

**A corpus study
of formulaic variation
and linguistic productivity
in early Greek epic**



Martina Astrid Rodda

Jesus College

University of Oxford

A thesis submitted for the degree of

Doctor of Philosophy

Hilary Term, 2021

A Pierangelo Agazzi
con grandissimo affetto

Preface

A person could do many things in three years (and a half). The *Mahabhārata* says of itself that it was composed in three years, and while the poet Vyāsa makes no claim to have created any new stories, the poem is said to encompass everything a wise person might want to know: “What is found here [...] is to be found elsewhere too, O bull-like heir of Bharata; but what is not found here is to be found nowhere.”¹ My thesis may not contain everything a wise reader might want to know, but my referencing is at least a little more consistent than Vyāsa’s.

This thesis is, as the date of its submission attests, a child of the COVID-19 pandemic. I was lucky enough to have completed any work that required being away from Oxford before the pandemic hit, and therefore my research has not suffered from it any more than all research has; most significantly, I have had very limited access to libraries for the final year prior to submission.² I therefore ask for lenience for any limitations in my bibliography and potential minor mistakes in my referencing. In particular, I have only checked the details of my bibliography entries against online library databases (primarily <http://solo.bodleian.ox.ac.uk/>), rather than against paper copies.

More technicalities: abbreviations of ancient works follow the OCD⁴ (<https://classics.oxfordre.com/staticfiles/images/ORECLA/OCD.ABBREVIATIONS.pdf>), with full titles provided where a work or author is not listed. When it comes to specialised language in oral-formulaic studies, I follow the British convention which favours the Latinate plural ‘formulae’ for ‘formula’ when it denotes the technical concept introduced by

¹ *Mahabhārata* I.56, translated by John D. Smith (Penguin Classics 2009).

² I have also been living in a pandemic hellscape which has damaged my mental health profoundly. As a disabled academic, it particularly matters to me that this is recorded somewhere in this thesis, and that I do not pretend that my academic work is not affected by my well-being, or that my well-being is not affected by everything going on in the world.

Milman Parry (but not when talking about mathematical formulas); I have not, however, changed any use of ‘formulas’ in quotes. I also do not distinguish between ‘formulaic’ and ‘formular,’ but rather use ‘formulaic’ throughout to refer to ‘something pertaining to formulae or to formulaic language;’ again, I have not edited any quotes that follow a different convention. As opposed to most oralist literature, I use singular ‘they’ for ‘Homer,’ or ‘the poet of the *Iliad/Odyssey*,’ given that we are talking about an underdetermined, essentially fictional figure who does not need to have a gender. This is not a statement on the actual gender of singers in archaic Greece, who can safely be presumed to have been mostly male; it is a way of distancing myself from the ways in which writers of English scholarly literature express the gender of even symbolic figures by implicitly accepting cis-patriarchal structures.

My DPhil was funded by the Clarendon Fund (reference number SFF1718_CB_HUMS_1128632), in partnership with Jesus College, Oxford. I am grateful to both my funding providers for their support. I also benefitted from an Enrichment Programme grant from the Alan Turing Institute (January-June 2019). In the final stages of my DPhil, I also received the support of Regent’s Park College, Oxford, where I am currently a stipendiary lecturer, and of Corpus Christi College through a non-stipendiary position.

I am thankful for all the professional support I have received throughout my DPhil, and particularly to the people I have collaborated with throughout the process: Alessandro Achille, Barbara McGillivray, and Hannah Greenstreet for the *Andromeda* project. I am also grateful to Marco Silvio Giuseppe Senaldi and Alessandro Lenci for their collaboration on the article that first introduced me to Distributional Semantics. Armand D’Angour, Constanze Güthenke, Lynn Robson, and Alison Rosenblitt have helped me at various stages of this project. Special thanks also go to Chiara Bozzone for encouraging me to start, and Emma Armitage for her invaluable work at the Turing Institute (and for the memes).

I had the chance to present my research at many seminars and conferences since 2017. I am grateful for all these opportunities, and in particular to the organisers of the Oxford Early Text Cultures seminar, the Oxford Graduate WiP seminar, the London ICS and Lyceum seminars, the London DigiClassics series, the Edinburgh WiP, and the Oxford Undergraduate Classics Society; and of the Cambridge AHRC Doctoral Conference in 2020, the 2019 New Ways of Analyzing Ancient Greek workshop (NWAAG¹) in Göttingen, the 2018 Oxford-Cambridge Philology Day, and the 11th Celtic Classic Conference in St Andrews in 2018.

My thanks also go to my examiners for Transfer and Confirmation of Status, Peter Barber, Gregory Hutchinson, Robin Meyer, and Luke Pitcher (in entirely accidental alphabetical order), who all provided extremely useful feedback. I am, of course, thankful in advance for the work of my doctoral examiners, Gabriel Bodard and Adrian Kelly. Last but very much not least (is this idiom modification allowed?), my supervisor Philomen Probert deserves my deepest gratitude – both for her contribution to the ideas developed in this work, which is substantial, whatever she might say, and for fostering the kindest, most supportive supervision environment for the past three and a half years. I could not have done this without your support.

I met many wonderful people in Oxford, all of whom deserve my thanks. Jason, Daria, Tom, and everyone at the WiP seminar; and Caio, who is one of the most brilliant people I have ever met. The Sanskrit reading group in its many iterations, together with Antonia and her feline companions, has been a source of great joy. The Jesus Peer Support crew has a special place in my heart, together with my dear Jesus friends, Anna, Huw, Jenyth, Lucy and Peter. A special thanks goes to my students at Jesus, Corpus Christi, and Regent's Park, for teaching me more than I have taught them. Also, if you have let me pet your cat, dog, or other pet at any point since October 2017, I promise I am also grateful to you.

Friends from farther away have held me together and sent me all the cute pictures and memes I could ever want, as well as fed me and taken me to the opera when visiting in person was an option: Sophie, Joe, Esther, Sasha, Miller, Chris, Justine, and of course Marie, who is the best cat-friend in existence, and Fabio, who is thankfully still here despite everything. Lena and Eli are not far away at all, but they are always the best genders in the chat. Alex deserves to be thanked separately, for all the Waterstones sessions, and for being one of the best teachers and people I know.

Every acknowledgments section ends with the people without whom I would not exist as the academic and person I am. My family, including all the cat-monsters (two of whom are as old as and definitely better than this thesis), and a very good if accidental dog – your unwavering support and occasional snarky comments (ciao Alberto) have meant everything. I am lucky enough to count two of my teachers among my most treasured friends as well as my mentors: Pierangelo Agazzi, to whom this thesis is dedicated, and Luigi Battezzato. Finally, of course, Camilla and Matteo, who I love very much, and who will react to this statement in two very different ways.

Contents

Part I: Introduction

I	The formula and its definitions: From Parry to 'post-oral' criticism	23
I.1	Why another literature review about formulae?	23
I.2	Repetition and oral composition: from the Parrys to Foley	24
I.2.1	The beginnings of formulae: Milman Parry's 'Studies in the Epic Technique of Oral Verse-Making'	26
I.2.2	Milman Parry's immediate successors: Adam Parry and Albert Lord	31
I.2.3	An oralist legacy: Foley	35
I.2.4	Criticism of Parryan oralist approaches	36
I.3	The 'soft Parryists' and the creation of structural formulae	44
I.3.1	Joseph Russo and the structural formula	45
I.3.2	Visser and Bakker and Fabbricotti: towards a cognitive definition of the formula	47
I.3.3	Nagler and other linguistic approaches	50
I.3.4	'Soft Parryism' as dissolution of the formula?	52
I.4	Conclusions: definition of formulae in this work	54
II	Models of formulaic change in early Greek epic	57
II.1	Approaching formulaic change through change in the 'everyday language'	57
II.1.1	Hoekstra and the first approaches to formulaic modification	59
II.1.2	Janko and the chronology of early Greek epic	62
II.2	A definitional issue: the formula as a flexible unity	66
II.2.1	Formulaic pairs: Hainsworth's definition	66

II.2.2	Some criticism of Hainsworth's definition	69
II.2.3	Beyond linguistic change and noun phrases.....	70
II.3	Modelling formulaic change as syntactic productivity	73
II.3.1	Bakker and the structure of Homeric discourse	73
II.3.2	A constructional approach to the evolution of formulae.....	76
II.3.3	Constructions in diachrony	78
II.3.4	What shall we do with a constructional approach? Some examples (and criticism).....	82
II.3.5	Linguistic change and constructions: an experiment in comparison	84
II.3.6	A critical evaluation of the constructional approach	88
Part II: Syntactic variation in Adjective + Noun phrases		
III	Modelling syntactic flexibility: The case of idioms	95
III.1	Idioms and formulae	96
III.1.1	Multiword expressions and formulae: an initial approach.....	96
III.1.2	Idiom flexibility and processing: decomposability	99
III.1.3	Idiom flexibility and processing: the idiom decomposition hypothesis.....	101
III.1.4	Idiom flexibility and processing: alternative hypotheses	103
III.1.5	Are formulae idioms?	105
III.2	Idiom flexibility in a Construction Grammar perspective.....	107
III.2.1	Idiomaticity in Construction Grammar	107
III.2.2	A corpus-based study of idiomaticity	111
III.2.3	Modelling formulaic flexibility as constructional variation.....	116

III.3	A pilot study: the flexibility of Adj + N pairs	118
III.3.1	Aims and context of the study	118
III.3.2	The target sample.....	122
III.3.3	The baseline sample and results.....	125
III.3.4	Further steps	129
IV	The syntactic flexibility of Adjective + Noun formulae.....	133
IV.1	Materials and methods.....	134
IV.1.1	The Diorisis digital corpus	134
IV.1.2	The sub-corpora	138
IV.1.3	Extracting the target forms	139
IV.1.4	Known issues with the corpus and data extraction	141
IV.1.5	Recording syntactic variation	146
IV.1.6	The baseline sample.....	149
IV.1.7	The data-gathering setup: a summary	150
IV.2	Statistical analysis.....	152
IV.2.1	Relative entropy	152
IV.2.2	Results	154
IV.3	Discussion of results	159
IV.3.1	The general picture	159
IV.3.2	Individual formulae: frequency effects	160
IV.3.3	Individual formulae: the quantifiers πᾶς and ἄλλος	161
IV.3.4	Individual parameters	164

IV.3.5	General conclusions.....	167
V	Comparing poetic grammars: A brief look at non-formulaic epic	169
V.1	The case for Apollonius Rhodius.....	170
V.1.1	Formularity in the Argonautica: a very brief overview.....	171
V.2	Syntactic variation.....	174
V.2.1	Gathering the data	174
V.2.2	Results and comparison with the baseline.....	177
V.2.3	Early epic vs Hellenistic epic: averages and frequency	181
V.2.4	Early epic vs Hellenistic epic: individual patterns.....	182
V.2.5	Early epic vs Hellenistic epic: types of variation.....	184
V.3	Discussion of results.....	186
Part III: Semantic variation in Verb + Object phrases		
VI	A Distributional Semantics approach to formulae: Methods and preliminary study	191
VI.1	Semantic variation in formulae	192
VI.1.1	Open slots and semantic variation	192
VI.1.2	Semantic variation and productivity	194
VI.2	Quantifying meaning: Distributional Semantics.....	199
VI.2.1	The distributional hypothesis	200
VI.2.2	Counting co-occurrences	205
VI.2.3	From matrix to vector space.....	207
VI.2.4	But how good is our model? The validation of DSMs	211

VI.2.5	Do you want to build a model? Count-based approaches vs predictive approaches	213
VI.2.6	Some applications of DSMs to ancient Greek.....	214
VI.3	Creating the vector space model	216
VI.3.1	Corpus and processing.....	217
VI.3.2	The comparison resources	219
VI.3.3	Evaluation	222
VI.4	Using the vector space model: a preliminary study.....	225
VI.4.1	The formulaic sample and distance measurement	226
VI.4.2	Competing hypotheses	228
VI.4.3	Results and discussion.....	230
VI.5	Conclusions.....	232
VII	Mapping semantic variation in transitive VP formulae	235
VII.1	Materials and methods.....	236
VII.1.1	Target formulae	236
VII.1.2	The target corpus.....	238
VII.1.3	Data collection and identification of formulaic pairs	242
VII.1.4	Frequency filtering and the list of target verbs.....	246
VII.1.5	The data for comparison.....	250
VII.1.6	The final TrV + Obj lists	253
VII.2	Analysis and results.....	254
VII.2.1	Setting up the analysis: research questions.....	254

VII.2.2	Setting up the analysis: the productivity and frequency variables	258
VII.2.3	Analysing the data: formulaic vs non-formulaic	261
VII.2.4	Analysing the data: frequency effects	266
VII.2.5	Analysing the data: productivity	268
VII.3	Discussion and conclusions.....	272
VII.3.1	Semantic range and frequency.....	273
VII.3.2	Semantic range and productivity	274
VII.3.3	Semantic range and formularity.....	275
VIII	Conclusions: How it started, how it is going, and some ideas for further work...	279
VIII.1	How it started: initial questions, theoretical grounding, and goals.....	279
VIII.2	How it is going: results and methodological advances.....	282
VIII.3	What we earned, with some ideas for further work.....	287
	Bibliography	289

List of figures

Figure 1: The life cycle of a construction (Bozzone 2014a, 91).	80
Figure 2: The productivity cline according to Barðdal (Barðdal 2008, fig. 7.2 p. 172). Constructions are more productive the closer to the line they fall in terms of type frequency and semantic coherence.	198
Figure 3: Scatterplot representing the relative positions of example target words in a two- dimensional space	208
Figure 4: The semantic space of Greek vocabulary in the AD corpus. From Rodda, Senaldi, and Lenci 2017.	216
Figure 5: Visualization of evaluation metrics according to the gold standard resources and different parameters in the semantic spaces. From Rodda, Probert, and McGillivray 2019.	224
Figure 6: Centroids of two example distributions of points	227
Figure 7: Distribution of cosine distances between the centroid and the object fillers for each formula.....	230
Figure 8: Box-and-whiskers plot of object similarities in formulaic (dotted) and non- formulaic (white) pairs.....	263
Figure 9: Median similarity in the target sample, ordered by frequency	267
Figure 10: Variance in the target sample, ordered by frequency	267
Figure 11: Median similarity in the target sample, ordered by productivity. The shading represents productivity 'bins.'	269
Figure 12: Variance in the target sample, ordered by productivity. The shading represents productivity 'bins.'	269
Figure 13: Average median similarity by productivity group	271
Figure 14: Average variance by productivity group.....	271

Abstract

This thesis addresses the general issue of formulaic variation in early Greek epic – not formulaic change, in that I do not adopt a diachronic perspective or seek to restore the original form of any formulae, but simply to describe how different variants of the same formulae coexist, and what their patterns of variation can tell us about formulaic language. I introduce the concepts of formula and formulaic variation in part I, which is functionally a two-part introduction (chapters I and II). Part II addresses the issue of syntactic variation, focusing on formulae composed by an adjective and a (common) noun; the behaviour of these formulae is analysed in a corpus made of Homer, Hesiod, and the major *Homeric Hymns*, and compared to the behaviour of similar structures in a non-formulaic 5th-century baseline corpus. The final chapter in part II (chapter V) analyses the behaviour of the same Adj + N pairs in a Hellenistic poem, the *Argonautica* of Apollonius of Rhodes. Part III focuses on semantic variation, and applies Distributional Semantics in order to study the range of meaning that objects of transitive verbs can take in Homer and Hesiod. Both Parts II and III have an introductory chapter of their own (chapters III and VI, respectively), in which the linguistic framework that will be used in the whole section is introduced and applied to a smaller-scale study.

Part I

Introduction

I.1 Why another literature review about formulae?

The aim of my thesis is to discuss (some aspects of) language variation in early Greek epic, and in the Homeric poems in particular. While my work does not, as this chapter and the next will show, subscribe to a specific definition of what formulae are in the epic tradition, our expectations of how flexible or rigid the language of early Greek epic can be is shaped on an essential level by our understanding of concepts like repetition, formularity, and the traditional use of language. I will discuss the extent to which repeated phrases undergo syntactic and semantic variation in early Greek epic; this is a question that cannot be separated from the definition of formula as, indeed, a repeated phrase. Formulaic variation, in particular, has been viewed alternatively as a crucial limit to oral-formulaic theory, to the point of its destruction, or as a perfectly natural and expected mechanism that only corroborates the framework of the theory itself. It is, therefore, necessary to discuss the origins and development of the concepts of orality and formularity in Greek epic, as they form the background for this entire work.

What this review does not aim to do is to introduce a methodological polemic by showing how previous definitions of the formula are inadequate or do not match the available data, or to offer a definition of the formula that subsumes, as it were, previous definitions and proves superior to them. My goal here is to clarify the history of the concept and my position with regard to that history; in practice, I will end up setting aside the issue of defining the formula as a concept, in favour of work that does not start from a definitional framework but aims to arrive at a description of formulaic behaviour that can then be used to emphasise empirical differences from non-formulaic phenomena. I will not adopt a dichotomous divide between formulaic and non-formulaic, but rather a framework in which

there is a gradient of more to less formulaic objects. This concept will be explored in more detail in the next chapter, which gathers views on formulaic variation and change, that is, on the formula as a fluid entity, not as a static object.

In practice, this means that some crucial developments in oral-formulaic studies, specifically the work of Hainsworth and the quantitative analyses of Hoekstra and Janko, will not be discussed until chapter II. Said chapter will also introduce some contemporary linguistic approaches to the formula, namely those of Bakker and of the most recent wave of Homeric scholars working under a Construction Grammar approach. The chapter which is now beginning, on the other hand, will mostly examine older work, some of which is outdated or at least has been re-examined critically since its publication. I have decided to keep this review separate from the specific aspects of formulaic studies that will be directly applied in my work, so that readers, especially those who are less versed in formulaic studies, can be aware of the historical context but can also feel free to set aside the finer details of this chapter in favour of the frameworks introduced in the following one.

I.2 Repetition and oral composition: from the Parrys to Foley

Discussions on the language of early Greek epic and what it can tell us about the history and composition of the Homeric poems do not start with Milman Parry. For our present purposes, however, and for the sake of brevity and relevance, it makes sense to begin with the introduction of the concept of the Homeric formula, which will be central to both this chapter and this whole thesis. Some overviews of the field of Homeric studies are available

both from traditionally 'Parryian' scholars³ and from the other side of the field, so to speak, that is from the German neoanalytic school.⁴

It is worth mentioning that in this chapter I will mostly avoid engaging with debates around the history of the composition and transmission of the Homeric (or for that matter the Hesiodic) poems; these questions will only be brought up when they are relevant to discussions about memorisation and reproduction vs recomposition in performance. This approach is not substantially different from that of traditional formulaic analysis: Milman Parry himself notoriously did not engage with the 'Homeric question' in its traditional form (how the Homeric poems were composed, whether their authorship was multiple, and whether different sections with different sources could be recognised, according to analytic and neo-analytic approaches).⁵ I will also not discuss the dialectal layering of the poems and the issue of phase theory, apart from some brief notes on Milman Parry's position.⁶

³ See A. Parry (1989), 199-209, who has something of an agenda in wanting to set up the work of the senior Parry as fully innovative and indeed genius-like in its isolation. More moderate is the first chapter of Foley (1988); the following chapters on Parry, Lord, and the state of the field up to the year of publication are also an excellent resource.

⁴ Latacz (1979) in particular makes the case for the parallels between Milman Parry's work and that of German scholars before him, setting an agenda that is the exact opposite of Adam Parry's.

⁵ Parry himself explicitly states his lack of interest on the very first page of "Studies in the Epic Technique of Oral Verse-Making. I. Homer and Homeric Style" (M. Parry 1971, 266, n. 1). See also A. Parry (1989, 106-7): "Parry avoided the old Homeric Question: was the *Iliad* [...], substantially as we now have it, the product of a single designer, or is our text some sort of composite to which many hands contributed? The proof of the traditional character of Homeric diction seemed to Parry to make this question almost otiose: even if one singer did put together our *Iliad*, his debt to the tradition was so great that the song could still be said to be a direct manifestation of the tradition...." This attitude was criticised by early reviewers (Shorey 1928).

⁶ On the role and perception of dialects in Greek literary history, see now the contributions in Willi (2019), as well as Cassio (2008).

I.2.1 *The beginnings of formulae: Milman Parry's 'Studies in the Epic Technique of Oral Verse-Making'*

Milman Parry, at the time a doctoral student in Paris, is universally credited with introducing the concept of the Homeric formula and of the traditional language of the Homeric poems.⁷ The two French doctoral theses (M. Parry 1928a; 1928b) in which his ideas were first developed are translated into English as the first two chapters of M. Parry (1971), the reference edition of his collected works.⁸ In what follows, however, I will focus on two English-language articles, both titled 'Studies in the Epic Technique of Oral Verse-Making' (M. Parry 1930; 1932, henceforth *Studies I* and *II*), which give a more developed version of the ideas developed in the *thèses*. These two articles are also the ones in which Parry draws the connection between formulaicity and the oral-traditional character of the Homeric diction, something that was not present in the French material.⁹ Due to Parry's untimely death, this is in practice the extent of his major contributions on formulaicity.

Studies I sets out the reasoning behind Milman Parry's definition of the formula and discusses formulaic systems, which are highly relevant to our interests in this section. For cognitive reasons, Parry argues, oral poetry cannot be composed by combining single words, but only through word-groups or phrases that are pre-made to fit the metrical structure when combined. In Homer, these word-groups are formulae, defined in what would become the classic Parryian way as "a group of words which is regularly employed under the same metrical conditions to express a given essential idea."¹⁰ By essential idea, Parry means the semantic nucleus of the expression, the one that remains after discounting all stylistic expansions (for instance, 'when it was morning' for ἦμος δ' ἠριγένεια φάνη

⁷ A. Parry (1989), 209-15, discusses Parry's predecessors (especially in France) in the analysis of Homeric language.

⁸ In this section I will be citing from the 1971 collected works volume. All original texts, including the French theses, are reproduced on the website of the Center for Hellenic Studies; the URL for each is provided in the respective bibliography entries.

⁹ As noted by Finkelberg (2004), 236.

¹⁰ M. Parry (1971), 272.

ροδοδάκτυλος Ἡώς). The second element that determines the nature of a formula is its usefulness to composition: “if a formula is used regularly there must be a steady need for it.”¹¹ The latter is the crucial distinction between formulae in an oral tradition and repeated phrases in non-oral poetry, which are not useful (for instance, αἶλινον αἶλινον εἶπέ in the parodos of Aeschylus’ *Agamemnon*) but are rather used for stylistic effect or as literary references. This usefulness is strictly connected to the metrical value of the formula;¹² formulaic variation which is limited to the change of endings without metrical variation is allowed, but any further level of change which would affect metrical shape presupposes a cognitive effort that is not congenial to the aims of oral poetry. A new formula, and indeed a new formulaic system, can instead be created to fill a metrical slot of a different shape (this happens all the time with name + epithet pairs).

Formulae are either isolated phrases or occur as part of a system of formulae “which express a similar idea in more or less the same words.”¹³ examples are ὀνείαθ' ἔτοῖμα προκείμενα (isolated) vs a system such as ὀλέκοντο δὲ λαοί / ἀρετῶσι δὲ λαοί / δαινυτό τε λαός. Note how the relationship between formulae within a system is regularly expressed by Parry as A ‘is like’ B, with the specification of what kind of ‘likeness’ is intended mostly left to examples; similarly, no explicit definition is provided of what counts as a ‘similar idea’ (or indeed an ‘essential idea’), despite this being a prominent part of Parry’s definition. This is particularly striking when a set of formulae like ὀλέκοντο δὲ λαοί / ἀρετῶσι δὲ λαοί / δαινυτό τε λαός, with meanings ranging from ‘being killed’ to ‘having a banquet,’ is presented as ‘expressing a similar idea.’

These systems of formulae, that is, groups of phrases with the same metrical value and some other similarity of “thought and words,” were consciously available as systems to the

¹¹ M. Parry (1971), 272.

¹² M. Parry (1971), 274: “The formula is useful only so far as it can be used without changing its metrical value.”

¹³ M. Parry (1971), 275.

oral poet, who must have known them “not only as single formulas, but also as formulas of a certain type.”¹⁴ Systems of formulae in oral poetry are characterised by length and thrift; length is the number of formulae in the system, while thrift consists in the avoidance of fully substitutable phrases, that is, phrases expressing the same essential idea in the same metrical shape. These traits are characteristic of oral poetry, and never appear at comparable rates in non-oral material; crucially, the oral poet does not look for new ways to express an idea for which they already have a formula at the ready. Parry illustrates this point through a detailed analysis of repetition and ‘pseudo-formularity’ (my own term) in non-epic poetry, with special attention dedicated to Greek tragedy. The use of repeated, ‘formulaic’ phrases in later epic (Apollonius and Theocritus) in particular is shown to be related to how much the poets wanted to echo the style of Homeric poetry: Apollonius avoids formulae,¹⁵ Theocritus uses them sometimes as a form of epic parody, but neither poet uses them in a way that is compatible with an oral tradition.

Having dealt with the issue of repetition in non-oral poetry, Parry continues with an analysis of repeated phrases in *Il.* 1.1-25 and *Od.* 1.1-25. He explicitly states that while it is easier to identify formulae based on repetition, this is not the most accurate criterion, as shown in the case of formulaic systems, whose similarity goes beyond ‘simple’ verbal and metrical identity. It is important to note here how Milman Parry himself clearly goes beyond his original definition of formula as a repeated phrase. In discussing individual formulae in the proems of the *Iliad* and the *Odyssey*, the issue of the meaning of epithets and other ‘decorative’ phrases emerges: formulae (especially epithets and metaphorical expressions) can have meaning, but their meaning does not “stand out by itself” but is smoothed out in the traditional diction¹⁶ – that is, no individual occurrence of a formula brings with itself a

¹⁴ M. Parry (1971), 275, for both quotations in this sentence. This concept of a formulaic system will be developed by Russo, on whose work see §I.3.1 below.

¹⁵ More on this in §V.1.1.

¹⁶ M. Parry (1971), 306.

unique, striking effect that is not felt in any other occasion in which the same formula appears. That said, intra-textual reference can enrich a formula's meaning: Parry mentions the parallel between *Il.* 1.33, ὥς ἔφατ'· ἔδδισεν δ' ὁ γέρων καὶ ἐπείθετο μύθῳ, referring to Chryses, and the same phrase used of Priam in *Il.* 24.571, where the second occurrence takes emotional force from the first.¹⁷

Finally, Parry takes a further step in the direction of higher schematisation of the concept of formula. Formulae fit specific metrical slots that are useful due to the shape of the hexameter verse and its internal subdivisions. This criterion can be extended to identify formulae that do not even share one word between them, but that follow the same pattern: “τεῦχε κύνεσσιν is like δῶκεν ἑταίρῳ.”¹⁸ This level of schematisation in the diction can only be explained as the effect of a long oral tradition, not the work of one poet: the system is too large and works too effectively for that to be possible. The other important argument for this is the dialectal and historical stratification in formulae, which is addressed in more detail in *Studies II* (summarised below).

Finally, Parry addresses the issue of ‘creativity’ and the creation of new formulae. A poet could in theory create new phrases to fit one of the systems in the poetic diction, but these phrases would not be ‘original’ in the usual sense of the word, as they would not be an act of independent creation in opposition to a received tradition. In any case, the pressure of oral performance will actively discourage poets from being creative; we can see the accidents caused by this pressure in the occasional appearance of unmetrical variants of an originally metrical formula. The creation of formulae that are only alike in sound and not in meaning is also a marker of orality. On the other hand, systems of formulae do develop from one original expression (which is usually impossible to identify), and formulaic

¹⁷ On this kind of intratextual reference, see most recently Currie (2016), who however adopts a Neo-Analyst approach that undervalues the importance of formularity and overstates the significance of repetitions: see the reviews by Rozier (2018) and Andersen (2019).

¹⁸ M. Parry (1971), 313; note again the definitional use of ‘is like.’

variation drives this development: “whenever [poets] could obtain a new formula by altering one which was already in use, they did so.”¹⁹ In any case, it makes no sense to talk about individual style in the *Iliad* and *Odyssey*, because “the poet is thinking in terms of formulas,”²⁰ and can only express the ideas for which a formula (or at least a system of formulae that can be varied or expanded) already exists.

Studies II is less concerned with the definition of formula, which is taken as received from the previous article, but rather with what oralist theory can tell us about the history of the language of archaic epic. Parry examines past theories about the formation of the Homeric language, starting from ancient biographies of Homer. Fick’s work²¹ is identified as the “needed if false step”²² which kickstarted the examination of dialectal stratification in Homer: Fick correctly noted how “Homer’s poetry can with no very great change be turned from Ionic into Aeolic, and [...] the non-Ionic forms are kept as a rule only when Ionic itself has no forms which could take their place.”²³ Indeed, oral poetry promotes the conservation of older forms and introduces new forms exclusively as a function of how useful they are to the poetic language. The poet who learns the language from their predecessor(s) can introduce only minimal changes in their own diction, but these changes accumulate in the course of the tradition. Two main sources of change can be posited: language change in the spoken language and contact between dialects.²⁴ Change in the poetic diction is automatic when no formulaic modification is required,²⁵ but avoided when it would prevent the new form from fitting the metre, at least until a new replacement formula is created and/or the old formula becomes so archaic as to be incomprehensible. Therefore, the

¹⁹ M. Parry (1971), 323. In this sense, Parry’s use of ‘is like’ to express relations between formulae can be seen as an attempt to mirror the cognitive process of the oral poet, who sees formulae (and new creations) as ‘like’ one another in an indeterminate, instinctual way.

²⁰ M. Parry (1971), 324.

²¹ Mainly Fick (1883).

²² M. Parry (1971), 327.

²³ M. Parry (1971), 327.

²⁴ M. Parry (1971), 331.

²⁵ This rule has at least one exception, that is, an archaic form can be retained against all other rules if it sounds desirable because of its own archaism.

traditional language does resist change, but only in a partial and localised way, conditioned by metrical structure and the restrictions of oral composition. The same principles apply when language change happens not due to diachronic evolution in the language of one region, but to the migration of the poetic tradition to regions with a different spoken dialect.

This fundamental introductory discussion is followed by a breakdown of the rules that govern linguistic change in epic diction²⁶ and a detailed discussion of dialectal layers in Homer and in Aeolic lyric poetry.²⁷ Both of these can be omitted here, as they are not relevant to our discussion of formulaicity; more details on dialects and their role in discussions of the diachronic evolution of formulae will be provided in §II.1. Parry concludes with a formulation of phase theory, according to which the epic tradition was originally developed in 'Arcado-Cypriot' (which would later be understood as Mycenaean, and indeed part of South Greek, after the deciphering of Linear B²⁸) and Aeolic, but was then taken up by Ionic singers who gave it the form in which it was transmitted to us.

1.2.2 Milman Parry's immediate successors: Adam Parry and Albert Lord

While Milman Parry's work was interrupted by his death, both his immediate collaborator, Albert Lord, and his son, Adam Parry, continued to develop aspects of formulaic theory. Their work mostly concerns the history of the composition of the poems, a topic that as before we are only going to touch upon cursorily; however, they both have points to make about formulae and particularly the ways in which they are used and learned. Specifically,

²⁶ M. Parry (1971), 339-42.

²⁷ M. Parry (1971), 342-60.

²⁸ See the first chapter of Haug (2002) for a critical discussion of 'Mycenaean' elements in the epic language. For a strongly contrasting perspective on some of these issues, see van Beek (2013). On 'South Greek' see Porzig (1954) and Risch (1955; 1979).

both start moving past the older Parry's work in directions that resonate with present-day linguistic approaches, especially language acquisition and variation.

Adam Parry's contributions, like his father's, are mainly in article form.²⁹ 'The language of Achilles' is perhaps the most famous one, and the one which gives the title to the collected works volume.³⁰ Its main theme is the aesthetic value of the formula: in Adam Parry's view, the crucial element of Homeric style is that it provides an established expression, "a single best way,"³¹ to describe every element of reality. "The style of Homer emphasizes constantly the accepted attitude toward each thing in the world, and this makes for a great unity of experience."³² This vaguely Sapir-Whorfian interpretation of formulaic language³³ can be used to further literary interpretation: Achilles, in particular, is characterised by his inability to use the formulaic language in a way that fully maintains the identity of word and substance. Misuse of the epic language, therefore, is in itself an element of the epic tradition, not just in a linguistic but in an interpretative sense as well.

Similar issues are in the foreground in the posthumous piece on 'Language and characterization in Homer,'³⁴ which revolves around a question raised by the older Parry in *Studies I*: is the meaning of formulae dependent on the individual words that make them up, or is it a feature of the formula as a whole, "so that these phrases are in operation equivalent to single words?"³⁵ Parry discusses how, while formulae are effectively used as lexical units (which would exclude the possibility that the words that make them up have individual meaning), the lexical meaning of elements such as epithets remains in fact

²⁹ Once again, I will cite from the collected works (A. Parry 1989), while also providing references to the original publications.

³⁰ Originally A. Parry (1956).

³¹ A. Parry (1989), 2.

³² A. Parry (1989), 3.

³³ The Sapir-Whorf hypothesis states, broadly speaking, that the language an individual speaks shapes their ability and willingness to analyse reality along specific lines. See Crystal (2008), s.v. 'Sapir-Whorf hypothesis.'

³⁴ Originally A. Parry (1972).

³⁵ A. Parry (1989), 302. See also the famous idea, articulated by Lord (2000), 25, that singers do not have any concept of what a word is, as opposed to a multi-word utterance.

present. The poet can and does care about the intersection of form and meaning, and can depart from traditional form for poetic effect. In linguistic terms, Parry's discussion of whether parts of formulae have individual meaning has many interesting parallels with discussions about whether the meaning of idiomatic phrases can be broken down into constituent parts, a topic that will be addressed at length in §III.1 of this thesis; while Adam Parry does not articulate the parallel (indeed, most of the literature on idioms postdates his death), this is one more previously unnoticed link between traditional formulaic studies and studies of idiomaticity.

Adam Parry's two other articles on Homeric topics, respectively a review of Fitzgerald's translation of the *Odyssey*³⁶ and the much better-known 'Have we Homer's *Iliad*?',³⁷ only touch briefly on issues that are relevant to this thesis. The review of Fitzgerald introduces a concept that is interestingly close to that of formulaic resonance, that is, the idea that the meaning of a formula is also a function of all the traditional contexts in which it occurs.³⁸ Parry makes this point in noting that in the mythical/fantastical sections of the *Odyssey*, the same formulae are used which appear in non-mythical contexts to describe real-life elements or interactions between 'real' people (not gods, monsters, or dead souls). This confers a realistic tone to the fantastical narrative that is at least partially dependent on the other traditional attestations of these formulae.³⁹ 'Have we Homer's *Iliad*?' on the other hand is mostly a response to Kirk (1962), who had argued that the Homeric poems as a whole were transmitted orally for a lengthy period of time before textualization.⁴⁰ Parry counters that under this hypothesis it would be very unlikely for the *Iliad* and the *Odyssey* to be the work of one poet, or one poet each – but, in the younger Parry's view, the *Iliad*

³⁶ Fitzgerald (1961). Parry's article originally appeared as A. Parry (1962).

³⁷ Originally A. Parry (1966).

³⁸ See Graziosi and Haubold (2005; 2015) and Kelly (2007; 2012). This concept is also interestingly close to the broader framework of Distributional Semantics which will be introduced in §VI.2 and applied in §VII.

³⁹ A. Parry (1989), 70.

⁴⁰ On issues of transmission and textualization, see the works cited in n. 53.

must be the work of a single poet, or it would not be the *Iliad* we have, with the same characterisation and succession of scenes and use of language. It is unreasonable to think that an oral tradition would have preserved the text faithfully through memory, as “bards in general do not memorize:”⁴¹ it is this latter point that will be relevant to our discussion of formulaic variation in §IV, which will provide some solid supporting evidence against memorisation (although this applies to short phrases, not to narrative structures).

Albert Lord’s book, *The Singer of Tales*, is another cornerstone of the development of oral-formulaic theory.⁴² Lord addresses in detail, often through comparison with the Yugoslavian bards whose works he had documented and studied together with Milman Parry, the practicalities of oral composition, which he defines as an interplay of memorisation and improvisation, never proceeding by rote, but always within a codified system of rules.⁴³ Chapter II in particular defines and discusses composition in performance, which combines singing, performing, and composing, as well as constant interaction with the audience. In this chapter Lord also discusses the acquisition of poetic language, which happens in three stages: first the bards listen passively to others singing; then they familiarise themselves with the forms of the song (metre, rhythm, formulae) by repeating what they have heard, and they begin to perform after having learned their first song (this does not involve stopping the learning process); finally, the culmination of the training is being able to compose one’s own song(s). This process as sketched out by Lord is regularly cited in discussions of the parallels between natural language acquisition and the acquisition of the performance language.⁴⁴

Chapter III of *The Singer of Tales* is the one that most directly addresses the definition of formula. While repeating Milman Parry’s definition, Lord warns of the risk in thinking about

⁴¹ A. Parry (1989), 126.

⁴² I quote from the second edition, Lord (2000).

⁴³ Lord (2000), 5.

⁴⁴ The latest contribution in this line is Bozzone (2016).

the formula “as a tool rather than as a living phenomenon of metrical language.”⁴⁵ Formulaic language has a grammar which can be viewed as made of paradigms of formulae, instead of verbal/nominal ones like ‘everyday’ linguistic grammar:⁴⁶ the singer will choose the appropriate formulae from the set of available ones, depending on the metrical, semantic, and syntactic context. Like everyday language, formulaic language is not based on rote memorisation, and admits irregularities and exceptions. The key process of creating new formulae consists in substituting new metrically compatible words in an existing formula pattern. At the same time, however, Lord maintains that in an individual’s perception only frequently repeated phrases become formulae, while less frequent ones become less fixed: “The phrases for the ideas most commonly used become more securely fixed than those for less frequent ideas, with the result that a singer’s formulas are not all of the same degree of fixity.”⁴⁷ In chapter IV, we will show how this prediction is in fact contradicted by the data, as more frequent formulae are more flexible than less frequent ones.

1.2.3 *An oralist legacy: Foley*

At this point in our review, we will abandon the history of traditional oralist approaches, which would take up far more space than necessary here. For a summary of the field after Parry and Lord, an excellent source is Foley (1990), who offers oral-formulaic theory readings of Yugoslavian, Old English, and Greek epic poetry. Foley’s approach to oral-formulaic theory emphasises the role of the poetic tradition as offering a network of references on which the poets can draw for connotative effect, in a way that is completely unlike any of the possibilities offered by deliberate allusion in written traditions.⁴⁸ Tradition

⁴⁵ Lord (2000), 30.

⁴⁶ Lord (2000), 35. This is in reference to a specific French structuralist approach to language, which views the opposition between paradigmatic choice and syntagmatic collocation as crucial in language production: see Sanders (2006).

⁴⁷ Lord (2000), 43.

⁴⁸ Foley (1990), 2. See also the works cited in n. 38.

is therefore an additional diachronic dimension that is characteristically and essentially present in oral texts. This focus on the literary implications of orality has been an extremely significant contribution to oral-formulaic approaches and is perhaps the one that remains most fertile nowadays.⁴⁹

1.2.4 Criticism of Parryan oralist approaches

Criticism of oral-traditional approaches to the language of Homer has been a component of Homeric studies since oral-formulaic theory itself made its first appearance.⁵⁰ Without aiming to offer a comprehensive review of the critiques that have been made of Parry's approach and of his followers, for most of this section I will focus on an area that has been particularly targeted by proponents of anti-oralist views, that is, the concept of formulaic economy – the principle that once a formula exists that expresses a given essential idea in a specific metrical shape, no formulae with the same meaning and metrical shape should be created and therefore show up in our texts.⁵¹ In particular, I will discuss how breaches to formulaic economy and other 'violations' of oralist expectations have been used to argue that the role of writing in the development of the Homeric poems, and its effect on their language, is much more significant than previously thought.⁵² I will not, on the other hand, address critiques of Parry and Lord that focus on the history of the composition of the poems, for the same reason why I did not discuss oralist theories of composition in any depth: the primary focus of this thesis is on language, and therefore I limit myself to linguistic arguments.⁵³

⁴⁹ See once again the introduction to Kelly (2007).

⁵⁰ See once again Shorey (1928).

⁵¹ See §1.2.1.

⁵² Another area on which there is active discussion is the extent to which Homeric enjambment corroborates the idea of oral composition or not. On this topic, see at least Higbie (1990), Clark (1994), and Dugan (2012), each of whom summarises previous work.

⁵³ On the composition of the Homeric poems from a decidedly non-Parryian, non-oralist perspective, it is impossible not to cite West (2011; 2014). For a strong counterpoint, see Nagy (2008; 2010),

The title of most recent proponent of a strong anti-oralist stance almost certainly goes to Rainer Friedrich, who, by his own definition, has set out to show how a traditionally 'Parryian' approach is insufficient to account for the language of the Homeric poems.⁵⁴ His counterproposal can be summarised by the notion of a 'post-oral Homer,' which gives the title to his second monograph on the topic (Friedrich 2019): the aesthetics of the Homeric poems, according to the author, are shaped by writing as much as by oral performance, and therefore we should favourably consider the idea that writing played a part in their creation.

Before expounding the concept of 'post-orality' in its entirety, Friedrich had questioned one of Parry's discoveries, the idea of formulaic economy, in a monograph dedicated to the topic (Friedrich 2007). In this book, Friedrich seeks to disprove the concept of formulaic economy or thrift, that is, the idea that there are no pairs of formulae that occupy the same metrical slot and express the same meaning, however this is defined. To contradict this principle, Friedrich draws up an extensive list of formulaic doublets, that is, formulaic phrases which do, in fact, occupy the same metrical slot and express the same meaning;⁵⁵ he then dedicates much of the book to discussing why a poet may want to choose one doublet rather than the other, either for the sake of variation, or to avoid stylistic infelicities, or, on the contrary, in order to choose an expression that is particularly felicitous in a certain context.

and other works by the same author. A small-scale comparison of the views of the two, in the form of a discussion on how the history of the poem should shape editorial criteria, can be found in Nagy (2000), West (2001).

⁵⁴ As Beck (2008) points out, Friedrich's choice of traditional Parryism, rather than later approaches to oral-formulaic theory, as his target is not anodyne, and leads him to construct his argument about orality vs writing in a less nuanced way than necessary. It is, however, useful to discuss Friedrich's objections here, after we have discussed Parry and Lord's theories.

⁵⁵ As one reviewer notes (Visser 2011, 3), Friedrich does not consider the issue of the relative frequency of these doublets: some of his cases are made of a strongly prevalent formulaic expression that is only neglected in favour of a doublet a couple of times.

Friedrich's target in chasing after doublets is both the Parry-Lord theory of oral composition and the very idea of an oral aesthetic. If the principle of thrift is disproved, then it is much less likely that the poems were composed exclusively orally, as economy, in Friedrich's understanding and others', is an essential instrument of orality; moreover, the lack of formulaic economy shows the higher aesthetic value of this hypothetically non-oral tradition.⁵⁶ Friedrich's conclusion is that the poet(s) of the *Iliad* and the *Odyssey* are operating with the tools of "a once broadly schematized oral-traditional diction,"⁵⁷ but are becoming increasingly less dependent on them outside of highly fixed environments such as speech introductions. This conclusion, in Friedrich's view, disproves Milman Parry's Theory of the Oral Homer (capitals Friedrich's own), in favour of a post-oral Homer as anticipated by Adam Parry,⁵⁸ whose diction is "one that has grown or evolved or derived from a purely oral style and would betray its provenance in the retention of oral features,"⁵⁹ but who had some level of command of writing.

Most reviews of Friedrich's first book have focused on his choice to engage with a rather outdated version of 'Parryism,' one which sees thrift as an all-powerful, inviolable rule, while the field of oral studies has by and large moved on from this concept and towards a more nuanced one;⁶⁰ Friedrich himself, after all, recognises that thrift is described by oralist scholars such as Nagy as "a trend or tendency" rather than an absolute principle.⁶¹ Both Christensen's and Minchin's reviews raise the objection that formulaic language, being a linguistic system, is unlikely to have been completely and inviolably systematic, but would

⁵⁶ Friedrich (2007), 138: "A breach of the economy is at the same time a break-away from the limitation and confinement that go with a schematized diction; and as such the frequency of breaches heralds a schema-free style with unrestricted poetic expression."

⁵⁷ Friedrich (2007), 139.

⁵⁸ A. Parry (1966), on which see §1.2.2 above. Friedrich does not award more than a footnote (n. 245 p. 144) to Finnegan (1977), who had also thoroughly anticipated this idea. (Finnegan's book should be cited from the more recent Finnegan 1992 edition, with an updated preface.)

⁵⁹ Friedrich (2007), 142.

⁶⁰ Beck (2008); Christensen (2009); Minchin (2009); Visser (2011).

⁶¹ Friedrich (2007), 140. In his response to Beck's review, Friedrich (2009) stated that he would give a more thorough picture of the field of oralist Homeric studies in his then-upcoming book, now published as Friedrich (2019), on which more will be said below.

rather only “inculcate in [its] speakers patterns of utterance and semantic combination,”⁶² making thrift only “a *general limiting principle* characterising poetic diction [...] that supported the poet as he composed in performance, making it easier for him to find the phrase that would complete his verse.”⁶³ Both of these definitions offer a strong basis for an oralist poetics that still acknowledges the existence of thrift as an overwhelming tendency (which could, I would add, be observed by quantitative measurement), but does not impose it as an absolute restriction. On the other hand, as Deborah Beck points out,⁶⁴ Friedrich’s own argument appears to undermine itself even beyond the issue of how rigidly we expect thrift to operate: if the use of formulaic doublets can be overwhelmingly explained by their higher suitability to a specific context, then there is something in the meaning of one half of the doublet that makes it more or less suitable to a context, and therefore the two halves of a doublet do not, in fact, have the same meaning – so they are not in fact doublets. This objection does not quite fit the traditional Parryian view of epithets as context-neutral, but it is not outside the boundaries of oralist scholarship from more recent decades. In this sense, Friedrich’s analysis of individual cases has a much wider interest even for oralist scholars, or scholars who are not convinced by his ultimate argument for post-orality.

Friedrich’s most recent book, on the other hand, takes up the argument for post-orality in full force, straight from his title, which evokes a *Postoral Homer* (Friedrich 2019). Friedrich’s primary target is once again the ‘Oral-Poetry School,’⁶⁵ represented chiefly by Milman Parry and Albert Lord; he does, however, engage with more recent oralist criticism, most of which has either been discussed above or will be discussed in §1.3. The book is divided into two parts: a critical review of oralist tenets, followed by a constructive part that aims to

⁶² Christensen (2009), 134.

⁶³ Minchin (2009), 11, emphasis original.

⁶⁴ Beck (2008).

⁶⁵ Friedrich (2019), 11.

interpret the *Iliad* in the light of ‘post-orality,’ that is, as a transitional text between oral and literate composition techniques.⁶⁶

The first part of the book is in itself divided between a review of oralist approaches to formularity and other key issues, including enjambment and type scenes, and a section in which Friedrich brings in some partial data from the *Iliad* to show that an oralist conception of the Homeric text is insufficient to explain its actual shape. According to Friedrich, Parry’s work in recognising elements of oral composition in the *Iliad* (in particular the principles of extension and economy) merely opened up a ‘New Homeric Question:’ what exactly are the roles of orality and literacy respectively in the genesis of the *Iliad* and the *Odyssey*? In particular, ‘Parryism’ can neither prove the orality of the Homeric poems nor provide an agreed-upon definition of the formula.⁶⁷ Instead, the Parryist school has by now “dissolved into a diffuse plurality without a unifying frame”⁶⁸ in the works of scholars ranging from Nagy to Foley to Oliver Taplin, each of whom has a different understanding of performance, textual formation, aesthetics.⁶⁹

After this premise, Friedrich offers a brief history of the concept of formula, starting from Milman Parry’s word-groups and formulaic systems characterised by extension and economy. Friedrich argues that despite this initial definition, Parry himself ends up with a “*tacit extension of the definition of the formula to include any verbatim repetition*,”⁷⁰ not just those that belong in systems with extension and economy; to these we must add ‘formulaic expressions’ that allow substitution of words and ‘analogical formulas’ made entirely of

⁶⁶ The *Odyssey* is not discussed, as its poet was “possibly literate from the beginning” (Friedrich 2019, 13).

⁶⁷ Friedrich (2019), 20-1.

⁶⁸ Friedrich (2019), 36,

⁶⁹ One might of course object that each of these scholars is addressing different issues – Nagy the textualization and transmission history of the poems, Foley the role and effect of formulaic language from an aesthetic and pragmatic point of view, and Taplin the performance context. Although there is of course some overlap and cross-fertilisation, it is in fact completely unsurprising that these scholars are providing different answers to different questions.

⁷⁰ Friedrich (2019), 67. When quoting from the book I keep the original italics, of which Friedrich is particularly fond.

variables, only maintaining the same metrical shape. A short summary of oralist definitions of the formula, from Parry to Nagler via Lord, Notopoulos, and Russo, follows, concluded by a useful summary in table form.⁷¹ Friedrich's main criticism of analogical formulae and other 'soft-Parryist' definitions (which I will discuss in the next section) is that their presence is not enough to determine orality, as they can also be found in literate poets (even Vergil). Indeed, metrically determined analogical patterns are simply intrinsic to hexameter poetry; if the metrical constraint is removed, we are just left, according to Friedrich, with an account of how language production works.⁷²

Friedrich then reaches the main point of the first half of the book, which is to reassess the results of the four 'tests of orality' devised by Albert Lord: the formula test, which is based on the density of formulae, the thematic test, based on the use of type scenes, the enjambment test, based on the frequency and type of enjambment, and the test by thrift, which analyses how far repeated expressions are part of a formulaic system which allows for no redundancy. Friedrich's criticism is that Lord's application of the four tests is insufficient, as his samples are too small and his argument tendentious; he then undertakes to prove that even by these tests, the *Iliad* does not have the characteristics of an orally-composed text.

The first test is the 'formula test,' which aims to measure the overall formula density of the Homeric epics. Friedrich briefly recognises that a formulaic analysis of the *Iliad* and the *Odyssey* as a whole exists (Pavese and Boschetti 2003),⁷³ but he disagrees with Pavese and Boschetti's definition of the formula as mere repetition of a sequence of words. Instead, he proposes a 'revised concept of the formula,'⁷⁴ which involves three levels of formularity:

⁷¹ Friedrich (2019), 79, Table b1. On Nagler and Russo see §I.3.1 and §I.3.3 below.

⁷² Friedrich (2019), 82. See §II.2.2 for more details of his criticism of Hainsworth's definition of the formula, and for a rebuttal to it. §V will address the claim that flexible formulae are simply common to all Greek hexameter poetry, not just early epic.

⁷³ This work will be discussed in more detail in §VII.1.3.

⁷⁴ Friedrich (2019), 97-9.

the “*schematized or systematic formula*,” which is Parry’s original definition of a formula that belongs to a system characterised by economy and extension; the “*plain or straight formula*,” that is, a phrase repeated at least three times even in the absence of a complete formula system; the “*formulaic expression*,” which is part of a set of phrases that share a metrical shape, at least one word, and some kind of essential meaning. The formulaic measure of a line or a passage consists in the count of morae in the line that are made of formulae under any of these definitions, combined with the total instances of enjambment. An analysis of some short passages of the *Iliad* shows that high formulaic density correlates with zero enjambment, and both of them are understood by Friedrich as being an oral trait; on the other hand, the frequency of formulae varies depending on the type of passage, with narrative having the highest density, speeches and similes the lowest.⁷⁵ There is no uniformity of oral diction throughout the Homeric epics; therefore, there is no point in testing their overall formularity, which depends on “narrative level and narrative material.”⁷⁶

Friedrich then moves on to other tests of orality, specifically the enjambment test and the thematic test (use of type scenes); the test of thrift is taken as already completed in the 2007 book. All these tests are based on the analysis of sections of text that are indeed larger and more representative than in previous work, but do not amount to a discussion of the whole available material, something Friedrich himself had flagged as necessary. To summarise their results briefly: the Homeric poems show higher rates of stronger necessary enjambment, where a phrase started in one line needs to be completed with material from the following line to be grammatical (as defined by previous scholars; Friedrich reuses statistics from previous work), than the oral Yugoslavian epics, which is taken to show that

⁷⁵ Friedrich acknowledges that this conclusion was anticipated by Griffin (1986).

⁷⁶ Friedrich (2019), 109.

oral and non-oral features co-exist in Homer.⁷⁷ Friedrich also discusses the use of parataxis and hypotaxis, coming to the exquisitely balanced conclusion that “both parataxis and hypotaxis are extensive in Homer but neither eclipses the other.”⁷⁸ As for type scenes, these are not very frequent in Homer, and most importantly tend to be adapted to the current situation, as opposed to their use in Yugoslavian epic, where they are more uncritically repeated.

According to Friedrich, then, Homeric poetry does not meet any of the oralist tests – its diction is not uniformly formulaic, and so on. Instead, in Homer elements of literate and oral composition coexist. From this comes the idea of Homer as a post-oral poet, along the lines of Adam Parry’s idea of the *Iliad* as a transitional text between oral and literary composition.⁷⁹ According to Friedrich, the virtue of post-oral (Homeric) technique is in “transforming the *technically useful into the aesthetically effective*,”⁸⁰ by using oral techniques for the sake of their artistic effect and blending them with non-oral ones as appropriate. So, formulaic repetition is used to connect passages and as an aesthetic choice; type scenes are adapted artistically to their appropriate context; similes offer a variety of images for the same ‘essential idea.’ The book ends with a discussion of narrative structure (largely embracing Neo-Analyst positions) and characterisation, without a proper conclusion on the theme of post-orality.

Both of Friedrich’s books seek to prove that elements in the technique of verse-making that had originally been highlighted as sure signs of the orality of the Homeric epics are actually not present to the extent that oralists want them to be. This leaves open the possibility of a different model of composition, the one that Friedrich describes as post-

⁷⁷ This argument, of course, is open to the same criticism as Parry’s comparison between the Homeric poems and Yugoslav epics: there is no reason for the latter to be the sole representatives of oral epic poetry, even less the measure by which oral poetry is defined.

⁷⁸ Friedrich (2019), 134.

⁷⁹ A. Parry (1966).

⁸⁰ Friedrich (2019), 184.

oral. While Friedrich's data is certainly a welcome addition to the study of formulaicity, his choice to limit himself to rebutting the 'original' oralist arguments while ignoring most of the work that has been done on the basis of and, most importantly, moving away from these original definitions ends up severely limiting the power of his conclusions, to the point that, as argued by Beck (2008), he occasionally ends up supporting oralist theses as they have been reformulated over the last half-century, rather than undermining them.

The following section of this chapter will provide a review of some approaches to the formula that deviate from the original arguments of Parry and Lord, and consequently go beyond Friedrich's analysis as well. While none of them will be applied wholesale in this thesis, we will employ some concepts drawn from these discussions.

1.3 The 'soft Parryists' and the creation of structural formulae

While the school of critics that follow Milman Parry's original definition of the formula as based on metrically identical repetition, much as Lord does, is sometimes known as the school of 'hard Parryists,' scholars who propose adaptations of the Parryian definition of formula are known as 'soft Parryists.'⁸¹ This section will discuss several examples of 'soft Parryist' approaches, which depart from Milman Parry's original work (while often still claiming to be interpreting his definition of the formula) to extend the concept of formula beyond metrical identity and, finally, beyond lexical repetition entirely. As we have discussed in the previous section, some of these aspects were in fact anticipated by Parry himself. I have gathered these approaches under the label of 'structural formulae,' as they all seek to identify the underlying structures of formulaic language in the mind of the poets.

⁸¹ The terms were first used by Rosenmeyer (1965), 297.

1.3.1 Joseph Russo and the structural formula

Joseph Russo is chronologically the first scholar to make a significant proposal to extend the concept of formula beyond Parry's definition (or what he interprets as such). His 1963 article (Russo 1963) builds on discussions of the 'natural' position of certain metrical word-types in the hexameter.⁸² Russo challenges any method of formulaic analysis that is "based on the *literal repetition* of a phrase or at least a word as the litmus test for the presence of a formula."⁸³ Parry himself did not require literal repetition as part of his definition; repetition is just a heuristic technique to begin to explore the formulaic system.⁸⁴ Russo, however, goes beyond Parry's vague 'is like' definition of formulaic similarity, and defines formulae as "localized phrases whose resemblance goes no further than the use of identical metrical word-types of the same grammatical and syntactic pattern."⁸⁵ In other words, Russo's formulae are based on the part of speech of their components: a verb + noun in the dative combination like τεῦχε κύνεσσιν is 'the same' formula as δῶκεν ἑταίρω, since it is made of the same parts of speech in the same tense or case, creating the same metrical shape.⁸⁶

Combining the results of previous metrical studies, Russo shows that most words show a marked preference for certain parts of the line, and there they also regularly co-occur with words of a certain metrical shape. This is then taken further: the positional preference also applies to grammatical forms, not just lexemes – that is, words in the dative of a specific metrical shape (not just a specific dative of that metrical shape) will tend to appear in the

⁸² O'Neill (1942); Porter (1951). It should be noted that, while most of Russo's argument relies on the fixity of metrical position for certain word-shapes in Homer, by his own admission this degree of fixity is higher in Callimachus (definitely not an oral poet), which somewhat undermines his following argument.

⁸³ Russo (1963), 236 (emphasis original).

⁸⁴ See above, §1.2.1.

⁸⁵ Russo (1963), 237.

⁸⁶ The 'same metrical shape' criterion applies to individual words in the formula as well as the whole formula. Russo speaks of parts of speech in the linguistic sense (nouns, adjectives, verbs, etc.), but in his analysis he seems to also include 'same case/tense/other inflectional features' as a requirement to create a formulaic system.

same metrical position. Russo then provides an analysis of a somewhat overused Homeric passage, *Il.* 1.1-7, focusing specifically on the phrases that were already marked as formulaic by Parry and Lord and highlighting the parallels with other expressions based on his 'structural formula' definition.⁸⁷ While the connection between formulae and preferred word position is due to the restriction imposed by formulaic patterns, Russo argues, these same formulaic patterns allow the poet space for creativity, as the formulaic store can be extended by analogy to existing phrases. This is indeed a very similar point to those made by both Milman Parry and Lord; Russo, however, does not require formulaic systems to be based on lexical identity, but only identity of metrical shape and part of speech.

Russo's second contribution (Russo 1966) sets out two issues. The first has to do with how to trace formulaic systems beyond verbatim repetition; the only way to do so is by close reading of the text and the drawing of qualitative parallels. The second is a response to Hainsworth,⁸⁸ who, in opposition to Russo's previous proposal, rejects the idea of moving beyond word repetition, and considers the line between repetition as a characteristic of oral style and repetition as mere repetition of words in a non-oral setting to be a matter of frequency, the threshold between the two being indeterminable. Russo proposes a sort of compromise, in which "the formula of exact repetition" coexists with "the formulaic pattern or structure, in which the same grammatical and metrical word-types are repeated, but with no constant words, all the terms being variables."⁸⁹

What is most relevant here is the development of an approach to formulae that is strongly based on qualitative observation of a formulaic system which, while being more strictly defined than Milman Parry's 'is like' approach, cannot be explored through any other tools than close reading and a deep knowledge of the Homeric style. On the other hand, we are

⁸⁷ Note how Russo's aim here is not to extend the percentage of formulaic coverage detected by Parry and Lord, but to highlight how each formula is part of a larger system.

⁸⁸ Hainsworth (1964). On the development of Hainsworth's ideas, see §II.2.1.

⁸⁹ Russo (1966), 226-7.

starting to see a glimpse of a more quantitative approach to Homeric formulae, which will be the centrepiece of Hainsworth's work (see §II.2), and which I will discuss in the following chapter, as it is much closer to the general approach of my thesis.

1.3.2 Visser and Bakker and Fabbricotti: towards a cognitive definition of the formula

In this section I will discuss some contributions that, like Russo's, also seek to expand our understanding of formula beyond repetition and towards a theory of the underlying structure of formulaic language.

Visser (1988) suggests that a better definition of the formula should be based on the "*lexical solidarities* between noun and epithet,"⁹⁰ that is, a level of expectation that is created by the existence of restrictions in the choice of epithets available for a certain noun, but which does not quite mean that a certain epithet obligatorily must follow that noun.⁹¹ The formula is therefore a unit of semantic value, provided by the noun, and prosodic structure, provided by the various epithets. The example he provides revolves around the range of metrically different verb forms and particles/connectives that are attested in killing scenes in the *Iliad*; Visser shows how in Homer this system contains no metrical doublets (phrases with the same meaning and same metrical shape), while in Quintus of Smyrna, the 4th-century AD author of the *Posthomerica*, several doublets and redundancies are present. Killing scenes are also structured around a meaningful nucleus and prosodically variable fillers: the names of the killer and the killed make up the inflexible semantic nucleus, while

⁹⁰ Visser (1988), 26

⁹¹ This sounds remarkably close to Hainsworth (1968), as will be discussed in §II.2.1. Visser does not reference Hainsworth here, although he is acknowledged in passing later in the article. Visser's approach, however, differs from Hainsworth's in emphasising the difference between a meaningful central element, which needs to remain constant, and peripheral elements that have metrical and stylistic function.

the killing verb and particles are the flexible fillers, to be chosen *ad hoc* depending on what section of the hexameter line needs to be completed.

Bakker and Fabbricotti (1991) more openly depart from the traditional Parryian concept of the formula, which they call “unwieldy ... as well as highly implausible,”⁹² as it relies too heavily on repetition and therefore on memorisation. More recent approaches, such as Visser’s, focus on the process of adaptation of what the poet wants to say to the constraints of the hexameter, a process that happens via the aid of the formula. The key property of Homeric diction in this framework is defined as ‘peripherality,’ that is, the process of addition of peripheral material to an essential nucleus.⁹³ This argument is illustrated through the example of δούρι φαεινῶ and related expressions: their function is “*not* ... to convey what they actually mean, viz. that someone is killed or wounded *by means of a spear* ... [but] to adapt the verb of killing or wounding to its metrical context, by giving it the appropriate length.”⁹⁴ The spear-expression is therefore the periphery to the nucleus provided by the verb of killing. Peripheral elements must be semantically neutral with respect to their context (which does not imply that they are meaningless, just that they are interchangeable in that context depending on metrical need), and take a range of metrical forms in order to fill different gaps in the line. Metrical diversity can be reached through mechanisms that are inherent to Homeric diction, like dialectal variation, or by the addition of even more extra elements (which are in themselves peripheral to the peripheral expression), or by choosing synonyms of different length. That said, peripheral expressions can also acquire more significance, especially when they are removed from their original context. All of this complicates the definition of meaningful vs meaningless in Homeric poetry, and turns it into something rather more gradient in nature (peripheral elements can

⁹² E. J. Bakker and Fabbricotti (1991), 63.

⁹³ See also E. J. Bakker (1988).

⁹⁴ E. J. Bakker and Fabbricotti (1991), 66-7. The authors obviously do not intend to imply that spears are not used for killing – they undoubtedly are –, but simply that transmitting this information is not the purpose of the formulaic expression.

be more or less meaningful); but it also heightens the poets' freedom, compared to a system in which formulae are ready-made building blocks that must be stuck together. The concept of gradient will be highly relevant to the rest of this thesis, although we will partially abandon Bakker and Fabbricotti's approach; the choice of phrases made up of one transitive verb and all its direct object in chapters VI-VII, however, owes much to their discussion of peripherality.

After Bakker and Fabbricotti, Visser's work on killing scenes is directly continued by Riggsby (1992), who extends the analysis to speech introductions.⁹⁵ Once again, the Homeric system is economical, with no unnecessary doublets, while in Apollonius of Rhodes this economy is violated. The position of the essential elements, or determinants (the names of the speaker and addressee), determines what space is left for the fillers, or variables.⁹⁶ There is a parallel between the determinant/variable polarity and the nucleus/periphery one in Bakker and Fabbricotti's work, although variables are not necessarily entirely peripheral and can also be subject to semantic or contextual restrictions (as in the example of introductions to different types of speech – angry speech, speech in front of a crowd, and so on, each of which requires specific expressions). In terms of acquisition, "frequently used combinations of nuclear noun and peripheral adjective are eventually learned as single units;"⁹⁷ Riggsby does not explicitly say whether this should predict higher or lower flexibility of those single units, but the implication seems to be, once

⁹⁵ On speech introductions, see Bozzone (2014a), which will be discussed in detail in §II.3.4, but also Beck (2009).

⁹⁶ Both Visser and Riggsby discuss the role of personal names, but interestingly neither of them considers the likely possibility that some of the names themselves were developed under the pressure of the hexameter form rather than independently of it, and therefore may not be essential or predetermined. This is even more true in killing scenes, where a lot of the victims are entirely irrelevant and only appear in order to be killed; it is easy to see the poet drawing on a reserve of names that not only could easily be accommodated to the hexameter form, but could also be adapted in length depending on the space available. An interesting discussion of character names, their importance, and the limits thereof can be found in Aloni (1986).

⁹⁷ Riggsby (1992), 113.

again, that more frequent units are more fixed and therefore less flexible, a hypothesis that we will directly address in §IV.3.

1.3.3 Nagler and other linguistic approaches

The last couple of approaches to be discussed are those that explicitly draw on or connect to the developments of generative linguistics from the 1950s onwards.

In a long article, Nagler (1967) calls for a theory of the formula that accounts for both the generative properties and the aesthetic value of formulaic language. He starts from pseudo-systems of formulae that seem to be held together by phonological properties rather than unity of essential idea: cases such as *πίονι δημῶ* / *πίονι δῆμω* ('rich fat' / 'prosperous people'), which also extend to sequences like *πίονα δημῶ* or to systems ending in *δῆμω* preceded by verbs of progressively longer length. These are not 'like' each other on any semantic level, and yet they 'sound like' part of the same system of formulae. Can formulae be determined exclusively by the metrical cola they fill and by their general shape, rather than by word identity? Nagler's attempt to "justify [the] subjective feeling of coherence" in these types of systems leads him to postulate the existence of 'families' of formulae that are only connected by "some central Gestalt—for want of a better term—which is the real mental template underlying the production of all such phrases."⁹⁸ This Gestalt is preverbal and has psychological reality in the unconscious mind of the poet.

Nagler emphasises the parallels that this idea shares with the concept of 'deep structure' vs 'surface structure' in Chomskyan linguistics (Chomsky 1957; 1965), but also with structural analysis of myth (Lévi-Strauss 1955). Each formulaic Gestalt is acquired through exposure to existing poetry; since it does not have a fixed verbal shape until the moment

⁹⁸ Nagler (1967), 280-1.

of utterance, it can be easily adapted to any context and can allow for phenomena such as the coexistence of different dialectal form, or forms from different stages of language development. The opposition between tradition and innovation becomes irrelevant: like in linguistic production, “all is traditional on a generative level, all original on the level of performance,”⁹⁹ that is, all linguistic production can be unique while still drawing onto a traditional system of language. While Nagler’s use of Chomskyian concepts is perhaps questionable, Kiparsky (1976) draws a series of connections between formulae and multi-word expressions that rely on the exploration of idioms and bound expressions in the linguistic study of his time. His approach will be best discussed in §III.1, where we will outline similarities and differences between formulae and idioms. Both Nagler’s and Kiparsky’s approaches were developed before the linguistic theory of Construction Grammar and other cognitive linguistic approaches reached their heyday, but they resonate in interesting ways with concepts such as constructions or archetypes, which will be discussed in §II.3.¹⁰⁰

Nagler himself continued to expand on formulaic systems in a related monograph (Nagler 1974). In this book, he explores criteria for resemblance between formulae, a task that he himself recognises as being somewhat nebulously defined: “if one does accept them all [the criteria for formulaic resemblance] impartially, the resultant formula hunting must quickly lead to embarrassingly rich and diverse ramifications.”¹⁰¹ His suggestion is, once again, that systems should be approached as open-ended, productive ‘families’ that rely on an underlying Gestalt or preverbal form, which is then adapted to each necessary context in turn, and realised on the basis of those.¹⁰² One of the consequences of this, according to Nagler, is that there is no dichotomy of ‘original’ and ‘derived’ in formulaic production, as

⁹⁹ Nagler (1967), 291.

¹⁰⁰ For an application of Nagler’s theories in a cognitive-linguistic perspective, see Coleman (2019), working on the Hebrew Psalter.

¹⁰¹ Nagler (1974), 3.

¹⁰² Nagler (1974), 11-14.

all verbal expressions are an adaptation of the preverbal underlying form, and only the latter can claim priority.

1.3.4 'Soft Parryism' as dissolution of the formula?

Much like the 'hard Parryist' approaches, and for not entirely different reasons, the 'soft Parryist' redefinition of the formula has exposed itself to critique. Finkelberg (2004) offers a perceptive summary of some issues. First, she posits, we have the 'traditional Parryist' equivalence between formularity and orality: what is formulaic is oral, which unfortunately also implies that if the text of Homer is not entirely formulaic, it is also not entirely oral. The 'soft Parryist' impulse to redefine the formula in more and more permissive ways is one possible strategy to cope with this problem, as it extends the definition to a point where it (potentially) covers all the available material. This, however, in practice leads to the dissolution of the Parryian concept of the formula as a unit that is based on meaning and metrical function. Finkelberg's critique of Nagler's 'preverbal Gestalt' is particularly sharp: Nagler's concept of the formula does not capture what for Parry are systems of formulae, but is mainly based on rhythm and sound.¹⁰³ On the other hand, the Visser/Bakker and Fabbriotti theory of the formula as a unit with a communicative function that is fulfilled by a nucleus and a periphery is, in Finkelberg's view, too closely aligned with the characteristics of non-poetic speech, does not actually describe all the potential combinations that are seen in Homer, and does not consider the elements of traditional referentiality that are characteristic of the Homeric text.¹⁰⁴

Finkelberg's criticism is a good summary of the reasons why the main body of my work does not, in fact, adopt any of the 'soft Parryian' definitions of formulae. While these can

¹⁰³ Finkelberg (2004), 238-40.

¹⁰⁴ Finkelberg (2004), 240-4.

be useful in pushing our thinking about formulaicity beyond repetition, and will in fact re-emerge in the chapters on the semantics of formulae (§VI-VII), they are also not particularly suited to linguistic work that aims to be quantitative, in that they do not offer a concept of the formula that is easily (or, sometimes, at all) definable and measurable. What we will adopt, on the other hand, is a series of flexible definitions that are mostly frequency-based and allow for variation, including changes to the metrical shape of formulae, but at the same time are still grounded in repetition and (to some extent) lexical identity. The most important of these definitions, proposed by Hainsworth in his doctoral thesis (published as Hainsworth 1968), will be introduced in the next chapter. As we will see, Hainsworth's definition of the formula is grounded in a measurable quality, that is, the association between the different words that make up a formulaic sequence.

Hainsworth's approach, along with others that will again be taken up in the next chapter, is highlighted by Finkelberg as a possible solution to the Parryian problem that was outlined at the beginning of her article. It is entirely possible to extend the concept of the formula to include formulaic variation; at the same time, we should also accept that not everything in Homer appears to be formulaic. Non-formulaic and non-oral do not have to be equivalent;¹⁰⁵ on the contrary, non-formulaic expressions are created to fill the gaps in the formulaic system, and are therefore complementary to it in serving the same function, that of smooth and effective composition in performance.¹⁰⁶ Combined, these factors create an optimal toolkit for language production under the constraints of oral epic poetry.

¹⁰⁵ As the author herself notes (Finkelberg 2004, 246), this line of reasoning is the exact opposite of the one adopted by Friedrich, on which see §1.2.4 above.

¹⁰⁶ Finkelberg (2004), 246-7.

I.4 Conclusions: definition of formulae in this work

In the end, the aim of this project is not to prove or disprove any definition of formula. Definitions are not something that can be proven or disproven, only shown to be more or less useful to conceptualise a specific phenomenon. Beyond this general observation, however, this thesis will focus on three things, none of which will lead us to make a statement about what a formula *is*.

First of all, we will show how we can describe the behaviour of some types of repeated structures in archaic Greek epic (pairs made of an adjective plus a noun or a transitive verb plus its accusative object¹⁰⁷), along the axes of syntactic and semantic behaviour respectively. Neither of these case studies assumes that the structures we are looking at are ‘formulae’ under any definition beyond the most basic idea of ‘a structure that is repeated in archaic epic;’ we will define a structure either by lexical identity, for adjective-noun formulae, or by repetition of the same syntactic head, for formulae involving a transitive verb.¹⁰⁸ The adoption of different definitions for each study will allow us to capture different aspects of the behaviour of these expressions, which would not otherwise be visible.

We will see that the range of behaviours displayed by these structures, when compared with similar structures in non-epic material, does show some patterns that can be used to understand what characterises a formula in its linguistic usage. This leads us to the second goal of this work, which is to understand what patterns of behaviour are more or less characteristic of (formulaic?) material in early Greek epic compared to (non-formulaic?)

¹⁰⁷ Incidentally, both of these structures meet Kiparsky (1976)’s working requirement that a (hypothetical) formula should be a constituent (see p. 80 of the article cited, and §III.1.1 for more details).

¹⁰⁸ As we will see in §VII.1.3, the definition of formula there is slightly more complex, and it also requires word repetition.

material in a baseline corpus. Our conclusions will offer some ideas about how far this characteristic behaviour can be connected to the issue of orality vs literate composition.

One important step in this work will involve a targeted comparison with the behaviour of the same kind of repeated structures in later Greek epic, specifically the work of Hellenistic poet Apollonius of Rhodes. The aim of this further comparison is to examine whether the behaviour of repeated structures in non-oral poetry, of which Apollonius is without doubt a representative, is analogous or substantially different compared to what we saw in archaic, oral material, and discuss what the significance of the similarity and differences is, in terms of style and of how Apollonius characterises his poetic language in relation to his predecessors. This has a bearing on the issue of orality vs literacy, but also on wider questions of style and of self-positioning on the part of Hellenistic poets facing the Homeric tradition. I do not aim to prove that the Homeric and Hesiodic poems must have been composed orally, and could not be immediately and/or partially postoral and still drawing onto an oral aesthetic – in the same way as I am not aiming to prove that a specific definition of formula is ‘right’ above all others. I merely aim to make some observations about a poetic grammar, observations that I believe are best explained by the influence and pressure of oral composition – whether we see this as an essential part of the creation of the Homeric and Hesiodic poems (and the Homeric Hymns), or as an element that had a strong importance in their immediate past.

In short, while this thesis will not offer a new definition of the Homeric formula, it should offer a new perspective on the behaviour of some repeated structures, how this behaviour differs (or does not) from similar structures in other kinds of texts, and why.

II Models of formulaic change in early Greek epic

The previous chapter has offered a survey of the concept of the formula in the Homeric tradition, its history and the criticism with which practically all proposed definitions were met. A trait that is common to all the definitions we surveyed is their synchronic nature: while formulae are often seen as part of a system, whether it is the economy-based system envisaged by Milman Parry or one characterised by a shared verbal Gestalt, these systems are supposed to exist in the mind of each poet – that is, they exist in their entirety at the same point in time.

This chapter will address the issue of formularity from a different perspective, that is, one that seeks to determine how formulae can change over time. We will start from the studies of morphological change that were popular from the 1960s to the 1980s, and then move on to examining approaches to formulaic change that go beyond morphology, and also beyond the idea that formulaic change is a merely automatic phenomenon caused by the pressure of an ‘everyday language.’ The central questions that we will be discussing are: what causes formulae to change? Is it not enough to have a structured system, and why should this be expanded? Is it possible to find a definition (or definitions) of the formula that organically allow for formulaic change, and that we can use (or adapt) for our study?

II.1 **Approaching formulaic change through change in the ‘everyday language’**

The language of Greek epic is widely known to preserve archaic forms that had dropped out of use in either all or many of the Greek dialects by the time the poems were widely circulating in the seventh and sixth century BC.¹⁰⁹ The artistic language (*Kunstsprache*) of

¹⁰⁹ See e.g. Hackstein (2010).

epic preserves word-forms that predate such phenomena as quantitative metathesis (which leads, through synizesis,¹¹⁰ from disyllabic *-āo*, *-āων* to monosyllabic *-εω*, *-εων* in first declension genitives, among other significant changes), the disappearance of digamma,¹¹¹ and the use of movable *v* to prevent hiatus or mark position.¹¹² Similarly, the archaic long datives in *-οισι/-αισι* coexist with the more recently created short forms in *-οις/-αις*,¹¹³ and so on. One of the challenges to which the oral-formulaic theory responded most enthusiastically and effectively in its first half-century of development was to explain how these forms came to appear together in the poetic language, and what the exact role of oral composition was in promoting the multiformity of the epic dialect.

In the formulaic system, the pressure of tradition, among other things, creates a resistance to linguistic change. In Parry's view, given that the primary utility of formulaic systems is to allow for easier metrical composition,¹¹⁴ new word forms created through phonological and/or morphological change in the everyday language are naturally allowed to enter the formulaic system only when they can be substituted within an existing formula without altering its metrical shape. Thus, the *Kunstsprache* evolves side by side with the everyday language of the singers only as long as the result does not lead to metrical irregularities.¹¹⁵ Subsequent research, however, has shown that this resistance is far from being as absolute as Parry envisaged it: in the formulaic system, new morphemes with different metrical shapes gradually replace forms that are no longer in use. Formulaic change of this kind can

¹¹⁰ But see Mendez Dosuna (1993), who argues that it might be wrong to separate synizesis from metathesis as two subsequent steps, and further on this topic Haug (2002), 108-10 and 129-33.

¹¹¹ Or, more accurately, the [w] glide (see Cassio 2008, 33, calling for terminological accuracy); I use 'digamma' as a shorthand.

¹¹² The latter is not simply an innovation, but a phenomenon that primarily belongs to Ionic dialects: see e.g. Janko (1982), 62-8.

¹¹³ There is some controversy on the chronology of these two forms based on data from Mycenaean (Ruijgh 1958); we will return briefly to this issue in §II.1.2.

¹¹⁴ See §I.2.1.

¹¹⁵ M. Parry (1971), 331-3: "The language of oral poetry changes neither faster nor slower than the spoken language, but in its parts it changes readily when no loss of formulas is called for, belatedly when there must be such a loss, so that the traditional diction has in it words and forms of everyday use side by side with others that belong to earlier stages of the language" (333).

be admitted without violating the principle of economy, that is, without creating formulaic doublets: change can affect either the metrical shape of a formula (while the meaning remains stable) or the meaning (within the same metrical shape); on the other hand, the creation of new formulae where both meaning and metrical shape remain stable compared to an existing one (doublets) is generally avoided.¹¹⁶

This fundamental idea – that the epic language is resistant to linguistic change but does not make it impossible, just slows down its penetration into an existing formulaic system – was an extremely fertile one. It did not take very long for scholars to realise that this slow percolation of innovative linguistic forms into the epic language could be used to trace the historical development of the formulaic systems themselves. This led to heightened interest in non-Homeric (or, in a diachronic perspective, post-Homeric) epic poetry, especially the Hesiodic poems and the tradition of the Homeric Hymns: these were seen as offering a window onto a later stage of the poetic tradition than the *Iliad* and the *Odyssey*, as well as one that had developed in a different area of Greece, especially in the case of Hesiod.¹¹⁷ This approach ultimately led, through the work of Arie Hoekstra, to Richard Janko's attempt at establishing a relative chronology of early Greek epic based on the diachronic study of morphological change within formulae.

II.1.1 Hoekstra and the first approaches to formulaic modification

Arie Hoekstra was the first scholar to publish a comprehensive study of formulaic change as a function of change in the spoken language of the singers. His contributions to the study of formulaic change in the 1960s focus on the impact of phonological innovations

¹¹⁶ For a summary of this issue, see E. J. Bakker (1997), 10-14, and Pavese (1972), 165-6 and 171-3.

¹¹⁷ These topics will be discussed below in §II.1.2, but see also Notopoulos (1960; 1962) for an initial expression of these ideas on epic chronology.

(metathesis, neglect of digamma, introduction of movable v to prevent hiatus) on formulaic usage. The main methodological assumption behind this approach is apparent in the reference to ‘formulaic prototypes’ in the title of Hoekstra’s main book on the topic:¹¹⁸ Hoekstra’s model of formulaic evolution presupposes the existence of an archaic prototype, which predates the entry of some specific linguistic innovation into the poetic *Kunstsprache*. After the non-poetic language has created new word forms, these can enter the poetic diction, giving rise to the attested formula; the latter will have a different metrical shape, and will therefore occur in different environments from the original formula.

Most of the time, while forms with archaic and innovative morphology (for instance, observed vs neglected initial digamma) coexist in different formulae, the same formulaic expression does not appear in both its archaic and innovative form. Therefore, when he encounters formulae that are only attested in a form that shows the effects of linguistic innovation, Hoekstra is prompted to reconstruct their archaising prototype, as the latter must have existed for the new formula to be developed. Thus, for instance, a line such as νῦν δ’ ἄγε νῆα μέλαιναν ἐρύσσομεν εἰς ἄλα δῖαν (*Il.* 1.141), which only scans if the etymological digamma in ἐρύσσομεν is neglected, could not be part of the poetic language when digamma was still pronounced; it must, therefore, go back to a prototype of the shape *νῆα πολυκλήϊδα φερυσσέμεν εἰς ἄλα δῖαν, with the original digamma observed.¹¹⁹ Note how, even if the first half of the reconstructed line is hypothetical, the change in metrical shape reverberates through the line, forcing the singer to adjust the whole structure to accommodate the innovation.

This approach leads Hoekstra to view formulaic variation exclusively as a sign of late composition, which required coping mechanisms to accommodate the external pressure of

¹¹⁸ Hoekstra (1965). Hoekstra (1969), a study of the diction of three major Homeric hymns (*Apollo*, *Aphrodite* and *Demeter*), lacks an overall synthesis of the data presented, and as such is of interest mainly for the study of the hymns themselves.

¹¹⁹ Hoekstra (1965), 60-1.

linguistic innovation, while the formulaic diction must have been more rigid in archaic times. Thus, for instance, since verb forms other than the third person singular of historical tenses are more difficult to place in front of words with an initial consonant (including digamma), Hoekstra suggests that “in the period prior to the loss of digamma the narrative aspect and, to some extent, the descriptive element of the style was more predominant than it is found to be in Homer”¹²⁰ – in other words, having more narrative and fewer speeches would have made composition easier, as it required more third-person-singular historic forms than any other verb forms. This approach becomes particularly problematic when Hoekstra is faced with the coexistence of innovative and archaic forms in attested phrases, such as πάντεσσι/πολέεσσι δ' ἀνάσσειν (with a long third-declension dative, archaic, but loss of digamma, innovative), or Χαρόποιό τ' ἄνακτος (archaic genitive in -οιο with loss of digamma): in all of these cases, a form with digamma could be reconstructed only if another linguistic innovation is admitted (*πᾶσιν/*πολέσιν δὲ φανάσσειν, *Χαρόπου τε φάνακτος, both suggested by Chantraine).¹²¹ These examples hint at an aspect of linguistic change in Homer that will be explored primarily by Richard Janko (see §II.1.2), that is, the idea that not all language change happens at the same rate, and that different phenomena can interfere, making it difficult to straightforwardly extrapolate chronological change from language change.

In contrast with this explicitly diachronic approach, Hoekstra is not overtly concerned with the question of whether formulaic change is the effect of a deliberate choice on part of the singer – a tool, as it were, of epic technique – or an automatic, inevitable consequence of the pressure of language change in the singer’s everyday dialect, where older forms were becoming obsolete and as such were perceived as less suitable even for poetic usage. On the one hand, expressions such as “linguistic innovations affected the development of epic

¹²⁰ Hoekstra (1965), 51-2.

¹²¹ See Hoekstra (1965), 55, on these formulae.

style and, in particular, entailed changes in and decomposition of the formulaic diction"¹²² seem to imply that linguistic change causes formulaic change, independently from the will of the singers; on the other hand, the evidence provided by Hoekstra can be used to argue that morphological change *enables* formulaic change, by creating new word shapes and therefore new formulaic shapes that can be used by the singers if they so desire. The issue of whether formulaic change is automatic or deliberate will be discussed in §II.2 when discussing the work that Hainsworth was conducting at the same time as Hoekstra; before we get to this, however, it is necessary to discuss the most impressive contribution on linguistic chronology in archaic Greek epic, that is, the work of Richard Janko.

II.1.2 Janko and the chronology of early Greek epic

Janko's 1982 book, a revision of his doctoral thesis, builds on Hoekstra's and Parry's principles¹²³ in an attempt to reconstruct the relative chronology of early Greek epic through a diachronic study of morphological change within formulae. In Janko's words, if linguistic innovations are entering the poetic language slowly but inevitably over time, "one expects old formulae and archaisms to diminish in frequency through the generations, as innovative phraseology and language creeps in; and if this could be quantified, it might provide a yardstick useful for assigning approximate relative dates to the poems."¹²⁴ The poems that Janko seeks to put in chronological order using this method are the *Iliad*, the *Odyssey*, the *Theogony*, the *Works and Days*, the *Catalogue of Women*, the *Shield of Heracles*, and the *Homeric Hymns to Demeter, Apollo* (divided into a Delian and a Pythian section), *Hermes*, and *Aphrodite*.

¹²² Hoekstra (1965), 131.

¹²³ Cf. the methodological statement in Janko (2012), 23-4.

¹²⁴ Janko (1982), 189.

Janko selects eight linguistic features: use of word-initial digamma, masculine a-stem genitive singulars, a-stem genitive plurals, o-stem genitive singulars, dative and accusative plurals of a- and o- stems, alternations between Z- and Δ-forms in the paradigm of Ζεύς, use of movable ν before consonants or to prevent hiatus.¹²⁵ For each feature, an archaic morph is opposed to a more recent one (for instance, -ᾶο vs -εω in a-stem genitive singulars, or long vs short dative plurals); their relative frequencies are then computed out of the total number of occurrences of each feature.¹²⁶ These frequencies are plotted onto a linear axis, common to all criteria, with the *Iliad* and *Theogony* as reference points;¹²⁷ most works show a visible cluster of results, along with some outliers.¹²⁸ Each cluster is expected to mark the relative position of the work within the chronology of early Greek epic, although Janko mitigates this with inter- and intra-textual considerations, such as signs of imitation between different works and references to historical events. The resulting sequence is *Iliad*, *Odyssey*, *Theogony* and *Catalogue of Women*, *Hymn to Aphrodite* and *Delian Apollo*, *Works and Days*, *Hymn to Demeter*, *Shield of Heracles*, *Hymn to Pythian Apollo*, *Hymn to Hermes*;¹²⁹ later, Janko himself suggested that the *Hymn to Aphrodite* should be moved to an early date, that is, “Homer’s time,” as indicated by the level of archaism in its diction.¹³⁰

As several critics have noted,¹³¹ an important premise of Janko’s chronological approach is that the epic diction not only evolves as a result of external pressure from the spoken language, but that this evolution proceeds at a constant pace through time and in the work

¹²⁵ The latter two are not taken directly into account for chronological purposes, since their variation is geographical rather than chronological; they are, however, used to discuss the place of composition of each of the Hymns (cf. Janko 1982, 62-8). Later work (Janko 2012) also considers other, less frequent criteria (γαῖα vs contracted γῆ, disyllabic ἐν-/ἐύ vs εὐ-/εὔ, the use of τέκος vs τέκνον).

¹²⁶ Janko (1982), 46-8.

¹²⁷ Janko (1982), 72-4 (Fig. 1 and 2).

¹²⁸ See the discussion in Janko (1982), 70-83 (the clusters for each work, as well as the errant results, are summarized in Table 24, p. 81).

¹²⁹ Janko (1982), 200 (Fig. 4). Janko also advances some hypotheses on the absolute date of the poems, depending on the date of the *Iliad* and *Theogony* and on the overall pace of the development of epic diction: cf. his Appendix E, pp. 228-31.

¹³⁰ Janko (2012), 21.

¹³¹ See e.g. the reviews by Fowler (1983), 346, and Hoekstra (1986), 163.

of different singers. No difference between sub-genres is admitted, not to mention the possibility that some sections of the epic poems, such as speeches or similes, might be recreated in performance more often and therefore be less conservative. Deliberate modification is also seen as a marginal phenomenon: while singers could certainly create new formulae, they were not able to control what word-forms they included in their new creations, and were therefore still contributing to the invisible trend towards linguistic innovation. A crucial exception to this is 'false archaism,' that is, the idea that late singers could deliberately employ archaic forms that were noticeably different, longer, and more 'epic-sounding' than the ones that were common in their vernacular.¹³² This, however, is according to Janko an individual phenomenon, limited to some criteria and still assumed not to affect the overall rate of development. Based on the assumption of a constant rate of change, Janko even proposes a relative chronology of the stages in which each archaic form was used exclusively, with no innovations making an appearance.¹³³ This last step is largely hypothetical; Janko, however, argues that chronological development should be assumed based on the fact that innovative forms tend to occur more frequently in modified formulaic phrases,¹³⁴ and that the linguistic shape of the Homeric poems is consistently more archaic than the Hesiodic corpus.¹³⁵ Both of these assumptions, however, are open to debate: as has been pointed out by others, Janko's distinction between archaic and innovative formulae relies on assumptions on the chronology of linguistic change (for instance, that formulae that feature neglected digammas must be more recent coinages);¹³⁶

¹³² Janko (1982), 76-9. An interesting complement to this idea can be found in the work of Mickey (1981), who highlights how metrical inscriptions often contain dialectal traits that are foreign to their area of provenance, in what seems to be an attempt to differentiate them from ordinary speech.

¹³³ Janko (1982), 88 (Fig. 3).

¹³⁴ Janko (1982), 42-6 (on digamma), 54-6 (on dative plurals).

¹³⁵ See also Janko (2012), 27-31.

¹³⁶ See Jones (2010), 298-305. The definition for Class D (innovative) formulae (Janko 1982, 44) mentions neglect of digamma as one of the criteria for categorisation.

and the issue of whether the Homeric poems are earlier than the Hesiodic ones is not entirely settled.¹³⁷

This is only one of several issues which undermine Janko's approach. There are specific methodological pitfalls connected with the use of statistical data to formulate a chronological hypothesis, and Janko does not test for the significance of the differences between each work, but only for the results that fall outside of each work's cluster.¹³⁸ There is also very limited evidence, outside of early epic itself, that the observed linguistic changes are purely the effect of chronological evolution: on the contrary, matters of dialect and geography play heavily into the matter. First there are cases where 'archaic' morphemes remain productive in some dialects (for instance, long datives in -οισι/αισι in Lesbian). Janko's model can, to a certain extent, accommodate these, but only by admitting splits in the epic tradition and/or local influences, such as conservative or innovative trends in the use of digamma and in the dative and accusative endings.¹³⁹ These *ad hoc* exceptions undermine the basic hypothesis that the observed changes in epic diction must be exclusively considered a sign of chronological evolution, and perhaps more importantly, they amount to an admission that the epic language does not, in fact, evolve at a constant pace everywhere. Another challenge comes from morphemes that are considered 'recent' because of their higher productivity in later Greek, but that may already be attested in Mycenaean texts, and as such can hardly be assumed to be innovations in the epic diction; they must, rather, be viewed as having at some point prevailed over their competitors.¹⁴⁰ In the end, the only forms for which innovation seems reasonably certain are genitives in

¹³⁷ Within the same collection of essays as Janko (2012), West (2012) restates the view that Hesiod was composing earlier than Homer.

¹³⁸ Janko (1982), 75-6. For a preliminary discussion of this issue see Rodda (2016); I plan to give a fuller account of my objections elsewhere.

¹³⁹ See Janko (1982), 172-80, on the *Hymn to Aphrodite*. Again, further discussion can be found in Rodda (2016).

¹⁴⁰ There is extended criticism of Janko's treatment of genitive, dative and accusative plurals in Jones (2010; 2012). See also at least Ruijgh (1958), on datives, and Morpurgo Davies (1964), 152-62, on accusatives, both of whom suggest that the forms that Janko considers innovative had always been in use.

-ou¹⁴¹ and neglected digamma. Unfortunately, the former can, in Janko's view, easily be manipulated to create 'false archaism;' as for the latter, philological uncertainties in the reconstruction of etymological digamma are not only inevitable, but can have a significant impact on statistical results.¹⁴²

II.2 A definitional issue: the formula as a flexible unity

Thus far we have discussed formulaic change as something that happens inevitably because of external pressure, mostly independently of the will of the singers; if singers are making conscious choices, these go in the direction of Janko's 'false archaism,' that is, they are an attempt at preserving a more archaic language style in the formulaic system. The other issue that arises from both Hoekstra's and Janko's approaches is how we can account for formulaic variation, that is, for the existence of similar formulae which are not, however, identical or metrically corresponding. Hoekstra and Janko sidestep this problem by working with an original formula containing archaic linguistic forms and its subsequent modifications, rather than with synchronic variation within a formula system.

II.2.1 Formulaic pairs: Hainsworth's definition

A completely different approach from a theoretical standpoint is Hainsworth's 1968 book on formulaic modification, based on his doctoral thesis. Hainsworth's study draws the focus back to the singers' choices, showing a marked preoccupation with deliberate modification and with the limits of epic diction in this respect. Starting with the observation that while

¹⁴¹ Haug's hypothesis (Haug 2002, 70-106) that genitives in -ou go back to an Aeolian (Lesbian) form does not impact the chronological hypothesis, as this is still an innovation compared to inherited -oto.

¹⁴² See Hoekstra (1986), 161. Textual issues can make the situation worse: examples from the *Hymn to Aphrodite* are discussed in Rodda (2016), 92-3.

formulae in the nominative case tend to be very rigidly fixed, inflected formulae show a higher degree of flexibility,¹⁴³ Hainsworth sets out to explore types of formulaic variation that go beyond simply inflecting a formulaic phrase with no impact on metre. Among these are the possibility that a formula may be moved to different positions in the line, modifications that change the formula's metrical shape, expansion through the addition of extra words (ranging from 'ornamental' adjectives to prepositions), separation between the words that make up a formula, and inversion of the word order.

If we want to be able to describe formulae as flexible structures, we need to come up with a definition that enables us to see all of these structures as fundamentally the same formula, or part of the same set of formulae, without relying on metrical identity. Hainsworth's solution is particularly germane to the aims of a corpus-based quantitative study. According to Hainsworth, what guarantees the identity of the formula throughout all these changes is the 'degree of mutual expectation' between its constituent words. Hainsworth is looking at pairs of common nouns and their recurring epithets (adjectives). These pairs are associated in the epic language in a way that goes beyond random chance occurrences: for instance, ships are so often *swift*, *black*, *concave*, the sea is so often *wine-dark*, wine is so often *flashy-sparkling* (αἶθοψ), that when we see either the noun or the epithet on their own, we, as informed listeners of epic, know that it is very likely that the other half of the pair will follow. In Hainsworth's words, for both listeners and singers of epic, "the use of one word [in these pairs] created a strong presumption that the other would follow," giving rise to a bond that is, among other things, non-directional and does not assume continuity between the elements (allowing for inversion and separation).¹⁴⁴ Not all of these connections are univocal, especially since some epithets can refer to a wider range of nouns: compare νήδυμος ὕπνος, where the adjective only appears in epic in agreement

¹⁴³ Hainsworth (1968), 23 and 29.

¹⁴⁴ Hainsworth (1968), 35-9 (quote from p. 36).

with this specific noun, with αἶθοπα οἶνον/αἶθοπι χαλκῷ, where the adjective can be applied to two different nouns, and finally a series like νηὶ μελαίνῃ/νυκτὶ μελαίνῃ/κῆρα μέλαιναν/..., where the adjective combines more freely with different nouns, all of which in turn have other common epithets.¹⁴⁵ Still, the fact that these pairs occur together with a frequency that is higher than could be expected if singers were able to pick from an open set of adjectives provides a foundation for considering these pairs as formulae, that is, structures that are bound together by an observable relationship.

That said, Hainsworth's study still assumes an original or primary form for each formula; in Hainsworth's case, this is defined as the form that is stored in the singer's memory, and which contains the minimum possible amount of modifications, especially in terms of additional material or separation. As opposed to Hoekstra's archaic proto-formulae, these original forms do not usually need to be reconstructed, but are attested in our text; they are often noticeably more frequent than the modified formulae, but not always, which inevitably introduces a degree of arbitrariness.¹⁴⁶ The process of formulaic modification is conceived as a deliberate one: when a singer wants to use a formula, they will 'naturally' think of it in its 'primary shape,' but will change said shape in case it does not fit the verse as is (for instance, if the position where the formula would normally occur is already occupied). This results in a collection of "transient" "secondary expression[s]," employed as "a response to an emergency in composition."¹⁴⁷

This view of modification as a poetic choice leads to another major difference from Hoekstra's approach: rather than as a sign of incompetence or degeneration of the epic tradition, the increase of formulaic modification is viewed as a sign of poetic skill.¹⁴⁸ At the

¹⁴⁵ This phenomenon is in fact entirely expected if we compare this type of flexible formula with bound expressions in everyday language, a topic that will be discussed at length in §III.1.

¹⁴⁶ Cf. Hainsworth (1968), 69.

¹⁴⁷ Hainsworth (1968), 61. This view of formulaic variation as a repair strategy leads to interesting parallels with e.g. the kind of variation studied by Baechle (2007) in tragic metre.

¹⁴⁸ Hainsworth (1968), 113.

same time, the role of traditional constraints such as metre and performance context in guiding formulaic modification is highlighted, as they are the primary factors that create formulaic change, replacing the external pressure of non-traditional language in Hoekstra's and Janko's models.

The idea that formulaic change is a deliberate process that happens due to the circumstances of composition, rather than merely reflecting independent evolution in the everyday language, makes Hainsworth's approach less suited to the diachronic study that was one of Hoekstra's main concerns. Indeed, Hainsworth's analysis is a synchronic and for the most part a descriptive one – a trait that brings it closer to the recent studies of formulaic productivity that will be discussed in §II.3 than to Janko's subsequent work on the chronology of early Greek epic.

II.2.2 Some criticism of Hainsworth's definition

The criticism of Hainsworth's approach advanced by Friedrich (2019) in the course of his review of past definitions of the formula is worth mentioning here, if only for the sake of accounting for and rebutting Friedrich's objections. Friedrich raises two main points, of which the first is an unfair representation of Hainsworth's approach, while the second is quite interesting for the aims of this thesis, as our work in chapters §III-IV will address this objection quite thoroughly. First of all, Friedrich states that Hainsworth's definition of the formula is so vague that it "does not require verbatim repetition:"¹⁴⁹ however, Hainsworth's definition does in fact rely on lexical identity. Even if the words that make up a formulaic pair are not repeated in the same order or the same case, they do need to be repeated as lexemes: Hainsworth's definition of the formula does not allow for lexical substitution. In

¹⁴⁹ Friedrich (2019), 87; the criticism of Hainsworth occupies pp. 87-92.

this respect, Hainsworth's formulae are more rigidly defined than all of the 'Soft Parryist' definitions discussed in §I.3, and even more than Parry and Lord's analogical formulae.

Friedrich's second objection is about whether Hainsworth's definition puts any constraints at all on formulaic change. According to him, Hainsworth's proposal does not respect oralist tenets, because by removing metrical constraints it takes away the usefulness of the formula for performance. If, however, we can show that Hainsworth's formulae are still subject to some constraints, then their usefulness in performance will be proportional to how much they reduce the range of options the poet is expected to choose from: in the same way as more fixed formulae in the style of Parry provide ready-made building blocks for the poet to use, Hainsworth's formulae provide a limited range of choices that can be adapted to the performance context, but is also not entirely open. This concept will be explored in detail in §III-IV, where we will find that Adjective + Noun pairs of the kind identified by Hainsworth do in fact show limitations to their range of variation, and that these limitations are organised in ways that help oral performance. We will also address Friedrich's objection that Hainsworth's definition of the formula can also be applied to expressions found in literate poets: §V will show that while a literate poet such as Apollonius of Rhodes does use a very limited number of Hainsworth-style Adjective + Noun pairs, he is also able to modify them in ways that go beyond the limited range observed for oral poetry.

II.2.3 Beyond linguistic change and noun phrases

Our criticism of Janko's approach in §II.1.2 has highlighted the inherent difficulties in tracking chronological development on the basis of linguistic change in a poetic language

that is highly conservative,¹⁵⁰ potentially split along geographical as well as chronological axes, and in a situation where reliable data on the evolution of the everyday language in archaic times are also lacking. Although Janko's study takes these issues to an extreme point, they apply to any work that assumes morphological change to be the main factor in the evolution of formulae.

They do not, on the other hand, detract from the interest of an approach such as Hainsworth's, which highlights the productivity of formulae as an essential trait of Homeric diction, rather than a largely mechanical phenomenon. Building on Hainsworth's approach in a study of the *Homeric Hymns*, Postlethwaite has argued that discussions on the oral character of Greek epic should actually not rely on measures of formulaic density, which are too dependent on the openness of the formulaic lexicon and on the potentially weak attestation of individual formulae – that is, as we only have a very small portion of the epic material that was circulating in archaic times, we are at very high risk of not recognising formulae or formulaic systems simply because we do not have enough attestations of them in our surviving texts.¹⁵¹ On the contrary, we should consider whether the patterns of formulaic modification in the text under review conform to established trends in oral poetry.¹⁵² Indeed, we should “[accept] mutation as a feature of oral style rather than as a particular mode of response to a difficulty during composition.”¹⁵³

Moreover, all the approaches discussed above focus on nominal morphology, and as such on change in noun phrases (specifically, noun-adjective combinations). Hainsworth gives reasons for this limitation, arguing that ‘sentence patterns,’ that is, formulaic structures

¹⁵⁰ Haug (2002), 22-3, also remarks on the potential pitfalls of tracking archaisms, some of which (e.g. the morpheme -φι) are still used productively in epic diction even after they have disappeared from everyday language, because they are metrically convenient. On the synchronic use of -φι see most recently Goldstein (2020).

¹⁵¹ Postlethwaite (1979), 5-6, especially n. 15.

¹⁵² Postlethwaite (1979), 16.

¹⁵³ Postlethwaite (1979), 17. Note how this differs from Hainsworth's view of ‘emergencies in composition,’ as discussed in §II.2.1.

involving verbs, being of superior length and complexity, are difficult to compare to noun-epithet formulae.¹⁵⁴ One crucial issue is the presence of open slots in verbal formulae (that is, subject or complement positions that can be filled by several different noun phrases, often formulaic in turn), which creates difficulties if a bond of mutual expectation between words is the main criterion used to define a formula.¹⁵⁵

In other words, it may be more productive to focus on describing *how* formulaic change works in a text, rather than looking for the reasons *why* it happens. One of the goals of this thesis is to consider whether there can still be a meaningful difference between formulaic and non-formulaic material, or at least material that we presume to be coming from an oral tradition and material that definitely comes from a literary one, even after we allow for formulae to be modified and for the definition of formula to shift closer to Hainsworth's, that is, to a largely usage-based one. At the same time, we will try to expand our targets beyond Noun + Adjective phrases in favour of the much higher variety that can be seen in Verb + Complement pairs.

The issue of the descriptive power of different frameworks, that is, their ability to give an exhaustive account of the data being analysed, as opposed to their explanatory power, that is, their ability to identify the causes of the observed phenomena, is something that will have to be considered when discussing more recent approaches to formulaic change. The main questions to keep in mind in moving to more recent studies of formulaic productivity, then, are how far these approaches provide the tools to describe formulae of different syntactic shape, and whether their descriptive power is matched by their explanatory power. Indeed, we will see that the main differences between the approaches outlined above and recent studies lie in the fact that the latter focus heavily on the descriptive side,

¹⁵⁴ Hainsworth (1968), 14-18, and, to a lesser extent, 115-16.

¹⁵⁵ This case is similar to situations where synonymy is exploited, e.g. to decline formulae (the case of πατρὶς ἄρουρα/πατρίδα γαῖαν/πατρίδος αἴης/πατρίδι γαίῃ), as discussed by Hainsworth on pp. 123-5. I plan to address expressions involving a verb, and what kind of mutual expectation and/or combinatory restrictions they lead to, in §VII.

and that they describe the behaviour of verb phrases and sentence-patterns better than they apply to the noun-epithet formula.

II.3 Modelling formulaic change as syntactic productivity

II.3.1 Bakker and the structure of Homeric discourse

In the last decades of the twentieth century, the main interest in studies of Homeric formulaicity appeared to move away from linguistic matters such as those discussed in §II.1, focusing instead on the effects of performance on text and culture, often through comparisons with other oral traditions.¹⁵⁶ While this work is not always language-oriented, the increased attention to the unique effects of oral transmission has led more recent studies of formulaic change to acknowledge the parallels between the language of archaic epic and the way natural language works. The idea of a *Kunstsprache* with unique mechanisms and restriction, evolving in parallel with but separate from the everyday vernacular, has thus been slowly replaced with a view of the Homeric language as a system the rules of which can be understood through the study of natural language, and even everyday speech.

Perhaps the most famous instance of this approach is Egbert Bakker's claim that "the *Iliad* and the *Odyssey* are not so much poetry that is oral as speech that is special, a matter of the special occasion of the performance."¹⁵⁷ In his 1997 book, Bakker focuses on discourse structure, a matter that may seem only tangentially related to this project; the interest of his approach, however, lies in how he highlights the connections between the constraints governing Homeric syntax and those that guide spoken language processing. Specifically, Bakker draws a parallel between the basic units of Homeric syntax and the intonation units

¹⁵⁶ Cf. e.g. Stolz and Shannon (1976), Nagy (1996), Foley (1999).

¹⁵⁷ E. J. Bakker (1997), 2.

observed by Wallace Chafe in his studies of spoken narratives.¹⁵⁸ In Chafe's interpretation, the length of these units and the structure of the transitions between them mirror the shifting focus of attention in the speaker's consciousness; as such, they are governed by cognitive constraints (specifically, the limits of our working memory) rather than syntactic ones.¹⁵⁹ In the same way, Homeric discourse is structured through a shifting flow of short units, often coextensive with a formula, each containing a single focus of attention; the relations among them, often signposted through particles, are according to Bakker to be understood in terms of discourse structure (for instance background and close-up devices) rather than 'conventional' syntax.

These 'stylized intonational units' in Homeric discourse are for Bakker nothing other than formulae.¹⁶⁰ Through an analysis of name-epithet formulae, Bakker outlines a model of the development of their discourse function: a 'formative' stage, when the formula is built to be used in connection with a specific function (naming a character in a specific context), is followed by 'routinization,' when the formula is used as an idiom in a recurring speech situation, and finally by 'deroutinization,' when the same expression is used in several different contexts, potentially losing its original function. For example, the original function of a name-epithet formula according to Bakker is to effect the 'epiphany' of a character onto the epic stage; in their routinized stage, however, name-epithet formulae appear after 'staging' formulae (introducing a character), whether the context requires an actual epiphany or not, until finally they come to be used in contexts that have nothing to do with

¹⁵⁸ E. J. Bakker (1997), 44-53, referring to Chafe (1994) and earlier work. Chafe's experiments are also known as 'pear story' experiments, as they elicited spontaneous speech by asking participants to describe a short video involving a boy and a basket of pears.

¹⁵⁹ A very clear explanation of Chafe's theories can be found in Bozzone (2014a), 154-61. On the connection between intonation units and working memory in particular, see 158-9: "Each intonation unit reflects a small package of information that the speaker is trying to communicate to his audience at one time; Chafe labels this package an *idea unit*. The content of each information package is variable, but constrained in terms of its maximum informational weight, which cannot outstrip the capacity of consciousness of both the speaker and the audience; at all times, an attentive speaker will try to match the content of his working memory with the presumed contents of the audience's working memory."

¹⁶⁰ E. J. Bakker (1997), 156.

staging, for instance, when a character who is already on the stage is mentioned again.¹⁶¹ For name-epithet formulae, then, the situation in Homer bears signs of all three different stages of usage, a trait that brings Bakker's model, despite its discourse-based focus, closer to studies of Homeric morphology, with their emphasis on the conservativeness and layering of the Homeric language. Similarly, Bakker's developmental model is inherently diachronic in nature, extrapolating a sequence of chronological development from the evidence that is attested synchronically in the Homeric poems.

Metre, in Bakker's view, also "result[s] from the rhythmical regularization and streamlining of speech units:"¹⁶² intonational units which tend to be more or less of the same size are deliberately stylised and turned into recurring rhythmical units, until rhythm itself becomes regular enough to impose a constraint on the flow of consciousness. The recurrence of routinized behaviour in speech, therefore, has two major outcomes: on the one hand, the 'grammaticalization' of intonation units into formulae, with consequent bleaching of their original pragmatic function; parallel to this, the regularization of the flow of the same intonation units into metre, which according to the author is likely to have been completed only with the transcription of the poems (and may never have been fully perfected, as evidenced by metrical irregularities and phenomena such as *diektasis*).¹⁶³ If taken in its strongest sense, this approach ultimately amounts to flipping Parry's view (among others) on its head: metre is a result of the specific characteristics of the Homeric language, rather than the main constraint determining its shape.

In this, Bakker's approach is representative of a tendency that we will find again when discussing more recent studies: it seeks to explain even the least familiar phenomena in the Homeric language, such as the role of metre, as arising from 'natural' constraints having to

¹⁶¹ E. J. Bakker (1997), 186-200.

¹⁶² E. J. Bakker (1997), 126.

¹⁶³ Cf. E. J. Bakker (1997), 208-9.

do with human communication and the poets' minds (in this case, the regularisation of a pattern that ultimately exists for cognitive reasons). These constraints are, most importantly, widely paralleled outside of early Greek epic. This is as far as it is possible to get from some of the theories discussed in the preceding section, according to which linguistic change in Homer is a unique process, determined by external factors such as the phonological and morphological evolution of the vernacular. On the other hand, the fact that Bakker's approach seems to lead to a more 'natural' view of how language works risks obscuring our lack of precise parallels for the language of Homer, whose nature of 'special speech' cannot be neglected; moreover, Bakker's theories are, by the author's own admission,¹⁶⁴ more easily applicable to an early, unattested stage of the language rather than to the one transmitted to us, where the metre is, in fact, rather rigidly fixed, and formulae seem to be prevalently used in their bleached, 'deroutinized' sense.

II.3.2 A constructional approach to the evolution of formulae

The tendency towards analysing the language of early Greek epic without resorting to the claim that it is an exceptional phenomenon is a productive one in contemporary studies. Along these lines, an attempt to give a complete account of the Homeric formula at all levels of abstraction through the linguistic framework of Construction Grammar has recently been made in a UCLA PhD dissertation by Chiara Bozzone, which contains, among other things, a model of formulaic evolution that is highly relevant to our discussion.¹⁶⁵ Bozzone's work is built around the idea that Homeric formulae can be fruitfully modelled as constructions, that is, "learned pairings of form with semantic or discourse function."¹⁶⁶

¹⁶⁴ E. J. Bakker (1997), 210.

¹⁶⁵ Bozzone (2014a), although the broad outline of this model was already introduced in Bozzone (2010). Antović and Cánovas (2016a) propose a very similar constructional model, with less detail in terms of case studies; they apply it to various oral traditions, in particular the Yugoslavian one. They do not seem to be aware of Bozzone's work, which predates theirs.

¹⁶⁶ Goldberg (2006), 5.

A constructional model of formularity allows us to subsume different definitions of the formula, from those based mainly on identical repetition to ‘structural formulae,’¹⁶⁷ as more or less rigidly schematised instantiations of the formula-construction: “constructions easily accommodate both the hybrid nature and the gradience that Parrian formulas exhibit, since constructions too are hybrid (a pairing of form and function) and can be gradient (in the sense that they can accommodate different levels of abstraction).”¹⁶⁸ The main result of this choice is that formulae can now be analysed as linguistic structures that are the opposite of unique (since, according to Construction Grammar, they form the basic structure of any language), and can in turn be described using a linguistic approach (a grammar, as it were, of the Homeric language), rather than be considered a feature of oral ‘style.’¹⁶⁹

While in Construction Grammar constructions can be used at all levels of linguistic analysis¹⁷⁰ (and, as such, it is in principle not surprising that this concept can also describe formulaic phrases), the most promising parallels for the language of archaic Greek epic turn out to be the elements of everyday speech that are in turn formulaic, that is, idioms and multi-word expressions.¹⁷¹ What was considered one of the characteristics of Homeric formulaic language, Bozzone argues, that is, the extent to which it relies on memorised expressions, is in fact a trait the epic poems share with everyday language: “formularity is a general hallmark of language production tout court, not just of oral traditions.”¹⁷² Among other things, accounts of the way oral poets learn formulae show striking parallels with

¹⁶⁷ As discussed in §I.3.

¹⁶⁸ Bozzone (2014a), 35. Note, however, that modelling the Homeric formula at different (discrete) levels of abstraction, as Bozzone does on p. 36, is not the same as providing a gradient model that allows for continuous variation (we will come back to this point in §II.3.6).

¹⁶⁹ On the dichotomy between style and grammar see Bozzone (2014a), 2-4.

¹⁷⁰ See again Goldberg (2006), 5-9.

¹⁷¹ Bozzone is not the first to have drawn the parallel between formulae and multi-word expressions: see notably Kiparsky (1976). We will return to this parallel in great detail in §III.1.

¹⁷² Bozzone (2014a), 25. The main reference here is to Sinclair’s definition of the ‘idiom principle’ in language use (Sinclair 1991, 109-15), and subsequent discussion of the extent to which ordinary language is made of prefabricated expressions (see e.g. Erman and Warren 2000).

some models of child language acquisition: the apprentice poet starts off by memorising and reproducing formulae, then moves on to a higher level of abstraction that allows him to create new expressions based on traditional patterns in a productive way, in much the same way as children first acquire language by memorising isolated constructions before they start abstracting syntactic patterns from them.¹⁷³ Studying specific constructions may even allow us to isolate the activity of individual poets: for instance, there is a rare speech introduction construction of the shape (υ)υ- προσέειπε(ν) υ-υ (at line-end), with the pre-verbal position filled by a subject or object pronoun and the post-verbal slot filled by a noun indicating either the speaker or addressee, which is only found in *Od.* 14-17 and may be the sign of the work of one individual poet, as are the many unique constructions found in *Il.* 23.¹⁷⁴

Il.3.3 Constructions in diachrony

Bozzone uses her framework to provide an illustration of the way the constructional approach can be used to model formulaic change as syntactic productivity in a natural language. The main criterion at play in her analysis is the interaction between the token and type frequency of competing formulae, that is, respectively, the number of times a specific formula-pattern appears in our text (independently from the exact shape of each specific instantiation) vs the number of different shapes the same formula-pattern can take.¹⁷⁵ Thus, for example, there are three types of speech introduction formulae with

¹⁷³ Bozzone (2014a), 72-7, with references to Lord (2000) on the apprentice singer (see §1.2.2) and the work of Michael Tomasello (see e.g. Tomasello 2003) and Alison Wray (see e.g. Wray 2002, 128-39) on child language acquisition.

¹⁷⁴ Bozzone (2014a), 77-83.

¹⁷⁵ On the difference between type and token frequency, cf. e.g. Bybee and Thompson (1997), 378: "Token frequency is the count of the occurrence in texts of particular words, such as *broken*, or *have*, or of specific phrases, such as *I don't think*. Type frequency, on the other hand, counts how many different lexical items a certain pattern or construction is applicable to." I will return to this definition many times in this thesis.

προσέειπε(ν), ‘they(sg.) addressed:’ one where the verb occupies the second and third foot (e.g. *Il.* 23.722, δὴ τότε μιν προσέειπε μέγας Τελαμώνιος Αἴας), one with the verb at line end (e.g. *Od.* 16.193, ἐξαυτίς μιν ἔπεσσιν ἀμειβόμενος προσέειπεν), and one with the verb in the fourth and fifth foot (e.g. *Od.* 16.166, στή δὲ πάροιθ’ αὐτῆς. τὸν δὲ προσέειπεν Ἀθήνη). Each of these types, in turn, can have several subtypes: for instance, type 1 occurs relatively frequently with αὖτε προσέειπε (subtype 1), πρότερος προσέειπε (subtype 2), or with an object Noun Phrase preceding the verb (subtype 3).¹⁷⁶

Bozzone’s theory starts from the premise that constructions with more types (and types with more subtypes) are more likely to lead to productive generalisations, and as such generate ‘healthy’ formulaic systems that will be used with even greater variety, while constructions with fewer subtypes are more fossilised and less likely to be extended to new cases. Formulae that fall into this second class tend to be used less and less frequently, until they are finally replaced by a more productive alternative; however, formulae with low type frequency but high token frequency are more resistant to replacement by productive patterns, because they are more solidly entrenched in the singers’ memory.¹⁷⁷ An example of this kind of pattern in a living language is the English simple past formation: the simple past of strong verbs has lower type frequency than the weak *-ed* past construction, but the token frequency of individual strong verbs varies; therefore, a strong verb with relatively low token frequency like *spit* is liable to have its strong past tense (*spat*) replaced by *spitted* by analogy with the weak verb pattern, while a high token frequency form like *went* is unlikely to be replaced by **goed* any time soon.¹⁷⁸ With time, it is assumed that the more archaic constructions will progressively lose ground, as types with lower token frequency

¹⁷⁶ All the examples are from Bozzone (2014a), 124-7.

¹⁷⁷ Cf. Bybee and Thompson (1997), 380-1 and 384, and the discussion of statistical pre-emption in Suttle and Goldberg (2011), 1240.

¹⁷⁸ The example is taken from Bozzone (2014a), 83-6. We will go back to this effect in §II.3.6.

disappear; meanwhile, newer, more productive constructions appear naturally within the language, until they finally replace the more archaic ones completely.

Bozzone's model of formulaic change is summarised in Figure 1. The diagram, which is meant to be read from the bottom left quadrant, shows how a compositional phrase can become a formula through chunking, that is, as it becomes perceived as a unit; this young formula then goes into a cycle of isolation and decay, which leaves it completely inflexible and therefore at risk of being replaced by a new, fully productive compositional sequence.

Fig 5.2 The Continuous Creation of Diction

	<i>Low token-frequency</i>	<i>Medium token-frequency</i>	<i>High token-frequency</i>
<i>Low type-frequency</i>	(4) Fossilized Formula (close to replacement)	← decay	(3) Resistant Fossilized Formula (opaque)
<i>Medium type-frequency</i>	↓ replacement ↓		paradigmatic isolation ↑
<i>High type-frequency</i>	(1) Compositional Sequence	→ chunking →	(2) Young Formula (transparent)

Figure 1: The life cycle of a construction (Bozzone 2014a, 91).

All stages of this cycle can (and do) coexist in our text. Bozzone's example is an analysis of speech introduction formulae in the *Iliad* and the *Odyssey*.¹⁷⁹ The token frequency of constructions using the verb προσέειπε(v) increases noticeably in the latter poem (64% when accounting for the difference in length), while the frequency of constructions involving προσέφη and προσηύδα remains strikingly stable (with an increase of less than 1% in both cases). To provide reasons for this difference in productivity, Bozzone breaks each construction down into several subtypes, based on the position of the verb in the line (see above for an example involving προσέειπε(v)), and computes the token frequency for

¹⁷⁹ Bozzone (2014a), 123-41.

each type as well as its subtypes, in order to single out the ones that are becoming more productive.

In the case of προσέειπε(v), for instance, the type 1 construction, with the verb in the second and third foot, shows an increase in frequency in only one of its subtypes, the one with αὐτε προσέειπε(v). This subtype is always followed by a name-epithet formula that fills the rest of the line; according to Bozzone, this restriction is a sign of a construction that is becoming more rigid and on its way to dying out, despite the slight overall increase of type 1 in the *Odyssey*. Type 2, on the other hand, with the verb at line end, shows a marked increase in a subtype, the ἀμειβόμενος προσέειπε(v) one, that allows for great flexibility in the rest of the line. This flexibility even increases in the *Odyssey* compared to the *Iliad*: in the earlier poem the first half of the line is always filled by a name-epithet formula, while in the more recent one the possible fillers are more varied. This trait points towards a healthy construction that is on its way to becoming more productive (indeed, type 2 shows an overall increase of 165% in the *Odyssey*¹⁸⁰). As for type 3, this one is entirely new to the *Odyssey*, as has been mentioned at the end of §II.3.2 above; it allows for great flexibility in the placement of subject and object and in their realisation as pronoun or NP, and therefore (in Bozzone's framework) it stands to reason that it should be highly productive.

While constructions with προσήυδα also fall into two types depending on the verb's position (with the verb occupying the third to fourth foot or at line end), both of which seem remarkably stable in all their subtypes, constructions with προσέφη, which can only occur in the third to fourth foot position, have to be broken down into subtypes depending on how the rest of the line is filled, as was the case with the προσέειπε(v) subtypes above. A look at these shows that the *Odyssey* privileges the ἀπαμειβόμενος προσέφη subtype,

¹⁸⁰ Which is not, incidentally, the highest increase: type 3 is new to the *Odyssey*, and as such has a theoretically infinite increase, which is by definition higher than any other (not a 100% increase, as erroneously reported in table 6.3, p. 125 – this would mean that the frequency of the construction has merely doubled).

while letting all the others decrease in frequency and die out; this, according to Bozzone, is the mark of “a construction that is quickly losing its vitality.”¹⁸¹

Productive constructions, then, are for Bozzone those which have productive types and subtypes; this statement is not entirely tautological, as productivity can vary substantially across types. It should be noted, however, that this does not, strictly speaking, amount to a shift in the analysis from token to type frequency, but rather to a breakdown of the token frequency analysis into further detail. Moreover, the classification into types and subtypes is, up to a point, arbitrary: one could in principle, for instance, classify speech introduction constructions into types according to the position of the verb in the line, rather than the specific verb used (although, to be sure, this classification allows for fewer generalisations, as the productivity of προσέειπε(ν) constructions at line end is much higher than that of προσήύδα constructions in the same position).

In any case, the main and most remarkable outcome of Bozzone’s analysis is that flexibility seems to be a desirable trait in formulaic usage: across the board, the constructions that are more productive are also the ones that allow for greater flexibility in the rest of the line. We will come back to this conclusion in §II.3.6, as it resonates well with some descriptions of the Homeric language, such as Hainsworth’s, and not with others.

II.3.4 What shall we do with a constructional approach? Some examples (and criticism)

According to Bozzone, this model of formulaic change can also offer novel solutions to some well-known issues in Homeric studies. For instance, the apparent exception to formulaic economy provided by the metrical doublets βοῶπις πότνια Ἥρη and θεὰ λευκώλενος Ἥρη at line-end is in fact an example of an unproductive pattern (the βοῶπις

¹⁸¹ Bozzone (2014a), 135.

πότνια one) in the process of being replaced by a newer, more productive one (θεὰ λευκώλενος); the latter shows parallels with other noun-epithet systems (e.g. θεὰ γλαυκῶπις Ἀθήνη, and other uses of λευκώλενος for female characters), and can be inflected in cases other than the nominative, which explains why it is more productive.¹⁸² Another controversial case, the metrically irregular formula (λυτοῦς) ἀνδροτῆτα καὶ ἥβην (Il. 16.857 and 22.363¹⁸³), which according to the standard historical view only scans if we assume a use of syllabic r (*ἀν[ɾ]τῆτα) that predates even the Mycenaean stage of the language, is shown to be a combination of two highly productive patterns, the λιπών/λυτοῦς + direct object pattern, and a construction involving abstract nouns in -της in the accusative or dative that can be schematised as ˘˘τητ˘ καὶ -x (e.g. φιλότητι καὶ εὐνῇ in Il. 3.455 etc.). Since productive constructions are, in Bozzone's view, young and healthy (pun not intended), ἀνδροτῆτα καὶ ἥβην cannot be the phonological fossil that it is widely supposed to be; it must be a much more recent coinage. Based on this argument, and on considerations about the unusual shape and derivation of the abstract ἀνδροτής,¹⁸⁴ Bozzone joins previous scholars in suggesting that the metrically and morphologically regular variant ἀδρότητα καὶ ἥβην, attested by Plutarch for Il. 16.857 and by some of the manuscripts for 22.363, should be adopted.¹⁸⁵

Regarding this last example, it must be objected that even if we accept the idea that the constructions involved in building ἀνδροτῆτα καὶ ἥβην are productive ones, that does not in principle tell us *how long* they have been productive, which would provide a *terminus post quem* for the moment the specific line at hand entered the language.¹⁸⁶ This is not an

¹⁸² Bozzone (2014a), 88-94.

¹⁸³ Not 16.858 and 22.364, as erroneously reported in Bozzone (2014a), 95.

¹⁸⁴ But see Probert (2006), 43-5 (with further references), for an argument that accounts for both the accent and the derivation.

¹⁸⁵ Bozzone (2014a), 94-113. A different proposal, with similar chronological subversion, has been made by Barnes (2011), who on the basis of Avestan parallels suggests that the phrase is a recent coinage based on a more archaic formula *ἄμ(β)ροτῆτα καὶ ἥβην.

¹⁸⁶ Bozzone's argument that "in order to prove that a construction is fossilized, we would have to show that it could not possibly have been created by the productive mechanisms of the language" (p. 96) is an example of the rule/list fallacy (Langacker 1987, 29-30): there is no principle that states

insignificant objection, considering how conservative the language of early epic is, and how innovative forms tend to enter the same inherited formulaic system. Indeed, the sequence ἀνδροτῆτα καὶ ἥβην may be an example of internal glossing of an obscure term, hinting at the fact that the first half of the formula may be archaic.¹⁸⁷ The fact that ἀνδροτῆτα is also attested in the equally unmetrical *Il.* 24.6, Πατρόκλου ποθέων ἀνδροτῆτά τε καὶ μένος ἥϋ, where it fills a different construction, also remains to be explained: even if we could show that this construction is also healthy and productive, should we assume that the text has been corrupted in both cases? In sum, the constructional approach does not provide a compelling reason to change the transmitted text, and any argument to this effect should rather be based on further evidence.

Il.3.5 Linguistic change and constructions: an experiment in comparison

While discussing the case study of ἀνδροτῆτα καὶ ἥβην, we have briefly come across one of the main issues with Bozzone's approach to constructional change. The problem lies in the fact that, while Bozzone's framework allows her to give a more comprehensive descriptive account of change in formulaic language (which is in itself, of course, a remarkable result), it is difficult to argue based on Bozzone's data that the change is due to chronological factors. The synchronic situation in one poem tells us little about the stages that preceded it (this is the case for ἀνδροτῆτα καὶ ἥβην, where it is impossible to establish whether the construction could already have been productive at a more archaic stage), and comparing two different poems is helpful only if we rely on independent assumptions about their relative chronology. In the discussion of speech introduction formulae, the increased

that a linguistic object is only memorised if it cannot be rule-generated, as production based on rules does not actually take functional priority over memorisation.

¹⁸⁷ On this practice see Graziosi and Haubold (2015), 7-10. This instance is peculiar in that what must be obscure is not the etymology of ἀνδροτής (which is in fact very transparent), but the exact meaning required: not 'virility' but something more specialised, such as 'youthful vigour.'

productivity of some constructions in the *Odyssey* is treated as an innovation only because of the assumed later date of this poem. This may seem reasonable,¹⁸⁸ but needs to be proven outside of the constructional analysis.

It should be said that Bozzone's approach does not explicitly claim to prove anything with regard to the relative chronology of the poems; however, the claim that syntactic flexibility is sought after and that the formulaic language is productive in a way that is comparable to the productivity of natural language relies on the idea that more flexible constructions are favoured and survive more easily *through time*. There are, in fact, some alarming signs that the phenomena observed for speech introduction constructions may be an idiosyncrasy in the style of the *Odyssey*, rather than the effect of chronological development: type 2 προσέειπε(v) constructions, for instance, are not in fact widely used in the *Homeric Hymns*, where type 1 predominates (and the Hesiodic corpus, according to Bozzone herself, has no instances of type 2 at all).¹⁸⁹ If the analysis had captured a widespread linguistic phenomenon in the epic tradition, rather than an idiosyncrasy of style, one would expect type 1 προσέειπε(v) constructions to decay and type 2 constructions to continue spreading in the texts that we have – at least, if the predictions made by the model based on type and token frequency are correct. As it stands, Bozzone's data seem in fact to speak *against* a chronological development, and in favour of the individuality of the style of each poem.

We do, however, have a competing model that makes explicitly chronological predictions: this is the morphology-based model, especially in Janko's version. Even though some of its results are not beyond dispute, it can be used to set up a relatively simple experiment which might at least provide independent confirmation of the existence of a chronological trend in Bozzone's data. The underlying hypothesis runs like this: if flexible, 'healthy'

¹⁸⁸ Although it does not account for the possibility of a parallel history of oral transmission for long enough that the language of the two poems could functionally be ascribed to the same time period.

¹⁸⁹ Bozzone (2014a), 130.

constructions are used productively, and if at the same time the epic language is slowly changing its morphological and phonological shape, with the singers introducing more recent forms every time the diction allows, it stands to reason that more productive constructions will be used more often to build new lines, and vice versa. One might then expect recent forms to be more frequently present in the immediate vicinity of a productive construction;¹⁹⁰ conversely, archaic forms should be used more sparingly.

One way to test this is to count the occurrences of innovative vs archaic word forms across Janko's six criteria (use of word-initial digamma, masculine a-stem genitive singulars, a-stem genitive plurals, o-stem genitive singulars, dative and accusative plurals of a- and o-stems)¹⁹¹ in more or less productive constructions – for instance, προσέειπε(v) type 2 formulae, which Bozzone considers to be the most healthy and flexible, vs the 'fossilised' formulae with προσέφη. If the hypothesis stated above holds, we expect to find more innovative forms and fewer archaic forms in προσέειπε(v) type 2 constructions, while the opposite distribution should be found in προσέφη constructions. The numbers for both constructions can be found in Table 1 below (the first form listed for each criterion is always the one Janko considers more archaic).¹⁹²

Table 1: Frequency of archaic vs innovative forms in speech introduction constructions

προσέειπε(v) type 2 constructions:							
	digamma obs./negl.	-ᾱο vs -εω	-ᾱων vs -εων	-οιο vs -ου	-αισι/οισι vs -αις/οις	-ᾱς/ους + V vs - ᾱς/ους + C	totals
archaic	0	0	0	0	4	0	4
innovative	7	0	0	0	0	0	7
							11

¹⁹⁰ Where by 'immediate vicinity' I mean the portion of the text that revolves around the construction – this is usually one line, but constructions with προσήύδα, for instance, can be argued to span two lines, as a separate one is required to accommodate the speaker's name. For this reason, they will not be considered in this section.

¹⁹¹ See above, §II.1.2, for Janko's criteria. I have not taken into account the forms of the name Ζεύς and the use of movable v, because I regard these elements as too largely influenced by dialect.

¹⁹² Data pulled from the online TLG on 11/11/2017. Etymological digamma is reconstructed following the list in Chantraine (2013), 117-54. Like Janko, I exclude cases where the difference is purely graphic, such as short datives followed by vowel (which could conceal an elided long form), as well as cases involving movable v before digamma.

προσέφη constructions:							
	digamma obs./negl.	-ᾱο vs -εω	-ᾱων vs -εων	-οιο vs -ου	-αισι/οισι vs -αις/οις	-ᾱς/ους + V vs - ᾱς/ους + C	totals
archaic	9	0	0	2	0	0	11
innovative	1	0	0	1	0	2	4
							15

Based on these results, there seems to be some level of correlation between more productive constructions and linguistically innovative forms: προσέφη constructions show a higher frequency of archaic forms across the board (with the exception of accusative plurals), while the situation for προσέειπε(v) type 2 constructions is not as straightforward and deserves further attention. Before moving forward, however, it must be noted that the numbers are very low, with some of Janko's criteria never making an appearance (mostly due to the fact that speech introduction formulae tend to involve nominatives and accusatives, both underrepresented among Janko's morphemes); the fact that the difference between the distributions is not significant may also be tied to the small sample size.¹⁹³

What we see in the (supposedly more productive) προσέειπε(v) type 2 constructions is a split between a criterion (initial digamma) which consistently shows innovative forms and another (dative plural) which consistently shows archaic ones. This situation becomes even more puzzling when we look at individual cases: all instances of neglected digamma come from formulaic lines of shape [-υ- μιν ἔπεσσιν / μ' ἔπέεσσιν / σφ' ἔπέεσσιν] ἀμειβόμενος προσέειπε(v),¹⁹⁴ while all long dative plurals come from lines of the shape [-υ-] μύθοισιν ἀμειβομέν[-ος/-η] προσέειπε(v). The two constructions are quite obviously

¹⁹³ The analysis was conducted on the totals for each table (archaic vs innovative), and gives a P-value of P=0.06, according to the two-tailed N-1 Chi squared test (Campbell 2007; performed with the calculator at <http://www.iancampbell.co.uk/twobytwo/calculator.htm>), or P=0.11, according to the two-tailed Fisher-Irwin test (performed with the GraphPad calculator at <https://www.graphpad.com/quickcalcs/>). As always, the smaller the sample size, the less consistent different significance tests tend to be.

¹⁹⁴ Square brackets indicate the range of variation observed within the same construction.

equivalent;¹⁹⁵ indeed, the first one looks like an innovative alternative to the second one, established after the disappearance of digamma.¹⁹⁶ The fact that the ἐπέεσσι/ἔπεσσι form appears more frequently in connection with the προσέειπε(ν) type 2 construction than the alternative with μύθοισιν may be seen as a sign of chronological evolution, even though the numbers are probably too low to allow us to draw any solid conclusion. At the same time, the fact that προσέειπε(ν) type 2 lines with μύθοισιν exist at all testifies to the conservative nature of the epic language, a trait that still needs to be accounted for in Bozzone's approach.

II.3.6 *A critical evaluation of the constructional approach*

While discussing the constructional approach to the evolution of the Homeric language throughout this section, several differences from the 'morphology-based' approaches discussed in §II.1 have been brought out; for ease of reference, they are summarised in Table 2 below. Some of these differences have been discussed by Bozzone herself, who emphasises the way her approach allows us to view the same kind of formulaic modification which Hoekstra, among others, had analysed as a sign of the decay of the formulaic system as "a testament to the vitality of the system – quite the opposite of decomposition."¹⁹⁷ Others, such as the emphasis on description in constructional analysis, have been repeatedly highlighted in the discussion above; the attempt to go beyond the descriptive level by combining constructional analysis and morphological data in §II.3.5 leads only to

¹⁹⁵ With the possible exception of the fact that the ἔπεσσι/ἐπέεσσι construction can more easily accommodate a clitic pronoun. Ancient Greek, however, can drop direct objects, so the pronoun is not necessary; in any case, we have one example of the μύθοισιν construction with a clitic (*Od.* 19.252, καὶ τότε μιν μύθοισιν ἀμειβομένη προσέειπε), and several with a non-enclitic pronoun at the start of the line (e.g. *Il.* 23.794, τὸν δ' Ἀχιλεὺς μύθοισιν ἀμειβόμενος προσέειπεν).

¹⁹⁶ Lines like *Od.* 24.393, μελιχίῳσ' ἐπέεσσι καταπτόμενος προσέειπεν, however, where the long dative is only allowed if neglect of digamma is assumed in ἐπέεσσι, make the connection between innovative morphology and new formulae less straightforward. As this line is ambiguous, in any case, it has not been counted in Table 1.

¹⁹⁷ Bozzone (2014a), 152.

partially satisfying results, but this may be due to sample size. The distinction between types and subtypes introduced by Bozzone, while it does enable us to capture finer-grained phenomena at work in the analysis of speech-introduction constructions, also plays a part in limiting her approach to the descriptive level: since the choice of what makes a type or a subtype is ultimately subjective, any attempt to provide explanations risks leading the researcher simply to choose the classification that leads to the desired predictions.

Table 2: Comparison between ‘morphology-based’ and constructional approaches to formulaic change

type of approach:	‘morphology-based’	constructional
<i>focuses on:</i>	NPs and nominal constituents	VPs and higher-level structure
	diachronic analysis and explanation of causes	synchronic analysis and description of phenomena
<i>sees formulaic variation as:</i>	decomposition of the formulaic system (Hoekstra)	a natural factor (productivity), sought after by the singers
	a solution to performance ‘emergencies’ (Hainsworth)	
<i>relies on external hypotheses about:</i>	linguistic change in the vernacular	relative chronology of the epics

Recasting formulaic change as syntactic productivity, rather than viewing it as a negative outcome of some external pressure on the epic language, is without doubt a valuable approach, in that it allows us to analyse the rules that govern Homeric language as closer to those that lie behind natural language. This, however, much like Bakker’s approach to Homeric discourse syntax as discussed in §II.3.1, carries the risk of losing sight of the differences between the situation in Homer and what is generally observed in natural language, or at least of designing a model that is less suited to capturing those differences.

The first sign of this issue can be found in Bozzone’s treatment of metre. While metrical elements are always integrated into the description of formulae, when discussing higher-level structures such as speech introduction constructions (which involve at least a full line), the metrical constraint is considered a negligible factor: “while it is surely true that verbs with different metrical shape combine with noun-epithet formulas of different metrical shape, it is also worth keeping in mind that all main characters come with noun-epithet

formulas of all shapes (feminine, bucolic, and hephthemimeral), so that none is banned from one of these constructions simply by virtue of meter.”¹⁹⁸ While this may be true for larger units, it is also quite obviously a consequence of the metrical constraint, not merely a feature of the epic language. It is likely, moreover, that this constraint would play a significant role in restricting formulaic flexibility in smaller settings: for instance, it is hard to neglect the importance of metre in determining the resistance of βοῶπις πότνια Ἥρη to inflection (while θεὰ λευκώλενος Ἥρη is readily inflected), an observation that undercuts the proposed explanation based on type and token frequency (summarised in §II.3.3). Once again, neglecting metre may appear to be a sensible choice at the descriptive level, but limits the explanatory power of the model.

The most important conclusion of the constructional study, however, is the idea that formulaic flexibility is sought after and exploited in the Homeric language just like it is in natural language. This conclusion, along with the parallels Bozzone observes with natural language acquisition, inevitably leads us to ask: if epic Greek poetic language works like any vernacular, then what traits make it ‘special,’ to use Bakker’s wording, or rather suited to its specific performance culture? Or, to phrase the question in a way that is more fertile for research: is it possible to isolate and describe these traits linguistically?

In the same way, when it comes to diachronic language change, the differences between the layering of productive and non-productive forms in epic Greek and analogous phenomena in everyday language still need to be accounted for. Is what makes the layering of forms in Homer so impressively varied and conservative merely a quantitative difference, while the underlying phenomena and causes are the same that are found in everyday language? This may well be the case, but if it is so, a way to measure such a difference would be a welcome find.

¹⁹⁸ Bozzone (2014a), 122.

Specifically, what is desirable is an approach that would allow us to quantify and compare the rate of linguistic change in Homeric poetry to the situation in other corpora, and to correlate this rate with other factors in order to draw up explanatory hypotheses. A natural point of comparison is provided by other ancient Greek corpora, such as later hexametric poetry, which is not widely regarded as formulaic.¹⁹⁹ It would be interesting to measure the difference between ‘natural’ linguistic usage in oral poetry and ‘artistic,’ deliberate variation in later poetry – or, indeed, test whether there is a measurable difference at all, and where it may lie. Another *desideratum* is a more complete description of formulaic flexibility as a gradient; specifically, something that would allow us to draw connections between the level of compositionality or fixedness of a Homeric construction (assuming that these range from more compositional to more rigidly formulaic) and the level of variation attested in our corpus. The constructional approach provides us with the conceptual tools to build this kind of model, but it is still necessary to move beyond the level of qualitative description.

¹⁹⁹ Bozzone (2014b) also briefly analyses speech introductions in Apollonius and Quintus Smyrnaeus in terms comparable with her descriptive analysis of Homer.

Part II

Syntactic variation in Adjective + Noun phrases

III Modelling syntactic flexibility: The case of idioms

One element that emerges clearly from the approaches surveyed in the previous chapter is a methodological trend towards describing formulae as linguistic structures, rather than a uniquely literary phenomenon. Bozzone's work, which describes formulae as constructions, is a thorough attempt in this direction, but still does not provide a straightforward answer as to what kind(s) of syntactic structure can provide a good point of comparison for formulae. In Construction Grammar(s), structures at any level of analysis can be described as 'constructions,' and as such no specific property can be postulated based on this definition alone.²⁰⁰

At first blush, the idea that formulae are some kind of linguistic structure may seem fairly obvious: formulae make up a (literary) language, and they should be describable in linguistic terms. This statement, however, hides at least one underlying assumption, that is, that the category 'formula' as a whole should correspond to *one* class of structure, rather than potentially covering a range of phenomena. It is useful to remember, once again, that there is an extremely wide variety of definition of formula in the literature,²⁰¹ and that the category itself was born as an instrument to describe a style of composition rather than a specific set of syntactic phenomena. As such, while we may want formulae to comprise a natural class – to share some set of properties that distinguishes them from non-formulae –, it is in fact extremely difficult to define the boundaries of this class.

The question that lies behind our efforts in this chapter is: is there a class of linguistic expressions that, while they may not correspond fully to formulae, can be productively compared to formulae? How far can the behaviour of non-formulaic linguistic expressions

²⁰⁰ See below, §III.2.1.

²⁰¹ See chapter I.

help us elucidate formulaic behaviour? It will not be a huge spoiler to say that the answer to these questions is positive: not only can we draw parallels between some classes of multi-word expressions and formulae, but methods that have been applied to the description of these expressions can be productively transferred to the study of formulae.

This chapter is not a renewed attempt to draw a straightforward correspondence between formulae and a specific, well-defined class of linguistic structures; in fact, the existence of such well-defined classes outside of formulaic languages is in itself theoretically questionable.²⁰² Rather, this section will start with a discussion of a class of multi-word expressions, that is, idioms, whose behaviour has parallels with that of formulae, but which also show characteristics that clearly set them apart from the object of this study. Specifically, while idioms, like formulae, show restricted semantic flexibility, they have also been characterised through their non-compositional semantics, a trait that is not shared by formulae. An overview of idiomatic expressions is necessary to establish whether and to what extent they are a useful *comparandum*, and how far methods that have been applied to the study of idioms are suitable for Homeric formulae.

III.1 Idioms and formulae

III.1.1 *Multiword expressions and formulae: an initial approach*

On a theoretical level, an attempt to draw connections between formulae and multiword expressions in natural languages, that is, expressions that are made of words that commonly occur together, was made by Paul Kiparsky as far back as 1976.²⁰³ The author starts from the consideration that formulae are naturally to be compared with “the *bound expressions* of ordinary language,” that is, expressions, ranging from collocations to idioms

²⁰² See §III.2.1 below.

²⁰³ Bozzone (2014a), 26-8, also gives a brief overview of Kiparsky’s proposal.

proper, that share (at least in part) the three properties of “*arbitrarily limited distribution*, *frozen syntax*, and *non-compositional semantics*.”²⁰⁴ In other words: there are linguistic entities that show non-random cooccurrence patterns (for instance, bound phrases such as *foregone conclusion* or *fraught with danger*²⁰⁵), limited syntactic and semantic flexibility (it is exceedingly difficult to accept a phrase like *?this conclusion was not foregone*, while *fraught with* never occurs with positively connotated nouns such as **fraught with good intentions*), and whose meaning is not (entirely) a function of the meaning of their parts;²⁰⁶ and all of these are potentially useful *comparanda* for formulae. Varying levels of these three features can be observed in both different linguistic structures and different kinds of formulae. According to Kiparsky, these parallels can be leveraged to achieve a more fine-grained classification of formulaic expressions: flexible formulae should be grouped together with modifiable structures such as “co-occurrence restrictions [...] between lexical items,” while fixed formulae are best analysed as fixed phrases parallel to the most rigid kinds of idioms.²⁰⁷

Since Kiparsky’s contribution, linguistic research on multi-word expressions has progressed significantly; most importantly, it has moved beyond the lexically-based paradigm adopted by Kiparsky, according to which bound expressions are largely comparable to compound words.²⁰⁸ New approaches, especially computational and corpus-based ones, have become widespread, and can provide the tools to transform Kiparsky’s theoretical proposal into a working model for the analysis of the Homeric language.

²⁰⁴ Both quotes are from Kiparsky (1976), 74. All emphasis original.

²⁰⁵ See Kiparsky (1976), 75.

²⁰⁶ We will come back to the latter two criteria in §III.1.1-III.1.5; the first criterion, non-random cooccurrence patterns, is quite transparently parallel to Hainsworth’s ‘mutual expectation’ (see §II.2.1).

²⁰⁷ Kiparsky (1976), 82.

²⁰⁸ Kiparsky (1976), 74: elements of bound expressions “function in certain respects like the parts of complex words usually do.”

The definition and behaviour of idioms, in particular,²⁰⁹ has been a source of debate in linguistics, both theoretical and applied, at least since the 1970s, when approaches arose that questioned the position of idioms in generative syntax. According to the traditional definition, an idiom is “a constituent or series of constituents for which the semantic interpretation is not a compositional function of the formatives of which it is composed;”²¹⁰ in other words, since idiomatic meanings rely on conventions that are different from those that govern the use of their constituents, the meaning of an idiom is not the sum of the meanings of its parts.²¹¹ Another way to view this is through the idea that idiomatic meanings rely on at least some level of figuration, that is, idioms map onto their meanings through some form of metaphor, hyperbole, etc.,²¹² although the strict divide between literal and figurative meanings has also been questioned.²¹³

The first half of this chapter (down to §III.2) will review different definitions of idioms, the most restricted class of multiword expressions; we will conclude with a discussion of how far idioms and formulae can be considered similar structures, before moving on to an approach that was originally applied to idioms, but can productively be transferred to the description of formulae.

²⁰⁹ And indeed, as we will see in §III.2, the necessity of theorising a separate class of idioms at all.

²¹⁰ Fraser (1970), 22.

²¹¹ Cf. Nunberg, Sag, and Wasow (1994), 492, and Titone and Connine (1999), 1663-4.

²¹² Nunberg, Sag, and Wasow (1994), 492-3. For discussions focusing specifically on these metaphorical relations, see Lakoff (1987), 447-53, and Clausner and Croft (1997). It is important to remember that not all metaphors are idioms: metaphors need at least to be consolidated into ‘common’ language use to become idiomatic.

²¹³ See Matlock (1998), 643-4 and 647, who argues that literalness is a gradient, rather than a straightforward black/white distinction, and points out that several idioms contain polysemous words, which complicates the notion that *one* literal meaning exists independently of the figurative one.

III.1.2 Idiom flexibility and processing: decomposability

According to the traditional generative view, the fact that the meaning of an idiom is not a function of its parts implies that idioms are stored and processed lexically, that is, as if they were compound words.²¹⁴ One problem for this paradigm is posed by the different degrees of syntactic flexibility that different kinds of idioms enjoy: there seems to be a difference between cases such as *lay down the law* (= *issue authoritative instructions*²¹⁵), which easily passivizes to *the law has been laid down*, and phrases like *shoot the breeze* (= *have a casual conversation*), which is hardly acceptable in the passive (**the breeze is/was shot*).

A widespread theory of idiomaticity analyses idioms into different classes depending on their degree of decomposability, or analysability.²¹⁶ This concept, first introduced in the work of Geoffrey Nunberg,²¹⁷ divides idioms into two classes, (semantically) nondecomposable and decomposable, depending on whether parts of the idiom's meaning can be mapped onto its constituents. Thus, for instance, an idiom such as *spill the beans*, meaning *reveal a secret*, is decomposable, because of the metaphorical relation connecting, respectively, *spill* to *reveal* and *beans* to *secret*; on the other hand, *kick the bucket*, meaning *die*,²¹⁸ is nondecomposable, because neither *kick* or *bucket* has a specific relationship with any part of the idiomatic meaning. Note that decomposability depends on whether *parts* of the idiom can semantically map onto *parts* of the idiomatic meaning, and has nothing to do with, for instance, how easily the metaphor behind an idiom can be understood; so, for instance, an idiom such as *see stars* = *be partially unconscious*, while it relies on real-world reference in a relatively transparent way (a person about to fall unconscious will often see

²¹⁴ See e.g. Fraser (1970).

²¹⁵ All idiom paraphrases have been taken from or checked against the *Oxford Dictionary of English Idioms* (Ayto 2009).

²¹⁶ Terminology in the literature varies; I have adopted 'decomposability' throughout as it is the least ambiguous term.

²¹⁷ Cf. Nunberg (1978), 117-35, and the seminal paper on this topic by Nunberg, Sag, and Wasow (1994).

²¹⁸ As noted in Wulff (2008), 25, this idiom, despite being one of the most used examples in the literature, is in fact rarely attested.

flashes of light), is still nondecomposable, as there is no specific meaning to be ascribed to either *see* or *stars*.²¹⁹ Explicability does not correlate reliably with idiom flexibility,²²⁰ even though the presence of an available metaphor schema does promote the formation of alternative idiomatic expressions for the same content (for instance *let the cat out of the bag* = *spill the beans*).²²¹ Explicability is also heavily influenced by real-world knowledge: for instance, the meaning of *carry coals to Newcastle* = *perform an unnecessary activity* is easily explicable for someone familiar with the city's mining history, but likely to be obscure to anyone else.²²²

Decomposability and compositionality rely on, so to speak, two symmetrical modes of semantic interpretation: while idioms are not compositional, that is, the meaning of their parts does not map onto the meaning of the whole expression, once the meaning of the whole expression is known it may be possible for it to map onto the meaning of the parts (through more or less transparent metaphorical relations). It must be noted, however, that *spill the beans* and *kick the bucket* are two extreme examples, falling at opposite ends of a continuum of decomposability that spans a range of more ambiguous cases.²²³ Some of these can be analysed as either decomposable or nondecomposable: for instance, an idiom such as *carry coals to Newcastle* is nondecomposable if taken to mean *perform an unnecessary activity*, but very easily decomposable if understood in the narrower sense of *bring something to a place where it is already plentiful*.²²⁴ Speakers' judgements for idioms

²¹⁹ There is definitely some variation between individual speakers here, to the extent that some may even see this phrase as having a literal meaning (a person who is about to faint may see flashes of light). As we will see in the course of this chapter, variation between speakers, to the extent that experimental results cannot be replicated, is a significant problem with the idiom decomposition hypothesis.

²²⁰ Reagan (1987), 431-6.

²²¹ Clausner and Croft (1997).

²²² Again, it is possible that some speakers may acquire the real-world knowledge from the idiom, which complicates the issue somewhat.

²²³ Gibbs and Nayak (1989), 106-7.

²²⁴ Both paraphrases are listed, for slightly different forms of the idiom, in the *Oxford dictionary of English idioms*, s.v. 'coal.' Espinal and Mateu (2010) discuss the case of an expression (*V one's [head/other body part] off*) which partly fits the criteria for decomposable idioms, partly for nondecomposable ones, especially as pertains to its aspectual analysis.

such as this one are not always reliable, as different speakers will tend to interpret the idiom in either the broader (nondecomposable) sense or the narrower (decomposable) one.²²⁵ Indeed, the main psycholinguistic evidence for the decomposable/nondecomposable distinction comes from a study performed by Gibbs and Nayak in the late 80s,²²⁶ where speakers were explicitly asked to rate whether parts of an idiom contributed to parts of the literal interpretation; subsequent experiments, however, have failed to reliably replicate these results,²²⁷ going so far as to suggest that they may be due to an experimental artefact in the original study.²²⁸ Since decomposability only exists as a psycholinguistic measure, evaluated through speakers' judgments, the unreliability of these judgements is obviously a problem for the general theory.

III.1.3 *Idiom flexibility and processing: the idiom decomposition hypothesis*

Despite these issues, there is some evidence in favour of a correlation between idiom decomposability and syntactic flexibility. The latter amounts to the extent to which an idiom can undergo "meaning-preserving syntactic operations,"²²⁹ that is, whether the result of a specific syntactic operation has the same meaning as the original idiom. Thus, for instance, *the law was laid down by Sam* maintains the same idiomatic meaning (in an appropriate context) as *Sam laid down the law*, while *the bucket was kicked by Kim* has a very

²²⁵ Decomposability judgments for 171 idioms have been systematically collected from speakers of American English by Titone and Connine (1994).

²²⁶ Gibbs and Nayak (1989), experiment 1.

²²⁷ Cf. Titone and Connine (1994) and Tabossi, Fanari, and Wolf (2008).

²²⁸ Tabossi, Fanari, and Wolf (2008), 316: "It is likely that the initial selection made by Gibbs and Nayak (1989) on the basis of their own intuitions led them to use idioms on which people agree, thus overestimating speakers' general ability to judge compositionality. However, with a larger sample of non-preselected expressions (80 vs. 40), we observed that speakers' agreement becomes less evident." The judgments collected in this article are sometimes strikingly inconsistent (see p. 322 for an example where an idiom received both the lowest and highest ratings on the scale); this may have to do with experimental design, as subjects were given an explicit definition of decomposability that may have been misunderstood (this issue is shared by Gibbs and Nayak's original study).

²²⁹ Cacciari and Glucksberg (1991), 231.

different meaning from *Kim kicked the bucket* (it can only be interpreted literally as *Kim booted the pail out of the way*).²³⁰

In the same study mentioned at the end of §III.1.2, Gibbs and Nayak tested the syntactic flexibility of decomposable vs nondecomposable idioms through five syntactic changes: present participle (*Sam is laying down the law*), adverb insertion (*Sam quickly laid down the law*), adjective insertion (*Sam lay down the office law*), passive (*the law was laid down by Sam*), and action nominalization (*Sam's laying down of the law*). Subjects were asked to rate the similarity between idioms that had undergone one of the syntactic changes and paraphrases of their idiomatic meaning.²³¹ According to Gibbs and Nayak's results, not only are decomposable idioms more flexible, but there is a difference in the level of disruptiveness of each syntactic change; specifically, all types of idiom can undergo syntactic changes that operate on the entire expression (for instance, present participle and adverb insertion), while changes that only have scope over individual parts (such as adjective insertion and passive) are only possible with decomposable idioms.²³² While Gibbs and Nayak's data are quite clear, their interpretation is not entirely: for instance, it is not particularly convincing for the passive operation to be considered to have scope only over part of the idiomatic expression.

²³⁰ All of these comments on idiomatic meaning must be taken with a grain of salt: it is always possible for speakers to be creative and exploit e.g. real-world context to preserve idiomatic meaning in a humorous or artistic way. A rather grim example might be someone announcing, after a person has passed away after a prolonged agony, that *the bucket has finally been kicked*: in this case, the real-world context of an impending death may be enough to constrain the interpretation of the phrase to the idiomatic meaning despite the forbidden syntactic operation.

²³¹ Again, this experimental design may have biased the subjects' ratings, as it draws attention to the idiomatic meaning while potentially obscuring the literal one: cf. Lebani, Senaldi, and Lenci (2015), 5.

²³² Gibbs and Nayak (1989), experiment 2. A subcategory of 'abnormally decomposable' idioms also shows limited syntactic flexibility compared to 'normally decomposable' ones. The traditional interpretation is that parts of these idioms map onto their meanings through highly conventional metaphors; therefore, they behave as if they mapped onto the relation rather than the meaning itself (Nunberg 1978, 127-9). This further distinction need not concern us here.

In any case, this evidence was taken as the basis of the so-called ‘idiom decomposition hypothesis,’ according to which “idioms are partially analyzable and speakers’ assumptions about how the meaning of the parts contribute to the figurative meanings of the whole determines the syntactic behavior of idioms.”²³³ This conclusion is not limited to syntactic flexibility: decomposable idioms are also more lexically flexible, that is, they maintain their meaning more easily after lexical substitution than their nondecomposable counterparts (*lay down the rules* has the same meaning as *lay down the law*, while *kick the pail* means, well, literally kick a literal pail).²³⁴ For decomposable idioms, whose parts have independent meaning, “the relation among these parts can be specified in conceptual structure;”²³⁵ therefore, syntactic operations that rely on moving around, emphasising, or modifying some parts of a phrase in its conceptual structure, such as left-dislocation (*the law, Sam laid it down*), passivation (*the law was laid down*), or the insertion of modifiers (*Sam immediately laid down the law*), are possible. The opposite holds for nondecomposable idioms, whose parts stand in no meaningful relation with each other. However, plausible motivation also plays a key role in judging syntactic flexibility: “a syntactic operation on an idiom will be acceptable if and only if it produces a comprehensible difference in interpretation, that is, a reasonable communicative intention for the syntactic operation can be inferred.”²³⁶ All of these issues related to communicative strategies complicate the apparently clear-cut results of the idiom decomposition hypothesis.

III.1.4 Idiom flexibility and processing: alternative hypotheses

In addition to the issues that were briefly mentioned in §III.1.3, there is the fundamental problem that Gibbs and Nayak’s apparently clear-cut results have not been reliably

²³³ Gibbs and Nayak (1989), 104.

²³⁴ Gibbs et al. (1989).

²³⁵ Jackendoff (1995), 152.

²³⁶ Cacciari and Glucksberg (1991), 233.

replicated. Subsequent experiments failed to find a significant difference between the syntactic flexibility of decomposable and nondecomposable idioms, and even evidence for the scale of disruptiveness of different syntactic operations is lacking.²³⁷ Similarly, experiments on the production of idiom blends (*we'll burn that bridge when we come to it*, a blend of *burning bridges* and *we'll cross that bridge when we come to it*) have found no difference between blends of nondecomposable and decomposable idioms, suggesting that these two classes are processed in similar ways. What does have an influence on the production of idiom blends is the syntactic structure of the original idioms, which hints that the role of syntax may be more significant in idiom processing than the distinction between decomposable and nondecomposable idioms.²³⁸ In short, decomposability may not be as decisive a factor as it appears from Gibbs and Nayak's experiments.

As an alternative to the classification into decomposable and nondecomposable idioms, the so-called 'configuration hypothesis' offers the idea that all idioms, independently of their level of decomposability, are processed literally until the so-called 'idiom key,' that is, the point where enough information is accumulated for the literal meaning to be ruled out, is reached.²³⁹ This theory is supported by evidence from semantic priming experiments.²⁴⁰ Experiments have compared predictable idioms, that is, phrases that have no possible completion other than the one that produces an idiomatic reading, and for which therefore the idiom key appears very early, to non-predictable idioms, that is, phrases that can be completed by producing an idiomatic or non-idiomatic reading, and for which the idiom key

²³⁷ Tabossi, Fanari, and Wolf (2008), experiment 2. Significant differences are observed only in the case of adverb insertion.

²³⁸ Cutting and Bock (1997).

²³⁹ On the configuration hypothesis and the evidence supporting it, see e.g. Cacciari and Tabossi (1988), Cacciari and Glucksberg (1991), Tabossi, Fanari, and Wolf (2005).

²⁴⁰ In these experiments, subjects are asked to respond to specific stimuli (e.g. press a button if a word on a screen has certain characteristics) after having been exposed to apparently unrelated material. The underlying idea is that subject responses will be faster if the introductory material is related to the stimulus, and delayed if it is unrelated. Priming can also influence responses: a classic example is the game where a speaker makes another repeat the word 'silk' several times, and then asks 'what do cows drink?' – leading in most cases to the nonsensical, but heavily primed, response 'milk!'

coincides with the end of the expression.²⁴¹ Semantic priming effects for words connected to the idiomatic meaning are observed if the target word is presented at any point during the elaboration of a predictable idiom, but only at the end of a non-predictable idiom. On the other hand, negative priming effects for words related to the literal meaning of an idiom can only be observed when the idiom is predictable; this is taken as evidence that literal meaning is always computed until the idiom key is reached (early for predictable idioms, but only at the very end for unpredictable idioms).

The configuration hypothesis emphasises the importance of predictability as a factor in idiom processing, which implies that the syntactic structure of unpredictable idioms, at least, is always computed up to the idiom key; it does not, however, account for syntactic flexibility as well as the idiom decomposition hypothesis does.

III.1.5 *Are formulae idioms?*

Regardless of the specific theoretical issues, the debate surveyed in this section has highlighted one aspect that clearly sets idioms apart from formulae. All the approaches discussed above, despite their differences, consider semantic noncompositionality an essential element in the definition of a class of idioms. Whether the meaning of an idiom is to be discussed in terms of its decomposability or its predictability, it is still not a function of the meaning of its parts, while the opposite is generally true of non-idiomatic utterances. This is quite transparently not true of formulae, which are (mostly) compositional: the meaning of, for example, θεῖος ἀοιδός is straightforwardly composed of the meaning of

²⁴¹ Compare e.g. Italian *non sapere che pesci prendere* (= *not knowing what to do*, lit. *not knowing what fish to catch*), which is unambiguously idiomatic after the word *pesci* is reached, to *fare un buco nell'acqua* (= *fail to obtain anything*, lit. *make a hole in the water*), which easily accepts non-idiomatic completions until the very last word, such as *fare un buco nell'albero* (= *make a hole in the tree*). (Examples are from Tabossi, Fanari, and Wolf 2005, slightly modified.)

θεῖος and the meaning of αἰδοός.²⁴² This, in turn, renders any discussion of decomposability entirely moot, as all formulae of this kind will be by definition decomposable. This is, in some sense, a fortunate circumstance for the researcher: since decomposability is generally regarded as a subjective measure, with no reality outside speakers' minds, it would be exceedingly hard to apply it the study of a dead language, where no speaker input can be sought.

There are only a few exceptions to this last point. A potential counterexample is provided by semantically opaque expressions such as νυκτὸς ἀμολγῶ, 'in the dead of night,' whose meaning appears to be taken as a unit in our transmitted texts.²⁴³ Similarly, counter-intuitive name-epithet combinations such as ἀμύμονος Αἰγίσθοιο, 'blameless Aegisthus' (*Od.* 1.25), may be another example of non-compositional semantics, where the meaning of the epithet simply disappears and the whole expression means 'Aegisthus' and nothing else. While these are somewhat marginal cases, their behaviour appears to reinforce at least one conclusion from the studies reviewed above, that is, that not all syntactic operations are equally disruptive when dealing with noncompositional expressions.²⁴⁴ There is good reason to think that this consideration should also apply to syntactic change in formulae. As an example, νυκτὸς ἀμολγῶ is not syntactically flexible, but can undergo basic expansion: it appears as ἐν νυκτὸς ἀμολγῶ in *Il.* 11.173 and as μελαίνης νυκτὸς ἀμολγῶ in *Il.* 15.324. This second example is broadly equivalent to cases of adjective insertion in Gibbs and Nayak,²⁴⁵ an operation that was generally permitted for decomposable idioms; indeed, while the precise meaning of the elements of the Homeric

²⁴² I am not considering here notions such as traditional referentiality, i.e. the idea that "the repetition of [formulae] breeds an association which adds a connotative level of meaning to the denotative level" (Kelly 2007, 4). The role of referentiality is crucial to the epic language, but does not affect the basic meaning of formulae to the same extent that noncompositionality is fundamentally important to idioms.

²⁴³ On this formula, see the summary and references in Vergados (2012), 226-7.

²⁴⁴ Remember how different syntactic changes led to very different results in Gibbs and Nayak's experiments.

²⁴⁵ The first one should probably be considered equivalent to adverb insertion, since νυκτὸς ἀμολγῶ alone already has locative meaning.

collocation is not transparent, its meaning is without doubt decomposable (with νυκτός meaning 'night' and ἀμολγῶ adding some kind of information regarding the hour).

All of this being said, we are left with the basic problem that most formulae are, at least to some degree, compositional: the meaning of *swift ships* is in fact a function of the meaning of *swift* and the meaning of *ships*, with some added meaning provided by the traditional function of the epithet (which means that ships can be swift even when they are anchored and not going anywhere, as the epithet refers to an essential characteristic). Are there any approaches to idiomatic behaviour that can account for these traits, rather than relying on the basic trait of noncompositionality as a unique and defining characteristic of idioms?

III.2 Idiom flexibility in a Construction Grammar perspective

Up to this point, the discussion of idioms, their behaviour and their relation to formulae has focused on the issue of semantic (non)compositionality. Not all theories of syntax, however, base their definition of idioms on this trait; indeed, not everyone agrees on the basic point that we should define a class of idioms that is qualitatively different from 'regular' phrases at all. In this section, I will discuss how this kind of theoretical approach allows us to draw more fruitful parallels between idioms and formulae, and what empirical instruments for the analysis of formulaic expressions we can gain from these parallels.

III.2.1 Idiomaticity in Construction Grammar

An essential trait of Construction Grammar, a framework developed in the 1980s, is that as opposed to generative approaches it does not postulate a separation between syntax and lexicon. Both words and syntactic constructions can be defined as pairs of form and

meaning or discourse function.²⁴⁶ These pairings are all, to varying extents, non-compositional, in that the form and/or function of a construction that is made of different elements is not just a function of the form and/or function of the elements themselves.²⁴⁷

In the cognitive-linguistic perspective that is typical of Construction Grammar, all constructions are noncompositional to a certain degree. In practice, “compositionality is not a binary, but a scalar concept:”²⁴⁸ all constructions fall along a continuum that runs from what other approaches would define as entirely compositional, that is, constructions that follow more general semantic rules, to strictly idiomatic, that is, constructions whose meaning is determined by more idiosyncratic rules.²⁴⁹

Flexibility, or variability, both on the lexical and syntactic level, can also be defined as a continuum, rather than a series of binary options. For example, we saw in §III.1.3 that some idioms ‘cannot,’ in a generative approach, be turned into the passive; however, we can always come up with marginal examples where a passive operation would be allowed, given context and a level of playfulness with the language. In a Construction Grammar approach, these idioms are seen as *very unlikely* to be used in the passive, while other idioms (and most non-idiomatic expressions involving transitive verbs) are much more commonly used with a passive verb. The constructionist view has the advantage of better accounting for corpus data, where linguistic structures can occur more or less frequently in a certain syntactic form, and where the range of variability may be wider than the limited set of operations discussed by the studies mentioned above.

²⁴⁶ Cf. e.g. Goldberg (2013), 17; Goldberg (1995), 7. On the use of CxG as an approach to Homeric formulae, see above, §II.3.2.

²⁴⁷ Cf. Goldberg (2006), 5: “Any linguistic pattern is recognized as a construction as long as some aspect of its form or function is not strictly predictable from its component parts or from other constructions recognized to exist.”

²⁴⁸ Wulff (2013), 280. This paper also provides a historical overview of approaches to idioms in CxG (pp. 274-9).

²⁴⁹ Wulff (2008), 18.

Consider an example that would not be considered an idiom in the traditional sense, such as the British English P N construction: [*be*] *in/at school, in prison, in hospital, on vacation, at work, etc.*²⁵⁰ The semantics of this construction is different from that of the corresponding P Det N construction: it is strictly associated with performing an activity connected to the location mentioned (for instance studying/learning *for school*) and does not always convey any information about the physical position of the referent (a pupil can be *in school* even when they happen to be absent for one day, although the same does not apply to *in hospital*). Therefore, from a syntactic standpoint, the construction does not accept the insertion of an article without a substantial change in meaning (visitors and staff can be *in/at the school*, but only pupils are *in/at school*). The productivity of the construction's noun slot is also restricted (one can be *at the computer*, but not **at computer*), as is the range of adjectival modification it allows (*in primary school, at private school*, but not **in excellent school*). In other words, the semantics of the English P N construction is also, like its syntax, not fully compositional: the additional information about 'performing the activity that is commonly performed in a specific setting' is not easily defined as a function of the meaning of the preposition compounded with the meaning of the noun. There is, however, no degree of 'figurativeness' involved in the construction's non-compositionality, which is partly why we do not consider it to be an idiom in traditional terms.

In this framework, idioms do not form a separate, anomalous class of (more or less) lexicalised multi-word expressions; the difference between them and other constructions is a matter of degree, rather than a qualitative split.²⁵¹ Getting rid of figurativeness and

²⁵⁰ Discussed in Goldberg (2013), 17-18. It should be noted that this construction shows a range of dialectal variation; my examples are in English as spoken in the South-East of England.

²⁵¹ Cf. e.g. Wulff (2013), 288: "What may license referring to some constructions as idioms and not others is merely a reflection of the fact that effects of idiomatic variation are *best observable* in partially schematic complex constructions—however, this does not make them *fundamentally* different in nature from other constructions." The conceptualisation of linguistic phenomena as a gradient, rather than a set of bounded categories, is connected with a usage-based perspective on language (Bybee 2013, 50); see §III.2.2 below.

metaphorical relations as an essential part of the definition of idiomaticity is part of this process. We have seen above an example of a construction that is non-compositional and limited in flexibility (that is, shows a higher level of idiomaticity than the average English construction) but does not involve any element of figuration. Constructions and idioms can still, of course, involve some level of figurativeness in their meaning (for instance, the relationship between *drive [OBJ] to distraction* and *drive [OBJ] to London* is metaphorical²⁵²), but there is no necessary connection between figurativeness and idiomaticity.

This approach also allows research on idioms to accommodate a wider range of phenomena.²⁵³ If necessary, we can still define a class of ‘core idioms,’²⁵⁴ that is, constructions with a high degree of semantic idiosyncrasy and very limited syntactic flexibility; it is important, however, to keep in mind that both these traits are still present to a varying degree throughout the constructional lexicon. The degree of correlation between semantic idiosyncrasy/noncompositionality and syntactic flexibility (plus, potentially, other relevant traits) can be explored on a more fine-grained level if these factors are conceptualised as scalar measures.²⁵⁵ There are noteworthy parallels here between the definition of idioms as a class vs part of a scalar continuum and the debate around the definition of formula in archaic Greek epic. We could indeed consider a category of ‘core formulae,’ that is, formulae that occur as part of a system which respects the principle of thrift, to be merely the top layer in a continuum of formularity that can also accommodate more abstract definitions, such as ‘structural formulae.’²⁵⁶

²⁵² See Goldberg (2013), 22.

²⁵³ On the range of expressions that do not fit under the traditional definition of idiom, but show restricted productivity and other similar traits, see also Jackendoff (1995).

²⁵⁴ The terminology here is borrowed from Wulff (2008), 2.

²⁵⁵ Wulff (2013), 278: “different factors like lexical specification, semantic irregularity, syntactic irregularity, and cognitive entrenchment prevail at all levels of the constructicon [lexicon of constructions, MAR] in *different shades of prominence and relative importance*” (emphasis original).

²⁵⁶ For a fuller exploration of this idea see Bozzone (2014a), 21-4 and 37-9.

By doing away with the treatment of idioms as a separate class from other expressions, and especially with the idea that this class is characterised by the presence of figurative meaning relations, constructionist approaches to syntax can provide us with the tools to discuss formulae (which may not form a natural class, and which certainly do not base their meaning on metaphorical relations) as idioms. The next step is to consider whether this can further our understanding of formulaic flexibility.

III.2.2 A corpus-based study of idiomaticity

Constructionist approaches also tend to subscribe to a usage-based perspective, that is, the idea that we learn language and its rules from exposure to the language itself in use, rather than by resorting to innate schemas: “experience with language creates and impacts the cognitive representations for language.”²⁵⁷ Corpus-based studies are an essential tool in this respect: by focusing on representative samples of actual language usage, they give researchers access to the input to which speakers are exposed, and therefore enable hypotheses on the generalisations on which their judgements are based.²⁵⁸ This is even more fundamental when direct access to the speakers and the gathering of new data are not a feasible route, as is obviously the case when working on a dead language.

When using corpus data to describe the idiomaticity continuum, however, one must face the obstacle presented by the fact that idiomaticity in itself is still a psychological measure, and can therefore only be measured through speakers’ judgements. It is assumed that speakers form their judgements on the basis of some characteristic(s) of idiomatic expressions, which should be observable in corpus data, even if we cannot know in advance which traits are relevant and to what extent they ultimately contribute to the positioning

²⁵⁷ Bybee (2013), 49.

²⁵⁸ Cf. e.g. Wulff (2013), 278. For a more nuanced discussion of quantitative approaches in corpus historical linguistics, cf. the Introduction to Jensen and McGillivray (2017).

of an expression along the idiomaticity continuum. These traits, once isolated from corpus data, can be used to define idiomaticity from a corpus-based (and usage-based) perspective.

The reasoning outlined above forms the basis of a corpus study conducted by Stefanie Wulff on the behaviour of English V NP idioms (of the type *take the plunge* or, to a lesser degree of idiomaticity, *see a point*). Wulff starts from the consideration that, even if we do not have access to the idiomaticity continuum directly, we can still measure what she calls the 'idiomatic variation continuum,' that is, the range of syntactic and semantic flexibility of the idioms observable in a corpus. The idiomaticity continuum is therefore a psychological, internal measure, while the idiomatic variation continuum is observable in our data: "the idiomaticity continuum is what speakers construct on the basis of the information about constructions that they consider salient to help them assign the construction a specific position on the continuum; the idiomatic variation continuum is defined by the range of values that constructions actually take with respect to their semantic and syntactic behaviour."²⁵⁹ The ultimate aim of Wulff's study is to investigate the correlation between the two continua, by assessing how well different measures of idiomatic variation can predict speakers' idiomaticity judgements.

Wulff analyses both the semantic and syntactic behaviour (broadly defined) of V NP idioms. On the semantic level, she provides a measure of compositionality based on distributional semantics, an approach that defines the meaning of words and expressions as a function of their collocates in their linguistic environment.²⁶⁰ In this approach, one can measure how closely the meaning of a construction as a whole is related to the meaning of each part by measuring how many of the collocates of the construction are shared by each of its

²⁵⁹ Wulff (2008), 4-5.

²⁶⁰ For an early formulation of this view, cf. Garvin (1962), 388: "Distributional semantics is predicated on the assumption that linguistic units with certain semantic similarities also share certain similarities in the relevant environments." Distributional semantics will be the main focus in §VI-VII of this dissertation, and will be discussed at length there.

constituents. Compositionality is therefore computed by summing the weighted contribution of the collocates of each part (that is, the verb and the noun) to the collocates of the whole V NP construction. To obtain a more reliable measure, Wulff considers both how many collocates of each part occur among the collocates of the whole construction and the reverse, that is, how many collocates of the construction occur among the collocates of each part; these are then summed to obtain a compositionality measure for the whole construction.²⁶¹

The formal flexibility of constructions is investigated through three sets of measures, which Wulff calls tree-syntactic, lexico-syntactic, and morphological flexibility. Tree-syntactic flexibility concerns “the flexibility with which a phrase occurs in different syntactic patterns;”²⁶² the set of patterns considered by Wulff includes declarative active (*make a point*) and passive (*the point is made*), relative clause both active and passive (*the point that they made/that was made*), to-clause (*the point to make/to be made*), interrogative (*did they make a point?/was a point made?*), and imperative (*make a point!*). Lexico-syntactic flexibility corresponds to modification through “the addition of some kind of material, be it internal or external modification;” Wulff keeps the insertion of material in the verb and noun slot separate. Finally, morphological flexibility involves “any kind of variation at the morphemic level,” that is, in person, tense, aspect, mood, voice (for the verb),²⁶³ number (for both verb and noun), presence of a negation and/or a determiner. Note that determiners and negations are considered to be morphemes, and therefore belong to the morphological flexibility class.

²⁶¹ Wulff (2008), 55-9. This is a relatively simplistic measure of compositionality, in that it does not account for the relation between the meanings of the collocates themselves; superficially different collocates can be semantically close to each other (see, again, the discussion of Distributional Semantics in §VI.2).

²⁶² All quotations in this paragraph are from Wulff (2008), 76. Examples are picked from the whole of Wulff’s chapter 4.

²⁶³ The fact that the passive voice is counted both as a syntactic operation and as a morphological parameter leads to some redundancy in the results (Wulff 2008, 154).

The flexibility of a construction along each of these axes, however, is not a self-sufficient measure, as Wulff notes: among other things, a construction will behave differently depending on the kind of variation being measured. Compare, for example, tense variation with a more invasive operation such as the passive imperative. Any V NP construction will be relatively likely to undergo tense variation; passive imperatives, on the other hand, are always going to be extremely rare, no matter the level of idiomaticity of the construction. Therefore, we can always expect to find more examples of an idiomatic construction in the past or future tense than in the passive imperative. Comparing the raw frequencies of these two flexibility measures will not provide any meaningful result, as one will always be lower than the other independently of idiomatic status. What matters is whether each idiom is more or less flexible than the average V NP construction, and whether some types of modification that are very common in non-idiomatic V NP pairs are avoided by idiomatic ones.²⁶⁴

Wulff measures this by comparing the construction's flexibility to a baseline level provided by the average flexibility of V NP constructions in a corpus reflecting contemporary English usage.²⁶⁵ We do not want to know how often *make a point* appears in the past tense or in the passive imperative; we want to know if, for any reason, it is significantly more unlikely to observe *made a point* than the average V NP phrase in the past tense, or indeed if *make a point* is somehow significantly more likely to occur in the passive imperative. To take this into account, flexibility is measured for each type of variation as the normalised squared sum of the deviations from the baseline flexibility of the construction class. This measure, since it involves a squared factor, only allows us to gauge how far each construction differs from the baseline, but not the direction of the variation (that is, whether the construction is more or less flexible than the baseline); for this reason, Wulff also introduces a measure

²⁶⁴ I use non-idiomatic and (most/more) idiomatic, as Wulff does, to refer to opposite ends of what we should still view as a continuum.

²⁶⁵ Wulff (2008), 77.

based on directional relative entropy, which also measures whether the deviation is towards more or less flexibility.²⁶⁶

This method allows Wulff to measure both the compositionality and flexibility of each V NP construction, compared to the baseline, along each of the axes discussed above; effectively, this means she is placing each construction somewhere on the idiomatic variation continuum. Different measures of flexibility are kept separate; each of them forms a scale on which each V NP idiom is situated, based on the frequency with which it undergoes the change in question compared to the baseline frequency of the change itself in the corpus.²⁶⁷ At this stage, Wulff is merely documenting the behaviour (range of flexibility) of each idiomatic construction in relation to different variation parameters, without attempting to produce a unified measure of idiomaticity or a single idiomatic variation scale. This is a consequence of the theoretical premise mentioned at the start of this section, that is, that while we expect there to be some kind of relation between idiomatic variation and idiomaticity, we do not yet know to what extent each idiomatic variation parameter contributes to idiomaticity judgements.

In order to address this question, the final step in the study is to explore the correlation between the variation as measured in the corpus and speakers' judgements collected through a questionnaire. Wulff performs a Principal Component Analysis of the 20 original variation parameters, and then a Multiple Regression Analysis with the speakers' idiomaticity judgements as the dependent variable and the idiomatic variation parameters as independent variables. The aim of the PCA is to isolate the factors that account for most of the variance of the overall distribution of constructions along the idiomatic variation continuum; the MRA, on the other hand, is used to identify which variation parameters are

²⁶⁶ Wulff (2008), 78-87. We will return to relative entropy in §IV.2.

²⁶⁷ See figg. 4.1-36 in Wulff (2008). Some aspects on this method will be reprised in §IV.3.

the best predictors of the speakers' judgements over whether a construction is an idiom.²⁶⁸ The results of the two analyses overlap strikingly (and satisfyingly): that is, the best predictors of perceived idiomaticity are also the factors that account for the most variance in the corpus data. This suggests that the "speakers' primary strategy is to concentrate on those distributional parameters which have the highest information value."²⁶⁹ speakers use those parameters in which the behaviour of idiomatic constructions deviates the most from non-idiomatic ones to construct their idiomaticity judgements. This confirms that the corpus measures evaluated in the study provide an adequate model of what guides speakers' judgements – that is, that the idiomatic variation continuum can be used to model the idiomaticity continuum.

III.2.3 Modelling formulaic flexibility as constructional variation

The final outcome of Wulff's study – that the idiomaticity continuum reflects the idiomatic variation continuum – is the most relevant to our discussion of Homeric formulaicity. Wulff's approach represents a promising quantitative method that can be used to evaluate formulaic flexibility along the three axes of syntax, morphology, and formulaic expansion (that is, tree-syntactic, morphological and lexical flexibility). Gathering data on the behaviour of Homeric formulae along these axes of variation would allow us to compare their variability to that of similar expressions in the language of later, non-formulaic texts. This comparison is not merely a descriptive endeavour, but can further our understanding of the factors involved in formulaic flexibility, much like the evidence reviewed in §III.2.2 regarding the correlation between variation parameters and idiomaticity judgements sheds light on the deeper cognitive significance of Wulff's results. Comparison with a corpus of later (Hellenistic and/or Imperial) epic poetry will also allow us to control for the potential

²⁶⁸ See Wulff (2008), 150-1 (PCA) and 157-8 (MRA).

²⁶⁹ Wulff (2008), 161.

influence of metre and genre, and therefore to be able to focus on the issue of oral-formulaic vs non-oral composition.

Wulff's approach must, of course, be adapted to an analysis of epic formulaicity. Formulae are not idioms, and they should not be expected to vary along the same axes;²⁷⁰ their formal characteristics (metre and syntactic shape) also pose challenges that need to be addressed. In particular, we may expect operations that affect syllable weight and metrical shape to be more disruptive in the context of hexameter poetry, a factor that is completely absent from the English data. From a theoretical standpoint, among Wulff's measures, compositionality is not relevant for Homeric formulae, for the reasons outlined in §III.1.5: beyond the issue of traditional referentiality, there is little question that the meaning of formulae is in fact a function of the meaning of their components. While other semantic measures, such as the flexibility of open slots in formulaic constructions, may still be worth exploring, in line with recent studies of the influence of semantic openness on constructional productivity, this discussion is best kept separate from the analysis of syntactic flexibility.²⁷¹ Wulff explicitly excludes the possibility of lexical substitution in idiomatic phrases; in her model, the use of a different lexeme is viewed as producing an entirely new construction.²⁷² In the Homeric language, however, the presence of metre provides an additional constraint that could justify a comparison between formulae with different lexical fillers of the same metrical shape.²⁷³

On the other hand, some of Wulff's remaining parameters bear remarkable similarities to elements that have already been considered in the study of Homeric formulae – in

²⁷⁰ Although we will see below that Wulff's measures provide a satisfactory level of coverage for the behaviour of N + Adj formulae.

²⁷¹ See Perek (2016) and, with less reliance on corpus data, Barðdal (2008). All of these issues will be addressed in §VI-VII.

²⁷² Wulff (2008), 76. This is the same choice made by Hainsworth (1968) with formulae.

²⁷³ This does not go quite as far as positing the existence of structural formulae (see §I.3), as most of the construction should still be kept constant (e.g. by comparing constructions with the same verb but differently filled noun slots); it is closer to the approach to formulae as partially lexically filled constructions outlined in Bozzone (2014a), 40-2.

particular, to Hainsworth's description of formulaic variation (see §II.2). Wulff's lexical flexibility measure corresponds neatly with what Hainsworth calls expansion (by means of epithets, that is adjectives, adverbs, articles, or a coordinated synonym); morphological variation translates to formulaic declension, a parameter considered by both Hainsworth and Hoekstra,²⁷⁴ although the level of possible variation in an inflected language like ancient Greek is much higher than in English.

Comparison with Hainsworth's work also draws attention to the main trait of Homeric formulae that Wulff's method, as it stands, fails to capture, that is, the presence and role of metre in licensing formulaic variation. Several of Hainsworth's flexibility measures (mobility within the line, word order inversion, and separation) are only significant because the formula is considered to be a metrical unit. One issue in accommodating this parameter into Wulff's model is that metrical constraints do not hold for similar constructions outside of the epic language; this makes it difficult to establish a baseline, although the possibility of comparison with later poetic traditions remains. Aside from this, however, compared with Hainsworth's approach the methods described in §III.2.2 have the advantage of allowing for more precise quantification, and therefore setting up a reliable comparison between corpora that can take us beyond the mere description of formulaic flexibility and towards a better understanding of the factors influencing it.

III.3 A pilot study: the flexibility of Adj + N pairs

III.3.1 Aims and context of the study

The parallels between Wulff's variability measures and Hainsworth's discussion of formulaic flexibility suggest applying the former to a dataset similar to Hainsworth's, as a

²⁷⁴ See above, §II.1.1 and II.2.1.

preliminary step before extending the study to a wider target sample as set out in §III.3.4. This will allow us to adapt Wulff's measures to the range of variability observed in the Greek epic corpus. Moreover, we will be able to start evaluating the variability of a sample of formulae of controlled but reasonable size, as well as assessing more clearly the level of sparseness of the data and their distribution between Homer, Hesiod and the Hymns. This preliminary study can be performed by looking up an existing list of noun-adjective pairs, therefore reserving the issue of corpus processing and data extraction for a later stage. By doing this, we will be able to focus on the characteristics of the dataset and the suitability of Wulff's measures while controlling for potential computational errors.

This study will be based on Hainsworth's list of formulae; while Hainsworth only focused on Homer, we will expand our corpus to include Hesiod (*Theogony* and *Works and Days*) and the major Homeric Hymns (*Demeter*, *Apollo*, *Hermes*, *Aphrodite*). Hainsworth's dataset, as listed in the Appendix to his book, is composed of all Adj + N (or N + Adj) pairs of metrical shape $\text{—}\cup\cup\text{—}\cup$ and $\cup\cup\text{—}\cup$ attested in the *Iliad* and the *Odyssey*, a large and flexible class of formulae that are particularly well suited to mobility.²⁷⁵ Expressions containing proper names and personifications are excluded, as are complex groups (that is, formulae comprising more than one adjective, or an adjective plus adverb or particle) and participial phrases.²⁷⁶ Hainsworth's formulae all come from Homer, which makes his sample inevitably Homer-centric: any expressions that are not attested in Homer but occur in Hesiod and/or the Hymns will not be included. Given how far the contribution of the Homeric poems outweighs the rest of the epic corpus, this is always a concern when analysing epic data (see below, §III.3.2).

²⁷⁵ Hainsworth (1968), 45. Postlethwaite (1979, see especially pp. 2-3) extends the study to the four major Homeric Hymns, including formulaic expressions (i.e. analogical collocations with one element in common with attested formulae) and non-formulaic pairs. His aim is to give a complete account of the behaviour of a syntactic pattern and what portion of its occurrences is formulaic, as well as measure the desire for variation (taken as a sign of literary composition).

²⁷⁶ Hainsworth (1968), 129.

The target sample for this study comprises those among Hainsworth's formulae which occur 10 times or more in the *Iliad* and *Odyssey*.²⁷⁷ We will be comparing these to a baseline sample of formulae that occur less than 10 times in the same two poems, as described in §III.3.3. Formulae containing a possessive adjective were excluded,²⁷⁸ leading to a total of 100 formulae for the target sample. Data were pulled from the TLG,²⁷⁹ using a proximity lemma search with a window of 5 words and no word order specified; results occurring in lines marked as spurious in the TLG text were excluded.

The following variability measures were assessed, and are presented here under Wulff's three categories. For each, an example of modified formula has been provided, together with Hainsworth's basic formula it derives from; note that the modified formulae may also occur elsewhere in the text, but only one example is provided here.

1. Tree-syntactic flexibility:²⁸⁰

- a. Relative clause: either the adjective or the noun (but not the whole expression) occurs in a relative clause. E.g. πάντες ὅσοι πάρος ἦσαν ἄριστοι for πάντες ἄριστοι (*Il.* 11.825).

2. Lexico-syntactic flexibility:

- a. Adjective insertion. E.g. ἀθάνατοι θεοὶ ἄλλοι (*Il.* 3.298 and *passim*) for θεοὶ ἄλλοι.
- b. Noun insertion; this includes coordinated nouns,²⁸¹ but not the addition of another Adj + N pair. So, for θυμὸς ἀγῆνωρ, μένος καὶ θυμὸς ἀγῆνωρ (*Il.*

²⁷⁷ αἶμα κελαινόν is mistakenly reported by Hainsworth as occurring 10 times, but the actual number is lower.

²⁷⁸ Since most possessive adjectives have the same metrical shape, and can be argued to be semantically equivalent, these expressions open up the possibility of lexical substitution.

²⁷⁹ <http://stephanus.tlg.uci.edu/>, last consulted May 2018.

²⁸⁰ Since there are no verbs in this dataset, only one measure of tree-syntactic flexibility is relevant here.

²⁸¹ Cases in which καὶ or τε are present have also been counted under 2.f.

20.174) counts as noun insertion, but μεγάλη τε βίη καὶ ἀγήνορι θυμῷ (Il. 24.42) does not.

- c. Modifying Noun Phrase. E.g. πάντες ἄριστοι Ἀργείων (Od. 4.272) for πάντες ἄριστοι.²⁸²
- d. Post-modifying relative clause. E.g. ἄλκιμον ἔγχος, ὃ οἱ παλάμηφιν ἀρήρει (Il. 3.338) for ἄλκιμον ἔγχος.²⁸³
- e. Preposition.²⁸⁴ E.g. εἰς οὐρανὸν εὐρύν (Il. 5.867) for οὐρανὸν εὐρύν.
- f. Adverb/particle insertion. E.g. πᾶν δ' ἦμαρ (Il. 1.592) for ἦματα πάντα (with separation and inversion).

3. Morphological flexibility:

- a. Negation. This never occurs in the target sample.
- b. Article. E.g. τὰ τεύχεα καλά (Il. 21.317) for τεύχεα καλά.²⁸⁵
- c. Comparative or superlative adjective. E.g. Od. 11.474, μεῖζον ἐνὶ φρεσὶ μήσεαι ἔργον, for μέγα ἔργον (with mobility and separation).
- d. Number (with or without metrical change). E.g. ὀξέα δοῦρα (Il. 5.495 and *passim*, no metrical change) for ὀξεί δουρί; μελαινάων ἐπὶ νηῶν (Il. 5.550 and *passim*, with metrical change) for νηὶ μελαίνῃ.
- e. Case (with or without metrical change). See the examples in 3.d.

Number and case variation have been assessed separately depending on whether the metrical shape of the formula remains unchanged. When both number and case change with metrical variation (a frequent occurrence), I have counted both changes when

²⁸² 'Post-modifying PP' in Wulff; the term was modified to account for the flexibility of Greek word order and the presence of dependent NPs in the genitive, as in the example above. We could easily call this 'PP insertion,' as indeed it falls in the same category as the previous two types.

²⁸³ It is irrelevant whether the content of the relative clause is itself formulaic. In this case, ἄλκιμα δοῦρε, τὰ οἱ παλάμηφιν ἀρήρει appears at Il. 16.139.

²⁸⁴ When there is ambiguity between preposition and preverb, I have generally erred on the side of the preposition.

²⁸⁵ The use of ὁ ἢ τό as an article in Homer is far less developed than in Classical Greek; on this issue, see Wackernagel (1924), xiv-xvi; Chantraine (2015)², §236-45, and the recent discussions by Manolassou and Horrocks (2007) and Guardiano (2013). We will return to this issue in §IV.3.3.

reverting to the original number or case does not create a metrically compatible form.²⁸⁶ A fourth category of variation was also introduced, based on Hainsworth's measures, in order to account for purely metrical variation:

4. Metrical flexibility:

- a. Mobility (to a different position in the line); only cases in which neither element appears in the usual position count.
- b. Inversion of the word order.
- c. Separation, independently of whether the intervening material is syntactically related to the formula. Runover formulae have not been counted when no additional material is introduced (that is, when the only separation is the verse-end). For an example of inversion and separation, see 2.f.

In formulae which show a high frequency of inflection and mobility, one could question which case or position in the line is to be considered 'primary,' or indeed whether one exists at all. For both case and position, I have considered the most frequent one as primary, pending further discussion in §III.3.2.

III.3.2 *The target sample*

Table 3 reports the total frequency of each type of modification across all Adj + N pairs that occur at least 10 times; both absolute numbers in each poem and frequency relative to the total number of occurrences of pairs in the target sample are reported. Note that occurrences of a formula that involve different modifications at the same time are counted

²⁸⁶ E.g. μελαινάων νηῶν (gen. pl.) counts as both number + metre and case + metre, because neither νηὸς μελαίνης (gen. sing.) nor νηυσὶ μελαίνας/-ησι (dat. pl.) have the same metrical shape as νηὶ μελαίνῃ (dat. sing.).

more than once (that is, once under each of the relevant categories); therefore, percentages do not add up to 100%.

Table 3: Distribution of Adj + N formulae occurring at least 10 times in Homer (target sample)

	<i>Il</i>	<i>Od</i>	<i>Theo</i>	<i>Erga</i>	<i>hDem</i>	<i>hAp</i>	<i>hHer</i>	<i>hAph</i>	total	% of target pairs
total pairs	1119	954	51	19	26	39	31	17	2256	
unmodified	365	304	14	3	10	8	8	4	716	31.74 %
<i>tree-syntactic</i>										
relative cl.	8	0	0	0	0	0	0	0	8	0.35 %
<i>lexico-syntactic</i>										
adjective ins.	128	135	2	4	3	6	3	0	281	12.46 %
noun ins.	52	43	5	0	1	4	0	1	106	4.70 %
modifying NP	121	66	10	1	3	1	5	0	207	9.18 %
post-modif. relative	7	8	0	0	0	0	0	2	17	0.75 %
preposition	204	169	9	3	3	11	0	1	400	17.73 %
adverb/particle	96	70	7	0	4	5	1	1	184	8.16 %
<i>morphological</i>										
negation	0	0	0	0	0	0	0	0	0	0.00 %
article	18	8	0	0	0	0	0	0	26	1.15 %
comparative/superlative	1	8	0	0	1	1	0	0	11	0.49 %
number, no metre change	28	44	2	0	2	1	0	0	77	3.41 %
number + metre change	58	70	1	4	1	8	2	1	145	6.43 %
case, no metre change	170	137	8	1	3	7	4	4	334	14.80 %
case + metre change	190	94	7	4	2	12	2	0	311	13.79 %
<i>metrical</i>										
mobility	330	317	18	9	6	19	11	7	717	31.78 %
inversion	189	100	6	2	3	10	7	2	319	14.14 %
separation	201	117	12	5	5	8	4	2	354	15.69 %

The first element to note here is that the frequency of formulaic modification is relatively high, with unmodified formulae making up less than a third of the dataset. When we come to individual types of modifications, we can already see that some appear more common than others. However, as discussed in §III.2.2, these frequencies mean little without comparison to a baseline sample: what matters is how far the pattern of attested variation deviates from the behaviour of other constructions – in our case, less frequent formulae. Moreover, some variation parameters may be affected by peculiarities of the Homeric language. For instance, variation involving the article is extremely rare in archaic Greek epic (see n. 285 above), but can be expected to be common in later texts.²⁸⁷ The same does not necessarily apply to other rare kinds of modification, such as the use of a post-modifying relative clause.

Another element that needs to be considered carefully is the pre-eminence of Homer in this dataset. Over 90% of formulaic pairs occur in the *Iliad* and *Odyssey*; this reflects both the composition of the corpus, where these two poems make up just less than 90% of the word count,²⁸⁸ and of our following Hainsworth's choice to only consider formulaic pairs that occur in the *Iliad* and *Odyssey* for our list of formulae. Therefore, these data should be taken to represent Homeric usage more than anything else.

This summary of results also ends up obscuring the remarkable range of variation in the behaviour of individual formulae. While some formulae are extremely rigid (for instance, ὄξει χαλκῷ appears 36 times, always unmodified), others occur much more frequently in modified form. Sometimes, modifications are consistent: ἄλκιμος υἱός (18x), for instance, is always accompanied by the name of the father in the genitive. This raises the theoretical issue of what is to be considered the 'primary' shape of a formula. Sometimes we should

²⁸⁷ We will come back to this expectation in §IV, when comparing Homeric data to data from a non-Homeric corpus.

²⁸⁸ The *Iliad* and *Odyssey* contain ca. 199.000 words, Hesiod ca. 13.000, and the four major Hymns also ca. 13.000 (rough numbers computed on the Perseus editions).

probably just assume a different underlying shape with an open slot, which Hainsworth's model cannot fully account for (as for ἄλκιμος υἱός, where what we have is probably [∪–∪∪]_{N.Gen} ἄλκιμος υἱός). However, formulae in which no one form is significantly more frequent than others are more problematic: for instance, of the 113 occurrences of νηυσὶ θεῶσιν, only 1 is unmodified. Formulae for ships, in general, consistently show high flexibility: the frequency of unmodified occurrences is 24% for ποντόπορος νηῦς (5/21), 20% for νηὸς εἴσης (4/20), 11% for νηὶ μελαίνῃ (10/88), and none for γλαφυραὶ νέες (out of 63). With 6/12 unmodified occurrences (50%), εὐεργέα νῆα is an exception; this, as we will immediately see, may be due to its lower frequency.

In light of these observations, we have now gained better insight into how to structure a larger case study: we want to be able to track the behaviour of individual formulae, and we want to be able to measure variation without resorting to the controversial concept of a 'baseline formula.' Both of these desiderata will be satisfied in §IV.

III.3.3 The baseline sample and results

The target sample alone, however, can only give us a limited perspective on the behaviour of Adj + N pairs. A preliminary comparison has therefore been set up with a baseline sample, composed of 322 formulae from Hainsworth's B and C categories (that is, excluding unique expressions) which occur less than 10 times in Homer (although some occur over 10 times in the whole corpus). The same restrictions and measures described in §III.3.1 apply.

Table 4: Distribution of Adj + N formulae occurring less than 10 times in Homer (baseline sample)

	<i>Il</i>	<i>Od</i>	<i>Theo</i>	<i>Erga</i>	<i>hDem</i>	<i>hAp</i>	<i>hHer</i>	<i>hAph</i>	total	% of baseline pairs
total pairs	662	592	39	24	12	32	20	17	1398	
unmodified	298	285	18	10	5	17	12	6	651	46.57 %

<i>tree-syntactic</i>										
relative cl.	1	2	0	0	0	0	0	0	3	0.21%
<i>lexico-syntactic</i>										
adjective ins.	42	58	1	1	0	3	0	1	106	7.58%
noun ins.	30	32	1	3	2	0	2	1	71	5.08%
modifying NP	81	53	6	2	1	6	1	0	150	10.73%
post-modif. relative	7	10	0	0	0	0	0	0	17	1.22%
preposition	60	42	3	3	2	3	1	1	115	8.23%
adverb / particle	53	56	2	6	3	2	3	2	127	9.08%
<i>morphological</i>										
negation	0	0	0	0	0	0	0	0	0	0.00%
article	2	2	0	0	0	0	0	0	4	0.29%
comparative / superlative	6	1	0	0	0	0	0	0	7	0.50%
number, no metre change	14	13	0	5	0	0	0	1	33	2.36%
number + metre change	20	7	2	0	0	0	1	0	30	2.15%
case, no metre change	46	33	4	3	0	2	2	1	91	6.51%
case + metre change	50	31	1	2	1	2	1	2	90	6.44%
<i>metrical</i>										
mobility	141	128	13	3	5	7	5	9	311	22.25%
inversion	56	62	4	3	1	0	1	1	128	9.16%
separation	69	60	2	4	3	1	1	1	141	10.09%

There are clearly differences between the behaviour of more frequent formulae (target sample) and that of less frequent ones (baseline sample). However, since we are dealing with statistical samples of different sizes, before drawing any conclusion we need to make sure that what appear to be differences in the distribution of formulae across types of variation are not in fact due to the different sample size. To do this, we perform a statistical significance test. The significance of the difference between the attested frequency of each type of modification in each sample was tested using a two-sided N-1 χ^2 test with a significance threshold of $P < 0.05$.²⁸⁹ Significant differences are marked as S in Table 5,

²⁸⁹ On the χ^2 test and its variants see Campbell (2007), including a reference to the online calculator used here (<http://www.iancampbell.co.uk/twobytwo/calculator.htm>). The test was conducted on

while non-significant differences are NS; P-values are always given, along with the direction of significant differences (B[aseline] > T[arget] or T[arget] > B[aseline]).

Table 5: flexibility in the target and baseline sample compared

	target sample frequency	baseline sample frequency	P values
unmodified	31.74%	46.57%	S (P < 0.0001); B > T
<i>tree-syntactic</i>			
relative cl.	0.35%	0.21%	NS (P = 0.5)
<i>lexico-synt.</i>			
adjective ins.	12.46%	7.58%	S (P < 0.0001); T > B
noun ins.	4.70%	5.08%	NS (P > 0.5)
modifying NP	9.18%	10.73%	NS (P = 0.12)
post-modif. relative	0.75%	1.22%	NS (P = 0.2)
preposition	17.73%	8.23%	S (P < 0.0001); T > B
adverb / particle	8.16%	9.08%	NS (P = 0.3)
<i>morphol.</i>			
negation	0.00%	0.00%	N/A
article	1.15%	0.29%	S (P = 0.005); T > B
comparative / superlative	0.49%	0.50%	NS (P > 0.5)
number, no metre change	3.41%	2.36%	NS (P = 0.07)
number + metre change	6.43%	2.15%	S (P < 0.0001); T > B
case, no metre change	14.80%	6.51%	S (P < 0.0001); T > B
case + metre change	13.79%	6.44%	S (P < 0.0001); T > B
<i>metrical</i>			
mobility	31.78%	22.25%	S (P < 0.0001); T > B
inversion	14.14%	9.16%	S (P < 0.0001); T > B
separation	15.69%	10.09%	S (P < 0.0001); T > B

Overall, the level of flexibility in the target sample is significantly higher than in the baseline (lower frequency of unmodified formulae); for individual parameters, the frequency of modified formulae tends again to be higher in the target sample, while in cases where this does not hold true (noun insertion, modifying NP, post-modifying relative, adverb/particle, comparative/superlative) the differences are far from significant. Since the only salient

2x2 tables with the total number of pairs in each sample as column totals and the total number of modified/unmodified pairs as row totals:

<i>adjective insertion</i>	target sample	baseline sample	totals
modified	281	106	387
unmodified	1975	1292	3267
totals	2256	1398	3654

difference between the two samples is the token frequency of individual formulae (that is, whether they occur more or less than 10 times), these results appear to be a token frequency effect: more frequent formulae are more easily subject to modification. Consider, for instance, the formulae for ships discussed in §III.3.2: it is easy to see how poets, faced with a need to mention ships frequently, would have had to employ related formulae in different metrical and syntactic environments, and would therefore have found it helpful to use more flexible forms. This interpretation is not radically different from Hainsworth's view that "secondary expression[s]" were created as "a response to an emergency in composition."²⁹⁰

Frequent repetition does not, therefore, lead to higher entrenchment of a fixed form, but rather exposes a formula to more occasions for change – or, if we want to view modified formulae as conceptually separate entities from their original form, it gives rise to more 'daughter expressions.' Accordingly, the variation parameters that give rise to significant differences are associated with declension and metrical change – two kinds of modification that will be frequently necessary when adapting to a new environment. Two other types of frequent modification dovetail perfectly with this interpretation: the addition of a preposition frequently co-occurs with change in case, and expansion by way of adjectives can be viewed as a form of less disruptive, almost ornamental modification.

Perhaps the most striking result, on the other hand, has to do with the addition of an article, where the observed difference is significant despite the very low absolute frequency: the higher pervasiveness of this modification in more frequently used formulae might be seen as a sign of historical change in progress. The fact that modification with an article only occurs in the *Iliad* and *Odyssey* may be seen as counterevidence to this claim, but the extreme rarity of the phenomenon makes it unlikely for it to occur in the shorter later texts.

²⁹⁰ Hainsworth (1968), 61.

However, 3 of the 12 phrases that occur with an article²⁹¹ contain the adjective ἄλλος; these alone make up almost half of the occurrences of definite article modification which appear in the *Iliad* (8/18), and the entirety of those which appear in the *Odyssey* (άνέρες ἄλλοι and ἄλλοι ἑταῖροι alone). Co-occurrence with contrastive adjectives like ἄλλος has been identified among the environments promoting the articular use of ὁ ἢ τό in Homer, as it makes the semantic definiteness of the noun particularly salient (due to the need to distinguish between ‘another’ and ‘the other’).²⁹² The presence of this feature may therefore have more to do with this semantic environment than with a frequency effect. We will come back to this insight in §IV.3.3, as data gathered through a larger-scale analysis also suggest that indefinite adjectives such as ἄλλος play a significant role in formulaic variation.

III.3.4 Further steps

While the results above give some interesting information on the influence of token frequency on formulaic flexibility, the conclusions of this preliminary experiment are limited by the choice of an arbitrary cut-off between the two samples;²⁹³ moreover, differences in behaviour between individual formulae have thus far only been discussed qualitatively, rather than assessed through quantitative analysis of the behaviour of individual formulae. Conceptually, we want to not just observe how frequent a certain type of variation is across a sample of formulae, but to be able to place every single formula in the sample on a scale of more-to-less flexible according to a specific parameter.

²⁹¹ ἄνδρας ἀρίστους, άνέρες ἄλλοι, δόρυ μακρόν, ἄλλοι ἑταῖροι, θεοὶ ἄλλοι, πᾶσι θεοῖσι, νηὶ μελαίνῃ, τεύχεα καλά in the target sample; ἔγχρῃ μακρῷ, πίονα μῆλα, αὐτὴν ὁδόν, παιδὸς ἀγαυοῦ in the baseline sample.

²⁹² Cf. Ambrosini (1991), 4; Manolossou and Horrocks (2007), 227-8.

²⁹³ Compare the case of εὐεργέα νῆα discussed in §III.3.2 for an example of how this may affect the results: while a phrase that occurs 12 times falls automatically into the target sample, it is reasonable for it to behave in a way that is closer to the baseline sample, as its frequency is much lower than other target formulae.

To do this, we also want to compare formulaic data to the average flexibility of similar constructions in a baseline corpus, rather than a baseline sample of other formulaic material. By doing this, we will be able to compare formulaic to non-formulaic material, and therefore to assess which types of variation are more characteristic of formulae vs non-formulae. Measuring flexibility against a baseline from the main corpus itself is problematic: since the language of early Greek epic as a whole is, if not formulaic, at least overwhelmingly made of formulae, this is akin to measuring baseline flexibility in a corpus composed entirely of idioms, leading to skewed results. The baseline corpus should therefore be a sample of non-epic texts, representing a synchronic situation as close as possible to the composition of archaic Greek epic. Due to the limitations of the attested material, this translates to a text or corpus of Classical texts, either from a monolithic dialectal tradition (for instance, an Attic corpus from the 5th or 4th century BC) or from one that shows closer similarities to the Homeric dialect (Herodotus is a potential candidate²⁹⁴). This choice should allow us to retrieve reliable results from a sufficiently large synchronic corpus,²⁹⁵ even if all the material will be chronologically later than the main sample. Both baseline and target corpus will need to be processed to allow for effective information retrieval; since we aim to gather data on syntactic flexibility, this will mainly involve Part-of-Speech tagging of all texts.²⁹⁶

Having outlined these desiderata (a target sample of formulae, a baseline sample, and processed corpora to perform the analysis), we can now move on to the larger-scale study of Adj + N formulae that will take up the whole of the next chapter. We will see that these

²⁹⁴ I owe this suggestion to Lucien van Beek.

²⁹⁵ Where ‘synchronic’ must unfortunately be understood to mean ‘spanning about a century.’

²⁹⁶ PoS tagging can be performed to a satisfying level of accuracy with the Classical Language ToolKit (Johnson and et al. 2014); extensive documentation on the toolkit’s functions is available at <http://docs.cltk.org/>. The texts forming the corpus are available in an open format, as opposed to the restricted, non-downloadable format of the TLG, from the Perseus Project digital libraries (<https://github.com/PerseusDL>). We will see in §IV.1 that the availability of processed corpora is essential to this work.

desiderata can be satisfied, and that a more detailed and fine-grained description of formulaic behaviour can lead us to interesting insights on what characterises formulaic usage in oral poetry. We will use these insights to set up a comparison with non-oral epic, specifically Apollonius of Rhodes, which will be discussed in §V.

IV The syntactic flexibility of Adjective + Noun formulae

This chapter is dedicated to an analysis of Adj + N phrases in early Greek epic, compared to the behaviour of the same type of phrases in a non-epic corpus from the 5th century BC. This comparison will allow us to identify some types of variation that are characteristic of formulaic phrases (or rather, the avoidance of which is characteristic of formulaic phrases), as opposed to common Adj + N pairs. To achieve this result, we will need a corpus on which to perform the analysis, for both the epic texts and our baseline comparison; this corpus needs to contain information about the part of speech of each word (whether it is an adjective, a noun, or something else that will not be part of our analysis), as well as some amount of morphological information. The corpus and data extraction will be described in §IV.1. We will also need a statistical method that allows us to compare the behaviour of Adj + N pairs between the target epic corpus and the baseline non-epic corpus, ideally without resorting to the assumption that each formula has a primary, unmodified shape (an assumption that can hardly hold for non-formulaic Adj + N pairs); this issue and its solution will be discussed in §IV.2.

At the end of this chapter, we will not just have a much more complete picture of how Adj + N pairs behave in early epic, and whether and how their behaviour differs from that of non-formulaic pairs; we will also be able to suggest some explanations for this behaviour, both in light of what we know of the mechanisms of oral composition and in light of the linguistic characteristics of formulae. Among other things, this will provide an initial response to the kind of criticism of the notion of flexible formula that was discussed in §I.3.4 and §II.2.2: we will show that even flexible formulae do show characteristic behaviours that set them apart from non-formulaic material, and that this behaviour is explainable as more than potentially random variation. The following chapter on Apollonius of Rhodes, on the other hand, will address the objection that this kind of Adj + N

combinations are a characteristic of hexameter poetry in general, rather than something that belongs more specifically to oral epic poetry.

IV.1 Materials and methods

IV.1.1 *The Diorisis digital corpus*

All stages of the analysis described in this chapter were performed on texts from the Diorisis Ancient Greek corpus.²⁹⁷ This corpus comprises 820 literary Greek texts, spanning chronologically from Homer to the 5th century AD, for a total of over 10 million words.²⁹⁸ The texts are mainly sourced from the Perseus Canonical Greek Literature repository (along with the Little Sailing and Bibliotheca Augustana digital libraries).

The Diorisis corpus files are in TEI-compliant XML format, where TEI stands for Text Encoding Initiative (<http://www.tei-c.org/>), a consortium that sets international guidelines for the digital representation of texts. The metadata (associated data that are not directly part of the text) for each file encode, among other relevant information, the text's date of composition (or the best approximation known) and genre, as well as the modern edition it comes from. The header of the *Iliad* file, for instance, contains details about its date ('-800,' that is 800 BC, obviously approximate), its genre ('Poetry'), subgenre ('Epic'), and the fact that the Diorisis text is based on the 1920 OCT edition by Monro and Allen (which is the one available in the Perseus files). This form of encoding, apart from maintaining the highest

²⁹⁷ Version 1 (Vatri and McGillivray 2018a), available online under a Creative Commons Attribution 4.0 International licence (CC BY 4.0). All information pertaining to the Diorisis corpus in this section comes from Vatri and McGillivray (2018b), to which the reader may refer for further technical details.

²⁹⁸ Note the size difference compared to previously-existing annotated corpora, such as the Perseus Ancient Greek Dependency Treebank (http://perseusdl.github.io/treebank_data; Bamman and Crane 2011) and the PROIEL treebank (<https://proiel.github.io>; Haug and Jøhndal 2008), both of which contain only a few hundred thousand tokens. The Diorisis corpus, however, does not contain information about syntactic relations. For more on Greek treebanks, see §VII.1.

level of transparency in tracing the source text, allows users to break down the corpus according to not just chronology, but genre and subgenre as well.

The text in each file is fully annotated,²⁹⁹ with information about lemma, part of speech, and morphological analysis appended to each word token. The position of each word in the text is also encoded in a unique way: each sentence, delimited by strong punctuation marks,³⁰⁰ is given a unique number, and each word is tagged with its position within the numbered sentence. Other textual elements are also included in the tagging system, such as location in terms of line or paragraph in the source edition, as well as whether the word is part of a quotation or is an editorial supplement. This makes it very easy to trace each token back to its original location.

The format of the morphological analysis is such as to be easily readable by a human, as well as by a machine. Each word form is associated to its lemma, that is, the corresponding dictionary entry, with its unique numeric ID; it also receives a Part of Speech tag, conventionally abbreviated as 'PoS tag.' Tagged Parts of Speech include noun, verb, adjective, preposition, etc., according to the function each word performs in the text. It is important to note that lemmatisation and PoS tagging were conducted separately, using two different tools (respectively, a custom-built dictionary and TreeTagger, which is a Natural Language Processing tool trained on Perseus Treebank and PROIEL materials): therefore, it is not the case that each dictionary entry is associated with a pre-defined word form (as it would be, for instance, if the entry for, as an example, *μῆνις* were automatically associated with the PoS-tag 'noun'). We will see below how this is used to double-check the lemmatisation for ambiguous cases, but first, here is an annotated example of how the

²⁹⁹ For a detailed discussion of how electronic corpus annotation works on different levels, the rationale behind it, and the challenges it brings, see Jensen and McGillivray (2017), 98-129.

³⁰⁰ I.e. full stop (.), middle dot (·), question mark (;). Punctuation was normalised before processing, to ensure that e.g. middle dots are encoded correctly.

first two words of the *Iliad* are tagged in Diorisis (note that word forms are in BetaCode to facilitate processing, but lemmas are in Unicode).³⁰¹

<code><sentence id="1" location="1.1"></code>	Start of new sentence tag, with unique ID number and text location
<code><word form="mh=nin" id="1"></code>	Start of new word tag: word form as it appears, with numeric ID (numbering restarts with each sentence)
<code><lemma id="67485"</code>	Start of lemma tag, with unique ID number for the corresponding dictionary entry
<code>entry="μήνις"</code>	The actual dictionary entry, i.e. the lemma
<code>POS="noun"</code>	Part of Speech tag as determined by
<code>TreeTagger="false"</code>	TreeTagger, with associated TT
<code>disambiguated="n/a"></code>	annotations ³⁰²
<code><analysis morph="fem acc sg"/></code>	Morphological analysis: feminine accusative singular
<code></lemma></code>	End of lemma tag
<code></word></code>	End of first word tag
<code><word form="a)/eide" id="2"></code>	Start of new word tag, with word form and numeric ID
<code><lemma id="1587"</code>	[See above for explanation]

³⁰¹ Spaces, line breaks, and other considerations of layout on a page are irrelevant in XML; the structure of a document is exclusively determined by the opening and closing of tags, each of which takes the format `<tag>content</tag>`. A tag can contain any number of other subtags (e.g. `<tag>content <subtag> subcontent </subtag> content </tag>`), but tag boundaries cannot overlap (it is impossible to have a structure of the type `<tag> content <subtag> subcontent </tag> more subcontent </subtag>`): the structure of an XML document is strictly hierarchical.

³⁰² These are irrelevant for our purposes, but more on them can be found in Vatri and McGillivray (2018b).

entry="ᾠείδω"

POS="verb"

TreeTagger="false"

disambiguated="n/a">

[The morphological tagger offers all possible

<analysis morph="imperf ind act 3rd sg analyses here, including inaccurate and
(epic doric ionic aeolic)"/> redundant ones; as we will see, this does not

<analysis morph="pres imperat act 2nd sg affect our use of the Diorisis data]
(epic ionic)"/>

<analysis morph="imperf ind act 3rd sg
(epic doric ionic aeolic)"/>

</lemma>

</word>

End of second word tag

...

[Each word in the rest of the sentence is
similarly tagged]

</sentence>

End of sentence tag

<sentence id="2" location="1.8">

Start of new sentence tag, with unique ID
number and text location

And so on.

In cases when a word form can be analysed as deriving from different lemmas, the output of the TreeTagger tool is used to solve the ambiguity wherever possible. To take an example from Vatri and McGillivray (2018b), the form πράξεις can represent either the nominative/accusative plural of the noun πράξις or the second person singular future indicative active of the verb πράσσω; in the Diorisis corpus, the token πράξεις has been associated to the lemma for either πράξις or πράσσω depending on whether the TreeTagger tool tagged it as a noun or verb. This is one of the chief advantages of

performing the PoS tagging separately from the lemmatisation. Sometimes, however, different lemmas belonging to the same PoS class are available; the most frequent one is then automatically chosen. This system guarantees some improvement in accuracy over lemmatisers that simply choose the most frequent lemma, regardless of part of speech (the overall accuracy of the TreeTagger model is around 91% according to the authors). On the other hand, when a word form admits several possible morphological analyses, all falling under the same lemma, these are always reported in full as possible alternatives; this can be observed in the tags for ἄειδε in the example above.

IV.1.2 The sub-corpora

Two different samples from the Diorisis corpus were extracted for the analysis conducted in this chapter.³⁰³ Since the metadata for each text contain information about its date of composition and genre, the selection has been made exclusively on the basis of this information, so as to guarantee full reproducibility and lack of bias: that is, the 5th-century baseline sub-corpus has been defined as containing all texts that are tagged with a date between 499 and 400 BC.³⁰⁴

The sub-corpora used in this chapter are:

- the target sub-corpus, or archaic epic sub-corpus, composed of the *Iliad*, *Odyssey*, *Theogony*, *Works and Days*, and the *Hymns to Demeter*, *Apollo*, *Hermes*, *Aphrodite*;

³⁰³ I will be using the term ‘sub-corpora’ for these, to avoid confusion with the main Diorisis corpus. For a discussion of how best to define a historical corpus, see Jensen and McGillivray (2017), 41-2. The sub-corpora used in my study are corpora in the sense of being (small) collections of machine-readable texts representing authentic ancient language usage, but their representativity and balance are limited by the way they are defined: for instance, the 5th-century baseline sub-corpus is somewhat representative of 5th-century literary language usage, but not of any other century or of e.g. epigraphic language, and is very strongly skewed towards Attic texts.

³⁰⁴ Checking the wordcount against the breakdown of the Diorisis corpus by century and genre provided in Vatri and McGillivray (2018b), Table 3, reveals that these are the dates used to define the 5th century BC, rather than 500-401.

- the 5th-century baseline sub-corpus, composed of all texts dated between 499 and 400 BC,³⁰⁵ that is, Aeschylus (*Supp.*, *Pers.*, *PV*, *Sept.*, *Ag.*, *Cho.*, *Eum.*), Andocides (*On the Mysteries*, *On His Return*, *Against Alcibiades*), Antiphon (*Prosecution of the Stepmother for Poisoning*, *Tetralogies 1-3*, *On the Murder of Herodes*, *On the Choreutes*), Aristophanes (*Ach.*, *Eq.*, *Nub.*, *Vesp.*, *Pa.*, *Av.*, *Lys.*, *Thesm.*, *Ran.*), Euripides (*Hec.*, *Supp.*, *HF*, *Ion*, *Tro.*, *El.*, *IT*, *Hel.*, *Phoen.*, *Or.*, *Bacch.*, *IA*, *Rhes.*³⁰⁶), Herodotus, Isocrates (*Against Euthynus*, *Against Callimachus*), Lysias (*On the Murder of Eratosthenes*, *On a Wound by Premeditation*, *For Callias*, *Against Andocides*, *Accusation of Calumny*, *Against Eratosthenes*, *For Polystratus*, *Defense against a Charge of Taking Bribes*, *Against Pancleon*, *On the Refusal of a Pension*, *Defense against a Charge of Subverting the Democracy*, *Against the Subversion of the Ancestral Constitution*), Pindar (*Ol.*, *Pyth.*, *Nem.*, *Isthm.*), Sophocles (*Trach.*, *Ant.*, *Aj.*, *OT*, *El.*, *Phil.*, *OC*, *Ichneutae*), Thucydides.

IV.1.3 Extracting the target forms

Adjective + Noun pairs were extracted from each sub-corpus through a Python script written especially for this purpose by Dr Alessandro Achille.³⁰⁷ The script takes as input a list of .xml files defined by the user; it reads the PoS tag and morphological analysis information encoded in the Diorisis corpus and retrieves all pairs composed of a word tagged 'noun' followed or preceded by a word tagged 'adjective' in the same case, gender

³⁰⁵ When only some of an author's works are listed, it is because some of them are dated to the 4th century (see e.g. Isocrates), with the exception of Euripides, half of whose production is simply missing from the Diorisis corpus.

³⁰⁶ The *Prometheus Bound* and *Rhesus* are tagged in Diorisis as composed by Aeschylus and Euripides, respectively, and are both assigned 5th century dates; they have therefore been included in the corpus, in keeping with my initial decision not to alter or challenge the information contained in the metadata. In fact, while *PV* is generally accepted to be a 5th-century tragedy, *Rhesus* is generally ascribed to the 4th century. On questions of authenticity and dating, cf. e.g. (on *PV*) Ruffell (2012), 10-16; (on *Rhesus*) Liapis (2012), lxxvii-lxxv, and Fries (2014), 22-42.

³⁰⁷ The script is available at <https://github.com/MartinaAstridRodda/dphil-thesis/blob/main/adjnoun.py>.

and number. The distance (window) between the two target words is customisable by the user; in this chapter, a window of 6 words (including the adjective and noun)³⁰⁸ was chosen as being sufficiently inclusive to capture most occurrences without leading to too many spurious hits. Sentence breaks are taken into account: only pairs in which both words occur in the same sentence are counted. The user can also filter the results by frequency, by setting a minimum number of occurrences for each pair in the sub-corpus; as we will see below, this chapter will only consider pairs occurring at least 20 times.

The output of the script is a list of Adj + N pairs with their respective frequencies and all corresponding occurrences; the latter are printed out along with their context, that is, all words in the window between the adjective and noun (or noun and adjective) and a number of words preceding and following the pair, which can again be defined by the user. Each occurrence is also associated with the corresponding sentence number, in order to make situations in which more than one pair occurs in the same sentence more immediately recognisable (these situations are often due to a false positive hit, see §IV.1.4 below).

Below is an example from a list of Adj + N pairs in the archaic epic sub-corpus, with a window of 6 words and a context of up to 3 extra words on each side. (The output has been converted from BetaCode to Unicode for easier readability.) Readers may well recognise the first occurrence as part of *Il.* 1.12-13; this also illustrates how line breaks that do not correspond to a sentence break are not taken into account. The low dashes highlight the target Adj + N in the output.

Pair ναῦς θοός: total 104³⁰⁹

Sentence 5: ὁ γὰρ ἦλθε _θοὰς_ ἐπὶ _νῆας_ Ἀχαιῶν λυσόμενός τε

Sentence 170: ἃ μοί ἐστι _θοῆ_ παρὰ _νῆϊ_ μελαίνῃ τῶν οὐκ

³⁰⁸ This means the widest possible bracket is Adj W W W W Noun (or Noun W W W W Adj).

³⁰⁹ See §IV.1.4 for an explanation of the difference compared to the count provided in §III.3.2 (where Hainsworth detected 113 occurrences of νηυσὶ θοῇσι).

Sentence 175: Ἀτρεΐδης δ' ἄρα _νῆα_ _θοήν_ ἄλαδε προέρυσσεν

All pairs retrieved by this Python script were examined by hand and categorised by type of syntactic modification, as detailed in §IV.1.5 below.

IV.1.4 Known issues with the corpus and data extraction

Just like any other work in corpus linguistics, the results of this analysis are constrained by the characteristics and the quality of the underlying corpus. This does not just mean that the structure of the corpus can limit the scope of the analysis, but also that it is by definition impossible to extract more information than is contained in the corpus itself – which also makes it very difficult to recover information that was lost during corpus processing.

When it comes to the data retrieval process described in this chapter, there are two main stages at which errors can occur. One is at the level of the corpus itself, where inaccuracies in the lemmatisation/PoS-tagging process can lead to loss of usable information. Some of these instances became visible during data analysis, and are discussed in this section. The second stage in which loss of information is possible is during the retrieval process through the Python script described in §IV.1.3. In this case, manual examination of the results largely allowed errors to be corrected.

Errors in the lemmatisation (or, more frequently, the PoS tagging) lead to both false negatives (actual Adj + N pairs remaining undetected) and false positives (two words that do not form an Adj + N pair being retrieved as one). If we take the archaic epic sub-corpus as a starting point, both types of errors become apparent when we take the list of Adj + N pairs retrieved by the Python script and compare it to a list of Adj + N pairs from the experiment performed in §III.3, which was based on Hainsworth's manually-gathered set of formulae.

For both sets, we introduced a frequency threshold: only pairs that occur at least 20 times are included. Pairs that meet the frequency threshold but involve a possessive adjective (such as πατήρ + ἐμός) have been excluded, as these adjectives have metrical equivalents that are also semantically compatible – forms of πατήρ + ἐμός are metrically exchangeable with, for example, forms of πατήρ + ἐός, depending on the intended referent of the possessive, but the script would see these as two entirely separate pairs.

The result of this comparison can be found in Table 6. Pairs are listed as two nominatives, in whichever order they first appear in the corpus, whether this is N + Adj (ναῦς + θεός) or Adj + N (πατρίς + γαῖα).

Table 6: Comparison between manual and automatic retrieval of Adj + N pairs in the archaic epic sub-corpus

Pair	Frequency in Hainsworth's list	Frequency retrieved by script
ναῦς + θεός	113	104
ναῦς + μέλας	88	75
πατρίς + γαῖα	86	88
θεός + ἀθάνατος		81
φίλος + υἱός	74	78
γλαφυρός + ναῦς	63	51
θεός + ἄλλος	55	51
φίλος + ἦτορ	54	55
ἦμαρ + πᾶς	53	51
πᾶς + θεός	46	49
οὐρανός + εὐρύς	43	40
φίλος + πατήρ		42
ἄνθρωπος + θνητός		39
ὠκύς + ἵππος	37	37
πότνια + μήτηρ	37	πότνια is PoS-tagged as noun
ναῦς + κόϊλος		36
φίλος + ἑταῖρος		40
κρατερός + ὑσμίνη		36
ὀξύς + χαλκός	36	35
μῶνυξ + ἵππος	34	μῶνυξ is PoS-tagged as noun
χάλκεος + ἔγχος	32	24
ἀγλαός + υἱός	32	32
φίλος + γαῖα	proper subset of πατρίς + γαῖα	32
κακός + πολύς	30	κακός is PoS-tagged as adjective
ὄδε + ἔργον	29	ὄδε is PoS-tagged as pronoun

θυμός + ἀγήνωρ	28	ἀγήνωρ is PoS-tagged as noun
δολιχόσκιος + ἔγχος		26
δόρυ + φαεινός	26	25
πᾶς + ἑταῖρος	26	29
θοῦρις + ἀλκή	25	θοῦρις is PoS-tagged as noun
αἶθοψ + οἶνον	25	οἶνον is parsed as neuter
μέγας + ἔργον	25	μέγας is PoS-tagged as adverb
ἀγλαός + δῶρον	24	23
αἰπύς + ὄλεθρος	24	25
μέγας + κύμα	23	μέγας is PoS-tagged as adverb
ὄδε + πᾶς	23	ὄδε is PoS-tagged as pronoun
άνήρ + ἄλλος	23	25
ἄνθρωπος + πᾶς		27
ὀξύς + δόρυ	22	21
μέγας + ὄρκος	22	μέγας is PoS-tagged as adverb
ἐρίηρος + ἑταῖρος		21
φαίδιμος + υἱός		21
τεῦχος + καλός	21	καλός is PoS-tagged as noun
ἄλλος + ἑταῖρος	21	22
ποντόπορος + ναῦς	21	some forms of ναῦς are lemmatised as (masculine) ναός ³¹⁰
λευκώλενος + θεά		20
πᾶς + μνηστήρ		20
ναῦς + ἴσος	20	some forms of ναῦς are lemmatised as (masculine) ναός
νηλής + χαλκός	20	νηλής is PoS-tagged as adverb
κήρ + μέλας	20	
οἶνοψ + πόντος	20	οἶνοψ is PoS-tagged as noun
total	1401	1349 (excluding φίλος + γαῖα)

Table 6 also contains a short explanation of the reason why each pair is missing from one of the two samples, wherever an explanation can be provided. For pairs that are missing from Hainsworth's list, the main explanation is metrical in nature: Hainsworth's sample only included formulae of metrical shape $-\cup\cup-x$ and $\cup\cup-x$. Note also the case of φίλος + γαῖα,

³¹⁰ This is likely to have impacted the frequency of all formulae involving ships, which do in fact show a lower frequency overall in the automatically retrieved list. Otherwise, the frequencies observed using the two methods are broadly comparable, after allowing for some reasonable margin of error in manual counts.

which must be excluded because all pairs containing the adjective φίλος also contain the adjective πατρίς, and will therefore be analysed under πατρίς + γαῖα.³¹¹

When it comes to data retrieval, on the other hand, the main source of error is created by individual false positives caused by the search parameters. This is best clarified through an example. In *Il.* 1.17-18, the sequence Ἀτρεΐδαι τε καὶ ἄλλοι ἐϋκνήμιδες Ἀχαιοί, ὑμῖν μὲν θεοὶ δοῖεν, etc., is mistakenly retrieved by the Python script as the first instance of the pair ἄλλος + θεός, as the adjective and the noun do occur within a window of six words of each other and their morphological analyses overlap. Technically, ἄλλοι is in the vocative case and θεοί in the nominative, but the Diorisis corpus provides all possible morphological analyses independently of context, so each of the two words is tagged as both nominative and vocative. This type of error is also ultimately a consequence of the characteristics of the corpus: the annotation does not encode any information about the syntactic relationship between words, so there is no way for an automatic processing tool to interpret ἄλλοι and θεοί correctly, as belonging to different NPs. False positives of this kind are easy to identify and were eliminated when each occurrence was manually checked to record syntactic variation data.

The only case in which this had consequences for the inclusion of an Adjective + Noun pair in the target sample is for the pair πᾶς + μνηστήρ, which was originally retrieved as having a frequency of 20 in the archaic epic sub-corpus, but was later excluded after 3 false positives were spotted upon manual examination of the occurrences, lowering the total to below the threshold. On a more general level, the issue of false positives means that the number of automatically-retrieved occurrences reported in Table 6 is not the actual

³¹¹ The latter takes precedence for two reasons: first, it is much more frequent; second, while all occurrences of φίλος + γαῖα also involve the adjective πατρίς, not all occurrences of πατρίς + γαῖα also involve φίλος.

number of N + Adj pairs observed in the archaic epic sub-corpus: manual screening revealed 26 false hits, bringing the actual number of occurrences down to 1306.³¹²

The percentage of false positives in the epic sub-corpus is therefore around 2%, making the precision of this measure (the proportion of retrieved pairs that are confirmed to be true positives) approximately 98%. Recall (the number of correctly retrieved pairs compared to the number of pairs actually present in the corpus) is on the other hand not easy to assess, because we do not *a priori* know the total number of pairs that are actually in the corpus. The only list available for comparison is the Hainsworth list, which in turn does not have a 100% recall rate: several pairs that were retrieved by the automated processing tool are not in Hainsworth, and among those that are in his list, the automated processing tool occasionally retrieved occurrences that were missed by human screening. The automated search, however, does show a decrease in recall when it comes to the number of types, that is, unique formulae, rather than tokens, that were detected: 34 Noun + Adjective pairs compared to the 38 that occur at least 20 times in Hainsworth's list.

The upshot of the discussion in this section should be that while errors in measurement that lead to false positives are easily noticeable and have been corrected *post hoc*, errors that bring about false negatives are impossible to recognise without running some kind of comparison with an existing dataset; even then, they are very difficult to correct systematically.³¹³ In other words, it is possible to guarantee the *quality* of the information contained in the data, by controlling the number of spurious hits; it is much more challenging to control the loss of *quantity* in the information available in the corpus. In any case, the question of whether or not we should use an automated processing tool is a moot point: automated processing becomes essential when approaching a sub-corpus that is not

³¹² 1349-26-17, since we have eliminated a further 17 occurrences of πᾶς + μνηστήρ.

³¹³ For this reason, and for greater comparability to the baseline sub-corpus for which no list exists, false negatives have not been corrected at all – no pairs that are contained in Hainsworth's list but not in the one retrieved by the script have been added, and so on.

the archaic epic one, where no list of Adjective + Noun pairs exists and the amount of material available is so large as to rule out manual screening.

IV.1.5 Recording syntactic variation

Occurrences of each Adjective + Noun pair were categorised according to a list of measures of syntactic and morphological variation that corresponds, by and large, to the one used in the previous chapter (see §III.3.1). Some of the measures, however, turned out to be unsuitable to the current analysis, and had to be removed or adapted to it. There are two main reasons for this. Firstly, metrical measures are clearly of interest when discussing the behaviour of phrases in an epic corpus, but not in a sample of prose texts, or even in a poetic corpus in a range of different metres. For this reason, measures having to do with position in the line had to be removed, as was the distinction between variations in case and number with or without metrical change.

Secondly, and most importantly, some measures were adapted in order to eliminate one of the main assumptions that were highlighted as problematic by the initial experiment conducted in §III.3, that is, the need to decide on a 'base form' for each formula at the data-collection stage, which often led to arbitrary choices. For the analysis in this chapter, the behaviour of each Adjective + Noun pair was first mapped in full, with no assumptions as to what an 'unmodified' form should look like.³¹⁴ This mainly affects the measures having to do with number, case, and word order, which were designed in chapter §III.3.1 so as to highlight deviation from a base form. Instead of measuring 'inversion,' at the data collection

³¹⁴ This is particularly important when moving from the target sub-corpus to a broader baseline, where it would be even more theoretically problematic to decide on a 'base' form: should this be the nominative singular? What is the 'basic' word order for an Adjective + Noun pair? On word order in non-epic texts, cf. Dik (1995; 2007); Devine and Stephens (1999); S. J. Bakker (2009); Goldstein (2015). The general consensus is that word order in ancient Greek appears to be pragmatically determined – whether this means that there is no underlying word order or whether pragmatics determines a transformation in which an underlying word order is modified. Both of these views would create issues for our study, as we are not accounting for pragmatics.

stage for the current experiment word order was assessed as Adj + N or N + Adj; similarly, the case and number for each occurrence were recorded separately.

Despite this, some assumptions remain built into the data collection phase, so to speak: the Adj + N/N + Adj phrase is still assumed to be the standard form for each pair, with any deviation from this being categorised as an addition (for instance, the creation of a prepositional phrase of the type P + Adj + N counts as 'added preposition'); similarly, it is assumed that noun and adjective should by default be adjacent, since separation is retained as a measure of variation. While these assumptions need to be kept in mind, the statistical analysis applied in §IV.2 is designed to compensate for them, so they will not ultimately affect the results.

The following is a list of the measures of syntactic variation used in this chapter, with details about the ones that were changed along the lines described above. Examples and details for all other measures can be found in §III.3.1.

1. Tree-syntactic flexibility:
 - a. Relative clause.
2. Lexico-syntactic flexibility:
 - a. Adjective insertion;
 - b. Noun insertion (as before, this applies to one or more coordinated nouns, but not the juxtaposition of another adjective + noun pair);
 - c. Modifying NP/PP;
 - d. Post-modifying relative clause;
 - e. Preposition;
 - f. Adverb insertion;
 - g. Particle insertion.

Note that the last two measures were grouped together as ‘adverb/particle insertion’ in the previous chapter, but have been split to highlight the difference in behaviour between adverbs and particles.³¹⁵

3. Morphological flexibility:

- a. Negation;
- b. Article;
- c. Comparative or superlative adjective;
- d. Number: each occurrence was marked as either ‘singular’ or ‘plural/dual’;³¹⁶
- e. Case: each occurrence was marked as either ‘nominative/vocative’,³¹⁷ ‘genitive,’ ‘dative,’ ‘accusative.’³¹⁸

4. Metrical flexibility:

- a. Word order: each occurrence was marked as either Adj + N or N + Adj;
- b. Separation.

Mobility (position within the line) is not considered, for the reasons detailed above. All formulae in the sample are tagged with all the modifications that are present in each instance, and classed by number, case, and word order. Table 7 offers an example of how the phrase *θοῶς ἐπὶ νῆας Ἀχαιῶν*, a subtype of *ναῦς + θοός* which occurs 10 times, is tagged;³¹⁹ the complete datasets are available as an Excel file at https://github.com/MartinaAstridRodda/dphil-thesis/blob/main/syntax_dataset.xlsx.

³¹⁵ On the linguistic categorisation of particles as discourse markers, see Bonifazi, Drummen, and de Kreij (2016), §1.3, as well as further references in Battezzato and Rodda (2018), 20-22. Adverbs and particles have different PoS tags in Diorisis, respectively ‘adverb’ and ‘particle.’

³¹⁶ The dual is not counted separately due to its unstable status in non-Attic dialects.

³¹⁷ Due to morphological overlap outside of second-declension singular forms.

³¹⁸ Neuter nominative, vocative, and accusative are grouped together under ‘nominative/vocative.’ Note the difference with respect to the nominative/vocative problem mentioned in n. 317: neuter non-oblique cases are the same independently of number, so they can be grouped together for all occurrences of neuter adjective + noun pairs, while a separate measure for the vocative would only apply to a small subset of words, and only for some of their forms.

³¹⁹ The number 10 in the table records the number of times the *ναῦς + θοός* formula occurs in this specific shape.

Table 7: Example of the tagging system for *θοάς ἐπὶ νῆας Ἀχαιῶν*

	Tree-syntac tic	Lexico-syntactic										
Form	Relati ve clause	Adject ive inserti on	Noun insertion	Modifi ng NP/PP	Post-modifying relative clause				Preposit ion	Adverb inserti on	Particle insertio n	
θοάς ἐπὶ νῆας Ἀχαι ῶν				10					10			
	Morphological							Metrical				
	Negati on	Article	Compara tive	Number		Case			Word order		Separat ion	
				S	P/ D	N/ V	G	D	A	Adj N		N Adj
					10				10	10		10

All of the different variants of the *ναῦς* + *θοός* pair are tagged in the same way, and numbers are then combined in an overview table that charts how many times each type of variation occurs for each formula. These individual tables will be in turn combined and reported in Table 9.

IV.1.6 The baseline sample

An identical tagging system is applied to the baseline sub-corpus. The Adj + N pairs in this sub-corpus are extracted on the same principles as described in §IV.1.3, including the same window size and frequency threshold; the resulting list of 33 pairs, with their respective frequencies, is provided in Table 8. Once again, the automatic processing results in a small number of false positive hits, which have also been recorded in the table.

Table 8: List of Adj + N pairs in the baseline sub-corpus

Pair	Retrieved total	False hits	Actual total
πολύς + χρόνος	101	2	99
πᾶς + ἄνθρωπος	58	2	56
ἀγαθός + ἀνὴρ	57	1	56
ἀνὴρ + πᾶς	45	3	42
ἀνὴρ + φίλος	43	2	41
στρατός + ναυτικός	41		41
πᾶς + λόγος	39	4	35
δεξιός + κέρας	36		36

άνήρ + ἄριστος	36	1	35
τοιόσδε + τρόπος	36		36
μακρός + χρόνος	32		32
πᾶς + χρόνος	32	4	28
τοιούτος + τρόπος	32		32
στρατός + πεζός	31		31
πᾶς + πρᾶγμα	31	1	30
ἄλλος + ἄνθρωπος	30	1	29
ἄλλος + στρατία	30	1	29
στρατός + πᾶς	29	3	26
χρόνος + ὀλίγος	29		29
ὀπλίτης + χίλιοι	29	2	27
πᾶς + τρόπος	28		28
ἡμέρα + εἷς	27	1	26
άνήρ + ἄλλος	27	1	26
ἄλλος + λόγος	25	1	24
πολύς + λόγος	25	2	23
πᾶς + γῆ	25	2	23
ὅσος + χρόνος	24		24
άνήρ + σοφός	24	1	23
άνήρ + δόκιμος	24		24
πᾶς + στρατία	23		23
πολύς + στρατία	22	2	20
πᾶς + ναῦς	22		22
άνήρ + κακός	22	1	21

IV.1.7 The data-gathering setup: a summary

Sections IV.1.3-IV.1.6 have described a data-gathering setup that was informed by, and therefore satisfies, the desiderata laid out at the end of the previous chapter. The loss of some variation measures and the changes in others are due to differences in the type of analysis to be performed, as well as the materials involved, compared to the preliminary analysis in §III.3. This analysis was instrumental in highlighting which measures of variation needed to be adapted, as well as the implications of some of Hainsworth's assumptions, most notably the idea of a 'base form;' we will make some final steps in removing this assumption in the data analysis section (§IV.2).

The use of a publicly available corpus, which fits the purposes of this study well but was not set up with this analysis in mind, guarantees both maximum reproducibility of the results and a lack of study-dependent bias in the data.³²⁰ The use of an automated data collection tool captures pairs that were not in Hainsworth's list (see Table 6), and allows for the extension of the analysis to different sub-corpora, for which a similar pre-existing list is not available. Both these choices, however, are compensated for by the loss of some data, mainly due to the fact that the lemmatisation and PoS-tagging in the Diorisis corpus are not 100% accurate (as is inevitable in an electronically processed corpus). Without electronic tools, however, it would be impossible to collect and process the amount of data involved in analysing Adjective + Noun pairs outside of the epic sub-corpus, a further step which was highlighted as essential at the end of §III.3.

Finally, the data collection process was structured so as not to assume a simple 'base form' for any of the pairs, but to record all instances of variation that occur. This significantly increases the descriptive power of the model, and leaves open the possibility for the actual base form for each pair to be chosen based on relative frequency, without assuming that each formula should have originated as an Adjective + Noun pair with no additions. It is also possible to do away with the idea of a base form entirely in the data analysis, as will be discussed in §IV.2.1.

In short, most of the points laid out in §III.3.4 have been addressed by the setup of the upcoming analysis, and the negative trade-off in accuracy while using a pre-existing electronic corpus is amply compensated by the benefits in feasibility and reproducibility of the analysis. What remains constant between the two analyses is the use of Adjective + Noun pairs as a target, rather than structures involving verbs, which would allow for more cases of tree-syntactic variation after Wulff's model. This would require a corpus with more

³²⁰ On this issue, see Jensen and McGillivray (2017), 154.

extensive syntactic tagging than what is provided in Diorisis; it also would not address the objections to Hainsworth's approach that were discussed in §II.2.2 and at the start of this chapter. Questions of syntactic dependency will be addressed in the part of this thesis dedicated to semantic analysis (Part III).

IV.2 Statistical analysis

IV.2.1 *Relative entropy*

We are now left with the task of comparing the behaviour of 32 Adj + N formulae (ranging from 104 to 20 occurrences, for a total of 1306 tokens) to that of 33 baseline phrases (ranging from 101 to 21 occurrences, for a total of 1077 tokens). We therefore need to look for a statistical measure that can compare the distribution of all forms of a formula, for each of the categories of variation listed above, with the average distribution of forms in the baseline corpus for the same category of variation. This measure also needs to avoid the assumption of a baseline form for each formula, that is, it cannot assume the primacy of either the unmodified or modified form for each type of variation.

Wulff's study, once again, provides us with a statistical measure that meets these requisites. This is known in information theory as Shannon entropy.³²¹ Entropy, a concept introduced by analogy with physics, measures how regularly distributed the possible states of a system are.³²² In our case, the system in question is an Adj + N pair, and its possible states are the ways it can appear according to a specific type of modification – for instance, whether it appears with a preposition or not. The system in our example has two possible states: +P and -P. It can be the case that each of these states has more or less the same likelihood of occurring in our data; in other words, the two states are more or less equiprobable. A

³²¹ Wulff's very brief explanation of entropy as a statistical measure can be found in Wulff (2008), 83-6.

³²² The entropy of a system with n possible states is $H(X) = -\sum_{i=1}^n P(x_i) \log_2 P(x_i)$.

system with this configuration will be quite difficult to predict: for every occurrence of the Adj + N pair I observe, if I were to guess whether it occurs with a preposition or not, I would have a 50/50 chance of being right. This hypothetical, very unpredictable system has high entropy. But let us assume that our Adj + N pair behaves similarly to the Homeric $\nu\eta\upsilon\sigma\iota$ $\theta\omicron\tilde{\eta}\sigma\iota$, which as we saw in §III.3.3 overwhelmingly occurs with a preposition (the probability of state +P is much higher than that of state -P); in this case, I would have a much easier job guessing whether I will see it with a preposition or not. This second system is much more predictable; it has much lower entropy. Crucially for our purposes, it does not matter whether the state that occurs more often is the one where something is added (+P) or the one where it is not (-P); what matters is that the system is dramatically skewed towards one outcome.

Specifically, for the purposes of this study, I will use relative entropy, which takes into account the number of possible state of the system and varies between 0 and 1. The relative entropy of a system will be higher the closer to equiprobable (equally likely) its possible states are: for instance, the relative entropy of a fair coin toss (two equiprobable states) is 1. Most of our parameters of variation (see §IV.1.5) only have two possible states, as in the example above; these two states can be with extra material/without extra material, singular/plural, Adj N/N Adj order. Only case can occur in four different possible states (N/V, A, G, D).

Our general expectation, much like with idioms, is for formulae to be less flexible than their non-formulaic counterparts: in other words, we expect formulae to mostly occur in one possible state (form), while non-formulaic Adj + N phrases in the baseline should not be skewed towards one form. Baseline phrases are our fair coin toss, while formulae are a coin toss in which the pressure of oral tradition has (we expect) skewed the outcome towards one of the possibilities. When comparing the behaviour of a formula to the average behaviour of constructions in a baseline corpus, therefore, we expect the relative entropy

of the formula to be generally lower than the baseline. This expectation, however, may or may not hold depending on (a) the individual formula being considered and (b) the type of modification.

IV.2.2 Results

Table 9 on the next four pages contains the relative entropy for each of our 32 Adj + N formulae, compared with the average relative entropy of all Adj + N phrases in the baseline corpus, as reported in the first line. Results that are not in line with our initial prediction that $H_{\text{baseline}} > H_{\text{formula}}$ are highlighted by shading the corresponding cell.

Note that a relative entropy value of 0 means that all occurrences show only one of the possible states (for instance, all occurrences of this specific formula are in Adj + N order, with none being N + Adj), while a relative entropy value of 1 means that occurrences are equally divided between all possible values (for instance, out of 8 occurrences, 2 are in the nominative/vocative, 2 in the genitive, 2 in the dative, and 2 in the accusative). All of these results will be discussed in detail in §IV.3.

Table 9: Relative entropy values for each Adj + N pair in archaic epic

		tree-syntac tic f.	lexico-syntactic flexibility							morphological flexibility								metrical flexibility			
formula	n. of tokens	relativ e clause	adjecti ve insertio n	noun inserti on	modifiy ng NP/PP	post-modif. rel. cl.	preposi tion	adverb	partic le	negati on	article	comp. / sup.	number		case				word order		separ ation
													S	PI	N/ V	G	D	A	Adj N	N Adj	
baseline values	1077	0.057	0.555	0.111	0.476	0.133	0.613	0.184	0.745	0.275	0.838	0.094	0.825		0.978				0.949		0.959
average H _{formulae}	1306	0.017	0.421	0.145	0.182	0.079	0.274	0.008	0.321	0.019	0.084	0.000	0.257		0.467				0.377		0.441
ναῦς θοός	104	0	10	1	11	0	66	1	4	0	0	0	60	44	6	0	44	54	35	69	43
relative entropy		0.000	0.457	0.078	0.487	0.000	0.947	0.078	0.235	0.000	0.000	0.000	0.983		0.627				0.921		0.978
πατρίς γαῖα	88	0	52	1	0	0	61	0	1	0	0	0	88	0	0	3	7	78	85	3	3
relative entropy		0.000	0.976	0.090	0.000	0.000	0.889	0.000	0.090	0.000	0.000	0.000	0.000		0.305				0.215		0.215
θεός ἀθάνατο ς	81	0	17	3	0	5	21	0	14	0	0	0	1	80	18	2	56	5	68	13	23
relative entropy		0.000	0.741	0.229	0.000	0.334	0.826	0.000	0.664	0.000	0.000	0.000	0.096		0.374				0.635		0.861
υἱός φίλος	77	0	13	2	29	6	1	0	3	0	0	0	74	3	35	0	0	42	70	7	1
relative entropy		0.000	0.655	0.174	0.956	0.395	0.100	0.000	0.238	0.000	0.000	0.000	0.238		0.497				0.439		0.100
ναῦς μέλας	75	0	28	4	4	0	41	0	8	0	1	0	58	17	12	1	40	22	12	63	4
relative entropy		0.000	0.953	0.300	0.300	0.000	0.994	0.000	0.490	0.000	0.102	0.000	0.772		0.754				0.634		0.300
ἥτορ φίλος	55	0	0	9	9	0	0	0	10	0	0	0	55	0	55	0	0	0	55	0	13
relative entropy		0.000	0.000	0.643	0.643	0.000	0.000	0.000	0.684	0.000	0.000	0.000	0.000		0.000				0.000		0.789
ἄλλος θεός	46	2	23	2	0	0	2	0	15	0	3	0	4	42	37	0	7	2	24	22	19
relative entropy		0.258	1.000	0.258	0.000	0.000	0.258	0.000	0.911	0.000	0.348	0.000	0.426		0.431				0.999		0.978

		tree- syntac tic f.	lexico-syntactic flexibility							morphological flexibility								metrical flexibility			
formula	n. of tokens	relativ e clause	adjecti ve insertio n	noun inserti on	modifiy ng NP/PP	post- modif. rel. cl.	preposi tion	adverb	partic le	negati on	article	comp. / sup.	number		case				word order		separ ation
													S	PI	N/ V	G	D	A	Adj N	N Adj	
ἡμᾶρ πᾶς	51	0	1	0	0	0	0	0	2	0	0	0	2	49	51	0	0	0	3	48	3
relative entropy		0.000	0.139	0.000	0.000	0.000	0.000	0.000	0.239	0.000	0.000	0.000	0.239		0.000				0.323		0.323
ναῦς γλαφυρ ός	51	0	4	0	0	0	43	0	1	0	0	0	7	44	5	0	24	22	5	46	37
relative entropy		0.000	0.397	0.000	0.000	0.000	0.627	0.000	0.139	0.000	0.000	0.000	0.577		0.682				0.463		0.848
πᾶς θεός	41	2	9	1	0	2	2	0	13	0	1	0	0	41	17	0	22	2	29	12	22
relative entropy		0.281	0.759	0.165	0.000	0.281	0.281	0.000	0.901	0.000	0.165	0.000	0.000		0.611				0.872		0.996
φίλος πατήρ	41	0	1	3	4	0	5	0	5	0	0	0	41	0	7	11	16	7	17	24	13
relative entropy		0.000	0.165	0.378	0.461	0.000	0.535	0.000	0.535	0.000	0.000	0.000	0.000		0.955				0.979		0.901
ἐταῖρος φίλος	38	0	1	0	0	2	4	1	8	0	0	0	21	17	8	12	9	9	28	10	26
relative entropy		0.000	0.176	0.000	0.000	0.297	0.485	0.176	0.742	0.000	0.000	0.000	0.992		0.991				0.831		0.900
οὐρανός εὐρύς	40	0	1	1	0	0	8	0	2	0	0	0	40	0	5	0	0	35	0	40	0
relative entropy		0.000	0.169	0.169	0.000	0.000	0.722	0.000	0.286	0.000	0.000	0.000	0.000		0.272				0.000		0.000
ἄνθρωπ ος θνητός	39	0	1	12	0	0	2	0	17	6	0	0	0	39	3	23	10	3	38	1	8
relative entropy		0.000	0.172	0.890	0.000	0.000	0.292	0.000	0.988	0.619	0.000	0.000	0.000		0.761				0.172		0.732
ὥκὺς ἵππος	37	0	3	0	2	2	1	0	0	0	0	0	0	37	11	4	0	22	31	6	1
relative entropy		0.000	0.406	0.000	0.303	0.303	0.179	0.000	0.000	0.000	0.000	0.000	0.000		0.657				0.639		0.179
ναῦς κόϊλος	36	0	4	0	3	0	35	0	0	0	0	0	9	27	1	0	17	18	35	1	35
relative entropy		0.000	0.503	0.000	0.414	0.000	0.183	0.000	0.000	0.000	0.000	0.000	0.811		0.577				0.183		0.183

		tree- syntac tic f.	lexico-syntactic flexibility							morphological flexibility									metrical flexibility		
formula	n. of tokens	relativ e clause	adjecti ve insertio n	noun inserti on	modifiy ng NP/PP	post- modif. rel. cl.	preposi tion	adverb	partic le	negati on	article	comp. / sup.	number		case				word order		separ ation
													S	PI	N/ V	G	D	A	Ad j N	N Adj	
ὕσμίνη κρατερός	36	0	1	0	0	0	32	0	0	0	0	0	26	10	1	3	12	20	36	0	0
relative entropy		0.000	0.183	0.000	0.000	0.000	0.503	0.000	0.000	0.000	0.000	0.000	0.852		0.721				0.000		0.000
χαλκός όξύς	35	0	0	0	0	0	0	0	0	0	0	0	35	0	0	0	35	0	35	0	0
relative entropy		0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000		0.000				0.000		0.000
υιός ἀγλαός	32	0	1	0	30	0	0	0	0	0	0	0	32	0	24	0	0	8	32	0	0
relative entropy		0.000	0.201	0.000	0.337	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000		0.406				0.000		0.000
πᾶς ἐταῖρος	27	0	9	1	0	1	0	0	5	0	0	0	0	27	8	2	1	16	24	3	9
relative entropy		0.000	0.918	0.229	0.000	0.229	0.000	0.000	0.691	0.000	0.000	0.000	0.000		0.711				0.503		0.918
πᾶς ἄνθρωπος	25	0	4	1	0	0	9	0	6	0	0	0	0	25	3	9	5	8	23	2	17
relative entropy		0.000	0.634	0.242	0.000	0.000	0.943	0.000	0.795	0.000	0.000	0.000	0.000		0.944				0.402		0.904
ἐγχος δολιχόσ κιος	26	0	2	0	0	0	0	0	0	0	0	0	26	0	26	0	0	0	25	1	1
relative entropy		0.000	0.391	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000		0.000				0.235		0.235
ἄλλος ἀνὴρ	24	0	10	1	4	1	0	0	4	0	4	0	13	11	18	4	1	1	19	5	14
relative entropy		0.000	0.980	0.250	0.650	0.250	0.000	0.000	0.650	0.000	0.650	0.000	0.995		0.562				0.738		0.980
φαεινός δόρυ	25	0	0	0	0	0	0	0	0	0	0	0	25	0	0	3	22	0	3	22	0
relative entropy		0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000		0.265				0.529		0.000

		tree-syntac tic f.	lexico-syntactic flexibility							morphological flexibility									metrical flexibility		
formula	n. of tokens	relativ e clause	adjecti ve insertio n	noun inserti on	modifiy ng NP/PP	post- modif. rel. cl.	preposi tion	adverb	partic le	negati on	article	comp. / sup.	number		case				word order		separ ation
													S	PI	N/ V	G	D	A	Adj N	N Adj	
αἰπύς ὀλεθρος	25	0	3	0	0	0	0	0	0	0	0	0	25	0	10	0	0	15	25	0	0
relative entropy		0.000	0.529	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000		0.485				0.000		0.000
ἐγχος χάλκεος	24	0	1	0	0	0	0	0	0	0	0	0	24	0	24	0	0	0	23	1	1
relative entropy		0.000	0.250	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000		0.000				0.250		0.250
ἀγλαός δῶρον	23	0	0	0	1	0	0	0	0	0	0	0	0	23	23	0	0	0	23	0	0
relative entropy		0.000	0.000	0.000	0.258	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000		0.000				0.000		0.000
ἄλλος ἑταῖρος	21	0	10	1	0	2	0	0	4	0	8	0	0	21	9	7	2	3	20	1	12
relative entropy		0.000	0.998	0.276	0.000	0.454	0.000	0.000	0.702	0.000	0.959	0.000	0.000		0.888				0.276		0.985
ἐρίηρος ἑταῖρος	21	0	2	0	0	0	0	0	0	0	2	0	1	20	11	0	0	10	20	1	0
relative entropy		0.000	0.454	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.454	0.000	0.276		0.499				0.276		0.000
ὄξύς δόρυ	21	0	1	0	0	0	0	0	0	0	0	0	13	8	10	0	11	0	20	1	1
relative entropy		0.000	0.276	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.959		0.499				0.276		0.276
υἱός φαιδίμο ς	21	0	0	1	10	0	0	0	1	0	0	0	21	0	13	0	0	8	20	1	1
relative entropy		0.000	0.000	0.276	0.998	0.000	0.000	0.000	0.276	0.000	0.000	0.000	0.000		0.479				0.276		0.276
λευκώλε νος θεά	20	0	0	0	0	0	0	0	0	0	0	0	20	0	20	0	0	0	0	20	0
relative entropy		0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000		0.000				0.000		0.000

IV.3 Discussion of results

IV.3.1 *The general picture*

If we only look at the first lines of our table (the average entropy for all baseline phrases and the average entropy for all formulae), the overall picture seems deceptively clear-cut: formulae are always less flexible than their baseline counterparts, with the exception of cases in which a noun is added to the basic Adj + N phrase, a type of modification under which formulae show slightly (but only very slightly) higher flexibility on average.

As is often true, however, the average values actually hide quite a significant amount of variation. The main advantage of our fine-grained analysis is that we are able to look into this variation and identify the roles played by individual variables in determining syntactic flexibility. In the following sections, we will start by examining variation in individual formulae; we will then try to determine which parameters of variation are more characteristic of formulaic behaviour, in that they show much higher or much lower levels of flexibility in formulae compared to baseline values.

As a first step, however, the behaviour of our formulae can be divided into four classes depending on how many parameters show $H_{\text{formula}} > H_{\text{baseline}}$. 11/32 formulae, about a third of the total, show no parameters where $H_{\text{formula}} > H_{\text{baseline}}$; 4/32 show $H_{\text{formula}} > H_{\text{baseline}}$ in one parameter; 13/32, another generous third, show higher entropy in two to four parameters; and finally, only 4/32 (1/8) show higher entropy in at least five parameters (that is, in one out of three of the available types of variation). ἄλλος + θεός and ἄλλος + ἀνὴρ are the two most flexible pairs, with entropy higher than the baseline in six parameters. It must be kept in mind throughout the discussion that follows, then, that for most of the available types of flexibility formulae do behave more rigidly than non-formulaic pairs; this overall observation stands. However, the crucial ways in which it is nuanced will be important in our final discussion of what characterises formulaic behaviour.

IV.3.2 *Individual formulae: frequency effects*

In terms of the behaviour of individual formulae, one specific tendency may be visible just by looking at the pattern of shaded cells (where $H_{\text{formula}} > H_{\text{baseline}}$) in Table 9. The table is ordered by decreasing token frequency; the further down the table one moves, the more uncommon shaded cells appear to be. Since higher entropy translates to a higher degree of flexibility, higher frequency formulae seem to be more flexible than lower frequency ones.³²³

Quantitative analysis confirms this first impression. We saw that, on average, only one parameter of variation showed higher flexibility in the formulaic sample, but in fact just over half of the individual formulae (17 out of 32) show higher flexibility than the baseline in at least two parameters. These, however, are not equally spread across the formulaic sample: of the formulae that occur more often than the median token frequency (which is 36), 12 out of 15, so all but 3, show higher flexibility than the baseline in at least two parameters; on the other hand, the moment we dip (equal or) under the median, the proportion is reversed to 5 formulae out of 17.³²⁴ (We will see in the next section that 4 of these 5 lower-frequency formulae share another characteristic trait which promotes their flexibility despite their being less common.)

What we find, then, is exactly the same frequency effect observed in §III.3.3; the conclusions from that chapter hold even when we introduce the comparison with a

³²³ It is true that the fewer tokens we have in a distribution, the more easily these can appear to be skewed in a particular direction (it is more likely for a coin tossed four times to appear unbalanced than it is if we toss it fifty times). However, by setting a fairly high frequency threshold (20 occurrences), we should have limited this issue to a reasonable degree.

³²⁴ Counterintuitively, the median token frequency does not exactly split the sample in half, as the two formulae that are ranked 16th and 17th by token frequency occur the same number of times. So, if we compute the median token frequency (which will be equal to the token frequency of both of these pairs), and then set a threshold based on it, both pairs number 16 and 17 will fall on one side of it.

baseline. The increase in flexibility that comes with formulaic frequency is even more striking in light of these results: not only are more frequent formulae more flexible compared to less frequent formulae, but they can even be more flexible than non-formulaic expressions, at least under some kinds of operations. This goes against the ideas of Albert Lord among others, who argued that more frequent formulae should become more entrenched in the singers' memories and therefore be repeated without variation.³²⁵ Therefore, even accounting for the fact that, as discussed in §IV.3.1, formulaic flexibility remains overall relatively low, 'higher resistance to modification' (or, in other words, the tendency to appear overwhelmingly more frequently in a specific form) can hardly be *the* characterising trait of formulae: among the most frequent formulae are noun-epithet pairs for ships, such as ναῦς + θεός or ναῦς + μέλας, which are strikingly recognizable as Homeric phraseology. We will see in §IV.3.4 that a much more promising way of distinguishing between formulae and non-formulaic phrases lies in the fact that they behave differently depending on what *type* of modification we are looking at, rather than avoid *all* types of modification indiscriminately.

IV.3.3 *Individual formulae: the quantifiers πᾶς and ἄλλος*

Before moving on to individual parameters, however, we must address one more trait that correlates strikingly with formulaic flexibility. We have seen in the previous section that only 5 formulae out of the 17 that occur 36 times or less (the median frequency) show higher entropy than the baseline in at least two parameters. These formulae are πᾶς + ἑταῖρος (27x, 3 high-entropy parameters), πᾶς + ἄνθρωπος (25x, 4 parameters), ἄλλος + ἀνὴρ (24x, 6 parameters), ἄλλος + ἑταῖρος (21x, 5 parameters), and υἱός + φαίδιμος (21x, 2 parameters). Apart from the last one, all of these have a strikingly high number of flexible

³²⁵ Lord (2000), 43, as discussed in §1.2.2.

parameters, and also contain a quantifier (πᾶς or ἄλλος). Similarly, if we look at the 4 formulae that show at least 5 flexible parameters, ἄλλος + ἀνὴρ and ἄλλος + ἐταῖρος are joined by ἄλλος + θεός (46x, 6 parameters) and πᾶς + θεός (41x, 6 parameters), that is, two formulae with the same noun and two different quantifiers. Formulae with ἄλλος and πᾶς, in short, appear to be much more flexible than other pairs; the only exception is the pair ἥμαρ + πᾶς, which almost always occurs in the form ἡμετα πάντα.

We have already discussed in §III.3.3 how phrases with ἄλλος are connected with the introduction of the article in archaic Greek, that is, with a significant example of linguistic innovation; it is possible to think that πᾶς also promotes a similar kind of development, especially since the article can also be used in a distinctive function with πᾶς in Attic Greek (compare πᾶς ἀριθμός ‘every number’ with πᾶς ὁ ἀριθμός ‘the whole number’).³²⁶ It is worth noting, however, that only one of these phrases, ἄλλος + ἐταῖρος, actually shows higher flexibility than the baseline in the use of the article. This could be, of course, simply a consequence of the fact that article usage is much more developed in our Attic baseline corpus; but again, this does not account for the increased flexibility of phrases with ἄλλος and πᾶς along other axes of variation, nor for the fact that two of these phrases (not counting ἥμαρ + πᾶς), πᾶς + ἐταῖρος and πᾶς + ἄνθρωπος, never occur with the article at all.

Table 10 below summarises the distinctive traits of formulae with ἄλλος and πᾶς, as discussed in this section. Note that, while trends in flexibility are broadly similar, formulae with πᾶς are on average slightly less flexible than formulae with ἄλλος, and that they do not seem to accept the article as easily (with the exception of πᾶς + θεός).

³²⁶ See van Emde Boas et al. (2019), §29.45. I follow their terminology in calling πᾶς and ἄλλος ‘quantifiers.’ None of the other Greek quantifiers they list (ὅλος, μόνος, ἕτερος) appears in the formulaic sample.

Table 10: Behaviour of formulae with ἄλλος and πᾶς

Formula	Frequency	N. of parameters where $H_{\text{formula}} > H_{\text{baseline}}$	Is H for the article higher than the formulaic average?
ἄλλος + θεός	46	6	yes
ἡμαρ + πᾶς	51	0	no
πᾶς + θεός	41	6	yes
πᾶς + ἐταῖρος	27	3	no
πᾶς + ἄνθρωπος	25	4	no
ἄλλος + ἀνὴρ	24	6	yes
ἄλλος + ἐταῖρος	21	5	yes, even $> H_{\text{baseline}}$

The increased rate of variation could also be explained by the particular semantics of phrases with ἄλλος and πᾶς. Both these adjectives are semantically very ‘light,’ and they can potentially combine with pretty much any noun;³²⁷ they may therefore have been less likely to trigger the ‘mutual expectation’ response, as they co-occur more freely outside of a specific formulaic pair (indeed, just our sample already shows three different phrases with ἄλλος and four with πᾶς, while the only other adjectives that occur with more than one noun are ὀξύς and φίλος, each with two possible combinations). It is therefore possible that phrases involving them were perceived as more generic, less formulaic, in some way. In this sense, the peculiar behaviour of these formulae can be explained within Hainsworth’s approach to formularity as a wider concept, as a weakening of mutual expectation, without recurring to external pressure from the everyday language.

Another element that goes in favour of this interpretation is that the presence of ἄλλος and πᾶς does not appear to influence the behaviour of baseline Adj + N pairs (which, in Hainsworth’s interpretation, are not subject to the same bond of mutual expectation). There are 14 pairs in the baseline sample that contain either ἄλλος or πᾶς; their average entropy is higher than the entropy for the whole baseline in 6 parameters (including article usage), but lower in the remaining 9. This also means that the higher frequency of ἄλλος and πᾶς in our baseline sample (14 pairs out of 33, instead of 7 out of 32) did not skew the results in either direction.

³²⁷ Much like ‘stop-words’ in semantic analysis (see §VI.2.2).

In any case, our data show two separate tendencies in determining which formulae are more or less flexible: one of these is frequency-based and independent from content or semantics, while the other is triggered by specific words. Both promote flexibility in the relevant formulae. Frequency of repetition, therefore, is once again not the only trait at play in determining formulaic flexibility.

IV.3.4 *Individual parameters*

Looking at the behaviour of individual formulae tells us something about the factors that affect formulaic variation. This is, however, not the only question with which we started this project. Another equally important question was of the sort that is asked by Wulff (2008) in her study of English idioms: what are the types of variation that are most characteristic of a formula? In other words, are formulae regularly less flexible than non-formulaic phrases in all parameters, with no distinction between these, or are there parameters along which formulae are more likely to vary? If the latter is the case, these parameters will be less indicative of formularity, as fixity along these parameters is not a necessary characteristic of a formula; on the other hand, if there are any parameters along which formulae are regularly very rigid, those will be the ones we should be looking for when defining *formulaic* behaviour, specifically.

Wulff's categorisation, together with Hainsworth's work, gives us a way to group our types of variation, so we are not limited to discussing the behaviour of individual traits. Note that an entire class of Wulff's parameters, the ones she calls 'lexico-syntactic,' map straightforwardly onto what Hainsworth calls 'expansion;' intuitively, these are the parameters that involve adding something to the formula. Among these, the addition of a preposition is partly exceptional, in that it also determines potential changes in the case of the noun and adjective.

Table 11 below summarises the results in Table 9, organised by type of variation. A detailed description of each type can be found in §III.3.1 and IV.1.5.

Table 11: Formulaic behaviour by type of variation

Wulff's category ³²⁸	Type of variation	Baseline entropy	Number of formulae where $H_{\text{formula}} > H_{\text{baseline}}$ (out of 32)
tree-syntactic flexibility	relative clause	0.057	2 (6.25%)
lexico-syntactic flexibility	adjective insertion	0.555	10 (31.25%)
	noun insertion	0.111	14 (43.75%)
	modifying NP/PP	0.476	5 (15.63%)
	post-modifying relative clause	0.133	8 (25.00%)
	preposition	0.613	7 (21.88%)
	adverb	0.184	0
	particle	0.745	4 (12.50%)
morphological flexibility	negation	0.275	0
	article	0.838	1 (3.13%)
	comparative/superlative adjective	0.094	0
	number	0.825	5 (15.63%)
	case	0.978	1 (3.13%)
'metrical' flexibility	word order	0.949	2 (6.25%)
	separation	0.959	5 (15.63%)

There are two things to keep an eye on here. One is, of course, which parameters allow for more formulae to violate the tendency that has $H_{\text{baseline}} > H_{\text{formula}}$. The other is which parameters have higher or lower baseline entropy rates: a parameter with low entropy in the baseline will be something for which not much variation is present anyway, and therefore will be less characteristic even though the rate of formulaic variation is even lower. Note, once again, that by choosing entropy as a flexibility measure we have entirely sidestepped one of the main flaws in Hainsworth's design, that is, that a formula that *always* occurs with some kind of expansion would appear to be extremely flexible, even if it was always the same addition, because Hainsworth measures flexibility as deviation from a 'base form,' which is the shortest variant of the formula available. With entropy, on the other hand, formulae that always occur with an expansion will still rate very low on the entropy scale, as their behaviour is very predictable.

³²⁸ Except for 'metrical' flexibility, which was added based on Hainsworth's work.

Let us discuss each category of variation based on these two criteria. Tree-syntactic flexibility, which is just a small remnant of what is a much wider category in Wulff's work, is not very helpful either way: formulaic flexibility is very low, for sure, but baseline flexibility is also close to 0, which makes this type of variation too rare to be useful. The next class, lexico-syntactic flexibility, is more interesting, as it contains all the parameters with the highest number of high entropy formulae, going all the way up to 14 for the addition of a noun. On the other hand, none of these types of variation are exceedingly common in the baseline corpus. We have, therefore, a category of phenomena where formulae are allowed more flexibility than for other categories, and where baseline phrases themselves are not too flexible; in other words, the behaviour of the two classes, formulaic and non-formulaic, is most similar here. Lexico-syntactic flexibility, or Hainsworth's expansion, is the category of variation that is the *least* characteristic of formulaic behaviour.

Things are very different when we look at the other two categories of variation, morphological and 'metrical' flexibility. With the exclusion of negation and comparative/superlative grade of the adjective, which have low entropy in the baseline corpus and extremely low entropy in the formulaic one, all other parameters (use of the article, inflection in case and number, changes in word order, and separation with any intervening material) show very high flexibility in the baseline corpus, while the flexibility of formulae is hardly ever higher than that of the baseline. These categories are, therefore, the most characteristic of formulaic behaviour: while 'ordinary' (non-formulaic) language does not show any preference for specific forms, formulae generally have a marked preference for one option, or a restricted number of options. The idea that morphological and metrical flexibility are the areas in which formulae are more rigid, both in themselves and compared with the baseline, also fits the expectation that operations that change the internal shape of a formula, either by inflecting it or changing the position of its parts, will be more difficult to accommodate in poetry.

IV.3.5 General conclusions

We started our analysis of syntactic flexibility with the goals of exploring the behaviour of formulae on a fine-grained level, by looking into individual variation and variation along single parameters, and of overcoming the limitations in Hainsworth's experimental design, in particular as it comes to the choice of a 'baseline formula.' We have already discussed in detail how the latter was addressed in §IV.2.1.

The data discussed in this chapter give us a very detailed picture of the behaviour of a sizable sample of Adjective + Noun formulae. When it comes to individual tendencies, we have noted how both frequency and the presence of quantifiers promote syntactic flexibility independently of each other; we have hypothesised that quantifiers may have that effect because they can combine with many different nouns, and therefore lead to a less strong 'bond of mutual expectation' within the parts of the formula. Moreover, we have established which parameters are more characteristic of formulaic behaviour: morphological and metrical flexibility are the areas in which the behaviour of formulae is most different from that of baseline phrases, while types of variation that fit in Hainsworth's category of expansion (flexibility by way of addition of material) are less distinctive.

It is worth spending a moment discussing how these findings – particularly the frequency effects and the different kinds of modifications – relate to our expectations regarding how formulae are used and perceived. As we observed at the end of chapter III.3, evidence that frequency promotes formulaic flexibility shows that formulae are not memorised as pre-made constituent blocks associated with a metrical position; this kind of scenario would lead to either higher entrenchment, that is, lower flexibility, of frequent formulae, or to a more or less constant rate of variation, as blocks are re-used without change most of the

time but may admit the occasional variant. This is further evidence of the kind of formulaic system envisaged by Hainsworth, where modification is an effect of performance and a useful tool for language use.

On the other hand, when we look at individual axes of variation, we find that those modifications that most substantially disrupt the metrical shape of a formula (inflection, including changes in number, and changes in word order, including separation) are the least common, while the most common class is the one where peripheral material is added to³²⁹ an existing nucleus. This restores some credit to Parry's (and, to an extent, Hainsworth's) intuition that formulae do have an essential shape, a primary form, that can be modified but opposes substantial resistance to this modification.

What remains clear is how, by exploring in detail how variation happens in our epic corpus, we can expand our knowledge of the inner workings of formularity. We can also address the objection that by admitting flexibility as an inherent trait of formulae we risk bringing them too close to non-formulaic phrases that occur with a similar frequency: we have shown that, contrary to this idea, the behaviour of formulae remains qualitatively and quantitatively characterised as different from similar structures in non-oral texts. An objection that remains to be addressed is whether this type of recurring, partially variable phrases is a characteristic of Greek hexameter poetry, and epic in particular, rather than oral epic poetry in general. The next chapter, on recurring Adj + N pairs in Apollonius of Rhodes, is dedicated to answering this question.

³²⁹ Or indeed subtracted from, since our flexibility measure does not tell us in which direction the modification occurred: a hypothetical formula with an entropy of 1 for adjective expansion is either an Adj + N base to which a second adjective gets added 50% of the time, or an Adj + Adj + N base which has lost one of the two adjectives in half of its occurrences.

V Comparing poetic grammars: A brief look at non-formulaic epic

The previous two chapters were dedicated to exploring the limits of formulaic variation on the level of syntax broadly conceived, from inflection to insertion of new material. The most salient variation factors were identified through the comparison with a baseline sub-corpus of non-epic material from the 5th century BC. These results answered some of our questions about the differences in behaviour between the formulaic language of archaic Greek epic and (on some level, limited by the fact that we mostly have access to literary texts) 'natural' linguistic usage. We saw that, while formulae are more flexible than some traditional definitions want them to be (even when maintaining lexical identity between their constituent parts), they still markedly differ from non-formulaic structures, especially when it comes to transformations that more invasively change their metrical shape, such as inflection in case and number, changes in word order, and separation through the insertion of intervening material.

Non-epic language, however, is not the only possible point of comparison for archaic epic. We also have a body of epic poetry that spans centuries after Homer, and which shares a metre and, broadly speaking, a series of topics and generic concerns with archaic epic. Comparing the results we obtained from Homeric and Hesiodic material in the previous chapters to the results of the same methods applied to non-archaic epic allows us to run a more controlled kind of experiment: rather than looking for a baseline sample for comparison that is as different as possible from the epic sample we were working on, we will be working with texts that share a genre, a metre, and literary models with archaic epic (or rather, for which archaic epic is the literary model),³³⁰ but which were not produced

³³⁰ On style and models in Hellenistic epic see at least Fantuzzi and Hunter (2004), 246-82, who illustrate the ways in which Hellenistic authors situate themselves within the epic tradition while at the same time distancing themselves from an aesthetic of repetition and drawing on lyric poetry and other sources.

within an oral tradition, even in the broadest understanding of the term.³³¹ In this chapter, I aim to try and isolate which (if any) effects, among the ones we observed in the previous experiments, can be more likely identified as the result of oral composition, and which (if any) are instead a characteristic of the epic genre broadly understood.³³² Where the latter is the case, we will also have to ask whether a feature was adopted merely through influence from the earlier tradition, or if it is useful to epic composition in some other way.

V.1 The case for Apollonius Rhodius

This chapter will focus on a single poem, Apollonius Rhodius' *Argonautica*.³³³ Composed in the 3rd century BC, whatever the much-debated details of their composition,³³⁴ the *Argonautica*'s four books follow the epic journey of Jason and the Argonauts to Colchis, their conquest of the Golden Fleece, and their return to Iolkos, a journey which reaches the dimensions of a veritable Odyssey in one book.

There is one straightforward reason to focus on the *Argonautica* as a good point of comparison for archaic Greek epic: it is the only complete surviving Hellenistic epic poem. Callimachus' *Hymns*, which would provide another suitable body of texts to work with, are shorter, and self-consciously belong to the hymnic genre; this is without even addressing the issue of hymns 5 and 6, *For the Bath of Pallas* and *To Demeter*, both of which are in a Doricising dialect that deliberately departs from the Homeric hymnic tradition, with the

³³¹ The *Homeric Hymns* likely originate from a different, much later setting than the Homeric poems, but they adopt the same traditional language and the same ways of modifying it (Postlethwaite 1979; Janko 1982; Faulkner 2011).

³³² The latter is the view advanced by Friedrich (2019), on whose criticism of Hainsworth's flexible formulae see §II.2.2.

³³³ For a general introduction to Apollonius' work see Fantuzzi and Hunter (2004), 89-132, as well as the individual contributions in Papanghelis and Rengakos (2008).

³³⁴ For a very concise and useful summary of the issues around dating and composition, see Hunter (2015), 1-3, with further references.

Bath of Pallas not even being in hexameters, but in elegiac couplets.³³⁵ It is therefore preferable to focus exclusively on the *Argonautica*, to avoid introducing too many confounding factors into the data.³³⁶

The second reason to focus on Apollonius in particular is his highly individual use of Homeric language and of formulaicity as a device, which will be discussed in detail in §V.1.1. Like other Hellenistic poets, Apollonius has a very self-conscious attitude to the language of archaic, and especially Homeric, epic.³³⁷ While this makes him an interesting target for comparison, as he deliberately departs from Homeric style and formulaic repetition, it also reinforces the decision to focus on the *Argonautica* alone rather than extend the comparison to much later poetry. Some other likely candidates for this comparison would be Quintus Smyrnaeus' *Posthomerica* and/or Nonnus' *Dionysiaca*; both of these poets, however, have their own highly individual relationship to Homeric language,³³⁸ and extending the comparison to two poems from a whole six centuries later than Apollonius would highly confound the data. I chose therefore to limit the comparison to the *Argonautica* for this chapter, while leaving the possibility of extending the study to other epic poems, considered separately, to further research.

V.1.1 *Formulaicity in the Argonautica: a very brief overview*

Questions of formulaicity in Apollonius have been discussed by Marco Fantuzzi in a couple of detailed studies, whose conclusions I will summarise in the rest of this section.³³⁹ This

³³⁵ For a general introduction to Callimachus' *Hymns*, see Haslam (1993) and Stephens (2015). On the *Bath of Pallas* in particular, see Hunter (1992).

³³⁶ Milman Parry himself compared name-epithet formulae in Homer to Apollonius' usage: M. Parry (1971), 165-72.

³³⁷ On the relationship between Apollonius' style and his epic models in broader literary terms, see Hunter (1993), 101-51.

³³⁸ See Maciver (2018) on Quintus, Hopkinson (1994) and Geisz (2017) on Nonnus.

³³⁹ Fantuzzi (1988, 7-47; 2008). On formulaicity in Apollonius see also Ciani (1975), Ruiz Vizcaya (2001), who only considers speech introduction/conclusion formulae and is therefore of limited

will be helpful in understanding how, when critics talk about ‘formulariness’ in Apollonius, they are in fact referring to a phenomenon that is very different from his archaic models. In this chapter, we will explore how the linguistic effects of this usage (may be expected to) differ from what we saw in previous chapters.

First of all, Fantuzzi sets Apollonius’ usage in context by tracking the evolution of formulariness from a necessary aid in oral composition to a stylistic choice, “namely a form of imitation of the Homeric text as regular code and model,”³⁴⁰ a shift that he identifies as already present in early post-Homeric epic, the *Homeric Hymns* in particular. Formulaic repetition is later highlighted as potentially unpleasant in 5th-century comic sources. The Hellenistic editors of Homer had a complicated attitude to traditional epithets and formulae, whose function they did not always fully appreciate, to the extent that Zenodotus (and others) regularly excised repetitions, for example in messenger speeches.³⁴¹ This version of the Homeric text, with fewer formulae compared to how we think of it nowadays, may have influenced Apollonius; Fantuzzi, however, suggests that the Alexandrians’ textual choices themselves may have been influenced by the ‘desirable’ epic style of their times, of which Apollonius is one of the models.

In the 1988 study, Fantuzzi focuses on formulaic variation in Apollonius compared to the kind of variation identified by Hainsworth (1968) for Homer – an analysis that is of obvious relevance to our work.³⁴² One of the closest models for Hellenistic writers, Antimachus of Colophon, appears to either reuse Homeric formulae verbatim or construct new phrases based on Homeric lexicon and Homeric-sounding syntax; Apollonius, on the other hand, prefers to vary Homeric phrases through analogy or recombination of formulaic elements.

relevance to my discussion here, and briefly Bozzone (2014b), also on speech introductions with the verb προσέειπεν.

³⁴⁰ Fantuzzi (2008), 222.

³⁴¹ On these edits by Hellenistic philologists, see further Montana (2014); Montanari (2014); Nünlist (2014).

³⁴² Fantuzzi (1988), 31-41, also offers brief samples of ‘formulaic analysis’ of Apollonius’ text, but these only cover a few ten-line extracts from each book and are therefore unusable for our study.

An example is the regular presence of speech introduction phrases, which are however introduced by a much more varied range of *verba dicendi*, and never exactly repeat Homeric usage. Apollonius does not reuse the same phrase to describe the same situation (not to mention type scenes), and avoids repetition in reported speeches. Even “para-formulaic expressions,” that is, recurring ways to introduce a hero, an object, or a theme,³⁴³ are always carefully varied: for instance, the Golden Fleece is described alternatively as κῶας/δέρος χρύσειον/χρύσεον, thus maximising the potential for different combinations, which is further expanded by playing with word order and position in the line. While Fantuzzi uses this as an example of non-Homeric use of ‘para-formulae’ in Apollonius, a system of this kind is actually in line with Milman Parry’s requirements: it expresses the same idea, the Golden Fleece, in several ways that can fill different metrical slots.³⁴⁴ This sort of issue is yet another reason why it is important to look beyond which expressions are repeated, in order to see whether the behaviour of these expressions on a syntactic level is parallel to or departs from what we found out about Homeric formulae in §IV.

In terms of Hainsworth’s descriptors, Fantuzzi already notes the frequent use of hyperbaton and mobility, especially in Noun + Adjective pairs; if we remember that resistance to word order variation and separation were among the traits that most characterised formulaic N + Adj pairs compared to baseline usage (§IV.3.4), we can already see how this observation is promising for our argument. Another of Fantuzzi’s points is that the frequency of Homeric ‘re-use’ declines through the *Argonautica*, and that the intended effect of Homeric phrases appears to develop into something different from simple traditional reference as the poem progresses: while their first usage often directly calls upon a Homeric reference, if they are repeated later they seem to only draw attention to Apollonius’ previous usage.³⁴⁵ Generally, all these examples illustrate the idea that

³⁴³ Fantuzzi (2008), 232.

³⁴⁴ I owe this point entirely to Philomen Probert.

³⁴⁵ Fantuzzi (2008), 235.

Apollonius' reuse of traditional epic language is a conscious choice rather than 'subconscious,' 'natural' usage. Apollonius appears to seek a "para-formulaic effect"³⁴⁶ that is only initially based on an echo of Homeric usage (whose memory, transmitted through oral/aural performance and reception, is indeterminate and imperfect³⁴⁷), but then continues to be built through intra-textual repetition and variation, which can be more accurately tracked thanks to the written medium.

V.2 Syntactic variation

To gain a fuller picture of syntactic variation in Apollonius compared to the one already offered by Fantuzzi, I start by repeating the analysis that was run in chapter IV on Noun + Adjective pairs. As we will see in this section, a number of adjustments need to be made because of the smaller size of the *Argonautica* (5,835 lines as opposed to the over 30,000 of the epic sub-corpus described in §IV.1.2) and because of Apollonius' stylistic choices in terms of repetition, as discussed in §V.1.1.

V.2.1 Gathering the data

To gather data on the behaviour of Adj + N pairs,³⁴⁸ I ran the same Python script that was used in §IV.1.3 on the *Argonautica* file from the Diorisis corpus. The tagging system is the same as described in §IV.1.1 for the rest of Diorisis; the edition used is the one available in the Perseus files, that is, Mooney (1912).

³⁴⁶ Fantuzzi (2008), 241.

³⁴⁷ A case can definitely be made for written reception of Homer by this date, but Fantuzzi still consider Homeric references different from intratextual connections within the *Argonautica*.

³⁴⁸ I will not be using 'formula' for Apollonius' usage, for the reasons outlined in §V.1.1. That said, similar caution should be and has been used in §IV when discussing epic Adj + N pairs: under the approach adopted in this thesis, their position on the formulaic continuum is determined by their behaviour, not a characteristic assigned *a priori* by way of definition.

The first and most significant adjustment that needs to be made concerns the search parameters. In §IV, we used a window of maximum 6 words, inclusive of the noun and adjective, an additional context of 3 words on each side (again, inclusive), and a frequency threshold of 20 – that is, only Adj + N pairs that occur at least 20 times in the epic sub-corpus (and, separately, in the baseline sub-corpus) were retrieved by the script. Running the same analysis on Apollonius gives a predictable result: there are no Adj + N pairs that occur at least 20 times. Indeed, the most frequent recurring pair, πατήρ + έός, occurs 11 times (and, as we will see in a moment, must be discarded for other reasons).

The frequency threshold was therefore lowered to at least 5 occurrences, which gives us 13 Adj + N pairs, as reported in Table 12 below. Four of these involve possessive adjectives of the type έός / έμός, which were excluded in chapter IV as they are metrically and functionally interchangeable despite not being lexically identical: arguably, πατήρ + έμός and πατήρ + έός are the same formula, with variation depending exclusively on the referent of the possessive adjective, that is, on the sentence context. These four pairs were therefore excluded from the analysis. This leaves 9 pairs, one of which, άλλος + ήρωρ, turns out to only occur 4 times, not 5, due to a false positive retrieval by the script.³⁴⁹ In the end, this last pair was kept as part of the analysis, not least because it includes an indefinite adjective, which the results from §IV.3.3 indicate as a source of interesting phenomena. Since we will proceed to discuss the behaviour of each of these pairs on an individual basis, without drawing conclusions based on totals or averages, there is no particular issue with keeping one extra pair.

Table 12: Adjective + Noun pairs occurring at least 5 times in Apollonius Rhodius' Argonautica

Pair	Number of occurrences
πατήρ + έός	11
έός + χείρ	9
κῶας + χρύσεος	8
χείρ + δεξιτερός	8
τοῖος + μῦθος	8

³⁴⁹ This is the only false positive in this sample. On false positives and negatives, see §IV.1.4.

δέρος + χρύσεος	7
άνήρ + άλλοδαπός	5
ναῦς + θοός	5
πᾶς + ἥμαρ	5
πᾶς + ἐταῖρος	5
μήτηρ + ἐός	5
πατήρ + ἐμός	5
ἄλλος + ἥρως	4 (see above)

It is worth noting that, despite Apollonius' known avoidance of Homeric formulae, three of the pairs listed above (ναῦς + θοός, πᾶς + ἥμαρ and πᾶς + ἐταῖρος) are also in the list of early epic Adj + N pairs; ἄλλος + ἥρως is not present in the 'Homeric' list, but similar pairs such as ἄλλος + ἀνὴρ and ἄλλος + ἐταῖρος are. Other pairs, on the other hand, such as the two expressions denoting the Golden Fleece, are definitely characteristic of Apollonius.

Before we move on to the results of the analysis, something must be said about the implications of lowering the frequency threshold. While this is not really a choice, given how much smaller the Apollonian sample is, results in §IV.3.2 highlighted frequency as one of the factors that clearly influence formulaic behaviour, with more frequent formulae being consistently more flexible. If the same pattern survives in Apollonius's phrase usage, there is a risk that the lower number of pairs we have retrieved will limit the range of variation that we can observe. If we consider frequency to be a predictor of behaviour, depending on what shape the correlation between the two takes, there may be less of a difference between two phrases that occur 5 and 10 times than between two phrases that occur 50 and 100 times, despite the ratio being the same; unfortunately, however, the Apollonius sample limits us to the former scenario. Moreover, the range of variation in the frequency of target Adj + N pairs is smaller in Apollonius than in Homer: while the size of the early epic sub-corpus is 5-6 times the *Argonautica*, the most frequent pair in early epic, ναῦς + θοός, occurs 104 times, which is 13 times the number of occurrences for κῶας + χρύσεος. On account of all these factors, the following chapters will not focus on small differences in entropy values, but just on more significant broader tendencies.

V.2.2 Results and comparison with the baseline

The data analysis was also conducted with the same methods adopted in §IV: relative entropy scores were computed for each Adj + N pair under each type of variation,³⁵⁰ and then compared to the relative entropy scores of Adj + N phrases occurring at least 20 times in the baseline sub-corpus, as defined in §IV.1.2. The baseline used here is exactly the same as the one which was used for comparison with archaic Greek epic, that is, a set of 5th-century mixed-genre texts. The frequency threshold for Adj + N pairs in the comparison corpus was also left unchanged, as the much larger size of this dataset compared to the *Argonautica* makes it both impractical and inadvisable to lower the threshold all the way to 5. Moreover, to examine the differences in behaviour between the early epic sub-corpus and the *Argonautica*, it makes sense to compare them to the exact same benchmark, so that this factor remains constant and we can focus exclusively on the two target sub-corpora.

The results of the comparison are reported in Table 13, which follows the exact same format as Table 9 in §IV.3.1. The table also contains the average values for the early epic sub-corpus (average H_{formulae}), for ease of comparison. As was the case in Table 9, shaded cells indicate cases where the entropy for each formula is higher than the average entropy in the baseline for this specific parameter; additionally, cells with underlined values indicate where the entropy for a pair in Apollonius is higher than the average entropy for the archaic epic sample. As always, an entropy of 0 indicates that the pair only appears in one form in all its attestations pertaining to a specific parameter, while an entropy of 1 indicates that occurrences are split exactly equally between possible states.

³⁵⁰ For the definition and applications of relative entropy, see §IV.2.1.

Table 13: Relative entropy values for each Adj + N pair in Apollonius

		tree-syntac tic f.	lexico-syntactic flexibility							morphological flexibility								metrical flexibility			
formula	n. of tokens	relativ e clause	adjecti ve insertio n	noun inserti on	modifiy ing NP/PP	post-modif. rel. cl.	preposi tion	adverb b	partic le	negati on	article	comp. / sup.	number		case				word order		separ ation
													S	PI	N/ V	G	D	A	Adj N	N Adj	
baseline values	1077	0.057	0.555	0.111	0.476	0.133	0.613	0.184	0.745	0.275	0.838	0.094	0.825		0.978				0.949		0.959
average H _{formulae}	1306	0.017	0.421	0.145	0.182	0.079	0.274	0.008	0.321	0.019	0.084	0.000	0.257		0.467				0.377		0.441
average H _{Apollonius}	55	0.000	0.306	<u>0.188</u>	0.060	<u>0.080</u>	<u>0.371</u>	<u>0.415</u>	<u>0.335</u>	0.000	<u>0.191</u>	0.000	0.221		0.320				<u>0.772</u>		<u>0.715</u>
κῶας χρύσεος	8	0	2	0	1	0	1	1	1	0	0	0	8	0	0	0	0	8	6	2	8
relative1 entropy		0.000	<u>0.811</u>	0.000	<u>0.544</u>	0.000	<u>0.544</u>	<u>0.544</u>	<u>0.544</u>	0.000	0.000	0.000	0.000		0.000				<u>0.811</u>		0.000
χείρ δε- ξιτερός	8	0	0	0	0	0	1	1	1	0	0	0	8	0	0	4	2	2	5	3	4
relative entropy		0.000	0.000	0.000	0.000	0.000	<u>0.544</u>	<u>0.544</u>	<u>0.544</u>	0.000	0.000	0.000	0.000		<u>0.750</u>				<u>0.954</u>		<u>1.000</u>
τοῖος μῦθος	8	0	0	0	0	0	0	3	3	0	0	0	7	1	0	1	1	6	8	0	8
relative entropy		0.000	0.000	0.000	0.000	0.000	0.000	<u>0.954</u>	<u>0.954</u>	0.000	0.000	0.000	<u>0.544</u>		<u>0.531</u>				0.000		0.000
δέρος χρύσεος	7	0	0	0	0	0	0	0	0	0	0	0	7	0	0	0	0	7	4	3	3
relative entropy		0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000		0.000				<u>0.985</u>		<u>0.985</u>
άνήρ ἄλλοδα- πός	5	0	0	0	0	0	4	0	0	0	0	0	1	4	0	1	4	0	3	2	3
relative entropy		0.000	0.000	0.000	0.000	0.000	<u>0.722</u>	0.000	0.000	0.000	0.000	0.000	<u>0.722</u>		0.361				<u>0.971</u>		<u>0.971</u>
ναῦς θεός	5	0	0	0	0	0	1	0	0	0	0	0	5	0	0	0	0	5	1	4	1
relative entropy		0.000	0.000	0.000	0.000	0.000	<u>0.722</u>	0.000	0.000	0.000	0.000	0.000	0.000		0.000				<u>0.722</u>		<u>0.722</u>

		tree-syntac tic f.	lexico-syntactic flexibility							morphological flexibility									metrical flexibility		
formula	n. of tokens	relativ e clause	adjecti ve insertio n	noun inserti on	modifiy ng NP/PP	post- modif. rel. cl.	preposi tion	adver b	partic le	negati on	article	comp. / sup.	number		case				word order		separ ation
													S	PI	N/ V	G	D	A	Ad j N	N Adj	
πᾶς ἡμᾶρ	5	0	3	1	0	0	0	1	2	0	0	0	1	4	0	0	0	5	3	2	2
relative entropy		0.000	<u>0.971</u>	<u>0.722</u>	0.000	0.000	0.000	<u>0.722</u>	<u>0.971</u>	0.000	0.000	0.000	<u>0.722</u>		0.000				<u>0.971</u>		<u>0.971</u>
πᾶς ἐταῖρος	5	0	2	2	0	1	0	2	0	0	1	0	0	5	2	0	0	3	4	1	3
relative entropy		0.000	<u>0.971</u>	<u>0.971</u>	0.000	<u>0.722</u>	0.000	<u>0.971</u>	0.000	0.000	<u>0.722</u>	0.000	0.000		<u>0.485</u>				<u>0.722</u>		<u>0.971</u>
ἄλλος ἥρως	4	0	0	0	0	0	1	0	0	0	2	0	0	4	2	1	0	1	3	1	1
relative entropy		0.000	0.000	0.000	0.000	0.000	<u>0.811</u>	0.000	0.000	0.000	<u>1.000</u>	0.000	0.000		<u>0.750</u>				<u>0.811</u>		<u>0.811</u>

V.2.3 Early epic vs Hellenistic epic: averages and frequency

As we did in §IV.3, we can start out by looking at the averages in the table, in order to get something of a general picture before moving on to examining individual formulae. There is a very clear tendency here: while on average N + Adj pairs in Apollonius are still less flexible, and on occasion much less flexible, than pairs in the baseline sub-corpus, they tend to be more flexible than archaic epic formulae for most parameters (8 out of 15).³⁵¹ This result is much more striking given that the frequency of Adj + N pairs in Apollonius is much lower than in the archaic epic sub-corpus, where frequency correlated directly to flexibility. However, if we look at the averages from the perspective of the baseline corpus, we can see that while Apollonian Adj + N ‘formulae’ (recurring pairs) are more flexible than archaic epic formulae of the same type, they still display potentially formulaic behaviour, in that they show restrictions to their flexibility compared to the baseline.

A summary of the results is provided below in Table 14, which contains both the comparison with archaic epic and with the baseline. We will use this table to discuss some areas in which the behaviour of Adj + N pairs had shown interesting tendencies in archaic epic: the presence of frequency effects, the influence of indefinite adjectives, and the distribution of entropy between different types of variation.

Table 14: Summary of Adj + N pairs in Apollonius compared to the baseline and to archaic epic

Pairs (in decreasing order of occurrences)	H _{Apollonius} > H _{baseline}	H _{Apollonius} > H _{formulae}
κῶας + χρύσεος	3	6
χεῖρ + δεξιτερός	3	6
τοῖος + μῦθος	2	4
δέρος + χρύσεος	2	2
ἀνήρ + ἄλλοδαπός	3	4
ναῦς + θοός	1	3
πᾶς + ἥμαρ	6	7
πᾶς + ἑταῖρος	5	8
ἄλλος + ἥρως	2	5

³⁵¹ This count includes parameters in which variation is never attested in Apollonius; due to the very small size of the Apollonian sample, the fact that these kinds of variation never appear is less significant than it would be in the early epic sub-corpus, and is probably due to data sparsity. The fact that these three parameters are embedded relative clause, negation, and comparative, which had low rates of attestation both in the early epic sub-corpus and in the baseline, supports this.

There is a not insignificant problem with this sample, which is that the effect of frequency and the effect of indefinite adjectives, both of which were seen to promote syntactic variation in formulae in §IV.3.2-IV.3.3, cancel each other out: three out of five formulae with the lowest frequency (4 to 5 occurrences) contain the adjectives *πᾶς* or *ἄλλος*, and they also show higher rates of flexibility in more parameters than the rest of the sample. As we have seen in §IV.3.3, the effect of *πᾶς* and *ἄλλος* in promoting variation is strong enough to counteract the reduced rate of variation that comes with lower frequency.

As for the other two pairs with a frequency of 5, *ἀνὴρ* + *ἄλλοδαπός* and *ναῦς* + *θεός*, the former appears about as flexible as the top frequency pairs, when compared to the baseline, while the latter is not very flexible at all (we will return to this result in particular below). In short, when it comes to frequency effects, not very much can be said about Apollonius; this is due in part to the distribution of pairs containing indefinite adjectives, but also to the fact that having a very small sample, as anticipated in §V.2.1, ends up significantly flattening frequency-based differences.

V.2.4 *Early epic vs Hellenistic epic: individual patterns*

As for the effects of indefinite adjectives like *πᾶς* and *ἄλλος*, the results from Apollonius confirm the same tendency that we saw in archaic epic (§IV.3.3): pairs containing these adjectives are much more flexible compared to other pairs, even those with much higher frequency.³⁵² Once again, two of the pairs containing an indefinite adjective, *πᾶς* + *ἐταῖρος* and *ἄλλος* + *ἥρως*, are the only ones in which the article is used at all in the entire Apollonius dataset; *ἄλλος* + *ἥρως* occurs with the article in exactly half of its attestations (which is, admittedly, only 2 out of 4). As the use of the article is without doubt fully

³⁵² The rates of higher flexibility are comparable to the highest ones observed in Table 9 in §IV.3.1 for archaic epic: 5 or 6 parameters in which entropy is higher than the baseline.

developed in everyday language by the time Apollonius is writing,³⁵³ the avoidance of the article in all other Adj + N pairs is most likely perceived as a feature of the epic style by this point. This does not, however, explain why pairs containing indefinite adjectives are still much more flexible than any other type of adjective. A ‘natural’ influence from the development of the article is basically impossible at this point, and would still not explain the behaviour of these pairs in relation to other parameters; we also saw in §IV.3.3 that pairs with *πᾶς* and *ἄλλος* in the baseline sub-corpus were not particularly flexible in any interesting way. In the end, the pattern of behaviour in Apollonius appears to reinforce the alternative explanation suggested in §IV.3.3, that is, that the semantically ‘light’ value of these adjectives may promote flexibility, as pairs containing them are perceived as closer to non-formulaic pairs, rather than formulae to which restrictions in behaviour should apply.

Some additional remarks can be made about *πᾶς* + *ἥμαρ*, which also occurred in the early epic sample, and was actually the least flexible of Adj + N formulae containing a quantifier, with no parameters in which $H_{\text{formula}} > H_{\text{baseline}}$. More specifically, this pair overwhelmingly occurs as *ἥματα πάντα*, with no additions or variation (47 times out of 51, according to my own count). In Apollonius, on the other hand, this pair occurs in four different forms out of five total occurrences, and the only one that appears twice is the highly untraditional *ἥματα δὲ τρία πάντα* (the pair *πᾶς* + *ἥμαρ* is never associated with a number in the archaic epic sub-corpus). Is it a coincidence that one of the very few formulae that Apollonius takes from Homer is used in a completely different pattern? The fact that this happens again for another of the pairs that are attested in Homer says otherwise. If we look at *ναῦς* + *θοός*, we find that this is actually the least flexible of all pairs in the Apollonian sample: it appears four times in its simplest form, *νῆα θοήν*, and only once as *θοήν ἐπὶ νῆα*, which is still much less ambitiously varied than most Homeric examples. Indeed, in the archaic epic sample,

³⁵³ See Manollessou and Horrocks (2007).

ναῦς + θοός is one of the most flexible pairs (as well as the most frequent one – as we know, the two things are related), appearing with a preposition in over half of its 104 occurrences, and reaching entropy rates close to 1 in number, word order, and separation.³⁵⁴

In other words, Apollonius here seems to have taken two early epic formulae, one of which is remarkably rigid even for a formula, while the other is remarkably flexible, and have re-used them in the exact opposite pattern, with the flexible one becoming highly codified and the rigid one acquiring high flexibility. This does seem to be a conscious stylistic choice;³⁵⁵ but it is one that cannot be spotted just by looking at individual occurrences of these pairs in Homer and Apollonius, or by thinking about intertextual connections. This observation is an excellent example of what quantitative analysis can offer to traditional philological and literary comparison.

V.2.5 *Early epic vs Hellenistic epic: types of variation*

It is when looking at individual types of variation and their categories, like we did in §IV.3.4, that the difference between the behaviour of pairs in Apollonius and of formulae in archaic epic becomes especially striking. Table 15 below contains a summary of the behaviour of Adj + N pairs in Apollonius by type of variation, compared to the baseline only. This table should be read side by side with Table 11 in §IV.3.4, which contained the same comparison

³⁵⁴ The third pair that appears both in Homer and Apollonius, πᾶς + ἑταῖρος, appears to behave like other pairs containing πᾶς in both samples.

³⁵⁵ This claim is not out of bounds, even if Apollonius will certainly not have thought of these entities as ‘Homeric formulae:’ there is solid evidence that Hellenistic scholars of Homer such as Zenodotus considered repetition a characteristic of Homer’s language, and intervened on the text based on this principle (see e.g. Nünlist 2009, 312-14; Nünlist 2014, 729). Parry also discusses Apollonius’ attitude to Homeric epithets in M. Parry (1971), 120-4 (see also Nünlist 2009, 299-306). Several other examples of Apollonius’ tendency to vary Homeric expressions can be found in Fantuzzi and Hunter (2004), 266-70, as well as in the pieces discussed in §V.1.1 above.

for the formulaic sample; for ease of comparison, the last column of that table is reproduced here, under ‘percentage in the epic sample.’

Table 15: Behaviour of Adj + N pairs in Apollonius by type of variation

Wulff's category ³⁵⁶	Type of variation	Baseline entropy	Number of pairs where $H_{\text{Apollonius}} > H_{\text{baseline}}$ (out of 9)	Percentage in the epic sample
tree-syntactic flexibility	relative clause	0.057	0	6.25%
lexico-syntactic flexibility	adjective insertion	0.555	3 (33.33%)	31.25%
	noun insertion	0.111	2 (22.22%)	43.75%
	modifying NP/PP	0.476	1 (11.11%)	15.63%
	post-modifying rel. cl.	0.133	1 (11.11%)	25.00%
	preposition	0.613	3 (33.33%)	21.88%
	adverb	0.184	5 (55.56%)	0
morphological flexibility	particle	0.745	2 (22.22%)	12.50%
	negation	0.275	0	0
	article	0.838	1 (11.11%)	3.13%
	comp./sup. adj.	0.094	0	0
	number	0.825	0	15.63%
'metrical' flexibility	case	0.978	0	3.13%
	word order	0.949	4 (44.44%)	6.25%
	separation	0.959	5 (55.56%)	15.63%

In the epic sub-corpus, lexico-syntactic flexibility was the least significant category for the definition of formulaic behaviour, that is, the one where formulaic pairs behaved in the closest way to their baseline counterparts. This observation holds in Apollonius: rates of variation in the lexico-syntactic category range from medium to high (note that due to the small number of formulae under consideration, the difference between, for instance, 11% and 33% is less significant than it appears from percentages).

Morphological and ‘metrical’ flexibility, on the other hand, were the categories in which formulaic behaviour differed the most from the baseline. When it comes to morphology, early epic formulae are much less likely to vary in number and case (and even less to take an article, a negation, or a comparative/superlative). This tendency is upheld in the data from Apollonius: morphological flexibility is practically always lower in the *Argonautica* compared to baseline usage, although it is often higher than in archaic epic ($H_{\text{Apollonius}} >$

³⁵⁶ Again, except for ‘metrical’ flexibility, which was added based on Hainsworth’s work.

H_{formulae} in 33.33% of pairs for number, 44.44% for case). This confirms what we had noted in §V.2.3 when we looked at averages: Adj + N pairs in Apollonius are 'less formulaic' than in archaic epic, that is, they are more flexible in ways in which formulae are normally quite rigid, but they are still 'more formulaic' than in the 5th-century baseline, that is, they are less flexible than 'everyday' Adj + N pairs for the same parameters.

Where things look very different, on the other hand, is in the 'metrical' category of flexibility. Both variation in word order and separation is much more frequent in Apollonius than in archaic epic; indeed, entropy for around half the pairs is even higher than the baseline entropy, which, lest we forget, has values approaching 1, that is, is already extremely high. In other words, Apollonius seems to look for variation in word order and in the presence/absence of separation not just significantly more than his oral predecessors, but also noticeably more than would be expected outside of epic usage.³⁵⁷ It does not seem a stretch to consider this an intentional stylistic choice, given that position in the line and word order are two things that can very easily be manipulated consciously, and that hyperbaton is, after all, a rhetorical device. In other words, like in the case of the specific pairs discussed in §V.2.4, $\pi\tilde{\alpha}\varsigma + \tilde{\eta}\mu\alpha\rho$ and $\nu\alpha\tilde{\upsilon}\varsigma + \theta\omicron\omicron\varsigma$, we have here another situation where Apollonius intentionally departs from traditional formulaic language usage to look for unusual and striking patterns.

V.3 Discussion of results

In his book on post-orality in Homer, Rainer Friedrich argues that formulae defined as Hainsworth suggests, that is, as sequences that frequently occur together to the point that a listener expects the missing half of the pair to appear after one half is used, are too flexible

³⁵⁷ For separation in particular, it is also worth noting that the two pairs that show an entropy of 0, $\kappa\tilde{\omega}\alpha\varsigma + \chi\rho\upsilon\sigma\epsilon\omicron\varsigma$ and $\tau\omicron\tilde{\iota}\omicron\varsigma + \mu\tilde{\upsilon}\theta\omicron\varsigma$, do so because they *always* appear separated, rather than the other way around.

a concept, to the point that it is easy to find similar sequences in non-oral epic poets, including Ovid and Vergil, as well.³⁵⁸ This kind of criticism is particularly damaging for any description of the oral epic formula, as there is no justification for the concept if it is completely uncoupled from its function as a characteristic of oral composition – or, at the very least, of the specific kind of composition that led to the Homeric and Hesiodic tradition. If Hainsworth's pairs were simply a characteristic of hexameter poetry in general (not even just in Greek!), then they would belong to a completely different domain of description from Parryian formulae, and one that is arguably much less useful.

Friedrich's objection, however, fails to consider how these formulae, or high-frequency pairs, behave in archaic epic poetry compared to later texts. In this chapter and the previous one, we have shown quite clearly that the behaviour of recurring Adj + N phrases in early epic is characterised by the avoidance of some kinds of variation, in a way that can be said to be a trait of formulae. This is not a general phenomenon in hexameter poetry, either: the later epic of Apollonius shows different patterns of variation. Not only that, but these patterns – both in the case of individual formulae such as $\pi\alpha\tilde{\alpha}\varsigma + \tilde{\eta}\mu\alpha\rho$ and $\nu\alpha\tilde{\upsilon}\varsigma + \theta\epsilon\omicron\acute{o}\varsigma$, and when it comes to specific types of metrical variation – can be seen as at least a partially conscious attempt on Apollonius' part to distance his style from that of early epic, where limitations on the behaviour of formulae were part of a system that guaranteed a now-obsolete aid in composition.

These results should put to rest the objection that lies, ultimately, behind Friedrich's criticism: that nothing useful can be gained from a flexible definition of the formula that is based on the cognitive connection between frequently occurring pairs. At least for Adj + N combinations, this is clearly not true: by examining the behaviour of these pairs in detail, we were able to reach conclusions that are of interest not only to people studying formulaic

³⁵⁸ Friedrich (2019), 87-92.

variation, but also to scholars interested in Hellenistic poetry. There is, moreover, ample scope for the analysis in this chapter to be extended to other bodies of text, starting from the more obvious (Nonnus and Quintus of Smyrna) and perhaps reaching more far-off choices (what about echoes of epic usage in tragedy, or lyric poetry, or even Herodotus?³⁵⁹).

³⁵⁹ Not only does Herodotus make ample use of epic models (see e.g. Boedeker 2002), but Herodotean pairs are particularly frequent in the baseline sub-corpus used in this chapter and the previous one. The issue of oral style in Herodotus is also very interesting: see e.g. Slings (2002).

Part III

Semantic variation in Verb + Object phrases

VI A Distributional Semantics approach to formulae: Methods and preliminary study

In the previous two chapters, we have shown how the syntactic flexibility of formulae can be mapped using computational methods. This chapter and the following one offer a similar approach to semantic flexibility, that is, variation in the meaning, rather than the structure, of formulaic expressions. Much like syntactic variation, analysing semantic variation in detail contributes to our understanding of questions about the life cycle of a formula.

The starting point for this discussion is, again, the parallel between formulaic behaviour and the behaviour of idioms and other constructions characterised by limited linguistic flexibility.³⁶⁰ In section VI.1, we will see how meaning change and flexibility play an important role in the evolution of language usage, where semantic flexibility correlates positively with productivity in diachrony; similarly, we can expect the semantic openness of a formula to influence the vitality of its usage through time.

In order to apply computational methods to semantic behaviour, as we did to the analysis of the syntactic flexibility of formulae, a quantitative framework under which to discuss meaning is required. I will introduce this framework, Distributional Semantics, in section VI.2. Section VI.3 presents the process through which we built the first computational model which uses Distributional Semantics to assess the scope of linguistic variation in formulaic language.³⁶¹ I will then present a preliminary study that applies this model to the analysis of a very small sample of formulae involving transitive verbs (specifically, formulae of physical and intellectual interaction with an object, such as *holding*, *thinking*, and *seeing*). This study, while not designed to give results that can be generalised to a larger sample,

³⁶⁰ The key references remain Kiparsky (1976); Bozzone (2014a); Antović and Cánovas (2016b), as introduced in §II.3 and III.1.1.

³⁶¹ Material in sections VI.3 and VI.4 is reworked (and, in the latter section, significantly expanded) from Rodda, Probert, and McGillivray (2019). My contributions are clearly marked both in this chapter and in the article.

provides us with a set of factors we can keep in mind in the more extensive study of transitive verb formulae that will be the central focus for chapter VII. A small experiment is not in itself sufficient to test how well various factors can act as predictors of semantic variation; however, starting from a small-scale experiment will allow us to show how a distributional approach to meaning can be applied to epic formulae, rather than just introduced on an abstract level, as well as to identify more easily elements that are of interest and can be tracked in a larger-scale study.

VI.1 Semantic variation in formulae

VI.1.1 Open slots and semantic variation

In the previous two chapters, we have focused on internal variation in formulae: in Adj + N pairs, we have looked at variation that affects the adjective, the noun, or both (through, for instance, expansion and inflection). The formulae that we have considered can be expanded with additional material, but this is not by any means a necessity. In the next two chapters, on the other hand, we will focus on more complex formulaic phrases that include an open slot that needs to be filled with additional material.³⁶² For reasons that will be outlined in §VI.1.2, we will focus on the range of meanings that can be admitted in this peripheral material, as opposed to the internal range of syntactic flexibility that was explored when discussing Adj + N pairs.

The example I am going to start from is formulae involving a transitive verb and a direct object. Below are some examples of one such formula, the ἔχειν ἐν χειρὶ / ἐν χερσὶ ('having in one's hand/hands') formula:

³⁶² The phrases in question involve a transitive verb that has an open slot for a direct object. While this object could in theory be implicit (and recoverable from the context), in practice it almost never is, and therefore transitive verbs form an excellent starting sample for an analysis of formulae with open slots.

στέμματ' ἔχων ἐν χερσὶν ἐκηβόλου Ἀπόλλωνος (*Il.* 1.14)

'having in his hands the wreaths of far-shooting Apollo'

τόξον ἔχων ἐν χειρὶ παλίντονον ἡδὲ φαρέτρην (*Il.* 15.443)

'having in his hand the curved bow and the quiver'

ἔχε δ' ἀστεροπὴν μετὰ χερσὶν (*Il.* 11.184)

'and he had the lightning in his hands'

In the first two cases, the metrical structure of the core formula is the same; it determines the metrical structure of the open slot filler, although that filler is in itself not part of the transitive formula. We can even extend our focus to more complex variations of our original formula, such as in the third example, where the preposition and form of the verb are different, which consequently changes the metrical structure; examples of this kind will be counted as the same formula throughout this chapter. In this and the following chapter, therefore, we will define a formula in terms of meaning and phrase structure (verb phrase formed by a certain verb + a specific adjunct prepositional phrase or noun phrase³⁶³), rather than of metrical shape.³⁶⁴

From the examples above we can get a sense of the kinds of objects that this type of formula can accommodate (continuous underline in the Greek text and translation). These are mostly noun phrases, in this case denoting concrete objects, that can in turn be formulaic: τόξον [...] παλίντονον in *Il.* 15.443 is in itself an instance of the rare Adj + N pair παλίντονα τόξα (occurring twice in the *Iliad* and once in *H.H.* 27 to Artemis), inflected in the singular and with inversion and separation in order to fit the metrical shape of the object of ἔχων ἐν χειρὶ. The VP formula, moreover, can admit different levels of syntactic variation within itself: the verb is usually inflected, and the accompanying material can be

³⁶³ As we will see below, I will accept PPs with different prepositions and without a preposition as part of the same formulaic system in this chapter, as long as the noun remains the same.

³⁶⁴ This is somewhat similar to E. J. Bakker and Fabbriotti (1991)'s definition of formula through nucleus/periphery relations (see §1.3.2), although I am not formally wedded to this definitional approach.

declined as well (in the example above, there is variation between singular ἐν χειρί and plural ἐν χερσὶ). As we will see below, different VP formulae admit different levels of internal (syntactic) flexibility.

In this chapter, we will focus not on the range of forms taken by the fillers of transitive VP formulae, but on the range of meanings that can be accommodated by these same fillers. We will still maintain a quantitative approach to meaning variation, the theoretical basis of which will be explained in §VI.2. First, however, we need to discuss why it is of any interest to look at the range of meanings that can be admitted in a formula.

VI.1.2 Semantic variation and productivity

The reasons why semantic variation is of interest in a description of formulaic behaviour lie, once again, in the parallels between formulae and constructions in everyday language. These have been introduced in detail in §II.3 and §III.2; they rely on the concept that both formulae and constructions are pairs of form and function, and that for both of them apparently arbitrary restrictions to their form apply together with specialised nuances in function (idiomatic meaning in the case of constructions, formulaic resonance/traditional referentiality in the case of formulae, whose meaning is enriched by their role as elements of a traditional language³⁶⁵). In both cases, the degree to which form and function are specialised and restricted can be quantified along a continuum of behaviour, rather than divided more or less arbitrarily into qualitatively impermeable categories.

There is a significant amount of research on what causes constructions to be more or less productive through time, that is, to survive and expand in usage or to fossilise to a limited number of inflexible examples and eventually die out. Constructional productivity refers to

³⁶⁵ On these concepts see very briefly §I.2.3.

the ability of constructions to attract new words to be used within the construction itself. For instance, the extension of the *beat the hell out of* construction to more figurative verbs such as *argue* or even *type* is an example of constructional productivity in English. In Construction Grammar and similar usage-based approaches, productivity is a scalar concept, not a yes-or-no question: a construction can be more or less productive depending on how many new fillers it can attract.³⁶⁶ In a wider sense, productivity can apply to all kinds of linguistic elements, insofar far as they are used in a new (previously unattested) environment;³⁶⁷ so, for instance, the suffix *-gate* is well-known for having become productive after Watergate to refer to any type of scandal (such as Gamergate), and coinages with *-exit* (not just Brexit, but Frexit, Italexit, Irelexit) became popular after the coinage of Grexit in 2012 to refer to the possibility that several countries might leave the EU.³⁶⁸

There are some restrictions to any new coinage: it needs to be meaningful and make sense, and there should not be another construction that is already conventionalised and has the same function as the new coinage.³⁶⁹ This would pre-empt the new construction from spreading, as its function is already performed by an existing one. This mechanism has interesting parallels with the idea of formulaic economy, which discourages the creation of a new formula when a metrically equivalent one with the same meaning exists.³⁷⁰

Some research on the productivity of Homeric formulae has been done by Bozzone (2014a), who however focuses on a completely different class of formulae, with different

³⁶⁶ See e.g. Suttle and Goldberg (2011), p. 1238. For a detailed review of how productivity is defined in the literature cf. Barðdal (2008), 9-54.

³⁶⁷ See e.g. Harris (1954), 150.

³⁶⁸ Cf. OED, s.v. 'Grexit,' and the examples in Liberman (2015). Some of these coinages quite obviously survived longer than others, for non-linguistic reasons. (Ironically, this footnote was last edited on 01/01/2021, the first day after the end of the Brexit transition period.)

³⁶⁹ Suttle and Goldberg (2011), 1239. We will return to this point in the next chapter, when talking about constructional pre-emption and formulaic economy.

³⁷⁰ M. Parry (1971), 275-9. Just as it is well known that formulaic economy admits exceptions (Friedrich 2007), constructional pre-emption is a statistical tendency, not an absolute rule (Suttle and Goldberg 2011, 1240-2).

methodology, and reaches different results from the ones presented here.³⁷¹ Nevertheless, a significant commonality between my research and Bozzone's should be highlighted, in that we both ultimately seek for an explanation of what drives constructional productivity by looking at the peripheral elements of a formula – in her case, subtypes of speech introductions with the same governing verb (such as τὸν δ' ἀπαμβρόμενος προσέφη vs τὸν δ' ἄρ' ὑπόδρα ἰδὼν προσέφη). The latter are not quite open slot fillers on a syntactic level, in that verbs of speaking do not strictly require expanding with, for example, a participle; in the grammar of epic poetry, however, the speech-introduction line can be considered to be a construction with an open slot in need of filling.

There is some evidence from historical English and Icelandic linguistics that semantic flexibility influences the behaviour of constructions in everyday languages. Research has highlighted the driving role of semantic flexibility in promoting the productivity of constructions with open slots: the range of different meanings that a construction can accommodate affects the way it survives and spreads through time. In particular, the key factor seems to be not how widely spread the fillers of a construction are in a metaphorical semantic space, but how well they cover a range of semantic areas: constructions that either only accept a very limited range of fillers or have fillers that are very semantically incoherent (have meanings that are very different from each other, so they have nothing in common that can be generalised and made into a productive rule) tend to die out or fossilise to the same limited range of fillers, while constructions that have a robust range of fillers are more productive and tend to extend to even more fillers. New items are more likely to be introduced to a construction if the semantic space around these items is already densely populated.

³⁷¹ See §II.3.2-II.3.6 for a detailed discussion of her work and its results.

In a study of the argument structure of Icelandic verbs,³⁷² Jóhanna Barðdal has argued that there is an inverse correlation between the type frequency³⁷³ and semantic coherence of productive verb constructions, and that verb constructions for which this correlation does not hold are not productive. In other words, verb constructions are maximally productive when there is either a high number of diverse verbs that fit the construction (high type frequency, low semantic coherence), in which case speakers will feel free to extend the pattern to other, equally disparate verbs, or a small number of verbs that are closely related in meaning (low type frequency, high semantic coherence), in which case speakers will find a reason to extend the construction to verbs that are semantically similar to the original set. In the case of low type frequency and low semantic coherence, speakers will not find a way to extract a meaningful pattern that can be extended to new verbs (as is the case for the high semantic coherence situation), and will not see the pattern often enough to generalise it to any kind of new verb. The high type frequency, high semantic coherence case, on the other hand, is unlikely to occur (as a large enough class of verbs, in Barðdal's view, will naturally be less semantically coherent than a small group).

³⁷² Barðdal (2008).

³⁷³ In this case, type frequency is the number of different verbs a construction can occur with (as opposed to token frequency, which would be how often the construction occurs overall, including repetitions with the same verb). See §VII.1.4 for a further discussion and application of these terms.

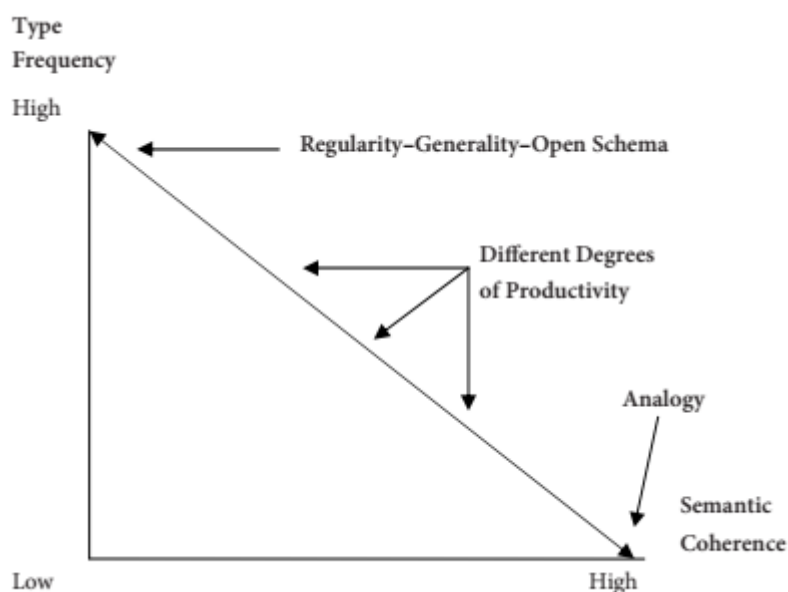


Figure 2: The productivity cline according to Barðdal (Barðdal 2008, fig. 7.2 p. 172). Constructions are more productive the closer to the line they fall in terms of type frequency and semantic coherence.

Barðdal's study is the first to highlight the role of meaning in driving the productivity of VP constructions. The main upshots of her argument are that productivity is a gradient, and that both productivity based on an open schema (low semantic coherence) and productivity based on high semantic coherence (analogy) are two sides of the same continuum.³⁷⁴

Florent Perek³⁷⁵ has taken a step forward in offering a quantitative assessment of semantic similarity that is based on Distributional Semantics, the same approach that will be used in this chapter and the next (see §VI.2 for more details). Perek uses a quantitative framework to study the evolution of verb fillers in the English V *the hell out of* NP construction (such as *my cat scared the hell out of me*) across eight decades, from the 1930s to the 2000s, in a historical corpus (the COHA, Corpus Of Historical American English). Perek's analysis of the way this construction accrues new fillers shows that if the semantic domain around a

³⁷⁴ A series of psycholinguistic experiments described in Suttle and Goldberg (2011) expanded the notion of semantic coherence introduced in Barðdal's work, showing how the determining factor in deciding whether a construction can be extended to new fillers is not simply semantic similarity, but rather the extent to which the attested instances cover the semantic area to which the new coinage would belong. This makes the semantic area 'acceptable' to speakers, and leads to the introduction of new coinages. This study, however, does not adopt a quantitative or gradient measure of semantic similarity.

³⁷⁵ Perek (2016).

potential new filler is already densely populated (that is, if many existing fillers are already attested in the same semantic domain) the probability of the new filler being coerced into the construction is higher. This is a quantitative, statistically robust confirmation of the hypothesis that coverage is a driving force in constructional productivity.

If semantic coverage influences the behaviour of constructions in everyday language, and if there is a parallel between epic formulae and constructions, the obvious question arises whether we should expect semantic variation to also have a role in influencing the behaviour of formulae. This chapter and the following are dedicated to addressing this question, by assessing the range of semantic variation in a sample of epic formulae and looking for other factors (syntactic variation and diachronic development in particular) that may correlate with this range. Although based on parallels with everyday languages we might expect semantic variation to have something to do with diachronic development, this is not necessarily the only factor at play.

VI.2 Quantifying meaning: Distributional Semantics

The central idea of this thesis' approach to the behaviour of formulae (in constructional terms) is that it is not necessary to work within a strict binary distinction between formulaic and non-formulaic language; rather, expressions can be described as ranging along a continuum of most to least rigidly defined, most to least 'formulaic' behaviour. A quantitative approach is perfect to describe this continuum, as it allows us to assess the full extent of the available data without having to rely on preconceived and largely arbitrary qualitative distinctions (such as formulaic vs non-formulaic). Quantitative analysis, therefore, allows us to fully describe formulaic behaviour.

This is a rather unproblematic statement, in principle, when it comes to syntactic flexibility, where the degree to which certain phenomena are present in a large sample can be

straightforwardly quantified, as we did in Part II. When looking at the ways in which semantic flexibility influences the behaviour of formulae, however, it becomes necessary to find a way to model meaning (semantic flexibility) as a factor in a quantitative framework, so that quantitative data can be compared to other quantitative data, and predictions can be made and tested. The theoretical framework that allows us to model meaning in precisely this quantitative way is called Distributional Semantics.³⁷⁶

VI.2.1 *The distributional hypothesis*

Distributional Semantics (DS) adopts an approach to the definition of meaning that is significantly different from that of formal semantics and other philosophical approaches which study the meaning of words either in relation to their referents or as a set of logical properties. It is also, crucially, different from structural (linguistic) semantics, which focuses on the semantic features of lexical items and how these features combine to create meaningful sentences.³⁷⁷ It comes, on the other hand, much closer to a traditional lexicographic (operative) definition of meaning, which seeks to define what a word means by offering examples of the contexts in which it may occur, as well as by associating it to its synonyms and/or cognates.

In Distributional Semantics, the meaning of a word is defined as a function of its collocates in a corpus: words that share a linguistic context are also related in meaning. In one of its earliest formulations, from the days in which the focus was mostly on developing strategies

³⁷⁶ In the next sections we will discuss the aspects of Distributional Semantics that are most relevant to the current project, which inevitably means leaving many current debates to the side. For a more up-to-date review of the field, see Fabre and Lenci (2015).

³⁷⁷ For these definitions cf. Crystal (2008), s.v. 'semantics'.

to retrieve information automatically from digital documents, the concept is introduced like this:³⁷⁸

Distributional semantics is predicated on the assumption that linguistic units with certain semantic similarities also share certain similarities in the relevant environments. If therefore relevant environments can be previously specified, it may be possible to group automatically all those linguistic units which occur in similarly definable environments, and it is assumed that these automatically produced groupings will be of semantic interest.

Note how formulations of the principles of Distributional Semantics emphasise *comparison* between word meanings, rather than the definition of individual meanings.³⁷⁹ The crucial advantage of a distributional approach is that this comparison happens in quantitative terms.

In order to understand the unique contributions of Distributional Semantics, some historical context on its development is helpful. The theoretical reasoning behind the approach is outlined by Zellig Harris in an article that is generally recognised as the first introduction of the ‘distributional hypothesis.’ In its strongest formulation, this linguistic approach states that “each language can be described in terms of a distributional structure, i.e. in terms of the occurrence of parts (ultimately sounds) relative to other parts, and [...] this description is complete without intrusion of other features such as history or meaning.”³⁸⁰ In other words, the way the elements of a language are arranged (all the way

³⁷⁸ Garvin (1962), 388. As the title of this article reveals, Distributional Semantics is a field that developed with a very strong connection to computational linguistics and the automation of Natural Language Processing.

³⁷⁹ Cf. also Harris (1954), 157.

³⁸⁰ Harris (1954), 146.

down to individual phonemes), which determines in what configuration they occur with each other, is a sufficient means of description of linguistic behaviour.

In this chapter we will limit ourselves to a much less ambitious use of the distributional hypothesis, the one that is limited to semantics (and widely applied in contemporary linguistics³⁸¹). In this approach, the meaning of an element (word) is defined as a function of the elements (words) that co-occur with it in an environment of pre-defined size. A much pithier and better-known formulation of this principle is the maxim by John Firth that “You shall know a word by the company it keeps!”³⁸² – that is, the environment in which a word occurs is the most relevant factor for its linguistic processing. The *distribution* of an element is defined as the sum of all its environments. In turn, the environment of each element is a list of the elements with which our target element co-occurs.³⁸³

The key definition of similarity in Distributional Semantics is based on substitutability:³⁸⁴ linguistic elements are substitutable for one another when they occur in the same or very similar contexts, however defined, and elements that are substitutable are considered semantically similar. Note that this does not straightforwardly map to a common-sense definition of synonymy: while synonyms can generally be substituted in the same context, considerations of style may intervene to exclude or favour one or the other. Moreover, the set of words that can be substituted in the same contexts often includes antonyms (*I have had many good teachers in my life vs I have had many bad teachers in my life*, and a vast set of similar possible utterances³⁸⁵), hypernyms/hyponyms, and other relations. In general, when

³⁸¹ See once again Fabre and Lenci (2015) for a review of the field.

³⁸² Firth (1957), 11.

³⁸³ Harris (1954), 146, who defines the environment as an ordered list. When talking about words and their meanings, working with an ordered list would mean taking into account the distance between each target word and its contexts in all computations; however, word order is commonly ignored by modern Distributional Semantics approaches, and the only limitation to the environment is its size (the window of collocations that is considered relevant to the meaning of each word).

³⁸⁴ Harris (1954), 159.

³⁸⁵ Note that *good* and *bad*, like many other pairs of antonyms, are not substitutable in specialised or idiomatic contexts, so there is still a difference in their distribution: e.g. groceries will *go bad*, but

discussing distributional results, it is not easy to keep synonymy and semantic relatedness apart;³⁸⁶ for this reason, we generally talk about (semantic) similarity, as an umbrella term that embraces all kinds of semantic relatedness.

In addition to this, the question of paradigmatic vs syntagmatic similarity arises.³⁸⁷ The examples above all involved words that are similar in the sense of being substitutable in the same context (paradigmatic similarity, which has to do with choice of one form among a set of available ones); however, distributional analyses can also capture another type of relation, the one that connects words that tend to occur together. The latter is called syntagmatic similarity, in that it has to do with the concatenation of words in linguistic production. While words like *good* and *bad* in the example above are paradigmatically similar, two words like *fixed-term* and *lecturer*, which form a (highly specialised) collocation, are syntagmatically similar despite not being at all substitutable.³⁸⁸

Both paradigmatically and syntagmatically similar words can appear related in the results of distributional studies, especially as the size of the context window varies. Larger windows lead to capturing syntagmatic similarities more easily, as words that occur close to each other will also occur in overlapping contexts, with only minor differences; the larger the context window, the less impact these minor differences will have, and the more similar words that occur close together will look. Since we are primarily interested in paradigmatic similarities, our work will only deal with fairly small context windows; it is, however, important to note that in corpus studies, the context for each word can vary significantly in size, ranging from a whole document to the immediate sentence context to specific

not **go good*, and a *bad rap* is not the opposite of a *good rap*, which is, at best, a positive example of a music genre.

³⁸⁶ See e.g. Curran (2004). For a dataset that was designed to encode specifically these types of relations see Baroni and Lenci (2011).

³⁸⁷ See e.g. Levy and Goldberg (2014). On the structuralist terminology used here, see e.g. Crystal (2008), s.vv. 'paradigmatic' and 'syntagmatic'.

³⁸⁸ Without defining this concept, we have employed it extensively in §III-IV when talking about the 'bond of mutual expectation' between parts of a recurring collocation or formula.

syntactic dependencies.³⁸⁹ This is another way in which the concept of ‘similarity’ is complicated in Distributional Semantics: different study setups and different approaches lead to capturing different kinds of relationship between words, each of which can be useful for different purposes.³⁹⁰

A distributional definition of meaning is especially useful when studying a dead language or an early stage of the language, as long as context is available, even where speakers’ judgements are not.³⁹¹ Computational corpus studies (and Natural Language Processing) are one of the two main fields in current Distributional Semantics, the other being language acquisition.³⁹² These interests came to Distributional Semantics through the developments of automated document search on the one hand, and of cognitive linguistics on the other: Distributional Semantics has been used as an approach to the questions of, respectively, how a computer can identify the documents that are most relevant to a query beyond simply looking up the words of the query,³⁹³ and how children learn the meaning of words without access to an explicit definition.³⁹⁴ What is common to all these applications is that

³⁸⁹ The latter is particularly interesting as it is not a linear context. For some examples of Distributional Semantics work that takes syntactic dependencies into account see Levy and Goldberg (2014) and Baroni, Dinu, and Kruszewski (2014), who also compare this approach to linear dependencies.

³⁹⁰ Capturing syntagmatic similarities is, for instance, quite useful in optimising document or web page searches: when searching for a specific topic, say ‘semantics,’ a user will probably also be interested in results which do not contain the query word but also words that frequently occur together with it, say ‘word meaning’ and ‘synonymy.’

³⁹¹ Or are only available through the proxy of ancient lexicography: see §VI.3.2 below. For an additional example of how Distributional Semantics can be applied to historical linguistics (in this case, the grammaticalization of a construction), see Jensen (2013).

³⁹² The connection between Distributional Semantics and corpus studies is mentioned in Harris (1954), 160-1, but otherwise Harris’ approach is neither corpus-oriented nor particularly interested in language acquisition.

³⁹³ So, for instance, while researching the development of caterpillars into butterflies, one might not want articles that use ‘coming out of its cocoon’ as a metaphor to appear among the most relevant results. See, again, Garvin (1962).

³⁹⁴ See e.g. Vinson, Andrews, and Vigliocco (2014), who however ultimately take a critical perspective on distributional models of language acquisition.

Distributional Semantics describes meaning as a variable that can be measured quantitatively;³⁹⁵ we will explain how in the next two sections.

VI.2.2 Counting co-occurrences

The first step towards building a Distributional Semantics model is fairly straightforward, given all that has been discussed above: the co-occurrence patterns for each word whose meaning we want to evaluate (henceforth ‘target word’) need to be extracted from the corpus.³⁹⁶ To do so, we also need to decide how large a context window we are interested in. In the example below, with three target words, the context window is two words on each side of the target (not including the target itself); more examples of window sizes are given in §VI.3.

... he held the beautiful **wreath** safely in his hands ...

... a sacred laurel **wreath** with beautiful ribbons tied around it...

... he held the curved **bow** safely in his hands ...

... a troop of Syrians with curved **bows** rode past them ...

... the strong belief that **philosophy** is the comfort of the soul ...

... Socrates was killed for teaching **philosophy** to beautiful youths ...

Before we count these co-occurrences, a decision needs to be made about what to do with so-called ‘stop-words,’ that is, words that perform a primarily syntactic function, co-occurring with other words independently of their semantic properties; they are therefore considered irrelevant to semantic modelling. Stop-words are mainly function words like

³⁹⁵ Harris himself introduces the probabilistic nature of distributional approaches in Harris (1954), 146-8.

³⁹⁶ We may want to extract co-occurrences for all the words in the corpus or a more restricted sample, as we will see below.

articles, prepositions and conjunctions, but stop-word lists normally also include common verbs like *be* and *have* (and, in English, auxiliaries and modals like *will* and *may*) as well as pronouns and possessive adjectives. Past studies of large English-language corpora have argued that stop-word filtering does not significantly affect the results when it comes to meaningful (non-stop) words, but it does reduce the number of words that need to be assessed, which makes processing easier and lighter.³⁹⁷

For ease of reading, the English examples above were given in their non-lemmatised form. A Distributional Semantics Model can be based on either a lemmatised text or an unlemmatised one; since, however, all of our experiments are conducted on a lemmatised corpus, from here on we will be dealing with lemmatised forms.³⁹⁸

Once we have decided whether to discard stop-words, we can tabulate the co-occurrences between our target words and their contexts. The example sentences above, after stop-word filtering, give us the following table of co-occurrences.

Table 16: Example of a cooccurrence table

		Context words							
		beautiful	safely	sacred	laurel	curved	ride	belief	teach
Target words	wreath	2	1	1	1	0	0	0	0
	bow	0	1	0	0	2	1	0	0
	philosophy	1	0	0	0	0	0	1	1

Even from this very small table, we can already see some patterns emerge: there is some similarity between *wreath* and *bow*, due to the fact that both are objects that can be held *safely*, and some between *wreath* and *philosophy*, which are connected by the word *beautiful*. These similarities are pretty marginal: bows and wreaths are not particularly

³⁹⁷ Bullinaria and Levy (2012).

³⁹⁸ On lemmatisation, a process that associates every word in a text or corpus to the dictionary lemma it comes from, see §IV.1.1.

similar, and neither are wreaths and philosophy. On the other hand, *bow* and *philosophy* are dissimilar in all ways.³⁹⁹

If we scale this example up to the size of a whole corpus, we will end up with an extremely large table of co-occurrences between targets and contexts. Often, as is the case for our study on ancient Greek as well, the sets of targets and contexts coincide, and are simply all the words in the corpus (excluding stop-words), or all the words that occur above a certain frequency.⁴⁰⁰ This oversized table is, in mathematical terms, a matrix, and will from now on be referred to as such.

VI.2.3 *From matrix to vector space*

Every row of our matrix now contains a series of numbers that correspond to the co-occurrences of the word that indexes the row with all the words that index the respective columns. This series of numbers is, in mathematical terms, a vector, and the n columns (contexts) are the dimensions of said vector. Thinking about the row indexed by each target word as the vector associated to this target word allows us to move from a numerical representation of the co-occurrences between words (a huge and unwieldy matrix) to a geometrical, that is spatial, one: each row of the matrix is an n -dimensional vector, and all the vectors together create a vector space model. If we accept the basic hypothesis that the pattern of co-occurrences for each word defines its meaning, then this meaning will be represented by the position of the word itself in the vector space model – which words are

³⁹⁹ It is worth saying here that the example sentences were hand-crafted to give exactly this pattern, and that this is not a reflection of any realistic corpus data. In the same way, the issue of *bow* being polysemous in (written standard) English is not considered here.

⁴⁰⁰ Rare words are problematic in that their co-occurrence patterns are extremely sparse.

close to each other and which are further away, for instance, indicates a relation between them.⁴⁰¹

The concept of a vector space model and how it can be used as a tool for semantic analysis can be more intuitively represented through a two-dimensional example. For this, let us take the first two columns of the example table of co-occurrences that we built in §VI.2.2.

Table 17: A minimal example of co-occurrence coordinates for a vector space model

	beautiful	safely
wreath	2	1
bow	0	1
philosophy	1	0

At this stage, we are not interested in what the contexts are any more: we just care about the fact that the word *wreath* is indexed by a vector $\mathbf{v}_w = (2, 1)$, *bow* by a vector $\mathbf{v}_b = (0, 1)$, and *philosophy* by a vector $\mathbf{v}_p = (1, 0)$. Each of these vectors has two dimensions,

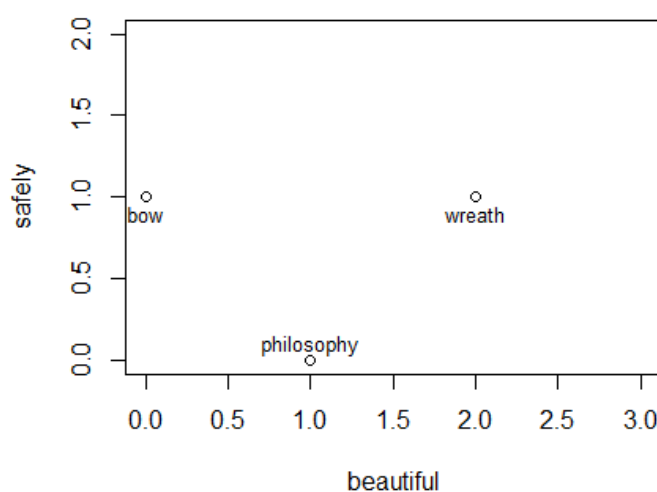


Figure 3: Scatterplot representing the relative positions of example target words in a two-dimensional space

corresponding to our two contexts. Now we can represent these vectors as coordinates of

⁴⁰¹ An extreme and very well-known example of this is the work by Mikolov, Sutskever, et al. (2013), where the accuracy of the representation was so high that, for instance, names of countries and their respective capital cities were always connected by vectors of the same size and direction (that is, by starting from the name of any country and metaphorically moving a certain distance in a certain direction within the semantic space you would end up at the multidimensional coordinates for the name of its capital city).

points in a Euclidean space, where the two dimensions are represented on the two axes x and y (Figure 3).

Now, we can measure the distance between these points (each of which is an underlying vector) in our Euclidean space. There are several ways to do so; Distributional Semantics analyses generally measure the cosine distance, that is, the cosine of the angle between each pair of vectors.⁴⁰² In our example, however, we will simply use the Euclidean distance between our points, as this is much easier to compute by hand.⁴⁰³ The Euclidean distance between two points p and q in two dimensions is $d(\mathbf{p}, \mathbf{q}) = \sqrt{(q_x - p_x)^2 + (q_y - p_y)^2}$, or in other words the hypotenuse of a triangle with points p and q as the vertices of the acute angles. The distances between our three words are tabulated below:

Table 18: Euclidean distances between example words

	wreath	bow
wreath	-	
bow	2	-
philosophy	1.41 $[\sqrt{2}]$	1.41 $[\sqrt{2}]$

In this example, *wreath* and *bow* are maximally distant, while *philosophy* is more similar to each of them. This is different from the assessment of similarity we had made by eyeballing our table at the end of §VI.2.2; adding more contexts (that is, dimensions) to the space might change these distances to approximate our initial assessment better.⁴⁰⁴

This is all well and good for a two-dimensional example, but the average vector space model for Distributional Semantics is built on millions or billions of word pairs. This creates some issues. First, most of these pairs will not co-occur anywhere in the corpus, especially as the

⁴⁰² Cosine distance consistently outperforms other distance measures when assessed against independent benchmarks such as language test data or hand-crafted databases: see Bullinaria and Levy (2007) and §VI.3.3 below.

⁴⁰³ The Euclidean distance is less accurate in representing semantic similarity, and therefore will not be used when we get to analysing corpus data. Cosine distances will be used instead.

⁴⁰⁴ In this case, the discrepancy is partially due to the fact that we did not initially consider that in our examples, *wreath* is much more strongly associated to *beautiful* than *philosophy* (2 co-occurrences vs 1, which makes a huge difference with values ranging 0:2); this negatively influences the similarity between *wreath* and *philosophy*.

context window becomes smaller; this means that most of the cells in the matrix will be empty (they will contain a zero), and that the vectors in the space will consequently be very sparse. Secondly, raw vectors do not take into account the difference in frequency among both target and context words, which is very problematic, given the unbalanced frequency distribution that is characteristic of linguistic data.⁴⁰⁵

To address this second issue, the raw scores are weighted in ways that give a better approximation of the actual importance of the association score between target and contexts, given their frequency and distribution in the corpus. The most common of these weighting schemes, and the one that is used in this thesis, is called Positive Point-wise Mutual Information, or PPMI.⁴⁰⁶ To address the issue of data sparsity, on the other hand, the number of dimensions of the raw vectors is reduced through dimensionality reduction techniques, the most common of which is called Singular Value Decomposition (SVD).⁴⁰⁷ Dimensionality reduction is widely recognised to improve the accuracy of vector space models by reducing noise and highlighting the contribution of the most significant dimensions.⁴⁰⁸ Optimal levels of dimensionality reduction for small and medium-sized corpora are around a few hundreds of dimensions.⁴⁰⁹

⁴⁰⁵ The frequency distribution of words in a corpus follows Zipf's Law, that is, the frequency of any word is (approximately) inversely proportional to its rank in a table ordered from most to least frequent: the first most frequent word will occur approximately twice as frequently as the second most frequent one, three times as frequently as the third, and so on. On this, see Zipf (1949) and Powers (1998).

⁴⁰⁶ Evert (2008).

⁴⁰⁷ Furnas et al. (1998). The kind of simplification that is involved in dimensionality reduction is akin to the one involved in the compression of digital images or files: redundant information is codified in a more efficient way, which makes it occupy less space.

⁴⁰⁸ Landauer and Dumais (1997), who also highlight the interesting parallels with cognitive processing of information by our own human brains.

⁴⁰⁹ The issue is different for larger corpora: Bullinaria and Levy (2012) show that models with several thousands of dimensions perform better on large, unprocessed corpora. This is most likely due to the fact that the size of the corpus counteracts the amount of noise present in it.

VI.2.4 But how good is our model? The validation of DSMs

The result of all these steps (corpus processing and stop-word filtering, co-occurrence count, context weighting, dimensionality reduction) is a relatively small matrix of dense vectors, which represents a multi-dimensional vector space – the Distributional Space Model, or DSM.⁴¹⁰ Since all the steps are designed to not warp the similarity between word-vectors too much, according to the principles we outlined above, the distance between the vectors in this DSM continues to be a relatively faithful representation of the semantic similarity between the corresponding words. This representation is measurable in quantitative terms, which was our goal from the start.

The process outlined in §VI.2.2-VI.2.3, however, involves a lot of choices at different stages: the size of the context window, the kind of dimensionality reduction and weighting that is applied, the measure of semantic similarity, even the choice of whether to eliminate stop-words or not. Is there a ‘best practice’ approach to all of these questions, a one-size-fits-all combination of parameters that gives the most accurate representation of semantic relationships between words in a corpus? The answer to this question is ultimately negative, for two reasons. Firstly, corpora are different from each other in many ways: in terms of their length (is this information scraped from tens of billions of web pages, or a small corpus of ancient texts which cannot be expanded?), of their content (tweets? Books from Google Books? Journal articles? Literary texts? Do these texts belong to one or more specialised genres, and how does this influence the results?), and of the linguistic information they contain (is the corpus richly processed and enriched, like Diorisis, which contain PoS tags and morphological parsing, or is it just the bare bones of a series of texts with no extra annotation?). Secondly, as we have seen in §VI.2.2, it is not at all easy to define what the ‘best’ results would be: do we want a model that captures paradigmatic or

⁴¹⁰ This is how a vector space model used in *Distributional Semantics* is generally called, and how it will be called from this point onwards.

syntagmatic similarity? Or do we want a model that can capture some specific semantic relationships, such as ‘x is the antonym of y’ or even ‘X-ville is the capital city of Y-land?’

It is for all these reasons that part of the process of building a DSM is validating its accuracy for the specific purposes for which it is being built.⁴¹¹ This is generally done by comparing the performance of the DSM to an independent benchmark database, which in turn contains some kind of information about synonymy or semantic relationships. This can be a dictionary, a hand-crafted database, or a set of judgements about which words are closer in meaning collected from native speakers. A particularly popular database is the TOEFL performance data, that is, the responses from test-takers in the part of the TOEFL language test that requires candidates to select the closest synonym of a target word from a set of potential words. The neatness of the TOEFL test benchmark lies in the fact that, since TOEFL scores for human candidates are available, it allows for a comparison between DSM performance and the performance of human second-language learners on a scored synonym identification task – that is, it allows researchers to find out if their language model is better or worse than the average non-English speaker at recognising related English words.

All of this can create some issues for DSMs built on a dead language, where judgements from native speakers are hard to come by. We will discuss these issues at length in §VI.3.2-VI.3.3, which are dedicated to the validation of a series of DSMs built on the Diorisis corpus of ancient Greek. As we will see below, the approach we took is to build several DSMs that use slightly different parameters (different context window, different dimensionality reduction), and see which of those obtained the best results in a comparison task with some baseline resources. The sections below will also summarise the ways in which the results

⁴¹¹ On the validation of semantics models and the questions addressed in this section, see Landauer and Dumais (1997); Curran (2004); Bullinaria and Levy (2007; 2012).

of this comparison task differ from results obtained in previous studies of large English-language corpora, especially the two landmark articles by Bullinaria and Levy (2007; 2012).

VI.2.5 Do you want to build a model? Count-based approaches vs predictive approaches

While the definition of the Distributional Space Model (DSM) as a vector space is the same no matter what approach is taken to its construction, the operative steps to the construction of a DSM described in the previous two sections are in fact only one of the possible ways in which a DSM can be built from a linguistic corpus. An important alternative, and one that has become more and more widespread in the very recent past, is that of so-called word embeddings or neural language models (predictive models).⁴¹² These models address the problem of weighting the contexts in a word vector in a different way from traditional DSMs. We have seen how in a traditional count-based DSM vectors for each target word are created from co-occurrence counts, and then the contexts are weighted through a series of mathematical operations that are the same for all vectors. Word embeddings, on the other hand, take both the unweighted word vectors and the sentences in which the target words appear in the corpus, and then apply a neural network to set the weights in each vector in such a way as to maximise the probability that the contexts which are actually observed in the corpus will occur, while lowering the probability of the ones that are not observed. This optimisation, and the specific weight measures that follow from it, will be individual to each vector. To achieve this weighting, word embedding models must be exposed to a training corpus, on which they will adjust their parameters, ‘learning’ the correct weights; they can then be run on the full corpus, with the ‘learned’ weights remaining unchanged. Predictive models are computationally heavier than

⁴¹² The two most widely used neural language models at the time of writing are probably word2vec (Mikolov, Chen, et al. 2013) and GloVe (Pennington, Socher, and Manning 2014).

traditional count-based models; the former have, however, been shown to outperform the latter in almost all tasks when trained on a sufficiently large corpus.⁴¹³

Even without going into more details about word embeddings, from the brief explanation above it should be clear that to make them work on an optimal level it is essential to have a training corpus of a sufficient size and balanced in a way that is representative of the larger corpus in both vocabulary and word association: so, for instance, a corpus of specialised technical language will be entirely useless when training a model that will then need to be used on a general-knowledge corpus, but will be perfect if the full corpus is only made of technical texts from the same domain. This exact requirement represents a thus-far insurmountable obstacle for the application of word embeddings to ancient Greek: while sampling a corpus of ancient texts to achieve a reasonable level of representativity would be very difficult but potentially possible,⁴¹⁴ the size of the full corpus is already very small compared to modern language corpora, which routinely total billions of words.⁴¹⁵ This would lower the performance of a neural language model to previously-untested levels. For this reason, all the semantic analysis in this thesis has been done using a count-based model, details for which will be given in §VI.3.1.

VI.2.6 Some applications of DSMs to ancient Greek

After having introduced the theory behind Distributional Space Models and before moving on to building a new one for ancient Greek, it is worth briefly summarising the (limited amount of) past work that has been done by applying Distributional Semantics approaches

⁴¹³ Baroni, Dinu, and Kruszewski (2014), from whom I have also borrowed the terminology of ‘count-based’ vs ‘predictive’ models.

⁴¹⁴ Do not forget that date, genre, dialect, metrical form would all be relevant parameters for this sampling.

⁴¹⁵ Cf. e.g. the Google News corpus used in Mikolov, Chen, et al. (2013) (6 billion words), or the five corpora used in Pennington, Socher, and Manning (2014), ranging from 1 to 42 billion(!) tokens. Diorisis, which is significantly larger than other processed corpora of ancient Greek, contains 10 million words.

to ancient Greek corpora. The size and availability of the corpora themselves has thus far posed a significant challenge to Distributional Semantics work on ancient Greek, although this can now be partially overcome thanks to the availability of large(r) open-source processed corpora such as Diorisis.

A DSM built on the TLG database has been used to show the effects of the spread of Christianity on the meaning of Greek words from the first century AD onwards.⁴¹⁶ The results of this experiment, published in Rodda, Senaldi, and Lenci (2017), show that the semantic shift of specific words connected to Christian theology, such as παραβολή, ἄγγελος, or even πατήρ and υἱός, can be traced in the distributional model. Another effect discovered by this study is the influence of technical language, according to which several terms tend to specialise towards a technical usage; rather than reflecting an actual shift in the language, this is more likely to reflect the different balance of the corpus in the AD centuries (where technical writings are predominant). The latter result highlights once again how tricky corpus balancing is when working on ancient Greek materials.

Another feature of the 2017 article was the visualisation of the DSM in two dimensions through t-SNE,⁴¹⁷ a technique that allows the user to compress several hundreds of dimensions without excessively distorting the space. The result is a graphic representation that can easily be interpreted by the naked eye (although, of course, this kind of intuitive analysis should always be taken with caution and backed up with robust quantitative data). An example from the original article is reproduced below.

⁴¹⁶ For examples of studies of diachronic meaning shifts in other languages, see Hilpert (2011); Hamilton, Leskovec, and Jurafsky (2016).

⁴¹⁷ Van der Maaten and Hinton (2008).



Other successful applications of DSMs to ancient Greek include the study of the semantics of verbs of motion and of lexical polysemy.⁴¹⁸ The latter study is particularly notable for not only comparing the results obtained through the model to the independent assessment of researchers, but for having the researchers record on what basis they disambiguated between different possible meanings of a word. This offers some insight into how the performance of a computational model can be compared to (self-reported) human reasoning.⁴¹⁹

VI.3 Creating the vector space model

The ultimate goal of this long account of vector space models is, of course, to devise our own application of a DSM to the analysis of formulaic behaviour. In order to do this,

⁴¹⁸ Respectively, Grewcock (2018) and McGillivray et al. (2019).

⁴¹⁹ Despite my best efforts, this whole section VI.2 remains terribly long and hard going. Here is a video of a kitten: https://github.com/MartinaAstridRodda/dphil-thesis/blob/main/rolli_with_pepper.mp4.

however, we first need to build the DSM itself, a task that poses plenty of challenges of its own, especially when it comes to evaluating the performance of the model. This section describes the procedure through which the DSM that will be used in the rest of this study was built and tested, in collaboration with Barbara McGillivray and Philomen Probert.⁴²⁰

VI.3.1 *Corpus and processing*

Once again, all the material for the present study (and for the next chapter, which will draw on the same DSM) was gathered from the Diorisis Ancient Greek Corpus.⁴²¹ The entirety of the corpus was used to build the vector space model, without chronological filtering. This decision was made both to compensate for the relatively small size of the corpus itself in comparison with modern language corpora, and on account of the nature of the ‘gold standard’ datasets that will be used to test the accuracy of the DSM itself and will be introduced in section VI.3.2. While the level of diachronic depth varies between the sources for these ‘gold standards,’ none of them provides a synchronic snapshot for any specific period, and there is therefore no reason to try to achieve this in the corpus.

Data was extracted through a Python script I wrote for this purpose.⁴²² The script extracts the collocates for each word in the corpus, that is, the words that occur together with the target word, within a window of co-occurrence defined by the user; in order to test the accuracy of different DSMs, we used windows of 1, 5, and 10 words on both sides of the target word (not including the target itself), and tested the accuracy of each.⁴²³ Context windows are sensitive to sentence boundaries: we only included words from the same

⁴²⁰ Full details are provided in Rodda, Probert, and McGillivray (2019).

⁴²¹ As described in Vatri and McGillivray (2018b). A full description of the corpus has already been given in §IV.1.1.

⁴²² The script is available in the usual thesis repository at https://github.com/MartinaAstridRodda/dphil-thesis/blob/main/corpusprocess_2.0.py.

⁴²³ The results are described in §VI.3.3. For the reasoning behind testing the accuracy of DSMs, see §VI.2.4.

sentence as the target word. Sentences are defined according to ancient Greek punctuation rules: a full stop (.), question mark (;), or high point (·) can all act as sentence boundaries. Sentence boundaries are already encoded in the Diorisis corpus; while they are, of course, the fruit of the editorial labour for each of the editions included in the Perseus Library source, they are fully independent of our judgement.

Below is an example of context windows of different sizes and how they interact with sentence boundaries.⁴²⁴

Window of 1 centred on προσμειδῖα:

έώρακας, ὧ Ἄπολλον, τὸ τῆς Μαΐας βρέφος τὸ ἄρτι τεχθέν; ὥς καλόν τέ
έστι καὶ προσμειδῖα πᾶσι καὶ δηλοῖ ἤδη μέγα τι ἀγαθὸν ἀποβησόμενον.

Window of 5 centred on προσμειδῖα:

έώρακας, ὧ Ἄπολλον, τὸ τῆς Μαΐας βρέφος τὸ ἄρτι τεχθέν; ὥς καλόν τέ
έστι καὶ προσμειδῖα πᾶσι καὶ δηλοῖ ἤδη μέγα τι ἀγαθὸν ἀποβησόμενον.

Window of 10 centred on προσμειδῖα (note the sentence boundaries):

έώρακας, ὧ Ἄπολλον, τὸ τῆς Μαΐας βρέφος τὸ ἄρτι τεχθέν; ὥς καλόν τέ
έστι καὶ προσμειδῖα πᾶσι καὶ δηλοῖ ἤδη μέγα τι ἀγαθὸν ἀποβησόμενον.

A frequency threshold, applied over the whole corpus, can also be set by the user; we built our semantic spaces respectively on words that occur at least 100 times in the corpus, at least 50 times, at least 20 times, and finally with no frequency threshold at all (including everything down to hapax legomena). All spaces were filtered for stop-words.⁴²⁵ Note that, as can be seen in the example above, stop-words are counted when determining the

⁴²⁴ The text in this example (Luc. DDeor 7.1 in the Perseus edition) is unlemmatized for ease of reading, which would not be the case in the actual experiment.

⁴²⁵ The list of stop-words used was compiled by Alessandro Vatri based on the Perseus Hopper Source, and is available at https://figshare.com/articles/Ancient_Greek_stop_words/9724613.

context window, but they are not added to the co-occurrence count; so if the context window is 5 on both sides and contains three stop-words, the number of co-occurrence pairs extracted by the script will be 7 (5+5-3).

The resulting combinations of parameters leads to twelve different semantic spaces to be assessed (three window sizes times four frequency thresholds). The vector space models were built using DISSECT (DIStributional SEmantics Composition Toolkit), a tool which creates count-based models (see VI.2.5).⁴²⁶ We used Positive Pointwise Mutual Information (PPMI) as our association measure, and applied Singular Value Decomposition (SVD) to reduce the resulting matrices to 300 latent dimensions.⁴²⁷

One vector space model was generated for each combination of parameters as detailed above: w1_t1, w1_t20, w1_t50, w1_t100, w5_t1, w5_t20, w5_t50, w5_t100, w10_t1, w10_t20, w10_t50, w10_t100, with *w* referring to the size of the context window on each side of the target word and *t* to the frequency threshold. The following section will describe the evaluation of how accurately these twelve semantic spaces represent meaning relations in ancient Greek.

VI.3.2 *The comparison resources*

Once the semantic spaces are built, the next step is assessing their accuracy in describing the language sample at hand. Similar assessments have been conducted on English-language corpora; in particular, Bullinaria and Levy (2007; 2012) argue that the best results are achieved with the smallest possible window size (1) and without filtering for frequency. The authors, however, are working on minimally processed corpora (stripped of sentence

⁴²⁶ DISSECT is described in Dinu, Pham, and Baroni (2013).

⁴²⁷ See §VI.2.3 on both of these concepts.

boundaries and unlemmatised); this is a very different situation compared to the richly annotated Diorisis corpus, on which we will indeed see different results.

The simplest way to do this would be to screen a sample of the results manually and formulate judgements of how well they match expectations.⁴²⁸ However, this method has several drawbacks: on the one hand, human brains are biased towards finding patterns in the existing data rather than considering other missing possibilities; on the other hand, semantic judgements of this kind are best gathered from a sizeable pool of native speakers, which is obviously not an option when working with ancient Greek. It is, therefore, methodologically sounder to refer to lists of synonyms in other, pre-existing resources, which can then be compared to the distributional synonyms (the words that share a closer proximity than others) in the semantic space(s); this is similar to the approach routinely taken to testing English-language corpora.⁴²⁹

In light of all this, the accuracy of each of the twelve spaces was assessed by comparing them to three different benchmarks, one drawn from ancient scholarship (Pollux's *Onomasticon*), one from modern-day lexicography (Schmidt's *Synonymik der Griechischen Sprache*⁴³⁰), and finally a resource derived from automatic Natural Language Processing (NLP), the Open Ancient Greek WordNet (AGWN⁴³¹). All these sources collapse information from different periods. Pollux' *Onomasticon*, composed in the late second century AD, gives us a wealth of information on Greek vocabulary, with particular but by no means exclusive emphasis on words used by Classical authors.⁴³² Somewhat similarly, Schmidt's dictionary of synonyms privileges classical and Homeric usage but also takes

⁴²⁸ A version of this procedure was adopted in Rodda, Senaldi, and Lenci (2017) to assess semantic change. For a methodologically more solid way to collect readers' judgements (on polysemy) that can be compared to the output of automated semantic analysis, see McGillivray et al. (2019).

⁴²⁹ See §VI.2.4.

⁴³⁰ Schmidt (1876-1886).

⁴³¹ Boschetti, Del Gratta, and Diakoff (2016).

⁴³² Pollux's definition of synonymy is remarkably similar to the distributional definition of the same concept: see the discussion by Philomen Probert in section 3.1.1 of Rodda, Probert, and McGillivray (2019).

Hellenistic sources into account, while AGWN includes data from dictionaries based on texts from a wide chronological span.

Using pre-existing resources instead of individuals' *post-hoc* judgements allows us to provide a robust assessment based on independent data. The characteristics of each source inevitably limit the absolute rate of accuracy obtained in the comparison with the semantic spaces, but they still allow us to compare the relative accuracy of different Distributional Semantics Models. The testing of the DSMs for this experiment was also the first time that such a comparison with pre-existing data, and especially ancient lexicographical resources, was attempted for ancient Greek.

A detailed description of the benchmark resources is provided in Rodda, Probert, and McGillivray (2019). Each of the resources is structured into lists of synonyms, collected with different goals in mind and using different criteria; for instance, Schmidt's definition of synonymy is based on a broad understanding of words belonging to the same semantic area, including etymologically related words,⁴³³ while AGWN is constructed on the basis of the Princeton English WordNet and structured into 'synsets' (synonym sets) that contain all the Greek words that share a specific English word in their dictionary definition.⁴³⁴

Despite these differences, each of the comparison resources gives us some lists of synonyms. All the lists available in Schmidt's dictionary and in WordNet were used, while for Pollux (where the search needed to be conducted manually) we operated by selecting 32 target words that appear in our semantic spaces as well as in the other two resources, and manually extracting the synonyms that Pollux lists for each of them. Additionally, for both Pollux and Schmidt's dictionary, only nouns were included in the synonym lists: this is

⁴³³ So, for instance, section 1 broadly gathers words for 'speaking,' ranging from verbs meaning 'to speak' (with different nuances, from λέγω 'to say' to λαλέω 'to chatter') to nouns for 'word' (λόγος, note the etymological relation) and 'voice' (φωνή).

⁴³⁴ The synsets are hierarchically ordered, as is characteristic of a WordNet database, so that broader meanings encompass narrower meanings.

because the case study (presented in §VI.4) focuses on nouns, and therefore we were particularly interested in assessing the accuracy of the model when it comes to this class.⁴³⁵

VI.3.3 Evaluation⁴³⁶

In order to compare the lists of lemmas in the benchmark resources to those in the semantic spaces, we extracted the nearest neighbours in each semantic space for each of the lemmas in the benchmark lists. Nearest neighbours are defined in the semantic space via the cosine distance between the vectors representing each word: the lower the cosine distance, the closer the two vectors are in the semantic space and the more closely related in meaning (according to the distributional hypothesis) the words are.

We then compared how many of the words in the list of nearest neighbours (we focused on the top 10 nearest neighbours for each word) also appear in the corresponding sets of synonyms in each of the benchmark resources. For instance, the nearest neighbours for *ικετεία* ‘supplication’ in the *w5_t50* space are *δέησις* ‘entreaty,’ *οἶκτος* ‘pity,’ *ικετεύω* ‘to beg,’ *ικεσία* ‘supplication,’ *ἐπικλάω* ‘to bend, move to pity,’ *ἐπήκοος* ‘listening,’ *ὑπερεῖδον* ‘looked over’ (past tense⁴³⁷), *αἴτησις* ‘request,’ *μνεία* ‘mention,’ and *λυταρής* ‘persisting.’ In AGWN, the synset for *ικετεία* also contains *σκήψις*, *παραίτησις*, *πρόφασις*, *λυτάρησις*, *γούνασμα*, and *ικεσία*.

Finally, the precision and recall of each DSM was measured across all lemmas.⁴³⁸ Precision is the ratio of shared neighbours (between the corpus neighbour-set and the benchmark

⁴³⁵ The complete lists of synonyms are available at https://github.com/alan-turing-institute/ancient-greek-semantic-space/blob/master/pollux_lexicon.txt and https://github.com/alan-turing-institute/ancient-greek-semantic-space/blob/master/schmidt_lexicon.txt.

⁴³⁶ This section is a very short summary of the findings in section 3.2 of Rodda, Probert, and McGillivray (2019), which is the work of Barbara McGillivray.

⁴³⁷ This word was clearly not lemmatised correctly. The presence of this sort of error is one of the reasons why it is necessary to test the accuracy of the semantic spaces.

⁴³⁸ For an explanation of the concepts of precision and recall, see §IV.1.4.

synset) among all the words in the corpus neighbour-set; recall is the same overlap between neighbours and synset, only as a proportion of the lemmas in the synset instead.⁴³⁹ These two measures assess how well the two sets match, by checking for both false positives (words that are in the neighbour-set but not among the expected synonyms gathered from the synset in the benchmark) and false negatives (words that are in the benchmark but do not appear as corpus neighbours). So, in the example of *ικετεία* above, with 10 words in the neighbour-set and 6 in the synset, and the only one in common to both sets being *ικεσία*, precision will be $1/10 = 0.10$ and recall will be $1/6 = 0.17$.

⁴³⁹ In mathematical terms, $Precision(l) = \frac{|synset(l) \cap neighbourset(l)|}{|neighbourset(l)|}$ and $Recall(l) = \frac{|synset(l) \cap neighbourset(l)|}{|synset(l)|}$.

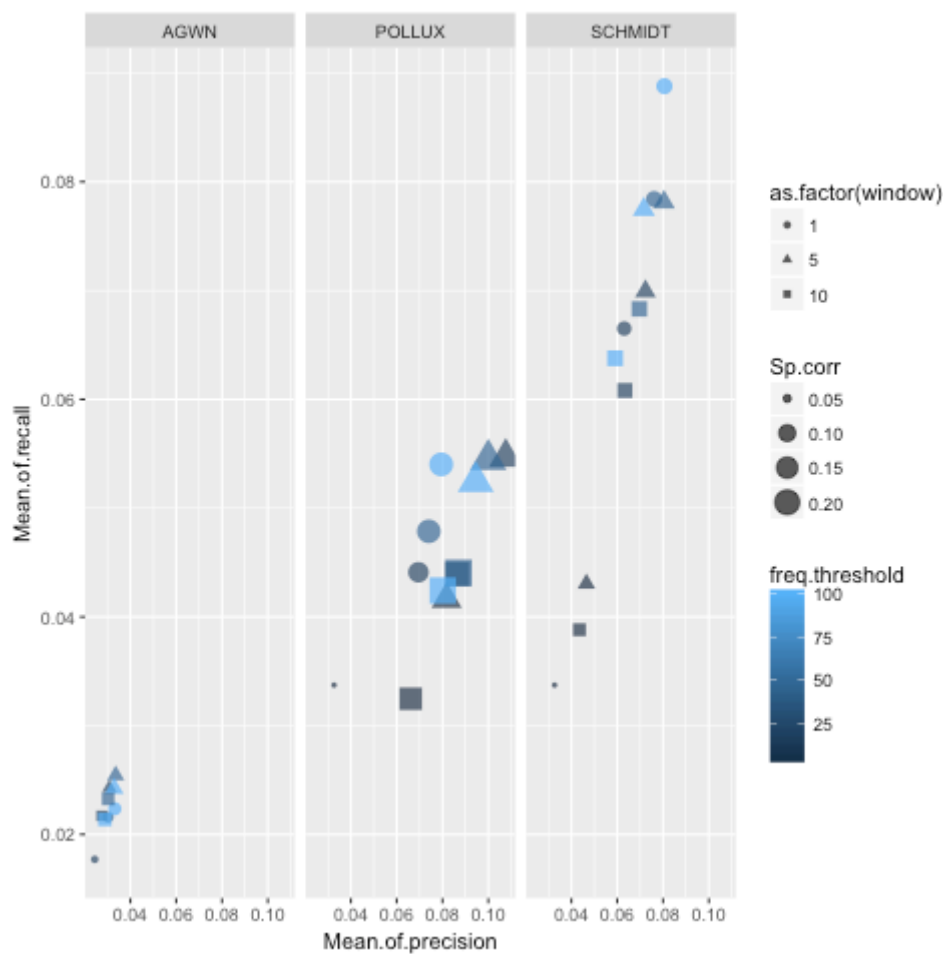


Figure 5: Visualization of evaluation metrics according to the gold standard resources and different parameters in the semantic spaces. From Rodda, Probert, and McGillivray 2019.

We also conducted a regression analysis and a correlation-based evaluation that takes into account the relative distances between words in the semantic spaces, as well as a ‘sum of inverse ranks’ correlation measure. All these analyses offer further quantitative insights into how closely the semantic spaces match the benchmarks, but as their results agree with those of the precision/recall measures, they are not discussed in detail here.⁴⁴⁰

Figure 5 shows the results of the comparison of all twelve semantic spaces against the three benchmark resources. Precision and recall are fairly low in absolute terms: this is partially to be ascribed to the nature of our benchmark sets. However, while we do not expect lists of synonyms generated in different ways and for different purposes to match

⁴⁴⁰ For more details, the interested reader should see Rodda, Probert, and McGillivray (2019).

perfectly, this analysis does give us a measure of how each vector space model behaves compared to the others.

The general trend is for both precision and recall to increase with higher frequency thresholds (lighter colours); the context window that performs best across all benchmark sets is 5 (triangles). Therefore, for the case study that follows, the vector space model with context window 5 and frequency threshold of 50 was chosen; it would theoretically be possible to push the threshold up to 100, but this option was discarded for fear of excluding too many lower-frequency words that are likely to come up in archaic Greek epic texts.

VI.4 Using the vector space model: a preliminary study

I now turn back to the issue that was initially sketched out in section VI.1: how can we devise a quantitative model of the semantic flexibility of a group of formulae? Our ultimate goal is to be able to test hypotheses about the role of semantic flexibility (and other factors) in the diachronic evolution of formulae; the experiment in this section, however, will have a more limited focus. It will introduce the method that will then be applied to a wider sample of formulae in chapter VII, but will only use it to discuss the behaviour a very small sample of formulae. I will then address how the semantic flexibility of these formulae correlates to a few potential explanatory factors, some of which will be discussed again in the larger study in the following chapter. While the results of the experiment in §VI.4.3 cannot be generalised to formulae in general, they are in themselves informative when it comes to the formulae analysed here, and they also serve to showcase the type of work that can be done on formulae in a Distributional Semantics perspective.

VI.4.1 *The formulaic sample and distance measurement*

For the following study, I will consider a small formulaic sample of transitive verb phrases. Three of these denote physical interaction with an object. They all involve the word χεῖρ, 'hand'.

1. χεῖρ + ἔχω, 'having in one's hand,' e.g. στέμματ' ἔχων ἐν χερσὶν ἐκηβόλου Ἀπόλλωνος (*Il.* 1.14);
2. χεῖρ + αἰρέω, 'picking up in one's hand,' e.g. δόρυ δ' εἵλετο χειρὶ παχείῃ (*Il.* 10.31);
3. χεῖρ + τίθημι, 'putting in/with one's hand,' e.g. ἀλόχοιο φίλης ἐν χερσὶν ἔθηκε / παῖδ' ἐόν (*Il.* 6.482-3).

The other two formulae denote two different means of perception and/or knowledge: intellectual and sensory. While they are less obviously related in form than the previous set of hand-formulae, part of the paradigm of the verbs involved shares the PIE root **wid-*, connected to both seeing and knowing.

4. εὖ εἰδέναι, 'knowing well,' e.g. τόξων εὖ εἰδώς (*Il.* 2.718);
5. ἰδεῖν ὀφθαλμοῖσιν, 'seeing with one's eyes,' e.g. ἦ μέγα θαῦμα τόδ' ὀφθαλμοῖσιν ὀρῶμαι (*Il.* 13.99).

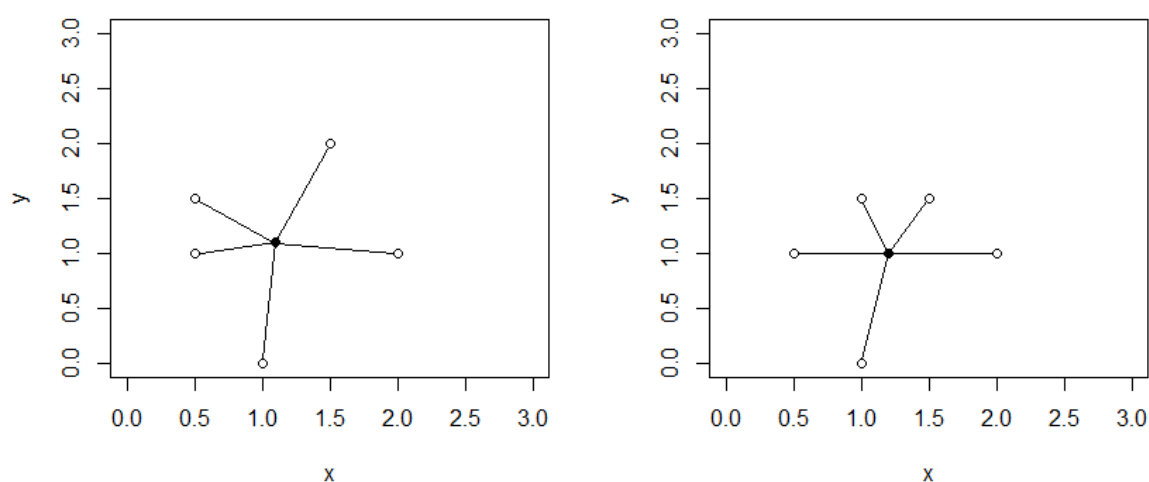
All occurrences of each of these phrases were extracted manually using the Perseus Scaife Viewer's lemma search function (to maintain uniformity of editions with the *Diorisis* corpus).⁴⁴¹ Their fillers were isolated and tabulated into a spreadsheet that also records textual location, word order, the form of the verb in the VP formula, the form of the associated NP/PP (case, number, presence or lack of a preposition), and the case taken by

⁴⁴¹ <https://scaife.perseus.org/search/>. I ran a proximity search for the key lemmas in each phrase: e.g. χεῖρ + ἔχω, or εὖ + εἶδον.

the direct object (this is relevant for the $\epsilon\tilde{\upsilon}$ εἶδέναι formula, which can take objects in the genitive or accusative).⁴⁴²

To assess the semantic similarity of fillers of each formula, a measure of distance was devised by calculating the cosine distance between each filler and the centroid of the whole distribution of fillers for the formula in question. Geometrically, the centroid (centre of

Figure 6: Centroids of two example distributions of points



mass) is the mean position of all the points in all the coordinate directions; Figure 6 shows two bidimensional examples. This measure allows us to compare the distribution of fillers-to-centroid distances between all formulae.

The other parameter to consider is the variance of each distribution, which measures how far the distances of all the individual fillers are spread out from their average value. This is a shorthand measure of how scattered the distribution of distances is.⁴⁴³ The two examples in Figure 6 have different variance: in the plot on the left, the distances between the points and the centroid is fairly uniform (lower variance), while in the one on the right some points

⁴⁴² So, for instance, the example in number 1 above is tagged as *ll. 1.14*, verb + noun [not vice versa], pres pc p m s nom, preposition ἐν, object lemma στέμμα, no obj. adjective and no anaphora. Anaphoric pronouns have been resolved by finding their referent and considering the latter as the object of the verb.

⁴⁴³ See e.g. Baayen (2008), 62.

are noticeably closer to the centroid and some are noticeably farther away than the others (higher variance).

Both the average and variance of the distribution of distances from the centroid for the fillers for each formula will be reported in the results section below. Before moving on to the results of the experiment, however, it is necessary to formulate some hypotheses about the behaviour of the formulaic sample; these hypotheses will then be tested against the data from the semantic space.

VI.4.2 *Competing hypotheses*

The five target phrases can be grouped in at least three different ways based on their salient characteristics other than semantic flexibility. These characteristics will constitute our potential correlates for semantic variation in the experiment.

Hypothesis 1: Semantic flexibility in object fillers correlates with the general meaning of a formula. According to this hypothesis, the formulae of physical contact (the three hand-formulae) should behave differently in terms of their semantic openness from the formulae of perception/knowledge. The grouping predicted by this hypothesis, therefore, is:

χείρ + ἔχω, χείρ + αἰρέω, χείρ + τίθημι vs εὔριδέναι, ἰδεῖν ὀφθαλμοῖσιν

Hypothesis 2: Semantic openness correlates with diachronic productivity, or at least productivity outside of Homer (whether this means diachronic evolution or belonging to a different branch of the tradition). The χείρ + ἔχω and ἰδεῖν ὀφθαλμοῖσιν formulae are attested in Hesiod and the *Hymns* as well as the *Iliad* and *Odyssey*, while εὔριδέναι is only attested in the Homeric epics. The other two hand-formulae only occur 2 and 1 times, respectively, in Hesiod and the *Hymns*, so while they appear outside of Homer, their attestation is sparser. The grouping predicted by this hypothesis is:

χείρ + ἔχω, ἰδεῖν ὀφθαλμοῖσ(ι) (+ χεῖρ + αἰρέω, χεῖρ + τίθημι) vs εὖ εἰδέναι

This grouping clearly differs from the one predicted by hypothesis 1, where χεῖρ + ἔχω and ἰδεῖν ὀφθαλμοῖσ(ι) are predicted to be polarly opposite.

Hypothesis 3: Semantic openness correlates with syntactic flexibility. We use syntactic flexibility as a broad concept, which encompasses different phenomena in different formulae. The εὖ εἰδέναι formula takes objects in either the accusative or genitive, with some differences in usage (the accusative is generally used for pronouns, the genitive for noun objects, although there are some exceptions⁴⁴⁴). On the opposite end of the spectrum, ἰδεῖν ὀφθαλμοῖσ(ι) only admits inflection in the verb; the hand-formulae show some flexibility in whether χεῖρ appears with or without a preposition, and which preposition is used. The grouping predicted by this hypothesis is therefore:

εὖ εἰδέναι (+ χεῖρ + ἔχω, χεῖρ + αἰρέω, χεῖρ + τίθημι) vs ἰδεῖν ὀφθαλμοῖσ(ι)

Unfortunately, this grouping is remarkably similar to the one predicted by hypothesis 2, with the only difference that we expect χεῖρ + ἔχω to pattern closer to εὖ εἰδέναι than ἰδεῖν ὀφθαλμοῖσ(ι).

One complication that will affect our assessment of results is the frequency of the target formulae. While all of them have a reasonably high number of attestations, χεῖρ + ἔχω (58x), ἰδεῖν ὀφθαλμοῖσ(ι) (52x), and to an extent εὖ εἰδέναι (41x) are much more frequently used than χεῖρ + αἰρέω (34x) and χεῖρ + τίθημι (28x). These differences in frequency will need to be taken into account when looking at the significance of the results.

⁴⁴⁴ These are *Il.* 13.665, κῆρ' ὀλοήν, 20.214, ἡμετέρην γενεήν, *Od.* 9.215, δίκας ... θέμιστας, 11.442, μῦθον ἅπαντα, 11.445, μήδεα, and 14.366, νόστον.

VI.4.3 Results and discussion

The distribution of cosine distances between the centroid and the object fillers for each formula is visualised in the boxplot in Figure 7. The thick black line represents the mean distance for each distribution, while the boxes and dashed lines represent its quantiles: the white box encompasses 50% of the points, the dashed lines 100%, excluding outliers (marked by an empty white dot).⁴⁴⁵

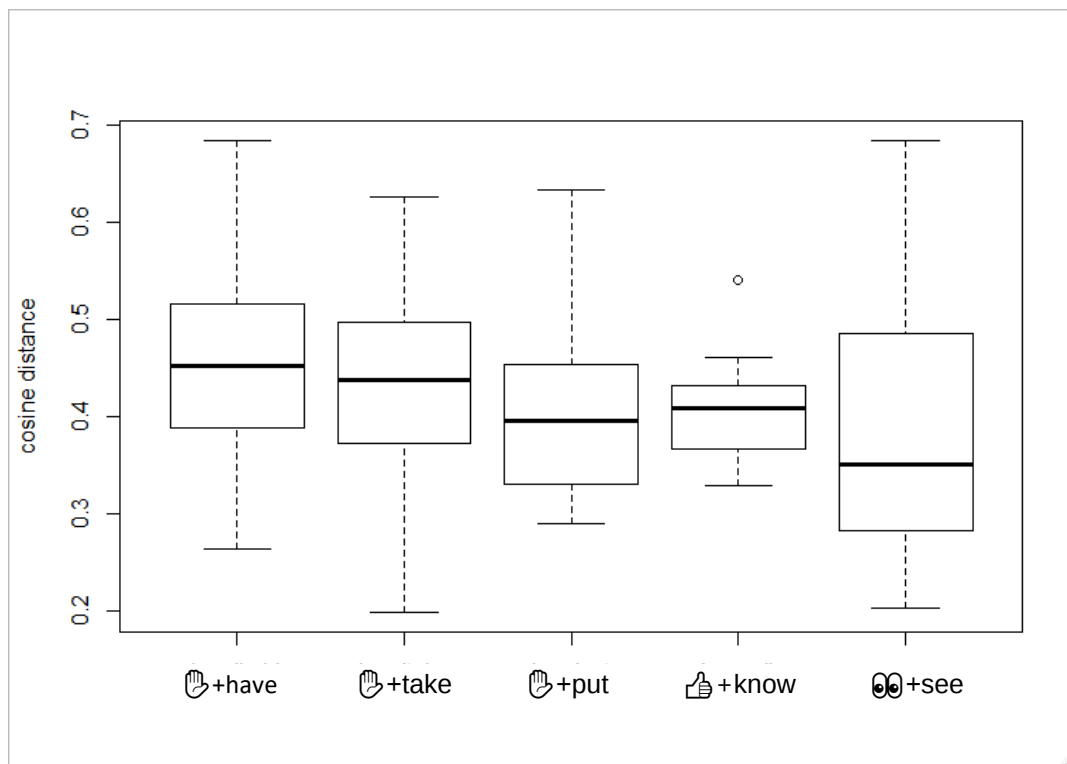


Figure 7: Distribution of cosine distances between the centroid and the object fillers for each formula

The two tables below report, respectively, the result of a pairwise significance test between the five formulaic samples, corrected for multiple comparisons,⁴⁴⁶ and the variance of each

⁴⁴⁵ Note, therefore, that the dashed lines are neither confidence intervals nor measures of variance, although the length of the boxplot for each distribution can be used as a visual shorthand for the latter – the longer the boxplot, the more scattered the results are, and therefore the higher the variance.

⁴⁴⁶ I applied a Bonferroni correction of $0.05/4$ (number of comparisons to be performed between 5 items) = 0.0125 significance threshold. See e.g. Baayen (2008), 106.

distribution. Significant differences in the first table are highlighted in bold; a marginally significant result is in italics.

Table 19: Pairwise comparison between the distribution of fillers in each formulaic sample

Pairwise Kolmogorov-Smirnoff test (Bonferroni-corrected)					
	✋+have	✋+take	✋+put	👍+know	👁+see
✋+have					
✋+take	0.7197				
✋+put	0.09325	0.3273			
👍+know	0.002604	0.07348	0.7056		
👁+see	0.01312	0.1847	0.2078	0.01199	

Table 20: Variance of each distribution of formulaic fillers

✋+have	✋+take	✋+put	👍+know	👁+see
0.0095	0.0124	0.0085	0.0026	0.0163

These results are less clear-cut than would be ideal, but some things can be said even about such a small sample. It helps to go through the hypotheses detailed in section VI.4.2 one by one and discuss how well they match the observed distribution.

Hypothesis 1 predicted a meaning-based grouping of formulae for physical contact vs formulae for perception. Something similar can be seen in the boxplot, where the physical contact formulae have higher mean distances (the bold line) than the perception formulae; and in the pairwise statistical comparison, where there is a significant difference between the χεῖρ + ἔχω formula and both perception formulae. No differences are significant when it comes to the other two hand-formulae, which may be due to their lower frequency.

Our initial prediction, however, turns out to be insufficient: there is also a significant difference between the two formulae of perception (εἶ εἰδέναι vs ἰδεῖν ὀφθαλμοῖς(ι)). There are, therefore, three meaning-based groups instead of two. This highlights the arbitrary nature of judgements about the general meaning of a formula, an observation that will need to be kept in mind in designing further experiments.

Hypothesis 2, diachronic (or cross-traditional) productivity, on the other hand, does not seem to be a particularly good predictor of behaviour. It is a very poor predictor of average distance, considering that the two most productive formulae, χεῖρ + ἔχω and ἰδεῖν ὀφθαλμοῖσ(ι), are on opposite ends of the boxplot. It is, however, worth noticing that the only formula that is not attested at all outside Homer, the εὔ ἰδέναί formula, has an impressively lower variance than the others: this means that its results are not just farther away from the centroid, but also more regularly spaced. On the other hand, both formulae that are very productive outside of the Homeric epics have higher variance. Does semantic coherence (fillers that are grouped in some kind of cluster, which will lead to a more irregular distribution) somehow promote flexibility, in a similar way as semantic density promotes the diachronic productivity of a construction (see VI.1.2)?

Unfortunately, variance also seems to correlate with syntactic flexibility (hypothesis 3). The most syntactically flexible formula, εὔ ἰδέναί, which can take objects in two different cases, has the lowest variance; the least flexible, ἰδεῖν ὀφθαλμοῖσ(ι) has the highest variance, while the other three formulae fall in between these two. It could be argued⁴⁴⁷ that a hypothetical construction with both high syntactic flexibility and high semantic flexibility would not be perceived as a formula, which would explain the correlation between results for the two hypotheses; but any conclusion drawn at this point will be premature, given the limited sample size.

VI.5 Conclusions

All in all, the results of the semantic analysis conducted in §VI.4 are anything but straightforward. Part of this is undoubtedly due to the small scale of the sample – an issue that will be addressed in the next chapter. Even this pocket-sized experiment, however, is

⁴⁴⁷ I owe this suggestion to a colleague at the London Universities Classics PGWiP.

helpful in terms of highlighting relevant parameters for a future analysis to be conducted on a much larger sample: we can consider both mean distance and variance in the semantics of filler objects as promising correlates of syntactic and semantic variation. If we want to continue to focus on syntactic variation, we also need a better definition of this parameter, one that can account for the conceptual difference between a formula that allows objects in different cases (variation in the dependency structure, as was the case for εὖ εἰδέναι) and a formula that only presents internal variation (such as the hand-formulae).

Another question that arises is whether focusing on the filler objects of transitive verb formulae may open ourselves up to the possibility that the meaning of these fillers is entirely context-determined, rather than having anything to do with the semantics of the construction they occur in. In other words, is there not a difference between the open slot in a construction like ___ *the hell out of* and the open slot in *take ___ in one's hands*? Is the filler in the latter not determined by some external concept – the reality of the object one needs to take in one's hands? This may be true in everyday language,⁴⁴⁸ but formulae are in this respect more rigid than constructions: the objects of the transitive verb formulae analysed above have to conform to a certain metrical shape and are themselves part of a traditional system of formulaic language. It is, however, not encouraging to see that the meaning of the constructions analysed in §VI.4 has such an important effect on the meaning of their fillers; we shall see if this tendency is upheld in a more representative sample.

The natural next step at this point is to extend our study to a sample of formulae that will not just be larger, but will be on some level representative of a class, in the same way as we expanded the initial experiment in chapter III. This will be the focus of the next chapter.

⁴⁴⁸ It is only partially true for idioms as well: there are restrictions, for instance, on the kinds of things one can take in one's hands (*situations* and *matters* are more likely than *kettles* or *branches*, even though the phrase can be used both literally and metaphorically).

VII Mapping semantic variation in transitive VP formulae

The goals of this study, as anticipated in chapter VI, are to assess the scope of semantic variation as fully as possible in a sample of epic formulae, and then to compare the results for this sample with a comparison corpus, parallel to what was done in chapter IV for syntactic variation. We will use Distributional Semantics as the framework within which to quantify semantic variation, as explained in the previous chapter. The target of our analysis is a sample of formulae made of a transitive verb and its direct object in the accusative (from here on, TrV + Obj formulae), selected exclusively on the basis of frequency. Specifically, we will be looking at the semantic range of the objects of this group of transitive verbs; we will answer questions about the semantic behaviour of these objects depending on formulaic status, frequency, and the productivity of the transitive verb phrases themselves. Some of these questions come from the literature on the history of English and Icelandic reviewed in §VI.1.2: does the semantic range of formulae influence their productivity in the same way it does for verbal constructions in these languages? Others come from our analysis in chapter IV, where frequency was a predictor of syntactic variation: does the same hold for semantic variation? And finally, is there any kind of relationship between formulaic status and semantic range in phrases with open slots that can be filled by different nouns? Do formulae have more tightly clustered or more sparsely scattered object fillers? In order to answer these questions, we will return to the method adopted in chapter IV, which involves comparing a formulaic corpus to a non-formulaic one in order to highlight any specific differences that can be used to characterise formulae.

VII.1 Materials and methods

VII.1.1 *Target formulae*

The choice of target formulae for this study differs from the one used in the small-scale study in the previous chapter in two substantial ways, beyond the fact that it is more comprehensive.

The first difference lies in the fact that we will be looking at TrV + Obj phrases in which the object is itself part of the phrase, as opposed to the previous study, where the objects were external to a formula made of a verb and a modifier (a prepositional phrase, a dative of instrument, or an adverb of manner). Like we did in §IV, we will talk about recurring phrases of this kind as formulae when they occur in early Greek epic; this is an assumption based on the guiding principle that things that frequently occur together in epic poetry can be defined as formulaic, in the way Hainsworth intended. We will also follow the formulaic analysis conducted by Pavese and Venti (2000) and Pavese and Boschetti (2003), which adopts similar principles (see §VII.1.3 below). Other than this, this chapter operates in a fairly theory-agnostic way: our choice to analyse phrases with the same verb and different (peripheral) objects can be compared to the views of formulae as made up of a nucleus and a periphery as discussed in §I.3.2, but this is neither a perfect match nor something that will significantly shape the argument here. Much more relevant is the comparison to studies of constructions with open slots in modern languages, as discussed in §VI.1.2.

Focusing on objects that are internal, rather than external, to a formula ensures that we are only considering formulaic material, which will be crucial when introducing a comparison to non-formulaic texts, as we did in chapter IV. Since we are aiming to test the hypothesis that formulaicity introduces some sort of restriction in the semantic scope of TrV + Obj phrases, we should consider phrases in which the object is also formulaic. The goals of this chapter, therefore, are different from those of the previous one, where we were aiming to

test the viability and limits of the Distributional Semantics approach and to explore potential explanatory factors for semantic flexibility. This difference in aims explains the difference in target sample (fully formulaic TrV + Obj phrases here vs formulaic TrV phrases with an object that is not part of the formula in chapter VI.4).

The second difference is that, as will be further detailed below, only formulae with an object in the accusative will be considered in this chapter. This choice was made in order to ensure a higher uniformity of the sample by removing a parameter for variation (the case of syntactic dependencies) whose influence (if any) would have been difficult to control for; it also has the useful side-effect of simplifying the data collection process and keeping the number of tokens manageable for a human reader. Excluding objects in cases other than the accusative, however, means abandoning one of the avenues of investigation that had opened up in the previous chapter, where a formula which admitted objects in two different cases showed interestingly lower levels of flexibility in meaning.

The kind of formulae we will be looking at in this chapter, therefore, are pairs such as μῆνιν ἄειδε, ἄλγε' ἔθηκε, or ψυχὰς ... προΐαψεν, to pick three examples from the very first sentence of the *Iliad*; note how neither the order of the two elements, verb and direct object in the accusative, nor their immediate contiguity are a relevant parameter for the analysis. Some of these formulae, like the κύδος / τεύχε' / ἄλγε' ἔδωκε system, have already been recognised as a textbook case of formulaic substitution;⁴⁴⁹ no data exists, however, about the semantic relations between the members of this substitution system. In this chapter, we are going to explore some of the rules that lead to the formation of formulae within this sort of substitution systems – what has been defined as “a grammar of composition, a transformational mechanism for the economical construction of phrases

⁴⁴⁹ See e.g. Fantuzzi (1988), 13.

with new meanings but analogous rhythmic structure and signifiers to traditional ones.”⁴⁵⁰

Specifically, we will be looking at the role of semantics in promoting or discouraging formulaic variation, in isolation and in comparison to non-formulaic material.

VII.1.2 The target corpus

Collecting data on transitive verb formulae creates very different practical requirements compared to our Adj + N study in chapter IV. Specifically, we need a corpus that contains information not only about the lemma, part of speech, and morphological parsing of each word, but also about the syntactic relation in which the word stands to other parts of the sentence. In more technical terms, we need a dependency treebank, that is, a corpus that encodes syntactic dependencies.⁴⁵¹ The Diorisis corpus does not contain this information, and therefore cannot be used for this new study. However, several dependency treebanks are available for classical languages. A very recent review of these resources is provided by Celano (2019), who surveys four different corpora: the Ancient Greek and Latin Dependency Treebank (Bamman and Crane 2011; Tufts University/Universität Leipzig), the Index Thomisticus Treebank (McGillivray, Passarotti, and Ruffolo 2009; Università Cattolica di Milano), the PROIEL Treebank (Haug and Jøhndal 2008; Universitetet i Oslo), and the SEMATIA Treebank (Vierros 2018; Helsingin Yliopisto).

Dependency treebanks essentially represent syntactic trees, in which each node can have one and only one head⁴⁵² but potentially several children (and, by implication, siblings). As

⁴⁵⁰ Fantuzzi (1988), 15: “[il meccanismo dell’analogia], che costituiva una sorta di grammatica della composizione, un meccanismo trasformatore per costruire economicamente frasi con significati nuovi ma strutture ritmiche e significati analoghi a quelli tradizionali.”

⁴⁵¹ More formally, “[a] dependency treebank is a corpus containing a symbolic representation of the syntax of one or more texts. It can be defined as a set of sentences parsed according to the linguistic formalism of dependency grammar” (Celano 2019). On treebanks as one of many possible forms of annotated corpus see also Jensen and McGillivray (2017), 115-119.

⁴⁵² By ‘head’ we mean here, as generally in syntax, the central element of a phrase, which determines all concord and government relations in the phrase or sentence (see e.g. Crystal 2008, s.v.).

Celano notes, this structure is perfectly parallel to that of an XML tree;⁴⁵³ for this reason, dependency treebanks are normally encoded as XML databases. The treebank used for our study is version 1.6 of the Ancient Greek and Latin Dependency Treebank (AGLDT), that is, the treebank associated with the Perseus Project.⁴⁵⁴ Version 1.6 is the most recent version to contain all of the *Iliad* and the *Odyssey* as well as Hesiod (*Theogony* and *Works and Days*); none of the treebanks available contain the *Homeric Hymns*, which therefore have to be excluded from this study.

AGLDT is manually annotated for syntactic dependencies by expert readers,⁴⁵⁵ while the morphological annotation is automatic in the first stage, and then checked by human annotators.⁴⁵⁶ Morphological information for each word is provided in the form of a nine-character string, in which each element represents one part of the analysis.⁴⁵⁷ Syntactic information, on the other hand, is contained in a separate tag, which can take a series of predefined values, from PRED = predicate to ExD = ellipsis; each word is also tagged with the ID for its syntactic head in the sentence structure.⁴⁵⁸ In addition to this, the tags for each word include its lemma, as well as a unique citation URN (Uniform Resource Name).⁴⁵⁹ The whole corpus is also broken down into sentences, much like Diorisis; the numbering of syntactic heads restarts with each sentence.

⁴⁵³ Celano (2019), 281, with some nuance on e.g. coordinate relations. On XML trees and their structure, see §IV.1.1.

⁴⁵⁴ https://github.com/PerseusDL/treebank_data, downloaded November 2019.

⁴⁵⁵ Specifically, the work of two annotators is checked by a third expert annotator, who reconciles any discrepancies (Bamman and Crane 2011, 4). The annotators are identified in the metadata of each .xml file, which makes the process completely transparent.

⁴⁵⁶ Celano (2019), 285.

⁴⁵⁷ Detailed guidelines for the morphological annotation are provided at https://github.com/PerseusDL/treebank_data/tree/master/v1/greek/docs (last accessed January 2021).

⁴⁵⁸ Detailed guidelines for the syntactic annotation are provided at https://github.com/PerseusDL/treebank_data/blob/master/v1/greek/docs/guidelines.pdf, with examples (last accessed January 2021).

⁴⁵⁹ See URI Planning Interest Group, W3C/IETF (2001).

Below is an example of a short clause, the first in the *Iliad*, as tagged in AGLDT 1.6. A brief explanation of the tagging system is provided on the side.

<code><word id="1"</code>	unique ID number for each word ⁴⁶⁰
<code>form="μῆνιν"</code>	word form as it appears in the text
<code>lemma="μῆνις1"</code>	dictionary lemma for the word
<code>postag="n-s---fa-"</code>	morphological analysis (PoS tag): Noun, Singular, Feminine, Accusative
<code>relation="OBJ"</code>	syntactic relation to the head: Object
<code>head="2"</code>	ID number of the head (here, ἄειδε)
<code>cite=</code>	unique citation URN
<code>"urn:cts:greekLit:tlg0012.tlg001:1.1"/></code>	
<code><word id="2"</code>	unique ID number
<code>form="ἄειδε"</code>	word form
<code>lemma="ἄειδω1"</code>	dictionary lemma
<code>postag="v2spma---"</code>	morphological analysis (PoS tag)
<code>relation="PRED_CO"</code>	syntactic relation to the head
<code>head="32"</code>	ID number of the head
<code>cite="..."/></code>	unique citation URN
<code><word id="3"</code>	(and so on for each word)
<code>form="θεᾶ"</code>	
<code>lemma="θεᾶ1"</code>	
<code>postag="n-s---fv-"</code>	
<code>relation="ExD"</code>	
<code>head="2"</code>	

⁴⁶⁰ This counter restarts at the start of each new sentence. Sentences are tagged separately, as shown for Diorisis in §IV.1.1.

```

cite="..."/>
<word id="4"
form="Πηληϊάδεω"
lemma="Πηληιάδης"
postag="n-s---mg-"
relation="ATR"
head="5"
cite="..."/>
<word id="5"
form="Ἀχιλῆος"
lemma="Ἀχιλλεύς1"
postag="n-s---mg-"
relation="ATR"
head="1"
cite="..."/>

```

Note in particular how dependencies are marked in the example above: the attributive genitive Ἀχιλῆος (ATR) has as its head the object noun μῆνιν (OBJ), whose head is in turn the main verb ἔειδε, tagged as PRED_CO as it is coordinate to another predicate in the same sentence (Διὸς δ' ἔτελείετο βουλή, according to the analysis in AGLDT).

The specific syntactic formalism that lies behind the AGLDT is of limited interest for our work,⁴⁶¹ which will only make use of an extremely limited part of it, but some features of the tagging system will be relevant to the data collection and should be briefly described. As we will only be looking at accusative direct objects of transitive verbs, the part of the AGLDT tagging system we will be interacting with, apart from the morphological tagging,⁴⁶²

⁴⁶¹ More details about it are provided in Bamman and Crane (2011).

⁴⁶² See the example above for how a noun in the accusative, μῆνιν, is PoS tagged.

is the set of OBJ tags. As detailed in the syntactic annotation guidelines,⁴⁶³ the OBJ tag is used for all obligatory arguments of verbs, independently of their case and part of speech; this includes indirect objects, complements, as well as infinitive verbs in accusative + infinite constructions, and even finite verbs in subordinate object clauses. In order to weed out some of the complexity that this would lead to, we have restricted our experiment to direct objects that are morphologically tagged as accusative, which includes nouns, pronouns, and adjectives (arguably only adjectives that are used as nouns).

The other element of the tagging system that creates some issues is the existence of a _CO tag, which is added to the syntactic tag for any element that is coordinated to another (see the example above, which contains a PRED_CO tag). The script that is described below in §VII.1.3 only retrieves objects that are tagged as OBJ, not OBJ_CO; this is to avoid situations in which, for instance, only one of two coordinated objects of a transitive verb is part of a formula with the transitive verb itself, which would introduce an unknown variable into the data.⁴⁶⁴

VII.1.3 Data collection and identification of formulaic pairs

A list of transitive verbs with their respective objects in the accusative was retrieved using a Python script that I wrote for the purpose.⁴⁶⁵ The script retrieves all verbs (elements which contain the PoS tag *v*) and their word ID; it then retrieves all objects in the accusative (elements tagged as OBJ whose PoS tag contains an *a*) which have as their head a word ID associated to a verb in the same sentence, and pairs them with their respective verb. The

⁴⁶³ As cited in n. 458, pp. 8-14.

⁴⁶⁴ The issue of predicates being tagged as PRED_CO is irrelevant, as we will soon see, since the data collection script does not consider the syntactic tag of the head verb, but only its morphological tag – i.e., the fact that it is a verb.

⁴⁶⁵ The script is available in the same repository as the others at https://github.com/MartinaAstridRodda/dphil-thesis/blob/main/extract_obj_1.0.py.

output is a Python dictionary in which the heads are word IDs corresponding to verbs, each mapped to the verb in question and its object(s). The script can be modified to retrieve either the word forms as they appear in the text or the corresponding lemmas. Below is an example of the output for the first sentence in the *Iliad*:

```
"2"
"ἄειδε"
"μῆνιν"
"12"
"ἔθηκε"
"ἄλγε'"
"19"
"προΐαψεν"
"ψυχὰς"
```

This output was printed as a series of 50 .csv files, one for each book of the *Iliad* and *Odyssey* and for each Hesiodic poem; the files were then manually cleaned up to get rid of any elements that were not TrV + Obj phrases.⁴⁶⁶

The next step, since our study will involve formulaic phrases, is to reduce the dataset to these only. To do this, it is methodologically sound practice to use an independent definition of formula, that is, one that was not related to our initial criteria for the retrieval of target phrases and, ideally, was not designed for this experiment. I chose to use the formulaic analysis by Pavese and Boschetti/Pavese and Venti in their editions of the *Iliad* and *Odyssey* (Pavese and Boschetti 2003), *Theogony* and *Works and Days* (Pavese and Venti 2000). I manually cross-referenced the cleaned-up list of TrV + Obj pairs retrieved by the

⁴⁶⁶ Due to the structure of the PoS tag system, the script also retrieves nouns and adjectives in the vocative as verbs (their PoS tag contains *v*), and objects that are adjectives or active verbs (both tagged with *a*). These elements have been filtered out manually.

Python script with each formulaic edition, marked the formulaic pairs, and pared down the list to formulae only.⁴⁶⁷

Pavese and Boschetti define a formula as “a group of words, composed of at least two lexemes, used at least twice to express a certain meaning with identical metrical value;” lexeme is used here to mean a content word, that is, not a function word.⁴⁶⁸ They further differentiate between formulae that are found at least twice in the Homeric poems, formulae that are only found once in Homer and once in Hesiod, once in Homer and once in the *Hymns*, and so on; all of these were counted as formulaic material in my data, with no distinction. Furthermore, I made no attempt to note whether the verb and object in each of my pairs were part of the same formula or not, as this would have created additional difficulties in the data collection; I considered it enough that both elements of the pair are formulaic.⁴⁶⁹

The method described above leads to a total of 6764 formulaic TrV + Obj pairs; Table 21 below shows the numbers of formulaic pairs for each book of each poem, as well as the frequency of formulae per one hundred lines for each book. (For the books marked with an asterisk, see note 469.) Perhaps the most notable feature here is the lower frequency of TrV + Obj formulae per line in Hesiod.

⁴⁶⁷ Datasets for this process are available at https://github.com/MartinaAstridRodda/dphil-thesis/blob/main/semantics_dataset.xlsx.

⁴⁶⁸ Pavese and Boschetti (2003), 24. Formulae in the *Iliad* and *Odyssey* were automatically identified using a computer script, while Pavese and Venti (2000) performed the comparisons manually.

⁴⁶⁹ Over the course of the comparison, I have noted a few discrepancies between the Pavese and Boschetti/Pavese and Venti formulaic editions and the Perseus files. In *Odyssey* 6, p. 156 of Pavese and Boschetti (2003) was never printed, but is instead substituted by a repetition of p. 150; this leads to the loss of ll. 40-80. Similarly, on pp. 468-70, a block of lines has been printed twice and read twice by the script, which means that everything in *Od.* 17.161-193 is considered formulaic. In both cases, the missing or misprinted lines have been excluded, with no further impact on the analysis: while the number of tokens of a specific TrV + Obj pair may have been marginally affected, we will ultimately work with object types, rather than tokens (see below, §VII.1.6). Finally, the AGLDT treebank file for the *Theogony* has some unexplained additional material inserted between ll. 929 and 930; I have not found any trace of this material either in Pavese and Venti (2000) or in the apparatus and notes of West (1966), and have therefore removed it.

Table 21: Number of formulaic TrV + Obj pairs in each book of the epic corpus

Work	Book	TrV + Obj formulae	Lines ⁴⁷⁰	Formulae/100ll.
<i>Iliad</i>	1	164	611	26.8
	2	141	877	16.1
	3	98	461	21.3
	4	127	544	23.3
	5	201	909	22.1
	6	112	529	21.2
	7	123	482	25.5
	8	161	565	28.5
	9	170	713	23.8
	10	129	579	22.3
	11	201	848	23.7
	12	76	471	16.1
	13	170	837	20.3
	14	117	522	22.4
	15	160	746	21.4
	16	199	867	23.0
	17	149	761	19.6
	18 ⁴⁷¹	149	617	24.1
	19	98	424	23.1
	20	105	503	20.9
	21	132	611	21.6
	22	108	515	21.0
	23	162	897	18.1
	24	211	804	26.2
	total <i>Iliad</i>	3463	15693	22.1
<i>Odyssey</i>	1	131	444	29.5
	2	129	434	29.7
	3	145	497	29.2
	4	223	847	26.3
	5	115	493	23.3
	6*	47	331	14.2
	7	80	347	23.1
	8	124	586	21.2
	9	140	566	24.7
	10	131	574	22.8
	11	149	640	23.3
	12	96	453	21.2
	13	106	440	24.1
	14	113	533	21.2
	15	147	557	26.4
	16	135	481	28.1
	17*	193	606	31.8
	18	88	428	20.6
	19	162	604	26.8
	20	101	394	25.6

⁴⁷⁰ The number of lines was manually computed on the Perseus online editions as published on <https://www.perseus.tufts.edu/hopper/>, since the metadata for AGLDT do not contain line numbers.

⁴⁷¹ Books 17 and 18 of the *Iliad* do in fact have the same number of TrV + Obj formulae – this is not a misprint.

	21	85	433	19.6
	22	133	501	26.5
	23	113	372	30.4
	24	148	548	27.0
	total <i>Odyssey</i>	3034	12109	25.1
<i>Theogony</i> *		143	1022	14.0
<i>Works and Days</i>		124	828	15.0
	total Hesiod	267	1850	14.4
overall total		6764	29652	22.8

VII.1.4 Frequency filtering and the list of target verbs

The list of formulaic TrV + Obj pairs was manually aligned with lemmatised data from the Ancient Greek and Latin Dependency Treebank (the same corpus the pairs were extracted from): each form as it appears in the corpus was paired with the corresponding dictionary entry, extracted from the 'lemma' tag in AGLDT. The example in §VII.1.3, from the first sentence in the *Iliad*, now looks like this:

"2"	"2"
"ἄειδε"	"ἄείδω1"
"μῆνιν"	"μῆνις1"
"12"	"12"
"ἔθηκε"	"τίθημι1"
"ἄλγε'"	"ἄλγος1"
"19"	"19"
"προΐαψεν"	"προιάπτω1"
"ψυχὰς"	"ψυχή1"

The numbers at the end of each lemma entry are used to distinguish between different meanings of the same word by the AGLDT lemmatising standard. Since we are not concerned with the issue of homonymy for this part of the work, and in order to ensure compatibility with the semantic spaces built from Diorisis (see chapter VI.3), these numbers

were completely stripped off when cleaning up the data. This means no attempt was made to disambiguate between homonyms.

The lemmatised data was used to tally up the frequency of each transitive verb in the formulaic list. Frequencies range from 335 occurrences of ἔχω (the second-most frequent word type⁴⁷² in the sample after the determiner ὁ, which occurs 882 times⁴⁷³) all the way down to *hapax legomena*, which as expected occur in large numbers.⁴⁷⁴ Based on these results, I selected a sample of 26 target verbs, consisting of all verbs that occur at least 50 times in formulae with a direct object in the accusative. Table 22 lists the target verbs in decreasing order of frequency:

Table 22: List of transitive verbs occurring at least 50 times in formulae in the epic corpus

	Verb	Meaning	Token frequency
1	ἔχω	have	335
2	πρόσφημι	speak to, address	202
3	αἰρέω	take	167
4	προσεῖπον	speak to, address	163
5	δίδωμι	give	157 (158*)
6	εἶπον	say	153 (154*)
7	προσαυδάω	speak to, address	142
8	φέρω	carry, bring	109
9	βάλλω	throw	108
10	εἶδον	see	102 (103*)
11	τίθημι	put, set, place	101
12	ικνέομαι	arrive, reach	100
13	οἶδα	know	87
14	ἀμείβω	exchange; reply	86
15	πάσχω	suffer	84
16	χέω	pour	83 (84*)
17	αὐδάω	speak, utter	77
18	ἀγορεύω	speak, say	71
19	φημί	speak, say	59
20	ἄγω	lead	57

⁴⁷² Types are distinct words that may occur several times in a corpus, while tokens are all the individual occurrences of a given type. For a more detailed discussion, see below.

⁴⁷³ Note the effects of Zipf's Law (see n. 405 in §VI.2.3): the second most frequent word type should have a token frequency of (approximately) half the first most frequent one. The curve here is even steeper than predicted by the law.

⁴⁷⁴ The raw data is available at https://github.com/MartinaAstridRodda/dphil-thesis/blob/main/semantics_dataset.xlsx.

21	ἐρύω	draw, pull; protect	56
22	λύω	loosen	54
23	τίκτω	give birth	54
24	ἵημι	send	53
25	ῥέζω	do	50 (52**)
26	ἀκούω	hear	50
total			2703

The verbs marked with an asterisk occur once in object position; this occurrence was of course discarded. One verb, ῥέζω, occurs twice in object position. This is possible because, under our definition of direct object, participles in the accusative are considered direct objects and extracted as such. Moreover, some of these verbs can occur with more than one object in the accusative: while coordinate objects were excluded as described in §VII.1.3, some verbs, such as προσαιδάω, can be construed with a double object in the accusative (for instance, in the formula μιν... ἔπεα... προσήυδα). Therefore, the number of occurrences of each verb does not exactly correspond to the number of accusative objects extracted for it. Since we will be operating with object types rather than object tokens (see below), this is only marginally relevant; still, object frequencies are provided in Table 23 for the sake of completeness. The order in which the verbs are presented is the same as in Table 22, for ease of reference; it will be kept constant throughout this chapter.

Table 23: Number of objects for each verb in formulae in the target corpus

	Verb	Token frequency⁴⁷⁵	Number of objects
1	ἔχω	335	337
2	πρόσφημι	202	205
3	αἰρέω	167	167
4	προσεῖπον	163	163
5	δίδωμι	157	157
6	εἶπον	153	167
7	προσαιδάω	142	213
8	φέρω	109	110
9	βάλλω	108	110
10	εἶδον	102	102
11	τίθημι	101	101
12	ικνέομαι	100	102
13	οἶδα	87	87

⁴⁷⁵ Reproduced from Table 22.

14	ἀμείβω	86	88
15	πάσχω	84	84
16	χέω	83	83
17	αὐδάω	77	78
18	ἀγορεύω	71	71
19	φημί	59	60
20	ἄγω	57	59
21	ἐρύω	56	56
22	λύω	54	54
23	τίκτω	54	54
24	ἴημι	53	54
25	ῥέζω	50	55
26	ἀκούω	50	50

For each verb, a full list of objects was created. Each list, of course, contains many repetitions: the beginning of the list for δίδωμι is provided as an example below.

Table 24: Example list of objects for δίδωμι (truncated)

δίδωμι
κόρη
γέρας
γέρας
κῦδος
πολύς
ἄλγος
πᾶς
ἐπίκουρος
εὖχος
ἵππος
εὖχος
θυγάτηρ

As per the smaller case study described in §VI.4, however, the next stages of the analysis will be done exclusively on object types, not tokens – that is, we will be looking at the relationship between the different nouns that occur as objects of a specific verb, regardless of whether any of these nouns occurs more than once, or indeed how often they occur.⁴⁷⁶

In order to obtain lists of object types, a word-counting script written for this purpose was

⁴⁷⁶ For the reasoning behind using types rather than tokens, see e.g. Barðdal (2008) as discussed in §VI.1.2.

run on each list of tokens; the script identifies and counts all occurrences of each word type, producing a list which was then ordered by decreasing frequency.⁴⁷⁷

At the end of this process, a list of distinct object types was created for each of the 26 verbs. To address the question of whether and how formulaic and non-formulaic language usage differ, we now need to set up some similar lists of objects that can act as a baseline for comparison.

VII.1.5 *The data for comparison*

This time, the data for comparison was not extracted from a corpus built *ad hoc*, but from a pre-existing database. As we will see, this has important and somewhat necessary advantages (the size of the experiment would have made it prohibitive to build and analyse a comparison corpus along the same lines as the target corpus), but it also limits our options in terms of the content of the database itself.

The database used is the ancient Greek verb valency lexicon built by Barbara McGillivray.⁴⁷⁸ This database, itself created from the Ancient Greek and Latin Dependency Treebank, contains information about the constructions of all verbs that appear in a set of texts ranging from early epic to the third century AD.⁴⁷⁹ Specifically, it contains data about the following texts (the ones reported in square brackets were excluded from the non-formulaic sample): Aeschylus, all the tragedies (including *Prometheus*); Aesop, *Fables* 1.1-1.50; Pseudo-Apollodorus, *Bibliotheca* 1.1-4; Athenaeus, *Deipnosophistae* 12-13; Diodorus

⁴⁷⁷ This script is available, like all others, at <https://github.com/MartinaAstridRodda/dphil-thesis/blob/main/wordcount%202.0.py>.

⁴⁷⁸ The database has been used in McGillivray and Vatri (2015), which however does not contain a detailed description. A detailed paper describing the database and giving an example of its application via an abbreviated version of the analysis in this chapter is currently in preparation by myself and Barbara McGillivray.

⁴⁷⁹ For a general discussion of valency lexica for ancient languages see Jensen and McGillivray (2017), 133-135.

Siculus, *Bibliotheca Historica* 11; Herodotus, *Histories* 1; [Hesiod, *Theogony* and *Works and Days*; Homer, *Iliad* and *Odyssey*; *Homeric Hymn to Demeter*]; Lysias, *Against Alcibiades* 1&2, *Against Pancleon*, *On the Murder of Eratosthenes*; Plato, *Euthyphro*; Plutarch, *Alcibiades* and *Lycurgus*; Polybius, *Histories* 1; Sophocles, *Ajax*, *Antigone*, *Electra*, *OT*, *Trachiniae*; Thucydides, *Histories* 1.

While this is something of a motley crew of texts, it fulfils both essential principles for a baseline sample: it represents a wide selection of characteristics (genres – both poetry and prose – and chronological range), and it was put together according to criteria that have nothing to do with the aims of the present work.⁴⁸⁰ It also contains no other epic material apart from Homer, Hesiod, and the *Hymn to Demeter*, which were excluded from our data as they are already part of the target corpus, or in the case of the *Hymn to Demeter* part of the archaic epic tradition. Table 25 contains information about the respective sizes of the target corpus and the baseline corpus, broken down by individual author; this information will become relevant when comparing results from the two, especially when it comes to the productivity of individual phrases. The source for these numbers is the word counts published on the Perseus website. For authors whose works are only partially included in the baseline corpus, numbers were obtained by pasting the relevant sections into a Microsoft Word document, stripping them of all non-Greek material (that is, all numbers and Latin letters) and counting the words there; these are marked with an asterisk in the table. The word count for tragedy and for Plato also includes character/speaker tags.

Table 25: Corpus sizes

Number of words ⁴⁸¹		
Baseline	Aeschylus	41,842
	Aesop ^{*482}	10,211
	Pseudo-Apollodorus [*]	1417

⁴⁸⁰ On these criteria see Jensen and McGillivray (2017), 127-129.

⁴⁸¹ <http://www.perseus.tufts.edu/hopper/collection?collection=Perseus:collection:Greco-Roman>, last consulted April 2020.

⁴⁸² The word count for the text of Aesop, which is not on the Perseus website (but is included in the AGLDT treebanks), should be considered approximate. Titles for the fables are included.

	Athenaeus* (Kaibel)	40,980
	Diodorus Siculus* (Loeb)	23,811
	Herodotus*	29,644
	Lysias*	6347
	Plato*	5479
	Plutarch	19,911
	Polybius*	26,509
	Sophocles	42,304
	Thucydides*	23,014
	Total	271,469
Target	<i>Iliad</i>	111,862
	<i>Odyssey</i>	87,185
	<i>Theogony</i>	7,040
	<i>Works and Days</i>	5,860
	Total	211,947

As the table shows, while the baseline corpus is slightly larger than the target one, the two corpora are fairly well balanced in size. We will, in any case, account for the different corpus sizes when discussing the productivity of verbs in each corpus (§VII.2.5).

I had access to data tables extracted from the valency lexicon.⁴⁸³ These tables contain information about all instances in which a specific verb appears in the corpus, and how it is construed; specifically, for our purposes, they list all the cases in which the verb is construed with the accusative and give the lemma of the accusative object. Effectively, they contain all the information needed to create a list of TrV + Obj pairs in the baseline corpus. This information was extracted from the table by hand (a much less cumbersome process than one might expect, due to the pre-processed and richly tagged nature of the data): this involved selecting all instances of an OBJ tag marked with an accusative case (for instance, active_OBJ[accusative]),⁴⁸⁴ independently of any other elements in the tag set, and then retrieving the lemma corresponding to the object in the accusative from the valency tables.⁴⁸⁵ Finally, as detailed in §VII.1.4, the list of object tokens for each verb was reduced to a list of object types (distinct objects) using the same word-counting script.

⁴⁸³ My thanks to Dr Barbara McGillivray for extracting the data and sending me the tables.

⁴⁸⁴ Note that, since both the data for the epic target corpus and the data for the valency lexicon come from the AGLDT corpus, the syntactic tags are the same.

⁴⁸⁵ All verbs were considered independently of their diathesis (active, middle, middle-passive; there is only one case of a passive verb tagged as having an accusative object, an occurrence of *τίκτω* in Plutarch, which I discarded as an obvious error).

VII.1.6 The final TrV + Obj lists

We now have two lists of object types for each of the 26 target verbs: one from our target sample, which is exclusively made of formulae in Homer and Hesiod, and one from our baseline sample, which ranges chronologically from Aeschylus to Athenaeus. Table 26 lists the length of these object lists for each verb (the order is still the same as in Table 22). *πρόσφημι* and *ἀγορεύω* have been excluded *a priori* because they never occur with a direct object in the baseline sample.

Table 26: Number of distinct objects (object types) for each verb in the target and baseline corpus

	Verb	Target sample	Baseline sample
1	ἔχω	91	480
2	πρόσφημι		
3	αἰρέω	56	80
4	προσεῖπον	13	4
5	δίδωμι	58 (59*)	87
6	εἶπον	28	46
7	προσαυδάω	23	3
8	φέρω	41	116
9	βάλλω	49	24
10	εἶδον	43	81
11	τίθημι	45	88
12	ἰκνέομαι	35	6
13	οἶδα	29	47
14	ἀμείβω	8 (9*)	20
15	πάσχω	6	37
16	χέω	13	13
17	αὐδάω	3 (4*)	7
18	ἀγορεύω		
19	φημί	4	17
20	ἄγω	33	91
21	ἐρύω	10	1
22	λύω	14	29
23	τίκτω	13	24
24	ἵημι	19	10
25	ρέζω	16	2
26	ἀκούω	11 (12*)	45

Lists marked with an asterisk also include one element tagged as ‘unknown’ in the Perseus lemmatising system (that is, an element that could not be recognised at any stage of the tagging process), which effectively needs to be subtracted from the count of actual objects.

Note how the number of tokens and the number of types need not be proportional in any way: verbs which had more tokens often have many fewer types than verbs which had fewer tokens. This explains why, if the verbs were ranked by the number of object types instead of by number of occurrences, the ordering in the table above would be quite different.

To ensure we have a large enough number of object types for statistical analysis, verbs which have fewer than 10 object types in either the target or baseline list were excluded. These verbs are italicised and grey-shaded in the table above; some of them are epic compounds such as *προσεῖπον* and *προσσυδάω*, while others are rare outside of early epic (*ἐρύω*, *ῥέζω*), or have changed their patterns of usage with direct objects (*ικνέομαι*). This brings the number of available verbs down to 15.

VII.2 Analysis and results

VII.2.1 Setting up the analysis: research questions

At this point in the experiment, we have two lists of accusative object types (one extracted from the formulaic corpus, as defined in §VII.1.2, and one from the comparison corpus, as defined in §VII.1.5) for each of the 15 verbs, for a total of 30 lists. We now move back to familiar territory, in that we are returning to the experimental setup described in §VI.4: we want to measure how tightly clustered (or how widely scattered) these objects are in the semantic space, or in other words what their range of meaning is, defined under a Distributional Semantics approach.

To obtain this measure, the same distance-measuring script described in §VI.4.1⁴⁸⁶ was run on each of our 30 lists of object types. The script, as before, returns the value of the cosine distance between each object and the centroid of the whole distribution of objects for the verb in question; in other words, it provides a distribution of the distances from the objects to their centre of mass. The characterising values of this distribution (mean, median, variance) can be used as measures of the semantic coherence of each set of objects. More importantly for our purposes, the two distributions of distances for each verb (formulaic and not) can now be compared with each other.

An extremely important element to note before we begin the analysis is that while the variance of each distribution of distances is straightforwardly interpreted (higher variance means a more irregular distribution of distances, lower variance a more regular one), the distances we are working with are *cosine* distances. The cosine distance is the cosine of the angle between two vectors in the multidimensional semantic space; as such, in a positive space⁴⁸⁷ such as that of the vector space model, it ranges between 0 and 1, with 1 being *maximum similarity* (parallel vectors⁴⁸⁸) and 0 being *maximum dissimilarity* (orthogonal vectors). Therefore, a median cosine ‘distance’ that is numerically higher means a *higher degree of similarity*, not the opposite – as would happen with other distance measures such as the Euclidean distance introduced as an example in §VI.2.3. This did not affect our discussion of results in the previous chapter, where we were just seeking to group formulae based on similarities and differences in their semantic behaviour, without quantifying the effects of said behaviour; in this chapter, however, we are seeking to draw correlations between specific parameters and the semantic openness of TrV + Obj formulae, so we will need to remember that higher cosine distance implies lower, not higher, semantic

⁴⁸⁶ Also available in the general thesis repository at https://github.com/MartinaAstridRodda/dphil-thesis/blob/main/proximity_1.4_centroid.py.

⁴⁸⁷ A space with no negative values.

⁴⁸⁸ The cosine of an angle of 0° (parallel vectors, maximum similarity) is 1; that of a 90° angle (orthogonal vectors, minimum similarity) is 0.

openness. To avoid the confusion that would predictably arise from this, from here on I will be using the term ‘cosine *similarity*’ rather than cosine distance: this should be helpful in remembering that higher distance equals higher similarity.

The comparison between distributions, in any case, needs to be guided by specific questions: just randomly comparing distributions will almost certainly lead to the discovery of at least *some* differences that are mathematically significant, but these may have no explanatory value, or we may not know how to explain them, unless we start from a specific hypothesis that has a bearing on these values. What we want to do is focus on whether certain factors, which we already have reason to believe may correlate with semantic coherence, affect the distributions.

The research conducted in chapter VI, together with the literature reviewed there, led us to highlight three key factors that may be relevant to a discussion of the semantic coherence of accusative objects:

1. Formulaic status: do the objects of formulaic TrV + Obj phrases (from the early epic sample) show a wider or more restricted semantic range than the objects of non-formulaic ones (from the baseline sample)?
2. Frequency: we have seen that frequency influences the syntactic flexibility of formulae (see §IV.3.2): more frequent formulae are more flexible, although this is not the only factor at play. Does frequency have a similar effect on the semantic openness of formulae? What about expressions in the baseline sample? If there is an effect of frequency, is it stronger in formulaic or non-formulaic usage?
3. Productivity: based on existing literature on verbal constructions in English and Icelandic (Barðdal 2008 and Perek 2016, as discussed in §VI.1.2), we expect the semantic openness of a construction to correlate positively with its productivity

through time: broadly speaking, more semantically open constructions continue to extend their usage to new fillers. Is this the case for our sample?

Note that for numbers 2 and 3 we expect the causal links to operate in opposite directions: we expect frequency to influence semantic openness, in the same way as we have seen it to promote syntactic openness in §IV.3.2, but we expect the semantic openness of each verbal construction, in turn, to influence its productivity in time. The latter hypothesis is based on external evidence from the behaviour of constructions in later languages.

In the light of these considerations, let us think about how we can encode these three variables for the statistical analysis.

1. Formulaic status as defined for the purposes of this analysis (on the basis of the formulaic editions of Homer and Hesiod) is a binary variable (formulaic or non-formulaic), and is in practice already encoded in our data: the 15 lists of objects that come from the target sample (formulaic TrV + Obj pairs in Homer and Hesiod) reflect formulaic usage,⁴⁸⁹ while the remaining 15 lists that come from the baseline sample reflect non-formulaic usage by definition. This can therefore easily be included in our analysis.
2. Frequency information for the target corpus is contained in Table 23. Note that (a) we are back to considering token frequency, not type frequency here, as token frequency was found to be the relevant predictor in the case of syntactic flexibility; (b) in order to exclude the influence of variable 1, formularity, we will only be looking at the effects of frequency in the target sample, not in the baseline. This is also in line with the results in §IV.3.2, where frequency effects were observed in the formulaic sample.

⁴⁸⁹ At least under Pavese and Boschetti/Venti's definition, which we agreed to adopt for this study: see §VII.1.3.

3. Productivity can be assessed by comparing the number of object types each verb has in the baseline corpus to the target corpus: as the baseline is, in its entirety, chronologically later than the target, we can see which verbs acquire more object types, which verbs maintain more or less the same number of them, and which verbs reduce their usage to a smaller number of objects. These categories can already be intuited by simply looking at Table 26, although numbers need to be adjusted to account for the fact that the two corpora are of slightly different sizes. In this case, we will be looking at whether semantic range in the target sample influences productivity in the baseline sample: we are looking for the effect of this parameter across time, rather than within one corpus.

In short, we can phrase our research questions (taking into account our previous assumptions on the direction of the correlation) as follows:

1. Are formulaic TrV + Obj pairs more or less semantically open than non-formulaic ones?
2. Are formulaic high-frequency TrV + Obj pairs more or less semantically open than formulaic lower-frequency ones?
3. Do formulaic TrV + Obj pairs which are more semantically open go on to be more or less productive in the non-formulaic sample?

VII.2.2 Setting up the analysis: the productivity and frequency variables

We can codify both the frequency and the productivity variables as discrete bins: that is, we can divide our sample of verbs into chunks depending on their frequency or productivity, respectively. As we will see below, this is easier for productivity than frequency.

In the case of productivity, we want to distinguish between those verbs whose number of object types *decreased* in the baseline sample (lower productivity), those whose number of object types *stayed more or less the same* (no change in productivity), and those whose number of object types *increased* (higher productivity). We can number these groups from 0 to 2. Table 27 contains the relative productivity scores for all 15 verbs; these scores are computed by dividing the number of object types by the size of each corpus, and then dividing the normalised number in the baseline corpus by the normalised number in the target corpus, which gives us a relative productivity ratio. Verbs with a productivity score below 1 (any reduction in the number of types) were assigned to group 0 = lower productivity; verbs with a score above 1.5 (significant increase in number of types) were assigned to group 2 = higher productivity, while verbs with a score of 1-1.5 were assigned to group 1 = little to no change.

The reader may notice that these groups are not symmetrical. The key issue here is how we define 'little to no change,' when of course none of the verbs have *the exact same* (relative) number of object types in the two corpora. I consider a reduction in the number of types, independently of size, to be more significant than a small increase in them, as it signals that the verb in question is becoming less productive; a small increase in productivity, on the other hand, can just be an effect of the wider selection of genres and topics present in the baseline corpus. While an increase of 1.5 times (the upper limit of group 1) is certainly not small, reducing the width of group 1 (for instance, limiting it to an increase of 1.25) would mean including an even higher range of variation in group 2, which already ranges from an increase of 1.5 times to 3.4.

The token frequency of each of our verbs, on the other hand, is listed all the way back in Table 22 and Table 23. The problem with frequency data is that it can be less easily reduced to discrete bins: frequency does not easily lend itself to qualitative distinctions in the same way as productivity does. In other words, while 'increased productivity' and 'decreased

productivity' are relatively uncontroversial classifications, 'higher' and 'lower frequency' would require a much more arbitrary definition.

What we can straightforwardly do, however, is order our verbs based on their token frequency in the target sample; indeed, we have already done so, and this ordering appears in all of the tables above. Frequency, therefore, will simply be encoded with values 1-15; to ensure that these values go in the same direction as the productivity variable, where the highest productivity was encoded as 2, and the lowest productivity as 0, the verbs will be numbered in decreasing order. In other words, ἔχω is now verb number 15, and ἀκούω is number 1.

Table 27: Values of the frequency and productivity variables for each verb

Frequency ranking	Verb	Normalised type frequency (target) ⁴⁹⁰	Normalised type frequency (baseline)	Relative productivity	Productivity group
15	ἔχω	0.429	1.768	4.118	2
14	αἰρέω	0.264	0.295	1.115	1
13	δίδωμι	0.274	0.320	1.171	1
12	εἶπον	0.132	0.169	1.283	1
11	φέρω	0.193	0.427	2.209	2
10	βάλλω	0.231	0.088	0.382	0
9	εἶδον	0.203	0.298	1.471	1
8	τίθημι	0.212	0.324	1.523	2
7	οἶδα	0.137	0.173	1.265	1
6	χέω	0.061	0.048	0.781	0
5	ἄγω	0.156	0.335	2.153	2
4	λύω	0.066	0.107	1.617	2
3	τίκτω	0.061	0.088	1.441	1
2	ἵημι	0.090	0.037	0.411	0
1	ἀκούω	0.052	0.166	3.194	2

Something that shows quite nicely in Table 27 is how productivity and frequency in the formulaic corpus are not at all correlated: each productivity group spans both higher and

⁴⁹⁰ As the numbers in this column and the next are extremely small (we are, after all, dividing the number of types, which has an order of magnitude of 10^2 , by the corpus size, which has an order of magnitude of 10^6), I have multiplied them by 1000 for the sake of readability. As both target and baseline data have been multiplied by the same factor, this makes for no change in the relative productivity scores. Numbers appear rounded to 3 decimal digits in the table, but all calculations were done on unrounded numbers.

lower frequency verbs, and the most and least frequent verb, respectively *ἔχω* and *ἀκούω*, are both in the highest productivity group. This is important, as it confirms that examining the two variables separately is not a redundant choice: each of them has its own individual behaviour, and they do not seem to influence each other.

VII.2.3 *Analysing the data: formulaic vs non-formulaic*

In order to answer our first question – whether there are significant differences in the semantic behaviour of the objects of formulaic vs non-formulaic transitive verbs –, we need to compare the two distributions of cosine similarities, formulaic and non-formulaic, for each verb. In order to do this, we will look at two values: the median cosine similarity from the centroid and the variance of the distribution. The median indicates how widely scattered the points (objects) are from the centroid of the distribution.⁴⁹¹ The variance measures the *range* of cosine similarities between the objects and the centroid: are they all fairly equidistant, or are some much closer and some much further away?⁴⁹² A higher variance means that objects are more irregularly spread out.

All of these characterising factors are represented in the chart in Figure 8,⁴⁹³ a box-and-whiskers plot of the distribution of cosine similarities for each verb, respectively in the baseline sample (white shading) and in the target sample (dotted shading). Box-and-whiskers plots, or boxplots for short, are an intuitive way to visualise how widely spread a distribution is. The line in the centre of each box is the median cosine similarity. The box and the whiskers represent the quantiles of the distribution: the box contains 50% of the values (25% on each side), while the whiskers extend over the remaining 50%. The whiskers

⁴⁹¹ I use the median, not the mean, as the former is less affected by the presence of any outliers. On the use of the median in exploratory data analysis, see e.g. Miller and Miller (2010), 155-160.

⁴⁹² For further examples and discussion see §VI.2.3.

⁴⁹³ The plots and median and variance values were obtained with Microsoft Excel.

therefore do not represent the variance, but they are a decent visual proxy for how widely spread the values are.

Table 28 contains the numeric data for each verb.⁴⁹⁴ For each pair, I have tested the significance of the difference between formulaic and non-formulaic objects using the Kolmogorov-Smirnov test, a non-parametric test that makes no assumption about the underlying normality of the distributions.⁴⁹⁵ Given the sparseness of the data, I have set two significance thresholds: highly significant (**, $p < 0.05$) and marginally significant (*, $p < 0.1$).

Table 28: comparison between formulaic and non-formulaic distributions of objects in the semantic space

Verb	Median similarity		Variance		Significance
	Formulaic	Baseline	Formulaic	Baseline	
έχω	0.427	0.482	0.0109	0.0100	* ($p = 0.084$)
αίρέω	0.407	0.455	0.0172	0.0123	($p = 0.240$)
δίδωμι	0.447	0.430	0.0077	0.0130	($p = 0.911$)
είπον	0.324	0.329	0.0116	0.0188	($p = 0.538$)
φέρω	0.422	0.468	0.0091	0.0086	($p = 0.414$)
βάλλω	0.446	0.416	0.0139	0.0147	($p = 0.716$)
είδον	0.371	0.445	0.0137	0.0103	($p = 0.106$)
τίθημι	0.410	0.439	0.0133	0.0093	($p = 0.723$)
οίδα	0.414	0.394	0.0070	0.0227	($p = 0.537$)
χέω	0.364	0.472	0.0096	0.0027	** ($p = 0.034$)
ἄγω	0.392	0.456	0.0095	0.0099	($p = 0.329$)
λύω	0.314	0.470	0.0080	0.0072	($p = 0.168$)
τίκτω	0.270	0.405	0.0122	0.0062	* ($p = 0.052$)
ἵημι	0.453	0.307	0.0104	0.0027	* ($p = 0.053$)
ἀκούω	0.413	0.375	0.0059	0.0118	($p = 0.787$)

⁴⁹⁴ As always, numbers in the table are rounded, but all calculations were made on unrounded data.

⁴⁹⁵ See Baayen (2008), 77-79. The test was performed in R (R Core Team 2017).

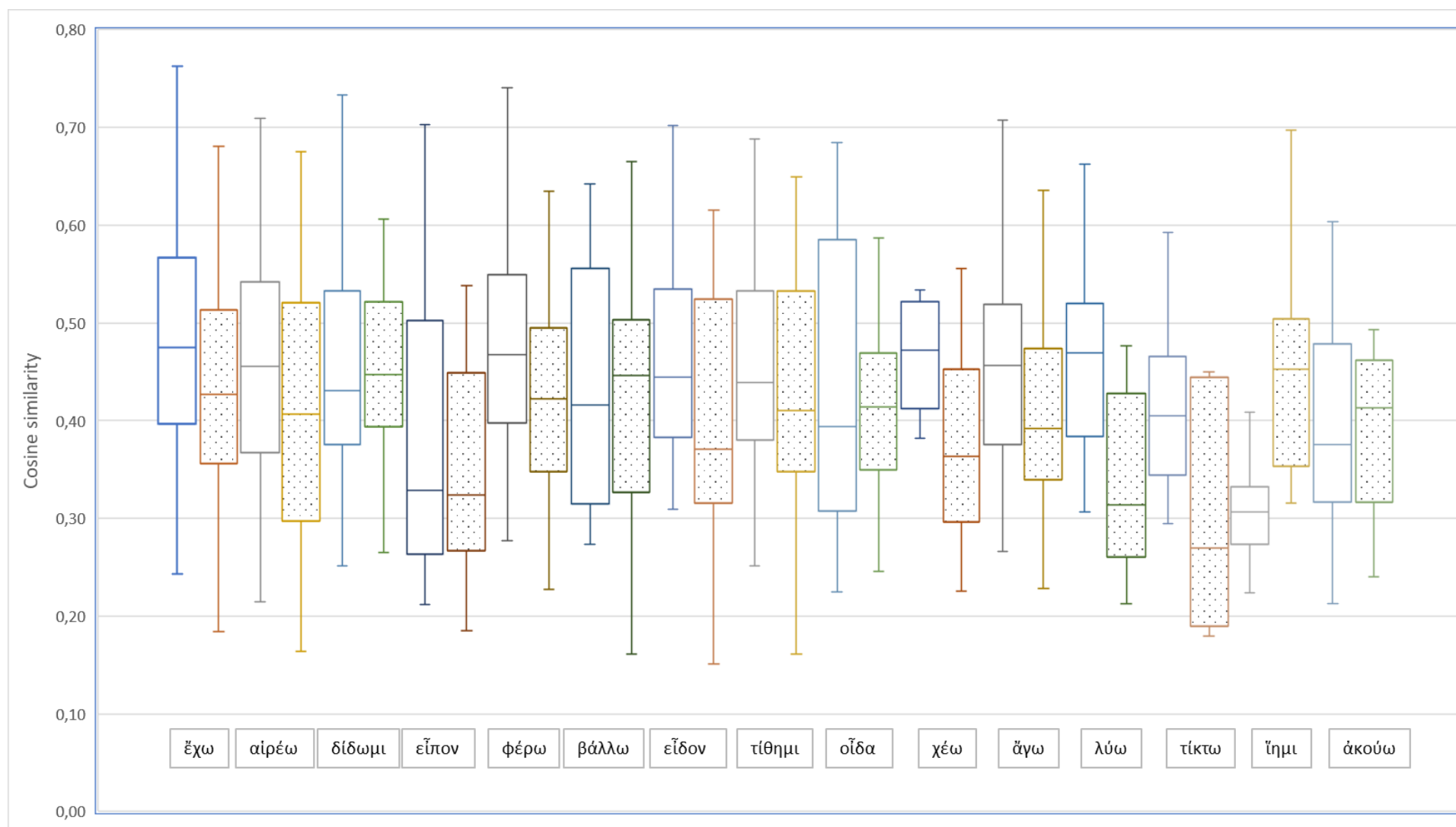


Figure 8: Box-and-whiskers plot of object similarities in formulaic (dotted) and non-formulaic (white) pairs

The last column in Table 28 shows that there are only four verbs (just over 1/4 of the total) for which a significant difference can be detected between the semantic range of formulaic and non-formulaic objects. One of them is our most frequent verb, ἔχω. For three of these verbs, ἔχω, χέω and τίκτω, the baseline sample (non-formulaic objects) has a higher median cosine similarity but lower variance: in other words, the non-formulaic accusative objects of these verbs are more tightly clustered around a central point and more regularly distributed. The last verb, ἴημι, shows the same tendency towards a lower variance in the non-formulaic sample, but it is the formulaic objects that show the higher median similarity compared to the non-formulaic ones. These differences are summarised for convenience in Table 29 below.

Table 29: Summary of significant differences between formulaic and non-formulaic objects

Verb	Median similarity	Variance
ἔχω	Baseline > formulaic	Formulaic > baseline
χέω	Baseline > formulaic	Formulaic > baseline
τίκτω	Baseline > formulaic	Formulaic > baseline
ἴημι	Formulaic > baseline	Formulaic > baseline

This is a small number of significant differences in a sample of 15 verbs, most of which do not show any difference in the range of their formulaic and non-formulaic accusative objects. We can, at least, try to use these significant results to see if the trend holds even when results are not significant. Median similarity is indeed higher in the baseline sample (as it was for three out of four of our significant cases) for two thirds of our verbs, and variance is higher in the formulaic sample for another two thirds; this, however, leaves one third of the sample going in the opposite direction, which gives us a far less than clear-cut result. While we can say that there is a potential tendency for median similarity to be higher and variance to be lower (objects to be more evenly spread in the semantic space) in non-formulaic objects, this is still a small effect, and most of the significant results are observed in lower-frequency formulae. For a discussion of this result, see below, §VII.3.3.

VII.2.4 *Analysing the data: frequency effects*

Our second potentially relevant variable is the frequency of formulae themselves. For this part of the analysis, we are merely looking at correlations between the frequency ranking of each formula in the target sample and the semantic flexibility of its objects; the latter is encoded by the same two parameters that were introduced in §VII.2.3 and provided in Table 28, that is, median cosine similarity and variance of the distribution.

An exploratory representation of the data is provided in the two bar charts below. These represent, respectively, the values for the median and the variance for each formula, in decreasing order of frequency. (These are the same values provided in Table 28, exclusively in the ‘Formulaic’ columns.) The charts offer an easy and straightforward way to look for statistical tendencies: the overall shape of each chart will indicate whether the median (or variance) tends to decrease with frequency, increases, or does not appear to follow any trend.

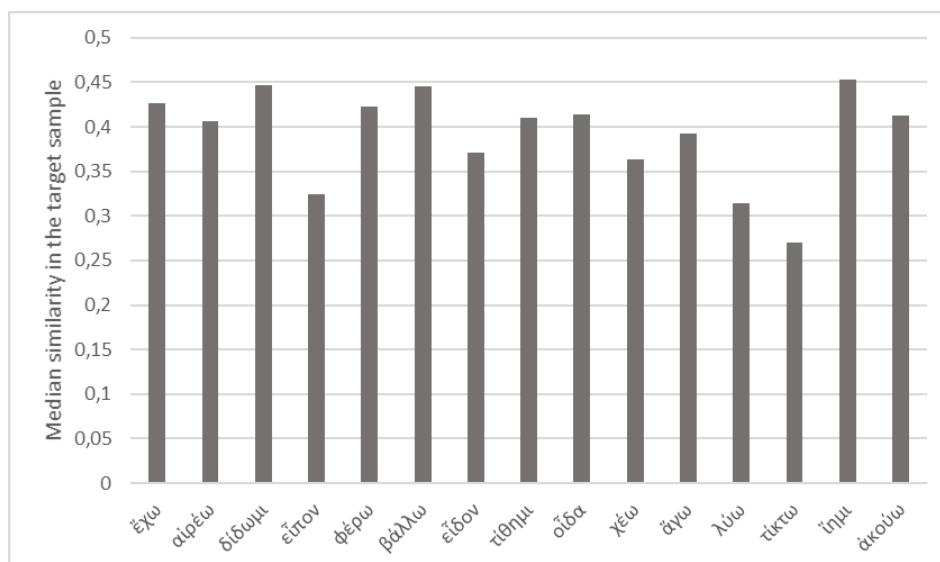


Figure 9: Median similarity in the target sample, ordered by frequency

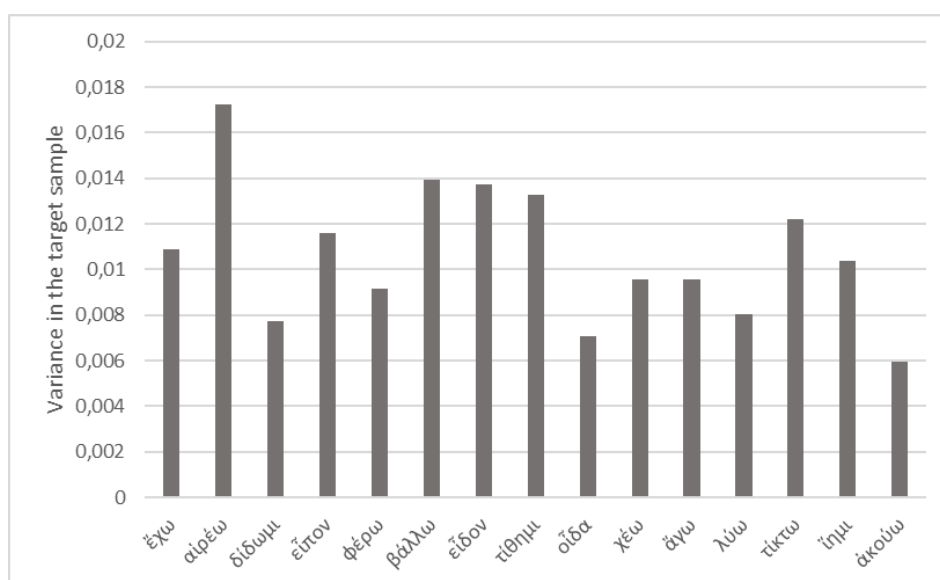


Figure 10: Variance in the target sample, ordered by frequency

Similarly to the results in §VII.2.4, there does not appear to be any meaningful trend or correlation. A quantitative check confirms this: the Spearman rank correlation test (a non-parametric correlation test, again making no assumptions about the shape of the distributions) between the frequency ranks (15-1) and either the median or the variance returns no significant results and extremely weak correlation coefficients (for median and variance, respectively: $p = 0.368$ and $p = 0.174$, $\rho = 0.250$ and $\rho = 0.371$). Unlike what we

could see for the syntactic flexibility of Adj + N formulae, therefore, there is no evidence of a frequency effect governing the semantic behaviour of TrV + Obj formulae.

VII.2.5 Analysing the data: productivity

To see if semantic flexibility has an effect on productivity, we will once again look at flexibility measures in the formulaic sample (that is the median and variance of object distributions for formulaic TrV + Obj pairs); this time, we will be cross-referencing them with the productivity 'bins' we defined in Table 27. We can also, however, rank the data in a similar way to what we did for the frequency analysis: the formulae in the two bar charts below have been sorted in decreasing order of productivity, while the bar colour indicates the group to which they belong (darker grey for group 2 down to lighter grey for group 0). After looking at the data for individual formulae, we will then move on to looking at averages by category.

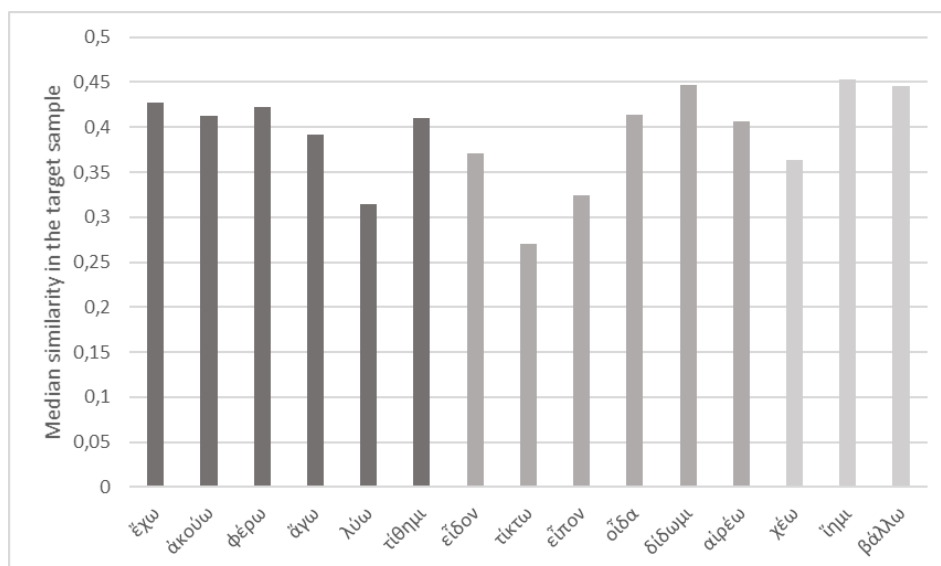


Figure 11: Median similarity in the target sample, ordered by productivity. The shading represents productivity 'bins.'

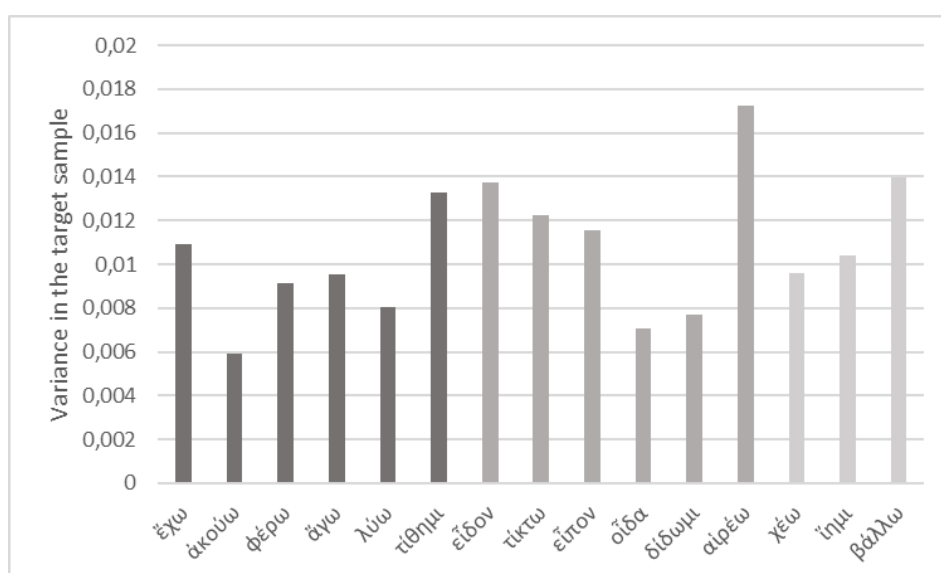


Figure 12: Variance in the target sample, ordered by productivity. The shading represents productivity 'bins.'

Once again, the bar charts do not look particularly encouraging. There certainly does not seem to be any correlation between median object similarity and productivity, and the same goes for variance, which varies wildly irrespectively of the ranking of formulae. Spearman correlation tests with median and variance as the independent variable and productivity as the dependent variable⁴⁹⁶ confirm this qualitative observation: there are no significant

⁴⁹⁶ For the reasons for this choice, see §VII.2.2.

results and very weak correlation coefficients (for median and variance, respectively: $p = 0.206$ and $p = 0.524$, $\rho = -0.346$ and $\rho = -0.179$).

Does this picture change at all if we look at the behaviour of groups of verbs rather than individual verbs? We would expect at the very least the behaviour of those verbs whose productivity increases in the baseline sample (group 2) to differ from that of verbs whose productivity has decreased (group 0). To address this question, we can turn to the averages for each group. The results for these are represented via the usual bar charts below. The scale of the y axis in these charts has been kept the same as the scale in the previous two charts: this will make it easier to compare the results and assess the magnitude of any differences.

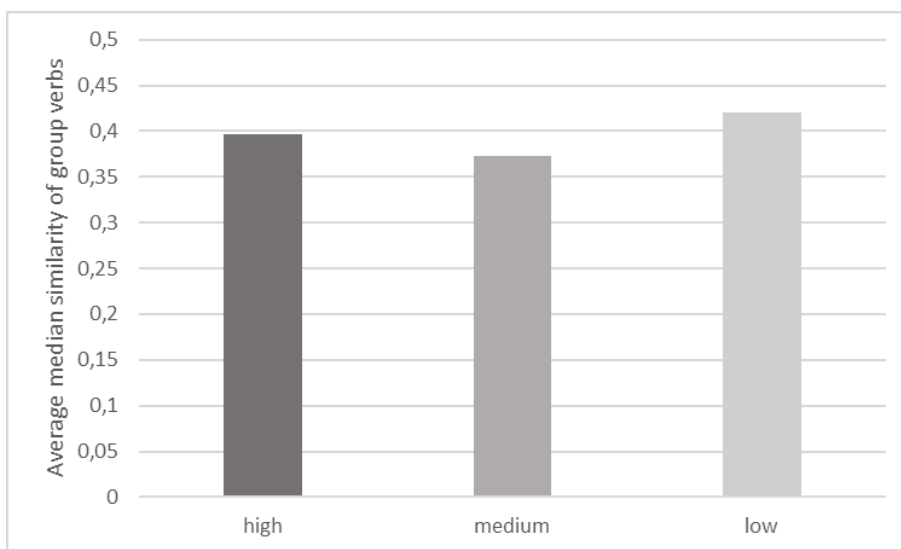


Figure 13: Average median similarity by productivity group

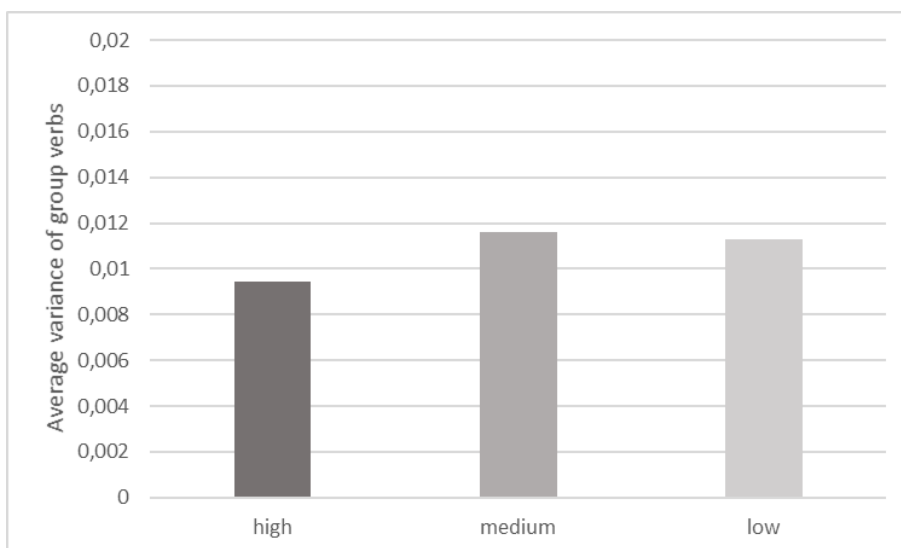


Figure 14: Average variance by productivity group

The group averages utterly fail to support the hypothesis we have just stated, that is, that there should at least be some difference between the highest and the lowest frequency group. Instead, we find that the only group which distinguishes itself among the results is the medium frequency one. In other words, there is no trend to be found anywhere in the results, exactly as we have seen when looking at individual verbs. In the same way as the frequency of a phrase does not seem to have any effect on its semantic range, semantic range and productivity do not appear to be connected.

VII.3 Discussion and conclusions

The picture from §§VII.2.3-VII.2.5 is made of very few significant differences, and none at all for two out of three of our initial hypotheses. The possibility that the data will not support the initial hypotheses is always present in a quantitative study; this answer is, in itself, not the same as ‘the data is meaningless and/or just made of random noise.’ Indeed, we have evidence from the small-scale study in §VI.4 that quantitative semantic data can in fact be used to compare epic formulae, although in a different kind of format, and the validation process for the DSM applied here shows that the data in it are worth considering. We are, in short, left with the task of making sense of the negative results in this chapter.

We had approached the data with a set of questions regarding formularity, frequency, and productivity; we expected each of these factors to potentially influence or be influenced by semantic openness, defined as the range of meanings shown by the objects of each verb in our target list, quantified via Distributional Semantics. The questions are repeated below from §VII.2.1, and the answers are summarised in Table 30.

1. Are formulaic TrV + Obj pairs more or less semantically open than non-formulaic ones?
2. Are formulaic high-frequency TrV + Obj pairs more or less semantically open than formulaic lower-frequency ones?
3. Do formulaic TrV + Obj pairs which are more semantically open go on to be more or less productive in the non-formulaic sample?

Table 30: Summary of results for the semantic analysis

Parameter	Observations
Formularity	For 4/15 verbs: significant differences and higher variance in F sample. This tendency is partially upheld in non-significant differences as well.
Frequency	No significant effect
Productivity (single)	No significant effect
Productivity (group)	No significant effect

We will address the outcomes of the data analysis in a slightly different order below, as the thorniest result is the one concerning formularity, and its implications deserve to be discussed at length after more briefly doing away with frequency and productivity.

VII.3.1 Semantic range and frequency

There is absolutely no evidence that frequency affects the semantic range of formulaic objects. This analysis, conducted exclusively on the formulaic sample, was aimed at comparing semantic and syntactic flexibility, given that frequency effects were clearly present in the syntactic flexibility data for formulaic Adj + N pairs. When discussing syntactic flexibility in the latter, we reached the conclusion that this was promoted by usage: formulae that are used more frequently need to be adapted to a wider range of contexts, which requires them to inflect, change their number, and so on.⁴⁹⁷ The most frequent types of change are those that do not affect the internal shape of a formula, but instead have to do with the addition of external material – that is, expansion is significantly more easily available than inflection, changes in word order, and separation.

We could have expected a similar trend for semantic variation, as verbs that are used more often may need to accommodate a higher number of objects with a higher range of meanings. This, however, does not show in our data, where semantic variation seems instead to happen independently of frequency. This hints at a higher entrenchment of TrV + Obj pairs as somewhat lexicalised items, which do not easily acquire a wider range of objects the moment they are exposed to a wider range of contexts.

⁴⁹⁷ See §IV.3.2.

VII.3.2 *Semantic range and productivity*

Our initial prediction for this part of the analysis was for semantic openness to have a positive effect of productivity, especially for more frequent verbs; this would have been in line with the findings of Barðdal (2008), as discussed in §VI.1.2, who reports that the most productive verbal constructions in her sample are those who have either high type frequency and high semantic openness or low type frequency and low semantic openness. Similarly, but with more detail about the shape of the semantic space, Perek (2016) notes that the best predictor of productivity is not whether a construction has a high number of filler types, but how well these types cover one or more areas in the semantic space. In other words, having a high number of fillers that are very irregularly scattered does not seem to trigger processes of analogical extension of a construction's usage.

The tendencies shown in our data (§VII.2.2), both for individual verbs and for groups of formulae, are not in line with these expectations. Verbs in the lowest and highest productivity groups have a similar median similarity and variance on average, indicating that there is no correlation between productivity and any of the parameters relative to semantic coherence. It should be noted that our analysis of the role of productivity does not control for the type frequency of each verb, and therefore our results do not address Barðdal's type frequency argument, only the role of semantic coherence.⁴⁹⁸ Perek's concept of coverage is also not fully addressed here, unless we view variance as an index of whether the objects in question are more or less regularly spread: higher variance, in this case, would indicate objects that are more irregularly scattered around the centroid. In relation to Perek's concept of coverage, verbs whose objects are more regularly distributed around their centroid, that is, verbs that have lower variance, cover the semantic space in a more

⁴⁹⁸ Neither do our frequency results, which are based on token frequency.

uniform way, and should go on to be diachronically more productive; once again, however, none of these tendencies are borne out by our data.

This result is interesting in itself, because it marks a difference between our TrV + N pairs and similar constructions in present-day languages. The question at this stage becomes whether this lack of similarity is due to the effect of formulaic usage (but, as we will discuss in more detail below, formulaic status does not correlate with semantic range apart from a very few cases), or whether this is a sign that the observations of Barðdal and Perek do not apply to ancient Greek. The latter hypothesis, however, cannot be proven by our study, which was focused specifically on formulaic pairs and only considered a portion of Greek texts; it would require an attempt to repeat Barðdal's and Perek's analysis on an ancient Greek dataset comparable to their original ones. A study of this kind could be a straightforward development of my work here, and could rely on similar methods (and indeed on the semantic spaces that have been built for this study and described in §VI.3). In this sense, while our results here do not align with any of our initial hypotheses, they do open up some interesting methodological avenues.

VII.3.3 Semantic range and formularity

The very few significant differences that exist between pairs of formulaic vs non-formulaic TrV + Obj constructions were summarised in Table 29. This summary is useful in making it clear that even these very few significant results are not entirely coherent: while the distances between formulaic objects are always more irregularly distributed than their non-formulaic counterparts (higher variance), our clearest indicator of semantic openness, the median similarity, behaves differently depending on the individual verb. Moreover, three out of the four verbs that show any significant differences are among the least frequent in our sample, which makes the effects observed even less salient for interpretation.

The lack of significant differences between most pairs of formulae vs non-formulae is quite striking. To understand the implications of this, we need to go back to the definition of formula that has been underlyingly present throughout this work, both in the syntactic and semantic analysis, that is, the one originally introduced by Hainsworth (1968). In §II.2.1, we discussed at length the implications of the idea that formulae are structures that are characterised by a “degree of mutual expectancy” between their constituent parts:⁴⁹⁹ on a cognitive level, what allows us to identify a formulaic structure is the fact that the elements that make it up occur together with a much higher frequency than we would otherwise expect from two (or more) random words or phrases. This definition allowed us to introduce a series of comparisons between the behaviour of formulae and that of bound expressions and idioms, which are also defined through a similar level of mutual expectation between their parts.

All of these definitions rely on a strong assumption of lexical identity: a formula is made of *the same words* occurring together, and while modifications like inflection and even the addition of new material is allowed under this definition, lexical substitution of one word or phrase for a synonym is not.⁵⁰⁰ This same principle of lexical identity, in a purely quantitative form, lies behind the definition of formula as a repeated meaningful phrase adopted by Pavese and Venti (2000) and Pavese and Boschetti (2003), and used to identify formulae in this chapter (see §VII.1.3). But what if some formulae allowed for one of their component parts to be substituted with a synonym, a word or set of words that are ‘close enough’ in meaning to the original formulaic element?

In the case of our target sample, that is, formulae made of a transitive verb together with its object, what we would see if lexical substitution were possible would be a series of

⁴⁹⁹ Hainsworth (1968), 36 (and further 35-39).

⁵⁰⁰ Further on all this see §III.1, which also discusses how this phenomenon is paralleled in idioms and bound expressions.

clusters of potential object types, that is, synonyms or words that are very closely related in meaning and which can be substituted for each other.⁵⁰¹ We would expect to see these clusters more often in formulaic rather than non-formulaic TrV + Obj pairs; we would, indeed, expect higher semantic similarity from formulaic TrV + Obj pairs, as the formulaic 'bond' would introduce restrictions on their semantic range. We do not see any of this in the data, where our few significant results mostly have a baseline similarity that is higher than the similarity in the formulaic sample, the opposite of what we would have expected if things were as outlined above. What we see, instead, is a situation where, while each *individual* TrV + Obj pair is a (lexically invariable) formula, these pairs do not generally appear to impose semantic restrictions on other formulae. There is, in short, no evidence of clusters of 'semantically related' formulae among the TrV + Obj pairs.

This is somewhat surprising, as every discussion of formulaic systems from Parry onwards⁵⁰² has stressed the fact that having a range of expressions that are similar in meaning but have different metrical shapes (a result which could be easily obtained by varying lexical items and using synonyms or near-synonyms) is in fact a desirable trait for oral composition, as it creates a system of phrases that can be used in (potentially) every possible context in the hexameter verse. Why do we not see this in our data? It is possible that things work differently for TrV + Obj expressions (Parry's conclusions, and others', are mostly based on either Adj + N pairs or speech introductions), and that our definition on formularity, based on simple repetition, does not capture some subtleties in the actual relation between verbs and objects which could help explain our results. It is also possible that a different approach to the data analysis would reveal a different pattern – for

⁵⁰¹ On synonymy in formulaic diction, the key references are Paraskevaides (1984), who follows a strictly Parryian approach and seeks (much like Friedrich 2007) to evaluate the actual extent to which formulaic economy is respected, and E. J. Bakker and van den Houten (1992), who argue that pairs of formulae that appear metrically and semantically equivalent can actually be shown to have subtle differences in their meaning and usage. None of this has much of a bearing here.

⁵⁰² See chapter I for much detail on this topic.

instance, if we set out to look for individual clusters of closely-related words among the objects of a formulaic verb, rather than measure their semantic proximity to a centroid in the semantic space. All of these avenues remain open for further analysis. What we can say for certain at the end of this chapter is that there appears to be more to be explored when it comes to the semantics of verbal constructions in early Greek epic.

VIII Conclusions: How it started, how it is going, and some ideas for further work

VIII.1 How it started: initial questions, theoretical grounding, and goals

This thesis started out with an object of study (linguistic variation in the language of early Greek epic, which is an oral tradition characterised by formulaic structures), a set of questions of interest, mainly related to what formulae can do (how do they behave? Is there any salient difference in how they behave compared to non-formulaic material? Do these differences tell us something about how formulae were processed or used to compose?), and, correspondingly, a set of questions we were not interested in, mainly related to what formulae are. To this we added a toolbox of linguistic comparisons (to idioms and to constructions on a more general level) that would give us an angle from which to approach these questions on a practical level, mainly through quantitative analysis. We also had some basic guiding principles that we wanted to respect: not to limit ourselves to a single ‘type’ of formula, but to extend the study to different cases; the application of different approaches; and a goal to offer the most complete account possible of each category, without limiting ourselves to a single set of examples, or even to one poem or one author, but involving the entirety of the tradition as it is extant or as close to that as possible.

These initial questions developed into a usage-based account of formularity – or rather, given that this is obviously a colossal endeavour, the beginnings of one. Of several potential aspects of formulaic behaviour, I chose syntax and semantics as the two areas I wanted to explore, abandoning morphology (which is by far the one that has been discussed in the most detail⁵⁰³) for the moment. I approached these two aspects in two different datasets (Adj + N and V + Obj pairs), using two different computational frameworks (quantitative

⁵⁰³ Most recently see McConnell (2019), whose treatment of morphological archaism brings some important advances, not least the use of rigorous statistical and linguistic analysis, into the discussion.

analysis of variation via an entropy-based measure, and the Distributional Semantics analysis of object fillers). Further work can undoubtedly be done on both of these datasets, and by extending the analysis to formulae I have not considered; and both of these frameworks can definitely be applied to different questions. Both my case studies also do not consider metre, a choice which was necessary to limit the work to a manageable size, but which of course creates a lacuna that could be filled by future studies.

My explorations have been guided by some key texts and theories. Hainsworth's work on defining the flexible formula (Hainsworth 1968) provides a way to describe formularity and indeed to isolate formulae as a target for study by using what is, after all, a usage-based and cognitive-based definition: the existence of associations between the words that make up a formula is based on the fact that they normally occur together in the specialised language of oral poetry, leading them to become an object of a definable kind in the minds of poets and listeners. Construction Grammar, which had already been applied to formulae by Bozzone (2014a) and Antović and Cánovas (2016a), provided the theoretical tools to translate this statement into one that is more understandable, and applicable, in linguistic terms: the association that connects the parts of a formula is a part of the grammar of oral poetry, together with some as yet undefined restrictions to its behaviour (the restrictions, of course, that I then set out to map). A Construction Grammar-based study of English idioms (Wulff 2008) offered the outline of a quantitative approach that gave us a detailed description of the behaviour of Adj + N pairs in early Greek epic poetry; compared to the behaviour of non-formulaic Adj + N pairs, this allowed us to isolate some traits that characterise formulae, and which therefore can be understood to be part of their cognitive definition in the minds of language users. This last claim, which for Wulff was based on the correlation between linguistic restrictions that apply to idioms and the judgement of native speakers, may seem an ambitious leap when it comes to a study of ancient Greek; however, by looking at how Apollonius of Rhodes, a poet who aims for a very different style from

Homer, manipulates exactly those traits that are most characteristic of formulae, we have shown how these linguistic restrictions were most likely perceived as characteristic of the oral-epic style.

The study of V + Obj pairs, on the other hand, was based on a convergence between the behaviour of constructions with open slots, as studied by Barðdal (2008) and Perek (2016), and the idea of formulaic systems as based on a nucleus (in our case, a transitive verb) and a periphery (in our case, the object), as outlined by Visser (1988) and Bakker and Fabbriotti (1991) among others. While this definition is a lot less primary for the study of formulaic semantics than Hainsworth's definition was for the syntactic study (partly due to the parallel use of the frequentist definition of formula by Pavese and Venti 2000 and Pavese and Boschetti 2003), it once again allowed us to look at Construction Grammar-based studies to find the quantitative methods needed to formalise and compare meaning relations between words. The application of Distributional Semantics⁵⁰⁴ allowed us to describe how the meaning of objects in transitive-verb formulae varies, and to compare this variation to another non-formulaic baseline. In this case, the results showed a much less salient role of semantic restrictions in formularity, a somewhat puzzling outcome that suggests the need for further exploration.

One of the broader goals of my work was also to offer a study of formulaic behaviour that, while drawing to a very significant extent on contemporary linguistic approaches and using computational methods to both collect and analyse the data, would be accessible and relevant to classicists interested in the question of formulaic language and variation. This thesis was written for a classically-minded audience; one of its intended results is to introduce some of these computational and linguistic approaches in a way that makes it clear that they can in fact be applied to known issues in classical studies, and that this

⁵⁰⁴ Which is of course not a branch of Construction Grammar, or related to it at all in principle, but has been fruitfully applied in Construction Grammar-based studies.

application is not insurmountably difficult. While my own background is a mix of classics (primarily) with some linguistics and programming on the side, one of my priorities has been to make sure that my results could be followed and ideally repeated without requiring detailed expertise in the computational methodologies I used.

VIII.2 How it is going: results and methodological advances

Both sections of this thesis, the one on syntax and the one on semantics, were divided into two parts: a preliminary study, which had the dual aim of testing the applicability of the methods and of allowing us to start formulating hypotheses on the behaviour of our target samples, and a full-scale study of a class of expressions, limited only by a frequency threshold. The preliminary studies did raise some interesting points of their own: the role of frequency in syntactic flexibility and the interesting behaviour of quantifiers were already clearly present in the results from chapter III, and chapter VI showed how we can define classes of formulae based on their meaning, with results corresponding at least partially to their semantic range.

The syntactic study in chapter III also introduced some key methodological points: the issue of how we can move beyond the idea of a ‘base form’ of each formula, a concept that is both hard to define and methodologically questionable, and still reach meaningful results on formulaic behaviour; and the importance of comparing epic data with a non-epic baseline (and, later, the benefits of this comparison when addressing changes in formulaic usage in non-oral poetry). Chapter IV was dedicated to an analysis of formulaic pairs that would satisfy these requirements, through the use of a baseline corpus from the fifth century BC and of a statistical measure, relative entropy, which would remove the need for a ‘base form’ for each pair. The results in this chapter were quite striking: Adj + N pairs in early Greek epic, while still on average less flexible than similar pairs in non-epic texts, are

particularly resistant to specific types of modifications (changes in case and number and those which affect their position in the line); this appears to be motivated by their usage in oral composition, which creates a need to have predictably-shaped building blocks with which to fill specific positions in the line.

Still, flexibility in general is an inherent trait of formulae, which do not become more entrenched with repetition; more frequent formulae are in fact more flexible, not less, as we would expect if they were simply memorised, like Lord and others argued they must have been.⁵⁰⁵ Can this result be integrated into literary discussions of memory and its role in the Homeric poems?⁵⁰⁶ Finally, formulae that involve adjectives with more general semantic properties, like indefinite adjectives, which combine more freely with a wider range of nouns, are more flexible and likely to be perceived as less ‘formulaic’ in nature. This result, again, may be of use in discussing the role of these adjectives in promoting the introduction of the article: the fact that formulae that involve the adjective ἄλλος are the place from which the article is introduced to the epic language may not just be due to the necessity to distinguish between ‘another’ and ‘the other,’ but may have been helped by the fact that ἄλλος was already not fully a part of the system of ‘generic epithets’ that refer to common nouns in early Greek epic.

Among other things, the results in this chapter show that flexible formulae are still recognisably formulae: they have not lost their identity because they are flexible, but rather they show a coherent pattern of restrictions. This, among other things, addresses the objections by critics such as Friedrich (2019), and also brings formulae even closer to linguistic accounts of idiomaticity, especially in Construction Grammar. We also obtained individual results on the behaviour of specific formulae, which can potentially be used in future studies of specific formulaic systems or semantic areas. Finally, the comparison with

⁵⁰⁵ Lord (2000), 43, on which see §1.2.2.

⁵⁰⁶ The key work on this topic being of course Minchin (2001).

Hellenistic epic (Apollonius of Rhodes) conducted in chapter V showed how some of these restrictions to formulaic behaviour are indeed characteristic of oral poetry, not of hexameter poetry in general; there is even some evidence that Apollonius is consciously manipulating some of these patterns of behaviour to distance himself from ‘Homeric’ style, an idea which fits very well with his poetic choices.

The results from the Apollonius chapter are particularly interesting in that they address one of the apparent risks of a study like mine, which has a heavy linguistic focus and applies very modern methodologies: the risk of adopting a completely etic approach, which imposes our frameworks on the data without any consideration of whether these frameworks would have had any relevance for the people who produced the data themselves. The phenomena that were described in chapter IV when it comes to variation in Adj + N formulae are not simply the application of a completely modern framework, which would not have had any relevance to users of the epic language (whether poets or audiences); on the contrary, they describe traits which would have been perceived by ancient Greek speakers attuned to traditional conventions, and which could be used to consciously break away from these same conventions, as Apollonius does.

The results on semantics in chapter VII are quite different. Contrary to the expectations we had, which were based on cross-linguistic comparison with constructions with open slots, there does not seem to be any link between the semantic range and productivity of verb phrases, at least when comparing usage in early epic with the baseline corpus. This is perhaps easily explained: if anything determined the productivity of verbs in later linguistic usage, it was probably not the way they were used in Homer and Hesiod.⁵⁰⁷ This may suggest that epic usage was perceived as separate enough from everyday linguistic usage to prevent this kind of influence; but to support this claim we would need some positive

⁵⁰⁷ This is not an issue for our syntactic analysis in §IV, where we did not assume that the usage of Adj + N pairs in the epic corpus would influence the 5th-century baseline corpus in any way.

evidence that semantic restrictions are in fact correlated to productivity in ancient Greek. To gather this evidence, a different study would be required, focusing on the development of constructional productivity and how it correlates to semantic range in different stages of a non-epic corpus. There is also no effect of frequency on semantic flexibility, or vice-versa, unlike what we had observed for syntactic flexibility in chapter IV.

As to the key question – how is semantic range related to formularity? –, the results from chapter VII are perhaps even more baffling. We cannot see much of an effect of formularity on semantic range, and what little we do have goes in an unexpected direction. We would expect the conventions of formulaic language to promote the creation of clusters of formulae with related meaning but different metrical shape, in order to provide maximum flexibility in terms of formulaic systems. This would lead to clusters of semantically related fillers, which would ideally show up in the results as a higher median similarity for the fillers of a specific verb in the formulaic sample compared to the baseline sample. What we find instead is that the *baseline* sample has consistently higher semantic similarity than the formulaic sample. (We also find very few significant differences between the two, but the tendency is at least upheld in most cases where differences are not significant.)

Is this parallel to constructional or idiomatic pre-emption, that is, the phenomenon through which the existence of an idiom with a certain meaning prevents the creation of a new idiom with a very similar one, because that communicative function is, in a sense, already fulfilled?⁵⁰⁸ This would be a great answer if it were not for the fact that, as we said above, all approaches to formularity assume that formulae with similar meaning but different metrical shape are actually useful and sought after. It is possible, of course, that clusters of meaning are present in our formulaic sample, but that the method we used to analyse the data did not quite succeed in capturing them: we only used a fairly primitive measure of

⁵⁰⁸ On this phenomenon see e.g. Suttle and Goldberg (2011).

semantic similarity, after all (the distance between all fillers of a TrV + Obj construction and their centroid in the semantic space). This may not be as suited to capturing clusters as other kinds of analysis, which could definitely be conducted on similar data, or as part of a completely new study.

In any case, this was the first time a Distributional Semantics approach was applied to the description of formulaic language. Bumps on the road are to be expected, but some results remain clear and can be productive for further work; not least, chapter VI shows that a reasonably accurate Distributional Space Model of ancient Greek can be built and used to compare the meaning of words in formulae, and this principle can be applied to a variety of different studies.

This observation leads us to the final important contribution of this thesis, which has to do with advances in method and the application of quantitative linguistic analysis to phenomena that go beyond the chronology of early Greek epic (see §II.1). My work has not just shown what these methods are useful for, and what they can tell us about formulaic language; it has also involved the creation of reusable scripts that are openly available.⁵⁰⁹ It also gave an example of the usefulness of new and sometimes previously-unused data sources (the Diorisis Ancient Greek Corpus, the valency lexicon used in chapter VI), and applied innovative statistical and quantitative approaches (entropy and the analysis of linguistic variation, Distributional Semantics) to the description of oral poetry. My hope is that these methods and materials can be put to use by future scholars in order to improve on the suggestions made in this thesis for a usage-based description of formulaic language.

⁵⁰⁹ Once again, all of this can be found at <https://github.com/MartinaAstridRodda/dphil-thesis>.

VIII.3 What we earned, with some ideas for further work

To scholars with an interest in archaic epic, language, and formularity, this thesis should have provided a good amount of food for thought. Beyond the results detailed above, the main takeaway from all this work is that variation phenomena in the language of Greek epic can be explored with the same tools that we use to explore everyday language usage, and indeed they show substantial parallels with the latter: the language of Homer is not so much of a *Kunstsprache* as to be subject to restrictions that are completely unique to it, and which require completely unique explanations. This conclusion should have a bearing on approaches to such phenomena as morphological or dialectal variation and chronological linguistic development, for which explanations have sometimes been advanced that would paint a completely idiosyncratic picture for the epic language. On the flipside, quantitative analysis also allows us to give a clear and well-defined picture of the areas in which the epic language is different from everyday language: our results highlight the traits that made it useful to singers, and at the same time the traits that would make formulae recognisable as formulae by an audience. This intersection of parallels with living languages tempered by a strong awareness of how the performance context and the audience's expectations define the poetic grammar is, after all, at the very root of research on orality and formularity, and is the most productive way in which this work should continue.

The analysis I have attempted could be extended by studying different samples, or by including different aspects of the material, such as meter, or different methodologies, such as a cluster-based statistical analysis; I have outlined some ways this could be achieved in §VIII.2. Two areas of particular interest are the analysis of recurring expressions in other traditions, especially those that still have a link with orality (what about early historiography, or even philosophical dialogue?), and the use of formulae in non-Greek

traditions, such as the Vedas or Old Babylonian poetry.⁵¹⁰ The latter of these could have the highly beneficial side-effect of stimulating the creation of curated corpora for these texts and for their respective languages.

These results could also be integrated with more literary accounts of phenomena such as traditional referentiality. Recent work on inter- and intratextuality in Homer has relied a lot less than would be possible on formularity;⁵¹¹ these discussions could be reassessed based on concepts of formulaic repetition and variation, with a better awareness of how the two interact. The issue of meaning in formulaic systems resurfaces in a different form: can the way a formula is varied, or the choice whether to coin a semantically similar formula, have a bearing on literary interpretation? Finally, for the more literary-minded, the results on Apollonius' self-conscious use of formulaic variation open up promising avenues towards the analysis of the reception of Homeric language in different poetic traditions.

⁵¹⁰ On parallels between Greek and Old Babylonian poetry, especially when it comes to orality, Ballesteros Petrella (in preparation) is of particular interest. For some work on formularity in the Hebrew Psalms, see Coleman (2019).

⁵¹¹ An example of this issue is Currie (2016).

Bibliography

- Aloni, Antonio. 1986. *Tradizioni arcaiche della Troade e composizione dell'Iliade*. Milano: Unicopli.
- Ambrosini, Riccardo. 1991. 'Osservazioni sulle funzioni referenziali dell'articolo nei poemi omerici'. In *Studi di filologia classica in onore di Giusto Monaco*, 1:3–16. Palermo: Università di Palermo. Facoltà di lettere e filosofia. Istituto di filologia greca. Istituto di filologia latina.
- Andersen, Øivind. 2019. 'Review of Bruno Currie: Homer's Allusive Art'. *Gnomon* 91: 97–103. <https://doi.org/10.17104/0017-1417-2019-2-97>.
- Antović, Mihailo, and Cristóbal Pagán Cánovas. 2016a. 'Construction Grammar and Oral Formulaic Theory'. In Antović, Mihailo, and Cristóbal Pagán Cánovas, eds. *Oral Poetics and Cognitive Science*, 79–98. *Linguae & Litterae* 56. Berlin; Boston (MA): De Gruyter.
- , eds. 2016b. *Oral Poetics and Cognitive Science*. *Linguae & Litterae* 56. Berlin; Boston (MA): De Gruyter.
- Ayto, John, ed. 2009. *Oxford Dictionary of English Idioms*. 3rd edition. Oxford: Oxford University Press. 10.1093/acref/9780199543793.001.0001.
- Baayen, R. Harald. 2008. *Analyzing Linguistic Data: A Practical Introduction to Statistics Using R*. Cambridge: Cambridge University Press. <https://doi.org/10.1017/CBO9780511801686>.
- Baechle, Nicholas. 2007. *Metrical Constraint and the Interpretation of Style in the Tragic Trimeter*. *Greek Studies: Interdisciplinary Approaches*. Lanham: Rowman & Littlefield.
- Bakker, Egbert J. 1988. *Linguistics and Formulas in Homer: Scalarity and the Description of the Particle Per*. Amsterdam: John Benjamins.
- . 1997. *Poetry in Speech: Orality and Homeric Discourse*. *Myth and Poetics*. Ithaca; London: Cornell University press.
- Bakker, Egbert J., and Florence Fabbricotti. 1991. 'Peripheral and Nuclear Semantics in Homeric Diction. The Case of Dative Expressions for "Spear"'. *Mnemosyne* 44 (1/2): 63–84.
- Bakker, Egbert J., and Nina van den Houten. 1992. 'Aspects of Synonymy in Homeric Diction: An Investigation of Dative Expressions for "Spear"'. *Classical Philology* 87: 1–13.
- Bakker, Stéphanie J. 2009. *The Noun Phrase in Ancient Greek: A Functional Analysis of the Order and Articulation of NP Constituents in Herodotus*. *Amsterdam Studies in Classical Philology* 15. Leiden; Boston: Brill.
- Ballesteros Petrella, Bernardo. In preparation. *The Divine Assembly: Comparative Studies in Early Greek and Babylonian Narrative Poetry*. *Oxford Classical Monographs*. Oxford: Oxford University Press.
- Bamman, David, and Gregory Crane. 2011. 'The Ancient Greek and Latin Dependency Treebanks'. In *Language Technology for Cultural Heritage: Selected Papers from the*

LaTeCH Workshop Series, 79–98. Theory and Applications of Natural Language Processing. Berlin; Heidelberg: Springer. doi: 10.1007/978-3-642-20227-8_5.

- Barðdal, Jóhanna. 2008. *Productivity: Evidence from Case and Argument Structure in Icelandic*. Constructional Approaches to Language 8. Amsterdam; Philadelphia: John Benjamins.
- Barnes, Timothy G. 2011. 'Homeric Ἀνδρωτῆτα Καὶ Ἥβην'. *The Journal of Hellenic Studies* 131: 1–13.
- Baroni, Marco, Georgiana Dinu, and Germán Kruszewski. 2014. 'Don't Count, Predict! A Systematic Comparison of Context-Counting vs. Context-Predicting Semantic Vectors'. In *Proceedings of the 52nd Annual Meeting of the Association for Computational Linguistics*, 1: Long Papers:238–47. Baltimore (MD): Association for Computational Linguistics. <https://doi.org/10.3115/v1/P14-1023>.
- Baroni, Marco, and Alessandro Lenci. 2011. 'How We BLESSed Distributional Semantic Evaluation'. In *Proceedings of the GEMS 2011 Workshop on GEometrical Models of Natural Language Semantics*, 1–10. Stroudsburg (PA): Association for Computational Linguistics. <http://dl.acm.org/citation.cfm?id=2140490.2140491>.
- Battezzato, Luigi, and Martina Astrid Rodda. 2018. 'Particelle e asindeto nel greco classico'. *Glotta* 94: 3–37.
- Beck, Deborah. 2008. 'Formular Economy in Homer. The Poetics of the Breaches. Hermes Einzelschriften, 100'. *Bryn Mawr Classical Review* 2008.10.27. <https://bmcr.brynmawr.edu/2008/2008.10.27/>.
- . 2009. 'Speech Act Types, Conversational Exchange, and the Speech Representational Spectrum in Homer'. In Grethlein, Jonas, and Antonios Rengakos, eds. *Narratology and Interpretation: The Content of Narrative Form in Ancient Literature*, 137–51. Trends in Classics. Supplementary Volumes 4. Berlin; New York: Walter de Gruyter.
- Beek, Lucien C. van. 2013. 'The Development of the Proto-Indo-European Syllabic Liquids in Greek'. Doctoral Thesis, Leiden: Leiden University. <https://openaccess.leidenuniv.nl/handle/1887/22881>.
- Boedeker, Deborah. 2002. 'Epic Heritage and Mythical Patterns in Herodotus'. In Bakker, Egbert J., Irene J.F. de Jong, and Hans van Wees, eds. *Brill's Companion to Herodotus*, 95–116. Leiden; Boston: Brill.
- Bonifazi, Anna, Annemieke Drummen, and Mark de Kreij. 2016. *Particles in Ancient Greek Discourse: Five Volumes Exploring Particle Use across Genres*. Hellenic Studies Series 74. Washington, DC: Center for Hellenic Studies. http://nrs.harvard.edu/urn-3:hul.ebook:CHS_BonifaziA_DrummenA_deKreijM.Particles_in_Ancient_Greek_Discourse.2016.
- Boschetti, Federico, Riccardo Del Gratta, and Harry Diakoff. 2016. 'Open Ancient Greek WordNet 0.5'. 30 May 2016. <https://dspace-clarin-it.ilc.cnr.it/repository/xmlui/handle/20.500.11752/ILC-56>.
- Bozzone, Chiara. 2010. 'New Perspectives on Formularity'. In *Proceedings of the 21st Annual UCLA Indo-European Conference*, 27–44. Bremen: Hempen.

- . 2014a. 'Constructions: A New Approach to Formularity, Discourse, and Syntax in Homer'. PhD diss., Los Angeles: UCLA: Indo-European Studies. <https://escholarship.org/uc/item/6kg0q4cx>.
- . 2014b. 'How Do Epic Poets Construct Their Lines? A Study of the Verb Προσέειπεν in Homer, Hesiod, Batrachomyomachia, Apollonius Rhodius, and Quintus Smyrnaeus'. Presented at the Annual Meeting of the American Philological Association, Chicago (IL).
- . 2016. 'The Mind of the Poet: Cognitive and Linguistic Perspectives'. In Gallo, Federico, ed. *Omero: Quaestiones Disputatae*, 79–106. Ambrosiana Graecolatina 5. Milano-Roma: Bulzoni.
- Bullinaria, John A., and Joseph P. Levy. 2007. 'Extracting Semantic Representations from Word Co-Occurrence Statistics: A Computational Study'. *Behavior Research Methods* 39: 510–26. <https://doi.org/10.3758/BF03193020>.
- . 2012. 'Extracting Semantic Representations from Word Co-Occurrence Statistics: Stop-Lists, Stemming, and SVD'. *Behavior Research Methods* 44: 890–907. <https://doi.org/10.3758/s13428-011-0183-8>.
- Bybee, Joan L. 2013. 'Usage-Based Theory and Exemplar Representations of Constructions'. In Hoffmann, Thomas, and Graeme Trousdale, eds. *The Oxford Handbook of Construction Grammar*, 49–69. Oxford; New York: Oxford University Press.
- Bybee, Joan L., and Sandra Thompson. 1997. 'Three Frequency Effects in Syntax'. *Annual Meeting of the Berkeley Linguistics Society* 23 (1): 378–88.
- Cacciari, Cristina, and Sam Glucksberg. 1991. 'Understanding Idiomatic Expressions: The Contribution of Word Meanings'. In Simpson, Gregory B., ed. *Understanding Word and Sentence*, 217–40. *Advances in Psychology* 77. Amsterdam: North-Holland.
- Cacciari, Cristina, and Patrizia Tabossi. 1988. 'The Comprehension of Idioms'. *Journal of Memory and Language* 27: 668–83. [https://doi.org/10.1016/0749-596X\(88\)90014-9](https://doi.org/10.1016/0749-596X(88)90014-9).
- Campbell, Ian. 2007. 'Chi-Squared and Fisher-Irwin Tests of Two-by-Two Tables with Small Sample Recommendations'. *Statistics in Medicine* 26: 3661–75.
- Cassio, Albio Cesare, ed. 2008. *Storia delle lingue letterarie greche*. Milano: Le Monnier Università.
- Celano, Giuseppe G.A. 2019. 'The Dependency Treebanks for Ancient Greek and Latin'. In Berti, Monica, ed. *Digital Classical Philology: Ancient Greek and Latin in the Digital Revolution*, 279–97. *Age of Access? Grundfragen Der Informationsgesellschaft* 10. Berlin; Boston (MA): De Gruyter Saur. <https://doi.org/10.1515/9783110599572-016>.
- Chafe, Wallace L. 1994. *Discourse, Consciousness and Time: The Flow and Displacement of Conscious Experience in Speaking and Writing*. Chicago; London: University of Chicago Press.
- Chantraine, Pierre. 2013. *Grammaire homérique*. Nouvelle édition revue et corrigée par Michel Casevitz. Vol. Tome 1. Phonétique et morphologie. Librairie Klincksieck. Série linguistique 3. Paris: Klincksieck.

- . 2015. *Grammaire homérique*. Nouvelle édition revue et corrigée par Michel Casevitz. Vol. Tome 2. Syntaxe. Librairie Klincksieck. Série linguistique 3. Paris: Klincksieck.
- Chomsky, Noam. 1957. *Syntactic Structures*. The Hague: Mouton & co. <http://hdl.handle.net/2027/mdp.39015004280957>.
- . 1965. *Aspects of the Theory of Syntax*. Special Technical Report (Massachusetts Institute of Technology. Research Laboratory of Electronics) 11. Cambridge (MA): MIT Press.
- Christensen, Joel P. 2009. 'Review: Rainer Friedrich. *Formular Economy in Homer: The Poetics of the Breaches*. *Hermes Einzelschriften* 100. Stuttgart: Franz Steiner, 2007.' *The American Journal of Philology* 130 (1): 131–35.
- Ciani, Maria Grazia. 1975. 'Ripetizione formulare in Apollonio Rodio'. *Bollettino dell'Istituto di Filologia Greca dell'Università di Padova* II: 191–208.
- Clark, Matthew. 1994. 'Enjambment and Binding in Homeric Hexameter'. *Phoenix* 48: 95–114. <https://doi.org/10.2307/1088310>.
- Clausner, Timothy C., and William Croft. 1997. 'Productivity and Schematicity in Metaphors'. *Cognitive Science* 21: 247–82. [https://doi.org/10.1016/S0364-0213\(99\)80024-X](https://doi.org/10.1016/S0364-0213(99)80024-X).
- Coleman, Stephen. 2019. 'The Psalmic Oral Formula Revisited: A Cognitive-Performative Approach'. *Biblical Interpretation* 27 (2): 186–207.
- Crystal, David. 2008. *A Dictionary of Linguistics and Phonetics*. 6th ed. Malden (MA); Oxford: Blackwell.
- Curran, James Richard. 2004. 'From Distributional to Semantic Similarity'. PhD diss., Edinburgh: Institute for Communicating and Collaborative Systems, School of Informatics. <https://era.ed.ac.uk/bitstream/handle/1842/563/IP030023.pdf>.
- Currie, Bruno. 2016. *Homer's Allusive Art*. Oxford: University Press.
- Cutting, J. Cooper, and Kathryn Bock. 1997. 'That's the Way the Cookie Bounces: Syntactic and Semantic Components of Experimentally Elicited Idiom Blends'. *Memory & Cognition* 25: 57–71. <https://doi.org/10.3758/BF03197285>.
- Devine, Andrew M., and Laurence D. Stephens. 1999. *Discontinuous Syntax: Hyperbaton in Greek*. New York; Oxford: Oxford University Press.
- Dik, Helma. 1995. *Word Order in Ancient Greek: A Pragmatic Account of Word Order Variation in Herodotus*. Amsterdam Studies in Classical Philology 5. Amsterdam: Gieben.
- . 2007. *Word Order in Greek Tragic Dialogue*. Oxford; New York: Oxford University Press.
- Dinu, Georgiana, Nghia The Pham, and Marco Baroni. 2013. 'DISSECT - DIStributional SEmantics Composition Toolkit'. In *Proceedings of the 51st Annual Meeting of the Association for Computational Linguistics: System Demonstrations*, 31–36. Sofia, Bulgaria: Association for Computational Linguistics. <http://www.aclweb.org/anthology/P13-4006>.
- Dugan, Kelly P. 2012. 'A Generative Approach to Homeric Enjambment: Benefits and Drawbacks'. MA thesis, Athens (GA): University of Georgia. https://getd.libs.uga.edu/pdfs/dugan_kelly_p_201208_ma.pdf.

- Emde Boas, Evert van, Albert Rijksbaron, Luuk Huitink, and Mathieu de Bakker. 2019. *Cambridge Grammar of Classical Greek*. Cambridge: Cambridge University Press.
- Erman, Britt, and Beatrice Warren. 2000. 'The Idiom Principle and the Open Choice Principle'. *Text. An Interdisciplinary Journal for the Study of Discourse* 20: 29–62.
- Espinal, M. Teresa, and Jaume Mateu. 2010. 'On Classes of Idioms and Their Interpretation'. *Journal of Pragmatics* 42: 1397–1411. <https://doi.org/10.1016/j.pragma.2009.09.016>.
- Evert, Stefan. 2008. 'Corpora and Collocations'. In Lüdeling, Anke, and Merja Kytö, eds. *Corpus Linguistics. An International Handbook*, 1212–48. Berlin: Mouton de Gruyter.
- Fabre, Cécile, and Alessandro Lenci. 2015. 'Distributional Semantics Today: Introduction to the Special Issue'. *Traitement Automatique Des Langues* 56 (2): 7–20.
- Fantuzzi, Marco. 1988. *Ricerche su Apollonio Rodio: Diacronie della dizione epica*. Filologia e critica 58. Roma: Edizioni dell'Ateneo.
- . 2008. "Homeric" Formularity in the Argonautica of Apollonius of Rhodes'. In Papanghelis, Theodore D., and Antonios Rengakos, eds. *Brill's Companion to Apollonius Rhodius*, 2nd rev. ed., 221–41. Leiden; Boston: Brill.
- Fantuzzi, Marco, and Richard Hunter. 2004. *Tradition and Innovation in Hellenistic Poetry*. Cambridge: Cambridge University Press.
- Faulkner, Andrew, ed. 2011. *The Homeric Hymns: Interpretative Essays*. Oxford: Oxford University Press.
- Fick, August. 1883. *Die homerische Odyssee in der ursprünglichen Sprachform wiederhergestellt*. Beiträge zur Kunde der Indogermanischen Sprachen: Supplementband. Göttingen: R. Peppmüller.
- Finkelberg, Margalit. 2004. 'Oral Theory and the Limits of Formulaic Diction'. *Oral Tradition* 19 (2): 236–52.
- Finnegan, Ruth H. 1977. *Oral Poetry: Its Nature, Significance and Social Context*. Cambridge: University Press.
- . 1992. *Oral Poetry: Its Nature, Significance and Social Context*. 1st Midland Book ed. Bloomington: Indiana University Press.
- Firth, John R. 1957. 'A Synopsis of Linguistic Theory 1930–1955'. In *Studies in Linguistic Analysis*, 1–32. Oxford: Blackwell.
- Fitzgerald, Robert, trans. 1961. *Homer: The Odyssey*. New York: Doubleday.
- Foley, John Miles. 1988. *The Theory of Oral Composition: History and Methodology*. Folkloristics. Bloomington (IN): Indiana University Press.
- . 1990. *Traditional Oral Epic: The Odyssey, Beowulf, and the Serbo-Croatian Return Song*. Berkeley; Los Angeles; Oxford: University of California Press. <http://ark.cdlib.org/ark:/13030/ft2m3nb18b/>.
- . 1999. *Homer's Traditional Art*. University Park (PA): Pennsylvania State University Press.
- Fowler, Robert L. 1983. 'Review of Homer, Hesiod and the Hymns by Richard Janko'. *Phoenix* 37: 345–47. <https://doi.org/10.2307/1088155>.

- Fraser, Bruce. 1970. 'Idioms within a Transformational Grammar'. *Foundations of Language* 6: 22–42.
- Friedrich, Rainer. 2007. *Formular Economy in Homer: The Poetics of the Breaches*. Hermes Einzelschriften 100. Stuttgart: Steiner.
- . 2009. 'Response: Friedrich on Beck on Rainer Friedrich, Formular Economy in Homer. The Poetics of the Breaches'. *Bryn Mawr Classical Review* 2009.02.29. <https://bmcr.brynmawr.edu/2009/2009.02.29/>.
- . 2019. *Postoral Homer: Orality and Literacy in the Homeric Epic*. Hermes Einzelschriften 112. Stuttgart: Franz Steiner Verlag.
- Fries, Almut. 2014. *Pseudo-Euripides, Rhesus*. Untersuchungen Zur Antiken Literatur Und Geschichte 114. Berlin: De Gruyter.
- Furnas, George W., Susan Dumais, Thomas K. Landauer, Richard A. Harshman, Lynn A. Streeter, and Karen E. Lochbaum. 1998. 'Information Retrieval Using a Singular Value Decomposition Model of Latent Semantic Structure'. In *Proceedings of SIGIR 1998*, 465–80. <https://www.microsoft.com/en-us/research/publication/information-retrieval-using-singular-value-decomposition-model-latent-semantic-structure/>.
- Garvin, Paul L. 1962. 'Computer Participation in Linguistic Research'. *Language* 38: 385–89. <https://doi.org/10.2307/410674>.
- Geisz, Camille. 2017. *A Study of the Narrator in Nonnus of Panopolis' Dionysiaca: Storytelling in Late Antique Epic*. Amsterdam Studies in Classical Philology 25. Leiden: Brill.
- Gibbs, Raymond W., and Nandini P. Nayak. 1989. 'Psycholinguistic Studies on the Syntactic Behavior of Idioms'. *Cognitive Psychology* 21: 100–138. [https://doi.org/10.1016/0010-0285\(89\)90004-2](https://doi.org/10.1016/0010-0285(89)90004-2).
- Gibbs, Raymond W., Nandini P. Nayak, John L. Bolton, and Melissa E. Keppel. 1989. 'Speakers' Assumptions about the Lexical Flexibility of Idioms'. *Memory & Cognition* 17: 58–68. <https://doi.org/10.3758/BF03199557>.
- Goldberg, Adele E. 1995. *Constructions: A Construction Grammar Approach to Argument Structure*. Cognitive Theory of Language and Culture. Chicago; London: University of Chicago Press.
- . 2006. *Constructions at Work: The Nature of Generalization in Language*. Oxford Linguistics. Oxford: Oxford University Press.
- . 2013. 'Constructionist Approaches'. In Hoffmann, Thomas, and Graeme Trousdale, eds. *The Oxford Handbook of Construction Grammar*, 15–31. Oxford; New York: Oxford University Press.
- Goldstein, David M. 2015. *Classical Greek Syntax: Wackernagel's Law in Herodotus*. Brill's Studies in Indo-European Language & Linguistics 16. Leiden; Boston: Brill.
- . 2020. 'Homeric -Phi(n) Is an Oblique Case Marker'. *Transactions of the Philological Society* 118: 343–75. <https://doi.org/10.1111/1467-968X.12192>.
- Graziosi, Barbara, and Johannes Haubold. 2005. *Homer: The Resonance of Epic*. London: Duckworth.
- . 2015. 'The Homeric Text'. *Ramus* 44: 5–28. <https://doi.org/10.1017/rmu.2015.14>.

- Grewcock, Rachel. 2018. 'Computational Semantics and the Syntax of Motion in Ancient Greek'. MPhil Thesis, Cambridge: University of Cambridge.
- Griffin, Jasper. 1986. 'Homeric Words and Speakers'. *The Journal of Hellenic Studies* 106: 36–57. <https://doi.org/10.2307/629641>.
- Guardiano, Cristina. 2013. 'The Greek Definite Article across Time'. *Studies in Greek Linguistics* 33: 76–91.
- Hackstein, Olav. 2010. 'The Greek of Epic'. In Bakker, Egbert J., ed. *A Companion to the Ancient Greek Language*, 401–23. Blackwell Companions to the Ancient World. Oxford: Wiley-Blackwell.
- Hainsworth, John Bryan. 1964. 'Structure and Content in Epic Formulae: The Question of the Unique Expression'. *The Classical Quarterly, New Series*, 14: 155–64.
- . 1968. *The Flexibility of the Homeric Formula*. Oxford: Clarendon Press.
- Hamilton, William L., Jure Leskovec, and Dan Jurafsky. 2016. 'Diachronic Word Embeddings Reveal Statistical Laws of Semantic Change'. In *Proceedings of the 54th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, 1489–1501. Berlin, Germany: Association for Computational Linguistics. <http://www.aclweb.org/anthology/P16-1141>.
- Harris, Zellig S. 1954. 'Distributional Structure'. *Word* 10 (2–3): 146–62. <https://doi.org/10.1080/00437956.1954.11659520>.
- Haslam, Michael W. 1993. 'Callimachus' Hymns'. In Harder, M. Annette, Remco F. Regtuit, and Gerry C. Wakker, eds. *Callimachus*, 111–25. *Hellenistica Groningana* 1. Groningen: Forsten.
- Haug, Dag T. T. 2002. *Les phases de l'évolution de la langue épique: Trois études de linguistique homérique*. *Hypomnemata* 142. Göttingen: Vandenhoeck & Ruprecht.
- Haug, Dag T. T., and Marius L. Jøhndal. 2008. 'Creating a Parallel Treebank of the Old Indo-European Bible Translations'. In *Proceedings of the Second Workshop on Language Technology for Cultural Heritage Data (LaTeCH 2008)*, 27–34. <https://www.hf.uio.no/ifikk/english/research/projects/proiel/Activities/proiel/publications/marrakech.pdf>.
- Higbie, Carolyn. 1990. *Measure and Music: Enjambement and Sentence Structure in the Iliad*. Oxford: Clarendon Press.
- Hilpert, Martin. 2011. 'Dynamic Visualizations of Language Change: Motion Charts on the Basis of Bivariate and Multivariate Data from Diachronic Corpora'. *International Journal of Corpus Linguistics* 16: 435–61. <https://doi.org/info:doi/10.1075/ijcl.16.4.01hil>.
- Hoekstra, Arie. 1965. *Homeric Modifications of Formulaic Prototypes: Studies in the Development of Greek Epic Diction*. *Verhandelingen Der Koninklijke Nederlandse Akademie van Wetenschappen, Afd. Letterkunde, Nieuwe Reeks* LXXI.1. Amsterdam: N.V. Noord-Hollandsche Uitgevers Maatschappij.
- . 1969. *The Sub-Epic Stage of the Formulaic Tradition: Studies in the Homeric Hymns to Apollo, to Aphrodite and to Demeter*. *Verhandelingen Der Koninklijke Nederlandse Akademie van Wetenschappen, Afd. Letterkunde, Nieuwe Reeks* LXXV.2. Amsterdam; London: North-Holland Publishing Company.

- . 1986. 'Review of Homer, Hesiod and the Hymns by Richard Janko'. *Mnemosyne* 39: 158–64.
- Hopkinson, Neil. 1994. 'Nonnus and Homer'. In *Studies in the Dionysiaca of Nonnus*, 9–42. Supplementary Volume (Cambridge Philological Society) 17. Cambridge: Philological Society.
- Hunter, Richard L. 1992. 'Writing the God: Form and Meaning in Callimachus, Hymn to Athena'. *Materiali e discussioni per l'analisi dei testi classici* 29: 9–34. <https://doi.org/10.2307/40236011>.
- . 1993. *The Argonautica of Apollonius: Literary Studies*. Cambridge: University Press.
- . 2015. *Argonautica. Book IV*. Cambridge Greek and Latin Classics. Cambridge: University Press.
- Jackendoff, Ray. 1995. 'The Boundaries of the Lexicon'. In Everaert, Martin, Erik-Jan van der Linden, André Schenk, and Rob Schreuder, eds. *Idioms: Structural and Psychological Perspectives*, 133–65. Hillsdale (NJ); Hove: Lawrence Erlbaum Associates.
- Janko, Richard. 1982. *Homer, Hesiod and the Hymns: Diachronic Development in Epic Diction*. Cambridge: Cambridge University Press.
- . 2012. 'Πρώτον Τε Καὶ Ὑστάτον Αἰὲν Ἀεῖδειν: Relative Chronology and the Literary History of the Early Greek Epos'. In Andersen, Øivind, and Dag T.T. Haug, eds. *Relative Chronology in Early Greek Epic Poetry*, 20–43. Cambridge; New York: Cambridge University Press.
- Jenset, Gard B. 2013. 'Mapping Meaning with Distributional Methods: A Diachronic Corpus-Based Study of Existential There'. *Journal of Historical Linguistics* 3: 272–306.
- Jenset, Gard B., and Barbara McGillivray. 2017. *Quantitative Historical Linguistics: A Corpus Framework*. Oxford Studies in Diachronic and Historical Linguistics 26. Oxford: Oxford University Press.
- Johnson, Kyle P. and et al. 2014. *CLTK: The Classical Language Toolkit*. <https://doi.org/10.5281/zenodo.593336>.
- Jones, Brandtly. 2010. 'Relative Chronology within (an) Oral Tradition'. *The Classical Journal* 105: 289–318.
- . 2012. 'Relative Chronology and an "Aeolic Phase" of Epic'. In Andersen, Øivind, and Dag T.T. Haug, eds. *Relative Chronology in Early Greek Epic Poetry*, 44–64. Cambridge; New York: Cambridge University Press.
- Kelly, Adrian. 2007. *A Referential Commentary and Lexicon to Homer, Iliad VIII*. Oxford: Oxford University Press.
- . 2012. 'The Audience Expects: Penelope and Odysseus'. In Minchin, Elizabeth, ed. *Orality, Literacy and Performance in the Ancient World*, 3–24. *Orality and Literacy in the Ancient World* 9. Leiden: Brill.
- Kiparsky, Paul. 1976. 'Oral Poetry: Some Linguistic and Typological Considerations'. In Stolz, Benjamin A., and Richard S. Shannon, eds. *Oral Literature and the Formula*, 73–125. Ann Arbor: Center for the Coördination of Ancient and Modern Studies, University of Michigan.

- Kirk, Geoffrey S. 1962. *The Songs of Homer*. Cambridge University Press.
- Lakoff, George. 1987. *Women, Fire, and Dangerous Things: What Categories Reveal about the Mind*. Chicago; London: University of Chicago Press.
- Landauer, Thomas K., and Susan T. Dumais. 1997. 'A Solution to Plato's Problem: The Latent Semantic Analysis Theory of Acquisition, Induction, and Representation of Knowledge'. *Psychological Review* 104: 211–40. <https://doi.org/10.1037/0033-295X.104.2.211>.
- Langacker, Ronald W. 1987. *Foundations of Cognitive Grammar*. Vol. 1: Theoretical prerequisites. Stanford (CA): Stanford University Press.
- Lebani, Gianluca E., Marco Silvio Giuseppe Senaldi, and Alessandro Lenci. 2015. 'Modeling Idiom Variability with Entropy and Distributional Semantics'. In *Proceedings of the 6th Conference on Quantitative Investigations in Theoretical Linguistics*. Tübingen: Universität Tübingen.
- Lévi-Strauss, Claude. 1955. 'The Structural Study of Myth'. *The Journal of American Folklore* 68 (270): 428–44. <https://doi.org/10.2307/536768>.
- Levy, Omer, and Yoav Goldberg. 2014. 'Dependency-Based Word Embeddings'. In *Proceedings of the 52nd Annual Meeting of the Association for Computational Linguistics*, 2: Short Papers:302–8. Baltimore (MD): Association for Computational Linguistics. <https://doi.org/10.3115/v1/P14-2050>.
- Liapis, Vayos. 2012. *A Commentary on the Rhesus Attributed to Euripides*. Oxford: Oxford University Press.
- Liberman, Mark. 2015. 'Grexicography'. *Language Log* (blog). 22 June 2015. <https://web.archive.org/web/20190918151123/https://languagelog.ldc.upenn.edu/nll/?p=19609>.
- Lord, Albert Bates. 2000. *The Singer of Tales*. 2nd ed. Cambridge, MA; London: Harvard University Press.
- Maaten, Laurens van der, and Geoffrey Hinton. 2008. 'Visualizing Data Using T-SNE'. *Journal of Machine Learning Research* 9: 2579–2605.
- Maciver, Calum. 2018. 'Program and Poetics in Quintus Smyrnaeus' Posthomerica'. In Simms, Robert, ed. *Brill's Companion to Prequels, Sequels, and Retellings of Classical Epic*, 71–89. Brill's Companions to Classical Reception 15. Leiden: Brill.
- Manolessou, Io, and Geoffrey C. Horrocks. 2007. 'The Development of the Definite Article in Greek'. *Studies in Greek Linguistics* 27: 224–36.
- Matlock, Teenie. 1998. 'What Is Missing in Research on Idioms?' *The American Journal of Psychology* 111: 643–48. <https://doi.org/10.2307/1423557>.
- McConnell, Thomas. 2019. 'Chronology, Dialect and Style: Studies in the Distribution of Linguistic Archaisms in Early Greek Hexameter Poetry'. D.Phil. Thesis, Classical Languages and Literature, Oxford: University of Oxford.
- McGillivray, Barbara, Simon Hengchen, Viivi Lähteenoja, Marco Palma, and Alessandro Vatri. 2019. 'A Computational Approach to Lexical Polysemy in Ancient Greek'. *Digital Scholarship in the Humanities* 34 (4):893–907. <https://doi.org/10.1093/llc/fqz036>.

- McGillivray, Barbara, Marco Carlo Passarotti, and Paolo Ruffolo. 2009. 'The Index Thomisticus Treebank Project: Annotation, Parsing and Valency Lexicon'. *TAL – Traitement Automatique Des Langues* 50 (2): 103–27.
- McGillivray, Barbara, and Alessandro Vatri. 2015. 'Computational Valency Lexica for Latin and Greek in Use: A Case Study of Syntactic Ambiguity'. *Journal of Latin Linguistics* 14: 101–26. <https://doi.org/10.1515/joll-2015-0005>.
- Mendez Dosuna, Julián. 1993. 'Metátesis de cantidad en jónico-ático y heracleota'. *Emérita* 61: 95–134. <https://doi.org/10.3989/emerita.1993.v61.i1.459>.
- Mickey, K. 1981. 'Dialect Consciousness and Literary Language: An Example from Ancient Greek'. *Transactions of the Philological Society* 79: 35–66. <https://doi.org/10.1111/j.1467-968X.1981.tb01186.x>.
- Mikolov, Tomas, Kai Chen, Greg Corrado, and Jeffrey Dean. 2013. 'Efficient Estimation of Word Representations in Vector Space'. *ArXiv:1301.3781 [Cs]*. <http://arxiv.org/abs/1301.3781>.
- Mikolov, Tomas, Ilya Sutskever, Kai Chen, Greg Corrado, and Jeffrey Dean. 2013. 'Distributed Representations of Words and Phrases and Their Compositionality'. *ArXiv:1310.4546 [Cs, Stat]*. <http://arxiv.org/abs/1310.4546>.
- Miller, James N., and Jane C. Miller. 2010. *Statistics and Chemometrics for Analytical Chemistry*. 6th ed. Harlow: Pearson Education.
- Minchin, Elizabeth. 2001. *Homer and the Resources of Memory: Some Applications of Cognitive Theory to the Iliad and the Odyssey*. Oxford: Oxford University Press.
- . 2009. 'The Homeric Formula'. *The Classical Review* 59: 10–11.
- Montana, Fausto. 2014. 'Hellenistic Scholarship'. In Montanari, Franco, Stephanos Matthaios, and Antonios Rengakos, eds. *Brill's Companion to Ancient Greek Scholarship*, 1:60–183. Brill's Companions in Classical Studies. Leiden; Boston: Brill.
- Montanari, Franco. 2014. 'Ekthesis. A Product of the Ancient Scholarship'. In Montanari, Franco, Stephanos Matthaios, and Antonios Rengakos, eds. *Brill's Companion to Ancient Greek Scholarship*, 2:641–72. Brill's Companions in Classical Studies. Leiden; Boston: Brill.
- Mooney, George W. 1912. *The Argonautica of Apollonius Rhodius*. London: Longmans, Green.
- Morpurgo Davies, Anna. 1964. "'Doric" Features in the Language of Hesiod'. *Glotta* 42: 138–65.
- Nagler, Michael N. 1967. 'Towards a Generative View of the Oral Formula'. *Transactions and Proceedings of the American Philological Association* 98: 269–311. <https://doi.org/10.2307/2935878>.
- . 1974. *Spontaneity and Tradition: A Study in the Oral Art of Homer*. Berkeley: University of California Press.
- Nagy, Gregory. 1996. *Poetry as Performance: Homer and Beyond*. Cambridge: Cambridge university press. <http://chs.harvard.edu/wa/pageR?tn=ArticleWrapper&bdc=12&mn=2064>.

- . 2000. 'Homeri Ilias. Recensuit / Testimonia Congessit. Volumen Prius, Rhapsodias I-XII Continens'. *Bryn Mawr Classical Review* 2000.09.12. <https://bmcr.brynmawr.edu/2000/2000.09.12/>.
- . 2008. *Homer the Classic*. Electronic edition <<http://chs.harvard.edu/CHS/article/display/3472>>. *Hellenic Studies* 36. Washington, DC: Center for Hellenic Studies, Trustees for Harvard University.
- . 2010. *Homer the Preclassic*. Electronic edition, <<http://chs.harvard.edu/CHS/article/display/4377>>. Washington, DC: Center for Hellenic Studies, Trustees for Harvard University.
- Notopoulos, James A. 1960. 'Homer, Hesiod and the Achaean Heritage of Oral Poetry'. *Hesperia: The Journal of the American School of Classical Studies at Athens* 29: 177–97. <https://doi.org/10.2307/147293>.
- . 1962. 'The Homeric Hymns as Oral Poetry: A Study of the Post-Homeric Oral Tradition'. *The American Journal of Philology* 83: 337–68. <https://doi.org/10.2307/292918>.
- Nunberg, Geoffrey. 1978. *The Pragmatics of Reference*. Bloomington (IN): Indiana University Linguistics Club.
- Nunberg, Geoffrey, Ivan A. Sag, and Thomas Wasow. 1994. 'Idioms'. *Language* 70: 491–538. <https://doi.org/10.2307/416483>.
- Nünlist, René. 2009. *The Ancient Critic at Work: Terms and Concepts of Literary Criticism in Greek Scholia*. Cambridge: Cambridge University Press.
- . 2014. 'Poetics and Literary Criticism in the Framework of Ancient Greek Scholarship'. In Montanari, Franco, Stephanos Matthaios, and Antonios Rengakos, eds. *Brill's Companion to Ancient Greek Scholarship*, 2:706–55. Brill's Companions in Classical Studies. Leiden; Boston: Brill.
- O'Neill, Eugene G. 1942. 'The Localization of Metrical Word-Types in the Greek Hexameter'. *Yale Classical Studies* 8: 103–78.
- Papanghelis, Theodore D., and Antonios Rengakos, eds. 2008. *Brill's Companion to Apollonius Rhodius*. 2nd rev. ed. Leiden; Boston: Brill.
- Paraskevaides, H. A. 1984. *The Use of Synonyms in Homeric Formulaic Diction*. Amsterdam: A. M. Hakkert.
- Parry, Adam. 1956. 'The Language of Achilles'. *Transactions of the American Philological Association* 87: 1–7.
- . 1962. 'Homer, The Odyssey'. *The Fat Abbott*, 48–57.
- . 1966. 'Have We Homer's Iliad?' *Yale Classical Studies* 20: 177–216.
- . 1972. 'Language and Characterization in Homer'. *Harvard Studies in Classical Philology* 76: 1–22.
- . 1989. *The Language of Achilles and Other Papers*. Oxford: Clarendon Press.
- Parry, Milman. 1928a. 'L'épithète traditionnelle dans Homère: Essai sur un problème de style homérique'. Dissertation (Docteur-ès-Lettres), Paris. http://nrs.harvard.edu/urn-3:hul.ebook:CHS_Parry.LEpithete_Traditionnelle_dans_Homere.1928.

- . 1928b. 'Les formules et la métrique d'Homère'. Dissertation (Docteur-ès-Lettres), Paris. http://nrs.harvard.edu/urn-3:hul.ebook:CHS_ParryM.Les_Formules_et_la_Metrique_d_Homere.1928.
- . 1930. 'Studies in the Epic Technique of Oral Verse-Making. I. Homer and Homeric Style'. *Harvard Studies in Classical Philology* 41: 73–148.
- . 1932. 'Studies in the Epic Technique of Oral Verse-Making. II. The Homeric Language as the Language of an Oral Poetry'. *Harvard Studies in Classical Philology* 43: 1–50.
- . 1971. *The Making of Homeric Verse. The Collected Papers of Milman Parry*. Edited by Adam Parry. Oxford: Clarendon Press.
- Pavese, Carlo Odo. 1972. *Tradizioni e generi poetici della Grecia arcaica*. Filologia e critica 12. Roma: Edizioni dell'Ateneo.
- Pavese, Carlo Odo, and Federico Boschetti. 2003. *A Complete Formular Analysis of the Homeric Poems*. 3 vols. Lexis Research Tools 2. Amsterdam: Hakkert.
- Pavese, Carlo Odo, and Paolo Venti. 2000. *A Complete Formular Analysis of the Hesiodic Poems: Introduction and Formular Edition*. Lexis Research Tools 4. Amsterdam: Hakkert.
- Pennington, Jeffrey, Richard Socher, and Christopher Manning. 2014. 'GloVe: Global Vectors for Word Representation'. In *Proceedings of the 2014 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, 1532–43. Doha, Qatar: Association for Computational Linguistics. <https://doi.org/10.3115/v1/D14-1162>.
- Perek, Florent. 2016. 'Using Distributional Semantics to Study Syntactic Productivity in Diachrony: A Case Study'. *Linguistics* 54: 149–88. <https://doi.org/10.1515/ling-2015-0043>.
- Porter, Howard N. 1951. 'The Early Greek Hexameter'. *Yale Classical Studies* 12: 1–63.
- Porzig, Walter. 1954. 'Sprachgeographische Untersuchungen zu den altgriechischen Dialekten'. *Indogermanische Forschungen* 61: 147–69. <https://doi.org/10.1515/9783110243024.147>.
- Postlethwaite, N. 1979. 'Formula and Formulaic: Some Evidence from the Homeric Hymns'. *Phoenix* 33: 1–18. <https://doi.org/10.2307/1087847>.
- Powers, David M. W. 1998. 'Applications and Explanations of Zipf's Law'. In *New Methods in Language Processing and Computational Natural Language Learning*, 151–60. <https://www.aclweb.org/anthology/W98-1218>.
- Probert, Philomen. 2006. *Ancient Greek Accentuation: Synchronic Patterns, Frequency Effects, and Prehistory*. Oxford Classical Monographs. Oxford: Oxford University Press.
- R Core Team. 2017. *R: A Language and Environment for Statistical Computing*. Vienna: R Foundation for Statistical Computing. <https://www.R-project.org/>.
- Reagan, Robert Timothy. 1987. 'The Syntax of English Idioms: Can the Dog Be Put On?' *Journal of Psycholinguistic Research* 16: 417–41.
- Riggsby, Andrew M. 1992. 'Homeric Speech Introductions and the Theory of Homeric Composition'. *Transactions of the American Philological Association* 122: 99–114. <https://doi.org/10.2307/284367>.

- Risch, Ernst. 1955. 'Die Gliederung der griechischen Dialekte in neuer Sicht'. *Museum Helveticum* 12: 61–76.
- . 1979. 'Die griechischen Dialekte im 2. vorchristlichen Jarhtausend'. *Studi micenei ed egeo-anatolici* 20: 91–111.
- Rodda, Martina Astrid. 2016. 'L'Inno Omerico ad Afrodite (V) e la cronologia relativa dell'epica greca arcaica'. In Di Donato, Riccardo, ed. *Comincio a cantare. Contributi allo studio degli Inni Omerici*, 83–101. Anthropoi. Studi e materiali di Antropologia storica del mondo antico 13. Pisa: ETS.
- Rodda, Martina Astrid, Philomen Probert, and Barbara McGillivray. 2019. 'Vector Space Models of Ancient Greek Word Meaning, and a Case Study on Homer'. *TAL – Traitement Automatique Des Langues* 60 (3).
- Rodda, Martina Astrid, Marco Silvio Giuseppe Senaldi, and Alessandro Lenci. 2017. 'Panta Rei: Tracking Semantic Change with Distributional Semantics in Ancient Greek'. *Italian Journal of Computational Linguistics* 3 (1): 11–24.
- Rosenmeyer, Thomas G. 1965. 'The Formula in Early Greek Poetry'. *Arion: A Journal of Humanities and the Classics* 4: 295–311.
- Rozier, Catherine. 2018. 'Review of Bruno Currie, *Homer's Allusive Art*. Oxford; New York: Oxford University Press, 2016. Xiii, 343. ISBN9780198768821 \$110.00.' *Bryn Mawr Classical Review*, January. <https://bmcr.brynmawr.edu/2018/2018.01.27/>.
- Ruffell, Ian. 2012. *Aeschylus: Prometheus Bound*. Companions to Greek and Roman Tragedy. London: Bristol Classical Press.
- Ruijgh, Cornelis Jord. 1958. 'Les datifs pluriels dans les dialectes grecs et la position du mycénien'. *Mnemosyne* 11: 97–116.
- Ruiz Vizcaya, Carmen. 2001. 'Usos formularios en Las Argonáuticas de Apolonio de Rodas'. Tesis de doctorado en Humanidades, Departamento de Filologías Integradas, Facultad de Humanidades, Universidad de Huelva: ProQuest Dissertations Publishing. <http://search.proquest.com/docview/305490189/?pq-origsite=primo>.
- Russo, Joseph A. 1963. 'A Closer Look at Homeric Formulas'. *Transactions and Proceedings of the American Philological Association* 94: 235–47. <https://doi.org/10.2307/283649>.
- . 1966. 'The Structural Formula in Homeric Verse'. *Yale Classical Studies* 20: 219–40.
- Sanders, C. 2006. 'Saussurean Tradition in 20th-Century Linguistics'. In *The Encyclopedia of Language and Linguistics*, 2nd ed., 769–74. Amsterdam; Boston (MA): Elsevier. <https://doi.org/10.1016/B0-08-044854-2/01385-7>.
- Schmidt, J. H. Heinrich. 1876-1886. *Synonymik der griechischen Sprache*. 3 vols. Leipzig: Teubner. <https://catalog.hathitrust.org/Record/008619115>.
- Shorey, Paul. 1928. 'Review of "Les Formules et La Metrique d'Homère" and "L'Epithète Traditionnelle Dans Homère" by Milman Parry.' *Classical Philology* 23: 305–6.
- Sinclair, John. 1991. *Corpus, Concordance, Collocation*. Describing English Language. Oxford: Oxford University Press.
- Slings, Simon R. 2002. 'Oral Strategies in the Language of Herodotus'. In Bakker, Egbert J., Irene J.F. de Jong, and Hans van Wees, eds. *Brill's Companion to Herodotus*, 53–77. Leiden; Boston: Brill.

- Stephens, Susan A. 2015. *Callimachus. The Hymns*. Oxford; New York: Oxford University Press.
- Stolz, Benjamin A., and Richard S. Shannon, eds. 1976. *Oral Literature and the Formula*. Ann Arbor: Center for the Coördination of Ancient and Modern Studies, University of Michigan.
- Suttle, Laura, and Adele E. Goldberg. 2011. 'The Partial Productivity of Constructions as Induction'. *Linguistics* 49: 1237–69. <https://doi.org/10.1515/ling.2011.035>.
- Tabossi, Patrizia, Rachele Fanari, and Kinou Wolf. 2005. 'Spoken Idiom Recognition: Meaning Retrieval and Word Expectancy'. *Journal of Psycholinguistic Research* 34: 465–95. <https://doi.org/10.1007/s10936-005-6204-y>.
- . 2008. 'Processing Idiomatic Expressions: Effects of Semantic Compositionality'. *Journal of Experimental Psychology: Learning, Memory & Cognition* 34: 313–27. <https://doi.org/10.1037/0278-7393.34.2.313>.
- Titone, Debra A., and Cynthia M. Connine. 1994. 'Descriptive Norms for 171 Idiomatic Expressions: Familiarity, Compositionality, Predictability, and Literality'. *Metaphor and Symbolic Activity* 9: 247–70. https://doi.org/10.1207/s15327868ms0904_1.
- . 1999. 'On the Compositional and Noncompositional Nature of Idiomatic Expressions'. *Journal of Pragmatics, Literal and Figurative Language*, 31: 1655–74. [https://doi.org/10.1016/S0378-2166\(99\)00008-9](https://doi.org/10.1016/S0378-2166(99)00008-9).
- Tomasello, Michael. 2003. *Constructing a Language: A Usage-Based Theory of Language Acquisition*. Cambridge (MA); London: Harvard University Press.
- URI Planning Interest Group, W3C/IETF. 2001. 'URIs, URLs, and URNs: Clarifications and Recommendations 1.0'. 2001. <https://www.w3.org/TR/uri-clarification/>.
- Vatri, Alessandro, and Barbara McGillivray. 2018a. 'The Diorisis Ancient Greek Corpus'. <https://doi.org/10.6084/m9.figshare.6187256.v1>.
- . 2018b. 'The Diorisis Ancient Greek Corpus'. *Research Data Journal for the Humanities and Social Sciences* 3: 55–65.
- Vergados, Athanassios. 2012. *The Homeric Hymn to Hermes: Introduction, Text and Commentary*. Berlin: De Gruyter.
- Vierros, Marja. 2018. 'Linguistic Annotation of the Digital Papyrological Corpus: Sematia'. In Reggiani, Nicola, ed. *Digital Papyrology II: Case Studies on the Digital Edition of Ancient Greek Papyri*, 105–18. Berlin; Boston (MA): De Gruyter.
- Vinson, David P., Mark Andrews, and Gabriella Vigliocco. 2014. 'Giving Words Meaning'. In Goldrick, Matthew, Victor S. Ferreira, and Michele Miozzo, eds. *The Oxford Handbook of Language Production*, 134–51.
- Visser, Edzard. 1988. 'Formulae or Single Words? Towards a New Theory on Homeric Verse-Making'. *Würzburger Jahrbücher für die Altertumswissenschaft* 14: 21–37. <https://doi.org/10.11588/wja.1988.0.26782>.
- . 2011. 'Review: Rainer Friedrich: Formular Economy in Homer. The Poetics of Breaches. Stuttgart: Steiner 2007 (Hermes. Einzelschriften. 100)'. *Gnomon* 83: 1–10.
- Wackernagel, Jacob. 1924. *Vorlesungen über Syntax, mit besonderer Berücksichtigung von Griechisch, Lateinisch und Deutsch*. Vol. Zweite Reihe. Basel: Birkhäuser.

- West, Martin L. 1966. *Hesiod: Theogony. Edited with Prolegomena and Commentary*. Oxford: Clarendon Press.
- . 2001. 'West on Nagy and Nardelli on West'. *Bryn Mawr Classical Review* 2001.09.06. <https://bmcr.brynmawr.edu/2001/2001.09.06/>.
- . 2011. *The Making of the Iliad: Disquisition and Analytical Commentary*. Oxford: OUP Oxford.
- . 2012. 'Towards a Chronology of Early Greek Epic'. In Andersen, Øivind, and Dag T.T. Haug, eds. *Relative Chronology in Early Greek Epic Poetry*, 224–41. Cambridge; New York: Cambridge University Press.
- . 2014. *The Making of the Odyssey*. Oxford; New York: Oxford University Press.
- Willi, Andreas, ed. 2019. *Formes et fonctions des langues littéraires en Grèce ancienne: neuf exposés suivis de discussions*. Entretiens sur l'antiquité classique 65. Vandœuvres: Fondation Hardt pour l'étude de l'antiquité classique.
- Wray, Alison. 2002. *Formulaic Language and the Lexicon*. Cambridge: Cambridge University Press.
- Wulff, Stefanie. 2008. *Rethinking Idiomaticity: A Usage-Based Approach*. Corpus and Discourse. Research in Corpus and Discourse. London; New York: Continuum International Publishing.
- . 2013. 'Words and Idioms'. In Hoffmann, Thomas, and Graeme Trousdale, eds. *The Oxford Handbook of Construction Grammar*, 274–88. Oxford; New York: Oxford University Press.
- Zipf, George Kingsley. 1949. *Human Behavior and the Principle of Least Effort: An Introduction to Human Ecology*. Cambridge (MA): Addison-Wesley Press.