



Universiteit
Leiden
The Netherlands

Tunen syntax and information structure

Kerr, E.J.

Citation

Kerr, E. J. (2024, September 4). *Tunen syntax and information structure*. LOT dissertation series. LOT, Amsterdam. Retrieved from <https://hdl.handle.net/1887/4054916>

Version: Publisher's Version

License: [Licence agreement concerning inclusion of doctoral thesis in the Institutional Repository of the University of Leiden](#)

Downloaded from: <https://hdl.handle.net/1887/4054916>

Note: To cite this publication please use the final published version (if applicable).

CHAPTER 3

Methodology

3.1 Introduction

This chapter presents the methodology chosen for the study of the interaction of syntax and information structure in Tunen. I start by discussing the conceptual framing of the study (section §3.2). I then discuss the data collection methods used for the fieldwork (section §3.3), the use of secondary data (section §3.4), and the analysis tools used (section §3.5). Next, I give information on the data organisation and how the data and metadata have been archived (section §3.6). I also discuss the limitations of the methodology used, including the impact the COVID-19 pandemic had on data collection (section §3.3.4) and the more general limitations of the framework followed (section §3.7).

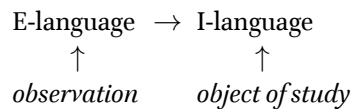
3.2 Conceptual framing

3.2.1 Positivism, parsimony, and falsifiability

As introduced in Chapter 2, the later chapters of this thesis are written within the formal syntax framework of generative linguistics, specifically the Minimalist framework (Chomsky 1993, 1995, 2000 *et seq.*), which builds on earlier generative theories developed as an alternative to structuralist and behaviourist approaches to language (Chomsky 1957, 1959, 1965, 1981, 1993; see e.g. Kertész 2017; Newmeyer 1980,

1986, 1996, 2022 for historical/historiographical overviews). Minimalism generally assumes a positivist epistemological stance in which language as an internalised cognitive system (*I-language*) is considered to exist in some form that is increasingly uncovered by means of scientific inquiry (Chomsky 1986; Isac and Reiss 2013). In contrast, the externalisation of language (*E-language*) — as mediated by various non-linguistic performance factors such as memory limitations and speech errors — is what we are able to directly encounter. In this way, linguistic research aims to uncover information about the internal rules of language — as understood to be part of cognition — working from E-language to I-language, as in (28) below.^{1,2}

(28) **Linguistic research workflow**



While positivism is a common epistemological stance within generative approaches to linguistics, often followed implicitly, a criticism is that such approaches may neglect the social aspect of language. It is well-known that language varies between speakers, with sociolinguistic variables such as gender, age, social class known to interact with how someone uses a language. Classical generative work is framed in terms of a monolingual speaker, when the reality for most of the world is one of multilingualism, with most speakers speaking multiple languages — as is true for Africa as a whole and the region around Cameroon more specifically (Di Carlo et al. 2016; Heugh 2019; Lüpke 2019). A large part of critique of generative approaches thus relates to the early presentation of Chomsky (1965) of an ‘ideal speaker-hearer’ within a monolingual speech community as the object of study:

“Linguistic theory is concerned primarily with an ideal speaker-listener, in a completely homogeneous speech-community, who knows its (the speech community’s) language perfectly and is unaffected by such grammatically irrelevant conditions as memory limitations, distractions, shifts of attention and interest, and errors (random or characteristic) in applying his knowledge of this language in actual performance.”

Chomsky (1965:3)

¹This is the type of linguistic research I pursue in this thesis, but study of E-language can also be considered to be linguistic research, under a more encompassing definition of ‘linguistic’ that includes the study of how language is manifested in society and in relation to cognitive factors.

²While I frame this discussion in terms of generative approaches, compare the Saussurean distinction between *langue* and *parole* and the broader scientific context of Cartesian dualism.

However, in the nearly 60 years that have passed since this statement was made, this narrow approach to linguistics has not only been rejected by non-generative theoreticians but has also been shifted within the generative paradigm (see Lohndal 2013; Lohndal et al. 2019; Scontras et al. 2015 for germane discussion). Contemporary work in the generative framework recognises the reality of multilingualism and increasingly embraces naturalistic field data as an important source of empirical foundations for linguistic theory. This thesis fits into this current approach to generative linguistics; in presenting generative analyses, I am not committing to the metatheoretical stance that was adopted in the early years of the framework. In order to give clarity on the combination of generative linguistics with field methods, I will explain in section §3.3 the specific approach to data collection I used in my fieldwork, discussing how I followed the workflow in (28) and weighed up the evidence for and against positing rules as part of Tunen speakers' I-language.

By adopting a Minimalist approach, I follow the logic that the linguist's role is to find the minimal rules that explain the data provided. This relates to the metatheoretical consideration of parsimony of analysis within scientific theory, as commonly formulated under the metatheoretical heuristic of Occam's Razor (29).

(29) **Occam's Razor**

"Entities are not to be multiplied beyond necessity." (Baker 2022:§2)

When applied to linguistics, this heuristic can therefore be expressed as the rule that, if two analyses both account for the data at hand, the one that requires the fewest number of rules is the one that should be chosen. I will use this argumentation thesis when deciding between competing analyses, specifically in terms of how Tunen's basic clausal syntax is built up (Chapter 6) and how discontinuous DPs should be modelled (Chapter 7). As I reflect on in Chapter 8, Occam's Razor crucially depends on the simpler analysis being able to account for the data, meaning that subsequent empirical investigation may lead to the simpler analysis no longer being sufficient on account of changes in the nature of the data to be accounted for (an aspect of theory formation sometimes termed *cyclicity*; see e.g. Kertész 2017). A goal of this thesis is to develop a clear formal model that can be tested in subsequent empirical investigation. I therefore mention potential alternative analyses that may end up being preferable if certain empirical facts do not hold.

One more general question is whether choosing the simplest set of rules commits the linguist to the statement that this set of rules is closest to how speakers/signers actually form and parse sentences. This is a much deeper point than I will go into in this thesis, but a basic argument can be made for the assumption that cognitive processes favour efficiency. In terms of IS, this idea of economy also

links to the idea introduced in Chapter 2 that conversational participants aim for a smooth transfer of information, with economy invoked as the null hypothesis. I also touch on the idea of economy in Chapter 7, in relation to iconicity of syntactic contiguity versus discontinuity.

Generative approaches aim for analyses to make predictions about the data which can be verified or falsified by further empirical investigation; falsifiability is a key metatheoretical criterion. I therefore highlight throughout the thesis which predictions the analysis presented makes, so that future research can test the predictions. If the predictions are accurate, this is evidence in favour of the analysis. If the predictions are inaccurate, then the analysis must be adapted to account for the new data. As this thesis is the first formal analysis of Tunen syntax and information structure, and as relatively little empirical work has been done on the effect of information structure on Tunen syntax, it is to be expected that some (or perhaps most) of the predictions will turn out not to hold. I therefore aim to make clear on exactly what empirical grounds and under which theoretical assumptions a particular analysis was made, so that the reader is free to consider the evidence and alternative approaches.

3.2.2 The use of grammaticality and felicity judgements

For the positivist framework adopted here, grammaticality judgements are crucial, as they are considered to give key insight into a speaker/signer's internal grammar, which is the object of study (28). The predictions of a model can be checked or falsified by grammaticality judgements, and so I present grammaticality judgements as evidence for and against certain rules being part of Tunen's grammar. Because of the focus on information structure, I also use felicity judgements.

There are several methodological advantages to using grammaticality and felicity judgements (Bower 2015; Vaux et al. 2006). The first methodological advantage is that elicitation enables the researcher to collect data on topics of interest within a limited timeframe available for fieldwork. This contrasts with corpus methods for natural speech data, where relevant examples may not show up simply due to the size of the corpus.

A second methodological advantage of grammaticality/felicity judgements is that they allow for negative evidence, i.e., confirmation of not only what is possible but also what is *not* possible. Such information on what is not possible — i.e., *negative evidence* — is not available from corpus methods, where the absence of evidence of a particular construction may give some indication of its acceptability but could equally be due to limited size of the corpus, as summarised in the aphorism that absence of evidence is not evidence of absence. The usefulness of

elicitation for providing negative evidence is one of the key arguments in favour of it as a methodological choice (Bower 2015; Vaux et al. 2006).

While grammaticality/felicity judgements are useful, they can vary both within and across speakers, for a variety of external reasons. It is therefore important that judgements are checked multiple times. Firstly, I checked judgements again with the same speaker, either later during the same session or in a different field session. Secondly, I checked judgements with different speakers, primarily in separate sessions but sometimes by having 2-3 consultants present in the same session (as indicated by '+' within the example presentation, e.g. [EE+EB 2228] indicates a form from a session conducted with both speaker EE and EB). In this study, I tested judgements multiple times during the same session with a consultant. I then reviewed the grammaticality judgements after the session and made hypotheses about particular grammatical rules to test in further sessions. For data that went against the hypothesis, I tested with the same speaker in a later session in order to see whether the judgement was stable. If so, I adjusted the hypothesis accordingly. For the key phenomena in this thesis, I tested with at least 3 speakers; for minor points, I tested with 2 speakers. Some points of interest to the reader may only have been tested by 1 speaker due to not being directly relevant to my research question at hand. I therefore used a fairly time-intensive methodology, with the advantage of improving the quality of the data in this thesis and the confidence with which I can make generalisations about Tunen grammar. The disadvantage is that this increased workload means I was not able to cover as many topics or work with as many different consultants as would be possible with a more shallow methodology. The choice to confirm judgements during elicitation sessions also means that elicitation recordings may be quite long for a relatively small number of sentences discussed.

One methodological tool was to purposefully mispronounce a given word or phrase, in order to test whether the consultants would correct me. If I was corrected, this provided evidence for the difference being meaningful in Tunen. For example, mispronunciations of [ɔ] as [o] were corrected by speakers while [p] instead of a speaker's [b] was not corrected. This provides evidence that /ɔ/ and /o/ are separate phonemes while [p] and [b] are allophones of the same phoneme /p/ (see Chapter 4 section §4.2).

The way one asks questions can influence the answer, which in certain epistemologies is interpreted as the answers arising through the research process rather than being part of an innate grammar. These issues affect the validity of the data collected and presented. This potential flaw is particularly important for the investigation of IS, where the judgement of interest may relate to felicity rather than grammaticality, and therefore be highly context-dependent. Because of this risk, I used various strategies designed to improve the reliability of the data. One re-

search methodology for fieldwork was to ask neutral questions rather than leading questions, so as to avoid priming a yes-response from the speaker due to revealing my own bias as the linguist asking the question. While this strategy may occasionally have the downside of making myself appear stupid³ and may lengthen the explanation, I followed it in order to improve the quality of the data. The idea is that the grammatical sentences are generally considered grammatical within the speech community. I present the examples with metadata of the speaker and the unique form ID in the database such that the interested reader is able to look up the form in the database and check the raw audio/video recording, if desired (for example, in order to check the pronunciation). While giving precise metadata information like this may suggest otherwise, the data selected for consideration are those that are representative of the data more broadly; any cases in which a form was unusual are indicated in the description in the main body of the text.

3.2.3 The use of natural speech data and corpora

Natural speech data refers to data from free-flowing speech. This allows for study of language as used by speakers and is particularly important for interactional speech and discourse-level phenomena such as referential encoding. Natural speech data can also be considered less corrupt than data obtained from elicitation or received only in transcribed form (Vaux et al. 2006:97-98; cf. Himmelmann 2012).

Natural speech texts can be divided into different genres. A basic distinction to make are between *monologic* and *dialogic* texts, the former referring to speech of a single speaker and the latter to speech of multiple speakers. Monologues are in turn dividable into stories, personal narratives, and instructional speech. Dialogues are dividable into semi-structured dialogues (based on prompts) and free dialogues (unprompted). I aimed for a mixture of these genres; 4 texts from different subgenres are provided in the Appendix as illustration of the natural speech dataset.

When natural speech data is compiled, it creates a corpus of data. This was done for my field recordings and was also done by Idelette Dugast in creation of her Dugast (1975) book, which forms a corpus of traditional folktales ('contes'), proverbs, and prophecies in Tunen. Such data allow for investigation of topics such as the frequency of particular words or constructions. When corpora are available from different time periods, they also allow for diachronic investigation. This is possible to a limited extent for Tunen given that Dugast's speakers were older men at time of recording, meaning that they were born over a century before my field

³By the end of the second fieldwork sessions, my ability to convincingly not know the answer to a question of interest was doubted by the consultants I worked with most closely, who sometimes joked that I already knew the answer.

recordings were made. Although this is a shallow time depth, the fact that this covers 3-4 generations allows for some consideration of diachronic change in Tunen. I discuss this in relation to the grammaticalisation of a specific indefinite determiner from the numeral ‘one’ in Kerr (2020) and in Chapter 4 section §4.3.11 and in terms of subject indexation in Chapter 8.

3.3 Primary data collection

3.3.1 Fieldwork stays

3.3.1.1 Overview

The primary source of data for this thesis was data collected via in-situ fieldwork in NdikiniMéki and Yaoundé, Cameroon. Fieldwork was conducted during two trips, one from March - June 2019 (3.5 months) and one from October 2021 - February 2022 (3.5 months), under research permits awarded by the Cameroonian Ministry for Scientific Research and Innovation (MINRESI; permit no. 90000061/MINRESI/B00/C00/C10/C12 and 000157/MINRESI/B00/C00/C10/C13 respectively). A small amount of data was also collected remotely in 2023. The delay of the second field trip and lack of further in-person field trips were due to the COVID-19 pandemic, which will be discussed further in section §3.3.4 below.

3.3.1.2 Consultants

I worked with approximately 11 consultants for linguistic study,⁴ of whom 8 were the principal consultants (i.e., consultants who were consulted in multiple sessions and whose data contributed directly to the topics explored in this thesis). Of these 8 consultants, 2 were women (AM, JO) and 6 were men (DM, EE, EB, EO, PB, PM), with an age range of 30-75 years old at time of recording (PB being born between 1980 and 1989; EE, EB, and JO 1970-1979; AM 1960-1969; PM 1950-1959; and EO and DM 1940-1949). All consultants spoke French as a second language, and some also spoke a neighbouring Bantu language as an L2. The consultants were educated to secondary education level. Many worked as agriculturalists, while JO worked as a secondary school teacher. Most consultants were literate in Tunen due to training from the Tunen literacy programme (*alphabétisation*) and consultant PM had some additional linguistic training through work on Cameroonian languages through the Cameroonian Association for Bible Translation and Literacy (CABTAL). All consultants spoke the Tɔ́bɔ́áŋɛ dialect (the standard dialect of NdikiniMéki), with the ex-

⁴I also worked with other consultants for recording of culturally-relevant practices such as songs.

ception of EO (who spoke Hiliŋ) and DM (who spoke Fombo), both of whom had lived in the Tɔ́bɔ́ŋɛ-speaking region for several decades.

3.3.1.3 Session format

Fieldwork sessions were conducted using French as a metalanguage. Translations are therefore provided in both the original French translation as agreed with the consultant and with an English translation later added by me. The French translations may be colloquial speech; they have been lightly edited for grammaticality but kept close to the original. The motivation for giving the French translation as agreed with consultants is to show the consultants' interpretation. For conversion of the examples used in this thesis for examples to be used for other purposes, e.g. language-learning resources, some further edits to the translations may therefore be desired for stylistic reasons.

The sessions were conducted in different environments. Many elicitation sessions were held at the house of consultant DM, either on the veranda outside (which had some background noise but had better light levels) or inside the house (where there was less background noise but lower light levels). Other sessions were held at consultants' houses, at an office space used by the CODELATU Bible translation team, and at the location of specific activities (e.g. in the town square when recording dancing at a funeral, and in JO's kitchen when recording cooking videos). The relevant information about location and presence of other people in a recording session is noted within a 'metadocumentation' file provided for each field session in the archival deposit.

Although this thesis centres on a research question related to syntax and IS, primary data collection was conducted following the principle that the data should be as high-quality as possible, in order to be of use to the broader scientific community and local community members. I therefore made both audio and video recordings. While video format is not key to the research question pursued here, the video format does allow for preservation of information conveyed by consultants via gesture, e.g. for iconic representation of tone height and deixis. The video files also provide an overview of the setting of each recording, and are key for capturing visual aspects of various cultural practices such as cooking and dancing.

During the sessions, I made handwritten notes in a notebook. In elicitation sessions, I made transcriptions by hand during the session, repeating words and sentences to check pronunciation and to ask the meaning. This means that the elicitation files have a relatively long duration for the number of utterances covered (versus an alternative workflow, where many utterances are elicited and transcriptions are only made at a later stage, using the recording). For natural speech data,

a recording was made and this was then played back on a laptop and transcribed with the help of a consultant in a future field session, due to the rapidity of natural speech preventing live transcriptions from being feasible. The natural speech transcriptions were generally made first in the same notebook and then typed up, and in some occasions (during the later fieldwork trips) were typed directly and/or written by the consultant before being checked together during the field session. The advantage of this method of transcribing together with a native-speaker consultant is that the transcriptions are of higher fidelity than transcriptions made without this help. These transcription sessions were also recorded — while not always strictly necessary, these recordings of transcription sessions nevertheless preserve interesting explanations of points of interest encountered in the natural speech data, as well as the elicitation of other relevant forms (e.g. the infinitival form of a verb first encountered in a natural speech context, or a judgement regarding whether a certain tone is obligatory, as will be seen to be key to the argumentation regarding biclausal clefts versus monoclausal focus constructions in Chapter 5).

3.3.2 Positionality

I conducted the research as an outsider to the Tunen-speaking community, being an L1 English speaker from the UK visiting as part of a research project based in the Netherlands. The project goals and methodology were discussed with consultants following procedure regarding informed consent for recordings. While many linguists working in Cameroon are associated with missionary organisations, I clarified my association with Leiden University and status independent of missionary linguists. I was invited to attend occasional meetings held by the Christian Tunen language committee CODELATU during my stay and found consultants through initial contact with CODELATU members, whose help I gratefully acknowledge.

Being an outsider linguist has methodological advantages and disadvantages (Bowerman 2015; Vaux et al. 2006). An advantage is the ability to ask many questions about the language without expectation that I knew the answers already, as I discussed in section §3.3 above. The clearest disadvantage is a lack of native-speaker knowledge about the language and culture. I therefore worked closely with consultants and in the second trip focused more time in training consultants, several of whom had already received basic linguistic training through the Tunen alphabétisation project of the Tunen language committee (CODELATU). Videos were made of cultural activities considered interesting, primarily on themes suggested by the consultants themselves. Several of these are currently untranscribed and therefore not covered in this thesis, but are included within the archival deposit in order to provide additional materials for further study.

3.3.3 Fieldwork questionnaires

As this project was conducted as part of a research project with research goals related to the interaction of syntax and IS in Bantu languages, and as the length of fieldwork was limited (a limitation exacerbated by travel restrictions related to the COVID-19 pandemic; section §3.3.4), fieldwork questionnaires were used to guide data collection and investigate the syntactic expression of IS. I explain each of these questionnaires below.

3.3.3.1 BaSIS project methodology

The research in this thesis was conducted as part of the Bantu Syntax and Information Structure (BaSIS) research project at Leiden University (NWO VIDI 276-78-001). As part of this project, a dedicated fieldwork methodology document was created (Van der Wal 2021) which follows the principle that investigation of IS requires contextualised data collection (Van der Wal 2016:259-260; see also Van der Wal et al. to appear). This methodology document was specifically targeted at the investigation of the interaction between syntax and IS in Bantu languages. It draws on previous sources including the Questionnaire on Information Structure (Skopeteas et al. 2006). Differences are that it is specifically tailored to Bantu languages (both in terms of linguistic topics covered and in terms of adaptation of example stimuli to local contexts) and that it includes a set of photos that replace the computationally-generated image stimuli from the QUIS. The latest version of this BaSIS methodology also includes tests for nominal licensing and the influence of Case.

A common task used for this thesis was to test for the felicity of a sentence in the context of a specific question (i.e., question-answer congruence), which was used to test for focus expression. More contrastive focus was tested using tasks relying on correcting someone's incorrect assumption (corrective focus) or contrasting something with something that was previously mentioned (contrastive focus). The questionnaire also covers topic expression; these are discussed in Chapter 5.

In my 2019 fieldwork, I used an early draft of the BaSIS project methodology that did not include the photo stimuli and lacked some of the tests from the 2021 version, notably the tests for nominal licensing. I therefore created some of my own hand-drawn picture stimuli for use in elicitation sessions and do not provide a full account of Case in Tunen (but see Chapter 6 and Chapter 8 for some initial hypotheses regarding argument licensing).

3.3.3.2 CHAOS/Co8 project questionnaire

In addition to the BaSIS project methodology, I collected data using the project questionnaire of the Consequences of Head Argument Order on Syntax (CHAOS) project at Potsdam University (subproject Co8 of SFB 1287 Limits of Variability in Language), where I was a short-term visiting PhD research fellow in September - December 2022. This project questionnaire was specifically developed to investigate syntactic correlates of OV versus VO word order. As Tunen is unusual for a Bantu language in having OV and not VO word order, this questionnaire allowed for more detailed investigation in the extent to which Tunen syntax resembles other OV or VO languages. Examples of tasks used for this thesis include investigation of word order in main versus embedded clauses and controlling for different predicate types and adverb types, as will be key to the investigation of Tunen's OV word order presented in Chapter 6.

During my 2021/22 fieldwork trip, I used a draft version of the methodology provided to me by Gisbert Fanselow in late 2021. On returning from Cameroon, I completed a 3-month visiting fellowship at Postdam University as part of the Co8 project. I then collected some additional data on Tunen using the updated version of the CHAOS/Co8 questionnaire via remote elicitation held over WhatsApp via audio calls (as discussed in section §3.3.4 below). These data are also transcribed and included within the corpus deposit. Data collected using the CHAOS/Co8 questionnaire are tagged as CHAOS within the Dative database and indicated by "CHAOSQuestionnaire" within the filename.

3.3.3.3 Other stimuli

Finally, a small number of examples were collected using the Max Planck Scope stimuli and by using materials I created myself specifically for my research purposes. The Max Planck Scope stimuli provide pictures that were used to disambiguate interpretations of sentences with multiple quantifiers.⁵ Data elicited using the Max Planck stimuli are tagged as *scope* within the Dative database. Within the comments field (part of the metadata), the picture ID number is given, e.g. 3/77 refers to picture number 3 of the set of 77 pictures.

My own materials include prompts for dialogue tasks and hand-drawn picture stimuli shown during sessions. These were presented to consultants with a prompt sentence in French or Tunen in order to record a (relatively) natural response.

⁵See https://www.eva.mpg.de/lingua/tools-at-lingboard/material_scope_fieldwork_project.php [accessed 02-2024].

3.3.3.4 Database coding

I aimed for the data collected in this study to be reusable for future research on the Tunen language. As different research questions require the researcher to be able to control for different variables, I set up the database to record key properties of each form. For example, each form is labelled for elicitation method, with the categories in Table 3.2 in the field `elicitation_method`.

<i>Label</i>	<i>Explanation</i>
elicitation_standard	Elicitation based on production or judgement of form in verbally-explained discourse context
elicitation_picture	Elicitation based on production or judgement of form in visually-prompted discourse context
elicitation_remote	Elicitation conducted remotely via WhatsApp call
monologue	Natural speech recording of a single participant
dialogue	Natural speech recording of multiple participants

Table 3.2: Encoding of elicitation type within Tunen Dative database.

This coding allows for example for the database to be searched only for natural speech examples (which can be achieved by a conjunction search for `monologue` OR `dialogue`). The data collection method can be checked for each example in this thesis by referencing the form from the database using the form UID indicated in square brackets by each example (e.g. [DM 166] refers to form id 166 in the database, by consultant DM). I also comment in-text on the nature of the data collection where relevant (for example, in mentioning constructions that occur frequently in dialogue but were not found in elicitation).

I also used the tag function in Dative to facilitate the analysis of the dataset, as I will discuss further in section §3.5.1 below. While I do include tags relevant to different parts of the grammar (e.g. ATR for examples of interest for the analysis of Tunen's ATR vowel harmony system), the current tagset prioritises variables related to syntax and IS, and I leave potentially interesting other variables (such as syllable structure) for future work.

3.3.4 Impact of COVID-19

This project was both directly and indirectly impacted by the COVID-19 pandemic. Firstly, data collection was directly impacted by the inability to travel in Spring 2020

as planned. The second field trip for this project was therefore significantly delayed, with travel to Cameroon not possible until October 2021. Because of this delay in the field trip, the second trip was also slightly shorter than planned and a third trip was not possible, meaning that some of the fieldwork stimuli were not covered and a smaller number of natural speech texts were transcribed than desired.

Secondly, the project was impacted indirectly, for example through restrictions on library access and through the delay in fieldwork leading to a more limited time period in which to conduct data analysis and write up the thesis. I acknowledge here the 3 months of additional financial support made available by the National Programma Onderwijs (NPO) COVID-19 funding of the Dutch government, which gave some compensation for these issues.

During the COVID-19 pandemic, a variety of linguists contributed to methodological advice on adapting linguistic fieldwork and language documentation workflows to the pandemic situation, specifically in how primary data collection could continue despite the international travel restrictions (see e.g. Sanker et al. 2021; Williams et al. 2021). For example, Williams et al. (2021) provide recommendations for remote fieldwork, using tools such as WhatsApp and Zoom for fieldwork sessions. Community-led approaches where community members make recordings themselves were also still possible during the pandemic.

Such community-driven initiatives rely on several factors: (i) prior contacts with the community, (ii) infrastructure such as phones, laptops, internet, and electricity, and (iii) technical knowledge. Unfortunately, the lack of infrastructure within the Tunen speech community and the fact that this thesis was not conducted within the context of a language documentation project meant that I could not conduct fieldwork remotely during the pandemic, as consultants did not have laptops or smartphones. When travel was again permitted, I was able to re-arrange my field budget so as to provide equipment to the Tunen speakers. I also re-scheduled my second field trip to have greater emphasis on training community members in transcriptions and technology use. This had the downside that it limited the amount of questionnaire material I could cover, but had the upside of contributing to a more sustainable working practice and allowing for remote communication since leaving Cameroon. By providing speakers with phones, headphones, memory cards, and audio recorders, I was able to keep in contact with the speakers after returning to the Netherlands in 2022. This thesis has been significantly improved by being able to conduct a small number of remote field sessions via WhatsApp checks for grammaticality judgements and WhatsApp audio calls for extra data collection.

In Sanker et al.'s (2021) study of the effect of remote recording set-up for phonetic measurements, three recommendations for best practice of remote recordings are provided: (i) record speakers locally when possible, (ii) keep recording

set-up consistent across recording sessions, and (iii) document the recording set-up (including recording devices, software programs, and settings) (Sanker et al. 2021:e379). Due to the absence of equipment and capacity for local recording management, recommendation 1 could not be followed. However, recommendation 2 and 3 were followed. All remote sessions were conducted via WhatsApp audio call with the speaker based in Cameroon. This audio call was played on speaker via a Samsung Galaxy A51 smartphone and recorded by a Zoom H5 audio recorder (the same recording device used for in-person sessions) at an equivalent recording distance. These decisions are documented in the database, and all remotely elicited material is tagged as `elicitation_remote` in the `elicitation_method` metadata field of the Dative database (Table 3.2). As non-locally recorded recordings are non-comparable with locally-recorded recordings even with the same speaker (Sanker et al. 2021), any researcher wishing to control for such small but statistically significant phonetic effects can filter the dataset by means of this tag (or, by identifying recordings made in 2023 as being remote).⁶ While the sound quality of the remote recordings is likely too poor to be of use for phonetic study, and no video was recorded for these remote sessions due to the poor quality of the internet connection, the recordings nevertheless help check the pronunciation and discussion of the semantics and pragmatics of each utterance, and were made and included in the archival deposit under the belief that a low-quality recording is preferable to no recording at all.

3.4 Secondary data

In addition to the collection of primary data through fieldwork with Tunen speakers, I was able to benefit from secondary data available on Tunen through previous sources (cf. Chapter 2 section §2.5.3.3).

3.4.1 Dugast (1975) text corpus

The most significant previous dataset available on Tunen is the Dugast (1975) text corpus. This book consists of a set of *contes* (folktales), proverbs, and prophecies. These are provided in Tunen (using Dugast's own transcription), with free translation in French, word-by-word glossing, and a small amount of commentary.

⁶Note that Sanker et al. (2021) did not investigate the effect of unstable internet connection on recording quality, which may further affect some measurements. This was an issue in some of the remote recordings due to unstable internet connection in Cameroon. I attempted to reduce the impact of this issue by asking the consultant to repeat their answers for clarification and to increase the chance of recording an uninterrupted utterance.

As the Dugast (1975) text corpus is not digitised, is presented in a difficult format for automatic digitisation via Optimal Character Recognition (OCR), and is under copyright until 70 years after Dugast's death, I generally limited my study of these texts to the folktales contained in Parts I and II (a total of approximately 60 texts). Parts I and II were chosen as an arbitrary sub-corpus of a great enough size to expect various forms of interest to appear while being small enough in size for manual searches to be feasible. Parts I and II were pre-selected before investigation in order to avoid cherry-picking, i.e., purposefully including or excluding a given text considered more useful for the research question at hand.

3.4.2 Other sources

There are a variety of other published works on Tunen available, as I discussed in Chapter 2 section §2.5. I am grateful to Maarten Mous for providing additional field-notes and unpublished material related to his work. I also received data from Ginger Boyd and Scott Isaac at SIL Cameroon, whom I thankfully acknowledge. In Ndikiniméki, Daniel Mbel shared a linguistic primer on Tunen with me and Pierre Molel shared resources developed by him together with members of the Cameroonian Association of Bible Translation and Literacy (CABTAL), including a translation of the New Testament (CABTAL 2019), a Tunen literacy primer, and a book of stories (see Chapter 2 section §2.5).

Where data from secondary sources are used, these are cited accordingly. I have lightly adapted the format of many of the examples in order to allow for easier comparison with the rest of the examples in this thesis, adapting some glosses for consistency. These adaptations are indicated by 'adapted' next to the citation, with optional extra explanation provided in a footnote.

3.5 Analysis

3.5.1 Dative/OLD

Field data were transcribed and annotated using the Dative database software, available at <https://www.dative.ca/>. Dative is a graphical user interface (GUI) to the Online Linguistic Database (OLD), an open-source piece of software. Dative was used instead of the more widely-used ELAN and FLEEx programmes in order for interoperability within the BaSIS project. As Dative was used by other members of the project, using it for Tunen meant that the datasets were more easily compared. An additional practical advantage was that the project members only needed to be trained in the use of one software. This facilitated comparative research within the

BaSIS project that constitutes work outside of this thesis (see e.g. Kerr et al. 2023 for a comparative study on Bantu word order and Kerr and van der Wal 2023 for a study on truth expression in Bantu languages).

The main benefit of using the Dative database for data organisation was the ability to search the entire corpus for the purposes of linguistic analysis, as mentioned in section §3.3.3. One benefit of Dative over other available software options is that it has a highly powerful and user-friendly search interface. The ability to search any level of the data (e.g. transcription, morpheme gloss line, translation, but also metadata fields) facilitated easy access of relevant examples. This was particularly useful for finding important data points, even if they had not originally been commented on as of interest. The tag field of Dative was particularly relevant here. This field allows the user to associate forms with a label, such as `question_polar` for polar questions. By developing a tagset related to the research questions of this thesis, I was able to code and find relevant examples. The ability to use conjunct search facilitated investigation of possibly interacting factors. For example, if looking into how negation is marked in relative clauses, a conjunct search can be made combining the tags `negation` and `relative`. Conjunct searches allowed also for cross-correlating tags with other query lines, e.g. searching for `focus_obj` in the tags line and `PRN` in the morpheme break line allows for investigation of the form of object pronouns in focal contexts.⁷ I provide the full tagset as part of the archival material of this thesis.

Finally, each form in Dative has a unique identifier (UID) which is permanent and unchangeable. Once a form was transcribed in Dative and the UID was generated, I noted this in my fieldnotes alongside the original transcription so that the original fieldnotes can be linked to the database. I also present all examples in this thesis with the UID of the corresponding form. By searching this number in the `id` metadata field of the Dative database, the example can be easily retrieved. For example, searching within Dative for ‘166’ within the `id` line provides form [DM 166] as a result. The recording information can be in turn extracted from the metadata associated with the form in Dative (see section §3.6); the form UID is used in this thesis as the minimal information needed to uniquely identify each form, with consultant initials provided to allow the reader to control for any potential variation conditional on dialect or other inter-speaker variables.

⁷Tags in Dative are provided at the level of the form, which corresponds in most cases in my corpus to a clause. Conjunct searches may return results that do not match the research question, for example in having focus on the object but a pronoun elsewhere in the clause. However, the search function still provides a short collection of forms that can easily be scanned manually for relevant examples.

3.5.2 Audacity and Praat

For acoustic analysis and manipulation of audio recordings, the programmes Audacity (Audacity Team 2023) and Praat (Boersma 2001) were used. While phonetics and phonology do not form the main focus of this thesis, and so no full acoustic analysis was conducted, some forms were checked in Praat in order to check the accuracy of transcriptions. The primary use of Praat was to investigate whether there was a prosodic break and to see the absolute realisation of the tones via the Praat pitch trace tool. I used Praat to check the absolute realisation of tonal downstep and the effect of utterance-final tone lowering rules discussed in Chapter 4 section §4.2, as well as checking the acoustic signal of minimal pairs elicited in the field-work (tagged `minimal_pair` within Dative), which all confirmed the previous studies regarding which phonemes are contrastive in Tunen (e.g. Mous 2003; Boyd 2015). Audacity was used to check the audio quality of recordings, to remove any unwanted content,⁸ and to play back recordings for transcription purposes with the Tunen consultants (as discussed in section §3.3.1).

All the field data collected are archived in audio/video format, meaning that in principle, all of the examples in this text can be traced to the original sound file, should further acoustic analysis be desired.

3.6 Data organisation and archiving

3.6.1 Filenaming and data organisation

All recordings made during this project were named according to filenaming convention agreed upon within the BaSIS project, as given in (30) below. As Tunen has the Guthrie code A44 (Chapter 2 section §2.5.2), the filename of all the recordings associated with this thesis therefore begins with *A44-*.

⁸ Removal of audio content was rarely necessary, but sometimes used, for example to remove an off-topic discussion about someone's family members that took place during one field session (removed in the interests of privacy). If something was removed from the middle of the recording in this way, this is noted in the metadocumentation text file for that field session.

(30) **BaSIS filenaming convention**

A44-20190313-EK-AM-SwadeshList-01_bkup.wav

A44	20190313	EK
Language UID (Guthrie code)	Date (YYYYMMDD)	Transcriber initials
AM	SwadeshList	01
Consultant initials	Brief description of subject matter	Part of session
_bkup	.wav	
(Optional extension to indicate version)		File extension

Within the archival deposit, the Tunen files are organised in the following way.

(31) Recordings > Consultant ID > Date

For example, the following file path leads to all files recorded on the 15th December 2021 with consultant PM.

(32) Recordings > PM > 20211215

Within each session folder, the audio and video files are provided in addition to a .rtf file named 'Metadocumentation' (e.g. A44-20211215-EK-PM-Metadocumentation.rtf). This metadocumentation file gives metadata information including session number date, time, location, consultant, recorder, and recording and file information, an overview of the contents of the recordings made in that session, and other notes relevant to the documentation. This information was then summarised by the general spreadsheet of metadocumentation for all the recordings, also available as part of the archive.

3.6.2 Archive

All the Dative (i.e. OLD) databases from this study are archived open access within The Language Archive (TLA) of the Max Planck Institute (MPI), as part of a joint collection from the BaSIS project members (see the TLA collection entitled 'Bantu Syntax and Information Structure (BaSIS)'; Table 3.4). These databases contain each form, with transcription line, line with morpheme breaks, glosses, free translation, and associated metadata (e.g. consultant initials, speaker comments, notes about discourse context, and so on).

In addition to the Dative transcriptions and metadata, I provide the raw audio and video files for the Tunen project within the Tunen sub-collection (Kerr in prep.). These recordings are made openly available under informed consent of the Tunen consultants, in order for other researchers and community members to be able to access them. In addition to the speech recordings analysed in this thesis, these recordings include some culturally-relevant materials, such as recordings of *eyganda* dancing (see Appendix Text 2) and some currently untranscribed Tunen songs (recorded at the UEBC church in Ndikiniméki and at Nefanté village).

The archival deposit of the Tunen database, audio/video recordings, and metadata is available from 4th September 2024 at The Language Archive as part of the BaSIS collection. The links to these collections are given in Table 3.4.

Collection:	Recordings and database on Tunen syntax and information structure
URL:	https://hdl.handle.net/1839/8fofia65-9a90-42e3-ae39-7bed93009b5a
Parent collection:	Bantu Syntax and Information Structure (BaSIS)
URL:	https://hdl.handle.net/1839/2acf92c5-e5db-445b-bcdb-3b13b7e58f3f

Table 3.4: Links to the Tunen archival TLA collection and parent BaSIS collection.

As discussed in section §3.6.1, The Tunen recordings are sorted by consultant and date. In order to reduce the time needed to locate a file, a document is provided that gives an overview of all the files in the Tunen collection. This document contains information such as the session number (1-76), the consultant(s), the speech type and subtype (e.g. *Natural speech*, *Dialogue*), the filenames associated with this session, and the respective Dative form UIDs. This document can therefore be used to look up the recording files and/or Dative database entries for a given example by using the consultant initials and Dative UID provided in square brackets alongside the examples in this thesis to check for the corresponding filename.

3.7 Limitations

The main limitations of the methodology used for this study is the size of the audiovisual corpus, with the Tunen Dative database containing c. 3300 utterances, the majority of which at the clause level but also containing individual lexical items. I

conducted 76 field sessions with 11 consultants, of whom 8 were the principal consultants, and have 11 fully-transcribed natural speech texts. As I discussed in section §3.3.4 above in the context of limitations caused by the COVID-19 pandemic, the current database contains a more limited set of natural speech examples than planned (although nevertheless contains multiple examples of texts of several genres). Other material that was collected is not yet transcribed and/or translated, and therefore was excluded for the analysis in this thesis, but can serve as useful material to explore in further work.

Another limitation of the methodology employed is that it was focussed primarily on what is and is not possible in Tunen. This means that the data are less well-suited to answering research questions related to frequency and preference. Such research questions are better answered with a larger corpus, especially with a bigger focus on naturalistic speech of different genres. I will come back to this question of frequency in Chapters 5-7 when looking at word order variation, and in Chapter 8 when drawing conclusions about the extent of influence of information structure on Tunen syntax.

3.8 Summary

This chapter has described the methods employed in collecting and analysing the data presented in this thesis, considering the value and limitations of elicitation and natural speech data for answering this thesis' main research questions related to the interaction of syntax and information structure in Tunen, as contextualised within the framework of generative linguistics. All data collected in this project are available open access in *The Language Archive*, together with the Dative database and metadocumentation, as available in Kerr (in prep.).