

Summing Up and Looking Ahead

This chapter will draw together a series of major findings from the chapters that preceded it. Rather than simply listing once again the conclusions drawn in the previous chapters, our goal here is to identify major findings that we feel are important for the field of learner corpus research, as well as for researchers interested in SLA and use in general.

Our first point returns to a motivation for this book introduced in Chapter 1. We wanted to work with MDA, but flex it in such a way that we could look at coherent linguistic units that would be problematic for standard MDA to deal with; namely, the discourse unit and the turn. While the technique we use, short-text MDA, had been shown to work on short sequences of data, most notably Twitter data,¹ we wanted to subject our hypothesis that MDA could work on texts of varying length to a severe test.² Our data represented three challenges to the technique. Firstly, we wanted the technique to work successfully at different levels. Rather than simply work with texts of a similar length – for example, tweets – we wanted to use the same technique to study small textual units of different lengths. Secondly, we wanted to use the technique to look at data in which different levels of proficiency were recurring features, with the associated challenge of differing levels of proficiency in any one file and across files (e.g. Learner A in a B2 test achieves three different grades: A in conversation, B in discussion and D in interaction; Learner B in a B1 test achieves three different grades: B in conversation, A in discussion and A in interaction). This diversity could, in principle, have had an impact upon the clusters of features which come together to compose a particular function, in turn obscuring our view of that function. Finally, we wanted to use short-text MDA to reveal differences between different situational contexts of use—that is, between tasks

¹ Twitter, a short messaging service, was rebranded as X in 2023.

² In so doing, we were following principle 34 of McEnery and Brezina (2022: 101); that is, we were putting the hypothesis to a severe test.

and casual conversation. In doing so, we wanted the technique to work well across contexts as well as across mixed L1/L2 and L1/L1 interactions in different situations. How did short-text MDA deal with these challenges?

Our answer to that question should be framed by the studies undertaken – we used the technique across three corpora and across two levels, the uppermost micro-structural level (the turn) and a macro-structural level (the discourse unit). In addition, within the corpora related to the GESE exam, we carried out the macro-structural analysis across a range of tasks, with those tasks being distinct from the general conversational situation represented by the Spoken BNC 2014. In each case, we interpreted dimensions to discover whether we could reliably move from form to function across a range of dimensions which express the variation of usage in those corpora and in those situations. In all of the analyses, the application of the technique led to interpretable dimensions – the clustering of units (turns or discourse units) that had identifiable functions associated with them. In all cases, as our analysis proceeded from the first dimension, the cut-off point (i.e. the point at which the interpretation of dimensions was no longer possible) was clear. For example, for the discourse unit analysis of the TLC, the last interpretable dimension was five; beyond that, the dimensions were not interpretable. The classifications were further tested through replication (see footnote 3, in Chapter 3). Our analyses of the 100 prototypical examples for any given function arising from a dimension interpretation were undertaken in two stages. First, the top fifty were analysed. The categories derived from this analysis were then tested on the following fifty prototypical examples to assess the replicability of our findings. Each study replicated as it was expanded, and the analysis of the prototypes proved a good guide to the meaningful analysis of both the form and function of the interpretable dimensions in our data. The dimensions revealed at the micro- and macro-structural levels helped us gain insight into how the micro-structural functions and the macro-structural functions interacted and, more generally, how pragmatics within the conversation at the macro-structural level operated. The exploration of the TLC Short function also showed the value of multiple explorations of what appeared to be the same function in different contexts – in that case, our wider exploration of the Short function led us to better understand how that function was operating and to propose a change of the categorisation of the function to Discourse Management.³

³ Note that the Short interpretation was initially considered to be replicated in the TLC discourse unit study. However, the subsequent analyses of Dimension 1 in both the TLC L1 and BNC analyses

So our categorisations, while successful, were also undertaken in a spirit of seeking falsification and reformulation—a process that clearly served us well in the case of the TLC Short function.

These observations allow us to reflect on the differences that the short-text MDA technique revealed in the repertoire of functions used in each of our three corpora. The differences are relatively few and none of the differences appear to be material for the operation of short-text MDA. Our use of short-text MDA in this book has shown it, repeatedly, to be a robust analytical procedure. However, the severity of our test was such that the limits of short-text MDA began to show with the analysis of discourse units. Notably, as texts get longer, the presence or absence of linguistic features, especially high-frequency features (e.g. general nouns, general verbs, positive interjection), is, as expected, less meaningful. This is because such high-frequency features are present in nearly all texts. The gradual preference for presence in the analyses arises from the greatest influence on the presence of features – the length of the text. As a result, in short-text MDAs of longer texts, most texts have many present linguistic features with few absences. This means that the variation between texts is no longer in the presences of features, but more specifically in the absences of features – what distinguishes a descriptive text is no longer its presence of BE as a main verb, attributive and predicative adjectives, *inter alia*, but, rather, its absence of past tense verbs and public verbs. Consequently, the MCA reveals dimensions largely comprising of the absence of linguistic features, rather than the presence of features. Interpreting the function of absent features is complicated – is it the opposite function of the presence of that feature? Or is it a different function altogether? For instance, the co-occurrence of a predicative adjective and BE as a main verb can be indicative of a stance-encoding function, so is the absence of both features an absence of such a function? Of course, it is convenient to think of language working like this, and indeed it often does. But sometimes absences can be random and not logically oppositional. Moreover, a dimension underpinned by an absence of a particular linguistic feature (e.g. negative interjection) can inherently bring noise into the analysis, as texts with those absences may be pushed towards the prototypical end of the dimension, but be functionally unrelated to the other texts where the absence of that feature is an intrinsic part

represented, in effect, a further replication of the classification. It was this that led to the reclassification. This shows that triangulation in corpus analyses of this sort (see Baker and Egbert (2016) for a fuller exploration of triangulation) represents a severe test in itself and, in this case, that led to a hypothesis (the interpretation of Dimension 1 of the TLC) being revised.

of a specific function. For the analysis of discourse units, there was, accordingly, a notable increase in the number of absences of features that were strongly associated with the dimension. Yet our results presented throughout demonstrate that this did not stop us from revealing communicative functions. But it did, in our experience, make the process of interpretation more difficult. Other researchers should thus consider whether short-text MDA is the most appropriate technique for their analyses. For us, we can say that short-text MDA suited our needs, but if our texts were on average slightly longer (e.g. >200-word tokens), then short-text MDA may not have been appropriate. So while traditional MDA may not be appropriate for texts shorter than 1,000 words, short-text MDA may not be appropriate for texts over 200 words in length. Consequently, a future direction for research in this area will be to formally test the limits of short-text MDA and traditional MDA and seek out new modifications of MDA to enable its application to any length of text, in particular this apparently difficult middle ground of texts that are too short for reliable MDA, but too long for a gainful use of short-text MDA.

A more positive note is that the facility that short-text MDA gave us to introduce supplementary variables meant that our analyses could easily be altered to reflect situational variation, with successful analyses taking into account different situations of production as well as other metadata variables (e.g. grade awarded, level of exam) without our use of the technique encountering any unpredictable difficulties. Of course, as we combined variables, we did encounter the problems outlined at the start of the book – data-sparsity issues in some combination of variables that generated results which, when considered in the light of the limited evidence available to interpret them, led to us setting those results aside. But overall, with this predictable exception, the technique worked well across the range of comparisons we made.

Finally, with regard to the varying levels of proficiency of the speakers in the corpus, this did not represent a challenge for the short-text MDA technique at all. It did, of course, mean that on occasion we saw functions used which were directly linked to level of proficiency. But that is a desirable outcome, as it meant that we were able to group the data in a way that was linguistically meaningful. What did not happen is that the technique failed to work because of grammatical infelicities in the language produced, whether that be because of the nature of spoken interaction or because of issues of proficiency. As a way of approaching interaction between heterogeneous speakers, the technique was successful in classifying data into functional groupings that gave us insight into the data.

To develop the last point, why might it be that grammatical errors, for example, do not cause the short-text MDA to encounter real problems in terms of producing meaningful functional classifications? As part of our answer to that we must first make an observation. While we have talked so far of the macro-structural level (discourse units) and micro-structural level (turns) being the units of organisation of our data and, hence, the level at which we assign functions, the truth of the matter is more subtle than that. MDA works on form-to-function relations in which the distribution of low-level features is used to make an inference about the function of a high-level feature – so, for example, we look at morphosyntax (the results of our part-of-speech tagging) and how the categories of our morphosyntactic description of the data distribute relative to one another within a linguistically meaningful unit (e.g. a discourse unit) to derive the function of that linguistically meaningful unit. This leads to two observations. Firstly, on the labelling of the high-level unit, there is a direct assertion that that unit is realised by a configuration of features at the micro-structural level (typically the word level) viewed through an aggregating abstraction—in this case, the limited range of labels arising from our part-of-speech tagging. So, rightly in our view, the micro-structural level is ineluctably bound to the macro-structural level as, through the micro-structural, the macro-structural is realised.

Secondly, the shift to the level of linguistic abstraction allows this technique to work well – imagine if we did not have that abstraction. If we carried out our analysis on word form only, the scale of data we would need would be much greater but also the linguistic explanation of how the macro arises from the micro would be much more difficult. The abstraction helps us to produce and explain our form-to-function mapping. So, the real problem presented to short-text MDA by learner language is probably at the level of how that process of assigning the categories of the abstraction to the data works. If our part-of-speech tagger handles our data poorly, producing a linguistic description of it that is inaccurate or random, then real problems would arise. We do not see this in our data. For both L1 and L2 speech (at least at the levels studied) we saw no evidence that the data was so ill-formed that our process of morphosyntactic analysis was reduced in value or credibility to the extent that meaningful functional groupings did not emerge from the data. In other words, if the part-of-speech tagging on which the short-text analysis rests is reliable enough, then it follows that the form-to-function mapping arising from it is likely to be reliable also. Our analyses provide abundant evidence that this is the case.

In that context, our work on meshing the turn and discourse unit levels in our analysis can now be slightly reframed. Rather than exploring how the micro- and macro-levels interact, we are looking at how functions, both arising from a low micro-structural level (the word), can be realised at different levels: the uppermost micro-structural (turn) and the macro-structural (discourse unit). This is a subtle shift of perspective as it means that, when we are seeing how the two mesh, we can operate on a presumption that such a meshing should occur as the functions at those levels are projected from a common base – morphosyntax. Accordingly, our classifications represent a spectrum of interrelated classifications, arising from a common starting point, but capable of being viewed from distinct perspectives. In that context, the linguistic plausibility of the units of observation that we use is crucial – we could land upon an arbitrary segmentation of the text, for example, every twenty words, and let our technique run on that basis. It would be very hard to defend that as linguistically meaningful. It is in that context that we appealed to a well-established unit that linguists have found plausible – the turn – and decided to look at that as well as another unit that has been posited yet has remained more elusive – the discourse unit. We also applied other optics to the observations, but all of them were linguistically plausible; that is, we introduced variables into our observation of the data which, *a priori*, have a linguistically plausible role to play in a situated view of language production—notably, task. It is in that context, taking a primarily linguistic approach to the segmentation of our texts, that we use short-text MDA to produce linguistic insights into our data. So, the success of short-text MDA, from our perspective, is not simply that it produces results. It is that in using it in a way that forces a linguistic categorisation onto the data (the different units we use as the basis of our analysis, the morphosyntactic categories we apply to the data) and works with a linguistically plausible hypothesis – that there is a form-to-function relation that links the micro to the macro – we derive from our data linguistically meaningful insights. MDA has been doing this successfully with longer texts for many years. However, because of the use of MCA in short-text MDA, we can now see that such insights may be derived from short linguistic units of varying length.

The strength of the technique has been seen in the repertoire of functions that we have been able to derive from each corpus. Those in turn, because they were extracted on the same basis—that is, using the same morphosyntactic features and procedure for form-to-function mapping—have been able to be compared across the corpora. That was the central

concern of Chapter 7. It is the capacity to make these comparisons, focused on linguistically plausible functional analyses, that enabled us in that chapter to explore the impact of situation on the repertoire of functions used and the possible role that L1 and L2 use may play in the composition of that repertoire.

Being able to use the same technique to look at different levels, as was done in Chapters 2–7, has another benefit. Given that the same micro-structural features (a fixed set of morphosyntactic features) and dimension-clustering technique was being used to look at both turns and discourse units, the possibility could have arisen that the two analyses would prove to be identical – a little like having a block of cheese. Think of the discourse unit as the block of cheese. If we take a slice from the cheese, its nature does not vary; It is still cheese. Likewise, if we looked at the turns within a discourse unit which had a specific function, it may be that we would see within it smaller constituent units with the same function. The short-text MDA allowed us to see that while sometimes that happens, it does not always. The picture is mixed. Likewise, the later initial analysis of the meso-structural level suggested that, in one discourse function at least (Discourse Management), it was possible to see function shifting coherently at this meso-level. So, the technique has given us an insight into discourse units – they are composed of turns, but those turns may, or may not, have the same function as the macro-structure of which they are part. Of course, in interaction we should expect to see some diversity – even though the speakers may be acting cooperatively, they still have individual perspectives and may, on occasion, have competing goals. They also, in the TLC, have markedly different levels of proficiency. But what we saw went beyond what these sources of variation might produce – some of the macro-structures seemed to be composed of assemblages of micro-structures—turns—which, while functionally heterogeneous, worked together to produce a macro-structure with a distinct function.

There is another sense in which our study could have been different. As noted, throughout the study of discourse units in spoken interaction, we are looking at a linguistic unit which is co-constructed. A possible point of departure for this book would have been to focus exclusively on L2 speech, ignoring the contributions made by the L1 interlocutor. This would certainly have simplified the analysis – we could have focused on the L2 production alone and have avoided some questions, such as those explored in Chapters 5 and 6, to a large extent. Tempting as it may have been to do this, it would have been the wrong decision – the study of spoken interaction cannot reasonably be reduced to the contributions

of just one speaker. While people may do this for ease, our view is that it is wrong to ignore the dialogic nature of spoken interaction, precisely because the macro-structure within which the language of the interlocutors is produced is co-constructed. The creation of a discourse unit typically represents a collaborative act. To look at only one set of contributions to that collaborative act is to misunderstand it. As shown in the analyses in Chapters 3 and 4 in particular, the speakers work together to co-produce discourse units. There is a process of negotiation and, at times, direction in the interaction, aimed at producing discourse unit functions. Consider the following example:

- (86) s: if you travel inevitably you have to eat </pause> so that's one thing like
 I I I mean personally I love my like you know my food
 s: and then there's some places that you know for example I went to
 Barcelona and I
 s: wanted to get African food and
 s: they said to me if you wanna get African food you need to travel this far
 s: so I had no choice than to eat like what they had available so

These are the examinee contributions from one discourse unit in the Conversation task in file 45 of the TLC L1. This seems coherent and, looking at the turns, one might conclude that it is indeed possible to study the turns of just one speaker and get a good sense the functions of discourse in spoken interaction. This appears to be something akin to an Informational Narrative. However, our view of the interaction is changed when we look at the discourse unit's first turn, which is produced by the examiner:

- E: I can see <unclear=what> you're saying but erm I'm not sure I entirely agree
 because a lotta people they they go on holiday they erm go to maybe a beach
 resort with people from their own country and </pause> I don't really think
 that they learn anything about the local culture

The student is trying to persuade the examiner that they are wrong, not simply provide them with information in narrative form. They are doing that to achieve a rhetorical purpose. Studying the examinee's turns out of context can distort our view of the data, then, and consequently may distort our interpretations.

So, while we might have been able to strip out the examiner speech from the TLC, if we had done so we would have misunderstood what was happening in the interaction. To expand on the example given, if we had stripped away the scaffolding that Information-Seeking turns from the examiner provided for some students, we would not have understood the motive for the production of some turns by a learner, nor would we have understood

who was guiding the student towards the production of a specific discourse unit function. The examiners influenced the students' utterances and without a clear sight of that interaction we may even have found ourselves wondering why some poorly graded students received poor grades – we need to see that their productions were scaffolded to understand that.

This is not just true of learner data, for even in everyday conversations, as we saw in Chapter 7, L1 discourse units are similarly co-constructed with considerations related to pragmatics serving as prominent a role in L1/L1 interaction as they do in L1/L2 interaction. To move to a simile again, looking at a discourse unit is not like looking at a bag filled with red and white balls, where we may choose to study either the red or the white balls separately and experience no difficulty. It is much more like a house built of red and white bricks, with two builders, each laying a brick of one colour, working together in real time to construct something where the decisions of one builder impact on the choices available to the other. As with all building projects, sometimes things go wrong, but by and large pragmatics, and in particular the cooperative principle, operates to ensure that the speakers work together to achieve an outcome – and that is a discourse unit function that they produce as a collaborative act, guiding the choice of function through their choice of turns with specific functions that they contribute.

When co-constructing a specific discourse unit, the speakers are mindful of context. Their choices are situated. In the TLC and TLC L1 we gained powerful evidence, in our comparison of them in Chapter 7, that it is the situation in which the speakers operate, more than their status as an L1 or L2 speaker, that militates in favour of them choosing to co-construct specific discourse unit functions in relation to specific tasks. Yet even in the Spoken BNC 2014, what we see, we would argue, is a set of discourse functions that are suitable for that situation (i.e. casual conversation at home). So, the speakers in our data are mindful of situation, and that produces what one may either view as a constraint on the range of discourse functions they might draw upon, or an encouragement to produce certain types of discourse functions appropriate to the situation. It is notable that, at the B1/B2 level explored in this book, the functions chosen by L2 speakers in the situation of the GESE exam are almost identical to those chosen by L1 speakers in the same situation. Situation, not language proficiency *per se*, is the driver of behaviour here.

This allows us to turn to a question that may have occurred to some readers. Has this book told us about L2 speech (and L1 speech), or has it simply told us a lot about a single exam—the GESE exam? Moreover,

given that the presence of the examiner is a feature shared by both the TLC and the TLC L1, are we seeing congruence in function because the examiner, not the situation more generally, is forcing that congruence? Let us deal with these questions in reverse order. We have gathered plentiful evidence, as the book has proceeded, of autonomy in both L2 and L1 speakers. For example, both select narratives, though the task they are engaged in does not require them to do so. While we may, as noted, see the examiner prompting the production of a narrative sometimes, the majority case is that the examinee chooses and relates the narrative. We have some functions which are often mainly produced either by the examiner (e.g. Informative and Instructive) or the examinee (e.g. Informational Narrative). On the other hand, there are also functions which are more equitably split between the two (e.g. Seeking and Encoding Stance). If we continue with our focus on Dimension 4 in the TLC discourse unit analysis, on the one hand, we see two functions—one of which is overwhelmingly examiner speech and another which is overwhelmingly student speech. On the other hand, there are also discourse unit functions that appear to be more evenly shared between the two (e.g. Seeking and Encoding Stance). This strongly suggests that the two speakers—the examiner and the student—are performing different roles in the discourse, not that the examiner is an ever-present force nudging and controlling the student to produce specific functions continually. While we saw a minority of cases where a specific function was elicited in our qualitative study, more typically what we saw in the quantitative analyses, when we explored examples, was an examiner who was directive in introducing tasks, but then minimally invasive otherwise, unless they needed to act as a cooperative listener to scaffold the L2 speakers' output. What appears to be the greater force is the situation itself—specifically, the task—which elicits the same behaviours from the L1 and L2 examinees and is clearly an important variable in the studies of task in Chapters 2 and 3. Returning to the first question, there is no doubt, given the importance of the task, that the exam itself is an important framing for our observations; the tasks require certain functions to be performed. However, we find many of those functions in conversational English, so if we view the exam as a construct – a model – of conversational English, then we may say that while it does not exhibit the full range of discourse functions of conversational English, it does seem to require many of them. Hence, in the sense that the exam allows the assessment of students producing discourse units with functions that are well attested in conversational English, as a construct it is successful and, by extension, that construct allows us to

make claims about the performance of the students in conversation with L1 speakers beyond the exam.

Our focus in the chapter to this point has been upon the discourse unit – but what of the micro-structural level? It is without question that if we explored the TLC corpora at the micro-structural level, differences would emerge. As discussed in Chapter 5, sometimes when we see identical functions across the TLC and TLC L1, we see a smaller set of features creating that function in the L2 data compared to in the L1 data. In this regard, we can see the types of differences, related to grammatical proficiency, upon which a lot of learner corpus studies focus. But the study here shows that they obscure significant similarities. In pragmatics, the broad set of principles controlling interaction seem to be just as in evidence in the TLC data as they are in the TLC L1 data. Also, in terms of functions, we see congruence between the situated language use of both sets of speakers in both corpora. The point about pragmatics is particularly interesting. A question that could arise from the observation of grammatical differences between the TLC and TLC L1 functions could point towards some form of difference arising from the fact that the learners are still acquiring the grammar of the L2. The congruence at the functional level, we would argue, arises from the principles of pragmatics, and we see it most clearly in the discussion of grade 6 in Chapter 4, Section 4.2. Students with difficulties do not fail to communicate. A cooperative interlocutor helps them. While the degree of cooperation may be heightened by a learner who struggles to perform a specific function, that cooperative behaviour is a natural part of human communication and might explain why even students who struggle can contribute discourse units that they would otherwise fail to produce – because their conversational partner is naturally cooperative. So, communicative competence trumps proficiency in such a context, as the principles of pragmatics permit such an outcome.

So if, from the perspective of pragmatics, some of our findings are, in hindsight, easy to explain, then the appearance of some functions was a surprise. This book did not set out to be an investigation of narrative. We began with a bottom-up investigation of our data and, from that, a strong focus on narrative emerged. This may not be predictable from the perspective of current learner corpus research or pragmatics, but it is more explicable from the perspective of work on first language acquisition (FLA) and SLA, where the salience of narrative is more notable in the literature, as shown in Chapters 8 and 9. Narrative, in various forms, is present at the micro-structural (turn) level and the macro-structural (discourse unit) level. Narratives were not explicitly called for in either the

TLC or Spoken BNC 2014 datasets. They occurred largely spontaneously, though our qualitative study in Chapter 9 does show that, sometimes, narratives may be prompted. However, this should not surprise us because, as discussed earlier in this chapter, one speaker may influence the function selected by another speaker. What emerges from our analysis is that narrative is an important affordance in L1 and L2 communication. It certainly represents, in our opinion, an area of learner corpus research which is nearly absent from the literature, even though narrative is clearly an important affordance in both L1 and L2 conversation. Further, our qualitative study showed that narrative may be subject to selection pressure by cultural background. While learner corpus research is often concerned with cross-linguistic interference at the micro-structural level, cross-cultural interference at the macro-structural level is a research area that is almost virgin territory. Our hope is that this book, by demonstrating the importance of narrative and the potential it brings to study cross-cultural selection effects, may stimulate further work on this topic. Likewise, for studies in L1 and L2 research focused on narrative, the use of corpora in this book to explore narrative in spontaneous speech shows that studies of narrative carried out on small numbers of speakers and limited contexts in those fields could also be approached in larger datasets with more speakers using suitable corpora and techniques like short-text MDA. We would not, however, advocate simply replacing the methods used in those studies with corpus-based MDA. But we would argue that supplementing current approaches with our approach may produce a more rounded and scalable, approach to the study of narrative.

This book represents a step in a new direction. While the further steps are many, they are also promising. For example, a corpus of L2 casual conversations, outside of the exam context, would be of value. This would allow us to further test and, potentially, critique the construct that the GESE exam represents. If we were to build such a corpus, we could design it to cover interactions with both L1 and L2 speakers. Gathering it would certainly be possible – it could be produced using some of the techniques that made the construction of the Spoken BNC 2014 (Love et al., 2017) and LANA-CASE (Hanks et al., 2024) possible. While a challenging prospect, such a corpus is not an impossibility, and it would have the virtue of letting us explore spoken learner language in a range of contexts in which the learner happens to use their L2. This would almost certainly extend the scope of the study of situated language use by L2 speakers beyond what is possible using the TLC. In doing so, we hypothesise that the range of macro- and micro-structural discourse functions

evidenced in learner speech would expand, possibly covering some of those which were shown to be unique to the Spoken BNC 2014 in Chapter 7, but also perhaps to reveal new functions which are not yet clear in any of the corpora studied in this book.

Looking at the meso-level and seeing what patterning emerges when we view the discourse units through the functions of the turns that compose them is also another fruitful avenue for future research to pursue. We have seen this in practice but have not studied it systematically at scale in this book. This would be an obvious study to undertake, though we anticipate that this would not be a trivial undertaking, and that a study of similar length to this book would be the result.

Another future avenue of research that we will conclude by mentioning relates to expanding the range of proficiency levels represented in the corpora available to us, as researchers interested in learner discourse. At the time of writing, the team that wrote this book has just completed a corpus of exams from the grades of the Trinity GESE exam, allowing us to begin the exploration of the A1/A2 grades of the CEFR scale. That data will, once again, challenge short-text MDA and require a great deal of analysis to understand. However, while we have preliminary results, we must now conclude this book and leave those results for future analyses.