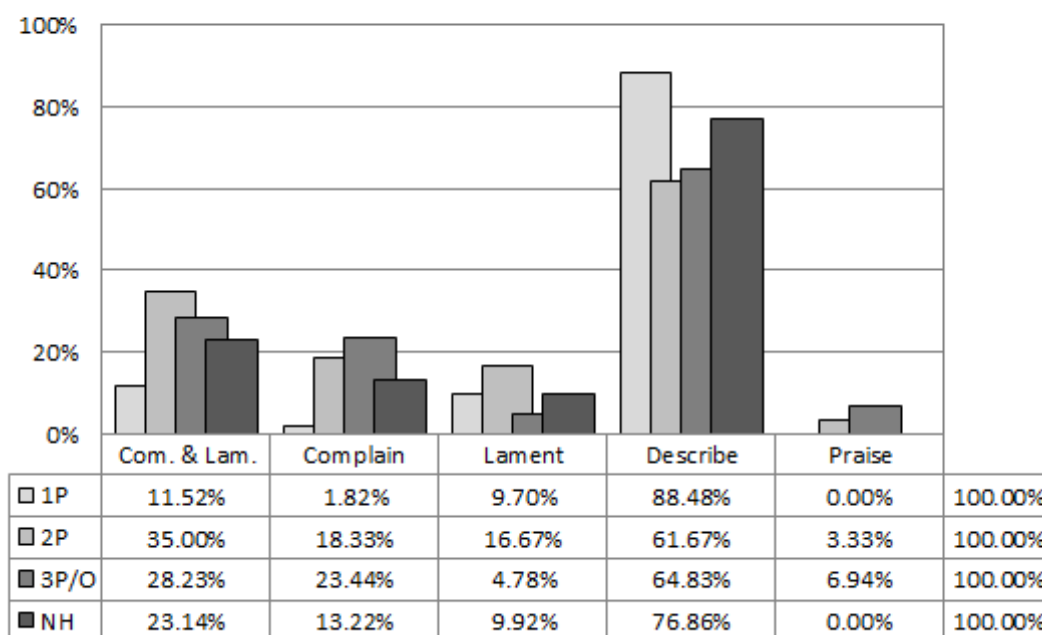


Figure 5.2: Function of *always* + progressive by subject

Namely, of the four subject types, 1Ps are the most likely, by far, to be described neutrally (88.43% of the 1P subject tokens were coded as Describe), and least likely to be complained about (only 1.82%), while 3P/Os were the most likely to be complained about (23.44%) and second least likely to be described neutrally (64.83%). In general, 1Ps stand out from the other subjects, while 2Ps and 3P/Os are more similar; they are nearly equally likely to be described neutrally (61.47% and 64.83% of 2P and 3P/O tokens, respectively), and second-most and most likely to be praised (6.94% and 3.33%, respectively), while NHs and 1Ps are never praised.

5.4.3 Genre

Genre effects were not strong, even between the Written and Spoken language (see Fig. 5.3). The percentages of written and spoken tokens coded as Describe were nearly identical (71.56% and 71.68%, respectively), and the percentages coded as Praise (4.28% and 3.54%) were quite similar too. The noteworthy exception to this pattern, though, is Academic Spoken language (MICASE), in which Describe was more common than in the other genres (by a difference of just over 12%), and Complain and Lament less common. Additionally, it contained no instances of Praise (which accounted for at least a small portion of each of the other genres). Overall, the *always* + progressive phrases in MICASE were mostly neutral.

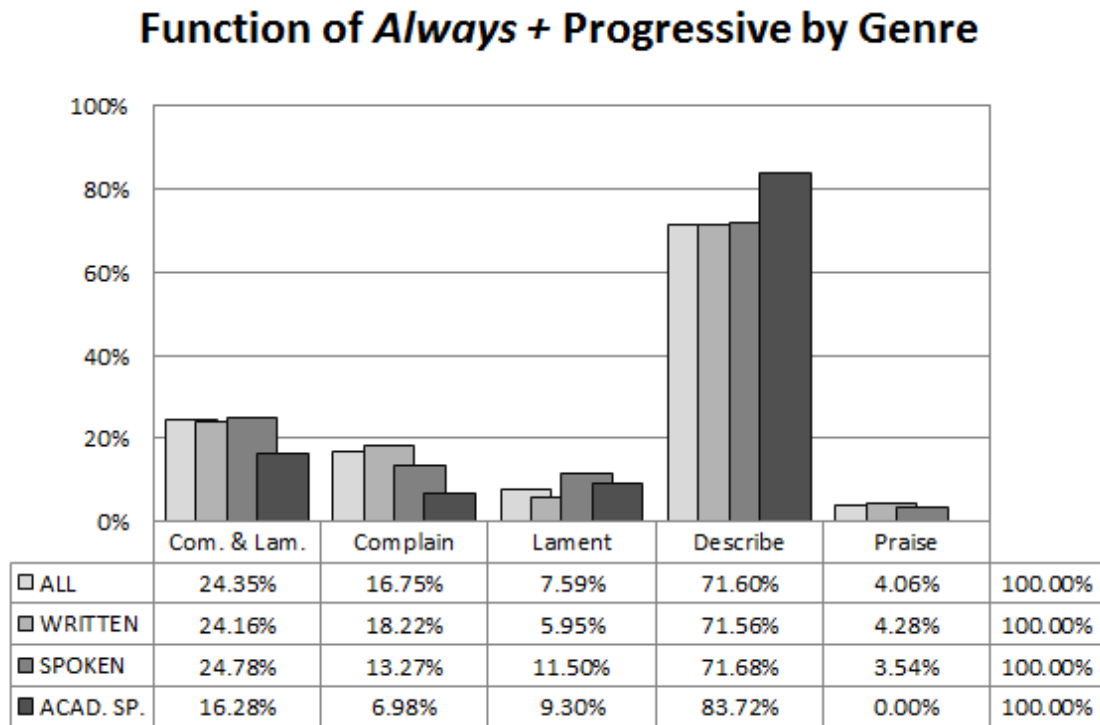


Figure 5.3: Function of *always* + progressive by genre

5.4.4 Differences Between *Always* and *Never*

The key findings which differentiate the *never* data from the *always* data are: (a) over three quarters of its verbs have a future meaning (something which does not occur in the *always* data)

and, in the raw data, the *go*-future is extremely common; (b) a fifth function, Vow, appears only in the *never* data and accounts for nearly a third of the tokens; and (c) the verbs *come* or *go* appear in over half the *never* tokens. Additionally, the *never* data has fewer Describe tokens than the *always* data, and an inverse distribution of Complain and Lament tokens.

Future Meaning: Of the 53 *never* X-ing tokens which made it through the quality control process, during which instances of the *go*-future were removed (along with fragments, gerunds, tenseless clauses, hypotheticals, modals, imperatives, duplicates, and repeats), 75.47% were found to carry a future meaning nevertheless, e.g., “I’m never getting another perm” and “I packed up a bag and decided I was never coming back” (both from COCA). This can be compared to the non-future meaning of the progressive in “It’s never taking into account ... how Haiti works” (also from COCA). In contrast, no tokens in the finalized set of *always* tokens carried a future meaning (even though this is possible, e.g., “This ice cream is great! I’m always ordering this flavor, from now on”). A second indicator of the prevalence of future meaning among *never* tokens involves those tokens which did not pass the quality control check: For every *never* + progressive token retained and analyzed, 6.83 were discarded for involving the *go*-future, while this same figure for *always* is only 0.26.

Vow function and other differences: Another distinguishing feature of the *never* data is its numerous instances of a fifth function, Vow, whereby the speaker³⁶ makes a vow or promise, e.g., “I hate weddings, and I’m never getting married” (COCA). This accounted for just under a third (32.08%) of the *never* tokens (see Fig. 5.4). The frequency of Vows appears to largely account for the smaller percentage of Describe tokens in the *never* data (47.17%, versus 71.60% of the *always* tokens). Finally, *always* and *never* are opposites with respect to Complain and Lament: In the *always* data, complaints are 2.21 times as common as laments while, in the *never* data, laments are 3.50 times as common as complaints (see Fig. 5.4).

³⁶ Vows can also be uttered by someone other than the person making them (though this is rare), as in third person omniscient narration (“So that’s what she wanted—to force him into a confession of his sins. Well, she was never getting that from him” (COCA)) or indirect quotations (“when he left he had said he was never coming back” (COCA)).

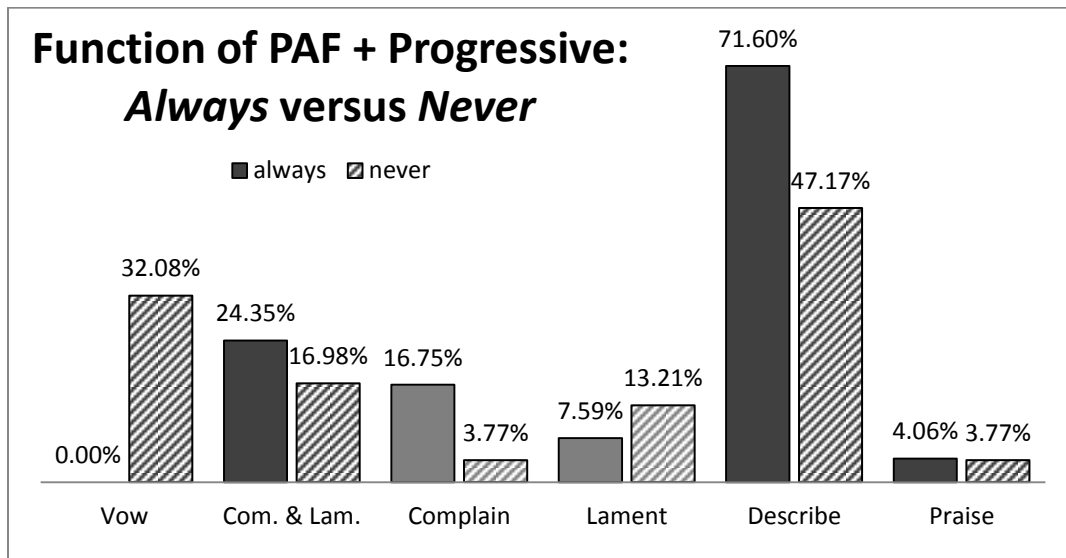


Figure 5.4: Function of PAF + progressive, *always* versus *never*

Common verbs: The *never* data was also distinguished from the *always* data by the extremely frequent appearance of *come* and *go*, with one or the other appearing as the verb in 54.72% (29 out of 53) of the *never* tokens, but only 3.40% (26 out of 764) of the *always* tokens (see Fig. 5.5). Moreover, in 26 of the 29 *never coming/going* tokens, the verbs were followed by *back* (22 tokens) or *home* (four tokens), versus only two of the verbs in the 26 *always coming/going* tokens (both of which contained *back*).

Percent of PAF + Progressive Tokens Containing <i>Come/Go (Back/Again)</i>		
	<i>always</i> (n=764)	<i>never</i> (n=53)
<i>coming/going</i> (all)	3.40% (n=26)	54.72% (n=29)
<i>coming/going back/home</i>	0.26% (n=2)	49.06% (n=26)

Figure 5.5: Percent of PAF + progressive tokens containing *come/go (back/again)*

In addition, nearly half (12/26) of the *never coming/going back/home* tokens were vows or laments, e.g., the vow “I’ll drown if I need to, but I’m never coming back” (said by a person expressing solidarity for an area succumbing to a rising sea level), and the laments “The man was never coming home” (in reference to the speaker’s father, who the speaker wished were not

in prison) and “Our sister is never coming back” (said of someone who had died) (all from COCA). Laments occurred with *always* as well (e.g., “he was always throwing up” (MICASE), in reference to Darwin’s miserable time aboard the H.M.S. Beagle), but only a little over half as often as with *never*, and they generally depicted situations less dire and/or less permanent.

5.5 Discussion

This five-part section addresses (1) the overall distribution of functions; (2) the findings relating function and grammatical subject; (3) the findings relating function and genre; (4) the differences between *always* and *never*, including in terms of tense-aspect; and, (5) the need to be cautious in making generalizations across words, word classes, or contexts.

5.5.1 Distribution of Functions

In this section, I first discuss the finding that the functions of *always* + progressive tokens were largely neutral, and then the finding that negative functions outnumbered positive functions. As will be discussed shortly, it is true that the construction is well-suited to convey complaints. Neutral functions are the most prevalent nevertheless, though, because the construction is equally well-suited to serve these multiple neutral functions which figure prominently in our lives. And part of its ability to do this stems from the flexible nature of *always* itself.

One basic traditional lexical distinction is that between content and function words. Content words, nouns being a prototypical example, “carry the primary communicative force of an utterance, are open or productive classes, and are variable in form (inflected)” while function words, such as prepositions, “express grammatical meaning (by relating sentence parts) ... and are generally invariable in form” (Brinton, 2000, p. 118). It is surprisingly difficult to find, in formally published work, a clear statement about which category adverbs of frequency fall into,

but they seem to be function words.³⁷ They are invariable in form and, like prepositions and other function words, they are dependent on other words (they do have meaning on their own but, being adverbs, they require a verb), and can be used in an endless variety of contexts, with any verb, any subject, etc. This, alone, makes it unlikely that any adverb of frequency would be restricted to mainly negative (or mainly positive) uses.

At the same time, the specific semantic content *always* does possess allows *always* + progressive to transcend the uses and meanings typically associated with the progressive (which can also indicate, as *always* can, recurring and thus habitual actions, but whose core meaning is that action is in progress) and extend into territory more commonly associated with the present simple. That is, much like *always* + present simple, *always* + progressive can indicate, without conveying one's affective stance, simply that something is/was (permanently, habitually, or at least usually) true. This is somewhat explicitly addressed in Quirk et al. (1985, p. 199), who write about the "habitual progressive".³⁸ While this usually refers to temporary habits, "in combination with *always*, *continually*, or *forever*, the progressive loses its semantic component of 'temporariness'" and thus can be used to refer to long-term or permanent habits (ibid.). Because of this, *always* paired with the progressive can act like the present simple and convey "general timeless truths, such as physical laws or customs" (Celce-Murcia & Larsen-Freeman, 1999, p. 113, on the present simple; cf. Sinclair, 1990, p. 247, on the present simple expressing "general truths" and "regular or habitual actions"). Thus, the definition I give here for the Describe function, "To state facts or habitual states of affairs", is quite similar to others'

³⁷ Many (e.g., Denham & Lobeck, 2012, p. 146; Shi, Werker, & Cutler, 2006, p. 188) say adverbs are content words. Others are clear that only some adverbs are (Brinton, 2000, p. 118). Still others write about function or content words without defining them (e.g., Li, Zhang, Luo, & Wu, 2014; "frequency adverbs" does appear in one of their diagrams (p. 1889, Fig. 2) of the "taxonomy of function words", though). Online resources are more straightforward, e.g., FLESL.net (flesl.net/Grammar/Grammar_Glossary/adverb.php, accessed May 14, 2015) explicitly says that adverbs of frequency are function words, and Pronuncian (pronuncian.com/Lessons/Default.aspx?Lesson=58, accessed May 14, 2015) says they are not content words. More often, though, in explanations of content versus function words in both published work and online resources, adverbs of frequency are simply not mentioned.

³⁸ Cf. "The professor is typing his own letters while his secretary is ill. [The habit is temporary]" to "The professor types his own letters. [The habit is permanent]" (Quirk et al., 1985, p. 199). Alternatively, it can imply "that every event in a sequence of events has duration/incompletion" as in "Whenever I see her, she's working in the garden".

descriptions of the present simple (and habitual progressive).

The four specific sub-functions of Describe are, likewise, closely related to uses typical not of the progressive but of the present simple: The first sub-function, *explain how things work*, is an instance of a “general truth” (which includes, sometimes, “physical laws”). The second, to *state truths/facts about human nature* is, likewise, simply another, this time human-focused, instance of a “general truth”. The third, *state facts about members of a particular group, especially a professional group*, is most closely related to truths involving “customs.” The fourth, *describe/analyze, objectively, the states or behavior typical of characters in novels*, may be a special case of the second. By dint of being in novels or movies rather than real life, fictional characters are incapable of doing anything other than what they always have done (and always “will do,” in future readings/viewings), so facts about those characters’ actions, dispositions, etc., are also timeless truths about (particular) human beings.

Technically, the positive and negative functions (Complain, Lament, and Praise), as well, are sub-functions of Describe, as is Vow. Depending on whether the subject is human or not; on whether the situation (if it is negative) involves sympathy, empathy, or blame; on whether the speaker is describing a process; etc., descriptions can function as complaints, laments, or praise. Likewise, a vow is just a statement about what will be true in the future (understood to be an assertion of one’s intent to ensure this). Far more often, though, *always* + progressive statements that describe things are simply neutral.

It is crucial to note here that “neutral” does not mean “unimportant.” In writing and in academic spoken English, in particular, but also in our daily lives, description is of the utmost importance: We narrate events, describing main actions and background actions; we share facts about the world around us and the things and people in it, and explain how those things/people or the world itself works; and we share all kinds of other information as well, both objective and subjective, and about things both real and imagined. Even instances of description offered for entirely phatic purposes (such as one’s response to “how was your day?”) are important and, in

fact, are some of the key activities humans use language for. It is not surprising, then, that so many of the tokens were coded as serving such purposes.

This argument is in line with a functionalist, “language as action” approach. As explained in §§2.1 and 5.2, “language as action” is defined as a focus on “what, and how, people *do* things with language” (Neville & Rendle-Short, 2007, p. 30.1). In other words, it calls for us to focus on the language user, as a person, who is:

a member of a particular linguistic community who, in speaking (and, indeed, in signing or writing), attempts to achieve a particular interactional goal or set of goals using particular linguistic and non-linguistic strategies. Interactional goals include attempts to elicit information or action on the part of the hearer, to provide information, to establish interpersonal rapport (e.g. when ‘passing the time of day’) and so on. The linguistic strategies used to achieve these goals might include the use of speech acts (requesting, informing, promising, thanking and so on), choices over words and grammatical constructions, intonation structures, choices over conforming or not conforming to discourse conventions like turn-taking and so on. (Evans & Green, 2006, p. 110)

In sum, one quintessential purpose of language is to do things (or get things done), and this means that one should be able to use the assumption that people need to and often do succeed at this (i.e., we use language effectively) to guide us in analyzing corpus findings.

Viewing language as action helps us make sense of the predominance of neutral uses of *always* + progressive functions (as the construction works well to perform numerous different important neutral functions that we commonly need). It can also help us understand the other major aspect of the pattern: Instances of negative functions are more common than instances of positive functions. Recall that Complain and Lament accounted for 24.35% and 16.98% of the *always* and *never* tokens, respectively, while Praise accounted for only 4.06% and 3.77%, respectively. I spend the rest of this section introducing a functionalist explanation for this pattern, an explanation I expand on in the following two sections.

One thing we use language to do—and which calls for the very opposite of emotionally neutral statements—is to “express our appraisal of and attitude towards whoever we are addressing and what we are talking about” (Halliday, 1985, p. 29). While we can certainly

appraise people and things positively (as in praise) in addition to appraising them negatively (as in complaints) or seeking sympathy/empathy (as in laments), from a practical standpoint, if the status quo suits us, we are not very motivated to speak up. In fact, doing so is a waste of energy. In contrast, when things are not going well and we wish to bring about change, we may attempt to do so through complaints, etc. As Baumeister, Bratslavsky, Finkenauer, & Vohs (2001) put it, “bad events signal a need for change, whereas good ones do not” (p. 357).

This argument is closely related to the negativity bias, a highly generalized cognitive bias the existence of which has been demonstrated empirically by, e.g., Peeters & Czapinski (1990), Taylor (1991), Cacioppo, Gardner, & Berntson (1999), and others; see Jing-Schmidt (2007) for a fuller review. This bias is defined as “an automatic tendency to pay significantly more attention to unpleasant than pleasant information,” as indicated by, for example, people’s behavior after exposure to negative versus positive events (p. 418), or patterns of event-related brain potentials (“depictions of the electrical activity on the scalp that results from the neutral processing of a given stimulus”) (Smith, Cacioppo, Larsen, & Chartrand, 2003, p. 172). Closely related to or even inherent in the negativity bias is “loss aversion” (Tversky & Kahneman, 1991), a well-studied phenomenon whereby “losses loom larger than corresponding gains” (p. 1039). If the aversion to loss outweighs the motivation to seek and pursue gains of equivalent strength, this too would result in us being more concerned with potentially bad things than with potentially good things, which might be reflected in a greater likelihood to discuss such things.

A posited reason for the negativity bias, pervasive across cultures and history (Rozin & Royzman, 2001, p. 296), is that it is evolutionarily adaptive (Pratto & Oliver, 1991; Baumeister et al., 2001). It is useful to seek good things but we need, first and foremost, to avoid harm, especially fatal harm. This means “individuals who are attuned to preventing and rectifying bad things should flourish and thrive more than individuals oriented primarily toward maximizing good things” (ibid., p. 357). In other words, in terms of staying alive, there is “little incentive to continue seeking further benefits or advances” if we are currently satisfied (ibid.).

The effect on language of the negativity bias and/or loss aversion on what we attend and react to is complicated. Most literature on the negativity bias also addresses the Pollyanna effect, the “universal human tendency to use evaluatively positive (E+) words more frequently and diversely than evaluatively negative words” (Boucher & Osgood, 1969, p. 1; see also Zajonc, 1968; Matlin & Stang, 1978). However, and as many have pointed out, the negativity bias is not disproven by positive words and, in fact, even predicts them. The negativity bias is an inherited cognitive disposition based on survival, but there also are social/cultural, i.e., learned, reasons to come off as more positive, polite, and pleasant (Jing-Schmidt, 2007, p. 423-424)—concerns expanded on in the next section. In other words, avoiding risks related to negative words is perfectly in line with being vigilant about harmful things.

In addition, a greater frequency of positive words in language could simply reflect “the same basic fact about the world, the dominance of positive experiences” (Rozin & Royzman, 2001, p. 297). Despite our being constantly on the lookout for bad things, they are actually rare. And if “a positive state of affairs is seen as the default case,” then “(events triggering) negative emotions are seen as more divergent from the norm, more marked, and thus more worthy of comment” (Claridge, 2011, p. 82)—even if we suppress such comments for social reasons.

There are two other important points to consider regarding the Pollyanna effect which show that it does not contradict the negativity bias. First, negative words and negative situations or actions do not always overlap. For example, the complaint “You’re always singing that same song, all day long” expresses anger and annoyance yet contains no negative words. In general, negative emotions, more than positive emotions, are expressed more indirectly (K. J. Anderson & Leaper, 1998, p. 439). Second, not all the functions coded here as negative involve socially harmful acts: Laments can be a way of offering and/or eliciting sympathy, and of rapport-building, and complaints about people not around to hear them are fairly “safe” as well. Thus, laments and complaints about non-present people are not things even a highly functional interpretation of the Pollyanna effect would necessarily predict that we avoid.

Also worth noting is that the literature on semantic prosody contradicts the Pollyanna effect literature. Semantic prosody (SP) is “a form of meaning which is established through the proximity of a consistent series of collocates, often characterisable as positive or negative, and whose primary function is the expression of the attitude of its speaker or writer towards some pragmatic situation” (Louw, 2000, p. 57). For example, “break out” (intransitive) is said to have negative SP because it is so often followed by noun phrases referring to negative things such as disorder and epidemics (Sinclair, 1990, xi). As it turns out, negative SP is much more common than positive (Louw, 2000, p. 52)—a pattern which coheres with the finding of K. J. Anderson & Leaper (1998, p. 439) that negative emotions are expressed (however indirectly) more often than positive ones. As Louw puts it, in studies of semantic prosody, people “emerge just as we might have expected: selfish, comfort-loving, nea-phobic [sic] complainers” (2000, p. 65).

Putting a more understanding spin on the situation, Louw also observes—much like Baumeister et al. (2001) regarding the negativity bias—that “in the same way that unrequited love forms most of the subject matter for the greatest love poetry ... contented human beings utter much less than discontented ones” (p. 52). Regardless of whether our individual words, themselves, are more often negative (as SP studies suggest) or positive (as the Pollyanna effect suggests), discontent motivates us to take action.

5.5.2 Grammatical subject

The relationship between grammatical subject and function in the *always* tokens was as follows: First, with one minor exception, the four subjects follow the same pattern, whereby Describe is most common, then Complain, then Lament, then Praise. Second, 1Ps are most likely to be described neutrally and least likely to be complained about, while 3P/Os are most likely to be complained about and second-least likely (second to 2P) to be described neutrally. Third, 2Ps are most likely to appear in laments, while 3P/Os are least likely, and, fourth, 3P/Os are most likely to be praised, and 2Ps second-most likely, while NHs and 1Ps are never praised.

One way of summarizing this data is to say that people say emotionally non-neutral things about other human beings more often than about themselves or non-human objects. This is consonant with the “language as action” interpretation that speakers do this to bring about change (cessation of the irritating action, and/or sanctioning of the actor) or to accomplish other important social functions (such as expressing pity or approval). If humans are greatly concerned with getting things done, we should not expect them to spend an excessive amount of time complaining about or praising themselves; they are already intimately familiar with their own nature, habits, and opinions, so there is less motivation to share these. Instead, the negative and positive functions are more about other people, who can hear (or hear about) our complaints, and maybe even react to them, and who can sympathize/empathize with us.

Especially interesting to analyze from a functional perspective are the results regarding the Complain function because, in them, we see evidence of the interaction of multiple social and practical concerns. We complain more about other people than ourselves or non-human subjects, but we do not complain the most to the person irritating us; complaints were more frequent among 3P/O tokens (23.44%) than 2P tokens (18.33%). This might seem impractical, since the guilty party might never learn of our feelings, and since the people we complain to more often probably cannot help. However, the tendency to complain about 3P/Os more than 2Ps is precisely what should be expected in terms of social awareness (and is also possibly more effective, if not in terms of changing the complaint-worthy situation then at least in other important ways). Three especially relevant concepts here are face, empathy, and gossip.

The concept of face is integral to Politeness Theory (P. Brown & Levinson, 1978, 1987, etc.). Positive face is “the positive consistent self-image or ‘personality’ (crucially including the desire that this self-image be appreciated and approved of) claimed by interactants” (P. Brown & Levinson, 1987, p. 61), while negative face is, essentially, “freedom of action and freedom from imposition” (ibid.). Complaining directly to someone about his/her habitual actions is a threat to both types of face: It is a threat to positive face because it is criticism and a threat to negative

face because the implication is that one should change one's behavior to suit the speaker. This effect is exacerbated in the case of *always* + progressive complaints due to their typically involving exaggeration as well. Exaggeration “emphasize[s] the problem or failings of the addressee” (or deliberately distorts them, actually), and thus constitutes a bigger “personal challenge” and is more upsetting (Legitt & Gibbs, 2000, p. 7).

In other words, it is rude to complain about someone to his or her face, even if that is the most direct and effective way of getting that person to realize that his or her actions are annoying, so—as anyone who has ever seen or written an anonymous passive-aggressive note knows—we engage in all manners of avoiding such confrontations. This is, in fact, an instance of the negativity bias in action. We are highly social beings, so the pressure to behave in socially acceptable ways is quite powerful (Baumeister et al., 2001, p. 361), and we see socially awkward and/or harmful situations as one of those things to be sensitive to and avoid. Moreover, being such socially-sensitive creatures, we likely recognize that complaints containing 2Ps are not necessarily effective anyway. The affront to face could cancel out any willingness, on the guilty party's part, to cooperate, making the situation lose-lose: The social relationship is damaged and the irritating behavior has not ceased.

This is not to say, though, that direct criticism is never effective. In fact, a functionalist would have to assume that it quite often is, given that 2Ps are the second-most likely subject type to be complained about. Somewhat paradoxically, the fact that such complaints are face-threatening is precisely what can (potentially) make them effective. And, as stated earlier, they are especially upsetting and effective if they involve exaggeration (Legitt & Gibbs, 2000, p. 7)—as *always* statements often seem to. In fact, rude complaints involving overstatement elicited more negative feelings in people even than rude complaints involving sarcasm (ibid.). This unpleasantness and heightened “emotional impact” can make complaints “very memorable,” with “long-ranging effect[s]” (Claridge, 2011, p. 142). In other words, complaining directly to someone's face can certainly work, and work very well, even; it is just that one obtains those

benefits (which are not even guaranteed) at a social cost.

In contrast, third parties are fairly safe to criticize. One could argue that such complaints are unlikely to evoke change, but, actually, this is not always the case (as is discussed below). Moreover, criticism of third parties can be rewarding even without evoking change. Conveying opinions and emotions, along with sharing information about other people, is as important in human interaction as getting things done—or is, in fact, one of the things that needs to be done. Particularly relevant here are empathy and gossip.

Sharing one's problems or concerns (a situation relevant to complaints and also laments) with a trusted interlocutor can serve an extremely important social or phatic purpose, and even benefit our physical bodies. It has been known for decades that “social ties and social support are positively and causally related to mental health, physical health, and longevity” (Thoits, 2011, p. 154; see also Cohen & Janicki-Deverts, 2009, p. 375). In particular, “ventilation” (“venting”, in layman's terms, about our problems), and the empathy and validation that one may receive for doing this, can “reduce the distressed person's psychological and affective arousal directly” and even “restore the person's sense of self-worth” (Thoits, 2011, p. 154; see also Thoits, 1986; and Coates & Winston, 1983). Thus, complaints and laments are motivated by more than just the desire to change the situation—as is evident from the simple fact that we commonly carry on about things that we know are impossible to change.

Another reason for complaining is that one might simply wish to engage in gossip, an activity understood to be highly socially valuable, and integral to human interaction. Gossip is pervasive in all human societies, accounting for as much as two thirds or more of daily talk (Emler, 1994; Dunbar, Duncan, & Marriott, 1997; Dunbar 2004), and has been shown to serve important social functions benefitting both groups and individuals (Kniffin & Wilson, 2010; Dunbar, 2004; Beersma & Van Kleef, 2012, etc.) such as bonding, sharing of crucial social information, “group protection” and sheer “social enjoyment” (ibid., p. 2640). That is, it is even pleasurable in its own right. And negative gossip, in particular (which could include complaints),

may be especially beneficial. In one study (Grosser, Lopez-Kidwell, & Labianca, 2010), people who engaged in both negative and positive gossip (rather than one or the other, or neither) were viewed by co-workers as exercising more informal influence in the workplace (ibid.), and negative gossip was more likely to occur between co-workers who were also personal friends than it was between co-workers with only a professional relationship (ibid.)

Many researchers, and particularly Dunbar, who famously relates gossip to grooming in primates (1993, 1996, 2004), study gossip from an evolutionary perspective. Gossip has likely served important functions in human social life since the earliest days of language, and that is unlikely to change. In sum, for reasons related to both empathy and gossip, complaints (and laments) are useful even if they do not stop the problematic situation.

At the same time, sometimes complaints about third parties can do exactly that, in addition to bringing about the benefits described above. As shown by Grosser, Lopez-Kidwell, & Labianca (2010), mentioned above, gossip is related to power and influence. This applies to not just the gossiper but also the person gossiped about: Negative gossip often makes its way back to and has repercussions for that person, and the ensuing shame or damage to reputation could lead him/her to cease doing what evoked the gossip. This exemplifies the “policing function” of gossip (Dunbar, 2004, p. 103), which can keep people in line with arbitrary social norms and also serve the extremely important, socially and evolutionarily, function of protecting one’s social group against “free-riders,” who sneakily “take the benefits of sociality but decline to pay all of the costs” and are “extremely destructive for societies” (p. 106) (see also Giardini & Conte, 2012, on cheaters). Finally, and sadly, some gossip is purely malicious, purposely aimed at harming a (possibly innocent) victim, and may well succeed in doing so.

The benefits of complaints and laments discussed above are not limited to tokens involving human grammatical subjects. Recall that 13.22% of the *always* tokens containing NH subjects were coded as Complain, and 9.92% were coded as Lament. Complaints and laments do not require a human subject in order to facilitate bonding, and can convey crucial information

either way—including important social information. By this I mean that third party humans can be the implied recipient of empathy or disapproval even if they do not appear as grammatical subjects. For example, “the company is going bankrupt!” could, depending on context, imply concern and sympathy for employees, or scorn directed at the CEO.

The findings regarding laments merit further explanatory comments, as laments follow a slightly different pattern than complaints (and praise). Laments were most likely to contain 2P subjects (not 3P/O, like complaints), closely followed by NH and then 1P, and least likely to contain 3P/Os. It might seem odd that we lament the situations of non-human objects slightly more, even, than our own situations, and that we (seemingly heartlessly) do not often lament the situations of third party humans. However, for at least three reasons, the lament findings are difficult to interpret without conducting an intense qualitative analysis of all tokens.

First, a lament’s grammatical subject is not necessarily the recipient of sympathy. “Dogs are always biting her” has a non-human grammatical subject but expresses sympathy for a third party human, as does “the past was always waiting” (COCA) (where *the past* refers to memories of the woman’s deceased husband, which would constantly spring up and upset her). Second, unlike complaints, laments involving 1Ps are hardly “ineffective”. In fact, if one’s goal is to elicit (rather than express) empathy or sympathy, this type of lament is highly motivated, whereas 3P/Os would likely not be present to benefit. Third, 2P subjects in laments may be generic *you* or refer to people in a particular group or profession—which means that some laments containing 2Ps function like laments containing 1Ps. For example, when late-night comedian Conan O’Brian says, about his profession, “You can get into a funk, and beat up on yourself. Because you’re always thinking, ‘Maybe I can get one more moment like that out of my career.’ And you’ll walk across [broken] glass to get it” (COCA), he is talking about himself as much as about other people in show business, and not about his interlocutor.

The explanations offered thus far have focused, approximately, on why non-neutral functions are motivated with 2P and 3P/O subjects, but another approach is to ask why the

neutral function, Describe, is motivated with 1P subjects. One could say that we are narcissistic, and love talking about ourselves. A more sympathetic take, though, is that we are talking about what we happen to know best: We are experts on ourselves and our own experiences. In making neutral statements containing 1P subjects, we may be explicitly sharing information about ourselves, or simply providing the background information of a larger story (cf. the example earlier of the person who mentioned that he was “always trying to catch little critters and show the kids” (COCA) not to boast or carry on about himself but simply to set up a narrative). In other words, just as Rozin & Royzan (2001, p. 297) argue that the predominance of positive words reflects the fact that bad things are rare, it is possible that we share more facts about ourselves simply because we are the topic we are best-equipped to discuss.

The important thing to take from all of this is that it appears that the grammatical subject is an important variable in linguistic data, and should be taken into special consideration whenever possible. Admittedly, one must tread carefully when interpreting correlations between function and grammatical subject; the explanations above are plausible but difficult to prove empirically. However, if the purpose of language is to help us “achieve a particular interactional goal or set of goals” (Evans & Green, 2006, p. 110), then it matters, a lot, in terms of social and other practical matters, who you are talking to and who/what you are talking about. At the very least, a basic awareness that sentences involving first person subjects can differ greatly from sentences involving other humans, or non-human subjects, is crucial.

5.5.3 *Genre*

In this section I first explain why the effect of genre on the function of *always* + progressive tokens was not especially strong, and then why academic spoken language was the exception to this. The comments are kept brief, however, because much of what can be said on this matter already has been, in discussions in chapter 4 about the considerable overlap between written and spoken language (§§4.2 and 4.3), and about the formal, written-like nature of academic

spoken language, especially (§4.5.1). Since exaggeration and complaints are linked, the same information that shed light on how exaggeration is or is not used in those genres can help us understand patterns in those genres involving complaints, laments, praise, and vows.

One reason genre did not have a strong effect on function is that the genres analyzed represent a continuum rather than highly distinct categories. First, some of the spoken language analyzed here has written-like elements. For example, as was shown in chapter 4, in terms of exaggeration COCA's spoken language lies somewhere between, on the one side, COCA's written language and MICASE's academic spoken language and, on the other, the fully casual language of SBC and Swb. And the MICASE language in that study more closely resembled COCA's written language than it did casual conversation (which is not surprising, given its speakers' daily immersion in highly formal academic written language). This is important because (due to their size) the corpora containing written-like spoken language are over-represented in the analysis of *always/never* + progressive: Together, the unscripted TV/radio dialog in COCA and the academic spoken language in MICASE contributed nearly all (214) of the 226 spoken language *always* tokens; SBC and Swb, combined, contributed only 12.

In addition to the spoken language studied here containing written-like elements, the reverse is true as well: COCA's written language data contains spoken-like elements. For example, because it includes fiction, it contains things like dialogues between characters, or interview transcripts in magazines. This, too, narrows the gap between written and spoken genres, and further illustrates the need for careful study of sub-genres (taking into account their specific contexts and characteristics) alongside the study of more general corpora.

The main effect of genre that we do see on function is that the Academic Spoken category stands out for the Describe function being more common in it than in the other genres. As Chafe & Danielewicz argue, in characterizing genres, we must look at "the context of language use, the purpose of the speaker or writers, [and] the subject matter" (1987, p. 84). In this light, the prevalence of description in MICASE makes perfect sense. The purpose of the language spoken

in universities is to facilitate learning. As such, MICASE contains many examples of *always* + progressive used in a neutral way, such as to present eternal truths about science (e.g., “the genes are always being transcribed”) or society (e.g., “We're always [trying] to define ourselves with [a] sense of responsibility”), or to analyze literary characters (facts about whose personalities are, in a sense, timeless).

Academic spoken language contained not only the most Describe tokens of any genre, but also the fewest complaints, fewer laments than spoken language overall, and no instances of praise. We can attribute this to the impersonal nature of academic spoken language, which “has a subject matter which excludes much talk about oneself” (Chafe & Danielewicz, 1987, p. 107) and thus is less likely to involve emotional content such as complaints or laments. Overall, the findings show once again how context, function, and subject matter interact in complex ways and affect the characteristics of genres.

5.5.4 Differences Between Always and Never

One goal of this study was to determine if it was accurate to, in describing the typical function(s) of PAF + progressive, group *always* and *never* together. Although they were similar in some ways (for both words, Praise constituted the smallest percentage of tokens, about 4%, followed by the negative functions, and Describe the largest), they were also very different. The key distinguishing features of the *never* data were the prevalence of (1) future meaning, (2) the verbs *come* and *go*, and (3) a fifth function, to make a vow. Also, Laments were more common than Complaints in the *never* data, while the opposite was true of the *always* data.

The findings regarding the *go*-future and progressive future are, arguably, the most significant, as they show that it may not be possible to directly compare the *always* + progressive and *never* + progressive data sets. Even with instances of the *go*-future (which was far more common with *never* than with *always*) eliminated, three quarters of the *never* +

progressive tokens still carried a future meaning (i.e., they were instances of the progressive future, as in “I’m flying to Budapest tomorrow”). This difference is linked to the other significant finding, the discovery of a Vow function in the *never* data. The future meaning is what licenses the Vow function, as a vow is essentially a statement about what will be the case in the future, and the making of a promise that this statement is true.

Why the *never* tokens so often had this future meaning in the first place, though (given that the *always* tokens did not, even though they could have) is harder to say. It might be that it is easier to state with accuracy and honesty that which will never happen than that which will always happen. Consider laments about deceased loved ones, for example, like “She’s never coming back.” It is certain that these people will never be seen again, never come home again, etc. In contrast, it is harder to say with certainty what we or other living people will always do (or even frequently do), as people and their habits may change.

One unexpected consequence of this study was the realization that, due to these differences regarding tense-aspect, *always* + progressive and *never* + progressive do not (always) have antonymous meanings. Present tense utterances containing these adverbs seem to be antonymous (cf. “The neighbors here always make noise” versus “The neighbors here never make noise”), and this is possible with the progressive as well (cf. “they’re always looking for [teaching assistants]” (MICASE) and “they’re never looking for teaching assistants”). However, in most situations, the opposite of *always* + progressive is not *never* + progressive but *never* + simple. For example, it seems that the most natural opposite of “She’s always picking on them” (MICASE) is “She never picks on them,” rather than “She’s never picking on them.”

Thus, a statement true of one PAF when used with one particular tense-aspect may not be true of that PAF’s antonym in the same context, because the syntactically parallel utterance involving the antonym may differ in meaning, or might never be used. Relevant here is Tao’s (2003) corpus study of *remember* and *forget*. Tao points out numerous differences between the two, including in terms of tense-aspect (*forget* permits a greater variety of tense-aspects) and

grammatical subjects (*forget* appears almost exclusively with 1P subjects, while both 1P and 2P subjects are common with *remember*). In addition, each has characteristics and uses that the other does not. For example, the negative imperative “Don’t forget!” is common but the parallel “Don’t remember” never appears. In a study that compared *remember* and *forget* strictly in utterances with parallel syntax, many of these details would have been missed.

Antonyms pairs are, by definition, extremely similar in meaning, differing in just one key aspect (Jones et al., 2012, p. 3). For example, *tall* and *short* are both adjectives, and both about height, but refer to opposite ends of the scale. However, given that even near-synonyms differ in important ways (e.g., see Partington, 1998, on *complete*, *pure*, and *absolute*, and Partington 2004 on *entirely*, *completely*, *totally*, and *utterly*), it stands to reason that generalizing across antonyms is risky. I develop this argument further in the following section.

5.5.5 On Generalizations

In determining if PAF + progressive was frequently used with functions involving negative, unpleasant situations, this investigation uncovered important differences between *always* and *never*, along with correlations between the grammatical subject (human or not human, and first, second, or third person) and, to a lesser extent, genre. A lesson to take from this is the need to be wary of broad characterizations of words or categories of words. Generalizations are necessary and useful, but, when antonyms are involved or important semantic distinctions such as human versus non-human or other versus self, it pays to proceed cautiously. We must analyze corpus data rather than invented data and, moreover, we should be careful even with authentic utterances, as I explain below.

It is imperative to begin with authentic data. As Sinclair writes, “there is no justification for inventing examples” (1999, p. xi) if you intend to derive specific information from those examples. After all, any patterns “emerging from a set of constructed examples could not, of course, be trusted” (ibid.). In contrast, “a set of real examples may show, collectively, aspects of

language that are not obvious individually” (ibid.). A more subtle issue, though, is that even example sets consisting of authentic utterances can be unrepresentative.

For example, of the 74 example sentences³⁹ in Celce-Murcia & Larsen-Freeman’s section entitled “Form, Meaning and Use of Preverbal Adverbs of Frequency” (1999, pp. 504-511), 71 contained human subjects. It is hard to determine how representative this is, since the examples are of various syntactic varieties, but it seems skewed toward human subjects. Unless these proportions of human subjects to other human subjects accurately reflect overall patterns in the language, such example sets can be misleading. Thus, it is not enough for examples to be merely authentic; they must also be “carefully selected” (Sinclair, 1990, p. xi).

Finally, even generalizations based on sets of examples that are both authentic and representative can be erroneous, because such example sets are still incomplete. For example (and even though the grammars’ PAF sentences, overall, seem to contain an overly high percentage of human subjects), the PAF + progressive sentences are quite representative of what was found in my own data regarding grammatical subjects. Of the relevant sentences in Carter & McCarthy (2006, p. 47), Quirk et al. (1985, p. 199), Celce-Murcia & Larsen-Freeman (1999, p. 117, 510), and Sinclair (1990, pp. 249, 253) combined, 80% (12 out of 15) contained human subjects, 60% (nine out of 15) contained 3P/O human subjects, and 20% (three out of 15) contained non-human subjects. This quite closely matches the proportions in my own 764 *always* + progressive tokens (84.16% human, 54.71% human 3P/O, and 15.84% non-human).

These sentences are representative in terms of the frequency of 3P/O human subjects, and tokens containing 3P/O subjects are indeed the most likely to be complaints. Yet, it does not follow from this (as the examples tempt us to conclude) that complaints of the form 3P/O + *always* + progressive are common. On the contrary, 3P/O subjects in *always* + progressive

³⁹ I included in this analysis all declarative example sentences involving a grammatical subject and an adverb of frequency, exclusive of examples following “cf.” or multiple rephrasings in brackets. For example “Is Mr. Franks strict? Yes, he often is. (cf. Yes, he is often strict)” (p. 505) was counted as just one instance of a 3P/O subject, and “I {sometimes go / usually go} there” was counted as one instance of a 1P subject.

tokens are far more likely to be described (64.83%) than complained about, and complaints (involving all subject types) account for only 16.75% of the *always* + progressive tokens. These facts are not at all evident if we focus on sentences involving 3P/O subjects and complaints—which are, themselves, not actually common among instances of *always* + progressive.

On a similar note, it is problematic that, though the grammars' comments on PAF + progressive are mainly about PAFs (in general), most of their examples contain *always*. This was the case for three out of the four sentences in Sinclair's section on PAF + progressive (1990, p. 249), for example. The total token counts in the research conducted here indicate that *always* really is more common than *never*, so in that sense the examples are representative. However, it does not follow from that that they are representative of all PAF + progressive utterances in terms of function or other characteristics and tendencies.

Thus, while generalizations are useful and powerful, they should still be contextualized, when possible, and we must be alert to which factors are most influential in a given situation. In studying a quite different topic, intonation, Couper-Kuhlen & Selting (1996, pp. 19-20) warn us of the risk of attributing meaning to intonation when it actually comes from the lexico-syntactic content of an example. Likewise, we must be careful, in describing PAF + progressive, not to attribute to the construction as a whole characterizations that, actually, are more closely tied to particular items in it—items which may not be present in all instances.

5.6 Conclusion

This goal of this study was to see if claims that *always*/PAF + progressive is associated with negative functions or situations held true, the degree to which grammatical subject types and genre might be relevant to this, and if it was legitimate to extend claims about one PAF (*always*) to another (*never*), or if antonymous PAFs behave differently. It was found that, while was a correlation between third party human subjects and complaints (they were the subject type most likely to be complained about), the neutral function, Describe, predominated across genres and

subject types. It was also found that *always* and *never* differ in crucial ways. Various details regarding correlations between subject type and function, and genre and function, were explained from a functionalist, “language as action,” perspective, and also in terms of “face” (from Politeness Theory), empathy, and gossip.

Going beyond just the PAF + progressive construction, the biggest conclusion to be made here is that the grammatical subject involved in a construction (whether it is human or not, and first person, second person, or third person/party) may be crucially important to other aspects of the construction, such as its function, and that this is a natural consequence of practical and/or social considerations (such as how best to accomplish one’s goal, and concerns about rapport, face, and so on). Some of the conclusions presented here might seem obvious, in a sense (e.g., it is awkward to complain about a person to his/her face), but they bear mentioning because their implications are, nevertheless, rarely taken into consideration outside of studies specifically on point of view or empathy, despite their effects being so pervasive.

CHAPTER 6: TENSE-ASPECT

6.1 Introduction

Regarding tense-aspect, comprehensive grammars of English concur that *always*, and/or preverbal adverbs of frequency (PAFs) in general, are likely to appear with the simple and present perfect and unlikely to appear with the progressive, and that this is due to their function. That is, they commonly appear with aspects used to indicate habitual actions, general states of affairs, or repeated events. To test this claim, I conducted a large-scale analysis of the tense-aspects with which *always* and *never* appear (involving over 5,000 instances of the two adverbs from the four different corpora), as well an analysis of additional features (i.e., patterns regarding verbs and set phrases) which could prove relevant to tense-aspect. In doing this, I was also checking if it is legitimate to extend characterizations of *always* to *never* or if different adverbs, and especially antonyms, behave differently.

The analysis revealed similarities between the two adverbs but also striking differences. As expected, the past and present simple were the most common tense-aspects, together accounting for 71.05% of finite *always* and 57.92% of finite *never* tokens, and the progressive was rare with *always* and nearly non-existent with *never*. Also as expected, the present perfect was fairly common as well, accounting for nearly 16% and just over 25% of finite *always* and *never* tokens, respectively. Unexpectedly, though, the present simple was associated with *always* (appearing with *always* 2.5 times more often than with *never*), and the past simple was associated with *never* (appearing with *never* twice as often as with *always*). It was also found that copular *be* was 2.5 times more likely to be the main verb appearing with *always* than with *never*. Finally, qualitative analysis uncovered a number of set phrases involving *always* or *never*, most notably a multi-slotted idiom I call TV *As Always*, found at the conclusion of TV show guest segments, e.g., “Joan, Ron, as always, nice to have you” (COCA, Spoken).

In the analysis section, I discuss the tense-aspect patterns of *always* and *never* in light of more general tense-aspect patterns of English and also their top verbal collocates, showing that the present simple is more motivated with *always* because *always* tends to modify more enduring actions. Regarding set phrases, I focus on *TV As Always* and *TV Always*, showing that these have become idioms and that the former likely lacks tense-aspect.

6.2 Literature Review

One useful way to characterize a lexical item is “by the company it keeps” (Firth, 1957/1968, p. 179), i.e., its collocates, and one way to analyze verbal collocates is according to tense-aspect. Grammars of English that comment on the tense-aspect preferences of *always* and/or PAFs include the corpus-based Quirk et al. (1985) and Carter & McCarthy (2006), as well as Celce-Murcia & Larsen-Freeman (1999), who make frequent use of quantitative data. The consensus is that *always*/PAFs are unlikely to appear with the progressive aspect (Carter & McCarthy, p. 47) (which is rare in general; see Quirk et al., p. 198) and most likely to appear with the simple (Carter & McCarthy, *ibid.*, on *always*, and Celce-Murcia & Larsen-Freeman, p. 509, on PAFs), and, in particular, the present and past simple, as well as the present perfect (*ibid.*)

The posited reason for the link between PAFs and these tense-aspects is an overlap in function, i.e., both PAFs and the tense-aspects they commonly appear with are conducive to making statements about habitual actions and/or general states of affairs. For example, citing Praninskas (1975), Celce-Murcia & Larsen-Freeman (1999) explain that because PAFs “express approximately how many times a habitual action or condition is repeated” they “tend to co-occur with tenses that are used to express habitual action” too (p. 509). Likewise, Carter & McCarthy assert that *always* “refers to general states of affairs and to repeated events, and is therefore mostly used with present simple” (2006, p. 47). These observations are a helpful starting point, but they are yet to be proven empirically.

Three questions that one can ask are as follows: (1) Are the claims about the tense-aspect

preferences of *always* and/or PAFs true? (And, if so, can those claims be made more specific?) (2) Being antonyms, do *always* and *never* differ, regarding tense-aspect, in notable ways? (3) Speaking more broadly, to what extent can one generalize across word classes, or antonyms in the same word class? All of these questions are addressed below.

6.3 Method

To study the tense-aspect preferences of *always* and *never*, I collected a large number of tokens containing these lexical items from COCA, and all tokens from MICASE, Switchboard (Swb) and SBC, and coded them for tense-aspect. The tokens were also coded for verb type (copular *be* versus any other main verb). Because the prevalence of a particular phrase could influence the overall frequency of the tense-aspect in which it commonly appears, I also recorded recurring or otherwise noteworthy lexical/syntactic patterns.

6.3.1 Initial Collection of Tokens

The first step was to collect every token of *always* and *never*, along with its context, from MICASE (939 *always*, 676 *never*), Swb (359 *always*, 406 *never*), SBC (187 *always*, 222 *never*), and COCA 2010-2012 (24,058 *always*, 33,483 *never*). Every token in MICASE, Swb, and SBC was coded as Finite, Non-finite/Other, or Reject (see below). Because the number of COCA 2010-2012 tokens was so large, I obtained a random sample from each major category (the written categories Academic, Fiction, Magazine, and News, and the Spoken category): I separated the tokens for each adverb into the five categories, shuffled them into a random order, then coded tokens until I obtained 200 non-Reject tokens from each of the four written categories, and 400 from the spoken category, for a total of 1,200 tokens per adverb.

6.3.2 Quality Control Check of Tokens

Tokens that were rejected include the following:

1. **Duplicates and instant repeats:** Duplicates are instances in which the same line is repeated verbatim for a reason other than the writer/speaker truly repeating it. This happened, for instance, in magazines if the caption under a photo repeated a line already in the article. Instant repeats are instances when a speaker utters the exact or nearly exact same phrase twice in rapid succession or starts then restarts a token, repeating and finishing it. For example, “they’re always getting different - they’re always getting different angles” (COCA), was counted as one token, not two.
2. **Tense indeterminable due to grammar, incompleteness, transcription issues or non-standard gapping/elision:** The tense-aspect was ambiguous or indeterminable due to odd grammar (e.g., a mid-utterance change in syntax), incompleteness (e.g., due to the speaker cutting off his/her utterance or being interrupted), transcription issues (words that could not be heard, apparent mistakes in transcription), or extensive, non-standard gapping/elision.
3. **Citation forms:** The adverb is explicitly discussed as a word, e.g., “*Never* is a very dangerous word, Ivan” (COCA, Fiction).
4. **Part of a survey response:** The adverb constituted part of a Likert-scale survey response, as in the method section of a scientific article, e.g., “responses ranged from 1 (*never*) to 7 (*several times a day*)” (COCA, Academic).

6.3.3 Coding of Tense-Aspect

All non-rejected tokens were coded as Finite or Non-finite/Other. Tokens were coded as Finite if the verb modified by the *always* or *never* carried an unambiguous tense and aspect. Thus, all Finite tokens were further categorized as past, present, or future simple (*will*-future); past, present, or future progressive; past, present, or future perfect; or past, present, or future perfect progressive. The *go*-future (e.g., “this is never going to happen”) was considered separately, i.e., counted neither as future nor progressive. Finite tokens were also coded as containing either copular *be* or some other main verb.

Tokens were coded as Non-finite/Other if (the token did not meet the criteria for rejection and) the verb modified by the *always* or *never* was unable to take on the full range of tense-aspects (e.g., because it is in the infinitive form), or the tense was indeterminable for reasons other than those listed in the previous section. For example, in “As always, Bob went to bed at 9pm” the *as always* is compatible with the present simple (“as he always does”), past simple (“as he always did”), or present perfect (“as he always has”).

Ultimately, the results presented in this study focus on tokens of non-ambiguous tense-aspect; Non-finite/Other tokens are not further discussed aside from the analysis of certain set phrases. Regardless, in order to accurately identify them, it was necessary to develop a detailed coding schema of Non-finite Tokens. This schema is presented below:

1. **Imperatives:** An example would be “Always remember to wash your hands”.
2. **Verb Complements:** (a) **Bare Infinitive** (additionally coded for being after modal verbs versus after other verbs): Examples would include “I should always go” and “She made him never forget”;⁴⁰ (b) **To-Infinitive Complement:** These are instances in which the *always/never* comes after a verb that can take a *to*-infinitive complement; (c) **Ing-Complement:** Examples would include “he considered going”.
3. **Gerunds:** Examples would include “Never eating well is bad for you.”
4. **Subjunctive Mood:** Examples would include “I wish I were always faster”. The subjunctive was extremely rare, accounting for only 0.32% of all tokens (16 tokens).
5. **Isolated Instances:** These are instances in which the *always* appears alone, e.g., as a question, the reply to a question, as an entire statement.⁴¹
6. **Set Phrase:** These are instances in which the item appears in a set phrase that appears to be derived from or an elided version of a fuller utterance with tense-aspect, but the tense-aspect is determinable. This is the case with the “As always, Bob went to bed at 9 p.m.” example discussed above. Aside from *as always*, other such set phrases found included *as never before*, and *better late than never*. These are discussed in further detail in the results and discussion section.
7. **Others:** There are other cases in which lexical items which are derived from verbs and thus can appear with adverbs of frequency nevertheless cannot take a full range of tense-aspects. Examples include adverbial participles (e.g., “Having never received an answer, he gave up,” “I ate, never realizing ...” or “After jogging, I ...”), compound participial adjectives (e.g., “The well-worn blanket,” “the never-eaten food”), adverbs of frequency preceding attributive adjectives (e.g., “The never-cheerful clerk said ...”), reduced relative clauses (e.g., The man [who is/was] at the bank was nice”), including reduced relative clauses that are appositives (e.g., “Katherine, [who is/was] never [called] Kathy, has been

⁴⁰ With modals, this coding applied even if the modal appeared after the adverb (the logic being that, e.g., “I should always eat this” and “I always should eat this” do not differ in meaning in significant ways). With other verbs, placement drastically affects meaning (e.g., “She made him never forget” and “She never made him forget” mean different things, so the former would be coded as past simple and the latter as a non-finite bare infinitive.) However, no instances of *always* or *never* with bare infinitives following verbs other than modal verbs were actually found.

⁴¹ Arguably, the tense-aspect in these situations could be deduced (e.g., in “Are you hungry?” “Always!” the response could be said to mean “I am always hungry.”) However, this sort of extensive reconstruction is unreliable because it is ambiguous. Consider “Did you ever get caught?” “No.” “Never?” “Never.” The first and second *never* could mean “You never got caught?” and “I never got caught,” respectively, both in past simple, but the present perfect would work too. Another possibility is that the second *never* means “I affirm the (negative) statement you just made”. Given the uncertainty, I err on the side of caution and try to avoid assuming underlying syntax, morphology, and forms.

working hard” (COCA)), and so on.

6.3.4 Data Analysis

To better understand the tense-aspect preferences of *always* and *never*, I calculated, first, the frequencies of the tense-aspects each adverb appeared in as a percentage of all Finite tokens, in both the overall data and by genre. Second, I calculated for each adverb the rate at which copular *be* appeared as the main verb, in both the overall data and by tense-aspect. The latter analysis was conducted in order to uncover other potential ways in which *always* and *never* differ regarding verbs, and which might shed light on their tense-aspect preferences. Differences in tense-aspect and verb patterns between adverbs were tested using the chi-square test for statistical significance at the $p < .0001$ level. Third, I noted several recurring phrases containing the adverbs and considered these in terms of tense-aspect.

6.4 Results

The total number of *always* tokens analyzed in terms of tense-aspect, main verb (copular *be* versus other verbs), and other features was 2,562 (1,200 from COCA, 866 from MICASE, 332 from Swb, 164 from SBC) and the total number of *never* tokens was 2,389 (1,200 from COCA, 628 from MICASE, 364 from Swb, 197 from SBC). The total number of *always* and *never* tokens coded as finite were 2,224 and 1,875, respectively (see Fig. 6.1).

Total and finite tokens										
	COCA		MICASE		Swb		SBC		TOTALS	
	all	finite	all	finite	all	finite	all	finite	all	finite
<i>always</i>	1200	1032	866	749	332	303	164	140	2562	2224
<i>never</i>	1200	915	628	478	364	320	197	162	2389	1875
	2400	1947	1494	1227	696	623	361	302	4951	4099

Figure 6.1: Total and finite tokens

The key results were that (a), regarding tense-aspect, there was an association between *never* and the past simple, and *always* and the present simple, and an absence of major cross-

genre differences; (b) copular *be* appeared as the main verb more often with *always* than with *never*; and (c) several set phrases or idioms—namely TV *As Always*—were found in which the tense-aspect was ambiguous or potentially absent.

6.4.1 Tense-Aspect

The percentage of finite tokens for each adverb in each tense-aspect is shown in Fig. 6.2. Tense-aspects unattested for either adverb—i.e., past perfect progressive, future perfect progressive, and future perfect—are not depicted, nor are tense-aspects for which a single token was found—i.e., future progressive and present perfect progressive—but these tokens are factored into the percentages. The *go*-future is included as well, separately from the progressive. This is true for all of the following figures in this section also.

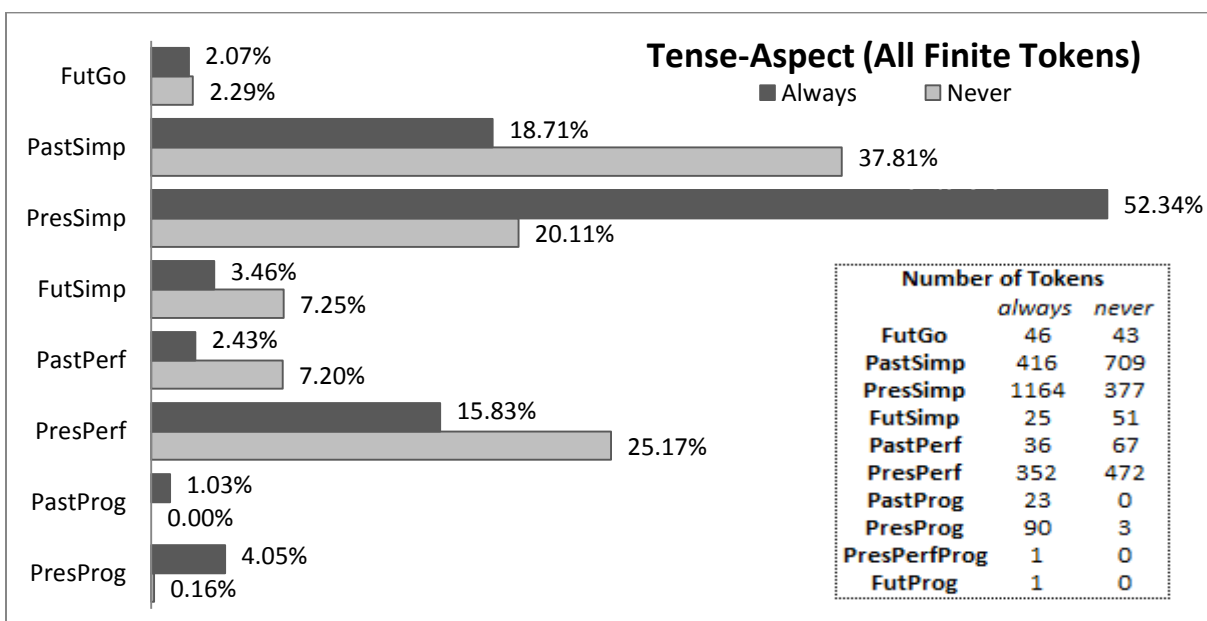


Figure 6.2: Tense-aspect (all finite tokens)

The most striking finding was a strong association between *always* and the present simple that contrasts with a strong association between *never* and the past simple (see Fig. 6.2). Specifically, 52.34% of *always* tokens were in the present simple tense-aspect, but only 20.11% of *never* tokens, and 37.81% of *never* tokens were in the past simple, but only 18.71% of *always*

tokens. In addition, the present perfect was somewhat more common with *never*, accounting for 25.17% of *never* tokens versus 15.83% of *always* tokens. Less strong patterns include the progressive being rare with both adverbs yet more common with *always* (90 tokens, 4.05% of finite tokens) than *never* (no instances), and the future simple and past perfect being over twice as common with *never* (but with a difference of only about 3% between adverbs). The data was also considered in terms of aspect, with all tenses combined (see Fig. 6.3), and in terms of tense, with all aspects combined (see Fig. 6.4).

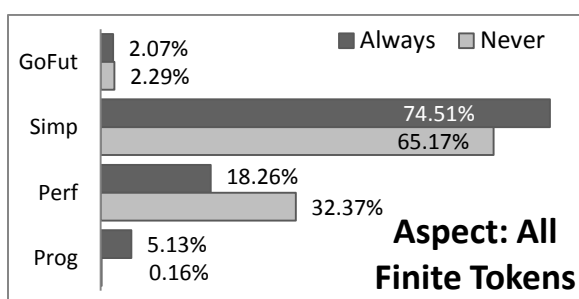


Figure 6.3: Aspect (all finite tokens)

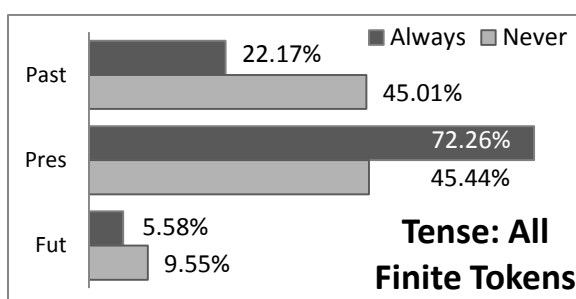


Figure 6.4: Tense (all finite tokens)

Surprisingly, the observed tense-aspect patterns differed very little between written and spoken language (COCA's news, academic, fiction, and magazine categories versus the academic spoken language of MICASE and casual spoken language of Swb and SBC) (see Figs. 6.5 and 6.6; whited out bars in the figures indicate the lack of statistical significance at the $p < .0001$ level). In fact, the general patterns described above are visible even in highly formal written language (COCA's academic written language) and the most casual spoken language (Swb and SBC) (see Figs. 6.7 and 6.8) (the only notable difference between the two being that the former lacks the usual gap between *always* and *never* in present perfect). Due to this, the rest of this chapter focuses on the combined data, and does not compare genres.

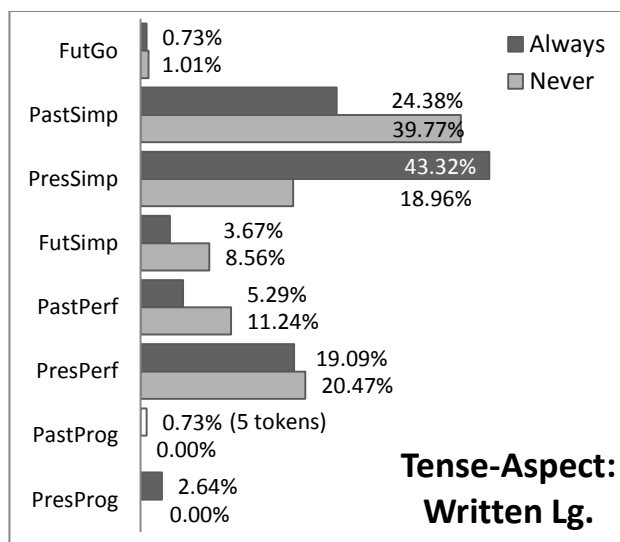


Figure 6.5: Tense-aspect, written language

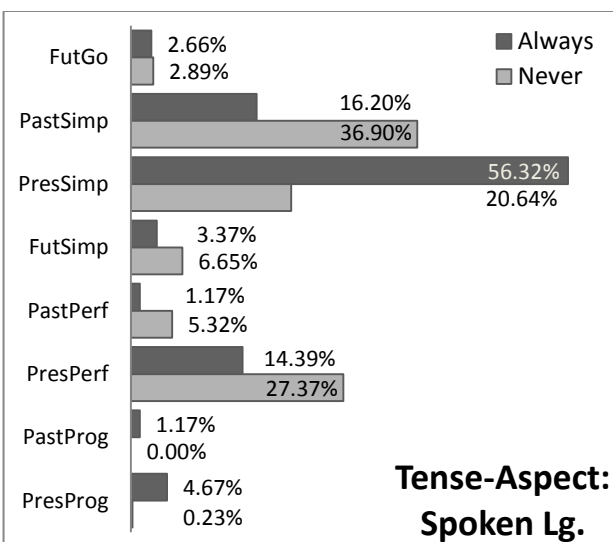


Figure 6.6: Tense-aspect, spoken language

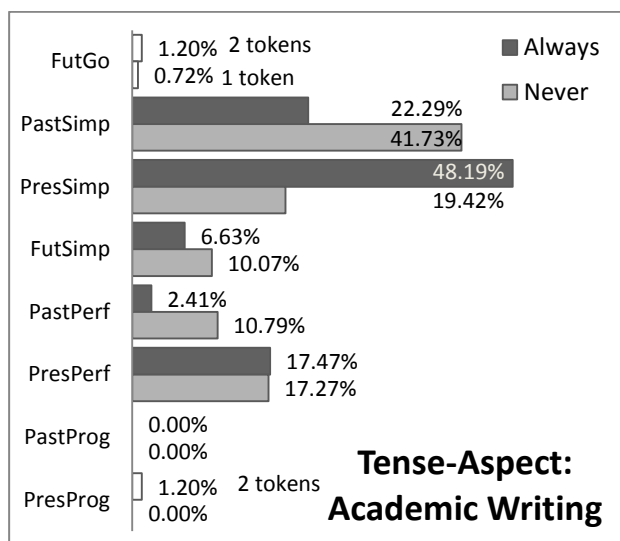


Figure 6.7: Tense-aspect, academic writing

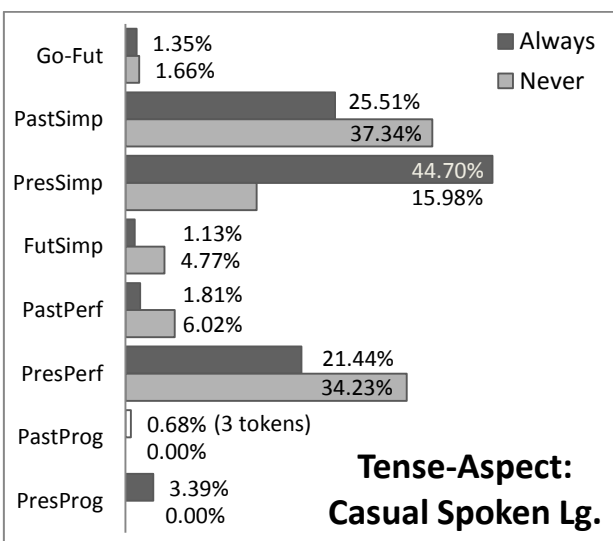


Figure 6.8: Tense-aspect, casual spoken language

6.4.2 Verbs

Every finite token was coded as containing copular *be* (e.g., *am* in “I am always happy”) or a different main verb (e.g., *run* in “He never runs” or “He is always running”, the latter containing auxiliary, but not copular, *be*). The resulting verb preferences of *always* and *never*, both as a whole and in terms of particular tense-aspects, were compared. Overall (see Fig. 6.9), copular *be* was 2.5 times more common with *always* than with *never*. These results were significant at the

$p < .0001$ level using the chi-square test.

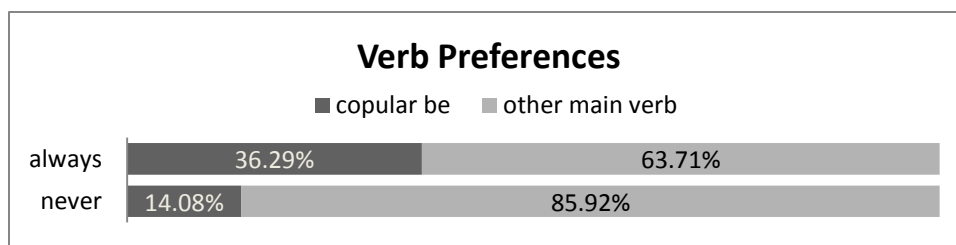


Figure 6.9: Verb preferences

This pattern is a very general one: For every tense-aspect combination for which statistically significant results at the $p < .0001$ level were obtained (i.e., all but future simple and present progressive; see Fig. 6.10), main verbs other than copular *be* are two to three times more common with *never* than with *always*—or, in the case of the past perfect, even slightly more than three times as common. This robust pattern indicates, once again, that *always* and *never* differ in ways beyond that directly predictable from their antonymous meanings.

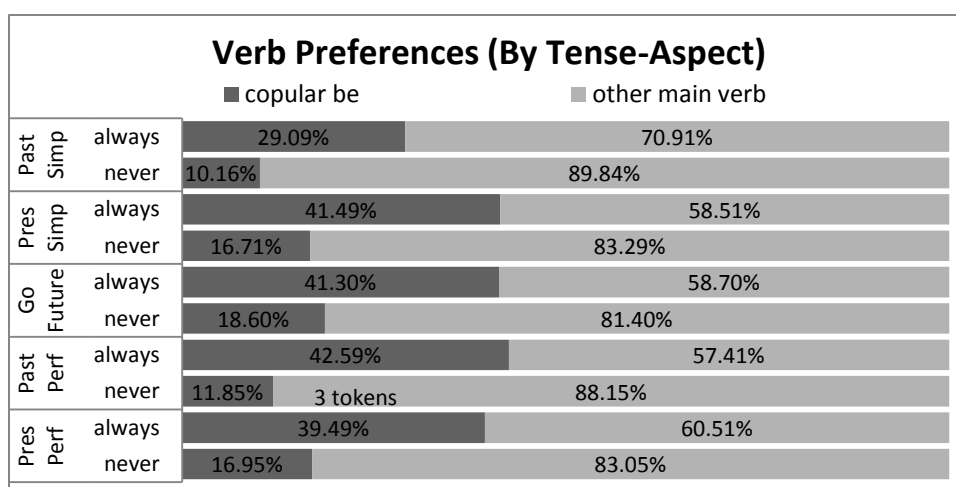


Figure 6.10: Verb preferences (by tense-aspect)

6.4.3 Set Phrases

During the coding process, frequent or otherwise interesting lexical patterns were noted. These appear to have become set phrases or idioms, at least one of which has arguably lost its tense-aspect. The most notable of these are (a) *as always* and *always* as used in TV shows (henceforth

TV *As Always* and TV *Always*), in addition to *as always* in other contexts; and (b) *never mind* (highly frequent among imperatives, and used in three ways). Other set phrases that appeared in the data were *as never before*, *never more*, and *better late than never*.

TV <i>As Always</i> : Examples from COCA (Spoken)	
1	HOST (DAVID): Mr. Daley, thank you as always. GUEST (Mr. DALEY): Thanks, David.
2	HOST: All right, guys, great panel tonight. As always and that's all the time we have left this evening.
3	HOST (BRET BAIER): Brit, as always. Thank you. GUEST (BRIT HUME): You bet, Bret.
4	HOST: Great. Regina Lewis, great insight, as always. Thanks so much for being here.
5	HOST: ... All right, Charles. Thanks, as always.
6	HOST: Michael Billy, I thank you as always.
7	ANCHOR: Thanks so much as always for watching.
8	ANCHOR: As always, appreciate you watching.
9	HOST: Lisa Bloom for us tonight. Lisa, as always, thank you. CORRESPONDENT (BLOOM): Thank you.
10	HOST (DAVID G.): David Plouffe, thank you as always. GUEST (DAVID P.): Thanks, David.
11	HOST: You are a winner, Eva Longoria. ... You knocked out of the park [sic] as always. Thank you. You're always such a good sport. GUEST (EVA L.): Thank you.
12	HOST: Thank you, Chad. Chad Myers as always right on top of things.
13	HOST (TODD): Joan, Ron, as always, nice to have you. GUEST 1 (RON CHRISTIE): Thanks for having us. GUEST 2 (JOAN WALSH): Thanks, Tony.
14	HOST: Peter, thanks. Good to see you as always. GUEST (PETER): You got it
15	HOST (ANDERSON COOPER): Fouad Ajami, appreciate it, as always. Paul Conroy, I appreciate you, interrupting your treatment to talk to us tonight ... Thank you, Paul.
16	HOST (RUSS M.): Alright, Doctor Jennifer Ashton, great advice. GUEST (Dr. ASHTON): Thank you, Russ. HOST: As always we appreciate it.
17	HOST (NEAL C.): Of course, our thanks as well to political junkie Ken Rudin. Appreciate it as always, Ken. GUEST (KEN R.): Thanks, Neal.

Figure 6.11: TV *as always*

The phrase *as always* occurs 40 times, 17 of which (all from COCA, Spoken) were found to be part of a conventionalized pattern used on TV shows, most typically by a host speaking to a guest (or anchor to a correspondent) at the end of the addressee's segment, and served as a way to thank the person and say goodbye (see Fig. 6.11, TV *As Always*).⁴² The utterance typically involves the addressee's name, *as always*, and a form of *thank you*, e.g., "Lisa, as always, thank

⁴² An additional five of the 40 instances of *as always* occurred in another, slightly different TV context: They were used to highlight what routinely happens on the show, e.g., "And a few words before we go, as always" or "As always, you can join the live chat now at AC360.com" (both from COCA, Spoken).

you” or “Mr. Daley, thank you as always” (both from COCA, Spoken). Of the 17 examples, 14 included the name of the addressee. (The other three involved pronouns; in two cases *you*, the audience, and, in another case, *guys* used as a plural second person pronoun addressing a panel.) In addition, 14 included *thanks* or *thank you* (and some of these included *appreciate* and/or a compliment as well), while the remaining three did not include *thanks* or *thank you* but did include *appreciate* or a compliment (e.g., “nice to have you” or “great advice”). All 17 occurred at the end of a segment, which sometimes also marked the end of the entire show.

TV <i>Always</i> : Examples from COCA (Spoken)	
1	HOST: Ray, thanks. Always good to have you here. GUEST (RAY M.): You're welcome.
2	HOST: Casey Jordan. Thanks a lot. GUEST (CASEY J.): Always great to be here. HOST: I appreciate it.
3	HOST (MICHEL M.): Welcome, Tavis, thanks so much for joining us. GUEST (TAVIS-SMILEY): Michel Martin, always an honor to be on with you.
4	HOST: Always good to see you, Chris. [CHRIS = correspondent]
5	HOST (BILL O'REILLY): George, always a pleasure, thanks very much. GUEST (GEORGE S.): Take care, Bill.
6	HOST: Well, Mark, thank you very, very much. It's always great to have you. GUEST (MARK L.): It's a great pleasure, my friend. HOST: Excellent writer, great thinker, great lawyer, good friend, Mark Levin.
7	HOST: Great advice. Doctor Holly Phillips. Thank you. ... We really appreciate it. GUEST (Dr. PHILLIPS): Always great to be here.
8	HOST: We always enjoy it when you're here, Chris. ... It is always great.
9	GUEST (CHRIS R.): Always great to see you.

Figure 6.12: TV *always*

Less commonly, *always* (without *as*) was used in the same television context and with a similar function. A typical example of this TV *Always* construction is “Ray, thanks, always great to have you.” These were coded as instances of the present simple with (an implied) copular *be*. (Unlike TV *As Always*, TV *Always* never appeared alone or with only with an addressee’s name, making its tense-aspect more evident.) Of the nine⁴³ examples of TV *Always* analyzed (see Fig. 6.12, TV *Always*), six included the addressee’s name (the three that did not were all in the

⁴³ Eight were found in the randomly selected COCA tokens but nine appear in the figure because one of the eight contained, within it, both a host-to-guest TV *Always* and a guest-to-host TV *Always*.

direction of guest to host rather than host to guest, and after the host had just used the guest's name); six involved *thanks* or *thank you*, and were in the direction of host to guest; seven lacked a grammatical subject (e.g., "Always great to be here," not preceded by *it's*); and in seven the noun or adjective (*great, pleasure, honor, good*) took a *to* complement (e.g., "Always a pleasure to see you" rather than "Always a pleasure"). Seven of the nine utterances occurred at the end of a guest's segment, and six of the nine were in the direction of host to guest.

Tokens containing the imperative *never mind* (56 tokens) accounted for 2.26% of the 2,389 Finite and Non-finite/Other tokens, with most of those tokens (39) coming from MICASE, and accounted for 67.45% of the 83 tokens, total, coded as imperatives. Ten of the 56 tokens were further coded as "isolated," meaning *never mind* was used as a standalone phrase to mean "disregard what I just said". Another ten were coded as "NP/Clause," meaning they were followed by a noun phrase or a clause beginning with *that* or *about* (e.g., "Never mind what put him there" and "never mind about the phone", both from COCA). This is simply a variant of the first use, meaning "disregard [thing/fact]". Finally, two were used to mean *let alone*. This conjunctive use of *never mind* appears in dictionaries, e.g., Merriam-Webster (2015), but is rare and, as is shown below, may not be used in entirely consistent ways.

One example of the *let alone* use, about a teacher's attempt to get a "selectively mute" child to speak, is "cajoling, doling out treats and praise without even winning a smile from Sheila, never mind an answer" (COCA). Here, the *never mind* is interchangeable with *let alone*: If Sheila would not even smile, then she certainly would not do what the teacher even more desperately wanted her to do, i.e., speak. The other example, about the length of science articles, is "After all, time and space, never mind attention span, are limited" (COCA) (in the context of how science writers struggle to convey scientific findings briefly yet accurately). This is harder to parse, because each of the noun phrases seems to involve different people: The writers deal with the limits on time and space, and the readers have the limited attention spans. It is possible, though, to produce a *let alone* version: "We have not successfully worked around the publisher

constraints on time/space, let alone the constraints on reader attention spans.”⁴⁴

6.5 Discussion

In sum, the results confirm the claims in Quirk et al. (1985), Carter & McCarthy (2006), and Celce-Murcia & Larsen-Freeman (1999) that, with *always* and *never*, the simple and present perfect are common and the progressive rare, but they also reveal a number of differences between the two items. Most notably, the present simple is more strongly associated with *always* than with *never*, while the reverse is true of the past simple. In addition, copular *be* is 2.5 times as common with *always* as with *never*, and the words appear in different set phrases, some of which appear to have lost their tense-aspect.

6.5.1 Tense-Aspect

In analyzing the tense-aspect findings, I begin with what *always* and *never* share in common. The rarity of the progressive and the frequency of the present simple with them is consonant not only with the *always* and *never* literature but also with more general studies of the distribution of tense-aspect. Analyzing the 40 million word Longman corpus, Biber, Johansson, Conrad, & Finegan found the progressive to be quite rare with respect to the simple (1999, p. 461). In conversation, the progressive and perfect aspects were about equally frequent, and each accounted for only roughly five thousand out of roughly 130 thousand verbs analyzed. (The remaining approximately 120 thousand were in the simple aspect; this is, again, consonant with the data presented here.) And in the fiction, news, and academic genres, the progressive was even rarer than the perfect. Finally, and as in the data here, Biber et al. (ibid.) find that, overall, the present perfect is much more common than the past perfect. (This was their finding for

⁴⁴ Admittedly, we should be wary of such an extensive transformation. However, it is unclear how else to understand the example, and the suggested understanding is somewhat supported by the fact that conjunctive *never mind* is found “especially in negative contexts” (Merriam-Webster, 2015). If so, it might be able to imply a negative context (“have not successfully”) even when one is not explicitly there.

conversation, written news, and academic prose, though not for fiction).

The findings of Biber et al. indicate that the reasons behind the frequency of the simple aspect and scarcity of the progressive noted here are more general, and thus unlikely to be uncovered through analysis of *always* and *never*. What is more interesting for the purpose of this study is not how the tense-aspect patterns of the two adverbs are similar (both to each other and to the pattern found in language in general), but how they differ, the most notable respect in which they do being their tense preferences in the simple aspect.

The difference in tense preferences of *always* and *never* regarding the present and past simple is striking: 52.34% of the verbs following *always* are in present simple and 18.71% in past, while 37.81% of the verbs following *never* are in the present simple and 20.11% in past. Because the pattern is complementary, it cannot be explained by general tendencies. The data for *always* are in line with the finding of Biber et al. (1999, p. 456) that in conversation present tense is much more common than past tense (with present tense accounting for roughly 104,000 out of roughly 152,000 verbs they analyzed). In written news and academic prose, as well (but not in fiction) present tense was more common than past (ibid.). The data presented here for *never*, in preferring past over present, contradict this general pattern.

Making sense of these differences requires more information. Because an understanding of the verbs commonly used with each word might be useful, I compared the 100 highest-ranked right-hand (one or two items to the right) lemma collocates of *always* to those of *never*. I did this using COCA's "compare" feature, which ranks collocates using a score calculated by dividing the number of times the collocate appears near the word being analyzed, e.g., *always*, by the number of times it appears near the other word, e.g., *never*. Because this method disregards collocates equally common for both words, it is very useful for teasing out differences.

Analysis of the strongest right-hand collocates reveals a trend regarding verbs of affection and verbs of cognition (Halliday, 1985, p. 111)—verbs that denote mental processes

such as *love*, *want*, and *hope*, and which I will refer to as *mental verbs*⁴⁵—that helps explain the association between *always* and present tense. Among the 100 strongest collocates of each adverb, mental verbs appeared 2.75 times more frequently with *always* than with *never* (22 instances versus eight) (see Fig. 6.13). In addition, they constituted six of the ten strongest collocates of *always*, but only one of the ten strongest collocates of *never* (and that one that did occur was part of a very common set phrase, *never mind*) (see Fig. 6.14; mental verbs are in bold). It is worth noting, further, that the scores are very high for approximately the first ten collocates and much lower for the rest (46.4 and lower for *always*, and 37.2 and lower for *never*), meaning that the first ten are particularly characteristic. The first collocate of *always*, in particular (*fascinate*, a mental verb) has an extremely high score, 603.8.

Mental Verbs (Totals)		
	<i>always</i>	<i>never</i>
of the first 100 collocates	22	8
of the first ten collocates	6	1

Figure 6.13: Mental verbs (totals)

Ten Strongest Collocates: Mental Verbs (bold) or Non-Mental Verbs											
ALWAYS (WORD 1)						NEVER (WORD 2)					
	WORD	W1	W2	W1/W2	SCORE		WORD	W2	W1	W2/W1	SCORE
1	[FASCINATE]	208	0	416	603.8	1	[CEASE]	464	1	464	319.7
2	[LURK]	64	0	128	185.8	2	[MIND]	3539	11	321.7	221.7
3	[PRIDE]	52	0	104	150.9	3	[CONVICT]	93	0	186	128.2
4	[PREFER]	168	2	84	121.9	4	[MASTER]	68	0	136	93.7
5	[INTRIGUE]	63	1	63	91.4	5	[WAVER]	251	2	125.5	86.5
6	[PUZZLE]	28	0	56	81.3	6	[RELINQUISH]	57	0	114	78.5
7	[REVOLVE]	21	0	42	61	7	[GRADUATE]	93	1	93	64.1
8	[AMAZE]	72	2	36	52.2	8	[DUPLICATE]	43	0	86	59.3
9	[CONSIST]	17	0	34	49.3	9	[REMARRY]	84	1	84	57.9
10	[TEMPER]	17	0	34	49.3	10	[ARREST]	127	2	63.5	43.8

Figure 6.14: Mental verbs (of ten strongest collocates)

Mental process verbs (among the top 100 collocates of *always* are *pride*, *prefer*, *distrust*, *admire*, *adore*, *cherish*, *revere*, *yearn*, *treasure*, and *favor*, for example) usually refer to what seem to be permanent/enduring states or situations. That is, our likes and dislikes, and what we

⁴⁵ Halliday uses a similar term, *mental process verbs*, but this includes verbs of perception such as *see* (p. 111), whereas my term *mental verbs* is not intended to include such verbs.

revere, yearn for, and treasure, are probably fairly stable. Thus, references to them are more likely to appear in the present simple than in the past simple, since the latter would suggest that the situation no longer holds. (The scarcity of the future, as a third option, is accounted for by the scarcity of the future tense in general.)⁴⁶

Not only did the list of the 100 strongest collocates of *never* include fewer mental verbs than that of *always* (again, seven compared to 22), but, also, several of the few it did contain refer to events that are punctual rather than durative. Consider the three highest-scoring collocates after *mind*, which were *dawn*, *forgive*, and *forget*. Forgetting is not generally seen as something that takes place gradually (even if this is the cognitive reality). By definition, things *dawn* on us suddenly, and forgiving is instantaneous too (assuming that true or prototypical forgiveness is full forgiveness), even if working up to it takes a long time. Thus, even when *never* is used with mental verbs, the present tense seems less motivated than it is with the mental verbs that appear with *always*. Moreover, many of the non-mental verbs among the collocates of *never* (which included, e.g., *abandon*, *arrest*, *baptize*, *consummate*, *graduate*, *marry*, *remarry*, *renounce*) are not only punctual, but probably happen only once in a person's life.

If or how the trend regarding mental process verbs helps explain the association between the present perfect and *never*, on the other hand, is unclear. It seems that the present perfect works equally well for punctual events and events with duration (e.g., “I’ve never asked her” versus “I’ve always loved it”), so we likely must seek an explanation elsewhere.

6.5.2 Verbs

Always tokens were, overall, 2.5 times more likely to contain copular *be* than *never* tokens were (36.29% of finite *always* tokens contained copular *be* but only 14.08% of finite *never* tokens).

⁴⁶ Biber et al. (1999, p. 456) found that the future accounted for fewer than 20,000 of roughly 152,000 verbs they analyzed. The relevant category in their study is actually “modal verbs.” Recognizing that “there is no formal future tense in English” and that, instead, “future time is typically marked in the verb phrase by modal or semi-modal verbs” (ibid.), Biber et al. do not include future tense as its own category. This is true but it was included as its own category here because I was interested in the future meaning, separate from other modal verb situations.

For each tense-aspect, as well, copular *be* was about two to three times more strongly associated with *always* than *never*. The reason for this pattern is yet to be determined, but it is significant for constituting yet another way in which *always* and *never* behave differently.

6.5.3 Set Phrases

As mentioned earlier, a number of set phrases or constructions were found: *Never mind*; *as never before*; *better late than never*; *never more (so) than*; *as always*, including TV *As Always*; and TV *Always*. TV *As Always* typically consisted, when used by a host, of “[name], as always, thank you” or “good to have you here, as always” or, when used by a guest, “a pleasure to be here, as always”. The more truncated TV *Always* typically consisted, when used by a host, of “[name], always good/great to have you” or, when used by a guest, “always a pleasure to be here”. Instances of TV *Always* were coded as present simple, and *never mind* as Non-finite and an imperative. The other set phrases were considered Non-finite/Other because their tense-aspects were ambiguous. I argue below that TV *As Always* and *Always* have become idioms (more specifically, according to the schema of Fillmore, Kay, & O’Connor (1988), encoding but perhaps becoming decoding, grammatical, formal idioms, with a clear pragmatic point), and that one facet of this idiomization is that the former appears to lack tense-aspect.

Idiomatic expressions or constructions are those not knowable based solely on knowledge of a language’s grammar and vocabulary; knowing those things, one could still be ignorant of “how to say [the idiom], ... what it means, or ... whether it is a conventional thing to say” (Fillmore et al., 1988, p. 504). A standard example of an idiom is *kick the bucket* when used to mean *die*. According to Fillmore et al. (ibid.), idioms can be decoding or encoding⁴⁷. Decoding idioms, *kick the bucket* being one example, must be learned whole because their meanings cannot be predicted from their parts (pp. 504-505). The meanings of encoding idioms, on the other hand, actually are knowable from their parts; what is unpredictable about them is just

⁴⁷ Technically, decoding idioms are encoding, too, “but there are encoding idioms which are not decoding” (p. 505).

their conventionality (p. 505). *TV As Always* and *TV Always* are—at least when enough of their parts (name, *thank you*, etc.) are present—encoding idioms. Any native or advanced speaker of English could understand them upon first encountering them, but they are idiomatic in that we could express gratitude toward TV show hosts/guests in a number of other ways which, by chance, have not become conventionalized, or not as strongly conventionalized.

Fillmore et al. explain, further, that idioms can be grammatical or extra-grammatical, substantive or formal, and with or without a pragmatic point (i.e., a specialized pragmatic or rhetorical purpose) (1988, pp. 505-506). If idioms are substantive, then “their lexical make-up is (more or less) fully specified”. These are the opposite of formal idioms, which are more like syntactic patterns with slots open to different lexical items. Despite the apparent heavy elision which obscures their tense-aspect, *TV As Always* and *TV Always* appear to be grammatical. They are also formal: The ordering of their parts and what fills the parts (different names, *good* versus *a pleasure*, etc.) can vary, and some parts are optional. Finally, they have a pragmatic point: They occur in a restricted setting (TV shows involving hosts or anchors and guests or correspondents) and allow guests or hosts to express appreciation or pleasure.

While these initially appear to be encoding idioms, they may be evolving toward being decoding. I present, below, several examples (1a-f) (these and all other examples in this section are from COCA’s spoken portion) in which *TV As Always* appears by itself, alone standing for the host’s “Bob, (it’s) good to have you, as always” (etc.) or guest’s “(it’s) a pleasure to be here, as always” (etc.). Admittedly, many instances of seemingly isolated *TV As Always* are simply discontinuous with a nearby *thanks* or *a pleasure to be here*, as in “HOST: Jean Chatzky, thank you. GUEST: Sure. HOST: As always.” In other cases, they could be derived via analogy from a contiguous instance of a host thanking a different guest, as in “Thanks—thanks for being here with us, David. And, Tony, as always”. Such explanations are inapplicable, however, to the following examples, all of which occur at the end of a guest’s segment:

- (1a) ANCHOR (COLMES): General, good to see you. Thanks for being here.
GUEST (General VALLELY): As always.
- (1b) HOST: Well, Daniel, delightful talking to you again.
GUEST (DANIEL P.): As always.
- (1c) HOST/ANCHOR: ... Donna, Ron, thanks so much.
GUEST (Prof. DONNA B.): Thank you.
GUEST (RON C.): As always.
- (1d) HOST: Doc, what a pleasure. Al—as always, good to see you.
DOC WILLOUGHBY: As always.
- (1e) HOST: ... Good to see you.
GUEST: Thank you.
HOST: Congratulations.
GUEST: Pleasure.
HOST: As always. Yeah. My hand. We've been doing this a long time, right?
GUEST: We have. It's been many, many years.

In (1a-d), the host sets up the typical host-version of TV *As Always* with statements such as “Good to see you” or “thanks,” and the guest responds “As always”. In such cases, the guest cannot be completing the host’s statement but must instead mean “(It’s) a pleasure to be here, as always” (etc.). Consider (1b): The guest must mean it is a pleasure for him to be there, not that his presence is a pleasure for the host. Finally, in (1e) we see the reverse situation: The guest says simply “pleasure”—a truncated version of “(It’s) my pleasure” or perhaps an even more truncated version of “(It’s) always a pleasure to be here” or “(It’s) a pleasure to be here, as always”—and the host responds “as always.” Presumably, he is expressing his own pleasure in having her there, rather than stating what is a pleasure for his guest. Interestingly, this extremely elided exchange occurs between a highly experienced actress, Jodie Foster, and a host who has apparently interviewed her many times before, prompting them to comment explicitly on how routine this type of interaction has become for them.⁴⁸

Finally, there is at least one instance (1f) in COCA of a host spontaneously using “[name] + *as always*” to thank the guest and end the segment, i.e., with no explicit thanking (or

⁴⁸ Their discussion about this goes on even longer, to include an awkward joke about how, because they have seen each other so many times, it might appear that they are dating: FOSTER: We have. It’s been many, many years. I’m not going to even say how many, because I don’t want to date you. HOST: You don’t want to date me? FOSTER: Well, that’s not what I meant. I mean, I’d love to date you, but I don’t want to date you. HOST: All right.

complimenting) or leave-taking beforehand to which it could be said to be “attached”, and without even an item like *alright* or *well* that would signal that he is wrapping up:

(1f) GUEST (ANDY R.): It happened once before. Who was that? [CROSSTALK] One of those people, yes. But I doubt if it’s going to happen this time, and if it does, why, it may force us to make change. We need serious changes in our electoral process. I mean, democracy is such a wonderful idea, and we made a mess of it.

HOST: You favor straight popular vote?

GUEST: Oh, I certainly do, yes.

HOST: Andy, as always. [END VIDEOTAPE]

While it is not necessary, in order for them to be idioms, for TV *As Always* and TV *Always* to have lost their (full range of) tense-aspect, (1a-f) suggest that TV *As Always* has. This is common with idioms, and also other verb-based constructions. At the very least, the tense-aspect in the given examples is ambiguous, as well as irrelevant to their function. I expand on this below and argue, further, that reconstructing the tense-aspect in these instances or insisting that it is present is not in keeping with an empirical, cognitive approach.

As stated earlier, formal idioms may allow different lexical items into certain slots of their structure. This is known as being “productive” (Fillmore et al., 1988, p. 507). In contrast, substantive idioms (ibid., p. 505) and less productive formal idioms are more restricted in their lexical content, syntax, and grammar, including tense-aspect. Consider the substantive idiom “kick the bucket”: It can take different tense-aspects (e.g., “he kicked the bucket” still means “he died”), but probably not the passive (“the bucket was kicked”) (Evans & Green, 2006, p. 13), and it cannot take other verbs or nouns.

Many English verb-based constructions, as well, cannot take on the full range of tense-aspects. Adverbial participles, for example (e.g., “waving happily” in “the boy, waving happily, stood on the sidewalk”) take the same form whether the sentence refers to the past, present, or future; only *stood* would change. Infinitive verbal complements (e.g., *to care* in “I want to care”) cannot be marked for tense, though they can be marked for aspect (Celce-Murcia & Larsen-Freeman, 1999). Appositives, as well (which can be seen as reduced forms of non-restrictive

relative clauses (ibid., p. 596)), could be said to lack tense-aspect because they lack even a verb. Consider “never Kathy” in “Katherine, never Kathy, has been working hard” (COCA), which seems to mean “who never goes by Kathy” (present simple) but could also mean “who has never gone by Kathy” (present perfect).

It appears that *as always*, in general, and TV *As Always*, in particular, are similar in this regard: They can be conceived of as being “derived from” a larger phrase with tense-aspect, yet that tense-aspect is not recoverable, and perhaps simply absent. Consider the following three (non-TV) examples of *as always* (all from COCA, Fiction):

- (2a) [Her daughter’s face] was heart shaped, as sweet as always, and was ...
- (2b) The boys followed Drew, as always, with grins of expectation and pride.
- (2c) She made no comment, as always, merely adjusting her own pace and drawing her mantle over her head against the sun.

Reconstructing the tense-aspect of *as always* in (2a) forces one to guess between at least two plausible options, past simple and past perfect (“as sweet as it always was” and “as sweet as it always had been”). And, regarding (2b), one has three options: past simple, present simple, and present perfect. Most interesting of all is (2c), though; the *as always* here seems out of place, preceded by a negative statement one would expect to find accompanied by *never* and *any* (“she never made any comment”). The paraphrase “She made no comment, as she always did” sounds odd and “... as she always didn’t” even odder.⁴⁹ There is another solution that maintains the *as* and *always* in sequence, but it calls for a whole new phrase, i.e., not “as (she) always (did)” but instead “as (was) always (the case)”.

A key assumption of cognitive linguistics is that explanations must be in accord with empirical evidence from other disciplines such as psychology (see the “Cognitive Commitment” in Lakoff, 1990, p. 40). The extensive reconstruction discussed above requires the positing of

⁴⁹ Note that this is where we might expect a negative counterpart of *as always*, i.e., *as never*, to appear: “She made no comment, as (she) never (did)”. Yet, isolated *as never* does not seem to exist as a valid option (though *as never before* is common). Thus, this (like the fact that the opposite of *always* + progressive appears to be *never* + present simple rather than *never* + progressive; see §5.5.4) constitutes an asymmetry regarding *always* and *never*.

multiple invisible underlying forms. Moreover, (2c) should cause us to revisit (2a-b) and wonder if those instances of *as always*, as well, have the underlying form “as was always the case”. Such a hypothesis would be more unified, but is unempirical; there is no easy way to prove these claims about which tense-aspect (and verb) underlies a particular *as always*.

Regarding TV *As Always*, one might argue that the tense-aspect is recoverable after all, and is (highly likely to be) the present simple (“It’s a pleasure to have you,” “It’s a pleasure to be here,” etc.). However, many instances in the corpus of TV *As Always* contain only a name, *as always*, and an expression of thanks or appreciation lacking a grammatical subject (i.e., *thank you* or *appreciate it*). While it does not seem odd to consider “Bob, I thank you, as always” and “Bob, I thank you, as I always do” to be the same thing, essentially, it is less clear if makes sense to think of “thank you, as always” (with no expressed grammatical subject) as a full clause with an elided subject rather than merely as one tenseless idiom (*thank you*, with the pragmatic point/function of thanking someone), followed by another tenseless idiom (*as always*). This is all the more true for *thanks*, which contains neither a subject nor object, nor the first person verb conjugation; *thanks* is technically a noun.

The tense-aspect of *never more so*, as well, can be reconstructed only in an uncertain way: “She also needed to exert control, and never more so than in the calculated way she presented herself to the world” could be construed as past simple or past perfect, and the same is true of *like/as never before*: “The willingness of scientists and policymakers to explore the crisis’s connections as never before ...” could be construed as past or present perfect. Still more ambiguous is *better late than never*, a fully substantive idiom which lacks a verb.

Finally, important from a functional perspective is the fact that not only is the tense-aspect of the given examples ambiguous; it is irrelevant. No one needs to know if a host’s *as always* is in the present simple or present perfect in order to understand that he/she is thanking the guest. If the tense-aspect is absent or ambiguous, and irrelevant, positing its existence is not a stance one should take without strong empirical motivation. Instead, we can simply say that

these are idioms that have become grammatically simplified over time.

6.6 Conclusion

This chapter presents an investigation of the tense-aspect preferences of the adverbs *always* and *never*. The literature claims that these words appear most frequently with the simple and present perfect. The results confirmed that there indeed is a strong association between the past and present simple and *always* (and *never*), and the claim of Celce-Murcia & Larsen-Freeman (1999, p. 509) that the present perfect is common as well. Furthermore, they are aligned with the tense-aspect patterns Biber et al. (1999) present for the Longman Corpus.

However, my study also found a hitherto unreported pattern, whereby the present simple is associated with *always*, and the past simple with *never*. I claimed that this is at least partly explained by the disparate tendencies of *always* and *never* regarding mental verbs, which appeared more frequently with *always*. Because these verbs generally refer to enduring states, they are more likely to be used in present tense. Finally, I argued that TV *As Always*, a particular instance of the general set phrase *as always*, is an idiom and may have—especially when it appears without a grammatical subject, and/or with *thanks*—lost its tense-aspect.

As in the previous chapters, the findings presented here indicate the efficacy of corpus methods in revealing the unique behavioral profiles of related words, in this case in terms of tense-aspect, verb preferences, and idioms. They also highlight, once again, the need for caution in making generalizations. We must not confuse what is typical of one very common lexical item in a class with what is typical of all/most lexical items in that class.

CHAPTER 7: MAXIMIZERS

7.1 Introduction

In this short, final study, I apply techniques similar to those used in the previous chapters to another set of words that refer to extreme ends of a scale, the maximizers *completely*, *totally*, *utterly*, and *fully*. Specifically, I analyze the scalability of their collocates, and their semantic prosody (SP, defined in §1.3 and below). Non-scalar lexical items are those which are extreme (extremity or totality is already inherent in their meaning, as in *obliterated*) or non-gradable (e.g., something is *dead* or *alive*, with no intermediate possibilities)⁵⁰. As noted by Altenberg (1991), both of these regularly appear with maximizers. Given the seeming redundancy or incompatibility inherent to such pairings, this phenomenon merits further investigation. The reason for analyzing SP, as well, is to verify the finding that *utterly* has especially negative SP (Greenbaum, 1970; Partington, 1993, 2004; Louw, 1993; W. Anderson, 2006), and to contribute to the growing body of corpus work on near-synonyms.

The results confirm that maximizers commonly appear with non-scalar collocates of both types, extreme and non-gradable. I explain this by relating the pairing of maximizers and extreme items to affect and emphasis, and the pairing of maximizers and non-gradable items to construal (Langacker, 1987; 2008, etc.). In addition, my study confirms that *utterly* often has negative SP. Finally, and unexpectedly, I found that *fully* is unique in its scarcity of non-scalar collocates and lack of semantically negative collocates.

7.2 Literature Review

Maximizers are words, such *absolutely*, which “denote the upper extreme of [a] scale”, as

⁵⁰ We actually can say things such as “She is a very alive person” or “Rembrandt is very dead” (an observation and examples I owe to Robert Kirsner (p.c.)). However, such cases are not (only) about being literally, biologically alive and breathing, but about something gradable; the former example is about being energetic, and the latter about being dead for a long time, or perhaps emphasizes the permanence and significance of Rembrandt’s condition.

defined by Quirk et al. (1985, p. 590), whose schema for contextualizing maximizers is adopted here (see Fig. 7.1). In this schema, *intensifiers* are words which scale quantities upwards or downwards “from an assumed norm” (ibid., p. 589, see also Bolinger, 1972, p. 17).⁵¹ Those which scale downwards are *downtoners* and those which scale upwards are *amplifiers* (or *amplifying intensifiers*). Amplifiers consist of *maximizers* (which include *always*), and *boosters*, the latter indicating a non-maximal increase in degree.

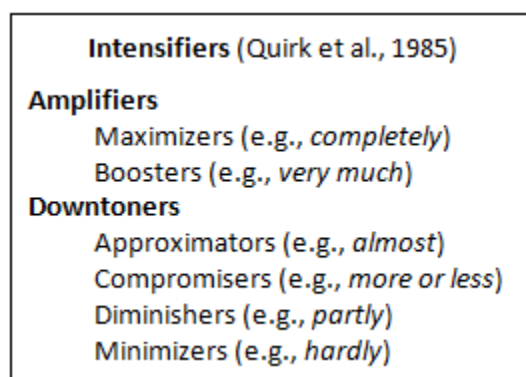


Figure 7.1: Intensifiers

Corpus techniques have been applied to maximizers (and other intensifiers) with useful results, such as the discovery of differences in the collocates and uses of near synonyms, as in Greenbaum (1970), Bäcklund (1973), Partington (1991, 1998, 2004), Louw (1993), Lorenz (1999), W. Anderson (2006), and Tao (2007) (discussed in more detail shortly). Maximizers and intensifiers are also the topic of Kirchner (1955), Greenbaum (1974), Allerton (1987), Altenberg (1991), Tagliamonte & Roberts (2005), and, overall, have been found to be “subject to a number of syntactic, semantic, lexical and stylistic restrictions affecting their use” (Altenberg, 1991, p. 142), and therefore highly amenable to corpus study.

Boosters and most downtowners “typically modify ‘scalar’ items, i.e. items that are fully gradable”, but maximizers “are typically used to modify ‘nonscalar’ items, i.e. items that do not

⁵¹ The term *degree modifiers* (Buchstaller & Traugott, 2006, p. 347) would be more fitting, but is less established. Another option is Bolinger’s *degree words* (1972), widely used to mean the same thing as Quirk et al.’s *intensifiers*. Technically, though, by *degree words* Bolinger means all words that involve degree in any way, including verbs, such as *maul* (p. 29), and nouns. In contrast, most works on the topic focus on adverbial or adjectival intensifiers (which Bolinger calls *degree adjectives*—another unideal term, since it is meant to include adverbs).

normally permit grading (e.g. *empty, impossible, wrong*) or already contain a notion of extreme or absolute degree (e.g. *disgusting, exhausted, huge, marvellous*, etc.)” (Altenberg, 1991, p. 129). This is true, for example, of *absolutely*, whose collocates commonly have “some kind of built-in superlative sense” (Partington, 1998, p. 57; also Partington, 1991).⁵² Maximizers also appear with scalar items, as in the London Lund Corpus (*ibid.*, p. 142; on London Lund, see Svartvik & Quirk, 1980) and COBUILD (Partington, 2001, p. 54; on COBUILD, see Sinclair, 1990).

This phenomenon is counter-intuitive, and very problematic for generative grammar, the approach dominant in linguistics for decades (see §2.1). As Bolinger puts it, maximizers and other degree modifiers are “an antidote to the overconfident description of language as a system” (1972, pp. 18–19), behaving in ways that bend and blur categories and thus showing that language is not “structured in an orderly manner, and reducible to rule” after all (p. 18). Some key assumptions of generative grammar are that (a) language adheres to the reductionist principle of avoiding redundancy, (b) meaning is fully compositional and (c) understood in terms of truth conditions, and (d) semantics and pragmatics are distinct.

If this is correct, one would predict that we use maximizers to add a notion of extremity to that which does not already contain such a notion and is gradable, e.g., “I was utterly disappointed” (as opposed to only slightly, or very). Their appearance with the opposite, i.e., extreme and/or non-gradable items, would appear on the generative view to be redundant and illogical, respectively: Using a maximizer with an extreme word is redundant because it adds no new information, and using a maximizer with a non-gradable word seems illogical due to the clash in semantics, since maximizers refer to the upper end of a scale.

In this study, maximizers were found modifying words such as *devastate*. This word means to “cause (someone) severe and overwhelming shock or grief” (Oxford Dictionaries, n.d., *devastate*, def. 1.1) and its etymology (Latin *devastat-*, from *devastare*) is literally *de-*

⁵² On a related note, in the Scottish Corpus of Texts & Speech (Corbett, 2004) consisting of Broad Scots and Scottish English, maximizers appear consecutively, e.g., “absolutely, totally, filled to the brim” (W. Anderson, 2006, p. 10).

‘thoroughly’ combined with *vastare* ‘lay waste’ (ibid.). It is further defined (ibid., def. 1) as synonymous with *destroy*, which itself can mean “defeat utterly” (Oxford Dictionaries, n.d., *destroy*, def. 1.2) or “reduce (an object) to useless fragments ... beyond repair or renewal; demolish; ruin; annihilate” (Dictionary.com, n.d., *destroy*, def. 1). In a semantic system that views meaning in terms of truth conditions, if the notion of an absolute degree is inherent to *devastate*, any maximizer preceding it adds no information and is therefore redundant.

Note that the redundancy of *totally devastated*, etc., is problematic for generative grammar even if the model can somehow correctly process such phrases. This is because “if claims of psychological reality are taken seriously, questions of economy assume the status of empirical issues, as opposed to methodological” (Langacker, 1988b, p. 129). That is, in generative grammar economy is not merely a preference, but a criterion by which claims are judged (Contini-Morava, 1995, p. 2)—which means redundancy calls the legitimacy of the model itself into question. Moreover, and no matter what one’s view of grammar, the redundancy is interesting in its own right because the reason for it is not immediately clear.

Another characteristic of maximizers I investigate is their semantic prosody (SP), defined as “a form of meaning which is established through the proximity of a consistent series of collocates, often characterisable as positive or negative, and whose primary function is the expression of the attitude of its speaker or writer towards some pragmatic situation” (Louw, 2000, p. 57). For example, the SP of “build up a” is positive (ibid., p. 52) because it is so often followed by nouns referring to positive things, such as *friendship* (p. 78). The study of SP generally calls for quantitative methods, as “they are a collocational phenomenon and [thus] preferably to be regarded as recoverable computationally from large language corpora rather than intuitively” (ibid., pp. 48-49). Though SP has been criticized (see Hunston, 2007; and Stewart, 2009) for being based on intuition at the coding stage, it has proven very valuable nevertheless (as in Berber-Sardinha, 2000, on equivalent items in English and Portuguese), especially on near-synonyms (as in Xiao & McEnery, 2006, and others, discussed below).

Work on near-synonyms includes Partington (1998), which describes differences in the collocates and uses of *sheer*, *complete*, *pure*, and *absolute*, and Partington (1991, 2004), which shows differences between *absolutely*, *perfectly*, *entirely*, *completely*, *thoroughly*, *totally*, and *utterly*. For instance, *absolutely* is the only word in the set that frequently appears in hyperboles (2004, p. 148). Others (Bäcklund, 1973; Lorenz, 1999; Tao, 2007), too, have found *absolutely* to be unique. Another item that differs from its near-synonyms is *utterly*, noted to have especially negative SP (Greenbaum, 1970; Louw, 1993; Partington, 1993, 2004; W. Anderson, 2006). In this study, I analyze the SP of *completely*, *totally*, *utterly*, and *fully*, seeing if the claims about *utterly* are correct and how it compares to the other maximizers. Just as it has been possible, in previous chapters, to find differences between the related but opposite items *always* and *never*, it should be possible to differentiate these four maximizers.

7.3 Method

In this study of *completely*, *totally*, *utterly*, and *fully* I determine, first, how commonly these words' collocates are non-scalar (extreme or non-gradable) and, second, their SP. The corpus used was COCA, in its entirety. The first step was obtaining, for each maximizer, both the 20 strongest (according to Mutual Information, or MI, score) and 20 most frequent collocates. MI score is a measurement of the "strength" of a collocation that takes overall frequencies into consideration (see §3.3 for the formula COCA uses). To weed out anomalously high-ranking collocates, i.e., those whose MI scores are high due only to their being very rare, overall, the minimum frequency was set to ten. The span consisted of one word to the right of the maximizer, and the collocates were restricted to verbs and past participles (achieved by restricting the collocates to verbs).

The advantages of using MI score are discussed in §3.3. Namely, MI scores are useful because they correct for the fact that, if we used only raw frequencies, some collocates would be highly ranked simply because they are common in the language overall. At the same time, due to

the importance of how frequently a token is encountered (because of the resulting degree of “entrenchment” (e.g., Bybee & Slobin, 1982; Langacker, 1987), it is prudent to sort collocates by frequency as well. In obtaining the most frequent collocates, I used the same settings described above (minimum frequency of ten, span of one word to the right, restricted to verbs).

Once the collocates were obtained, I categorized them, first, according to SP, using two categories, Negative or Other (the latter consisting of those which were positive or neutral, combined; this was because the claims being investigated pertain only to negative prosody, and because positive prosody is so rare in language). I coded collocates as Negative if they involved destruction (e.g., *destroy*), other unkind actions (e.g., *ignore*), lack (e.g., *lack*, *gone*), syntactic negation (e.g., *disregard*), and unpleasant states of mind or being (e.g., *devastated*); the rest I coded as Other. Second, I categorized the collocates according to whether they were extreme, i.e., “already contain a notion of extreme or absolute degree” (Altenberg, 1991, p. 129); non-gradable, i.e., “do not normally permit grading” (ibid.), or neither. For example, *obliterated* is extreme because it means something like “completely destroyed”, and *empty* and *gone* are non-gradable because they refer to “all or nothing” situations.

7.4 Results

Non-scalar collocates of both types (extreme and non-gradable, but especially the extreme type) were common with *completely*, *totally*, and *utterly*, and *utterly* did indeed have the most semantically negative collocates. Unexpectedly, though, *fully* stood out even more, both in terms of SP (negative SP was notably absent) and its collocates’ scalability.

7.4.1 Non-Scalar Collocates

As stated above, I collected the 20 strongest and 20 most frequent collocates of each of the four maximizers. Due to the minimum frequency requirement of ten, only 16 collocates of *utterly* were analyzed. I then determined, for each maximizer and type of collocate list, the percentage

of collocates that fell into each category: extreme, non-gradable, or neither. The results are displayed in Figs. 7.2 (strongest collocates) and 7.3 (most frequent collocates).

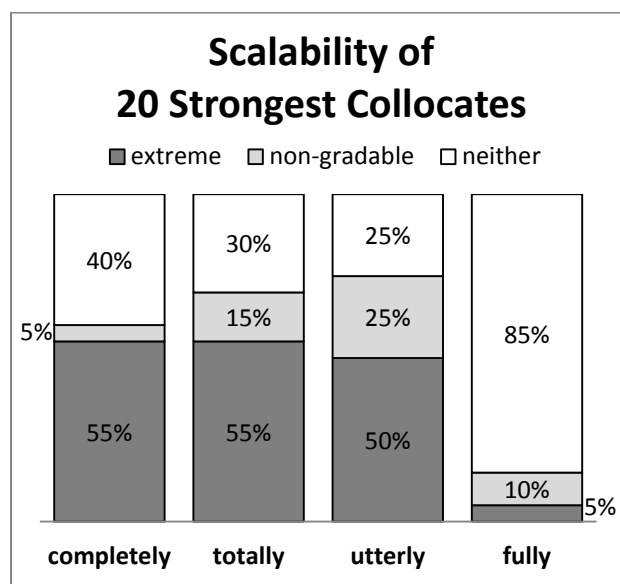


Figure 7.2: Scalability of strongest collocates

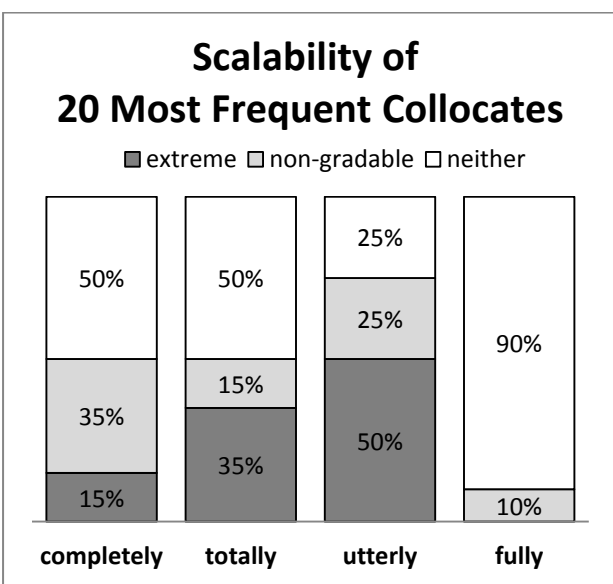


Figure 7.3: Scalability of most frequent collocates

As expected based on the literature, non-scalar collocates of both types were common. Together, they constituted 60% or more of the strongest and 50% or more of the most frequent collocates of *completely*, *totally*, and *utterly*. Specifically, they accounted for 60%, 70%, and 75%, respectively, of those maximizers' strongest collocates and 50%, 50%, and 75% of the most frequent. Extreme non-scalar collocates were more common than the non-gradable type, accounting, alone, for 50% or more of the strongest collocates of those three maximizers. *Fully* was different, though. Non-gradable collocates accounted for only 10% of both its strongest and its most frequent collocates, and extreme collocates accounted for 5% of its strongest collocates and none of its most frequent. This difference is quite striking.

While the results for *completely*, *totally*, and *utterly* are very similar, minor differences can be noted. For example, among the most frequent collocates, the percentage that are extreme is higher (35%) for *totally* than for *completely* (15%), and higher still for *utterly* (50%). Among the strongest collocates, though, this figure is more stable: 55% for *completely* and *totally*, and

50% for *utterly*. Finally, the percentage of non-gradable collocates—whether ranked by strength or frequency—is 15% and 25% for *totally* and *utterly*, respectively, and, for *completely*, varies from 5% (ranked by strength) to 35% (ranked by frequency).

Thus, of the four maximizers investigated, *utterly* most strongly fits the observation that maximizers appear with extreme and non-gradable collocates. The pattern is less robust for *totally* (yet still true at least 50% of the time), and a little weaker, still, for *completely*. With *fully*, on the other hand, these types of collocates are much less common.

7.4.2 Semantic Prosody

Regarding SP, *fully* stood out once again. None of its strongest (see Fig. 7.4) or its most frequent (see Fig. 7.5) collocates had negative SP. In contrast, the collocates of *completely*, *totally*, and *utterly* were quite likely to: 40%, 70%, and 81%, respectively, of the strongest collocates of these words had negative SP, and 60%, 60%, and 81% of the most frequent.

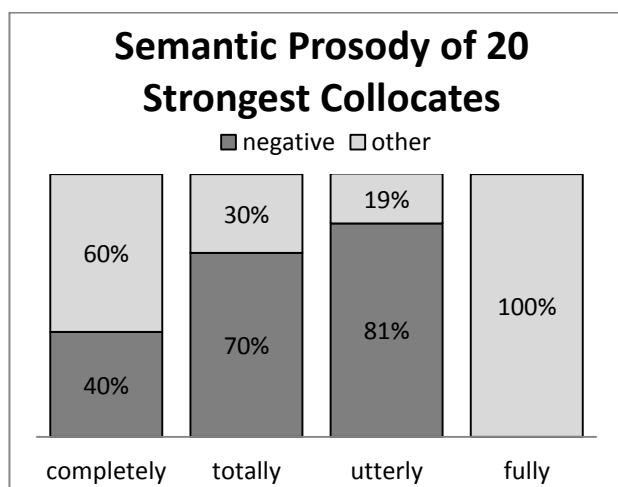


Figure 7.4: SP of strongest collocates

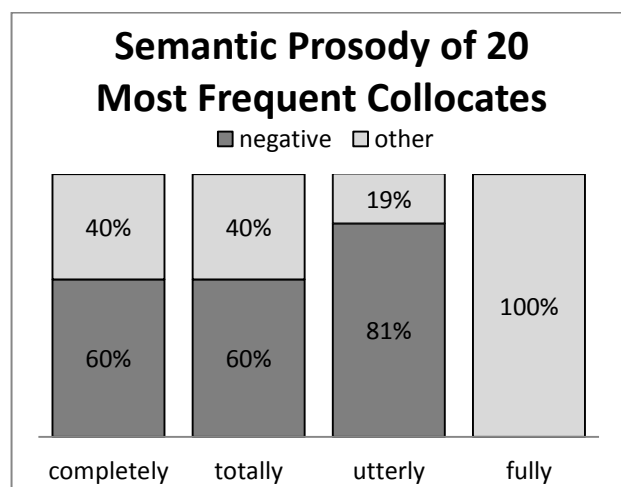


Figure 7.5: SP of most frequent collocates

7.5 Discussion

In this section, after briefly addressing the SP findings, I argue that the pairing of maximizers and non-scalar collocates is best explained from a cognitive perspective: Specifically, I posit that

the pairing of maximizers and extreme items conveys emphasis and affect, and that the pairing of maximizers and non-scalar items is a matter of construal.

7.5.2 *Extreme Collocates and Emphasis*

The pairing of maximizers with extreme words is, rather than being redundant, a means of (a) emphasizing and thereby strengthening the meaning of the extreme word and (b) expressing strong affect, especially negative affect. This phenomenon is, perhaps, best exemplified in the embarrassing stories section of magazines geared toward teenage girls, filled with short tales that end with comments like “I was completely mortified!”

In *completely mortified*, the maximizer strengthens the meaning of *mortified* beyond its normal, seemingly already maximally strong, meaning. And the heightened emotional valence comes from the “metaphorical mapping between the emotional domain and the linguistic” (discussed in §4.5.1), which causes us to associate intense emotions with lexical items indicating extremes (Jing-Schmidt, 2007, p. 433). This also can “boost the speaker’s illocutionary force”, “maximize dramatic effect”, “elicit attention”, and “establish rapport” (ibid., p. 428), goals all particularly relevant in the sharing of embarrassing stories. The reason this emotional valence is usually negative is simply that the collocates of maximizers so often have negative meanings—a pattern consonant with the tendency in language, overall, to express negative emotions more than positive emotions (K. J. Anderson & Leaper, 1998, p. 439).

The above argument is supported by the findings regarding *fully* (see Figs. 7.1 and 7.2): This maximizer was unique among those studied in that none of its collocates had negative meanings and, at the same time, very few (5% of the strongest and none of the most frequent) were extreme. This supports the claimed connection between extreme items (made even more extreme by being paired with a maximizer) and heightened negative affect.

Thus, the (supposed) redundancy is unproblematic for cognitive approaches, which do not rigidly adhere to reductionist principles anyway. More importantly, though, their rejection

of truth condition-based semantics and the semantics-pragmatics distinction leads to a much more inclusive view of meaning, in which lexical items can convey things such as “sensory, emotive, or kinesthetic sensations” and “a person’s awareness of physical, social and linguistic context” (Langacker, 1988a, p. 6). In sum, on this view—and as was just argued—maximizers that appear with extreme words are not even redundant after all.⁵³

That being said, we do avoid redundancy in some senses. For example, “John loves Mary and John protects Mary” becomes “John loves and protects Mary” (Osgood & Richards, 1973, p. 390). Notably, though, the repetition of near-entire clauses can and does occur if the longer (more salient, more marked) version conveys a pragmatic point. “John loves Mary and John protects Mary” might be useful, for example, if one wanted to emphasize all the wonderful, caring things John does for his beloved. From this we see that even what looks like highly uneconomical redundancy can serve useful pragmatic purposes.

7.5.2 *Non-Scalar Collocates and Construal*

This study confirmed that the pairing of maximizers with non-gradable (“all or nothing”) items, such as *gone*, is common. In the strongest case, that of *completely*, 35% of the top 20 most frequent collocates were non-gradable, yet such utterances are anomalous according to the criteria of compositionality and truth-conditional semantics: The semantic content of the maximizer (imposing gradation to the fullest degree) appears to contradict the semantic content of the item it modifies, which is non-gradable. The fact that maximizers also appear with gradable items complicates the matter further, as any characterization of maximizers has to account for that, as well. The notion of construal does precisely that.

Construal is “our manifest capacity for conceptualizing the same situation in alternative ways” (Langacker, 1988a, p. 4; see also Langacker, 1991, 2008, etc.), e.g., telling someone “I will

⁵³ The same is true of supposedly redundant pieces of para-linguistic or non-linguistic aspects of communication. For example, although raising my voice and pointing to my addressee in a two-person conversation as I say *you* adds nothing in terms of identifying the referent, it conveys anger, highlights the act of accusation, and so forth.

come to Chicago” (addressee as point of reference) as opposed to “I will go to Chicago” (speaker as point of reference) (Langacker, 1988c, p. 85). Both utterances have the same truth conditions, but “there is much more to meaning than ... truth conditions” (Langacker, 1988a, p. 32; cf. Diver, 1995, p. 75). Another example is our ability to perceive count nouns as mass nouns and vice versa. When we say “I’ll have two beers,” we are “mentally packaging the referents of mass nouns” into containers, and if we say “There was cat all over the driveway” we are construing the count noun, *cat*, as a mass (Pinker, 2007, p. 170). Altenberg (citing Cruse, 1986, and Allerton, 1987) argues that something similar is true of scalar and non-scalar items, i.e., that “words that are basically scalar can be reinterpreted as non-scalar and vice versa” (Altenberg, 1991, p. 142), and that this is why non-scalar adjectives are so often intensified by maximizers (ibid.)

I, as well, argue that non-gradable items can be construed as gradable, i.e., that the presence of a maximizer can cause us to construe the non-gradable item as possessing the scalar characteristics present in the maximizer. And, building on this further, I argue that the construal can happen in a general way and in a more specific way, involving parts and wholes. Regarding the more general way, consider *gone* and *lose* (as in *lose a game*). Generally, a person can only be *gone* or not *gone*, and can *lose* or *win* a game. However, using these words with a maximizer such as *totally* allows us to construe a person as *gone* in a more serious, permanent, and/or emotionally salient sense (e.g., she is not merely physically absent at the time of the utterance; she took all her belongings, too, and will never, ever return). Likewise, we can construe *lose* as losing to an embarrassing degree.

The second way to construe non-gradable items as gradable is to consider the non-gradable version to apply to a unified whole and the scalar version to apply to its parts. For example, a cake (construed as a unified whole) can only be *gone* or not *gone*. However, since the individual slices need not be in the same location, we can say the cake is *totally gone* when we want to stress that every slice and every last crumb is gone. Similarly, a gun construed as a whole can only be *loaded* (whether it contains one bullet or five) or not *loaded*, but a *fully*

loaded gun has a bullet in every chamber. With *lose*, as well, we have the option of construing the situation in terms of measurable pieces: Ultimately, we can only *win* or *lose* a game, but the difference in the number of points between the winners and losers can vary.

Lindner's (1981) discussion of the particle *up* in its completive sense (pp. 150-175) is also relevant here. Because "the notion of completion is related to the notion of reaching a goal" (p. 150), one aspect of Completive *Up* is "achievement of a goal state" (p. 158). For example, the difference between *harden* and *harden up*, *sweeten* and *sweeten up*, etc., is that something that has *hardened* could potentially harden further, while something that has *hardened up* has reached the maximal and/or desired hardness (pp. 158-159).⁵⁴ Though, in a sense, maximal hardness refers to an objective state, goals and desired states of hardness are concepts which exist only in someone's mind; in terms of observable objects and events, there is no difference between clay that has *hardened* and clay that has *hardened up*. Thus, this appears to be at least partly a matter of construal.⁵⁵ Though they do not involve the same exact construals, the flip in our construal of verbs that Completive *Up* can bring about is somewhat similar to the non-scalar to scalar flip in our construal of verbs that maximizers can evoke.

Another way to reach a goal is "by acting on the entire substance of an object" (Lindner, 1981, p. 150), as with "processes affecting an object's overall form" (p. 157). This relates closely to the parts versus whole type of construal I described earlier. There is a difference, Lindner writes, between *fixing a bike* (so it functions) versus *fixing up a bike* (adjusting various parts so that it is like new again) (ibid.). The two understandings of the bike, in this case, are something like "a whole" which can only be *fixed* or *not fixed* versus "something having many parts (all of which can be affected)". This difference is simultaneously a matter of gradability: By construing something as consisting of many parts, each of which may or may not be affected, we can

⁵⁴ Bolinger says something similar about *up* (1972, p. 30).

⁵⁵ Even in the absence of Completive *Up*, many words are ambiguous with respect to whether they are scalar or non-scalar. Consider sweetness, for example: We often consider foods to fall into two all-or-nothing categories, *sweetened* and *unsweetened*. Other times, we clearly think about degrees of sweetness, e.g., "You need to sweeten this coffee a little more," in which case the *a little more* signals gradability.

construe as gradable what otherwise would not appear to be.

Thus far, I have written as if emphasis and affect are relevant only for extreme words, and construal only for non-gradable words. This distinction should be taken with a grain of salt. Consider *completely obliterated*: One could explain this in terms of emphasis and/or negative affect but also parts and wholes. A car is either *obliterated* or not; it cannot be “partly obliterated.” Its parts, however, can each have their own separate status. Thus, if I say my car is *completely obliterated*, I may be expressing despair, and/or I may be construing the car as a collection of parts rather than a whole. Either or both of the two possibilities may be applicable to a given maximizer + non-scalar item pairing.

7.6 Conclusion

One goal of this chapter was to determine how frequently *totally*, *completely*, *utterly*, and *fully* appear with extreme or otherwise non-gradable collocates, resulting in combinations which seem redundant or illogical. I argued that maximizers that appear with extreme items can strengthen those items even beyond their supposed maximal meaning and/or add an emotional and usually negative valence, and that the appearance of maximizers with non-gradable items involves construal, in either a general (and often emotional) sense or a “part versus whole” sense. The other aspect of the maximizers I analyzed was their SP—which proved relevant to the above claim about emphasis and affect. The commonness, overall, of extreme collocates and collocates with negative/unpleasant meanings supports the claim that when maximizer + non-scalar item pairs contribute emotional valence it is usually negative.

More specifically, the only maximizer of the four, *fully*, with no negative/unpleasant collocates among its 20 strongest or 20 most frequent was also the only one that had almost no extreme collocates. Thus, while much of this chapter focused on the similarities between all or most of the maximizers (namely, their common appearance with items that are non-scalar and negative), the unique profile of *fully* is noteworthy as well. The two findings about *fully* shed

light on the workings of maximizers as a whole, as well as drive home the fact that it should not be assumed that even the closest of synonyms behave in the same way.

More generally speaking, the goal of this chapter was to show that—as with *always* and *never*—with maximizers, corpus techniques illuminate important details of meaning and use, details which cohere with a cognitive-functional approach. In the following chapter, the conclusion, I briefly revisit all four studies presented here and their findings, contextualizing them in terms of the literature and summarizing their significance.

CHAPTER 8: CONCLUSION

In this dissertation, I presented four studies, three on *always* and *never* (on their exaggerated versus literal use, their functions when they appear with the progressive, and their overall tense-aspect preferences), and one on a related phenomenon, maximizers. After a brief summary of the key findings and interpretations, I discuss what we can get out of such a study of multiple phenomena and what the methodology may imply for other, similar studies.

8.1. Summary of Findings and Interpretations

8.1.1 Exaggeration

The first of the four studies, on exaggeration, revealed a connection between less formal genres and exaggerated/non-literal use of *always* and *never*. This distinction existed not only between major genres such as written and spoken language, but also between sub-genres. For example, local news articles contained more exaggeration than international ones, humanities articles more than medical articles, and casual conversation more than unscripted TV/radio dialog. I explained these findings in terms of accountability and the functions associated with particular genres or sub-genres, especially the appropriateness of expressing emotion and/or displeasure. I offered evidence that accuracy is more expected in published, written language, and more in certain of its sub-genres than others, and that exaggeration and the expression of emotion, especially negative emotion, is most strongly associated with casual language.

8.1.2 Functions of *Always* + *Progressive*

The second study, on the functions of *always* + progressive, found that, contrary to the literature, this construction was not most commonly used for emotionally negative functions (Complain, Lament). Instead, its emotionally neutral function, Describe, was far more common.

I argued that this is explained by the fact that *always* is a very generic word, able to appear with any verb, and the fact that its meaning, when combined with that of the present progressive, enables us to do many of the same things we do with the present simple (e.g., to state general, objective facts). At the same time, *always* is quite useful for negative purposes: As part of the *always* + progressive construction, it contributes to complaints being highly effective (because they are exaggerated, emotional, and memorable). Though not as common as neutral functions, the negative functions, Complain and Lament, were more common with *always* + progressive and *never* + progressive than the positive one, Praise. I attributed this fact to the negativity bias in cognition and language, a bias which stems from highly practical concerns.

I also analyzed the effects on function of grammatical subject and genre, and found that *always* + progressive complaints were most common with third person human subjects (then second person, then first person, then non-human). I explained this in terms of politeness theory: Complaining directly to an offender might be effective, but it is also extremely face-threatening; it is safer to complain about people in their absence, thereby avoiding the social consequences of a face-threatening act and still reaping the benefits of seeking empathy and/or engaging in gossip. The genre results, as well, reflect practical considerations. For example, description is more common in academic spoken language because the main purpose of such speech is to explain and teach. I also noted reasons that some spoken genres are, in terms of function, more like written genres, and vice versa. Finally, I noted several differences between the adverbs *always* and *never*, which I discuss below.

8.1.3 Tense-Aspect

The third study, on tense-aspect, confirmed claims in the literature that *always* and *never* appear most often in the simple (both past and present) and present perfect. However, an interesting difference exists: While *always* is more strongly associated with the present simple than with the past (in line with overall tense-aspect patterns of English, as presented in Biber et

al., 1999), the opposite is true of *never*. I argued that this is connected to *always* appearing more often with mental verbs, which are more likely to be in present tense. Last, I discovered two idioms, *TV As Always* and *TV Always* (e.g., “It’s a pleasure to be here Bob, as always” and “Always a pleasure, Bob”) and showed that the former may lack tense-aspect.

8.1.4 Maximizers

The fourth study was on the scalability of the collocates of the maximizers *completely*, *totally*, *utterly*, and *fully*, and on these maximizers’ SP. Others have noted that maximizers appear with non-scalable (extreme or non-gradable) collocates, and that certain maximizers (i.e., *utterly*) have negative SP. I confirmed the commonness of both non-scalable and negative/unpleasant collocates, and explained these patterns from a cognitive perspective. I argued, first, that maximizers that appear with extreme words further strengthen those words and/or add an emotional, typically negative, valence. This claim is supported by the overall commonness, in this study, of extreme and negative collocates, and the fact that the only maximizer with no negative collocates, *fully*, had nearly no extreme collocates, either. Second, I showed that the appearance of maximizers with non-gradable words is an example of construal.

8.2 Implications and Significance

The studies presented in this work span several topics. Nevertheless, they are united at the theoretical level and in terms of recurring explanations for the findings. I begin with broad connections and implications, and then move to more specific ones.

8.2.1 A Cognitive-Functional Perspective

In the broadest sense, this has been an argument for a cognitive-functional approach to linguistics. One key feature of such approaches—and which distinguishes them from generative approaches—is the constant awareness that language is used by people, who have specific

cognitive abilities and limitations, and for particular purposes. In addition, cognitive linguistics holds the view that the meanings of lexical items or constructions are underspecified and sometimes, in large part, determined by things considered within generative approaches to be pragmatic rather than semantic. For example, forms can convey (in addition to their more straightforward meanings such as *dog* = ‘four-legged domesticated canine mammal kept as a pet and/or work animal’) concepts such as emotional stance, point of view, empathy, focus, and so on, and we correctly interpret forms based on immediate linguistic context, genres and the typical purposes and conventions of those genres, social rules and expectations, world knowledge, the relationships between the speakers, and so on.

I argued (in chapter 4), for example, that both accountability (the permanence of the printed word and pressures of the world of publishing) and genre characteristics (e.g., the fact that the expression of emotions is typical in casual conversation), or even factors as specific as reader education level or subject matter, likely affect whether one uses and/or interprets *always* and *never* literally. I also showed (in chapter 5) that the likelihood of *always* + progressive being used to describe, complain about, lament, or praise someone or something is related to genre: In academic spoken language, for example, description is especially common and praise absent, and in casual spoken language complaints and laments, combined, were more common (by a small margin) than in the written data and in all genres combined, and much more common than in academic spoken language.

Finally, in the study of maximizers (in chapter 7), I contrasted a generative grammar approach with a cognitive-functional approach and showed that the appearance of maximizers with non-scalar adjectives is problematic for the former but coherent with the latter, in which meanings can include notions of cognition and perception such as emotions and the selection of one construal of a situation over other possible construals.

In its method, as well, this work is aligned with cognitive-functional approaches. Because they value context as a crucial part of determining meanings, and reject the generative

distinction between “competence” versus (imperfect) “performance”, cognitive linguists characterize language based on authentic, contextualized instances of its use, such as entire conversations and journal articles. The corpora used for the studies here contain, in total, about 455 million words, with roughly 100 million of those coming from spoken language.

Moreover, the claims made here are motivated and corroborated by research in a variety of fields, such as psychology, cognition, and sociology. The negativity bias (relevant in chapters 4, 5, and 7), for example, is a highly generalized and evolutionarily adaptive cognitive tendency that has been studied empirically in the context of social interaction, emotion (anger, fear, disgust), neurology, information processing, facial recognition, learning, memory, and more (Jing-Schmidt, 2007). Politeness theory (see chapters 4 and 5) is borrowed from sociology and anthropology, and even communications research is called upon (in chapter 4), to confirm intuitions about the demographics of readers of different types of news stories.

8.2.2 Furthering Genre Studies

Genre studies have come a long way, beginning with recognizing the value of studying spoken language, and then seeing that spoken and written language are not homogenous categories but, instead, vary greatly depending on factors such as context, purpose, and subject matter (Chafe & Danielewicz, 1987, p. 84). Thus, an immense number of studies have been conducted of genres, e.g., conversation, fiction, news, academic writing, personal letters, and so forth (especially the work of Biber, such as Biber, 2006, Biber et al., 1999, etc.). Moreover, these studies are becoming increasingly specialized, as seen in studies of how academic writing differs by field (e.g., Biber, Conrad, & Reppen, 1998 on ecology and history articles, or any work based on the Michigan Corpus of Upper-level Student Papers (2009), which contains seven types of papers from 16 disciplines, coded for four levels of quality and eight textual features).

In this work (chapters 4 and 5), I noted differences between large categories such as spoken and written language, or newspapers and academic writing, and also between academic

journals in three fields (humanities, medical, and science-technology), three types of written news stories (local, national, and international), and casual conversation, unscripted TV/radio dialog, and academic spoken language. The differences between the three types of spoken language are especially interesting given the methodological question of how similar unscripted TV/radio dialog, which is fairly easily obtainable but not entirely natural, is to spontaneous and fully-natural conversation between acquaintances.

8.2.3 Characterizing Similar Words

Another general contribution of this work has been to show the value of corpus methods in revealing minor and major differences between words in the same class, i.e., antonyms (the adverbs of frequency *always* and *never*) and near-synonyms (the maximizers *completely*, *totally*, *utterly*, and *fully*). For example, although the simple aspect is common with both *always* and *never* (see chapter 6), the past simple is more strongly associated with *never* and the present simple with *always*. And when ExagQ+ and ExagQ- (see chapter 7), which are derived from *always* and *never* tokens, respectively, were calculated for various (sub-)genres, the resulting rankings were nearly identical but, at the same time, differed dramatically in that, for many (sub-)genres, ExagQ- was as much as 16 or even 18 times as great as ExagQ+.

Several differences between *always* and *never* when followed by the progressive aspect were identified as well (in chapter 5). First, the *go*-future and progressive future were prevalent in the *never* data but not the *always* data. This is related to a second difference, which is that nearly a third of the *never* + progressive tokens served a “to make a vow” function, while none of the *always* + progressive tokens did. Third, the adverbs differed regarding complaints and laments (in the *always* data, complaints were more than twice as common as laments while, in the *never* data, laments were 3.5 times as common as complaints). Fourth and finally, they differed regarding the prevalence of certain verbs: *coming/going* appeared in about 50% of the *never* + progressive tokens but in only 3.4% of the *always* + progressive tokens.

Differences were found, also, between the near-synonymous maximizers *completely*, *totally*, *utterly*, and *fully* (chapter 7). Namely, *fully* stood out in that none of its twenty strongest or twenty most frequent collocates had negative/unpleasant meanings (this figure was between 40% and 81% for the other maximizers). Moreover, only 15% of its strongest and 10% of its most frequent collocates were non-scalar (while these figures ranged from 40% to 75% for the other maximizers), and none of its most frequent collocates were non-gradable.

This work builds on similar corpus research (e.g., Partington, 2004; Tao, 2007), who show that even near-synonyms differ in meanings, usage patterns, and functions. As valuable and necessary as generalizations are, they must be made carefully. In characterizing categories of words, we should ensure that those characterizations are true of all or at least the majority of the most representative words in those categories, lest we learn that one or more has a very different profile. With antonyms, in particular, this is quite likely.

8.3 Further Research

Involving a network of methods, topics, and genres, this work could be expanded in a number of ways, so here I offer just three suggestions. First, Exaggeration Quotients, obtained via a very straightforward, quantitative method, and thus suitable for use on even the largest of data sets, could be calculated for other corpora, genres, and sub-genres, and compared to other, less automated measures of exaggeration or non-literal meaning (e.g., as found in Claridge, 2011). In addition, the formula for calculating Exaggeration Quotients could be enriched, such as by including more adverbs of frequency, and the qualifier *almost*. Or, the concept could be adapted for use with lexical items other than just adverbs of frequency.

Second, one could look into additional ways in which the spoken portion of COCA (unscripted TV/radio show dialog) differs from fully casual spoken language, both in terms of overall frequencies of particular tokens or constructions and in terms of the texts as a whole. For example, the interactions could be dissected into their typical stages and discourse moves. Close

analysis of televised interactions is nothing new, but explicit comparison of phenomena such as narratives, praise, greetings, closings, arguments, etc., in TV/radio shows with their equivalents in casual, spontaneously-occurring conversation would be.

Third, it would be helpful, in order to carefully evaluate the interpretations offered here of the results regarding maximizers, to conduct for tokens of maximizers appearing with non-scalar collocates a qualitative analysis similar to that conducted regarding the functions of *always/never* followed by the progressive. In this way, we could seek evidence that they are truly used for emphasis, to express emotion, and/or to evoke certain construals.

8.4 Closing Remarks

In a way, this has all been an ode to “common words,” which Sinclair (1999, p. 158) tells us are understudied despite their prevalence, and are much more than what they appear to be. On a similar note, Pennebaker (2011) writes that “pronouns, articles, prepositions, and a handful of other small, stealthy words ... make up almost 60 percent of the words [we] use” (p. ix) and serve as “powerful tools to excavate people’s thoughts, feelings, motivations, and connections with others” (p. xi). Thus, we study common words to learn not just about the words in question, but about cognition and communication itself.

WORKS CITED

- Allerton, D. J. (1987). English intensifiers and their idiosyncrasies. In R. Steele & T. Threadgold (Eds.), *Language topics: Essays in honour of Michael Halliday* (Vol. 2, pp. 15–31). Amsterdam, The Netherlands: John Benjamins.
- Altenberg, B. (1991). Amplifier collocations in spoken English. In S. Johansson & A.-B. Stenström (Eds.), *English computer corpora: Selected papers and research guide* (pp. 127–148). Berlin, Germany: Mouton de Gruyter.
- Always*, adv. [Defs. 1, 3]. (2015). *Oxford English Dictionary Online*. Oxford, England: Oxford University Press. Retrieved March 31, 2015 from <http://www.oed.com/view/Entry/5941>.
- Anderson, K. J., & Leaper, C. (1998). Emotion talk between same- and mixed-gender friends: Form and function. *Journal of Language and Social Psychology*, 17(4), 419–448.
- Anderson, W. (2006). “Absolutely, totally, filled to the brim with the Famous Grouse”: Intensifying adverbs in the Scottish Corpus of Texts and Speech. *English Today*, 22(3), 10–16.
- Anthony, L. (2011). AntConc (Version 3.3.0). Tokyo, Japan: Waseda University. Retrieved from http://www.antlab.sci.waseda.ac.jp/antconc_index.html
- Atkinson, J. M., & Heritage, J. (1984). *Structures of social action: Studies in conversation analysis*. Cambridge, England: Cambridge University Press.
- Austin, J. L. (1962). *How to do things with words*. Oxford, England: Clarendon Press.
- Bäcklund, U. (1973). *The collocation of adverbs of degree in English*. Stockholm, Sweden: Almqvist & Wiksell.
- Baker-Brown, G., Ballard, E. J., Bluck, S., De Vries, B., Suedfeld, P., & Tetlock, P. E. (1990). *Coding manual for conceptual/integrative complexity*. Unpublished manuscript, University of British Columbia, Vancouver, Canada and University of California, Berkeley.
- Baumeister, R. F., Bratslavsky, E., Finkenauer, C., & Vohs, K. D. (2001). Bad is stronger than good. *Review of General Psychology*, 5(4), 323–370.
- Beersma, B., & Van Kleef, G. A. (2012). Why people gossip: An empirical analysis of social motives, antecedents, and consequences. *Journal of Applied Social Psychology*, 42(11), 2640–2670.
- Benz, J. K., Tompson, T. N., & Rosenstiel, T. (2014). *The personal news cycle*. Arlington, VA: American Press Institute & Chicago, IL: Associated Press-NORC Center for Public Affairs Research. Retrieved from <http://www.mediainsight.org/Pages/the-personal-news-cycle.aspx>
- Berber-Sardinha, T. (2000). Semantic prosodies in English and Portuguese: A contrastive study. *Cuadernos de Filología Inglesa*, 9(1), 93–110.

- Besnier, N. (1993). Reported speech and affect on Nukulaelae Atoll. In J. H. Hill & J. T. Irvine (Eds.), *Responsibility and evidence in oral discourse* (pp. 161–181). Cambridge, England: Cambridge University Press.
- Biber, D. (1988). *Variation across speech and writing*. Cambridge, England: Cambridge University Press.
- Biber, D. (2006). *University language: A corpus-based study of spoken and written registers*. Amsterdam, The Netherlands: John Benjamins.
- Biber, D., Conrad, S., & Reppen, R. (1998). *Corpus linguistics: Investigating language structure and use*. Cambridge, England: Cambridge University Press.
- Biber, D., Johansson, S., Conrad, S., & Finegan, E. (1999). *The Longman grammar of spoken and written English*. Harlow, England: Pearson Education.
- Bolinger, D. (1972). *Degree words*. The Hague, The Netherlands: Mouton.
- Boucher, J., & Osgood, C. E. (1969). The Pollyanna Hypothesis. *Journal of Verbal Learning and Verbal Behavior*, 8(1), 1–8.
- Brinton, L. J. (2000). *The structure of modern English*. Amsterdam, The Netherlands: John Benjamins.
- British National Corpus. (2001). (Version 2, BNC World). Oxford, England: Distributed by Oxford University Computing Services on behalf of the BNC Consortium. Retrieved from <http://www.natcorp.ox.ac.uk/>
- Brown, K. (Ed.). (2006). *Encyclopedia of language and linguistics* (2nd ed.). Oxford, England: Elsevier.
- Brown, P., & Levinson, S. C. (1978). Universals in language usage: Politeness phenomena. In E. N. Goody (Ed.), *Questions and politeness: Strategies in social interaction* (pp. 56–311). Cambridge, England: Cambridge University Press.
- Brown, P., & Levinson, S. C. (1987). *Politeness: Some universals in language usage*. Cambridge, England: Cambridge University Press.
- Brugman, C. M. (1988). *The syntax and semantics of “have” and its complements*. (Unpublished doctoral dissertation). University of California, Berkeley.
- Buchstaller, I., & Traugott, E. C. (2006). *The lady was al demonyak*: Historical aspects of adverb *all*. *English Language and Linguistics*, 10(2), 345–370.
- Bybee, J. (2001). Frequency effects on French liaison. In J. Bybee & P. J. Hopper (Eds.), *Frequency and the emergence of linguistic structure* (pp. 331–360). Amsterdam, The Netherlands: John Benjamins.
- Bybee, J. (2002). Word frequency and context of use in the lexical diffusion of phonetically conditioned sound change. *Language Variation and Change*, 14, 261–290.

- Bybee, J. (2003). Mechanisms of change in grammaticalization: The role of frequency. In B. D. Joseph & R. D. Janda (Eds.), *The handbook of historical linguistics* (pp. 602–623). Oxford, England: Blackwell.
- Bybee, J. (2010). *Language, usage and cognition*. Cambridge, England: Cambridge University Press.
- Bybee, J., Perkins, R., & Pagliuca, W. (1994). *The evolution of grammar: Tense, aspect and modality in the languages of the world*. Chicago, IL: University of Chicago Press.
- Bybee, J., & Slobin, D. (1982). Rules and schemas in the development and use of the English present tense. *Language*, 58, 265–289.
- Bybee, J., & Thompson, S. (1997). Three frequency effects in syntax. In *Proceedings of the Twenty-Third Annual Meeting of the Berkeley Linguistics Society: General session and parasession on pragmatics and grammatical structure* (pp. 378–388). Berkeley, CA: Berkeley Linguistics Society.
- Cacioppo, J. T., Gardner, W. L., & Berntson, G. G. (1999). The affect system has parallel and integrative processing components: Form follows function. *Journal of Personality and Social Psychology*, 76(5), 839–855.
- Carter, R., & McCarthy, M. (2006). *Cambridge grammar of English*. Cambridge, England: Cambridge University Press.
- Cazes, J. (1983). Lest we exaggerate. *Journal of Liquid Chromatography*, 6(9), 1557–1558.
- Celce-Murcia, M., & Larsen-Freeman, D. (1999). *The grammar book: An ESL/EFL teacher's course* (2nd ed.). Boston, MA: Heinle & Heinle.
- Chafe, W., & Danielewicz, J. (1987). Properties of spoken and written language. In R. Horowitz & S. J. Samuels (Eds.), *Comprehending oral and written language* (pp. 83–113). New York, NY: Academic Press.
- Claridge, C. (2011). *Hyperbole in English: A corpus-based study of exaggeration*. Cambridge, England: Cambridge University Press.
- Coates, D., & Winston, T. (1983). Counteracting the deviance of depression: Peer support groups for victims. *Journal of Social Issues*, 39(2), 169–194.
- Cohen, S., & Janicki-Deverts, D. (2009). Can we improve our physical health by altering our social networks? *Perspectives on Psychological Science*, 4(4), 375–378.
- Conrad, S., & Biber, D. (2000). Adverbial marking of stance in speech and writing. In S. Hunston & G. Thompson (Eds.), *Evaluation in text: Authorial stance and the construction of discourse* (pp. 56–73). Oxford, England: Oxford University Press.
- Contini-Morava, E. (1991). Deictic explicitness and event continuity in Swahili. *Lingua*, 83, 277–318.

- Contini-Morava, E. (1995). Introduction: On linguistic sign theory. In E. Contini-Morava & B. S. Goldberg (Eds.), *Meaning as explanation: Advances in linguistic sign theory* (pp. 1–39). Berlin, Germany: Mouton de Gruyter.
- Cook-Gumperz, J., & Gumperz, J. J. (1981). From oral to written culture: The transition to literacy. In M. F. Whiteman (Ed.), *Writing: The nature, development, and teaching of written communication* (Vol. 1: Variation in writing: Functional and linguistic-cultural differences, pp. 89–109). Mahwah, NJ: Lawrence Erlbaum.
- Corbett, J. (2004). *Scottish Corpus of Texts & Speech*. University of Glasgow, Scotland. Available online at www.scottishcorpus.ac.uk.
- Couper-Kuhlen, E., & Selting, M. (1996). Towards an interactional perspective on prosody and a prosodic perspective on interaction. In E. Couper-Kuhlen & M. Selting (Eds.), *Prosody in conversation* (pp. 11–56). Cambridge, England: Cambridge University Press.
- Cruse, D. A. (1986). *Lexical semantics*. Cambridge, England: Cambridge University Press.
- Davies, M. (2008-). *The Corpus of Contemporary American English: 450 million words, 1990-present*. Available online at <http://corpus.byu.edu/coca/>.
- Davies, M. (2010-). *The Corpus of Historical American English: 400 million words, 1810-2009*. Available online at <http://corpus.byu.edu/coha/>.
- Davies, M. (2011-). *Google Books Corpus*. (Based on Google Books n-grams). Available online at <http://googlebooks.byu.edu/>.
- Davies, M. (2013). *Corpus of Global Web-Based English: 1.9 billion words from speakers in 20 countries*. Available online at <http://corpus2.byu.edu/glowbe/>.
- Davis, J. (1995). Italian pronouns and the virtue of relative meaninglessness. In E. Contini-Morava & B. S. Goldberg (Eds.), *Meaning as explanation: Advances in linguistic sign theory* (pp. 423–440). Berlin, Germany: Mouton de Gruyter.
- Davis, J. (2002). Rethinking the place of statistics in Columbia School analysis. In W. H. Reid, R. Otheguy, & N. Stern (Eds.), *Signal, meaning and message: Perspectives on sign-based linguistics* (pp. 65–90). Amsterdam, The Netherlands: John Benjamins.
- Davis, J. (2004). Revisiting the gap between meaning and message. In R. Kirsner, E. Contini-Morava, & B. Rodriguez-Bachiller (Eds.), *Cognitive and communicative approaches to linguistic analysis* (pp. 155–176). Amsterdam, The Netherlands: John Benjamins.
- Deese, J. E. (1964). The associative structure of some common English adjectives. *Journal of Verbal Learning and Verbal Behavior*, 3, 347–357.
- Denham, K., & Lobeck, A. (2012). *Linguistics for everyone: An introduction* (2nd ed.). New York, NY: Cengage Learning.
- Destroy [Def. 1]. (n.d.). *Dictionary.com Unabridged*. Retrieved March 11, 2015 from <http://dictionary.reference.com/browse/destroy?s=t>.

- Destroy* [Def. 1.2]. (n.d.). *Oxford Dictionaries*. Retrieved March 11, 2015 from www.oxforddictionaries.com/us/definition/american_english/destroy#destroy__4.
- Devastate* [Defs. 1 and 1.1]. (n.d.). *Oxford Dictionaries*. Retrieved March 11, 2015 from www.oxforddictionaries.com/us/definition/american_english/devastate.
- Diver, W. (1987). The dual. In *Columbia University working papers in linguistics* (Vol. 8, pp. 100–114). Columbia, NY: Columbia University.
- Diver, W. (1995). Theory. In E. Contini-Morava & B. S. Goldberg (Eds.), *Meaning as explanation: Advances in linguistic sign theory* (pp. 43–114). Berlin, Germany: Mouton de Gruyter.
- Du Bois, J. W. (2007). The stance triangle. In R. Englebretson (Ed.), *Stancetaking in discourse: Subjectivity, evaluation, interaction* (Vol. 164 of *Pragmatics and Beyond*, pp. 139–182). Amsterdam, The Netherlands: John Benjamins.
- Du Bois, J. W., Chafe, W. L., Meyer, C., Thompson, S. A., Englebretson, R., & Martey, N. (2000–2005). *Santa Barbara Corpus of Spoken American English*, Parts 1–4. Philadelphia: Linguistic Data Consortium.
- Dunbar, R. I. M. (1993). Coevolution of neocortical size, group size and language in humans. *Behavioral and Brain Sciences*, 16, 681–735.
- Dunbar, R. I. M. (1996). *Grooming, gossip, and the evolution of language*. Cambridge, MA: Cambridge University Press.
- Dunbar, R. I. M. (2004). Gossip in evolutionary perspective. *Review of General Psychology*, 8, 100–110.
- Dunbar, R. I. M., Duncan, N. D. C., & Marriott, A. (1997). Human conversational behavior. *Human Nature*, 8, 231–246.
- Dutta-Bergman, M. J. (2004). Complementarity in consumption of news types across traditional and new media. *Journal of Broadcasting & Electronic Media*, 48(1), 41–60.
- Emler, N. (1994). Gossip, reputation, and social adaptation. In R. F. Goodman & A. Ben-Ze'ev (Eds.), *Good gossip* (pp. 119–140). Lawrence, KS: University Press of Kansas.
- Evans, V., & Green, M. (2006). *Cognitive linguistics: An introduction*. Mahwah, NJ: Lawrence Erlbaum.
- Fellbaum, C. (1995). Co-occurrence and antonymy. *International Journal of Lexicography*, 8(4), 281–303.
- Fillmore, C. J., & Kay, P. (1993). *Construction grammar*. Unpublished manuscript, University of California, Berkeley.
- Fillmore, C. J., Kay, P., & O'Connor, M. C. (1988). Regularity and idiomaticity in grammatical constructions: The case of *let alone*. *Language*, 64(3), 501–538.

- Firth, J. R. (1968). A synopsis of linguistic theory 1930-1955. In F. Palmer (Ed.), *Selected papers of J. R. Firth 1952-59* (pp. 168-205). Harlow, England: Longman. (Original work published in F. R. Palmer (Ed.). (1957). *Studies in linguistic analysis: Special volume of the Philological Society* (pp. 1-32). Oxford, England: Blackwell).
- Francis, W. N., & Kučera, H. (1964). *Brown Corpus: A standard corpus of present-day edited American English*. Providence, RI: Department of Linguistics, Brown University. Available online at <http://icame.uib.no/brown/bcm.html>.
- Giardini, F., & Conte, R. (2012). Gossip for social control in natural and artificial societies. *Simulation: Transactions of the Society for Modeling and Simulation International*, 88(1), 18-32.
- Givón, T. (1973). The time-axis phenomenon. *Language*, 49, 890-925.
- Givón, T. (1979). *On understanding grammar*. New York, NY: Academic Press.
- Godfrey, J. J., Holliman, E. C., & McDaniel, J. (1992). SWITCHBOARD: Telephone speech corpus for research and development. In *The Institute of Electrical and Electronics Engineers International Conference on Acoustics, Speech and Signal Processing* (Vol. 1, pp. 517-520).
- Goldberg, A. E. (1995). *Constructions: A construction grammar approach to argument structure*. Chicago, IL: University of Chicago Press.
- Greenbaum, S. (1970). *Verb-intensifier collocations in English: An experimental approach*. The Hague, The Netherlands: Mouton.
- Greenbaum, S. (1974). Some verb-intensifier collocations in American and British English. *American Speech*, 49, 79-89.
- Gries, S. T., & Otani, N. (2010). Behaviorial profiles: A corpus-based perspective on synonym and antonymy. *International Computer Archive of Modern and Medieval English*, 34.
- Grosser, T. J., Lopez-Kidwell, V., & Labianca, G. (2010). A social network analysis of positive and negative gossip in organizational life. *Group & Organization Management*, 35(2), 177-212.
- Halliday, M. A. K. (1985). *An introduction to functional grammar*. London, England: Arnold.
- Halliday, M. A. K., & Matthiessen, C. M. I. M. (2004). *An introduction to functional grammar* (3rd ed.). Oxford, England: Routledge.
- Hardy, J. A., & Römer, U. (2013). Revealing disciplinary variation in student writing: A multi-dimensional analysis of the Michigan Corpus of Upper-level Student Papers (MICUSP). *Corpora*, 8(2), 183-207.
- Hartung, W. (1996). Die Bearbeitung von Perspektiven-Divergenzen durch das Ausdrücken von Gereiztheit. In W. Kallmeyer (Ed.), *Gesprächsrhetorik* (pp. 118-189). Tübingen, Germany: Narr.

- Heim, I. (1982). *The semantics of definite and indefinite noun phrases*. (Unpublished doctoral dissertation). University of Massachusetts, Amherst.
- Heim, I. (1983). File change semantics and the familiarity theory of definiteness. In M. Barlow, D. Flickinger, & M. Wescoat (Eds.), *Proceedings of the Second West Coast Conference on Formal Linguistics* (pp. 164–190). Stanford, CA: Stanford University Linguistics Department.
- Heim, I. (1990). E-type pronouns and donkey anaphora. *Linguistics and Philosophy*, 13, 137–177.
- Heine, B., & Kuteva, T. (2002). *World lexicon of grammaticalization*. Cambridge, England: Cambridge University Press.
- Hopper, P. J. (1991). On some principles of grammaticalization. In E. C. Traugott & B. Heine (Eds.), *Approaches to grammaticalization* (Vol. 1, pp. 17–37). Amsterdam, The Netherlands: John Benjamins.
- Huffman, A. (2001). The linguistics of William Diver and the Columbia school. *Word*, 52, 29–68.
- Hunston, S. (2007). Semantic prosody revisited. *International Journal of Corpus Linguistics*, 12(2), 249–268.
- Iyengar, S., Hahn, K. S., Bonfadelli, H., & Marr, M. (2009). “Dark areas of ignorance” revisited: Comparing international affairs knowledge in Switzerland and the United States. *Communication Research*, 36(3), 341–358.
- Jing-Schmidt, Z. (2007). Negativity bias in language: A cognitive-affective model of emotive intensifiers. *Cognitive Linguistics*, 18(3), 417–443.
- Jones, S. (2002). *Antonymy: A corpus-based perspective*. New York, NY: Routledge.
- Jones, S. (2006). A lexico-syntactic analysis of antonym co-occurrence in spoken English. *Text & Talk*, 26(2), 127–244.
- Jones, S. (2007). “Opposites” in discourse: A comparison of antonym use across four domains. *Journal of Pragmatics*, 39(6), 1105–1119.
- Jones, S., & Murphy, M. L. (2005). Using corpora to investigate antonym acquisition. *International Journal of Corpus Linguistics*, 10(3), 401–422.
- Jones, S., Murphy, M. L., Paradis, C., & Willners, C. (2012). *Antonyms in English: Construals, constructions and canonicity*. Cambridge, England: Cambridge University Press.
- Jones, S., Paradis, C., Murphy, M. L., & Willners, C. (2007). Googling for opposites: A web-based study of antonym canonicity. *Corpora*, 2(2), 129–154.
- Justeson, J. S., & Katz, S. M. (1991). Co-occurrences of antonymous adjectives and their contexts. *Computational Linguistics*, 17(1), 1–19.

- Keage, H. A. D., Muniz, G., Kurylowicz, L., Van Hooff, M., Clark, L., Searle, A. K., ... McFarlane, A. (2014). Age 7 intelligence and paternal education appear best predictors of educational attainment: The Port Pirie Cohort Study. *Australian Journal of Psychology*, 1–9.
- Kearns, K. (2000). *Semantics*. New York, NY: St. Martin's Press.
- Keenan, E. L. (1971). Quantifier structures in English. *Foundations of Language*, 7(2), 255–284.
- Keenan, E. L. (2006). Quantifiers: Semantics. In K. Brown (Ed.), *Encyclopedia of language and linguistics* (2nd ed., Vol. 10, pp. 302–308). Oxford, England: Elsevier.
- Kempson, R. M. (1977). *Semantic theory*. Cambridge, England: Cambridge University Press.
- Kilgariff, A. (1996). Why chi-square doesn't work, and an improved LOB-Brown comparison. *Association for Literary and Linguistic Computing-Association for Computers and the Humanities Conference*, June 1996, Bergen, Norway.
- Kirchner, G. (1955). *Gradadverbien: Restriktiva und Verwandtest im heutigen Englisch*. Halle, Germany: Max Neimeyer Verlag.
- Kirsner, R. S. (2002). The future of minimalist linguistics in a maximalist world. In W. H. Reid, R. Otheguy, & N. Stern (Eds.), *Signal, meaning and message: Perspectives on sign-based linguistics* (pp. 339–372). Amsterdam, The Netherlands: John Benjamins.
- Kirsner, R. S. (2011). Instructional meanings, iconicity and l'arbitraire du signe in the analysis of the Afrikaans demonstratives. In B. de Jonge & Y. Tobin (Eds.), *Linguistic theory and empirical evidence* (pp. 97–137). Amsterdam, The Netherlands: John Benjamins.
- Kniffin, K. M., & Wilson, D. S. (2010). Evolutionary perspectives on workplace gossip: Why and how gossip can serve groups. *Group & Organization Management*, 35(2), 150–176.
- Lakoff, G. (1987). *Women, fire and dangerous things: What categories reveal about the mind*. Chicago, IL: University of Chicago Press.
- Lakoff, G. (1990). The invariance hypothesis: Is abstract reason based on image-schemas? *Cognitive Linguistics*, 1(1), 39–74.
- Lambrecht, K. (1994). *Information structure and sentence form: A theory of topic, focus, and the mental representation of discourse referents*. Cambridge, England: Cambridge University Press.
- Langacker, R. W. (1987). *Foundations of cognitive grammar* (Vol. 1: Theoretical Prerequisites). Stanford, CA: Stanford University Press.
- Langacker, R. W. (1988a). An overview of cognitive grammar. In B. Rudzka-Ostyn (Ed.), *Topics in cognitive linguistics* (pp. 3–48). Amsterdam, The Netherlands: John Benjamins.
- Langacker, R. W. (1988b). A usage-based model. In B. Rudzka-Ostyn (Ed.), *Topics in cognitive linguistics* (pp. 127–161). Amsterdam, The Netherlands: John Benjamins.

- Langacker, R. W. (1988c). A view of linguistic semantics. In B. Rudzka-Ostyn (Ed.), *Topics in cognitive linguistics* (pp. 49–90). Amsterdam, The Netherlands: John Benjamins.
- Langacker, R. W. (1991). *Foundations of cognitive grammar*. (Vol. 2: Descriptive application). Stanford, CA: Stanford University Press.
- Langacker, R. W. (2008). *Cognitive grammar: A basic introduction*. Oxford, England: Oxford University Press.
- Leech, G. (1974). *Semantics*. Harmondsworth, England: Penguin.
- Leech, G., Garside, R., & Bryant, M. (1994). CLAWS4: The tagging of the British National Corpus. In *Proceedings of the 15th International Conference on Computational Linguistics* (pp. 662–628). Kyoto, Japan.
- Legitt, J. S., & Gibbs, R. W. (2000). Emotional reactions to verbal irony. *Discourse Processes*, 29(1), 1–24.
- Lewis, D. (2002). Adverbs of quantification. In P. Portner & B. Partee (Eds.), *Formal semantics: The essential readings* (pp. 178–188). Oxford, England: Blackwell. (Original work published in E. Keenan (Ed.). (1975). *Formal semantics of natural language: Papers from a colloquium sponsored by the King's College Research Centre, Cambridge* (pp. 3–15). Cambridge, England: Cambridge University Press.
- Lindner, S. J. (1981). *A lexico-semantic analysis of English verb particle constructions with 'out' and 'up'*. (Unpublished doctoral dissertation). University of California, San Diego.
- Li, D., Zhang, X., Luo, D., & Wu, X. (2014). Learning the taxonomy of function words for parsing. In *Proceedings of COLING 2014, the 25th international conference on computational linguistics: Technical papers* (pp. 1886–1896). Dublin, Ireland.
- Linell, P. (2005). *The written language bias in linguistics: Its nature, origins and transformations*. London, England: Routledge.
- Link, K. E., & Kreuz, R. J. (2005). Do men and women differ in their use of nonliteral language when they talk about emotions? In H. L. Colston & A. N. Katz (Eds.), *Figurative language comprehension. Social and cultural influences* (pp. 153–179). Mahwah, NJ: Lawrence Erlbaum.
- Lorenz, G. (1999). *Adjective intensification—learners versus native speakers: A corpus study of argumentative writing*. Amsterdam, The Netherlands: Rodopi.
- Louw, B. (1993). Irony in the text or insincerity in the writer: The diagnostic potential of semantic prosody. In M. Baker, G. Francis, & E. Tognini-Bonelli (Eds.), *Text and technology: In honour of John Sinclair* (pp. 157–174). Amsterdam, The Netherlands: John Benjamins.
- Louw, B. (2000). Contextual prosodic theory: Bringing semantic prosodies to life. In C. Heffer, H. Sauntson, & G. Fox (Eds.), *Words in context: A tribute to John Sinclair on his retirement*. Birmingham, England: University of Birmingham.

- Matlin, M. W., & Stang, D. J. (1978). *The Pollyanna principle: Selectivity in language, memory, and thought*. Cambridge, MA: Schenkman.
- McEnery, A., & Wilson, A. (1996). *Corpus linguistics* (2nd ed.). Edinburgh, Scotland: Edinburgh University Press.
- McGlone, M. S., & Reed, A. B. (1998). Anchoring in the interpretation of probability expressions. *Journal of Pragmatics*, 30(6), 723–733.
- Mettinger, A. (1994). *Aspects of semantic opposition in English*. New York, NY: Oxford University Press.
- Michigan Corpus of Upper-level Student Papers*. (2009). Ann Arbor, MI: The Regents of the University of Michigan.
- Muehleisen, V. (1997). *Antonymy and semantic range in English*. (Unpublished doctoral dissertation). Northwestern University, Evanston, IL.
- Murphy, M. L. (1994). *In opposition to an organized lexicon: Pragmatic principles and lexical semantic relations*. (Unpublished doctoral dissertation). University of Illinois at Urbana-Champaign.
- Murphy, M. L. (2003). *Semantic relations and the lexicon: Antonyms, synonyms and other semantic paradigms*. Cambridge, England: Cambridge University Press.
- Murphy, M. L. (2006). Antonyms as lexical constructions: Or, why paradigmatic construction is not an oxymoron. *Constructions*, SV1–8/2006 (available at www.constructionsonline.de).
- Nakao, M. A., & Axelrod, S. (1983). Numbers are better than words: Verbal specifications of frequency have no place in medicine. *The American Journal of Medicine*, 74(6), 1061–1065.
- Never mind*. (2015). *Merriam-Webster*. Retrieved April 7, 2015 from www.merriam-webster.com/dictionary/never_mind.
- Nevile, M., & Rendle-Short, J. (2007). Language as action. *Australian Review of Applied Linguistics* (special thematic issue *Language as action: Australian studies in conversation analysis*, edited by J. Rendle-Short & M. Nevile), 30(3), 30.1–30.13.
- Ochs, E., & Schieffelin, B. (1989). Language has a heart. *Text*, 9(1), 7–25.
- Osgood, C. E., & Richards, M. M. (1973). From Yang and Yin to *and* or *but*. *Language*, 49(2), 380–412.
- Otheguy, R. (1980). Meaning in the grammar and in the lexicon. In *Columbia University working papers in linguistics* (pp. 6–11). Columbia, NY: Columbia University.
- Palmer, F. R. (1976). *Semantics*. Cambridge, England: Cambridge University Press.

- Paradis, C., & Willners, C. (2006). Antonymy and negation—The boundedness hypothesis. *Journal of Pragmatics* (special issue *Processes and products of negation*), 38(7), 1051–1080.
- Paradis, C., Willners, C., & Jones, S. (2009). Good and bad opposites: Using textual and experimental techniques to measure antonym canonicity. *The Mental Lexicon*, 4(3), 380–429.
- Partington, A. (1991). *A corpus-based study of the collocational behavior of amplifying intensifiers in English*. (Unpublished masters thesis). University of Birmingham, England.
- Partington, A. (1993). Corpus evidence of language change: The case of the intensifier. In M. Baker, G. Francis, & E. Tognini-Bonelli (Eds.), *Text and technology: In honour of John Sinclair* (pp. 177–192). Amsterdam, The Netherlands: John Benjamins.
- Partington, A. (1998). *Patterns and meanings*. Amsterdam, The Netherlands: John Benjamins.
- Partington, A. (2001). Corpora and their uses in language research. In G. Aston (Ed.), *Learning with corpora* (pp. 46–62). Houston, TX: Athelstan.
- Partington, A. (2004). “Utterly content in each other’s company”: Semantic prosody and semantic preference. *International Journal of Corpus Linguistics*, 9(1), 131–156.
- Peeters, G., & Czapinski, J. (1990). Positive-negative asymmetry in evaluations: The distinction between affective and informational negativity effects. *European Review of Social Psychology*, 1, 33–60.
- Pennebaker, J. W. (2011). *The secret life of pronouns: What our words say about us*. New York, NY: Bloomsbury Press.
- Pepper, S., & Prytulak, L. S. (1974). Sometimes frequently means seldom: Context effects in the interpretation of quantitative expressions. *Journal of Research in Personality*, 8(1), 95–101.
- Pew Research Center. (2007). *What Americans know: 1989-2007. Public knowledge of current affairs little changed by news and information revolutions*. Washington, D.C. Retrieved from <http://www.people-press.org/2007/04/15/public-knowledge-of-current-affairs-little-changed-by-news-and-information-revolutions/>
- Pew Research Center. (2012). *Trends in news consumption: 1991-2012. In changing news landscape, even television is vulnerable*. Retrieved from <http://www.people-press.org/2012/09/27/section-3-news-attitudes-and-habits-2/>
- Pinker, S. (2007). *The stuff of thought: Language as a window into human nature*. London, England: Penguin Books.
- Portner, P., & Partee, B. H. (2002). Introduction. In P. Portner & B. H. Partee (Eds.), *Formal semantics: The essential readings* (pp. 1–16). Oxford, England: Blackwell.
- Praninskas, J. (1975). *Rapid review of English grammar*. Englewood Cliffs, NJ: Prentice-Hall.

- Pratto, F., & Oliver, P. J. (1991). Automatic vigilance: The attention-grabbing power of negative social information. *Personality & Social Psychology*, 380, 380–391.
- Quirk, R., Greenbaum, S., Leech, G., & Svartvik, J. (1985). *A comprehensive grammar of the English language*. London, England: Longman.
- Reid, W. H. (1977). The quantitative analysis of a grammatical hypothesis. In *Columbia University working papers in linguistics* (pp. 58–77). Columbia, NY: Columbia University.
- Reid, W. H. (2002). Introduction: Sign-based linguistics. In W. H. Reid, R. Otheguy, & N. Stern (Eds.), *Signal, meaning and message: Perspectives on sign-based linguistics* (pp. ix–xxi). Amsterdam, The Netherlands: John Benjamins.
- Reid, W. H. (2006). Columbia School and Saussure's *langue*. In J. Davis, R. J. Gorup, & N. Stern (Eds.), *Advances in functional linguistics: Columbia School beyond its origins* (pp. 17–40). Amsterdam, The Netherlands: John Benjamins.
- Reppen, R. (2010). Building a corpus: What are the key considerations? In A. O'Keeffe & M. McCarthy (Eds.), *The Routledge handbook of corpus linguistics* (pp. 31–37). London, England: Routledge.
- Roget, P. M. (1972). *Roget's thesaurus of English words and phrases*. London, England: Penguin.
- Rooth, M. (1985). *Association with focus*. (Unpublished doctoral dissertation). University of Massachusetts at Amherst.
- Rozin, P., & Royzman, E. B. (2001). Negativity bias, negativity dominance, and contagion. *Personality and Social Psychology Review*, 5(4), 296–320.
- Sampson, G. (1992). Probabilistic parsing. In J. Svartvik (Ed.), *Directions in corpus linguistics* (pp. 425–447). Berlin, Germany: Mouton de Gruyter.
- Schegloff, E. A., & Sacks, H. (1973). Opening up closings. *Semiotica*, 8, 289–327.
- Searle, J. R. (1962). Meaning and speech acts. *The Philosophical Review*, 71(4), 423–432.
- Searle, J. R. (1969). *Speech acts: An essay in the philosophy of language*. Cambridge, England: Cambridge University Press.
- Shi, R., Werker, J. F., & Cutler, A. (2006). Recognition and representation of function words in English-learning infants. *Infancy*, 10(2), 187–198.
- Simonton, D. K. (2006). Presidential IQ, openness, intellectual brilliance, and leadership: Estimates and correlations for 42 U.S. chief executives. *Political Psychology*, 27(4), 511–526.
- Simpson, R. C., Briggs, S. L., Ovens, J., & Swales, J. M. (2002). *The Michigan Corpus of Academic Spoken English*. Ann Arbor, MI: The Regents of the University of Michigan.

- Sinclair, J. M. (1989). Uncommonly common words. In M. Tickoo (Ed.), *Learners' dictionaries: The state of the art* (pp. 135–152). Singapore: SEAMEO Regional Language Centre.
- Sinclair, J. M. (Ed.). (1990). *Collins COBUILD English grammar*. London, England: Collins.
- Sinclair, J. M. (1991). *Corpus, concordance, collocation*. Oxford, England: Oxford University Press.
- Sinclair, J. M. (1996). The search for units of meaning. *Textus*, IX, 75–106.
- Sinclair, J. M. (1999). A way with common words. In S. Johansson, H. Hasselggård, & S. Oksefjell (Eds.), *Out of corpora: Studies in honour of Stig Johansson* (pp. 157–179). Amsterdam, The Netherlands: Rodopi.
- Smith, N. K., Cacioppo, J. T., Larsen, J. T., & Chartrand, T. L. (2003). May I have your attention, please: Electrocortical responses to positive and negative stimuli. *The Cognitive Neuroscience of Social Behavior*, 41(2), 171–183.
- Stewart, D. (2009). *Semantic prosody: A critical evaluation*. London, England: Routledge.
- Stubbs, M. (1995). Collocations and semantic profiles: On the cause of the trouble with quantitative studies. *Functions of Language*, 2(1), 23–55.
- Stubbs, M. (1996). *Text and corpus analysis*. Oxford, England: Blackwell.
- Stump, G. (1981). *The formal semantics and pragmatics of free adjuncts and absolutes in English*. (Unpublished doctoral dissertation). Ohio State University, Columbus.
- Stump, G. (1985). *The semantic variability of absolute constructions*. Dordrecht: D. Reidel.
- Suedfeld, P., Tetlock, P. E., & Streufert, S. (1992). Conceptual/integrative complexity. In C. P. Smith, J. W. Atkinson, D. C. McClelland, & J. Veroff (Eds.), *Motivation and personality: Handbook of thematic content analysis* (pp. 393–400). Cambridge, England: Cambridge University Press.
- Svartvik, J., & Quirk, R. (1980). *A corpus of English conversation*. Lund, Sweden: Gleerup.
- Tagliamonte, S., & Roberts, C. (2005). So weird, so cool, so innovative: The use of intensifiers in the television series 'Friends.' *American Speech*, 80(3), 280–300.
- Tannen, D. (1982). Oral and literate strategies in spoken and written narratives. *Language*, 58(1), 1–21.
- Tao, H. (2003). A usage-based approach to argument structure: 'Remember' and 'forget' in spoken English. *International Journal of Corpus Linguistics*, 8(1), 75–95.
- Tao, H. (2007). A corpus-based investigation of 'absolutely' and related phenomena in Spoken American English. *Journal of English Linguistics*, 35(1), 1–21.
- Taylor, S. E. (1991). Asymmetrical effects of positive and negative events: The mobilization-minimization hypothesis. *Psychological Bulletin*, 110(1), 67–85.

- Tetlock, P. E. (1985). Integrative complexity of American and Soviet foreign policy rhetoric: A time-series analysis. *Journal of Personality and Social Psychology*, 49(6), 1565–1585.
- Tetlock, P. E. (1986). A value pluralism model of ideological reasoning. *Journal of Personality and Social Psychology*, 50(4), 819–827.
- Tetlock, P. E., Hannum, K. A., & Micheletti, P. M. (1984). Stability and change in the complexity of senatorial debate: Testing the cognitive versus rhetorical style hypotheses. *Journal of Personality and Social Psychology*, 46(5), 979–990.
- Thoits, P. A. (1986). Social support as coping assistance. *Journal of Consulting and Clinical Psychology*, 54(4), 416–423.
- Thoits, P. A. (2011). Mechanisms linking social ties and support to physical and mental health. *Journal of Health and Social Behavior*, 52(2), 145–161.
- Tobin, Y. (Ed.). (1988). *The Prague School and its legacy: In linguistics, literature, semiotics, folklore, and the arts*. Amsterdam, The Netherlands: John Benjamins.
- Tomasello, M. (1998). Introduction: A cognitive-functional perspective on language structure. In M. Tomasello (Ed.), *The new psychology of language: Cognitive and functional approaches to language structure* (Vol. 1, pp. vii–xxiii). Mahwah, NJ: Lawrence Erlbaum.
- Tomasello, M. (2003). Introduction: Some surprises for psychologists. In M. Tomasello (Ed.), *The new psychology of language: Cognitive and functional approaches to language structure* (Vol. 2, pp. 1–14). Mahwah, NJ: Lawrence Erlbaum.
- Tommasi, M., Pezzuti, L., Colom, R., Abad, F. J., Saggino, A., & Orsini, A. (2015). Increased educational level is related with higher IQ scores but lower g-variance: Evidence from the standardization of the WAIS-R for Italy. *Intelligence*, 50(0), 68–74.
- Traugott, E. C. (1989). On the rise of epistemic meanings in English: An example of subjectification in semantic change. *Language*, 65, 31–55.
- Tuckman, B. W. (1966). Integrative complexity: Its measurement and relation to creativity. *Educational and Psychological Measurement*, 26(2), 369–382.
- Tversky, A., & Kahneman, D. (1991). Loss aversion in riskless choice: A reference-dependent model. *The Quarterly Journal of Economics*, 106(4), 1039–1061.
- Von Fintel, K. (1995). *A minimal theory of adverbial quantification*. Unpublished manuscript, Massachusetts Institute of Technology, Cambridge, MA.
- Wallsten, T. S., Fillenbaum, S., & Cox, J. A. (1986). Base rate effects on the interpretations of probability and frequency expressions. *Journal of Memory and Language*, 25(5), 571–587.
- Willners, C. (2001). Antonyms in context: A corpus-based semantic analysis of Swedish descriptive adjectives. In *Travaux de l'Institut de Linguistique de Lund 40*. Lund, Sweden: Department of Linguistics, Lund University.

- Wittgenstein, L. (2009). *Philosophical investigations*. (P. M. S. Hacker & J. Schulte, Eds., G. E. M. Anscombe, P. M. S. Hacker, & J. Schulte, Trans.) (4th ed.). Oxford, England: Wiley-Blackwell (Originally published as Wittgenstein, L. (1953). *Philosophical investigations*. (G. E. M. Anscombe, Trans., G. E. M. Anscombe & R. Rhees, Eds.). Oxford, England: Basil Blackwell.).
- Xiao, Z., & McEnery, A. (2006). Near synonymy, collocation and semantic prosody: A cross-linguistic perspective. *Applied Linguistics*, 27(1), 283–305.
- Zajonc, R. B. (1968). Attitudinal effects of mere exposure. *Journal of Personality and Social Psychology*, 9(2, Special Supplement, Part 2), 1–27.
- Zubin, D. (1979). Discourse function of morphology: The focus system in German. In T. Givón & C. Li (Eds.), *Syntax and semantics 12: Discourse and syntax* (pp. 469–504). New York, NY: Academic Press.