

Inflectional classes without class features

JAN DON^{1*}, FENNA BERGSMA², ANNE MERKUUR² and MEG SMITH³

¹ Open University, Netherlands

² Fryske Akademy, Netherlands

³ Universiteit van Amsterdam, Netherlands

Received: October 31, 2022 • Revised manuscript received: May 31, 2023 • Accepted: May 31, 2023

Published online: August 10, 2023

© 2023 Akadémiai Kiadó, Budapest



ABSTRACT

In this paper, we propose a comprehensive account of the paradigms of Frisian verb-classes. Verb-classes in Frisian are an example of a more general phenomenon of inflectional classes that we encounter in many natural languages across the major word classes. Members of different inflectional classes show different paradigms. Traditionally, inflectional classes have been analyzed using class-features (see e.g., Marzi et al. 2020). However, such features suffer from being *ad hoc* devices that seem to have no other function in the grammar than to code this difference. In the present analysis we propose that the verb stems from different classes show a difference in size. Using phrasal spell-out, we will show that these stems differ in the amount of morpho-syntactic structure that they may realize, rendering class-features superfluous.

KEYWORDS

verbal inflection, germanic, Frisian, syncretism, verb class

1. INTRODUCTION

The aim of this paper is to account for the differences between Frisian verbal classes without making use of language-specific class features. The major divisions in the Frisian verbal system are between regular (weak and strong) and irregular verbs (Dyk 2020). At least within the regular weak verbs, there is a division between class I and class II, which both have their own inflectional paradigms. We analyze the difference in verb classes as a size difference between the verb stems of the different classes. We focus on Frisian regular verbs, and not yet

* Corresponding author. E-mail: jan.don@ou.nl

consider any irregular verbs in this paper.¹ Interestingly, we show that the same technique can be applied to both weak and strong regular verbs.

As an example of the difference in class, consider the weak verbs, *bakke* [bakə], a so-called class I verb, and *wurkje* [wœrkjə], belonging to class II.

(1)

class I	Pres.	Past	class II	Pres.	Past
1 sing.	bak	baktə	1 sing.	uœrkjə	uœrkə
2 sing.	bakst	baktəst	2 sing.	uœrkəst	uœrkəst
3 sing.	bakt	baktə	3 sing.	uœrket	uœrkə
1 plur.	bakə	baktən	1 plur.	uœrkjə	uœrkən
2 plur.	bakə	baktən	2 plur.	uœrkjə	uœrkən
3 plur.	bakə	baktən	3 plur.	uœrkjə	uœrkən

(Data from Merkuur 2021, 55)

Two observations can immediately be made. First, verbs from different classes have different paradigms both in the present and in the past tense (e.g., the suffix *-je* shows up in the paradigm of class II but never in class I, conversely, *-te* shows up in class I but never in class II). Second, the paradigms partly overlap in the sense that some affixes (such as *-n* in the past plural, *-t* in 3rd person, and *-st* in 2nd person) are shared by both classes. Ideally, any theory of these paradigms, should account for these observations.

From a theoretical perspective, morphological classes have been taken by some as an argument for the existence of a separate module of the grammar where words are built (e.g., Aronoff 1994). The line of reasoning of these proponents of a separate word-formation component starts from the observation that elements belonging to different morphological classes syntactically behave in the same way. For example, there are no syntactic rules in Frisian that only pertain to class I or class II-verbs. Furthermore, there is no phonological difference between elements of different morphological classes either. That is, we cannot attribute the different paradigms to any phonological properties of the stems (Merkuur 2021, 39). Given this lack of phonological or syntactic reasons to separate members of different morphological classes, there must be some feature that is only relevant to morphology proper that distinguishes the two. Or, put differently, since there is neither a syntactic nor a phonological rule referring to morphological class, such information (coded as ‘class features’) should be confined in its own domain i.e., in the domain of morphology.

However, since the rise of Distributed Morphology (Halle & Marantz 1993, 1994; Marantz 1997), there is a strong trend in modern linguistics to break down any walls that separate morphology from syntax. The attractiveness of such proposals lies in the requirement of having only a single module (or even stronger, a single operation) to account for all structure-building in the grammar, rather than having more places, or different rule-types, that have the power to construct forms (Marantz 1997; Hauser et al. 2002).

In this paper, we will follow this alluring idea, and show that the existence of morphological classes is not an argument for a separate word-building module of the grammar. We will show that class-features (e.g., Marzi et al., 2020) can be eliminated and that there is an alternative available in

¹Although the inflectional paradigms of irregular verbs show more anomalies and less similarities, they also show (traces of) class differences (as historically they developed from regular verbs).



the form of different sizes of lexical items. The theoretical idea of phrasal spell-out is one of the key ingredients of the present proposal. Rather than, as is customary in Distributed Morphology, spelling out terminal nodes of the syntactic tree, we will assume a nanosyntactic framework in which phrases are the target of spell-out (Caha 2009; Starke 2011, 2014a, 2014b, 2014a). Caha, De Clercq & Wyngaerd (2019) proposes that the difference between two classes of adjectives in Czech, is the result of a difference in size between the adjectival bases. Caha (2021) uses the same analytical idea to account for the differences in Russian nominal declensions. Here, we also follow this theoretical idea and show that it can be successfully applied to different morphological classes in a different language, eliminating class-features completely. In this way, we hope that this paper not only contributes to an understanding of the Frisian data, but also to a growing body of literature that argues for phrasal spell-out (Weerman & Evers-Vermeul 2002; Neeleman & Szendrői 2007).

During the paper, we also signal a problem. It turns out that the second person singular in the verbal paradigm of Frisian induces a so-called ABA-pattern. Such patterns have been assumed to be impossible in natural language (see e.g., Bobaljik 2012; Caha 2009), although they do occur. We will show that, indeed, the analysis employed here can not account for this Frisian ABA-pattern. We suggest a solution in the form of an adaptation of the spell-out algorithm, as earlier proposed by Blix (2022), who argues that certain ABA-patterns can be accounted for with the aid of so-called complex left branches. Because the present paper concentrates on accounting for the class difference, we postpone the analysis of the Frisian ABA pattern to future work and leave the second person singular out of the present analysis.

This paper starts out with a brief theoretical introduction in Section 1, explaining the main ingredients for our nanosyntactic analysis. In Section 2, we present the relevant Frisian data for weak verbs and analyze these data using the theoretical tools presented in Section 1. In Section 3 we turn to the strong verbs and show that, although some paradigms require an adapted version of the theory, a possibility that goes beyond the scope of this paper, at least for some verbs, the presented analysis for weak verbs directly accounts for their paradigms. In Section 4 we briefly evaluate and conclude.

2. SOME THEORETICAL BACKGROUND: CLASS AS SIZE DIFFERENCE

Our analysis is embedded in nanosyntax (Starke 2009, 2014; Caha 2009). Here, we focus on the spell-out algorithm, in particular the way that different root sizes may lead to different inflectional paradigms. We refrain from a broader introduction to the framework. We refer the interested reader to Starke (2014a), Caha (2019) and Baunaz & Lander (2018) for such broader introductions to the nanosyntactic framework. We have tried to present our analysis in such a way that also those readers that are not familiar with the framework should be able to understand the core of the analysis.

A central idea of nanosyntax is that the syntactic structure is built up by stepwise merging of features. After each merger, the grammar tries to spell-out the resulting structure with an entry from the lexicon.² Often, there is no lexical entry available that can spell-out the resulting

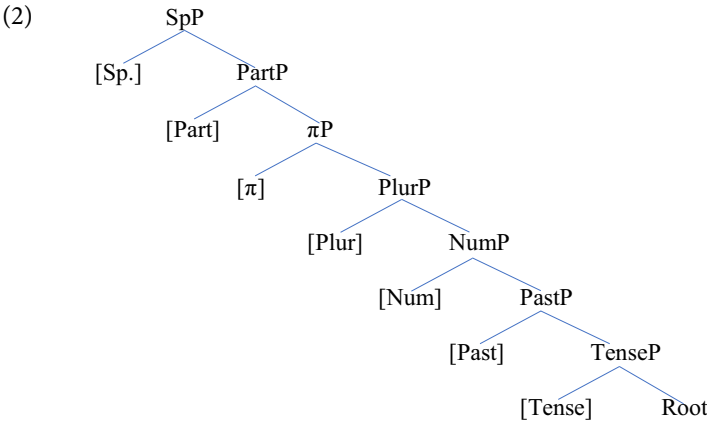
²A caveat is in order here. We do not imply that each merger is to be equated with a phase in the sense of phase-theory (Chomsky 2008). Here we intend that the relevant lexical items are somehow activated and put in memory before being realized.



structure directly. Therefore, in many cases one or more movement operations are necessary to adapt the structure in such a way to allow matching.

Each syntactic feature is a terminal node that heads its own projection. These features come in a specific hierarchical order (often referred to as the functional sequence or, in short, fseq.), which is assumed to be universal. There are more such functional sequences, depending on word-class. These features merge stepwise on top of a root (see the structure in (2)). Of course, not every feature of the hierarchy is always present. The syntactic use of the item determines which features will be present. E.g., in a third person singular present tense, the features [past], [plural] and [speaker] will be absent. However, each feature that is part of the structure will need to be spelled out by some element of the vocabulary (i.e., by some morph in the terminology of Haspelmath 2020).

We assume the functional sequence represented by the structure in (2) (Starke 2020).



In contrast to what is customary in earlier proposals (such as Distributed Morphology), nano-syntax makes use of phrasal spell-out. This means that phrases such as TenseP, PastP are the targets of spell out, rather than the terminal nodes (such as [Tense], [Past], etc.).

We make use of a simplified version of the spell-out algorithm proposed by Starke (2018):³

- (3) Spell-out algorithm (simplified version)
- Merge F (F being a feature on the fseq.) and
 - a. Spell out FP
 - b. If (a) fails, attempt any of the rescue strategies below (in the order given), and retry (a), until spellout is successful:
 - (i) move the spec of the complement of F
 - (ii) move the complement of F.

³It would require too much of the available space here to fully explain the spell-out algorithm and its technical details. We refer the interested reader to Caha (2019, 105 ff.) for a detailed exposition.

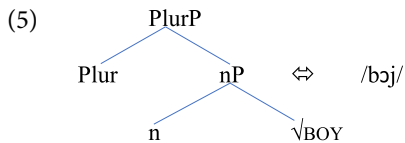


Let us illustrate the workings of this algorithm with a small example in which we build the plural of the English noun *boy*.⁴ As a first step, the root merges with a feature [n] that heads an nP.

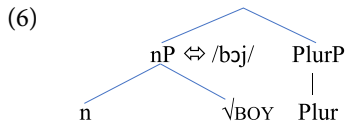


After this merger, clause (a) of the Spell-out algorithm applies to the nP. To fulfill this task, we need to consult the lexicon and see whether there is a lexical item that corresponds to the nP. The English lexicon contains a form /bɔj/ that has exactly the form of the tree in (4). So, spell-out is successful for this tree, indicated by the double-arrow in (4).

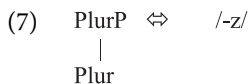
As a second step, we merge the plural feature, leading to the structure in (5):



Now, we try to spell out PlurP, but since the lexicon does not contain a special form for the plural of this noun, the (b)-clause of the Spell-out algorithm applies. Therefore, we first try to move the specifier of the complement of Plur, but since there is no specifier in this complement, clause (b) (ii) applies, and we move the whole complement, resulting in the structure in (6):



Again, we consult the lexicon to see whether there is something stored that may realize the PlurP. One of the items that the lexicon contains is the suffix /-z/, which is stored as a treelet (7):



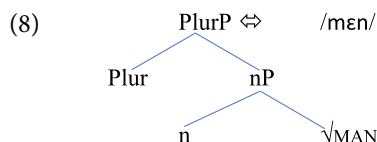
The Spell-out of the PlurP in (6) can be successfully completed since the PlurP is realized as /-z/, and the remaining nP as /bɔj/.

Note that in the case of a root such as $\sqrt{\text{MAN}}$ (/mæn/) (rather than $\sqrt{\text{BOY}}$), there would be a special lexical item available (the ‘suppletive’ form *men*) realizing the corresponding

⁴Note that we simply assume here a nominal fseq. in which plural is the highest head. This does not reflect any claim we make about the structure of the nominal fseq., it merely serves our illustration of the Spell-out algorithm.



structure in (5). For this structure, English has a ‘bigger’ lexical item that not only covers the nP containing this root, but also the PlurP:



There is one more theoretical element of the nanosyntactic approach that we need to explain. Suppose we would have chosen to build the plural of the root $\sqrt{\text{SHEEP}}$. As with $\sqrt{\text{MAN}}$, there is a special form (*sheep*) for the plural (corresponding to the structure in (8)), which is also used for the singular. How can we have this ‘bigger’ special form also realizing the ‘smaller’ singular structure?

Nanosyntactic theory relies on the Superset Principle, proposed by [Starke \(2009\)](#):

(9) Superset Principle ([Starke 2009](#))

A lexically stored tree matches a syntactic node iff the lexically stored tree contains the syntactic node.

That is, if the lexical tree of $\sqrt{\text{SHEEP}}$ contains the ‘singular’ nP, then, according to the Superset Principle, this tree may be used to spell-out the singular form. This, of course, raises a new question: why is it that in the case of $\sqrt{\text{MAN}}$ it is the smaller lexical item (/mæn/) that wins over the bigger one in the singular? The answer comes from a second principle, which is a version of the so-called Elsewhere Condition ([Kiparsky 1973](#), 18):

(10) Elsewhere Condition ([Kiparsky 1973](#))

In case two rules, R1 and R2, can apply in an environment E, R1 takes precedence over R2 if it applies in a proper subset of environments compared to R2.

This condition ensures that the smaller (i.e., more specific) lexical tree will be chosen in case two lexical items compete for insertion.

Before we turn to a detailed analysis of the Frisian verb classes, we are now in the position to explain the gist of our analysis. We will argue that roots from class I are bigger in that they spell out one feature more than the roots from class II (see the paradigms in (1)). As a result, spell-out calls for an extra morph for roots from class II to realize the feature that is ‘covered’ by class I roots. We will demonstrate that this small difference in size and its consequences for spellout will explain the difference between paradigms in (1). Furthermore, as we will demonstrate in due course, the overlap between the two paradigms also follows since a class II root plus the extra morph equals the size of a class I root. Thus, to account for the two different paradigms, we need only to understand the size difference between lexical roots, and we do not need to revert to language-specific morphological features. The class-variation follows from general syntactic principles that are part of UG. In a DM-account (see [Merkuur 2021](#)), by contrast, one would have to rely on language-specific features that refer to the class of the verb, triggering different inflectional affixes. In the nanosyntactic approach, there is no such need since the



class-differences fall out from a universally available mechanism i.e., so-called size differences between lexical items.

In our discussion of the paradigms, we encounter a problem concerning the second person singular forms (ending in *-st*). These forms challenge the ordering on the fseq. of the second person vis-à-vis the first and third person. Consider the singular second person past tense forms in (1). Here, we see that there is a syncretism between first and third person to the exclusion of the second. This is an instantiation of a so-called ABA-pattern given the hierarchy in (2). Because such ABA patterns are assumed to be impossible in natural language (see e.g., Bobaljik 2012), ideally, an elegant analysis would resolve this problem for the Frisian data. On closer inspection, however, it turns out that certain ABA-patterns do occur. Proper treatment of such patterns requires an adaptation of the spell-out algorithm. Here, we will limit ourselves to the elimination of class-features, postponing discussion of the second person singular to future work.

The step-wise merger and application of the spell-out algorithm is a rather tedious and space-consuming way of presenting our analysis. Fortunately, it is possible to translate them to transparent lexicalization tables that we will use to exemplify our analyses.

The first row of the table (apart from the last column that contains the resulting phonological form, added for convenience) gives the fseq., starting from the root in the leftmost cell.

The leftmost cell of the second row contains the syntactic specification of the given form (here: 1st person plural past). The other cells of this row contain the morphs that realize the features from the fseq. with which they are aligned in the table which they spell out. Note that a single affix (e.g., *-n* in Table 1) may realize more than a single feature (due to phrasal spell-out). However, the following rules in the lexicalization tables must be obeyed, since they directly reflect the way the spell-out algorithm works:

- Each affix has a maximal size which corresponds to its lexical specification.
- Each affix has a minimal size which corresponds to the lowest feature in the hierarchy that it may lexicalize (it may ‘shrink from the top’).
- Affixes cannot skip features in the hierarchy that are part of their lexical representation.

The first statement simply states that an affix cannot be used to spell out features that are not part of its lexical specification. Note that this element of the theory contrasts with systems that rely on underspecification of affixes (such as DM, Halle & Marantz 1993).⁵ The second statement is a direct consequence of the Superset Principle (Starke 2009). This implies that not all features contained in a lexical item always must be present in the syntactic representation for the

Table 1. 1st plural past tense of weak class I verb

VerbType	Synt. Repr.	Verb-root	Tense	Past	Num	Plur.	π	Sp.	Phon. form
Class I	1st plur. past	bak		-tə		-n			baktən

⁵We refer the interested reader to Caha (2020) for an extensive comparison between DM and nanosyntax.



item to be inserted. Specifically, lexical items may ‘shrink from the top’ to their minimal size i.e., their foot, but features at the bottom or in the middle of the lexical specification cannot be left out or skipped. We will see the workings of the Superset Principle in discussing the different Frisian forms, to which we will now turn in Section 2.

3. WEAK VERBS

The forms that we use as our data represent what one might call ‘standard Modern Frisian’. This is not the only form spoken in Fryslân (a province in the Northern part of the Netherlands) today; there is quite some variation among speakers. Ignoring this variation for now, we limit the discussion to these standard forms (taken from a recent study by Merkuur 2021).

Frisian verbs belong to either one of two classes: class I or class II. The paradigms of these classes are exemplified in (1), using *wurkje* ‘to work’ and *bakke* ‘to bake’ as our examples.

Let us start by looking at the plural forms in the present tense. The present tense plural forms of the two verb classes differ only in the occurrence of [j] (*bakke* versus *wurkje*). This [j] must be the realization of an extra morpheme that is lacking in the class I verbs.⁶ Since the plural present tense forms are for the rest identical, the logic of our approach only allows the conclusion that this morpheme realizes a feature that is covered by the stem of class I verbs. Hence, class I stems are bigger than class II stems. The lexicalization table below (Table 2) gives our analysis of the plural present tense forms of the regular weak class I and II verbs.

Specifically, we claim that the stem of class I verbs covers the Tense-feature which is not covered by the stem of class II verbs. Since all features need to be realized, an immediate consequence is that another affix must be recruited to spell out Tense in class II verbs. This affix is the suffix that surfaces as [j]. Thus, we propose that the lexical representation of the stems of *bak* (class I) and *wurk* (class II) is as in (11):

- (11) a.

TP ⇔ /bak/
├── T
└── √BAK
- b. √WURK ⇔ /værk/

Table 2. Plural present tense of weak class I and class II verbs

VerbType	Synt. Repr.	Verb-root	Tense	Past	Num	Plur.	π	Part.	Sp.	Phon.form
Class I	3 rd plur. pres.	bak				-ə				bakə
	2 nd plur. pres.	bak				-ə				bakə
	1 st plur. pres.	bak				-ə				bakə
Class II	3 rd plur. pres.	værk	-j			-ə				værkjə
	2 nd plur. pres.	værk	-j			-ə				værkjə
	1 st plur. pres.	værk	-j			-ə				værkjə

⁶Note that this [j] cannot be part of the root since it occurs in all class II verbs.



The syncretism between the three person-forms in the plural is easily accounted for by assuming that the suffix [ə] is not only specified for plural but also for all person features. Its lexical representation is given in the first line of (12).⁷ Here we also give our first approach towards the representation of the suffix -j.

- (12) *Frisian lexicon of verbal suffixes (first version)*
- ə

:

Num – Plur. – π – Part. – Sp.
- j

:

Tense

The superset principle ensures that the ‘plural’ suffix not only realizes the first-person plural (its ‘largest’ coverage) but may also be recruited (in case more specific affixes are lacking) to realize second and third-person plural.

As the next step in our analysis, let us turn to the past tense forms before we come back to the singular present tense forms. Consider Table 3 below. First, the stems *bak* (class I) and *wurk* (class II) are fixed in size, *bak* covering one more feature than *wurk*. Second, the morpheme *-n* should have the same size in both paradigms, with the same foot. This leaves the morphemes *-ə* and *-tə* to be of different sizes. The morpheme that surfaces as *-ə* has its foot in Tense because that is where the class II stems end. Consequently, this morpheme must be the same morpheme as [j] in Table 2. Because there is only a single morpheme between *wurk* and *-n*, and there is already a morpheme that starts at Tense, it must be the same.

In fact, it makes a lot of sense that the morpheme that surfaces as *-ə* here must be identified with the morpheme *-j* in Table 2. Following Merkuur (2021), we assume that, underlyingly, this affix consists of the abstract vowel [I], which reduces to schwa in stressless position, whereas it shows up as a glide [j] elsewhere. We describe the phonology of the suffix informally with the following rules:

Table 3. 3rd plural past tense of weak class I and class II verbs

VerbType	Synt. Repr.	Verb-root	Tense	Past	Num	Plur.	π	Part.	Sp.	Phon.form
Class I	3 rd plur. past	bak		-tə		-n				baktən
	2 nd plur. past	bak		-tə		-n				baktən
	1 st plur. past	bak		-tə		-n				baktən
Class II	3 rd plur. past	ʋærk		-I		-n				ʋærkən
	2 nd plur. past	ʋærk		-I		-n				ʋærkən
	1 st plur. past	ʋærk		-I		-n				ʋærkən

⁷We use linear notations here rather than trees for reasons of space. We assume that the reader may easily translate these linear notations to treelets.



- (13) a. /I/ is a vowel before C, a glide elsewhere:

/I/	→	i	/____C
/I/	→	j	/ Elsewhere

- b. Reduction:

/i/	→	ə
		[–stress]

Interestingly, in some versions of presentday Frisian, the [I] is realized as [i], rather than in its reduced version.⁸ We take this as further evidence for the correctness of our analysis.

Consequently, the suffix /I/ is larger than we thought before. It also realizes the features Past and Num. One may ask at this point of our analysis whether the suffix *-n* could have a different foot. The answer is negative, since it cannot foot in Num because *-ə* does, and there cannot be identical representations for different morphemes. Furthermore, the morpheme *-ə* cannot have a foot higher in the structure, because then there are no features for *-te* to realize.

Both in class I and in class II, the forms for 1st, 2nd, and 3rd person are syncretic in the plural of the past tense. As we have seen above, the way that nanosyntax deals with these syncretisms is by assuming that the syncretic affix is lexically specified as containing the full hierarchy of features. Because of the superset principle, it may ‘shrink’ from the top, and therefore, those syntactic structures that contain a subconstituent of the lexical structure, may also be realized by this affix. Apparently, the suffix *-n* spells out any phrase containing plural and one or more person features. However, it does not have to realize all its person-features. It may ‘shrink’ from the top. Therefore, it can do first, second and third person (in the absence of more specific affixes).

We update our lexicon of affixes in (14):

- (14) Frisian lexicon of verbal suffixes (*second version*):

<i>-ə</i>	:	Num – Plur. – π – Part. – Sp.
<i>-I</i>	:	Tense – Past – Num
<i>-te</i>	:	Past – Num
<i>-n</i>	:	Plur. – π – Part. – Sp.

Let us now turn to the singular forms, again starting with the past tense, leaving out the second person forms. Our analysis of these forms is summarized in Table 4. We use three dots to indicate that these features are realized by the suffix in a column further to the right. E.g., in the first row the suffix *-te* realizes the features Past-Num- π . Similarly, in the third row *-I* realizes the features Tense-Past-Num- π -Sp.

These tables show that the only step we have to make to include these forms in our analysis is extending the lexical representation of both *-te* and *-I*. If we include the person features in their representations, the past tense forms in Table 4 immediately follow. So, again updating our lexicon, gives us the affixes in (15):

⁸The first author of this paper is a speaker of such Frisian dialect.



Table 4. Past tense forms of weak class I and class II verbs

	Verb-root	Tense	Past	Num	Plur.	π	Sp.	Phon.form
3 rd sing. past	bak		...			-tə		baktə
1 st sing. past	bak		...			-tə		baktə
3 rd sing. past	værk		...			-I		værkə
1 st sing. past	værk		...			-I		værkə

- (15) Frisian lexicon of verbal suffixes (*third version*):
- ə

:

Num – Plur. – π – Part. – Sp.
- I

:

Tense – Past – Num – π – Sp.
- tə

:

Past – Num – π – Sp.
- n

:

Plur. – π – Part. – Sp.

Finally, let us turn to the singular forms of the present tense (again leaving out second person forms). As in the plural (Table 2), we expect the same suffixes to occur in class I and class II, since the stem of class II equals the stem + I of class I. The first person is zero-marked in class I verbs. This could be taken as an indication that the verb-root realizes all person and number features. However, if that were the case, we would expect a syncretism with the second and third person, which is clearly not what we find. Therefore, for now we are forced to postulate a zero-affix that realizes the number and person features in the present first person singular.⁹

In the third person, we encounter a new suffix *-t* that only covers Num and π . The analysis, asking for two new suffixes, *-t* and *-0*, is straightforward, as presented in Table 5.

This leads to a final update our lexicon, adding the two present tense suffixes *-0* and *-t*:

Table 5. Singular present tense forms of weak class I and class II verbs

	Verb-root	Tense	Past	Num	Plur.	π	Part.	Sp.	Phon.form
3 rd sing. pres.	bak			...		-t			bakt
1 st sing. pres.	bak			...			-0		bak
3 rd sing. pres.	værk	-I		...		-t			værkət
1 st sing. pres.	værk	-I		...			-0		værkjə

⁹We are fully aware that postulating a zero-affix is an analytical last resort. We are therefore committed to find another solution in which the stem that covers the root-features also shows up in the first person singular of the present tense. However, since we focus on the elimination of class-features here, for which this discussion is not relevant, and since this analytical step would require so-called complex left branches (Blix 2022), we leave this issue here without further discussion.



(16) Frisian lexicon of verbal suffixes (*final version*):

-ə	:	Num – Plur. – π – Sp.
-I	:	Tense – Past – Num – π – Sp.
-tə	:	Past – Num – π – Sp.
-n	:	Plur. – π – Part – Sp.
-0	:	Num – π – Sp.
-t	:	Num – π

Before we move on, let us briefly zoom in on the 1st singular present tense form of class II verbs (værkjə). The suffix *-I* is realized as a palatal glide here since it does not precede a consonant. This renders the phonological form /wærkj/ but since that is not a viable phonological word in Frisian, we assume that the phonology inserts a final schwa as a kind of repair, leading to the surface form [wærkjə] (see for similar analyses: Visser 1992; Merkuur 2021). This epenthetic segment is therefore not part of the morphology of Frisian.

At this point of our analysis, one may appreciate multiple reasons for which we consider ‘size’ a better explanation than class features for the paradigmatic differences between class I and class II verbs. Firstly, the current analysis makes predictions as to where the paradigms differ and where they are the same. The morpheme /I/ sometimes, together with the verbal stem (class II), realizes the same features as the larger class I stem. In such cases, we see the same morphemes showing up in both classes higher up in the tree (more to the right). If the morpheme *-I* realizes more features, then we see different morphemes higher up in both classes. As far as we are aware, an account using class features cannot make such specific empirical predictions as it can be used to label any difference.

Secondly, as indicated above, using size rather than featural make-up to code the class-difference, we obviate the need for language-specific features and make use of a universally available mechanism to code this lexical difference.

Thirdly, another theoretical advantage of the current approach is that we may use the same technique to account for the differences between the paradigms in the weak verbs, and those we see in the strong verbs. That is, a theory relying on features to distinguish between verb classes, would have to invoke two features: [$\pm$ strong] and [I/II class]. In our proposal only size differences of lexical items are used to derive the different verbal paradigms. The central point of our section 3 is that a substantial subset of the strong verbs can be analyzed within our framework with no additional theoretical apparatus.

4. ‘STRONG’ VERBS

Next to the weak class I and class II verbs discussed above, Frisian hosts strong verbs that, although they seem less regular at first sight, can also be classified as either class I or class II. By way of example, we give the paradigms of *wurde* (‘to become’) (class I) and *hingje* (‘to hang’) (class II). Next to class I and class II weak and strong verbs, there is also a set of truly irregular verbs. We will refrain from discussing those latter verbs, although in principle we see no problem for the analytical tools used to capture these too. We leave this for future research.

The strong verbs are characterized by having multiple stems i.e., *wurd*, *waard* for the verb ‘to become’, and *hing*, *hong* for the verb ‘to hang’ (Table 6). In the past tense, in contrast to the weak



Table 6. Present tense of strong class I verb

	Verb-root	Tense	Past	Num	Plur.	π	Part.	Sp.	Phon.form
3 rd plur. pres.	vurd				-e				vurdə
2 nd plur. pres.	vurd				-e				vurdə
1 st plur. pres.	vurd					-e			vurdə
3 rd sing. pres.	vurd			...		-t			vurt
1 st sing. pres.	vurd			...			0		vurt

verbs, they do not show the suffixes, *-I* or *-te*. The situation in Frisian is not dissimilar from what we observe in English strong verbs that also have multiple stems. Below, we investigate to what extent the same theoretical notion ‘size difference’ that successfully may replace class features, may also, replace the notions ‘strong’ and ‘weak’.

More specifically, the general idea of our analysis is that the stems that only occur in the past tense (such as *waard* ‘became’), realize the feature [past].

(17)	class I	Pres.	Past	class II	Pres.	Past
	1 sing.	vurd	vard	1 sg.	hɪŋjə	hɔŋ
	2 sing.	vurdst	vardst	2 sg.	hɪŋəst	hɔŋst
	3 sing.	vurdt	vard	3 sg.	hɪŋət	hɔŋ
	1 plur.	vurdə	vardŋ	1 plur.	hɪŋjə	hɔŋŋ
	2 plur.	vurdə	vardŋ	2 plur.	hɪŋjə	hɔŋŋ
	3 plur.	vurdə	vardŋ	3 plur.	hɪŋjə	hɔŋŋ

Let us first have a look at a strong verb from class I, the verb *wurde* ‘to become’. The present tense forms of this verb, as is true for all strong verbs of this class, are identical to the paradigm found in the weak class I verbs. We present the relevant lexicalization table for the present tense of *wurde* below. The table shows that *wurde* behaves parallel to *bakke* in the present tense.

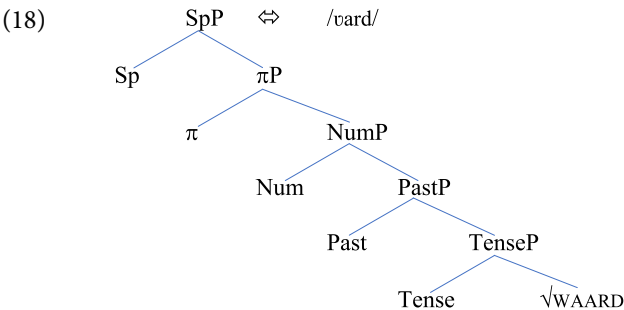
The past tense, however, clearly differs. The absence of the *-te* suffix in the past tense in strong verbs suggests that the stem *ward* covers the features that are covered by stem + *te* in the weak verbs. Table 7 makes this idea concrete.

Table 7. Past tense forms of strong class I verb

	Verb-root	Tense	Past	Num	Plur.	π	Part.	Sp.	Phon.form
3 rd plur. past			vard		-n				vardŋ
2 nd plur. past			vard		-n				vardŋ
1 st plur. past			vard			-n			vardŋ
3 rd sing. past			...			vard			vard
1 st sing. past			...				vard		vard



The whole past tense paradigm results from this step without further ado. The paradigm follows from the lexical representation of the affixes in (16) and the assumption that the deviant stem has the lexical representation, given in (18) below.



In principle, the stem *waard* /vard/ may be used to realize the present tense form as well. However, the Elsewhere Condition will select the ‘smaller’ form *wurd* since it has less superfluous features.

Turning to class II strong verbs, with the verb *hingje* as an example, the picture is slightly different: the present tense behaves as in the regular cases. The verb *hingje* is class II and behaves just like *wurkje* in the present tense.

Now, if we turn to the past tense and take the same analytical step as we did for the strong verbs of Class I, we encounter a serious problem.

Comparison of Table 9 with Table 8 visualizes that point:

Table 8. Present tense of strong class II verb

	Verb-root	Tense	Past	Num	Plur.	π	Sp.	Phon.form
3 rd plur. pres.	hing	-I			-e			hmjə
1 st plur. pres.	hing	-I			-e			hmjə
3 rd sing. pres.	hing	-I		...		-t		hmjət
1 st sing. pres.	hing	-I		...		0		hmjə

Table 9. Past tense forms of strong class II verb

	Verb-root	Tense	Past	Num	Plur.	π	Part.	Sp.	Phon.form
3 rd plur. past			həŋ		-n				həŋŋ
2 nd plur. past			həŋ		-n				həŋŋ
1 st plur. past			həŋ			-n			həŋŋ
3 rd sing. past		...				həŋ			həŋ
1 st sing. past		...					həŋ		həŋ



At first sight, this step nicely derives the past tense forms without any additional assumptions. However, in this case, the form *hong* would be chosen over the form *hing+I* in the present tense. The reason is that *hing+I* is bimorphemic (rather than the monomorphemic *wurd* in Class I strong verbs) and therefore, the spell-out algorithm used in this paper, chooses *hong* over *hing+I* in the present tense. The structure containing the root plus the Tense-feature is a subtree of the lexical structure of *hong*. Therefore, without any movement operations, this stem will be able to spell out that structure, incorrectly predicting that *hong* is used in the present tense. We would need to make use of complex left branches (Blix 2022) to be able to deal with these cases, a step that we do not want to take yet, since we cannot oversee all consequences of such a farreaching move. We leave the Class II strong verbs as a problem for now.

5. EVALUATION AND CONCLUSION

We have argued that in Frisian verbal stems from class I spell out one feature more than verbal stems from class II. As a result, the stems from class II call for an extra affix to realize this feature. This small difference in size between the roots explains further differences in the paradigms of class I and II verbs. Furthermore, the overlap between the paradigms is expected within this approach since a class II root plus the extra affix equals the size of a class I root. In this way, we do not need to revert to language-specific morphological features. The class-variation follows from general syntactic principles that are part of UG. In the nanosyntactic approach, there is no need for special class-features, but morphological classes fall out from a universally available mechanism i.e., so-called size differences between lexical items.

Furthermore, the same technique of size-differences allows us to bring previously considered irregularities (such as the strong verbs) within the same system without any *ad hoc* devices or special rules.

The key ingredient of the nanosyntactic analysis presented here is the possibility for one and the same morpheme to realize a stretch of the fseq. rather than just a single feature, or a feature-bundle (i.e., phrasal spell-out). Interestingly, the nanosyntactic approach puts a severe constraint on possible syncretisms: features involved should be (structurally) adjacent on the fseq. (see Caha 2009 for extensive discussion of case-syncretisms). Since morphemes realize a stretch of features, they can differ in size. In line with earlier proposals, such as Caha, De Clercq & Wyngaerd (2019), we have shown that these size differences are the key to a better understanding of morphological verb classes in Modern Frisian. Making class-features follow from size differences, mitigates an important argument for separating morphology from syntax. If the defended approach can be generalized to other such features, there is all the more reason to try to unify these grammatical components.

Finally, we observe that Frisian also challenges the ordering of the second person vis-à-vis the first and third person. The observed syncretism between first and third person to the exclusion of the second, seems to be an instantiation of an ABA-pattern that the theory excludes. In future work, we aim to show that a solution proposed by Blix (2022) proposing complex left branches plus an adaptation of the spell-out algorithm, is needed to deal with these ABA-patterns. Furthermore, these future analyses should provide an account for all strong and irregular verbs as well.



ACKNOWLEDGEMENT

The authors would like to thank the audience of the International Morphology Meeting (Budapest, Sept 2, 2022) for their comments and discussion.

REFERENCES

- Aronoff, Mark. 1994. *Morphology by itself*. Cambridge, MA: MIT Press.
- Baunaz, L., Lander, E. 2018. Nanosyntax: The basics. In: Baunaz, L., Haegeman, L., De Clercq, K., Lander, E. (Eds.), *Exploring nanosyntax*. Oxford University Press, New York, NY, pp. 3–56. <https://doi.org/10.1093/oso/9780190876746.003.0001>.
- Blix, H. 2022. Interface legibility and nominal classification: A nanosyntactic account of Kipsigis singulatives. *Glossa: A Journal of General Linguistics* 7. <https://doi.org/10.16995/glossa.5825>.
- Bobaljik, Jonathan. 2012. *Universals in comparative morphology*. Cambridge, MA: MIT Press.
- Caha, Pavel. 2009. *The nanosyntax of case*. Doctoral dissertation. Tromsø University, Tromsø. <https://munin.uit.no/handle/10037/2203>.
- Caha, P. 2019. In: *Case competition in Nanosyntax: A study of numeral phrases in Ossetic and Russian*. Language Science Press, Berlin. <https://docplayer.net/195577209-Case-competition-in-nanosyntax.html>.
- Caha, Pavel. 2020. Nanosyntax: Some key features. Manuscript. Masaryk University, Brno. <https://ling.auf.net/lingbuzz/004437>.
- Caha, Pavel. 2021. Modeling declensions without declension features: The case of Russian. *Alternatives to Laboratory Animals: ATLA* 68. 385–425. <https://doi.org/10.1556/2062.2021.00433>.
- Caha, P., De Clercq, K., Wyngaerd, G.V. 2019. The fine structure of the comparative. *Studia Linguistica* 73. 470–521. <https://doi.org/10.1111/stul.12107>.
- Chomsky, Noam. 2008. On phases. *Current Studies in Linguistics Series* 45. 133.
- Dyk, Siebren. 2020. Verbs. *Taalportaal*. Retrieved from <https://taalportaal.org/taalportaal/topic/pid/topic-13998813311890721> (accessed 27 October 2022).
- Halle, Morris and Alec Marantz. 1993. Morphology and the pieces of inflection. In K. Hale and S.J. Keyser (eds.) *The view from Building 20: Essays in linguistics in honor of Sylvain Bromberger*. Cambridge, MA: MIT Press. 111–176.
- Halle, Morris and Alec Marantz. 1994. Some key features of distributed morphology. In A. Carnie and H. Harley (eds.) *Papers on phonology and morphology*. (MITWPL 21). Cambridge, MA: MIT. 275–288.
- Haspelmath, M. 2020. The morph as a minimal linguistic form. *Morphology* 30. 117–134. <https://doi.org/10.1007/s11525-020-09355-5>.
- Hauser, Marc D., Noam Chomsky and W. Tecumseh Fitch. 2002. The faculty of language: What is it, who has it, and how did it evolve? *Science* 298. 1569–1579.
- Kiparsky, Paul. 1973. “Elsewhere” in phonology. In: Anderson, S., Kiparsky, P. (Eds.), *A Festschrift for Morris Halle*. Holt, Rinehart and Winston, New York, pp. 93–106.
- Marantz, Alec. 1997. No escape from syntax: Don’t try morphological analysis in the privacy of your own lexicon. In A. Dimitriadis, L. Siegel, C. Surek-Clark and A. Williams (eds.) *Proceedings of the 21st*



- Annual Penn Linguistics Colloquium: Penn Working Papers in Linguistics 4. Philadelphia, PA: Penn Graduate Linguistics Society, University of Pennsylvania. 201–225.
- Marzi, Claudia, James P. Blevins, Geert Booij and Vito Pirrelli. 2020. Inflection at the morphology-syntax interface. In V. Pirrelli, I. Plag and W.U. Dressler (eds.) *Word knowledge and word usage: A cross-disciplinary guide to the mental lexicon*. Berlin & Boston, MA: De Gruyter Mouton. 228–294.
- Merkuur, A. 2021. In: *Changes in modern Frisian verbal inflection*. LOT Publications, Amsterdam. <https://www.lotpublications.nl/changes-in-modern-frisian-verbal-inflection>.
- Neeleman, Ad and Kriszta Szendrői. 2007. Radical pro-drop and the morphology of pronouns. *Linguistic Inquiry* 38. 671–714.
- Starke, Michal. 2009. Nanosyntax: A short primer to a new approach to language. *Nordlyd* 36. 1–6.
- Starke, Michal. 2014a. Towards elegant parameters: Language variation reduces to the size of lexically stored trees. In C. Picallo (ed.) *Linguistic variation in the minimalist framework*. Oxford: Oxford University Press. 140–152.
- Starke, Michal. 2014b. Cleaning up the lexicon. *Linguistic Analysis* 39. 245–256.
- Starke, Michal. 2018. Complex left branches, spellout, and prefixes. In: L. Baunaz, K. De Clercq, L. Haegeman and E. Lander (eds.) *Exploring nanosyntax*. Oxford: Oxford University Press. 239–249.
- Starke, Michal. 2020. UM – Universal Morphology. Talk at NELS 51, UQAM, Montréal, November 6, 2020.
- Visser, Willem, 1992. Oer-je en-JE. *De morfology en fonology fan it einichste wurddiel-je*. *Tydskrift foar Freyske Taalkunde* 7.
- Weerman, Fred and Jacqueline Evers-Vermeul. 2002. Pronouns and case. *Lingua* 112. 301–338.

