

1

Introduction

In short, this book explores how and to what extent morphosyntactic variability in traditional British English dialects is structured geographically. Taking an interest in the forests rather than in individual trees, the study is concerned with *aggregate* dialectal (that is, morphosyntactic) variability among the measuring points investigated. We address this variability by establishing text frequencies of dozens of morphosyntactic features in a major naturalistic dialect corpus that covers dialect speech in over thirty counties all over Great Britain. Utilizing state-of-the-art dialectometrical analysis methods and visualization techniques, the study is original both in terms of its fundamental research question (“What are large-scale patterns of grammatical variability in traditional British English dialects?”) and in terms of its methodology (CORPUS-BASED DIALECTOMETRY).

1.1 Rationale, method, and objectives

The study proceeds from the fact that we know next to nothing about aggregate morphosyntactic variability in British English dialects. While it is known that “every corner of the country demonstrates a wide range of grammatically non-standard forms” (Britain 2010, 53), we note that the bulk of the literature on dialect grammar consists of atomistic single-feature studies, and the handful of studies that have taken an aggregate approach typically focus on lexis and, in particular, phonology, but *not* morphosyntax. In this connection, it should also be noted that there is an oft-implicit notion in large parts of the dialectological community that morphological and (particularly) syntactic variation is not really patterned geographically. For example, Lass (2004, 374) contends that “English regional phonology and lexis ... are generally more salient and defining than regional morphosyntax” (for similar views, see Wolfram and Schilling-Estes 1998, 161 and, in the realm of German dialectology, Löffler 2003, 116).

To address these gaps and prejudices, the study utilizes a methodology we take the liberty to dub CORPUS-BASED DIALECTOMETRY. As a branch of geolinguistics, DIALECTOMETRY proper is concerned with measuring, visualizing, and analyzing aggregate dialect similarities or distances

2 Introduction

as a function of properties of geographic space; for seminal work, see Séguy (1971) (the paper that sparked the dialectometry enterprise); Goebel (2006), Bauer (2009), and Goebel (2010) (the “Salzburg School of Dialectometry”); and Nerbonne et al. (1999), Heeringa (2004), and Nerbonne (2006) (the “Groningen School of Dialectometry”). Whereas practitioners of traditional DIALECTOLOGY are dedicated to the study of “interesting” – typically phonological or lexical – dialect phenomena, one feature at a time in a handful of dialects at most, dialectometrical inquiry endeavors to identify “general, seemingly hidden structures from a larger amount of features” (Goebel and Schiltz 1997, 13). This means that dialectometricians put a strong emphasis on quantification, cartographic visualization, and exploratory data analysis to infer patterns from feature aggregates. Empirically, the bulk of the dialectometrical literature draws on linguistic atlas material as its primary data source. For example, Goebel (1982) investigates joint variability in 696 linguistic features that are mapped in the *Sprach- und Sachatlas Italiens und der Südschweiz* (AIS), an atlas that covers Italy and southern Switzerland; Nerbonne et al. (1999) analyze aggregate pronunciational dialect distances between 104 Dutch and North Belgian dialects on the basis of 100 word transcriptions provided in the *Reeks Nederlands(ch)e Dialectatlassen* (RND). Some dialectometricians have also relied on dialect dictionaries (for example, Speelman and Geeraerts 2008). Against this backdrop, Leinonen (2008), Grieve (2009), Heeringa et al. (2009), and Auer et al. (forthcoming) are rare examples of dialectometrical-geolinguistic work which bases claims about aggregate accent differences on the analysis, auditory or acoustic, of actual speech samples. In any case, given that most dialect atlases – the *Survey of English Dialects* is a good example – and dictionaries focus on lexis and pronunciation at the expense of syntax and morphology, it should surprise nobody that much of the dialectometrical literature drawing on such material is biased towards lexis and pronunciation at the expense of morphological and, especially, syntactic variation (but see Spruit 2005, 2006; Spruit et al. 2009 for some recent atlas-based yet syntax-centered dialectometrical work).

In an attempt to overcome this bias, the present study seeks to marry the qualitative-philological jeweler’s-eye perspective inherent in the analysis of naturalistic corpus data with the quantitative-aggregational bird’s-eye perspective that is the hallmark of dialectometry. This synthesis is desirable for two principal reasons. First, *multidimensional objects, such as dialects, call for aggregate analysis techniques*. That rigorous dialectology requires aggregation (in dialectological parlance, *bundling*) is by no means a new insight. Back in 1933 already, Bloomfield argued that

a set of isoglosses running close together in much the same direction – a so-called *bundle* of isoglosses – evidences a larger historical process and offers a more suitable basis of classification than does a single isogloss

1.1 Rationale, method, and objectives 3

that represents, perhaps, some unimportant feature. (Bloomfield [1933] 1984: 342)

The point is that so-called “single-feature-based studies” (Nerbonne 2009, 176), with their atomistic focus on typically just one feature, are fine when it is the features themselves that are of analytic interest. They are woefully inadequate, however, when it comes to characterizing multidimensional objects such as dialects or varieties (or relations between them). Outside linguistics, this sort of inadequacy is well known: Taxonomists, for instance, typically categorize species not on the basis of a single morphological or genetic criterion, but on the basis of many; economists assess the economic climate not on the basis of individual macroeconomic indicators (e.g. unemployment), but also consider inflation, GDP per capita, interest rates, and so on. The problem with single-feature-based studies – in linguistics as well as everywhere else – is that feature selection is ultimately arbitrary (see Viereck 1985, 94), and that the next feature down the road may or may not contradict the characterization suggested by the previous feature. Thus, there is no guarantee that different dialects will exhibit the same distributional behavior in regard to different features; isoglosses do not necessarily overlap (again, see Bloomfield [1933] 1984: 329). In addition, individual features may have fairly specific quirks to them that are irrelevant to the big picture. This is why “[s]ingle-feature studies risk being overwhelmed by noise, i.e., missing data, exceptions, and conflicting tendencies” (Nerbonne 2009, 193). For these reasons, the aggregate perspective – in Goebel’s parlance, “the synthetic interpretation” of linguistic data (Goebel 2006, 415) – is called for when the analyst’s attention is turned to the forest, not the trees. Aggregation mitigates the problem of feature-specific quirks, irrelevant statistical noise, and the problem of inherently subjective feature selection, and thus provides a more robust linguistic signal.

Second, *compared to linguistic atlas material, corpora yield a more realistic linguistic signal*. Atlas-based dialectometry typically aggregates observations such as “in the Yorkshire dialect, the lexeme *bus* is typically pronounced /bʊs/,” while corpus-based (that is to say, frequency-based) approaches seek generalizations along the lines of “in Nottinghamshire English, multiple negation is twice as frequent (6 occurrences per ten thousand words) in actual speech than in Yorkshire English (3 occurrences per ten thousand words).” The atlas-based method has undeniable advantages: We emphasize, in particular, a fairly widespread availability of data sources and superb areal coverage. By contrast, dialect corpora are a rarer species, and their areal coverage is typically inferior to dialect atlases. Having said that, as a data source, corpora appear to have two major advantages over dialect atlases. First and foremost, the atlas signal is categorical, exhibits a high level of data reduction, and may hence be less accurate than the corpus signal, which can provide graded frequency information and which is hence a more suitable

4 Introduction

method to get a handle on continuous linguistic variation (Holman et al. 2007; Anderwald and Szmrecsanyi 2009; Grieve 2009; Wälchli 2009). This highlights the most crucial difference between atlas-based and corpus-based dialectometry: *corpus-based dialectometry is frequency-based dialectometry in its purest form*.¹ The point is that although the exact cognitive status of text frequencies is admittedly still unclear (for example, we do not know about the precise extent to which corpus frequencies correlate with psychological entrenchment; see Arppe et al. 2010; Blumenthal 2011), we do claim that text frequencies better match the perceptual reality of linguistic input than discrete atlas classifications; this is true even though some varieties of atlas-based dialectometry derive – with considerable computational effort – some form of commonness weighting, for instance at the phonetic segment level, from the atlas signal. Second, we note that the atlas signal is non-naturalistic, meta-linguistic, and competence-based in nature. It typically relies on elicitation and questionnaires, and is analytically twice removed, via fieldworkers *and* atlas compilers, from the analyst – a limitation that is particularly acute when the atlas-based analysis is based on so-called “interpretive maps” (as opposed to “display maps”; see Chambers and Trudgill 1998, 25). By contrast, text corpora provide more direct, performance-based access to language form and function, and may thus yield a more realistic and trustworthy picture (see Chafe 1992, 84; Leech et al. 1994, 58). The well-known major intrinsic drawback of the corpus-based method is that it is unable to deal with rare phenomena (see Penke and Rosenbach 2004, 489; Haspelmath 2009, 157–158), and syntactic phenomena are a good deal rarer than, for example, phonetic phenomena (Chambers and Trudgill 1991, 291), which is why more text is needed to study (morpho)syntax than pronunciation. But then again, it is arguable whether phenomena that are so infrequent that they cannot be described on the basis of a major text corpus should have a place in an aggregate analysis at all.

Adopting a corpus-cum-aggregation approach exactly along these lines, this study taps the *Freiburg Corpus of English Dialects* (FRED), a sizable dialect corpus that samples old-fashioned dialect speech all over Great Britain (though we would like to mention right at the outset that the study also draws on two smallish reference corpora sampling Standard British and American English for benchmarking purposes). FRED is by design heavily biased towards elderly speakers with a working-class background – so-called NORMs (*non-mobile old rural males*) (see Chambers and Trudgill 1998, 29). The majority of the interviews in the corpus were conducted in the 1970s and 1980s, and most of the speakers were born around the beginning of the twentieth century. Dialect speech and dialect variability mirrored in

¹ This is why the present study’s methodology is somewhat similar to the pioneering, frequency-based dialectometry approach of Hoppenbrouwers and Hoppenbrouwers (1988, 2001); see Heeringa (2004, 16–20) for a discussion in English.

1.1 Rationale, method, and objectives 5

FRED therefore reflect an intermediate stage between the state of affairs represented in the *Survey of English Dialects* of the 1950s and 1960s and present-day dialect variability in Great Britain. In short, then, the present study is overwhelmingly concerned not with geographic variation in modern or mainstream or urban dialects, but with TRADITIONAL DIALECTS, which Peter Trudgill defines as follows:

Traditional Dialects are what most people think of when they hear the term dialect. They are spoken by a probably shrinking minority of the English-speaking population of the world, almost all of them in England, Scotland and Northern Ireland. They are most easily found, as far as England is concerned, in the more remote and peripheral rural areas of the country, although some urban areas of northern and western England still have many Traditional Dialect speakers. (Trudgill 1990, 5)

We hasten to add that the present study also investigates some vernacular varieties spoken in Wales and the Scottish Highlands, which are rather young and thus cannot count, strictly speaking, as “traditional” (cf. Trudgill 2004a, 15). Yet for the most part it is traditional dialects in Trudgill’s sense that really take center stage in this book. That these traditional dialects are dying out fast should be an added incentive to document them as best we can.

How does the study proceed empirically? In keeping true to the spirit of dialectometrical analysis, the goal was to base the analysis of dialect variability on as many morphosyntactic features as possible. The study thus defines a fairly comprehensive catalogue of fifty-seven features. These are essentially the usual suspects in the dialectological, variationist, and corpus-linguistic literature. The catalogue spans eleven major grammatical domains: (i) pronouns and determiners (e.g. non-standard reflexives), (ii) the noun phrase (e.g. preposition stranding), (iii) primary verbs (e.g. text frequencies of the verb TO DO), (iv) tense and aspect (e.g. the present perfect with auxiliary BE), (v) modality (e.g. text frequencies of epistemic/deontic MUST), (vi) verb morphology (e.g. non-standard weak past tense and past participle forms, such as *goed*), (vii) negation (e.g. *never* as a preverbal past tense negator), (viii) agreement (e.g. non-standard WAS), (ix) relativization (e.g. the relative particle *what*), (x) complementation (e.g. unsplit *for to*), and (xi) word order and discourse phenomena (e.g. lack of auxiliaries in *yes/no* questions).

Crucially, the book is not concerned with basing its empirical investigation on the mere presence or absence of individual features in particular locations. Instead, we seek to exploit the corpus material to the fullest and take an interest in graded text frequencies, feature by feature, in the interview material sampled in FRED. So, on the basis of fifty-seven feature frequencies extracted from texts from thirty-four measuring points (which yields a total of $34 \times 57 = 1,938$ continuous feature frequencies as data points), the study pursues two analytical avenues. For one thing, we utilize the well-known

6 Introduction

Euclidean distance metric to derive a measure of aggregate morphosyntactic distance between measuring points. The resulting distance matrix is then subjected to a range of state-of-the-art dialectometrical analysis techniques – various cartographic projections to geography (many of which heavily rely on color coding to capture the inherent dimensionality of morphosyntactic variability) as well as a number of correlative techniques. Second, we rely on Principal Component Analysis to embark on an analysis of feature interdependencies for the sake of exploring linguistic structure in the aggregate view and uncovering the layered nature of joint morphosyntactic variability in Great Britain.

On the interpretational plane, this book explores joint morphosyntactic variability in Great Britain as a function of geographic space. It is a time-honored axiom in dialectology and geolinguistics that geographic distance should predict linguistic distance (see Nerbonne and Kleiweg 2007, 154 and Chapter 5 for a discussion). Nonetheless, this axiom has not yet been put to a systematic test outside the realm of atlas-based dialectometry; additionally, the geolinguistic underpinnings of dialectal morphosyntactic variability are generally underresearched. Thus, the set of more specific research questions that guide the analysis in this book can be succinctly summarized as follows:

- (i) Does a frequency-derived measure of morphosyntactic variability in traditional British English dialects exhibit a geographic signal?
- (ii) If there is a geographic signal, exactly how are morphosyntactic distances and similarities distributed? Specifically: Are we rather dealing with a dialect continuum scenario or with a dialect area scenario?
- (iii) Do feature subsets make a difference, and what is the extent to which individual features gang up to create areal (sub)patterns?

In a nutshell, this book will suggest that aggregate morphosyntactic variability in British English dialects is indeed geolinguistically significant. More precisely, the distribution of morphosyntactic distances is such that Great Britain's morphosyntactic dialect landscape is neither a flawless dialect continuum, nor is it perfectly organized along the lines of dialect areas. Instead, we are dealing with a hybrid type, and with significant differences between the dialect network in England and the dialect network in Scotland. Lastly, we shall see that there are a number of feature bundles which create layers of geolinguistically conditioned morphosyntactic variability. The most important of those is a bundle comprising a comparatively large number of well-known dialect features which creates a very substantial South–North continuum.

1.2 Previous big-picture accounts

Naturally, we would like to validate our findings against as much previous scholarship along the lines of ours as possible. Alas, big-picture accounts

1.2 Previous big-picture accounts 7

based on the study of feature bundles (as opposed to so-called single-feature studies) in English dialectology are rare, and such accounts as exist are typically not dedicated to morphosyntactic variability. Be that as it may, this section canvasses the literature for aggregate approaches to *structural* variability in British English dialects, be it phonetic/phonological or morphological in nature (we ignore research on lexical variability on account of the fact that lexis is arguably “the least structured and most fluctuating plane of language” [Viereck 1986a, 725]). As much of the relevant research is empirically based on the *Survey of English Dialects*, we begin by adding a few introductory remarks about this particular atlas project. In addition to this line of more traditional dialectological inquiry, we also review in this section an experiment that investigates the perceptual side of dialect variability in Great Britain (Inoue 1996).

1.2.1 On the Survey of English Dialects

A project headed by Harold Orton and Eugen Dieth, the *Survey of English Dialects* (henceforth: SED) (Orton and Dieth 1962) was conducted between 1948 and 1961, primarily in rural England. The target informants were NORMS in 313 localities all over England, interviewed by nine trained fieldworkers. The SED database principally consists of the so-called *Basic Material*, which details responses to the extensive SED questionnaire that comprises no less than 1,326 questions and is organized into nine “books” (“The Farm,” “Animals,” etc.). After the publication of the *Basic Material* in the 1960’s, a number of linguistic atlases interpreting the SED material were published: *A Word Geography of England* (WGE) (Orton and Wright 1974), the *Linguistic Atlas of England* (LAE) (Orton et al. 1978), the *Atlas of English Sounds* (AES) (Kolb 1979), the *Structural Atlas of the English Dialects* (SAED) (Anderson 1987), and the *Computer Developed Linguistic Atlas of England* (CLAE) (Viereck et al. 1991). Though the LAE and the CLAE specifically cover lexical, pronunciatonal as well as morphological and syntactic variability, it seems fair to say that syntactic variability especially is generally given rather short shrift – not only in the interpretation atlases, but also in the SED questionnaire itself (on this point, cf. Shorrocks 2001, 1557). For contemporary evaluations of the SED endeavor, see Goebel and Schiltz (2006, 2356–2357), Sanderson and Widdowson (1985), and Viereck (1988, 267–268).

1.2.2 Nineteenth-century accent differences: Alexander Ellis’ survey of English dialects (1889)

We begin this literature review with a precursor of sorts to the SED. In the nineteenth century, Alexander J. Ellis, a gentleman scholar of independent means, conducted a rather monumental survey of dialect pronunciations

8 Introduction

in England, Wales, and Scotland (see Shorrocks 1991; Francis 1992; Ihalainen 1994 for contemporary reviews). The resulting database – deriving from a translation task (a “comparative specimen”), a shorter “dialect test,” and a “classified word list” – is documented in *The Existing Phonology of English Dialects*, the fifth volume of Ellis’ series *On Early English Pronunciation* (Ellis 1889). Ellis gathered data on 1,454 locations, an endeavor which took more than twenty years (see Ellis 1889, xvii–xix) and left us with one of the first dialect surveys “based on rich, systematically collected evidence” (Ihalainen 1994, 232).

The appendix of Ellis (1889) features two synopsis maps to project groupings of dialect features to geography – the first of their kind (Sanderson and Widdowson 1985, 36). The criteria used by Ellis to group dialects are pronunciational. Thus in the maps we find forty-two Districts, “in each of which a sensible similarity of pronunciation prevails” (Ellis 1889, 3), as well as six major dialect areas (“divisions,” in Ellis’ parlance): (1) Southern dialects, (2) Western dialects, (3) Eastern dialects, (4) Midland dialects, (5) Northern dialects, and (6) Lowland (Scots) dialects. In addition, the maps detail what Ellis refers to as “varieties, or parts of Districts separately considered” (Ellis 1889, 3), and “Ten Transverse Lines” (1889, 3), which map isoglosses of particular accent features (e.g. the pronunciation of words such as *some* or *house*).

1.2.3 SED-based analyses 1: Trudgill’s (1990) division of traditional dialects

Trudgill (1990) is one of the best-known contemporary dialectological accounts of dialect differences in Great Britain. To establish traditional dialect areas, Trudgill considers accent differences only. Specifically, he bases his classification on “eight major features of English Traditional Dialects which we can use to divide the country up into different dialect areas” (1990, 32): The vowels in *long* (/læŋ/ vs. /lɒŋ/), *night* (/ni:t/ vs. /naɪt/), *blind* (/blɪnd/ vs. /blaɪnd/), *land* (/lænd/ vs. /lɒnd/), and *bat* (/bat/ vs. /bæt/); postvocalic /r/, as in *arm* (/a:rm/ vs. /a:m/); *h*-dropping, as in *hill* (/hɪl/ vs. /ɪl/); and voicing, as in *seven* (/sevn/ vs. /zevn/). On the basis of the regional distribution of these features according to the SED, Trudgill presents a composite map that defines thirteen traditional dialect varieties (Northumberland, the Lower North, Lancashire, Staffordshire, South Yorkshire, Lincolnshire, Leicestershire, the Western Southwest, the Northern Southwest, the Eastern Southwest, the Southeast, the Central East, and the Eastern Counties) plus Scots. These thirteen varieties are grouped into six major dialect areas: (1) Scots, (2) Northern dialects, (3) Western Central (Midlands) dialects, (4) Eastern Central (Midlands) dialects, (5) Southwestern dialects, and (6) Southeastern dialects. As for higher-level groupings, Trudgill offers that “the major division of English dialects is into dialects of the North and dialects of the South, and that the

1.2 Previous big-picture accounts 9

boundary between them runs from the Lancashire coast down to the mouth of the River Humber” (1990, 35).

1.2.4 *SED-based analyses II: The Salzburg School of Dialectometry*

In the course of the past fifteen years, Hans Goebel and his collaborators have marshalled the full range of dialectometrical analysis techniques developed by the Salzburg School of Dialectometry (henceforth S-DM) to analyze the SED material, as available through the CLAE database (Viereck et al. 1991; see also next section). True to S-DM’s spirit of generating a large number of colorful maps, Goebel and his collaborators’ foray into dialect variability in England has left us with a sizable body of cartographic projections, mostly based on lexical and morphosyntactic SED material: similarity maps (Goebel 1997a, 2001, 2007), so-called parameter maps (Goebel and Schiltz 1997; Goebel 2001, 2007), dendrogrammatic cluster maps (Goebel 1997b, 2007), honeycomb maps (Goebel 2007), beam maps (Goebel 2007), and proximity profiles (for instance, Goebel 2007). In short, Goebel and his co-workers show that there is a fairly clear North–South split in the data in that, for instance, all Southern SED localities have their similarity minima in the North, and vice versa (Goebel 1997a). Second, beyond this very robust split, dialect area boundaries are, according to Goebel and Schiltz (1997) and also Goebel (1997b), largely in accordance with Peter Trudgill’s dialect division (see Section 1.2.3). Third, as for dialect integration, the South of England is on the whole more integrated than the North. In terms of morphosyntax specifically, though, it appears that there are actually two fairly well-integrated areas: “One is located in the South and centred near Salisbury; the second area lies further northeast in Nottingham/Lincolnshire” (Goebel and Schiltz 1997, 18).

1.2.5 *SED-based analyses III: Bamberg-type dialectometry*

In the 1980s and 1990s, a Bamberg research team headed by dialectologist Wolfgang Viereck generated a number of basic dialectometrical accounts of dialectal variability in England, all drawing on the SED-based *Computer Developed Linguistic Atlas of England* (CLAE) (Viereck et al. 1991) compiled in Bamberg. Some of the principal findings of this research include the following. Viereck (1986b, 243) reports that, all in all, there is substantial agreement between the geography of lexical, phonological, and morphological variability in dialectal England. Viereck (1997) and Händler and Viereck (1997) utilize a so-called “gravity center approach” (an actually fairly simple method designed to detect the geographic center of particular linguistic variants) and find that there is “a rather clear linguistic divide between the southeast and the southwest of England” (Viereck 1997, 3). What is more, Händler and Viereck (1997, 34) offer that “[t]he North is ... a

10 Introduction

dialectal area that is more strongly shaped lexically” than morphosyntactically. Viereck (1986b), mentioned above, is also concerned with establishing dialect divisions, relying on forty-five vocalic and consonantal features as well as fifty-three morphological features. Viereck (1986b, 243) distinguishes between the following five major dialect areas in England: (1) Northern dialects, (2) Lincolnshire plus East Anglia, (3) the West Midlands, (4) the Southeast of England, and (5) the Southwest of England.

1.2.6 SED-based analyses IV: *Shackleton (2007)*

In a paper entitled “Phonetic Variation in the Traditional English Dialects” (Shackleton 2007), Robert G. Shackleton Jr. offers a detailed account of aggregate phonetic variability in English English dialects whose “results largely corroborate standard characterizations in the literature” (2007, 87). The paper features cartographic visualizations in the spirit of the Groningen School of Dialectometry and an array of parametric and non-parametric statistical analysis techniques (regression analysis, Multidimensional Scaling, Cluster Analysis, and Principal Component Analysis). Empirically, Shackleton (2007) draws on the original SED database as well as on the SED-derived *Structural Atlas of the English Dialects* (SAED) (Anderson 1987) to construct two different datasets. The feature-based dataset rests on a set of fifty-five words (and their pronunciation) in English English dialects, while his variant-based dataset “summarizes over 400 responses, grouping them into 199 variants of thirty-nine phonemes or combinations of phonemes” (2007, 36 and 38). We may summarize the insights afforded by Shackleton (2007) as follows: To begin with, phonetic variation in traditional English English dialects is, on the whole, “not very systematic, but instead tends to involve largely uncorrelated variations that, in some areas, coalesce into patterns that appear more systematic” (2007, 42). Second, a small number of variants accounts for most usage, in accordance with Kretzschmar and Tamasi’s (2002) *A-curve* principle (2007, 40). Third, correlating as-the-crow-flies distances with phonetic distances, Shackleton finds that both feature-based linguistic distances and variant-based distances correlate quite robustly with geographic distances ($r \geq .7$) (2007, 47–48). Shackleton’s more sophisticated subsequent regression analysis shows that geographic distance and dialect area membership explain circa 77 percent of the variability in linguistic distances (depending on the dataset used) – and in regression analysis, geographic distance alone accounts for more than half of the variability (2007, 61). Thus, for explaining phonetic distances it is apparently geographic separation between the dialect locations that counts – rather than dialect area affiliation (2007, 64). This comparatively weak explanatory potency of dialect area membership notwithstanding, Shackleton applies multiple cluster analyses to partition the SED locations into the following seven major regions (2007, 52–53 and 88–89): (1) the Far