

18

The emergence of grammatical structure from inter-predictability

John Mansfield and Charles Kemp

Introduction

How did grammatical structures evolve to be the way they are? Presumably, they have developed slowly over thousands of years of human language learning and use, but what forces in particular have shaped outcomes such as a noun phrase schema, Dem-Num-Adj-N? More importantly, we might ask how human language has arrived at its global distribution such that Dem-Num-Adj-N is common to many distinct lineages, Dem-Num-N-Adj somewhat less common, and N-Num-Dem-Adj occurs hardly at all (Dryer 2018).

There are two main approaches to explaining these phenomena: theories that involve a fixed underlying syntactic structure, around which specific languages ebb and flow (e.g. Cinque and Rizzi 2009; Chomsky and Lasnik 2015 [1993]), versus theories where there is no fixed underlying structure, only flow (e.g. Hopper 1987; Bybee 2001; Hopper and Traugott 2003). This chapter belongs firmly in the second camp. Although we cannot demonstrate the actual evolutionary development of Dem-Num-Adj-N, in English or any other language, we propose some fundamental principles that we believe to have played a key role in the evolution of such structures. We believe that these principles are active in human cognition when

individuals acquire and process language, but modelling actual acquisition and processing is beyond our reach in this study. Instead, we demonstrate the operation of our principles in a relatively simple computational implementation. We firstly articulate the core principles and secondly present a preliminary implementation using artificial language data.

This chapter is a tribute to Professor Jane Simpson, who may or may not agree with any of its conclusions. One of Professor Simpson’s major contributions to grammatical theory is her study of ‘templatic’ structure in morphology (Simpson and Withgott 1986), which we here integrate into a more general proposal about linear grammatical schemas, encompassing both morphology and syntax. Simpson and Withgott pointed out a range of idiosyncratic morphological constructions that defy the main tenets of phrase-structure grammar and should therefore be treated as an alternative type of grammatical computation. In this chapter, we do not dwell on the question of morpheme-specific idiosyncrasies, but instead take the general idea of linear schemas as a point of departure for a model of how grammatical structures can be explained without any recourse to a fixed hierarchy of grammatical categories. In fact, while templates or linear schemas are often regarded as ‘arbitrary’ or ‘stipulative’ analyses that fall outside of principled linguistic theory, we will argue that, on the contrary, linear schemas can be explained based on very sound principles of linguistic cognition.

Recent research has shown striking correspondences between linear proximity of words or morphemes and their statistical inter-predictability, measured as mutual information (Futrell 2019). For example, a grammatical category order such as Dem-Num-Adj-N, familiar from languages such as Ket (1), corresponds with corpus data showing that adjectives are the most inter-predictable with nouns, numbers less so, and demonstratives least of all (Culbertson, Schouwstra and Kirby 2020).

- (1) Ket
- | | | | |
|--------------|------------|-------------|-------------|
| <i>kin'e</i> | <i>qo·</i> | <i>aqta</i> | <i>dεʔŋ</i> |
| DEM | ten | good | people |
- ‘these ten good people’ (Rijkhoff 2004, 126)

This chapter proposes an explanation for why linear proximity and inter-predictability should be correlated, with holistic retrieval of complex symbols being the key connection. We test this proposal by implementing a computational model and showing that it generates similar grammatical sequences to those found in natural languages, starting from input data

that have no intrinsic ordering. The model is not a realistic simulation of any human language, since there is no situation in which a language user must devise linear sequencing naively. Our model begins with completely random ordering, while users of natural language are exposed to particular sequencing patterns, which they reproduce with occasional modifications. Rather than modelling actual language use or acquisition, our more modest goal in this chapter is to show that inter-predictability can, *in principle*, account for grammatical linear ordering, and to propose an explanation for this in efficient chunking.

Since our model does not reflect many properties of actual human language, rather than calling it a 'language' or a 'grammar', we refer to it more modestly as an 'encoder'. The encoder does not require any intrinsic ordering principles that reference specific grammatical categories, but instead achieves natural-language-like sequences based on independently evidenced principles of linguistic cognition. We propose three fundamental principles that drive linear ordering (these will be explained in more detail below):

- i. *Chunking*: To reduce the number of symbols processed, any group of symbols that is highly inter-predictable can be stored and retrieved as a single complex symbol;
- ii. *Adjacency*: When a message is output as a linear sequence, complex symbols are preferably linearised with their sub-parts adjacent to one another;
- iii. *Structural priming*: When a message is output as a linear sequence, associative memory residues from previous messages make the grammatical category sequence of the new message more likely to match those of previous messages.

If a grammatical structure can be shown to emerge from these cognitive principles, then there is no need to propose an extrinsic explanation involving a fixed underlying syntax of grammatical categories. Our claim is not that *all* grammatical structure can be explained by principles I–III – for example, it seems likely that basic constituent order (SOV, etc.) requires other principles. But we expect that many linear structures are at least partially explained by these principles. In this study, we support this claim by applying the encoder to NP-internal categories {N, Adj, Num, Dem}. The encoder succeeds in generating NP-internal orderings that bear a striking resemblance to natural languages; however, there is nothing about the system that is specific to these grammatical categories, and therefore it should be applicable to a range of other linear structures.

Before describing the encoder, we briefly reflect on how Professor Simpson's work provokes consideration of linear versus hierarchical structure and summarise recent research that demonstrates the widespread correspondence of inter-predictability with linear ordering.

Linear versus hierarchical, stipulative versus principled

Simpson and Withgott argue that some grammatical structures are better analysed as 'linear arrays of morphemes', rather than hierarchical structure (Simpson and Withgott 1986, 156; see also Perlmutter 1971). They use the term 'templates' for the schematic generalisations that can be made over such morpheme arrays. A template is fully schematic if it refers only to grammatical categories, e.g. TAM-Subj-Obj, where all morphemes belonging to these categories are linearised in that way. But templates may also be 'partially schematic', where some positions are filled by grammatical categories and others by specific morphemes, or by heterogeneous sets of morphemes.

Partially schematic templates allow for idiosyncratic dependencies between morphemes, in a way that does not reflect classic hierarchical phrase-structure representations (see also Rumsey, this volume). Simpson and Withgott present a range of examples from Warumungu, Warlpiri and French. For example, Warumungu has pronominal clitic clusters with idiosyncratic ordering depending on exactly which subject/object combinations are involved (Simpson and Withgott 1986, 161).¹ Subsequent research has demonstrated many other examples of idiosyncratic sequencing in morphology (e.g. Stump 1997; Nordlinger 2010b), including paradigmatic inconsistencies that can, for example, place 1SG.S in one linear position, but 2SG.S in a quite different position (Crysmann and Bonami 2016; Mansfield et al. 2022).

By contrast, hierarchical phrase-structure grammar is a fully schematic model of grammatical category relations: the rules generalise across all members of grammatical categories, as opposed to, say, having lexically specific phrasal rules. There is a fixed underlying structure of grammatical categories, and

1 This could be taken to imply that *=ngki* is an 'inverse' marker, but in fact it only appears in this specific clitic cluster, which undermines any analysis as a 'recurrent partial' with an independent meaning.

in some versions this hierarchy is fixed for all human languages (Chomsky and Lasnik 2015 [1993]). Since the clitic clusters in Warumungu instead involve idiosyncratic dependencies between morphemes, these are claimed to belong to another type of grammar – the template. The structure of templates is considered to be language-specific and stipulative – i.e. it is free from the constraints of general grammatical principles (Good 2016). By association, linear sequences are often seen as stipulative and language-specific, in contrast to hierarchical structure, which is principled and universal. Indeed, some influential universalist models take hierarchy to its logical limit by proposing that all structure is strictly binary-branching (Kayne 1981, 1994). The slogan ‘structures not strings’ captures the view that hierarchical grouping is the proper subject of syntactic theory, while linear sequences are superficial and/or unimportant (Everaert et al. 2015; Kulmizev and Nivre 2022 *inter alia*).

In contrast to the widespread view that linear schemas are stipulative and unprincipled, in this chapter, we will argue that linear grammatical schemas are based on very robust, general principles. Linear schemas found in particular languages are not failures of linguistic theory, but rather can be incorporated into a stochastic model of symbolic encoding and linearisation that predicts constrained variation in grammatical structure. The principles that we will evoke involve efficient processing and associative memory, both of which are extensively evidenced in human language use. However, we do not claim that *all* grammatical structures are linear. Hierarchical structure is clearly motivated for structures with unbounded recursion (Chomsky 1957), and perhaps even more essentially, those parts of grammar that involve recurrent expandable nodes, such as the multiple NP positions in clause structure (Wells 1947). Rather, our claim is that hierarchical grouping should be proposed only where it is clearly motivated, while other structures can be satisfactorily modelled as linear arrays.² This is similar to the position argued by Culicover and Jackendoff (2005), who note that hierarchical structure adds complexity to syntactic representations and to derivational processes, and should therefore be posited only as necessary.

2 More specifically, if we assume that grammar has both linear schemas and hierarchical nesting, this implies a model of linear schemas in which some positions are themselves filled by schemas. This is not mentioned in Simpson and Withgott’s proposal, but is later suggested by Good, who calls these ‘elastic’ positions (Good 2016, 57).

There is an interesting parallel between our proposal and theories of hierarchical syntax, since our model does in fact involve recursive chunking of linguistic symbols (see details below). But the important difference is that our hierarchical structures are a format for storing and retrieving specific linguistic symbols, as opposed to phrase-structure grammar, which involves a hierarchy of grammatical categories. We return to this difference in the discussion section below.

Inter-predictability and linear proximity

Several recent studies have revealed a striking relationship between linear proximity and statistical inter-predictability, evidenced in both syntax and morphology. Futrell, Hahn and colleagues have identified various correlations, both with respect to specific constructions such as adjective stacking (e.g. *beautiful green shirt* versus *?green beautiful shirt*) (Hahn et al. 2018), and more widely with respect to the proximity of words in phrase structure, and the proximity of affixes to the stem in word structure (Futrell 2019; Hahn, Degen and Futrell 2021; Hahn, Mathew and Degen 2022). Pairs of morphemes (either words or affixes) are *inter-predictable* when they tend to co-occur in the same phrases, rather than independently. For example, *man* + *big* tend to co-occur, as do *house* + *large*, showing that the occurrence of these nouns and adjectives is not statistically independent (Biber, Conrad and Reppen 1998, 46). The crucial observation is that inter-predictable morphemes tend to occur close together in linear sequences, a principle that Futrell labels *information locality* (Futrell 2019). Futrell, Hahn and colleagues propose an explanation for information locality based on the idea that predictability helps us process the incoming linguistic signal (Hale 2001; Levy 2008), but our memory of contextual predictors decays over time. Therefore, the benefit of predictive processing is maximised when highly inter-predictable symbols have minimal gaps between them (Futrell 2019; Hahn, Degen and Futrell 2021). They call this the ‘memory–surprisal tradeoff’. Other relevant strands of theory include dependency locality (e.g. Hawkins 2004; Futrell, Levy and Gibson 2020; Jing, Blasi and Bickel 2022), and uniform information density (Jaeger 2010; Jaeger and Buz 2017), which may or may not align with information locality.

In this study, we propose a different explanation for information locality, where storage and retrieval of complex symbols is the cognitive mechanism responsible. Let us call this new proposal *informational chunking*. The key

insight of informational chunking is anticipated in some earlier studies of English morphology (Hay 2002; Plag and Baayen 2009), arguing that stem-affix adjacency is explained by certain combinations being more likely to be holistically processed. This holistic processing is, in turn, related to statistical association measures revealed in corpus data. The same principle has also been implemented in one previous computational model, the Chunk Based Learner (McCauley and Christiansen 2019), which simulates children's syntactic production in such a way that predictable word combinations are output as holistic chunks. We build on these earlier studies by connecting them back to information theory, in a model that is intended to generalise over morphology and syntax. Alongside the memory–surprisal tradeoff, dependency locality and uniform information density, our proposal follows a blossoming of research on the question of how functional constraints and cognitive biases shape grammatical structure. Comparing the merits and predictions of these theories would be of great interest, but we must leave such a comparison to future work.

One of the recent studies demonstrating information locality is Culbertson and colleagues' (2020) cross-linguistic investigation of {N, Adj, Num, Dem} sequencing in noun phrases (Culbertson, Schouwstra and Kirby 2020). This is interesting because there is disagreement in the literature about whether noun phrases should be treated as linear or hierarchical structures (e.g. Culicover and Jackendoff 2005; Alexiadou, Haegeman and Stavrou 2008; see further references below). While the elements {N, Adj, Num, Dem} show a diverse range of sequences, most languages conform to a general constraint, whereby linear proximity to the noun is ranked as $\text{Adj} \geq \text{Num} \geq \text{Dem}$ (the adjective must be at least as close as the number, and the number must be at least as close as the demonstrative). Of the 24 possible permutations of the categories {N, Adj, Num, Dem}, there are eight orderings that satisfy this ranking constraint, and these account for 83 per cent of languages in a typological sample of 576 languages (Dryer 2018, 799). Culbertson and colleagues calculate inter-predictability statistics from syntactically annotated corpora and find that the ranking stated above is reflected in the degree to which each of these modifier classes is inter-predictable with nouns, and this is true for all 24 languages tested (Culbertson, Schouwstra and Kirby 2020, 697). This provides a striking example of information locality, i.e. a correlation between linear proximity and inter-predictability.

Information locality in noun phrases connects with an earlier suggestion by Dryer (2009, 197), who proposes that word classes are closer to the noun if they ‘denote more inherent properties of the referent’. Adjectival properties such as colour and size tend to be permanent properties of referents, and thus co-occur more predictably with noun labels, whereas numerosities and deictic markers are ephemeral (most entities can freely appear in different numerosities, or different relations to the speaker/addressee). Thus, our experience of the world and the semantic basis of grammatical categories present a plausible account of why some word classes have greater statistical inter-predictability than others (Culbertson, Schouwstra and Kirby 2020, 702). Nonetheless, in our view, it is the statistical aspect of this that is the proximal cause of grammatical linearisation, and thus the main focus of this study.

Culbertson and colleagues interpret NP linear schemas as evidence of an underlying nested conceptual representation [Dem [Num [Adj [N] Adj] Num] Dem] (see also Rijkhoff 2004, 218). The nested conceptual representation generates isomorphic word orders, explaining the strong bias towards certain NP-internal sequences across languages (Culbertson, Schouwstra and Kirby 2020, 709). This explanation is compatible with hierarchical syntactic accounts in which the linear ordering of NP-internal categories is modelled as an underlying structure along the lines of Figure 18.1 (e.g. Cinque 2005; Alexiadou, Haegeman and Stavrou 2008).

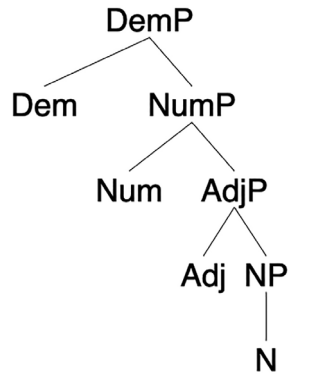


Figure 18.1. Hierarchical syntax tree for the elements {N, Adj, Num, Dem}.

Source: Authors.

In this chapter, we take Culbertson and colleagues’ NP findings as the basis for a model of how linear sequences of grammatical categories can be explained by inter-predictability. But rather than treating word order as mediated by a fixed underlying hierarchical structure, we instead propose a mechanism that more directly generates word ordering from inter-predictability. A position closer to ours is held by Dryer (2009, 2018), who treats the NP as a linear structure, and discusses a range of functional motivations for

linearisation, including the abovementioned ‘inherent properties’ theory. But our proposal does not even strictly require a linear schema (though we don’t rule it out as an efficient abstraction); instead, our linearisation procedure works directly by priming from previous exemplars (cf. Ambridge 2020a, 2020b). In the implementation described below, messages are simply linear sequences of symbols, and consistent grammatical category sequences are an emergent property of inter-predictability of symbols, together with associative memory.

Generating {N, Adj, Num, Dem} ordering from cognitive and communicative principles

Our encoder begins with no inherent ordering of grammatical categories and develops schematic ordering based on independently evidenced cognitive and communicative principles. Repeated from above, these are:

- i. *Chunking*: To reduce the number of symbols processed, any group of symbols that is highly inter-predictable can be stored and retrieved as a single complex symbol;
- ii. *Adjacency*: When a message is output as a linear sequence, complex symbols are preferably linearised with their sub-parts adjacent to one another;
- iii. *Structural priming*: When a message is output as a linear sequence, associative memory residues from previous messages make the grammatical category sequence of the new message more likely to match those of previous messages.

Each of these is a sound working assumption, supported by extensive previous research. *Chunking* follows from information-theoretic research on effort reduction (e.g. Shannon 1948; Zipf 1949; Aylett and Turk 2006; Grünwald 2007; Levshina 2022), though as we explain below, our application refers to effort reduction in the retrieval of symbols, rather than their external articulation (Seržant and Moroz 2022). *Adjacency* is apparently so obvious that it is simply assumed in psycholinguistic research on holistically memorised complex symbols. In most or all studies, the stored complex symbols are output as contiguous units, and this goes equally for memorisation of morphologically complex words (e.g. Baayen and Schreuder 2003; Marslen-Wilson 2007; Kuperman, Bertram and Baayen 2010),

or multi-word phrases (e.g. Arnon and Snider 2010; Pijpops, De Smet and Van de Velde 2018; Contreras Kallens and Christiansen 2022). Adjacency is also central, though somewhat more flexibly operationalised, in the large body of research on corpus collocations (e.g. Sinclair 1991). *Structural priming* and persistence are extensively evidenced in language processing (e.g. Pickering and Branigan 1998; Szmrecsanyi 2006; Van Gompel and Arai 2018). With respect to the ordering of grammatical categories, priming effects are found when speakers select among variable word orders for a sentence (Hartsuiker, Kolk and Huiskamp 1999; Fukumura and Zhang 2023) or variable morphological orders for a word (Mansfield, Stoll and Bickel 2020; Mansfield et al. 2022). Consistent linear positioning of elements from the same grammatical category is tacitly assumed in most theories of grammar, but is stated explicitly by Dik (1997) as the *Principle of Functional Stability*.

We will show that a model implementing principles I–III produces similar distributions of {N, Adj, Num, Dem} sequences to those observed in natural languages. We will also see how the relative weighting of these principles affects the emergence of grammatical ordering and the distinctive roles played by each principle. This suggests that at least some of the linear grammatical schemas observed in natural languages need not be seen as ‘stipulative’, but on the contrary can be motivated by fundamental, well-evidenced principles. In the course of explaining the model and illustrating its output, potential interpretations in terms of human language cognition will become apparent. The model is implemented as a single Python script of modest complexity (< 500 lines of code), which is available at the code repository for this study.³

Methodological preliminaries

The model described here is a simple ‘message encoder’, a proof-of-concept using artificial language data. A follow-up study will present more complex implementations using natural corpus data. Some features here are deliberately simplified, allowing us to focus on the main question of how the inter-predictability of symbols might give rise to ordering.

3 John Mansfield and Charles Kemp, ‘Explaining Grammatical Schemas with Inter-predictability’, Zenodo, 6 February 2023, zenodo.org/record/7601752.

The input data used for the encoder is an artificially generated dataset of co-occurring atomic concepts, assigned to categories {N, Adj, Num, Dem}. These form a set of ‘messages’, where a message contains one concept from each category, in unordered sets such as {*dogs, brown, three, these*} and {*cats, black, four, those*}. These sets of concepts are the semantic content of the messages. The encoder builds a linear string of symbols to encode the content of each message.

Our message set is generated using a stochastic process, designed so that the inter-predictability by grammatical category is similar to that found in Culbertson, Schouwstra and Kirby (2020), namely: $N;Adj \geq N;Num \geq N;Dem$. Following Culbertson and colleagues, we formulate inter-predictability as mutual information (MI), which can be calculated for each specific combination of symbols, and averaged across each modifier category Adj, Num, Dem. We generated one set of approximately 5,000 training messages, which is used by the encoder to estimate the pointwise mutual information (PMI) of each symbol combination, and a second list of 500 test messages, which the encoder must linearise. (The exact numbers of messages are variable due to the stochastic nature of the generation process.) The data has more possibilities for nouns, adjectives and numbers, since these are large open classes, and more constrained possibilities for demonstratives, since this is a small closed class. There are eight nominal concepts: *dogs, cats, rabbits, birds, chairs, tables, cars, wheels*; eight adjectival concepts: *brown, black, white, grey, big, small, old, new*; eight numerical concepts: *two–nine*; and two deictic concepts: *these, those*. For technical details of how the datasets are created, see the script `generate-MI-distros.py` at the code repository.⁴

The most important property of the input data is that it reflects the kind of inter-predictability relations found in NPs in natural language corpora. Over 10 runs of the message generator, the mean MI figures and standard deviations for each modifier category are: $PMI(N;Adj) \mu = 0.77, \sigma = 0.13$; $PMI(N;Num) \mu = 0.33, \sigma = 0.06$; $PMI(N;Dem) \mu < 0.01, \sigma < 0.01$. The absolute values of these figures are smaller than those reported for natural language corpora by Culbertson, Schouwstra and Kirby (2020, 705), but this is simply an effect of the smaller sets of distinct symbols. The important thing for our model is the ranking of PMI by grammatical category.

⁴ Mansfield and Kemp, zenodo.org/record/7601752.

Aside from inter-predictability ranking, the messages we generate are very unlike natural language, because they always consist of the categories {N, Adj, Num, Dem}. Obviously, humans usually use only a subset of these conceptual categories to form a referential expression, and the consistent use of all four types was implemented purely to simplify the model. A follow-up study will extend the model to allow the more flexible message content of natural corpora.

The efficiency trade-off in Informational Chunking

Before commencing the proper ‘encoder’ phase, there is a preparatory phase in which all atomic concepts are stored in memory as linguistic symbols, such that for each concept, e.g. *cat*, there is an associated symbol *cat*. More importantly, this memorisation process represents the statistical inter-predictabilities between all symbols, as calculated from the ~5,000 messages in the training data. Once all symbols and their inter-predictabilities are stored in memory, the encoder proper works through the ~500 messages in the test data, using two distinct steps to encode each as a linear string: first, symbols required for expressing the concepts are retrieved from memory, either as simple or complex symbols; second, the resulting set of simple and complex symbols is output in a linear sequence. In this section, we describe the preparatory phase of symbolic memorisation and the symbolic retrieval process, which we call ‘chunking’.

As mentioned above, the symbolic memorisation phase reads all the messages in the training data, and uses these to calculate inter-predictability for all symbols. This is formulated as pointwise mutual information (PMI) for each specific symbol combination, e.g. $\text{PMI}(\text{cats}; \text{black}) = 1.26$, or $\text{PMI}(\text{dogs}; \text{three}) = 0.31$. The use of PMI has a long tradition as an association measure in corpus linguistics (e.g. Church and Hanks 1990). PMI is typically measured between pairs of variables, but it can also be applied recursively to calculate the inter-predictability between a binary combination and a third simple variable, e.g. $\text{PMI}((\text{cats}, \text{black}); \text{three})$, or similarly between a ternary combination and a simple fourth variable.⁵

Once all symbols have been memorised and their inter-predictabilities calculated, the encoder works through the ~500 messages in the test data, encoding each into a linear message sequence. As outlined above, the

5 We might expect these multivariate PMIs to be less important than binary PMIs, but the encoder nonetheless calculates ternary and quaternary PMIs, since human language does show some effects of holistic storage for larger complex symbols (e.g. Arnon and Snider 2010).

semantic content of each message includes nominal, adjectival, numeric and deictic concepts, for example *{cats, black, three, these}*. The encoder retrieves the associated symbols from memory, either as simple or complex symbols. For example, a set of symbols might be retrieved as *{(cats, black), three, these}*, i.e. with one binary chunked symbol and two simple symbols. We represent chunks with round brackets, and the complete set of retrieved symbols with curly brackets. Retrieving a chunked complex symbol reduces the number of symbols that need to be retrieved, and this can be interpreted as ‘effort reduction’. One might think of this like retrieving clothes from a cupboard before getting dressed: if a particular belt tends to go with a particular pair of jeans, it may be more efficient to store and retrieve these already conjoined; a pair of socks that always go together should be stored and retrieved as a single bunched-up unit.

The storage of complex chunks does not prevent simple symbols from being accessed individually: for example, *(cats, black)* does not delete *cats* or *black* from memory. We also assume that the semantics and grammatical category membership for parts of complex symbols remains identical with simple symbols, setting aside the fact that in natural language, complex symbols may gradually become semantically or grammatically dissociated from their parts (Brinton and Traugott 2005). The storage of complex symbols adds to the complexity and size of our lexical storage, as a trade-off against the benefit of more efficient retrieval (Wray 2002, 2017). This type of efficiency reflects the concept of chunking used in research on linguistic production planning (Blumenthal-Dramé et al. 2017), and theories of the mental lexicon (e.g. ten Hacken 2019).

The connection we draw between efficient retrieval and mutual information is different from previous information-theoretic work on language processing (Gibson et al. 2019). Most information-theoretic work focuses on the number of symbols used in message *transmission* (e.g. articulatory reduction or word length), while we apply the same mathematical principles to efficient retrieval of symbols from memory (cf. Seržant and Moroz 2022), which we call *informational chunking*. Since natural language exhibits pervasive variability, we model informational chunking as a variable process, using a random error term drawn from a Gaussian distribution. The PMI of a symbol combination, plus a random error term, is tested against a chunking threshold. For example, the chunk threshold can be set at 1 bit, and the encoder might retrieve the concepts *{cat, black}* with $\text{PMI}(\text{cats}; \text{black}) = 1.26$, and a random error $\epsilon = -0.04$, which would yield a result of 1.22 and thus pass the threshold for a complex chunk.

The encoder begins by testing binary combinations, potentially returning binary chunks like (*cats*, *black*); but the process is recursive, also testing for larger chunks like ((*cats*, *black*), *three*) or ((*black*, *three*), *cats*). At the maximum, all four symbols can be chunked into a single complex symbol, such as (((*cats*, *black*), *three*), *these*); however, PMI naturally tends to reduce as symbols grow more complex.

If it is easier to retrieve fewer symbols, we might wonder why not chunk the semantic content of every message as a single complex symbol? Intuitively, the problem here would be the exponential proliferation of complex symbols that need to be accurately retrieved. This assumes that there is some upper limit on the range of distinct symbols that can be stored and retrieved (Kirby et al. 2015). Constrained by the symbol proliferation problem, the encoder should perform chunking selectively to gain savings on the symbol count of individual messages, while also limiting the range of complex symbols it retrieves. The degree of proliferation could be formulated simply as the number of distinct (simple or complex) symbols retrieved by the encoder over its entire repertoire of messages. However, a more nuanced formulation is the *entropy* of symbol retrieval (Shannon 1948), which reflects both the number of distinct symbols and whether some particular symbols are more frequently retrieved (less entropy) or whether all symbols have close to equally frequent retrieval (more entropy). In information-theoretic terms, we can think of symbol retrieval as having a *channel capacity*, which limits the per-symbol entropy at which retrieval can accurately occur (Cover and Thomas 2002). The encoder thus has conflicting constraints to reduce the number of symbols retrieved per message, while limiting the retrieval entropy per symbol. Formulating the problem in these terms is similar to the idea of two-part codes from the Minimum Description Length literature, which balances the length of a set of encoded messages against the length of the ‘code book’ used for coding these messages (Grünwald 2007).

Figure 18.2 illustrates the superior efficiency of Informational Chunking, in comparison to purely random chunking based only on the Gaussian error term. Both versions were run using a range of chunking thresholds, reflected in their spread of average symbol counts (x-axis). The Informational method demonstrates superior efficiency, as represented by its outputs being closer to the bottom-left quadrant of the graph. This indicates that it achieves fewer symbols per message, at a lower encoder entropy (y-axis). Informational chunking uses a smaller set of complex symbols, such as (*cats*, *black*), which are frequently repeated. By contrast, random chunking uses a large range of complex symbols, many of which appear only once. The latter scenario results in higher encoder entropy.

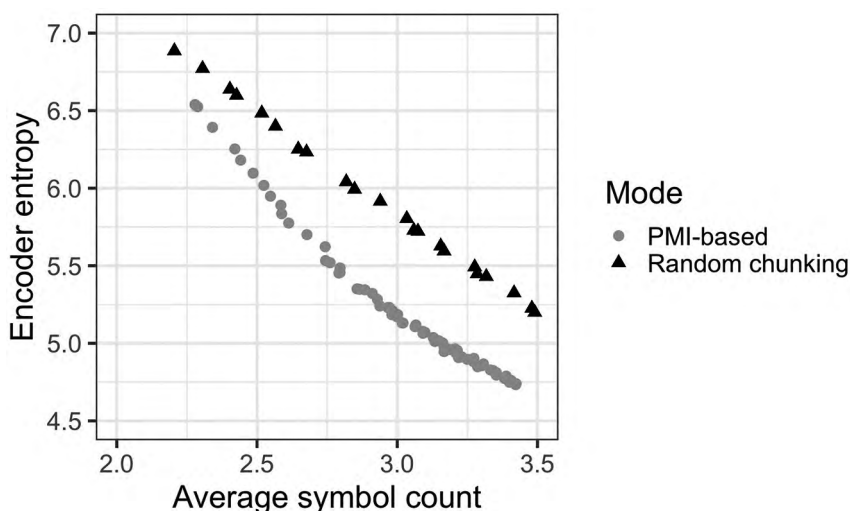


Figure 18.2: Trade-off between symbol count per message and encoder entropy.

Note: Illustrating results from Informational Chunking versus an encoder where chunking is purely random.

Linearisation via contiguity and consistency parameters

The second step of string encoding involves linearisation of the retrieved symbols. As mentioned above, each set of symbols begins with no intrinsic ordering, but consistent grammatical category sequences emerge from the application of our two remaining fundamental principles: (II) complex symbols are preferentially transmitted in a contiguous fashion; and (III) previously used grammatical category sequences are preferentially reused. Each of these principles is implemented as a parameter that can be applied at various weights. Our model, therefore has two ‘linearisation’ parameters, *contiguity* and *consistency*, in addition to the chunking threshold parameter described above. Like chunking, linearisation is a fundamentally variable process (see Wilmoth et al., this volume). In the encoder’s initial state, all the 24 possible permutations of the categories {N, Adj, Num, Dem} are equally probable as transmission sequences; but on each message transmission, these probabilities are modified by the two linearisation parameters.

The contiguity parameter applies if the message includes a complex symbol (a chunk) of any size, boosting the probability of sequences in which the parts of the complex symbol are contiguous. For example, if a message

contains the binary complex symbol (*cats, black*), then a boost is applied to the probability of schemas such as *cats-black-three-these*, *black-cats-three-these*, *three-cats-black-these* ... etc. If a message contains a ternary complex symbol, such as ((*cats, black*), *three*), the boost requires contiguity of both the inner binary complex, as well as contiguity of this binary with the third term. This would therefore apply a boost to sequences such as *cats-black-three-these*, *three-cats-black-these*, but would not apply to *cats-three-black-these*, where all elements of the complex symbol are contiguous, but the inner binary element (*cats, black*) is discontinuous.

The contiguity parameter is a probabilistic bias, rather than a hard constraint, for two reasons. First, varying the weight of the parameter helps us to understand its effect on sequencing outcomes. Second, while we assume that complex symbol contiguity is generally preferred, natural languages do sometimes allow discontinuous linearisation of signs that are presumably stored and retrieved as holistic symbols. One clear example is English expletive infixation, as in *fan-fucking-tastic* (McCarthy 1982); other examples are West Germanic lexicalised particle verbs, e.g. English *I picked the children up*, or German *Er schlägt das Wort im Wörterbuch nach*, 'He looks up (lit. hits to) the word in the dictionary'.⁶ It seems that in human languages, complex symbol contiguity is strongly preferred, but is still a soft constraint.

The contiguity parameter is applied as a multiplier effect on the probabilities of possible linear outputs. As mentioned above, the model begins with all possible sequences as equiprobable, which is implemented as a pool of candidate grammatical sequences that initially contains one token of every possible sequence, i.e. all 24 permutations Dem-Adj-Num-N, Dem-Num-Adj-N, Dem-N-Adj-Num ... Linearisation of the current message works by making a random selection from the pool, and using the grammatical category sequence to linearise the symbols, e.g. Dem-N-Adj-Num → *these-cats-black-three*.

As we will see below, the consistency parameter will gradually change the contents of the category ordering pool. But the contiguity parameter acts as an ephemeral multiplier. For example, in the pool's initial state, all grammatical sequences have probability $1/24 = 0.04$. If the symbols to be linearised are {(*cats, black*), *three, these*}, and the contiguity multiplier is 5,

6 We are grateful to a reviewer for pointing out the relevance of particle verbs, and to Christian Döhler for the German example.

the sequence will be drawn from a temporarily modified version of the pool, with 5 tokens of each sequence that satisfies {N, Adj} adjacency (e.g. Adj-N-Dem-Num, Dem-Adj-N-Num, etc.). There are 12 such sequences, and this gives a probability of $5/72 = 0.07$ to each of the preferred sequences, and $1/72 = 0.01$ to each of the dispreferred sequences. This contiguity boosting applies only to the current message, with its symbol-specific chunking properties, rather than changing the grammatical probabilities of the candidate pool.

The consistency parameter, which prefers repeated use of the same grammatical sequences, incrementally changes the candidate pool as the encoder works its way through the ~500 messages. Once a linear sequence has been selected for each message (via the contiguity parameter, as described above), this selection has a permanent effect on the candidate pool. For example, if we are at the initial state with a single token for each sequence, and the first message is encoded as *these-cats-black-three*, the sequence Dem-N-Adj-Num will be boosted in the candidate pool for subsequent messages. With a consistency multiplier of 10, this sequence will now have a probability of $10/33 = 0.30$, while all other sequences will have their probabilities accordingly reduced to $1/33 = 0.03$. This is now the candidate pool for encoding the next message, which will again apply the contiguity parameter to modulate the probabilities according to the specific chunking of the next message. This is a rich-get-richer system, which tends to converge on a single grammatical structure.

Notice that the contiguity parameter is an effect on individual messages, based on their specific symbolic PMIs, whereas the consistency parameter is more systematic, leading gradually to generalised effects on all messages. The interaction of these parameters captures the way that natural languages generally favour consistent sequencing by grammatical category, but at the same time do show some item-specific ordering (Simpson and Withgott 1986; Mansfield, Stoll and Bickel 2020; Mansfield et al. 2022). Also note that the effect of the multipliers is quite large at the initial stages of the encoder, when the candidate pool contains few tokens, but becomes smaller as the system encodes more messages and the overall size of the pool increases. If the consistency parameter is strong enough, the model gradually converges on one grammatical sequence.

The model converges on structures that are similar to natural languages

The effect of the two linearisation parameters in concert is that complex symbols tend to be contiguous in transmission sequences, and the grammatical categories that are most often chunked into complex symbols will tend to be adjacent not just in those messages where the chunking applies, but also in other messages (cf. Culbertson, Schouwstra and Kirby 2020, 710). However, the two principles can be in competition if, say, categories {N, Adj} overall tend to have higher PMI on average, but in a specific message, say {N, Num} has a higher PMI. This will be illustrated below.

The main test of our model is whether it tends to produce the same NP-internal orderings that are observed in natural languages (Dryer 2018). Inspection of the encoder's output reveals that it does indeed produce typologically preferred grammatical category sequences, at a broad range of parameter settings. For example, with contiguity multiplier = 20 and consistency multiplier = 5, running 20 iterations of the encoder converged on typologically preferred sequences 17 times (for details on 'convergence', see below), and on typologically dispreferred sequences just 3 times. This shows that for these parameter weights, the encoder succeeds in producing typologically preferred sequences most of the time. But it is also capable of producing typologically dispreferred sequences, just as these do occur occasionally in natural languages.

(2) *Sample of 20 converged sequences (* = typologically dispreferred)*

Dem-Adj-N-Num	*Dem-N-Num-Adj
Num-N-Adj-Dem	Num-N-Adj-Dem
Dem-Adj-N-Num	Adj-N-Num-Dem
Dem-Adj-N-Num	Adj-N-Num-Dem
Num-N-Adj-Dem	N-Adj-Num-Dem
Adj-N-Num-Dem	N-Adj-Num-Dem
Adj-N-Num-Dem	*Dem-N-Adj-Num
Adj-N-Num-Dem	Adj-N-Num-Dem
Num-N-Adj-Dem	Num-N-Adj-Dem
Num-Adj-N-Dem	*N-Adj-Dem-Num

The typologically preferred sequences represent a clear probabilistic bias in natural languages, rather than a hard constraint (cf. Bickel 2015). As noted above, in Dryer's (2018) sample of natural languages, 83 per cent exhibited one of the typologically preferred sequences. To further test whether the encoder produces a similar probabilistic bias, it was iterated 20 times each over a range of parameter weights for the contiguity and consistency parameters, with the outcome of interest being what proportion of the iterations converged on a typologically preferred linear schema. The chunking threshold was also tested at seven different levels, but this did not significantly affect the results and therefore is not reported below.⁷ The results illustrated here group over seven chunking threshold values, each iterated 20 times, giving 140 iterations at each pair of contiguity and consistency parameter values.

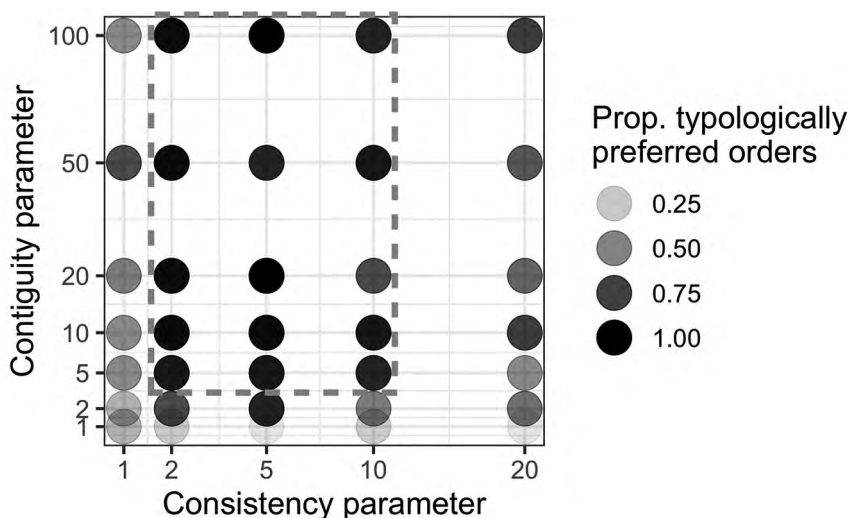


Figure 18.3: Sequencing of {N, Adj, Num, Dem}, as manipulated by the contiguity parameter and the consistency parameter.

Note: Darker shading indicates a higher proportion of iterations converging on one of the eight typologically common category sequences identified by Dryer (2018). The 'sweet region' is indicated by the dashed rectangle.

Source: Authors.

⁷ The chunking threshold described in the section above was tested across the range of values that resulted in a moderate degree of chunking (average symbol count > 2). Different chunking threshold values within the range used made no noticeable difference to the results.

Figure 18.3 illustrates the proportion of typologically preferred sequences output at each pairing of contiguity and consistency parameters.⁸ The figure shows that there is a parameter range in which the encoder produces a similar probabilistic bias to that observed in natural languages. In particular, when the consistency multiplier is in the range of 2–10, and the contiguity multiplier is in the range of 5–100, the encoder produces typologically common orders in 91 per cent of its iterations (percentages for specific parameter combinations range from 70–100, $\sigma = 7$). This is similar to the 83 per cent of natural languages that exhibit these orderings. We call this parameter range the ‘sweet region’, indicated in Figure 18.3 by the dashed rectangle.

Figure 18.3 shows that the encoder fails to converge frequently on typologically common orders if the contiguity parameter is too low (< 5), indicating that the complex symbol contiguity principle is a requirement for typologically preferred sequencing in this model. The consistency parameter must have a weight above 1 (since multiplication by 1 has no actual effect), as otherwise the encoder tends not to converge on any particular category sequence at all. The encoder also struggles to converge on typologically preferred sequences if the consistency parameter is *too high* (> 10). This is because the encoder becomes over-eager to converge on a category sequence, and if it happens to encounter messages that have uncharacteristic chunking properties early in the message repertoire, it may converge on these sequences before more characteristically chunked messages have a chance to take effect.

The lack of harmonic dependencies

As outlined previously, within the sweet region, the encoder favours the same eight sequences as natural languages. However, natural languages also exhibit further biases among these eight possibilities. In particular, Dryer’s (2018) sample shows that sequences are especially favoured where the noun is at one edge, and the other elements are sequenced in order of average PMI, i.e. N-Adj-Num-Dem or Dem-Num-Adj-N. Taking the N as the head and the others as dependents, we can say that these two particular sequences achieve a fully harmonic direction of dependencies.

⁸ We first tested parameter values heuristically to identify the range in which results were sensitive to parameter changes, namely contiguity multipliers in the range of 1–100, and consistent sequencing multipliers in the range of 1–20.

The encoder does *not* favour the two fully harmonic sequences among the eight frequent sequences, but instead most favours Adj-N-Num-Dem and Num-N-Adj-Dem, accounting for 64 per cent of convergence sequences in the sweet region. What these two sequences have in common is that the two categories that have the highest average PMI with the noun, namely {Adj, Num}, are both adjacent to the noun. These sequences are presumably favoured by the encoder because they satisfy the complex symbol contiguity principle both for symbols such as (*cats*, *black*) and symbols such as (*wheels*, *four*).

This discrepancy between the encoder's output and natural languages suggests that our model is missing a crucial principle, namely harmonic dependency ordering (Greenberg 1963; Culbertson and Newport 2015; Jing, Blasi and Bickel 2022). Harmonic ordering is satisfied by the sequences most favoured in natural languages, N-Adj-Num-Dem and Dem-Num-Adj-N, but not by the sequences most favoured by the encoder. We address this discrepancy in our follow-up study using natural language corpora.

Language-internal variation

The studies of NP ordering mentioned above focus on which ordering dominates in each language, while setting aside the question of language-internal variation. However, studies that do report on variable ordering of NP elements suggest that this is quite common cross-linguistically (Rijkhoff 2004; Talamo and Verkerk 2022). Even English, which is usually treated as having a fixed order Dem-Num-Adj-N, in fact exhibits some variation:

- (3) a. *The battery support and the four impressive wheels*⁹
 b. *You already get some quite impressive four wheels for that*¹⁰

We might suppose that examples such as (3b), which breaks the supposed rules of the English NP template, arise due to the complex symbol contiguity principle being triggered by informational chunking of (*wheels*, *four*), while examples such as (3a) are produced by the dominance of the consistent sequencing principle.¹¹

9 affordablemedicalusa.com/blog/4-wheel-compact-travel-scooter-ensures-an-easy-driving-experience-with-adjustable-setup/.

10 Florian Buechting, 'Driving the Ferrari 488 with Pushstart in Maranello', flyctory.com, 29 June 2021, flyctory.com/2021/06/29/driving-the-ferrari-488-with-pushstart-in-maranello/.

11 Another potential influence in this example is the elaboration of the adjective into a heavier element, *quite impressive*. We thank James Gray for pointing this out.

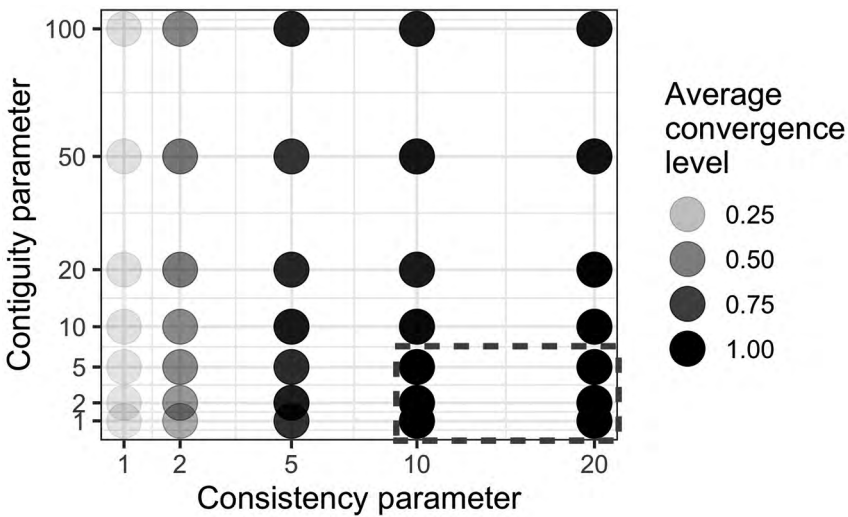


Figure 18.4: Degrees of convergence in categorical ordering, as manipulated by the contiguity parameter and the consistency parameter.

Note: Darker shading indicates a higher degree of convergence, i.e. the degree to which one particular category sequence dominates in the last 50 messages of each iteration of the encoder.

As mentioned above, the encoder’s linearisation algorithm is probabilistic and has an initial state where all sequences are equally probable. Language-internal variation is therefore inherent to the model. However, the consistency parameter tends to gradually bestow dominance on particular category sequences, or indeed on just one sequence. In the results above, the encoder is said to ‘converge upon’ particular category sequences, and we evaluate this based on whichever sequence is most frequent in the last 50 messages of the repertoire. Thus, there are various degrees of convergence, as illustrated in Figure 18.4. This figure shows the same range of linearisation parameter values as above, with 140 iterations at each pair of values; but here, for each, the dot shading represents the average degree of convergence in the iterations.¹² As we might expect, increases in the consistency parameter lead to greater convergence; but there is also a mild effect from the contiguity parameter, which slightly reduces convergence at its highest values (though this is difficult to distinguish in the figure shading). This is presumably because a strong contiguity parameter favours more message-specific sequencing, e.g. *cats-black-three-these* with N-Adj contiguity versus

¹² The figure again collapses weightings of the chunking threshold, though in this case there did appear to be a very small effect, with more chunking slightly reducing the degree of convergence.

wheels-four-white-those with N-Num contiguity. The dashed rectangle in the lower-right corner indicates parameter values at which mean convergence reaches 100 per cent, i.e. completely fixed category ordering in the last 50 messages. We call this the ‘fixation region’. However, we do not assume the fixation region to be an important outcome for validation of the model, as mentioned previously, we do not know to what extent NP-internal ordering is rigid in natural languages.

In this section we have reflected on free variation, but the principles in our model might also produce outcomes where specific lexical combinations have their own fixed ordering, which are not consistent by grammatical category (e.g. *black-cats*, *dogs-brown*). For this to occur, the structural priming mechanism would need to take into account specific lexical items, either instead of, or as well as, grammatical category labels (Pickering and Branigan 1998). We therefore expect that our principles of linear ordering are quite compatible with the kinds of idiosyncratic templates described by Simpson (see above), and can do so based on general principles rather than stipulative exceptions.

The model fails without informational chunking

To validate the claim that principles I–III can explain the sequencing of grammatical categories, we should be able to show that typologically preferred orders *fail* to emerge when any of the principles is absent from the model. This has already been done for the linearisation principles (II and III) in the previous section, where Figure 18.3 illustrated that if the parameter weight for either contiguity or consistency is too small, then the encoder fails to converge on typologically preferred category sequences.

We can additionally confirm that informational chunking is required to generate typologically preferred category sequences. Figure 18.5 illustrates the output of the encoder when it is run with the same overall degree of chunking, but with purely random as opposed to Informational Chunking. The figure shows no evidence of a bias towards typologically preferred sequences, the proportion of which mostly hovers around 30–40 per cent, which is what we might expect from a random output, given that these orders constitute 33 per cent (8/24) of possible orders. Furthermore the results lose any structural regularity with respect to the two linearisation parameters.

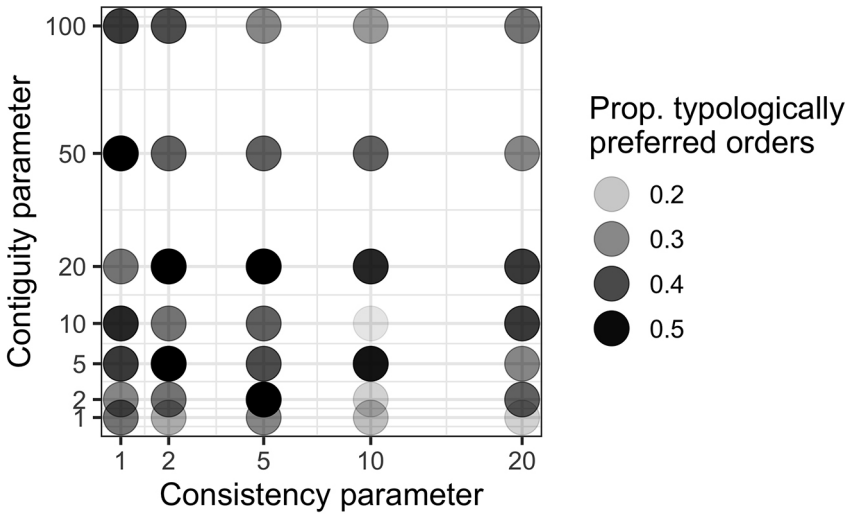


Figure 18.5: Output sequences when the encoder is run with random chunking, as opposed to PMI-driven chunking.

Source: Authors.

What have we done?

The encoder illustrated in this chapter suggests that the NP-internal grammatical category orderings for {N, Adj, Num, Dem} can be explained as an emergent effect of independently evidenced principles of linguistic cognition. Most importantly, the encoder is not designed around any specific grammatical categories and does not involve any ‘hard-coding’ of specific grammatical orderings. Rather, the orderings emerge from the inter-predictability of grammatically categorised symbols in the input data. We used artificial input data to capture the relevant inter-predictability distribution, which corpus research has shown to be relatively constant in every one of 24 languages studied (Culbertson, Schouwstra and Kirby 2020). Like Culbertson and colleagues, we assume that these symbolic inter-predictability patterns arise from semantic properties of adjectives, numbers and demonstratives, and how these reflect human experiences of the world. For example, adjectival qualities are more predictably associated with objects, while numerosities are less predictably associated with objects.

The principles implemented in the encoder are not specific to NP-internal grammar. Rather, they are intended to apply to any grammatical structure in which linear proximity corresponds to degrees of inter-predictability.

As recent studies have shown, this property of ‘information locality’ appears in a wide range of structures, and, therefore, the encoder modelled here has broad relevance to grammatical theory.

By formulating an explicit model of how inter-predictability can produce grammatical structure, this study advances the cause of usage-based or ‘emergent’ grammar. Much work in this tradition has discussed ‘entrenchment’ of linguistic structures from repeated use (e.g. Bybee 2006; Schmid 2016; Diessel 2019). However, implementation of formal models has been scarce in this tradition (though see Bod 1998; Steels and de Beule 2006), and much of the literature focuses on specific lexical constructions rather than general grammatical patterns. The current study advances the field by addressing how specific messages interact with generalised grammatical schemas, and by testing an explicit model of such interaction.

The encoder illustrated above is also relevant to theoretical debates about hierarchical structure versus linear sequences. Simpson and Withgott (1986) proposed a role for linear schemas (templates) in grammar. The encoder formalises this proposal and expands it beyond idiosyncratic morphological constructions. We have shown that our encoder is capable of converging upon consistent grammatical category sequences, while also noting that variation is inherent to the model. We also briefly observed that a similar encoder might produce lexically idiosyncratic sequences, using a modified version of the structural priming algorithm.

To reiterate, we do not claim that there is no hierarchical structure at all in syntax: unbounded recursive structures such as relative clauses and reiterated complex units like the NP itself, both suggest a hierarchical structure in which sets of signs are grouped together. But the assumption that all linguistic symbols fit into an underlying binary hierarchy seems unwarranted and requires more complex structure than necessary whenever a linear schema will do just as well (Culicover and Jackendoff 2005, 111). Our encoder does, in fact, produce hierarchical structures, but these are not ‘syntactic’ in any conventional sense. Instead, the encoder applies recursive grouping to symbols in a storage-and-retrieval system, similar to the concept of a ‘mental lexicon’. The hierarchical structures involve specific symbols, rather than generalised grammatical categories. Grammatical categories only play a role in the linearisation process, where structural priming favours the re-use of previous category sequences. Therefore, generalised grammatical category schemas, i.e. morphosyntactic structure, are produced as the aggregate outcome of specific symbol combinations. It emerges iteratively and is never more than a distributional property of a probabilistic system.

The question of morphology versus syntax has not been addressed in this chapter, though it has hovered in our peripheral vision. Simpson's templates were intended to be morphological and have subsequently been taken up primarily within morphology (Stump 1997; Nordlinger 2010a), though see also Good (2016). Our suggestion is that principles of linear ordering do *not* respect the purported morphology/syntax distinction – if indeed such a distinction can be rigorously maintained (Haspelmath 2011; Tallman 2020). Future development of the encoder could attempt to capture the grammatical properties that are usually referenced to distinguish morphology from syntax. For example, it has been shown that selectional restrictiveness, one of the core properties involved in 'bound morphology', is associated with inter-predictability (Mansfield 2021). This would contribute to a wider project that uses information theory to investigate the cluster of properties that tend to distinguish morphology from syntax (Ackerman and Malouf 2013; Blevins 2016). While attempts to formulate morphology/syntax as a categorical binary have been inconclusive, information theory may present a fresh approach based on gradient properties.

Finally – what's missing? Several limitations were noted above, for example, the use of artificial rather than natural corpus data and the unrealistic restriction that each message consists of exactly four symbolic types, {N, Adj, Num, Dem}. The encoder illustrated here is merely a proof-of-concept, and its promising results should be validated by further research. But there are other, more fundamental ways in which the encoder does not accurately represent human language. The encoder models a 'closed system' that outputs a fixed message repertoire, with no interaction between agents or generations of learners. It is like a lonely individual who starts out with a proto-language and evolves into grammatically structured utterances by speaking into the void. Of course, natural languages have not evolved in that way, but rather through cultural evolutionary processes shaped by the exchange of messages between many individuals, each of whom goes through an acquisition process. Where our model has a very simple pool of remembered grammatical sequences, natural language instead evolves by the gradual accretion of group memory over generations. There is also the matter of where grammatical categories such as {N, Adj, Num, Dem} come from, how they should be identified and whether they are even theoretically defensible (Anward 2000; Croft 2001; Bisang 2010; Kenesei 2020). All these will be challenges for future related research.

References

- Ackerman, Farrell and Robert Malouf. 2013. 'Morphological Organization: The Low Conditional Entropy Conjecture'. *Language* 89, no. 3: 429–64.
- Alexiadou, Artemis, Liliane Haegeman and Melita Stavrou. 2008. *Noun Phrase in the Generative Perspective*. Berlin: De Gruyter Mouton. doi.org/10.1515/9783110207491.
- Ambridge, Ben. 2020a. 'Abstractions Made of Exemplars or "You're All Right, and I've Changed My Mind": Response to Commentators'. *First Language* 40, no. 5–6: 640–59. doi.org/10.1177/0142723720949723.
- Ambridge, Ben. 2020b. 'Against Stored Abstractions: A Radical Exemplar Model of Language Acquisition'. *First Language* 40, no. 5–6: 509–59. doi.org/10.1177/0142723719869731.
- Anward, Jan. 2000. 'A Dynamic Model of Part-of-Speech Differentiation'. In *Approaches to the Typology of Word Classes*, edited by Petra M. Vogel and Bernard Comrie, 3–46. Berlin: De Gruyter Mouton. doi.org/10.1515/9783110806120.3.
- Arnon, Inbal and Neal Snider. 2010. 'More than Words: Frequency Effects for Multi-Word Phrases'. *Journal of Memory and Language* 62, no. 1: 67–82. doi.org/10.1016/j.jml.2009.09.005.
- Aylett, Matthew and Alice Turk. 2006. 'Language Redundancy Predicts Syllabic Duration and the Spectral Characteristics of Vocalic Syllable Nuclei'. *The Journal of the Acoustical Society of America* 119, no. 5: 3048–58. doi.org/10.1121/1.2188331.
- Baayen, R. Harald and R. Schreuder, eds. 2003. *Morphological Structure in Language Processing*. Berlin: De Gruyter.
- Biber, Douglas, Susan Conrad and Randi Reppen. 1998. *Corpus Linguistics: Investigating Language Structure and Use*. Cambridge: Cambridge University Press.
- Bickel, Balthasar. 2015. 'Distributional Typology: Statistical Inquiries into the Dynamics of Linguistic Diversity'. In *The Oxford Handbook of Linguistic Analysis*, edited by Bernd Heine and Heiko Narrog, 901–23. 2nd ed. Oxford: Oxford University Press.
- Bisang, Walter. 2010. 'Word Classes'. In *The Oxford Handbook of Linguistic Typology*, edited by Jae Jung Song. Oxford: Oxford University Press.

- Blevins, James P. 2016. *Word and Paradigm Morphology*. Oxford: Oxford University Press.
- Blumenthal-Dramé, Alice, Volkmar Glauche, Tobias Bormann, Cornelius Weiller, Mariacristina Musso and Bernd Kortmann. 2017. 'Frequency and Chunking in Derived Words: A Parametric fMRI Study'. *Journal of Cognitive Neuroscience* 29, no. 7: 1162–77. doi.org/10.1162/jocn_a_01120.
- Bod, Rens. 1998. *Beyond Grammar: An Experience-Based Theory of Language*. Stanford: CSLI Publications.
- Brinton, Laurel J. and Elizabeth Closs Traugott. 2005. *Lexicalization and Language Change*. Cambridge: Cambridge University Press. doi.org/10.1017/CBO9780511615962.
- Bybee, Joan L. 2001. *Phonology and Language Use*. Cambridge: Cambridge University Press.
- Bybee, Joan L. 2006. 'From Usage to Grammar: The Mind's Response to Repetition'. *Language* 82, no. 4: 711–33.
- Chomsky, Noam. 1957. *Syntactic Structures*. The Hague: Mouton.
- Chomsky, Noam and Howard Lasnik. 2015 [1993]. 'The Theory of Principles and Parameters'. In *Syntax: An International Handbook of Contemporary Research*, edited by Joachim Jacobs, Arnim von Stechow, Wolfgang Sternefeld and Theo Vennemann, 506–69. Berlin: De Gruyter Mouton.
- Church, Kenneth Ward and Patrick Hanks. 1990. 'Word Association Norms, Mutual Information, and Lexicography'. *Computational Linguistics* 16, no. 1: 22–29.
- Cinque, Guglielmo. 2005. 'Deriving Greenberg's Universal 20 and Its Exceptions'. *Linguistic Inquiry* 36, no. 3: 315–32. doi.org/10.1162/0024389054396917.
- Cinque, Guglielmo and Luigi Rizzi. 2009. 'The Cartography of Syntactic Structures'. In *The Oxford Handbook of Linguistic Analysis*, edited by Bernd Heine and Heiko Narrog, 52–66. Oxford: Oxford University Press.
- Contreras Kallens, Pablo and Morten H. Christiansen. 2022. 'Models of Language and Multiword Expressions'. *Frontiers in Artificial Intelligence* 5. doi.org/10.3389/frai.2022.781962.
- Cover, Thomas A. and Joy A. Thomas. 2002. *Elements of Information Theory*. 2nd ed. London: Wiley.
- Croft, William. 2001. *Radical Construction Grammar: Syntactic Theory in Typological Perspective*. Oxford: Oxford University Press.

- Crysmann, Berthold and Olivier Bonami. 2016. 'Variable Morphotactics in Information-Based Morphology'. *Journal of Linguistics* 52, no. 2: 311–74.
- Culbertson, Jennifer and Elissa L. Newport. 2015. 'Harmonic Biases in Child Learners: In Support of Language Universals'. *Cognition* 139: 71–82. doi.org/10.1016/j.cognition.2015.02.007.
- Culbertson, Jennifer, Marieke Schouwstra and Simon Kirby. 2020. 'From the World to Word Order: Deriving Biases in Noun Phrase Order from Statistical Properties of the World'. *Language* 96, no. 3: 696–717.
- Culicover, Peter W. and Ray Jackendoff. 2005. *Simpler Syntax*. Oxford: Oxford University Press.
- Diessel, Holger. 2019. *The Grammar Network. How Linguistic Structure Is Shaped by Language Use*. Cambridge: Cambridge University Press.
- Dik, Simon C. 1997. *The Theory of Functional Grammar: Part 1, the Structure of the Clause*. Berlin: De Gruyter.
- Dryer, Matthew S. 2009. 'The Branching Direction Theory of Word Order Correlations Revisited'. In *Universals of Language Today*, edited by Sergio Scalise, Elisabetta Magni and Antonietta Bisetto, 185–207. Dordrecht: Springer. doi.org/10.1007/978-1-4020-8825-4_10.
- Dryer, Matthew S. 2018. 'On the Order of Demonstrative, Numeral, Adjective, and Noun'. *Language* 94, no. 4: 798–833. doi.org/10.1353/lan.2018.0054.
- Everaert, Martin B. H., Marinus A. C. Huybregts, Noam Chomsky, Robert C. Berwick and Johan J. Bolhuis. 2015. 'Structures, Not Strings: Linguistics as Part of the Cognitive Sciences'. *Trends in Cognitive Sciences* 19, no. 12: 729–43. doi.org/10.1016/j.tics.2015.09.008.
- Fukumura, Kumiko and Shi Zhang. 2023. 'The Interplay between Syntactic and Non-Syntactic Structure in Language Production'. *Journal of Memory and Language* 128: 104385. doi.org/10.1016/j.jml.2022.104385.
- Futrell, Richard. 2019. 'Information-Theoretic Locality Properties of Natural Language'. In *Proceedings of the First Workshop on Quantitative Syntax (Quasy, SyntaxFest 2019)*, 2–15. Paris: Association for Computational Linguistics. doi.org/10.18653/v1/W19-7902.
- Futrell, Richard, Roger P. Levy and Edward Gibson. 2020. 'Dependency Locality as an Explanatory Principle for Word Order'. *Language* 96, no. 2: 371–412. doi.org/10.1353/lan.2020.0024.

- Gibson, Edward, R. Futrell, S. T. Piantadosi, I. Dautriche, K. Mahowald, L. Bergen and R. Levy. 2019. 'How Efficiency Shapes Human Language'. *Trends in Cognitive Sciences* 23, no. 5: 389–407. doi.org/10.1016/j.tics.2019.02.003.
- Good, Jeff. 2016. *The Linguistic Typology of Templates*. Cambridge: Cambridge University Press.
- Greenberg, Joseph. 1963. *Universals of Language*. Cambridge: MIT Press.
- Grünwald, Peter D. 2007. *The Minimum Description Length Principle*. Cambridge: MIT Press. doi.org/10.7551/mitpress/4643.001.0001.
- Hahn, Michael, Judith Degen and Richard Futrell. 2021. 'Modeling Word and Morpheme Order in Natural Language as an Efficient Trade-off of Memory and Surprisal'. *Psychological Review* 128, no. 4: 726–56. doi.org/10.1037/rev0000269.
- Hahn, Michael, Judith Degen, Noah D. Goodman, Dan Jurafsky and Richard Futrell. 2018. 'An Information-Theoretic Explanation of Adjective Ordering Preferences'. In *Proceedings of the 40th Annual Meeting of the Cognitive Science Society*, 1766–71. London: Cognitive Science Society.
- Hahn, Michael, Rebecca Mathew and Judith Degen. 2022. 'Morpheme Ordering across Languages Reflects Optimization for Memory Efficiency'. *Open Mind: Discoveries in Cognitive Science* 5: 208–32. doi.org/10.1162/opmi_a_00051.
- Hale, John. 2001. 'A Probabilistic Earley Parser as a Psycholinguistic Model'. In *Second Meeting of the North American Chapter of the Association for Computational Linguistics*. ACL Anthology. aclanthology.org/N01-1021.
- Hartsuiker, Robert J., Herman H. J. Kolk and Philippine Huiskamp. 1999. 'Priming Word Order in Sentence Production'. *The Quarterly Journal of Experimental Psychology* 52, no. 1: 129–47. doi.org/10.1080/713755798.
- Haspelmath, Martin. 2011. 'The Indeterminacy of Word Segmentation and the Nature of Morphology and Syntax'. *Folia Linguistica* 45, no. 1: 31–80.
- Hawkins, John A. 2004. *Efficiency and Complexity in Grammars*. Oxford: Oxford University Press. doi.org/10.1093/acprof:oso/9780199252695.001.0001.
- Hay, Jennifer. 2002. 'From Speech Perception to Morphology: Affix Ordering Revisited'. *Language* 78, no. 3: 527–55. doi.org/10.1353/lan.2002.0159.
- Hopper, Paul. 1987. 'Emergent Grammar'. *Proceedings of the Thirteenth Annual Meeting of the Berkeley Linguistics Society* 13: 139–57. doi.org/10.3765/bls.v13i0.1834.

- Hopper, Paul and Elizabeth Closs Traugott. 2003. *Grammaticalization*. 2nd ed. Cambridge: Cambridge University Press.
- Jaeger, T. Florian. 2010. 'Redundancy and Reduction: Speakers Manage Syntactic Information Density'. *Cognitive Psychology* 61, no. 1: 23–62. doi.org/10.1016/j.cogpsych.2010.02.002.
- Jaeger, T. Florian and Esteban Buz. 2017. 'Signal Reduction and Linguistic Encoding'. In *The Handbook of Psycholinguistics*, edited by E. M. Fernández and H. Smith Cairns, 38–81. Hoboken: John Wiley & Sons. doi.org/10.1002/9781118829516.ch3.
- Jing, Yingqi, Damián E. Blasi and Balthasar Bickel. 2022. 'Dependency-Length Minimization and Its Limits: A Possible Role for a Probabilistic Version of the Final-over-Final Condition'. *Language* 98, no. 3: 397–418. doi.org/10.1353/lan.0.0267.
- Kayne, Richard S. 1981. 'Unambiguous Paths'. In *Levels of Syntactic Representation*, edited by Robert May and Jan Koster, 143–84. Berlin: De Gruyter Mouton. doi.org/10.1515/9783110874167-006.
- Kayne, Richard S. 1994. *The Antisymmetry of Syntax*. Cambridge: MIT Press.
- Kenesei, István. 2020. 'Life without Word Classes: On a New Approach to Categorization'. In *Syntactic Architecture and Its Consequences II: Between Syntax and Morphology*, edited by András Bárány, Theresa Biberauer, Jamie Douglas and Sten Vikner, 67–80. Berlin: Language Science Press.
- Kirby, Simon, Monica Tamariz, Hannah Cornish and Kenny Smith. 2015. 'Compression and Communication in the Cultural Evolution of Linguistic Structure'. *Cognition* 141: 87–102. doi.org/10.1016/j.cognition.2015.03.016.
- Kulmizev, Artur and Joakim Nivre. 2022. 'Schrödinger's Tree – On Syntax and Neural Language Models'. *Frontiers in Artificial Intelligence* 5. doi.org/10.3389/frai.2022.796788.
- Kuperman, Victor, Raymond Bertram and R. Harald Baayen. 2010. 'Processing Trade-Offs in the Reading of Dutch Derived Words'. *Journal of Memory and Language* 62, no. 2: 83–97. doi.org/10.1016/j.jml.2009.10.001.
- Levshina, Natalia. 2022. *Communicative Efficiency: Language Structure and Use*. Cambridge: Cambridge University Press. doi.org/10.1017/9781108887809.
- Levy, Roger. 2008. 'Expectation-Based Syntactic Comprehension'. *Cognition* 106, no. 3: 1126–77. doi.org/10.1016/j.cognition.2007.05.006.

- Mansfield, John Basil. 2021. 'The Word as a Unit of Internal Predictability'. *Linguistics* 59, no. 6: 1427–72. doi.org/10.1515/ling-2020-0118.
- Mansfield, J., C. Saldana, P. Hurst, R. Nordlinger, S. Stoll, B. Bickel and A. Perfors. 2022. 'Category Clustering and Morphological Learning'. *Cognitive Science* 46, no. 2: e13107. doi.org/10.1111/cogs.13107.
- Mansfield, John Basil, Sabine Stoll and Balthasar Bickel. 2020. 'Category Clustering: A Probabilistic Bias in the Morphology of Argument Marking'. *Language* 96, no. 2: 255–93.
- Marslen-Wilson, William D. 2007. 'Morphological Processes in Language Comprehension'. In *The Oxford Handbook of Psycholinguistics*, edited by M. Gareth Gaskell, 175–93. Oxford: Oxford University Press. doi.org/10.1093/oxfordhb/9780198568971.013.0011.
- McCarthy, John J. 1982. 'Prosodic Structure and Expletive Infixation'. *Language* 58, no. 3: 574–90.
- McCauley, Stewart M. and Morten H. Christiansen. 2019. 'Language Learning as Language Use: A Cross-Linguistic Model of Child Language Development'. *Psychological Review* 126, no. 1: 1–51. doi.org/10.1037%2Frev0000126.
- Nordlinger, Rachel. 2010a. 'Agreement in Murrinh-Patha Serial Verbs'. In *Selected Papers from the 2009 Conference of the Australian Linguistic Society*, edited by Yvonne Treis and Rik De Busser.
- Nordlinger, Rachel. 2010b. 'Verbal Morphology in Murrinh-Patha: Evidence for Templates'. *Morphology* 20, no. 2: 321–41.
- Perlmutter, David M. 1971. *Deep and Surface Structure Constraints in Syntax*. New York: Holt, Rinehart & Winston.
- Pickering, Martin J. and Holly P. Branigan. 1998. 'The Representation of Verbs: Evidence from Syntactic Priming in Language Production'. *Journal of Memory and Language* 39, no. 4: 633–51. doi.org/10.1006/jmla.1998.2592.
- Pijpops, Dirk, Isabeau De Smet and Freek Van de Velde. 2018. 'Constructional Contamination in Morphology and Syntax: Four Case Studies'. *Constructions and Frames* 10, no. 2: 269–305. doi.org/10.1075/cf.00021.pij.
- Plag, Ingo and Harald Baayen. 2009. 'Suffix Ordering and Morphological Processing'. *Language* 85, no. 1: 109–52.
- Rijkhoff, Jan. 2004. *The Noun Phrase*. Oxford: Oxford University Press.

- Schmid, Hans-Jörg. 2016. *Entrenchment and the Psychology of Language Learning: How We Reorganize and Adapt Linguistic Knowledge*. Berlin: De Gruyter Mouton.
- Seržant, Ilja A. and George Moroz. 2022. 'Universal Attractors in Language Evolution Provide Evidence for the Kinds of Efficiency Pressures Involved'. *Humanities and Social Sciences Communications* 9, no. 1: 1–9. doi.org/10.1057/s41599-022-01072-0.
- Shannon, Claude E. 1948. 'A Mathematical Theory of Communication'. *Bell System Technical Journal* 27, no. 3: 379–423.
- Simpson, Jane and M. Withgott. 1986. 'Pronominal Clitic Clusters and Templates'. In *The Syntax of Pronominal Clitics*, edited by Hagit Borer, 147–74. New York: Academic Press. doi.org/10.1163/9789004373150_008.
- Sinclair, John. 1991. *Corpus, Concordance, Collocation*. Oxford: Oxford University Press.
- Steels, Luc and Joachim de Beule. 2006. 'A (Very) Brief Introduction to Fluid Construction Grammar'. In *Proceedings of the Third Workshop on Scalable Natural Language Understanding*, edited by J. Allen, J. Alexandersson, J. Feldman and R. Porzel, 73–80. New York City: Association for Computational Linguistics. aclanthology.org/W06-3510.
- Stump, Gregory T. 1997. 'Template Morphology and Inflectional Morphology'. In *Yearbook of Morphology 1996*, edited by Geert Booij and Jaap van Marle, 217–41. Amsterdam: Kluwer Academic Publishers.
- Szmrecsanyi, Benedikt. 2006. *Morphosyntactic Persistence in Spoken English: A Corpus Study at the Intersection of Variationist Sociolinguistics, Psycholinguistics, and Discourse Analysis*. Berlin: De Gruyter Mouton. doi.org/10.1515/9783110197808.
- Talamo, Luigi and Annemarie Verkerk. 2022. 'A New Methodology for an Old Problem: A Corpus-Based Typology of Adnominal Word Order in European Languages'. *Italian Journal of Linguistics* 34, no. 2: 171–226.
- Tallman, Adam J. R. 2020. 'Beyond Grammatical and Phonological Words'. *Language and Linguistics Compass* 14, no. 2: e12364. doi.org/10.1111/lnc3.12364.
- ten Hacken, Pius. 2019. 'The Mental Lexicon in Jackendoff's Parallel Architecture'. In *Word Formation in Parallel Architecture: The Case for a Separate Component*, edited by Pius ten Hacken, 3–22. Cham: Springer. doi.org/10.1007/978-3-030-18009-6_2.

- Van Gompel, Roger P. G. and Manabu Arai. 2018. 'Structural Priming in Bilinguals'. *Bilingualism: Language and Cognition* 21, no. 3: 448–55. doi.org/10.1017/S1366728917000542.
- Wells, Rulon S. 1947. 'Immediate Constituents'. *Language* 23, no. 2: 81–117. doi.org/10.2307/410382.
- Wray, Alison. 2002. *Formulaic Language and the Lexicon*. Cambridge: Cambridge University Press.
- Wray, Alison. 2017. 'Formulaic Sequences as a Regulatory Mechanism for Cognitive Perturbations during the Achievement of Social Goals'. *Topics in Cognitive Science* 9, no. 3: 569–87. doi.org/10.1111/tops.12257.
- Zipf, George K. 1949. *Human Behavior and the Principle of Least Effort: An Introduction to Human Ecology*. Cambridge: Addison-Wesley Press.

This text is taken from *Projecting Voices: Studies in Language and Linguistics in Honour of Jane Simpson*, edited by Carmel O'Shannessy, James Gray and Denise Angelo, published 2025 by ANU Press, The Australian National University, Canberra, Australia.

doi.org/10.22459/PV.2025.18