

A word-based account of comprehension and production of Kinyarwanda nouns in the Discriminative Lexicon

Ruben van de Vijver und Emmanuel Uwambayinema

Article - Version of Record



Suggested Citation:

van de Vijver, R., & Uwambayinema, E. (2022). A word-based account of comprehension and production of Kinyarwanda nouns in the Discriminative Lexicon. *Linguistics Vanguard*, 8(1), 197–207.
<https://doi.org/10.1515/lingvan-2021-0160>

Wissen, wo das Wissen ist.



UNIVERSITÄTS- UND
LANDESBIBLIOTHEK
DÜSSELDORF

This version is available at:

URN: <https://nbn-resolving.org/urn:nbn:de:hbz:061-20251124-103841-4>

Terms of Use:

This work is licensed under the Creative Commons Attribution 4.0 International License.

For more information see: <https://creativecommons.org/licenses/by/4.0>

Ruben van de Vijver* and Emmanuel Uwambayinema

A word-based account of comprehension and production of Kinyarwanda nouns in the Discriminative Lexicon

<https://doi.org/10.1515/lingvan-2021-0160>

Received December 23, 2021; accepted August 2, 2022; published online December 8, 2022

Abstract: Are the cognitive units in the mental lexicon of Bantu speakers words or morphemes? The very small experimental literature addressing this question suggests that the answer is morphemes, but a closer look at the results shows that this answer is premature. A novel theory of the mental lexicon, the Discriminative Lexicon, which incorporates a word-based view of the mental lexicon, and is computationally implemented in the Linear Discriminative Learner (LDL) is put to the test with a data set of 11,180 Kinyarwanda nouns, and LDL is used to model their comprehension and production. LDL predicts comprehension and production of nouns with great accuracy. Our work provides support for the conclusion that the cognitive units in the mental lexicon of Kinyarwanda speakers are words.

Keywords: Bantu; discriminative lexicon; Kinyarwanda; mental lexicon; word-based morphology

1 Introduction

Bantu languages have complex gender systems (Güldemann and Fiedler 2021; Hyman et al. 2019; Katamba 2003) in which each noun is marked by a class marker. The nouns in each class are hypothesized to share a semantic property (e.g. “human being” or “animate”) or a grammatical function (e.g. “plural” or “diminutive”). For example, in Kinyarwanda (classified as J60 (Nurse and Philippson 2006)), which is spoken in Rwanda, Eastern Congo and Southern Uganda, the word *umuntu*, meaning ‘man’, is a noun of class 1 and *abantu* is its plural which is a class 2 noun. Noun classes in Bantu have been studied extensively from a historical and typological perspective (Güldemann and Fiedler 2021; Hyman et al. 2019; Katamba 2003; van der Wal 2015), but very few studies have addressed the question how Bantu nouns are represented in the mental lexicon (Ciaccio et al. 2020; Kgolo and Eisenbeiss 2015). Yet, the highly inflectional nature of Bantu languages (Nurse and Philippson 2006) can shed light on an important theoretical question concerning the mental lexicon: are the cognitive units in the mental lexicon words (Baayen et al. 2018, 2019; Blevins 2006, 2016a) or morphemes (Ciaccio et al. 2020; Goldsmith and Mpiranya 2018; Kgolo and Eisenbeiss 2015)?

We address the question of the cognitive units in the mental lexicon by computationally modeling comprehension and production of Kinyarwanda nouns. The highly inflectional nature of Bantu languages is well-suited to investigate this question. This is because such highly inflectional languages most closely adhere to the so-called morphemic ideal, according to which complex words are composed of unique and easily identifiable morphemes (Ainsworth 2019). Among Bantu languages, Kinyarwanda has a rather complex set of noun classes, because most noun classes are preceded by an extra vowel, often called the pre-prefix, with an ill-understood function (Rosendal 2006).

Our work is situated within the framework of the Discriminative Lexicon (Baayen et al. 2018, 2019), which espouses a word-based theory of morphology (Blevins 2016b). In the Discriminative Lexicon word forms are hypothesized to discriminate among meanings, and meanings discriminate among word forms. This theory is

*Corresponding author: Ruben van de Vijver, Institut für Linguistik und Information, Heinrich-Heine-Universität, Düsseldorf, Germany, E-mail: ruben.vijver@hhu.de. <https://orcid.org/0000-0003-0761-4649>

Emmanuel Uwambayinema, Institut für Linguistik und Information, Heinrich-Heine-Universität, Düsseldorf, Germany, E-mail: emuwa100@hhu.de

implemented computationally as a fully connected network with linear mappings (Baayen et al. 2018, 2019). To foreshadow our results, we can model comprehension and production of Kinyarwanda nouns well by only providing the model with information about word forms and their meaning, but without information about morphemes.

1.1 Experimental work on the mental lexicon in Bantu languages

Despite the fact that there are about 240 million Bantu speakers (Nurse and Philippson 2006), we found only two experimental studies that address the structure of the mental lexicon in Bantu languages. Ciaccio et al. (2020) and Kgolo and Eisenbeiss (2015) conducted masked visual priming experiments on the Bantu language Setswana.

Ciaccio et al. investigated whether there are priming effects for inflected prefixed words, such as *dikgeleke* ‘experts’ and *kgeleke* ‘expert’, and derived prefixed words, such as *bokgeleke* ‘talent’ and *kgeleke* ‘expert’, and for inflected suffixed words, such as *supile* ‘showed’ and *supa* ‘to show’, and derived suffixed words, such as *supega* ‘proven’ and *supa* ‘to show’. Ciaccio et al. (2020) couched their experiment in theories that explain visual masked priming effects on the basis of morphological decomposition (Grainger and Beyersmann 2017; Rastle and Davis 2008; Stockall and Marantz 2006).

The results showed a faster reaction time when prime and target were related through prefixation, but not when prime and target were related through suffixation. Ciaccio et al. (2020) conclude that these results are in agreement with morphological decomposition theories.

Two aspects of this interpretation are suprising though. The first is that if morphological decomposition is a universal mechanism, as Ciaccio et al. (2020) assert, the process should apply to both prefixes and suffixes. This is not the case. To explain this discrepancy the authors point out that many Setswana speakers are unfamiliar with written Setswana. However, it is unclear by which mechanism familiarity with orthography asymmetrically affects morphological decomposition.

The second is that Ciaccio et al. (2020) had to discard 36 of the 85 participants of the study (42.3%), because it was not clear whether they had understood the task. The excluded participants did not reach a 60% threshold of correct answers in the lexical decisions. As Ciaccio et al. (2020) write, this could be a consequence of many Setswana speakers not being used to reading Setswana, but it is not clear whether this applied to the excluded participants. And if it does apply to excluded participants, it means that the remaining participants had good reading skills, the acquisition of which also involves acquiring meta-linguistic knowledge (Dong et al. 2020), which may have affected their ability to isolate morphemes.

The second study addressing the structure of the Bantu mental lexicon is the one of Kgolo and Eisenbeiss (2015) which deals with deverbal nouns in Setswana. These are nouns that are derived from verbal roots by addition of a nominal prefix. They conducted two sets of visual masked priming experiments. One set contained prime target pairs in which the verb was related to a class 1 noun, for example *moroki* ‘tailor’ and the verb *roka* ‘to sew’. In another set the verb was related to a class 9 noun for example *mphe* ‘a gift’ and the verb *fa* ‘to give’. Class 1 nouns are morphologically more transparently related to their verbs than class 9 nouns.

Kgolo and Eisenbeiss expected either priming effects for class 1 and class 9 nouns of comparable magnitude, or, if priming is the result of semantic or formal overlap (in the sense of shared letters), that there should be less priming for class 9 nouns than for class 1 nouns. The results, however, corroborate neither of these expectations: they reported a stronger priming effect for class 9 nouns.

These results, too, are puzzling with respect to morphological decomposition. If it is a universal mechanism, why does it not apply across-the-board and why does it appear to affect morphologically transparent words less than morphologically nontransparent words?

Even though there are no experimental or computational studies yet that provide arguments in favor of a word-based view of the Bantu mental lexicon there are some considerations that favor such an account. One concerns the difficulty of identifying morphemes. Children acquiring Bantu never hear individual morphemes, so they have to isolate them by some mechanism. This, however, is not always possible, even in Bantu languages as Katamba (1978) shows. And even if we assume that this problem can be overcome, there is the conundrum that a child certainly sets out her presumed quest for morphemes by first storing whole words in her lexicon over which she may then generalize. This raises the question as to what happens to these stored

words once the morphemes are identified (Ambridge 2020; Baayen and Ramscar 2019)? From other languages, there is evidence that complex words are in fact retained in memory intact (Mitterer and Reinisch 2017; Moscoso del Prado Martin et al. 2004), which would make an analysis in terms of morphemes redundant. In short, it is worthwhile to investigate whether modeling comprehension and production of Kinyarwanda nouns is possible, if the model is only provided with information about whole words and their meanings.

1.2 The present study

Experimental evidence to support a morphological decomposition of nouns in Bantu is inconclusive. Moreover, there are some arguments to support a word-based view of the mental lexicon even for highly inflectional languages. We therefore set out to test the word-based view of the mental lexicon (Blevins 2016b), and in particular, we pursue the hypothesis of the Discriminative Lexicon (Baayen et al. 2018, 2019; Chuang et al. 2020) that comprehension is based on a linear mapping of the phonology of words onto their meaning and production is based on a linear mapping of the meaning of words onto their phonology. The Discriminative Lexicon theory has been computationally implemented as the Linear Discriminative Learner (LDL), a fully connected network of two layers, one for word form and one for meaning (Baayen et al. 2018, 2019; Chuang et al. 2020).

We use Kinyarwanda, of which the nominal morphology is to a large extent comparable to Setswana, except for the extra complication that Kinyarwanda's noun class markers are preceded by an additional pre-prefix with an ill-understood function (Rosendal 2006).

We relied on computational modeling since this allows us to consider nouns from all classes; as a result of the sheer number of words to be tested, an experiment would become prohibitively large. We will next introduce Kinyarwanda noun classes, and our data set, followed by an introduction to LDL. The results of the modeling are presented in Sections 4 and 5 concludes the paper.

2 Kinyarwanda noun classes

Rosendal (2006) distinguishes 16 noun classes, which are indicated by roman numerals following the tradition in Bantu linguistics. Examples of each noun class are given in Table 1. The class of a noun determines its

Table 1: Kinyarwanda noun classes (Rosendal 2006).

Class	Phonology	Semantics	Example	Gloss
1	umu	human beings	umuntu	man
2	aba	plural of class 1	abantu	men
3	umu	mass nouns, inanimates, animals	umusozi	mountain
4	imi	plurals of 3	imisozi	mountains
5	i(ri)	body parts, loan words	izuru	nose
6	ama	plurals of 5	amazuru	noses
7	iki	body parts (sing), animals, inanimates, loanwords	ikiganza	hand
8	ibi	plural of 7	ibiganza	hands
9	in	animals, inanimates	inzoka	snake, worm
10(a)	in	plurals of 9, 11	inzoka	snakes, worms
10(b)	in	plurals of 9, 11	inkingo	vaccines
11	uru	singular of class 10	urukingo	vaccine
12	aka	abstract nouns, inanimate nouns	akabago	period, stop
13	utu	plurals of 12	utubago	periods, stops
14	ubu	abstract pluralless nouns	ubutaka	earth
15	uku	paired body parts in singular	ukuboko	arm
16	aha	locations	ahantu	place, places

agreement pattern in a phrase (Katamba 2003). In Kinyarwanda, noun classes are usually preceded by a pre-prefix consisting of a single vowel (this vowel is not present in some contexts, for example after demonstratives). The function of the pre-prefix is unclear in Kinyarwanda (Rosendal 2006), even though it may have a number of functions in other Bantu languages (Katamba 2003). The locative meaning ‘on’ is expressed by a prefix *k* that precedes a noun class marker and its pre-prefix, and the meaning ‘in’ by the prefix *m*.

2.1 Kinyarwanda data set

We manually created a data set consisting of 11,180 inflected word forms of 1,493 different nouns, which were annotated for *lexeme*, and the grammatical functions *noun class*, *number*, *diminutive* and *locative* (Table 2). The word forms were written in Kinyarwanda orthography, to which we added information about vowel length (by adding a vowel symbol) and tones (by giving vowels with a high tone an acute accent). As all syllables in Kinyarwanda end in a vowel (Kimenyi 1979), we indicated syllable boundaries by adding a period after every short and long vowel.

Table 2: Examples from our data set for the word glossed as ‘ancestor’.

Orthography	Prosody	Gloss	Class	Number	Locative	Diminutive
umukurambere	umukúraambere	ancestor	1	sg	–	–
agakurambere	agakúraambere	ancestor	12	sg	–	dim
kumukurambere	kumukúraambere	ancestor	1	sg	on	–
mumukurambere	mumukúraambere	amcestor	1	sg	in	–
abakurambere	abakúraambere	ancestor	2	pl	–	–
udukurambere	udukúraambere	ancestor	13	pl	–	dim
kubakurambere	kubakúraambere	ancestor	2	pl	on	–
mubakurambere	mubakúraambere	ancestor	2	pl	in	–

The data set contains several homonyms. For example the word *uturence* means ‘foot’ and ‘sector’, with otherwise identical specifications for grammatical functions. There are 165 homonyms in the data set. Homonyms are common in any language, and may be distinguished on the basis of different phonetic details (Gahl 2008; Lohmann 2018), but such details are not available for our data set. These homonyms will have consequences for the way in which we assess the accuracy of our modeling. We will address these consequences in Section 3.

On the basis of our data set, we further created a data set in which the meanings are based on word embeddings. Word embeddings are representations of word meanings on the basis of the distribution of words in a corpus (Landauer and Dumais 1997). The idea behind this way of representing meanings is that words that occur in similar contexts tend to have similar meanings. The word embeddings for Kinyarwanda are described in detail in Niyongabo et al. (2020). We created this data set by selecting all words in our data set for which word embeddings are available. This was the case for 1,732 word forms.

3 Linear Discriminative Learning

Linear Discriminative Learning (LDL) is a computational implementation of the Discriminative Lexicon theory (Baayen et al. 2018, 2019; Chuang et al. 2020).¹ Comprehension and production are modeled by means of a fully connected network of two layers, one layer to represent the word forms and another one to represent the meaning.

¹ We used the implementation of LDL for the programming language *julia* in the package *JudiLing* (Luo 2021; Luo et al. 2021). A manual for the *JudiLing* package is available at: <https://megamindhenry.github.io/JudiLing.jl/stable/>. Upon publication our data set and scripts will be made available through OSF.

Table 3: Excerpt of the C matrix. Cues that are present in a word are indicated with 1, cues that are absent are indicated with 0.

	#u.mu	mu.ku	ku.ra	ra.mbe	mbe.re#	#a.ba	ba.ku	#mu.ba
u.mu.ku.ra.mbe.re	1	1	1	1	1	0	0	0
a.ba.ku.ra.mbe.re	0	0	1	1	1	1	1	0
mu.ba.ku.ra.mbe.re	0	0	1	1	1	0	1	1

The word form layer is a matrix in which each word is represented as a vector. The ngrams of a word are one hot encoded in the vector: A present ngram is coded as 1, an absent one as 0. This is illustrated in Table 3, for words in ngrams of bisyllables. The vectors of the ngrams of the word forms are stored in a matrix called C.

We used two kinds of ngrams for the word forms: bigrams of syllables and trigrams of syllables. We choose to rely on syllables because of their role in speech production and perception (for a recent excellent review of neural evidence see Poeppel and Assaneo 2020).

The meaning layer is a matrix in which the meaning of each word is represented as a vector. In order to do this, the meaning has to be represented numerically. The distribution of the meaning of the grammatical functions *noun class*, *number*, *locative*, *diminutive* and the *lexeme* was simulated by constructing values for each of the grammatical functions of each word form following Baayen et al. (2019). An excerpt of the S matrix is provided in Table 4. The specifications of each lexeme and grammatical function describe a distribution class (Blevins 2016a).

Table 4: Excerpt of the S matrix. The numbers in the columns of the semantic dimensions $S_1, S_2, \dots, S_n$ reflect the strength of their semantic features. For example, there is strong positive support for the feature S_2 of the word form *kwitoongo* and strong negative support for the feature S_4 of the word form *mwitoongo*.

Word	S_1	S_2	S_3	S_4	S_5	S_6	S_7
i.too.ngo	7.525	9.443	10.447	-9.735	-17.675	15.343	22.638
mwi.too.ngo	8.572	7.215	15.720	-17.294	-13.777	16.831	6.209
kwi.too.ngo	8.750	12.529	13.203	-10.380	-16.550	9.127	21.543
a.ga.too.ngo	-5.401	10.494	12.976	-12.926	-18.246	12.344	15.571
a.ma.too.ngo	3.785	1.387	21.928	-0.085	-20.093	15.833	-6.393
mu.ma.too.ngo	4.759	0.728	19.561	-2.401	-22.436	14.763	-5.941
ku.ma.too.ngo	4.672	4.424	19.690	4.887	-24.980	8.465	8.207

The meaning of a word can then be represented as the sum of these distributional vectors as illustrated in example 1. The vectors of the meaning of each word are stored in a matrix called S. Alternatively, the values in the meaning layer can also be derived from word embeddings (Landauer and Dumais 1997; Niyongabo et al. 2020). Simulated word meanings give the researchers tighter control over their data, but the meanings may not reflect the distribution of word meanings that arise from usage. The choice between these types of representation depends on a number of factors, one of which is whether word embeddings are available for a language, and another one is how such embeddings are derived (detailed discussion is provided in Heitmeier et al. 2021).

$$\overrightarrow{\text{umukurambere}} = \overrightarrow{\text{ancestor}} + \overrightarrow{\text{one}} + \overrightarrow{\text{singular}} + \overrightarrow{\text{no locative}} + \overrightarrow{\text{no diminutive}} \quad (1)$$

The C and S matrices are used to model comprehension by mapping C onto ..., since it answers the question which meaning is predicted by which word form, and to model production by mapping S on C, since it answers the question which word form is predicted for a meaning. The mappings are arrived at by transformation matrices F and G, which can be derived from C and S by solving equations(2) and(3).²

² A detailed description of the mathematics behind these matrix operations is provided in Baayen et al. (2018).

$$CF = S \quad (2)$$

$$SG = C \quad (3)$$

Because the matrices are large (the C and S matrices for this study have a dimensionality of $11,180^\circ \times 5,932$), it is not possible to solve these equations directly, but they must be estimated. The estimated F and G matrices can then be used to calculate the predicted matrices $\hat{S}$ and $\hat{C}$.

The word forms and the meanings of the predicted matrices are used to assess the accuracy of comprehension and production. For comprehension the vector of the meaning of a word in S is correlated with the predicted vector of meaning for that word from $\hat{S}$. The meaning with the highest correlation is selected as the recognized meaning, and if this is indeed the meaning of the word, the word form has been accurately comprehended. In case of homonyms we also counted a predicted form as correct if the meaning of a homonym was predicted. We did so, because LDL is a computational model and it has no further means to decide among the meaning of homonyms on the basis of the data set.

As for production, the JudiLing implementation of LDL offers two measures of accuracy: production (build) and production (learn). The accuracy of the production (build) is assessed by searching for a path from the ngram at the beginning of the word to the ngram at the end of the word. As there are many possible paths (many possible words), the algorithm limits its search, in our case 15 candidate words were considered. For each of these candidates, the correlation of their predicted semantic vectors with the one of the targeted word is assessed. The word that has the highest correlation with the targeted word is selected as the predicted word, and the word form is counted as accurate if the predicted word and the targeted word are identical.

The accuracy of production (learn) is assessed by establishing a path from the first ngram of the word to the last ngram of the word, for each position in the word, the support for all n grams given the C matrix is estimated. For each ngram at each position within the word a meaning is selected which has the highest support. This procedure also constructs several candidate words. For each candidate, the correlation with the semantics of the intended word is assessed and the word form with the highest correlation is selected as the predicted word. If the predicted word is identical to the intended word it is counted as accurate.

4 Results

How successful is a model of comprehension and production of Kinyarwanda nouns in a word-based view of the mental lexicon? To answer this question, we will first present the results of modeling all data, both with simulated vectors for meaning and with vectors derived from word embeddings. In Subsection 4.2, we will discuss the results of modeling held-out data.

4.1 Comprehension and production of all words

The accuracy of comprehension and production of the model trained with bigrams of syllables and simulated vectors for meaning is almost perfect (see Table 5).

Even though the model makes very few mistakes, it is instructive to have a look at them. Table 6 lists all errors. The errors that involve lexical meanings (Gloss) are a consequence of presenting the words in isolation.

Table 5: Accuracy for modeling based on bigrams of syllables.

Comprehension	99.9%
Production (build)	99.8%
Production (learn)	99.8%

Table 6: All comprehension errors (9) for the model based on bigrams of syllables.

Target	Predicted	Error
akabú	agatubú	Gloss
utwáaro	udutwáaro	Gloss
akáaka	agakáaka	Gloss
utubú	udutubú	Gloss
kumuri	kumubiri	Gloss, Noun Class, Number
mumuri	mumubiri	Gloss, Noun Class, Number
kubanyámujinyá	kubanyámusózi	Gloss, Noun Class
akara	agakara	Gloss
utunyámujinyá	utujinyá	Gloss, Noun Class

For example, the target *akáaka* means ‘small year’, whereas the predicted word *agakáaka* means ‘small grand parent’. It is difficult to imagine a situation in which the intended meanings of *akáaka* and *agakáaka* cannot be inferred from the context of the sentence or the discourse in which they occur. But it is easy to imagine that in isolation words can be misheard, especially if the difference in phonological form is so small.

Table 7 lists 10 production errors for the build algorithm with the highest support for the wrong semantics. There were 21 errors overall. Upon closer inspection of all errors, it turns out that for all erroneous predictions the winner was part of the 15 candidates the algorithm created. 13 errors involved homonyms or forms in which the singular form is the same as the plural form. The algorithm selected the winner correctly for one of them and for eight of the other forms the target was the second best prediction of the algorithm.

Table 7: Ten errors of the production build algorithm in which the predicted form had higher support than the target form.

Target	Predicted	Error
udusígísgí	udusígi	Omission
agasígísgí	agasígi	Omission
kudusígísgí	kudusígi	Omission
mudusígísgí	mudusígi	Omission
mugasígísgí	mugasígi	Omission
kugasígísgí	kugasígi	Omission
mushooza	mumushooza	Addition
munyamáanza	mumunyamáanza	Addition
udutóorero	udukóro	Replacement
uducíiro	udukóro	Replacement

Table 8 lists 10 production errors for the learn algorithm with the highest support for the wrong semantics. There are 22 errors in total. Inspection of the errors reveals that all targets were among the ten candidates. There

Table 8: Ten errors of the production learn algorithm in which the predicted form had higher support than the target form.

Target	Predicted	Error
mugasígísgí	mugasígi	Omission
uturééré	uduko	Replacement
utubago	uduko	Replacement
utubáandé	uduko	Replacement
mushooza	mumushooza	Addition
munyamáanza	mumunyamáanza	Addition
munyama	mumunyama	Addition
akáaka	akare	Replacement
udutóorero	utudíri	Replacement
mushiingano	mumushiingano	Addition

are 13 errors that involves homonyms or word forms that have the same form in the singular and plural. In all 13 pairs, the algorithm selected the correct form as winner once, and for seven forms the target was the second best prediction.

For the data set with words the meaning of which is based on word embeddings, as illustrated in Table 9, comprehension and the production data based on the learn algorithm are still good, but the production data based on the build algorithm are not good. The drop in performance is probably a result of the way in which the build algorithm predicts a word form for production: It lines up cues in such a way as to find a string such that each cue is a possible link to its preceding and following cue. After having constructed 15 such strings, it assesses the meaning of each string. Crucially, it does so without gauging the contribution of each individual cue. The learn algorithm, in contrast, gauges the support for each cue in each position in the word. With the larger full data set, the difference between these algorithms might not appear as striking, but with small data sets, the difference has dramatic consequences. The data set based on word embeddings is much smaller, which explains the drop in performance.

Table 9: Accuracy for modeling based on bigrams of syllables word embeddings.

Comprehension	97.5%
Production (build)	19.0%
Production (learn)	83.9%

We will now turn our attention to the model based on trigrams of syllables. The accuracy of its comprehension and production is perfect as is illustrated in Table 10. However, this could well be the result of overfitting, as there are more unique cues (for discussion see Heitmeier et al. 2021).

Table 10: Accuracy for modeling based on trigrams of syllables.

Comprehension	100%
Production (build)	100%
Production (learn)	100%

For the data set with words the meaning of which is based on word embeddings, as illustrated in Table 11, comprehension and the production data based on the learn algorithm are very good, but the production data based on the build algorithm less so, just as it was for the model based on bigrams of syllables.

Table 11: Accuracy for modeling based on trigrams of syllables representing meanings with word embeddings.

Comprehension	99.9%
Production (build)	46.0%
Production (learn)	84.5%

4.2 Comprehension and production of held-out words

How does the model fare with held-out data? We trained the model on 90% of the data and tested it on the remaining 10%. The accuracy for comprehension of the test set is excellent at almost 90%, and if we count as correct cases where the model understood a homonym the accuracy is 91%; the accuracy for production data

Table 12: Accuracy for held-out test data with modeling based on bigrams of syllables.

Comprehension	89.9% (91% homonyms)
Production (build)	43.2%
Production (learn)	85.3%

based on the learn algorithm is good at 85%, but the accuracy of the production data based on the build algorithm is not good (see Table 12).

The accuracy of the model based on trigrams of syllables on the 10% held-out data is unspectacular at about 61% for comprehension and at about 58% for production (learn). The accuracy for the production (build) is dismal. A model based on trigrams of syllables is very good at recognizing what it has already encountered (see Table 10, but not good at using its memory-stock to make predictions: the model overfits).

5 Conclusion

Are Bantu nouns represented in the mental lexicon in terms of morphemes, or as whole words? The evidence for morphological decomposition of Bantu nouns from priming experiments is inconclusive (Ciaccio et al. 2020; Kgoro and Eisenbeiss 2015), but there are arguments in favor of a central role for words in the mental lexicon (Ambridge 2020; Baayen and Ramscar 2019; Baayen et al. 2019; Chuang et al. 2020) from non-Bantu languages. The highly inflectional nature of Bantu languages is well-suited to test whether nouns are understood and produced on the basis of the phonology and semantics of whole words. This is because such highly inflectional languages most closely adhere to the so-called morphemic ideal, according to which complex words are composed of unique and identifiable morphemes (Ainsworth 2019). Among the Bantu languages, Kinyarwanda has additional complexity provided by pre-prefixes (Rosendal 2006). We reasoned that if comprehension and production of Kinyarwanda nouns can be modeled well without recourse to morphemes or other prespecified morphological units, it provides a strong argument in favor of a word-based account of the Kinyarwanda mental lexicon.

We found that LDL models comprehension and production of Kinyarwanda nouns successfully, both for the whole data set (see Tables 5 and 10) and for held out data (see Tables 12 and 13). It does so by relying only on word forms and meanings. Our results support a theory of the mental lexicon in which words are the central cognitive units, since we have not provided our model with information about morphemes.

In the errors that the model makes our modeling also showed that it is necessary to study words in context rather than in isolation. Context will help reduce ambiguities that are the result of homonyms that can easily be resolved by context, and agreement markers in Bantu sentences (van der Wal 2015) will further reduce any ambiguity.

The Discriminative Lexicon incorporates a discriminative learning perspective on language, and this could serve to explain the results of the experiments of Ciaccio et al. (2020) and Kgoro and Eisenbeiss (2015). In discriminative learning, learning is achieved by minimizing prediction errors (Ramscar et al. 2013; Rescorla and Wagner 1972). Ciaccio et al. (2020) found a priming effect for prefixes but not for suffixes. This is in agreement with the idea that order matters in error-driven discrimination (Hoppe et al. 2020): Cues predict

Table 13: Accuracy for held-out test data with modeling based on trigrams of syllables.

Comprehension	60.5% (61% homonyms)
Production (build)	6.2%
Production (learn)	58.1%

following outcomes. A prefix predicts whatever it prefixes, but a suffix is predicted by whatever precedes. In an experiment without any linguistic context, a word does not predict its suffix, but a prefix does predict its related unprefixated word. This then could translate in a difference in priming. This discriminative perspective would also offer an explanation for the behavior of class 9 nouns in the experiment of Kgolo and Eisenbeiss, who found the faster reaction times for class 9 targets than for class 1 targets. An explanation could be that the cues in the transparent class 1 targets overlap with the cues in the prime, this competition between similar cues for an outcome is more inhibiting than the competition between different cues and the outcome of class 9.

Our results provide an argument in favor of the word and paradigm model (Blevins 2016a), as incorporated in the Discriminative Lexicon, and highlight that even in highly inflectional languages such as Kinyarwanda reference to words suffices to model comprehension and production.

Acknowledgment: We want to thank two anonymous reviewers, and Harald Baayen, Yu-Ying Chuang, Xuefeng Luo, Jessica Nieder, Ingo Plag and Fabian Tomaschek for their constructive comments.

References

- Ainsworth, Zeprina-Jaz. 2019. The Veps illative: The applicability of an abstractive approach to an agglutinative language. *Transactions of the Philological Society* 117(1). 58–78.
- Ambridge, Ben. 2020. Against stored abstractions: A radical exemplar model of language acquisition. *First Language* 40(5–6). 509–559.
- Baayen, R. Harald, Yu-Ying Chuang & James P. Blevins. 2018. Inflectional morphology with linear mappings. *The Mental Lexicon* 13(2). 230–268.
- Baayen, R. Harald, Yu-Ying Chuang, Elnaz Shafaei-Bajestan & James P. Blevins. 2019. The discriminative lexicon: A unified computational model for the lexicon and lexical processing in comprehension and production grounded not in (de) composition but in linear discriminative learning. *Complexity* 2019. <https://doi.org/10.1155/2019/4895891>.
- Baayen, R. Harald & Michael Ramscar. 2019. Abstraction, storage and naive discriminative learning. In Ewa Dąbrowska & Dagmar Divjak (eds.), *Cognitive Linguistics - Foundations of Language*, 115–139. Berlin, Boston: De Gruyter Mouton.
- Blevins, James P. 2006. Word-based morphology. *Journal of Linguistics* 42(3). 531–573.
- Blevins, James P. 2016a. The minimal sign. In Andrew Hippisley & Gregory Stump (eds.), *The Cambridge handbook of morphology* (Cambridge Handbooks in Language and Linguistics), 50–69. Cambridge, UK: Cambridge University Press.
- Blevins, James P. 2016b. *Word and paradigm morphology*. Oxford, UK: Oxford University Press.
- Chuang, Yu-Ying, Kaidi Lõo, James Blevins & Harald Baayen. 2020. Estonian case inflection made simple a case study in word and paradigm morphology with linear discriminative learning. In Livia Körtvélyessy & Pavol Štekauer (eds.), *Complex Words: Advances in Morphology*, 119–141. Cambridge, UK: Cambridge University Press.
- Ciaccio, Laura Anna, Naledi Kgolo & Harald Clahsen. 2020. Morphological decomposition in Bantu: A masked priming study on Setswana prefixation. *Language, Cognition and Neuroscience* 35(10). 1–15.
- Dong, Yang, Shu-Na Peng, Yuan-Ke Sun, Sammy Xiao-Ying Wu & Wei-Sha Wang. 2020. Reading comprehension and metalinguistic knowledge in Chinese readers: A metaanalysis. *Frontiers in Psychology* 10. 3037.
- Gahl, Susanne. 2008. Time and thyme are not homophones: The effect of lemma frequency on word durations in spontaneous speech. *Language* 84(3). 474–496.
- Goldsmith, John & Fidèle Mpiranya. 2018. Learning Swahili morphology. In Galen Sibanda, Deo Ngonyani, Jonathan Choti & Ann Biersteker (eds.), *Descriptive and theoretical approaches to African linguistics: Proceedings of the 49th annual conference on African Linguistics*, 73–106. Berlin: Language Science Press.
- Grainger, Jonathan & Elisabeth Beyersmann. 2017. Edge-aligned embedded word activation initiates morpho-orthographic segmentation. In Brian H. Ross (ed.), *Psychology of learning and motivation*, 67, 285–317. Cambridge, MA: Elsevier.
- Güldemann, Tom & Ines Fiedler. 2021. More diversity enGENDERed by African languages: An introduction. *STUF - Language Typology and Universals* 74(2). 221–240.
- Heitmeier, Maria, Yu-Ying Chuang & R. Harald Baayen. 2021. Modeling morphology with linear discriminative learning: Considerations and design choices. *Frontiers in Psychology* 12. 4929.
- Hoppe, Dorothee B, Jacolien van Rij, Petra Hendriks & Michael Ramscar. 2020. Order matters! Influences of linear order on linguistic category learning. *Cognitive Science* 44(11). e12910.
- Hyman, Larry M., Florian Lionnet & Christophe Ngolele. 2019. Number and animacy in the Teke noun class system. In Samson Lotven, Silvina Bongiovanni, Phillip Weirich, Robert Botne & Samuel Gyasi Obeng (eds.), *African linguistics across the disciplines*, 89–102. Berlin: Language Science Press.

- Katamba, Francis. 1978. How agglutinating is Bantu morphology? *Linguistics* 16(210). 77–84.
- Katamba, Francis. 2003. Bantu nominal morphology. In Derek Nurse & Gérard Philippson (eds.), *The Bantu languages*. New York, NY: Routledge.
- Kgolo, Naledi & Sonja Eisenbeiss. 2015. The role of morphological structure in the processing of complex forms: Evidence from Setswana deverbative nouns. *Language, Cognition and Neuroscience* 30(9). 1116–1133.
- Kimenyi, Alexandre. 1979. *Studies in Kinyarwanda and Bantu phonology*, vol. 33. Linguistic Research.
- Landauer, Thomas K. & Susan T. Dumais. 1997. A solution to plato's problem: The latent semantic analysis theory of acquisition, induction, and representation of knowledge. *Psychological Review* 104(2). 211.
- Lohmann, Arne. 2018. Time and thyme are not homophones: A closer look at Gahl's work on the lemma-frequency effect, including a reanalysis. *Language* 94(2). e180–190.
- Luo, Xuefeng. 2021. *JudiLing: An implementation for discriminative learning in julia*. Germany: Eberhard Karls University of Tübingen MA thesis. Available at: https://github.com/MegamindHenry/JudiLing.jl/blob/master/thesis/thesis_JudiLing_An_implementation_for_Discriminative_Learning_in_Julia.pdf.
- Luo, Xuefeng, Yu-Ying Chuang & Harald Baayen. 2021. *JudiLing: An implementation in Julia of Linear Discriminative Learning algorithms for language models*. Available at: <https://github.com/MegamindHenry/JudiLing.jl/blob/master/docs/src/index.md>.
- Mitterer, Holger & Eva Reinisch. 2017. Surface forms trump underlying representations in functional generalisations in speech perception: The case of German devoiced stops. *Language, Cognition and Neuroscience* 32(9). 1133–1147.
- Moscato del Prado Martín, Fermín, Raymond Bertram, Tuomo Häikiö, Robert Schreuder & R. Harald Baayen. 2004. Morphological family size in a morphologically rich language: The case of Finnish compared with Dutch and Hebrew. *Journal of Experimental Psychology: Learning, Memory, and Cognition* 30(6). 1271.
- Niyongabo, Rubungo Andre, Hong Qu, Julia Kreutzer, and Li Huang. 2020. KINNEWS and KIRNEWS: Benchmarking Cross-Lingual Text Classification for Kinyarwanda and Kirundi. *arXiv*.
- Nurse, Derek & Gérard Philippson. 2006. *The Bantu Languages*. London & New York: Routledge.
- Poeppel, David & M. Florencia Assaneo. 2020. Speech rhythms and their neural foundations. *Nature Reviews Neuroscience* 21(6). 322–334.
- Ramscar, Michael, Melody Dye & Stewart M. McCauley. 2013. Error and expectation in language learning: The curious absence of “mouses” in adult speech. *Language* 89(4). 760–793.
- Rastle, Kathleen & Matthew H. Davis. 2008. Morphological decomposition based on the analysis of orthography. *Language & Cognitive Processes* 23(7-8). 942–971.
- Rescorla, Robert A. & Allan R. Wagner. 1972. A theory of Pavlovian conditioning: Variations in the effectiveness of reinforcement and nonreinforcement. In Abraham F. Black & William H. Prokasy (eds.), *Classical conditioning II: Current research and theory*, 2, 64–99. New York: Appleton Century Crofts.
- Rosendal, Tove. 2006. *The noun classes of Rwanda – An overview. Från urindoeuropeiska till ndengereko: Nio uppsatser om fonologi och morfologi*, 143–161. Göteborg, Sweden: University of Göteborg.
- Stockall, Linnaea & Alec Marantz. 2006. A single route, full decomposition model of morphological complexity MEG evidence. *The Mental Lexicon* 1(1). 85–123.
- van der Wal, Jenneke. 2015. Bantu syntax. *Oxford handbooks online*, 1–57. Oxford, UK: Oxford University Press.