

UC Berkeley

UC Berkeley Previously Published Works

Title

Comparing Fifty Natural Languages and Twelve Genetic Languages Using Word Embedding Language Divergence (WELD) as a Quantitative Measure of Language Distance

Permalink

<https://escholarship.org/uc/item/3dw6q6r1>

Authors

Asgari, Ehsaneddin
Mofrad, Mohammad RK

Publication Date

2016-04-28

Peer reviewed

Comparing Fifty Natural Languages and Twelve Genetic Languages Using Word Embedding Language Divergence (WELD) as a Quantitative Measure of Language Distance

Ehsaneddin Asgari and Mohammad R.K. Mofrad

Departments of Bioengineering
University of California, Berkeley
Berkeley, CA 94720, USA

asgari@ischool.berkeley.edu, mofrad@berkeley.edu

Abstract

We introduce a new measure of distance between languages based on word embedding, called word embedding language divergence (WELD). WELD is defined as divergence between unified similarity distribution of words between languages. Using such a measure, we perform language comparison for fifty natural languages and twelve genetic languages. Our natural language dataset is a collection of sentence-aligned parallel corpora from bible translations for fifty languages spanning a variety of language families. Although we use parallel corpora, which guarantees having the same content in all languages, interestingly in many cases languages within the same family cluster together. In addition to natural languages, we perform language comparison for the coding regions in the genomes of 12 different organisms (4 plants, 6 animals, and two human subjects). Our result confirms a significant high-level difference in the genetic language model of humans/animals versus plants. The proposed method is a step toward defining a quantitative measure of similarity between languages, with applications in languages classification, genre identification, dialect identification, and evaluation of translations.

1 Introduction

Classification of language varieties is one of the prominent problems in linguistics (Smith, 2016). The term language variety can refer to different styles, dialects, or even a distinct language (Marjorie and Rees-Miller, 2001). It has been a long-standing argument that strictly quantitative methods

can be applied to determine the degree of similarity or dissimilarity between languages (Kroeber and Chrétien, 1937; Sankaran et al., 1950; Krámský, 1959; McMahon and McMahon, 2003). The methods proposed in the 1990’s and early 2000’ mostly relied on utilization of intensive linguistic resources. For instance, similarity between two languages was defined based on the number of common cognates or phonological patterns according to a manually extracted list (Kroeber and Chrétien, 1937; McMahon and McMahon, 2003). Such an approach, of course, is not easily extensible to problems involving new languages. Recently, statistical methods have been proposed to automatically detect cognates (Berg-Kirkpatrick and Klein, 2010; Hall and Klein, 2010; Bouchard-Côté et al., 2013; Ciobanu and Dinu, 2014) and subsequently compare languages based on the number of common cognates (Ciobanu and Dinu, 2014).

In this paper our aim is to define a quantitative measure of distance between languages. Such a metric should reasonably take both syntactic and semantic variability of languages into account. A measure of distance between languages can have various applications including quantitative genetic/typological language classification, styles and genres identification, and translation evaluation. In addition, comparing the biological languages generating the genome in different organisms can potentially shed light on important biological facts.

1.1 Problem Definition

Our goal is to be able to provide a quantitative estimate of distance for any two given languages. In our framework, we define a language as a weighted

graph $\Omega_L(V, e)$, where V is a set of vertices (words), and $e : (V \times V) \rightarrow \mathbb{R}$ is a weight function mapping a pair of words to their similarity value. Then our goal of approximating the distance between the two languages L and L' can be transferred to the approximation of the distance between $\Omega_L(V, e)$ and $\Omega_{L'}(V', e')$. In order to approach such a problem firstly we need to address the following questions:

- What is a proper weight function e estimating a similarity measure between words $w_i, w_j \in V$ in a language L ?
- How can we relate words in V to words in V' ?
- And finally, how can we measure a distance between languages Ω_L and $\Omega_{L'}$, which means $D(\Omega_L, \Omega_{L'})$?

In the following section we explain how researchers have addressed the above mentioned questions until now.

1.1.1 Word similarity within a language

The main aim of word similarity methods is to measure how similar pairs of words are to each other, semantically and syntactically (Han et al., 2013). Such a problem has a wide range of applications in information retrieval, automatic speech recognition, word sense disambiguation, and machine translation (Collobert and Weston, 2008; Glorot et al., 2011; Mikolov et al., 2013c; Turney et al., 2010; Resnik, 1999; Schwenk, 2007).

Various methods have been proposed to measure word similarity, including thesaurus and taxonomy-based approaches, data-driven methods, and hybrid techniques (Miller, 1995; Mohammad and Hirst, 2006; Mikolov et al., 2013a; Han et al., 2013). Taxonomy-based methods are not easily extensible as they usually require extensive human intervention for creation and maintenance (Han et al., 2013). One of the main advantages of data-driven methods is that they can be employed even for domains with shortage of manually annotated data.

Almost all of the data-driven methods such as matrix factorization (Xu et al., 2003), word embedding (Mikolov et al., 2013a), topic models (Blei, 2012), and mutual information (Han et al., 2013) are based on co-occurrences of words within defined units of text data. Each method has its own convention for unit of text, which can be a sentence, paragraph or a sliding window around a word.

Using distributed representations have been one of the most successful approaches for computing word similarity in natural language processing (Collobert et al., 2011). The main idea in distributed representation is characterizing words by the company they keep (Hinton, 1984; Firth, 1975; Collobert et al., 2011).

Recently, continuous vector representations known as word vectors have become popular in natural language processing (NLP) as an efficient approach to represent semantic/syntactic units (Mikolov et al., 2013a; Collobert et al., 2011). Word vectors are trained in the course of training a language model neural network from large amounts of textual data (words and their contexts) (Mikolov et al., 2013a). More precisely, word representations are the outputs of the last hidden layer in a trained neural network for language modeling. Thus, word vectors are supposed to encode the most relevant features to language modeling by observing various samples. In this representation similar words have closer vectors, where similarity is defined in terms of both syntax and semantics. By training word vectors over large corpora of natural languages, interesting patterns have been observed. Words with similar vector representations display multiple types of similarity. For instance, $\overrightarrow{King} - \overrightarrow{Man} + \overrightarrow{Woman}$ is the closest vector to that of the word $\overrightarrow{Queen}$ (an instance of semantic regularities) and $\overrightarrow{quick} - \overrightarrow{quickly} \approx \overrightarrow{slow} - \overrightarrow{slowly}$ (an instance of syntactic regularities). A recent work has proposed the use of word vectors to detect linguistic changes within the same language over time (Kulkarni et al., 2015). The fact that various degrees of similarity were captured by such a representation convinced us to use it as a notion of proximity for words.

1.1.2 Word alignment

As we discussed in section 1.1, in order to compare graphs Ω_L and $\Omega_{L'}$, we need to have a unified definition of words (vertices). Thus, we need to find a mapping function from the words in V to the words in V' . Obviously when two languages have the same vocabulary set this step can be skipped, which is the case when we perform within-language genres analysis or linguistic drifts study (Stamatatos et al., 2000; Kulkarni et al., 2015), or even when

we compare biological languages (DNA or protein languages) for different species (Asgari and Mofrad, 2015). However, when our goal is to compare distributional similarity of words for two different languages, such as French and German, we need to find a mapping from words in French to German words.

Finding a word mapping function between two languages can be achieved using a dictionary or using statistical word alignment in parallel corpora (Och and Ney, 2003; Lardilleux and Lepage, 2009). Statistical word alignment is a vital component in any statistical machine translation pipeline (Fraser and Marcu, 2007). Various methods/tools has been proposed for word alignment, such as GIZA++ (Och, 2003) and Any-align (Lardilleux and Lepage, 2009), which are able to extract high quality word alignments from sentence-aligned multilingual parallel corpora.

One of the data resources we use in this project is a large collection of sentence-aligned parallel corpora we extract from bible translations in fifty languages. Thus, in order to find a word mapping function among all these languages we used statistical word alignment techniques and in particular Any-align (Lardilleux and Lepage, 2009), which can process any number of languages at once.

1.1.3 Network Analysis of Languages

The rather intuitive approach of treating languages as networks of words has been proposed and explored in the last decade by a number of researchers (i Cancho and Solé, 2001; Liu and Cong, 2013; Cong and Liu, 2014; Gao et al., 2014). In these works, human languages, like many other aspects of human behavior, are modeled as complex networks (Costa et al., 2011), where the nodes are essentially the words of the language and the weights on the edges are calculated based on the co-occurrences of the words (Liu and Cong, 2013; i Cancho and Solé, 2001; Gao et al., 2014). Clustering of 14 languages based on various parameters of a complex network such as average degree, average path length, clustering coefficient, network centralization, diameter, and network heterogeneity has been done by (Liu and Cong, 2013). A similar approach is suggested by (Gao et al., 2014) for analysis of the complexity of six languages. Although, all of the above mentioned methods have presented

promising results about similarity and regularity of languages, to our understanding they need the following improvements:

Measure of word similarity: Considering co-occurrences as a measure of similarity between nodes, which is the basis of the above mentioned complex network methods, is a naive estimate of similarity, (Liu and Cong, 2013; i Cancho and Solé, 2001; Gao et al., 2014). The most trivial cases are synonyms, which we expect to be marked as the most similar words to each other. However, since they can only be used interchangeably with each other in the same sentences, their co-occurrences rate is very low. Thus, raw co-occurrence is not necessarily a good indicator of similarity.

Independent vs. joint analysis: Previous methods have compared the parameters of language graphs independently, except for some relatively small networks of words for illustration (Liu and Cong, 2013; i Cancho and Solé, 2001; Gao et al., 2014). However, two languages may have similar settings of the edges but for completely different concepts. Thus, a systematic way for joint comparison of these networks is essential.

Language collection: The previous analysis was performed on a relatively small number of languages. For instance in (Liu and Cong, 2013), fourteen languages were studied where twelve of them were from the Slavic family of languages, and (Gao et al., 2014) studied six languages. Clearly, studying more languages from a broader set of language families would be more indicative.

1.2 Our Contributions

In this paper, we suggest a heuristic method toward a quantitative measure of distance between languages. We propose divergence between unified similarity distribution of words as a quantitative measure of distance between languages.

Measure of word similarity: We use cosine similarity between word vectors as the metric of word similarities, which has been shown to take into account both syntactic and semantic similarities (Mikolov et al., 2013a). Thus, in the weighted language graph $\Omega_L(V, e)$, the weight function $e : (V \times V) \rightarrow \mathbb{R}$ is defined by word-vector cosine similarities between pairs of words. Although word vectors are calculated based on co-occurrences of words

within sliding windows, they are capable of attributing a reasonable degree of similarity to close words that do not co-occur.

Joint analysis of language graphs: By having word vector proximity as a measure of word similarity, we can represent each language as a joint similarity distribution of its words. Unlike the methods mentioned in section 1.1.3 which focused on network properties and did not consider a mapping function between nodes across various languages, we propose performing node alignment between different languages (Lardilleux and Lepage, 2009). Consequently, calculation of Jensen-Shannon divergence between unified similarity distributions of the languages can provide us with a measure of distance between languages.

Language collection: In this study we perform language comparison for fifty natural languages and twelve genetic language.

Natural languages: We extracted a collection of sentence-aligned parallel corpora from bible translations for fifty languages spanning a variety of language families including Indo-European (Germanic, Italic, Slavic, Indo-Iranian), Austronesian, Sino-Tibetan, Altaic, Uralic, Afro-Asiatic, etc. This set of languages is relatively large and diverse in comparison with the corpora that have been used in previous studies (Liu and Cong, 2013; Gao et al., 2014). We calculated the Jensen-Shannon divergence between joint similarity distributions for fifty language graphs consisting of 4,097 sets of aligned words in all these fifty languages. Using the mentioned divergence we performed cluster analysis of languages. Interestingly in many cases languages within the same family clustered together. In some cases, a lower degree of divergence from the source language despite belonging to different language families was indicative of a consistent translation.

Genetic languages: Nature uses certain languages to generate biological sequences such as DNA, RNA, and proteins. Biological organisms use sophisticated languages to convey information within and between cells, much like humans adopt languages to communicate (Yandell and Majoros, 2002; Searls, 2002). Inspired by this conceptual analogy, we use our languages comparison method for comparison of genetic languages in different organisms. Genome refers to a sequence of nu-

cleotides containing our genetic information. Some parts of our genome are coded in a way that can be translated to proteins (exonic regions), while some regions cannot be translated into proteins (introns) (Saxonov et al., 2000). In this study, we perform language comparison of coding regions in 12 different species (4 plants, 6 animals, and two human subjects). Our language comparison method is able to assign a reasonable relative distance between species.

2 Methods

As we discussed in 1.1, we transfer the problem of finding a measure of distance between languages L and L' to finding the distance between their language graphs $\Omega_L(V, e)$ and $\Omega_{L'}(V', e')$.

Word Embedding: We define the edge weight function $e : (V \times V) \rightarrow \mathbb{R}$ to be the cosine similarity between word vectors.

Alignment: When two languages have different words, in order to find a mapping between the words in V and V' we can perform statistical word alignment on parallel corpora.

Divergence Calculation: Calculating Jensen-Shannon divergence between joint similarity distributions of the languages can provide us with a notion of distance between languages.

Our language comparison method has three components. Firstly, we need to learn word vectors from large amounts of data in an unsupervised manner for both of the languages we are going to compare. Secondly, we need to find a mapping function for the words and finally we need to calculate the divergence between languages. In the following section we explain each step aligned with the experiment we perform on both natural languages and genetic languages.

2.1 Learning Word Embedding

Word embedding can be trained in various frameworks (e.g. non-negative matrix factorization and neural network methods (Mikolov et al., 2013c; Levy and Goldberg, 2014)). Neural network word embedding trained in the course of language modeling is shown to capture interesting syntactic and semantic regularities in the data (Mikolov et al., 2013c; Mikolov et al., 2013a). Such word embedding known as word vectors need to be trained from a large number of training examples, which are ba-

sically words and their corresponding contexts. In this project, in particular we use an implementation of the skip-gram neural network (Mikolov et al., 2013b).

In training word vector representations, the skip-gram neural network attempts to maximize the average probability of contexts for given words in the training data:

$$\begin{aligned} \operatorname{argmax}_{v,v'} \frac{1}{N} \sum_{i=1}^N \sum_{-c \leq j \leq c, j \neq 0} \log p(w_{i+j}|w_i) \\ p(w_{i+j}|w_i) = \frac{\exp(v'_{w_{i+j}} v_{w_i})}{\sum_{k=1}^W \exp(v'_{w_k} v_{w_i})}, \end{aligned} \quad (1)$$

where N is the length of the training, $2c$ is the window size we consider as the context, w_i is the center of the window, W is the number of words in the dictionary and v_w and v'_w are the n -dimensional word representation and context representation of word w , respectively. At the end of the training the average of v_w and v'_w will be considered as the word vector for w . The probability $p(w_{i+j}|w_i)$ is defined using a softmax function. In the implementation we use (Word2Vec) (Mikolov et al., 2013b) negative sampling has been utilized, which is considered as the state-of-the-art for training word vector representation.

2.1.1 Natural Languages Data

For the purpose of language classification we need parallel corpora that are translated into a large number of languages, so that we can find the alignments using statistical methods. Recently, a massive parallel corpus based on 100 translations of the Bible has been created in XML format (Christodouloupoulos and Steedman, 2015), which we choose as the database for this project. In order to make sure that we have a large enough corpus for learning word vectors, we pick the languages for which translations of both the Old Testament and the New Testament are available. From among those languages we pick the ones containing all the verses in the Hebrew version (which is the source language for most of the data) and finally we end up with almost 50 languages, containing 24,785 aligned verses. For Thai, Japanese, and Chinese we use the tokenized versions in

the database (Christodouloupoulos and Steedman, 2015). In addition, before feeding the skip-gram neural network we remove all punctuation.

In our experiment, we use the word2vec implementation of skip-gram (Mikolov et al., 2013b). We set the dimension of word vectors d to 100, and the window size c to 10 and we sub-sample the frequent words by the ratio $\frac{1}{10^3}$.

2.1.2 Genetic Languages Data

In order to compare the various genetic languages we use the IntronExon database that contains coding and non-coding regions of genomes for a number of organisms (Shepelev and Fedorov, 2006). From this database we extract a data-set of coding regions (CR) from 12 organisms consisting of 4 plants (arabidopsis, populus, moss, and rice), 6 animals (sea-urchin, chicken, cow, dog, mouse, and rat), and two human subjects. The number of coding regions we have in the training data for each organism is summarized in Table 1. The next step is splitting each sequence to a number of words. Since the genome is composed of the four DNA nucleotides A,T,G and C, if we split the sequences in the character level the language network would be very small. We thus split each sequence into n -grams ($n = 3, 4, 5, 6$), which is a common range of n -grams in bioinformatics (Ganapathiraju et al., 2002; Mantegna et al., 1995). As suggested by (Asgari and Mofrad, 2015) we split the sequence into non-overlapping n -grams, but we consider all possible ways of splitting for each sequence.

Organisms	# of CR	# of 3-grams
Arabidopsis	179824	42,618,288
Populus	131844	28,478,304
Moss	167999	38,471,771
Rice	129726	34,507,116
Sea-urchin	143457	27,974,115
Chicken	187761	34,735,785
Cow	196466	43,222,520
Dog	381147	70,512,195
Mouse	215274	34,874,388
Rat	190989	41,635,602
Human 1	319391	86,874,352
Human 2	303872	77,791,232

Table 1: The genome data-set for learning word vectors in different organisms. The number of coding regions and the total occurrences of 3-grams are presented. Clearly, the total number of all n -grams ($n=3,4,5,6$) is almost the same.

We train the word vectors for each setting of n -grams and organisms separately, again using skip-

gram neural network implementation (Mikolov et al., 2013b). We set the dimension of word vectors d to 100, and window size of c to 40. In addition, we sub-sample the frequent words by the ratio 10^{-3} .

2.2 Word Alignment

The next step is to find a mapping between the nodes in $\Omega_L(V, e)$ and $\Omega_{L'}(V', e')$. Obviously in case of quantitative comparison of styles within the same language we do not need to find an alignment between the nodes in V and V' . However, when we are comparing two distinct languages we need to find a mapping from the words in language L to the words in language L' .

2.2.1 Word Alignment for Natural Languages

As we mentioned in section 2.1.1, our parallel corpora contain texts in fifty languages from a variety of language families. We decided to use statistical word alignments because we already have parallel corpora for these languages and therefore performing statistical alignment is straightforward. In addition, using statistical alignment we hope to see evidences of consistent/inconsistent translations.

We use an implementation of Anymalign (Lardilleux and Lepage, 2009), which is designed to extract high quality word alignments from sentence-aligned multilingual parallel corpora. Although Anymalign is capable of performing alignments in several languages at the same time, our empirical observation was that performing alignments for all languages against a single language and then finding the global alignment through that alignment is faster and results in better alignments. We thus align all translations with the Hebrew version. To ensure the quality of alignments we apply a high threshold on the score of alignments. In a final step, we combine the results and end up with a set of 4,097 multilingual alignments. Hence we have a mapping from any of the 4,097 words in one language to one in any other given language, where the Hebrew words are unique, but not necessarily the others.

2.2.2 Genetic Languages Alignment

In genetic language comparison, since the n-grams are generated from the same nucleotides (A,T,C,G), no alignment is needed and V would be the same as V' .

2.3 Calculation of Language Divergence

In section 2.1 we explained how to make language graphs $\Omega_L(V, e)$ and $\Omega_{L'}(V', e')$. Then in section 2.2 we proposed a statistical alignment method to find the mapping function between the nodes in V and V' . Having achieved the mapping between the words in V and the words in V' , the next step is comparison of e and e' .

In comparing language graphs what is more crucial is the *relative* similarities of words. Intuitively we know that the relative similarities of words vary in different languages due to syntactic and semantic differences. Hence, we decided to use the divergence between relative similarities of words as a heuristic measure of the distance between two languages. To do so, firstly we normalize the relative word vector similarities within each language. Then, knowing the mapping between words in V and V' we unify the coordinates of the normalized similarity distributions. Finally, we calculate the Jensen-Shannon divergence between the normalized and unified similarity distributions of two languages:

$$D_{L,L'} = JSD(\hat{e}, \hat{e}'),$$

where $\hat{e}$ and $\hat{e}'$ are normalized and unified similarity distributions of word pairs in $\Omega_L(V, e)$ and $\Omega_{L'}(V', e')$ respectively.

2.3.1 Natural Languages Graphs

For the purpose of language classification we need to find pairwise distances between all of the fifty languages we have in our corpora. Using the mapping function obtained from statistical alignments of Bible translations, we produce the normalized and unified similarity distributions of word pairs $e^{(k)}$ for language $L^{(k)}$. Therefore to compute the quantitative distance between two languages $L^{(i)}$ and $L^{(j)}$ we calculate $D_{L_i, L_j} = JSD(e^{(i)}, e^{(j)})$.

Consequently, we calculate a quantitative distance between each pair of languages. In a final step, for visualization purposes, we perform Unweighted Pair Group Method with Arithmetic Mean (UPGMA) hierarchical clustering on the pairwise distance matrix of languages (Johnson, 1967).

2.3.2 Genetic Languages Graphs

The same approach as carried out for natural languages is applied to genetic languages corpora. Pair-

wise distances of genetic languages were calculated using Jensen-Shannon divergence between normalized and unified similarity distributions of word pairs for each pair of languages.

We calculate the pairwise distance matrix of languages for each n-gram separately to verify which length of DNA segment is more discriminative between different species.

3 Results

3.1 Classification of Natural Languages

The result of the UPGMA hierarchical clustering of languages is shown in Figure 1. As shown in this figure, many languages are clustered together according to their family and sub-family. Many Indo-European languages (shown in green) and Austronesian languages (shown in pink) are within a close proximity. Even the proximity between languages within a sub-family are preserved with our measure of language distance. For instance, Romanian, Spanish, French, Italian, and Portuguese, all of which belong to the Italic sub-family of Indo-European languages, are in the same cluster. Similarly, the Austronesian languages Cebuano, Tagalog, and Maori as well as Malagasy and Indonesian are grouped together.

Although the clustering based on word embedding language divergence matches the genetic/typological classification of languages in many cases, for some pairs of languages their distance in the clustering does not make any genetic or topological sense. For instance, we expected Arabic and Somali as Afro-Asiatic languages to be within a close proximity with Hebrew. However, Hebrew is matched with Norwegian, a Germanic Indo-European language. After further investigations and comparing word neighbors for several cases in these languages, it turns out that the Norwegian bible translation highly matches Hebrew because of being a consistent and high-quality translation. In this translation, synonym were not used interchangeably and language usage stays more faithful to the structure of the Hebrew text.

3.1.1 Divergence between Genetic Languages

The pairwise distance matrix of the twelve genetic languages for n-grams ($n = 3, 4, 5, 6$) is shown in Figure 2. Our results confirm that evolutionarily closer species have a reasonably higher level of

proximity in their language models. We can observe in Figure 2, that as we increase the number of n-grams the distinction between animal/human genome and plant genome increases.

4 Conclusion

In this paper, we proposed Word Embedding Language Divergence (WELD) as a new heuristic measure of distance between languages. Consequently we performed language comparison for fifty natural languages and twelve genetic languages. Our natural language dataset was a collection of sentence-aligned parallel corpora from bible translations for fifty languages spanning a variety of language families. We calculated our word embedding language divergence for 4,097 sets of aligned words in all these fifty languages. Using the mentioned divergence we performed cluster analysis of languages.

The corpora for all of the languages but one consisted of translated text instead of original text in those languages. This means many of the potential relations between words such as collocations and culturally influenced semantic connotations did not have the full chance to contribute to the measured language distances. This can potentially make it harder for the algorithm to detect related languages. In spite of this, however in many cases languages within the same family/sub-family clustered together. In some cases, a lower degree of divergence from the source language despite belonging to different language families was indicative of a consistent translation. This suggests that this method can be a step toward defining a quantitative measure of similarity between languages, with applications in languages classification, genres identification, dialect identification, and evaluation of translations.

In addition to the natural language data-set, we performed language comparison of n-grams in coding regions of the genome in 12 different species (4 plants, 6 animals, and two human subjects). Our language comparison method confirmed that evolutionarily closer species are closer in terms of genetic language models. Interestingly, as we increase the number of n-grams the distinction between genetic language in animals/human versus plants increases. This can be regarded as indicative of a high-level diversity between the genetic languages in plants versus animals.

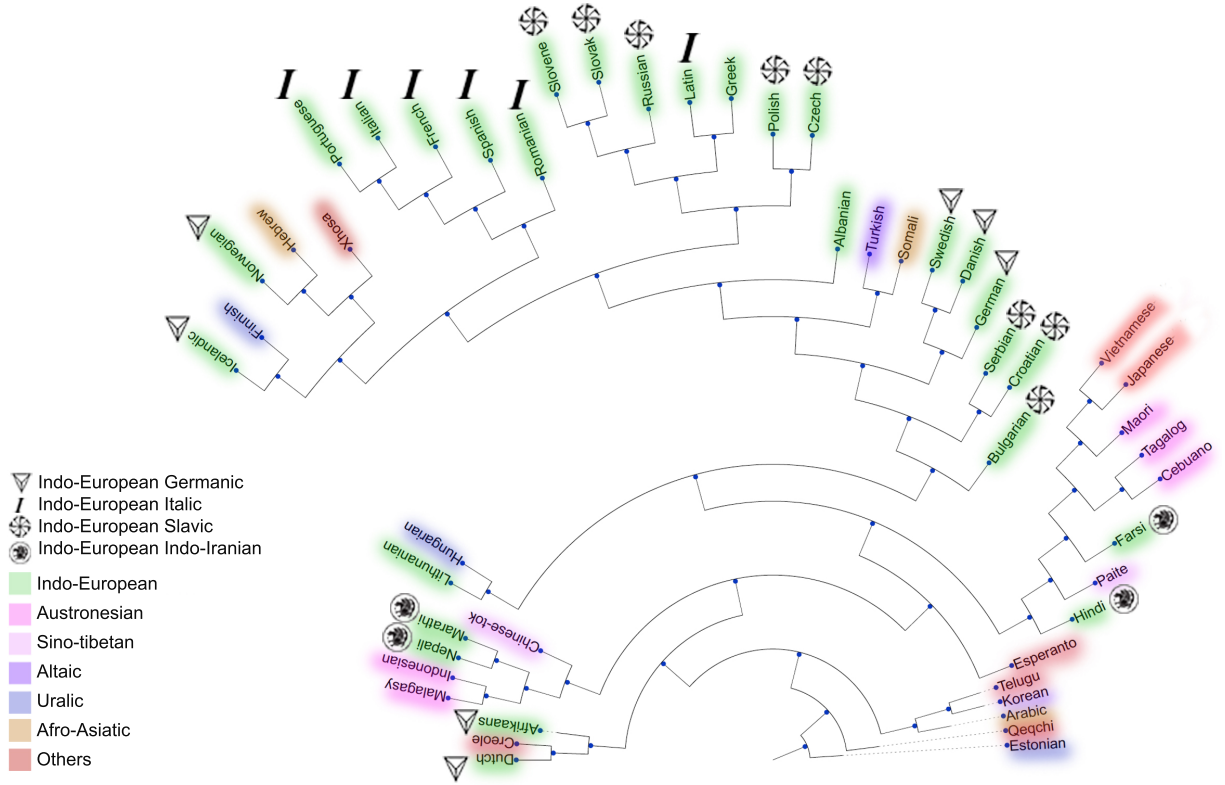


Figure 1: Hierarchical clustering of fifty natural languages according to divergence of joint distance distribution of 4097 aligned words in bible parallel corpora. Subsequently we use colors to show the ground-truth about family of languages. For Indo-European languages we use different symbols to distinguish various sub-families of Indo-European languages. We observe that the obtained clustering reasonably discriminates between various families and subfamilies.

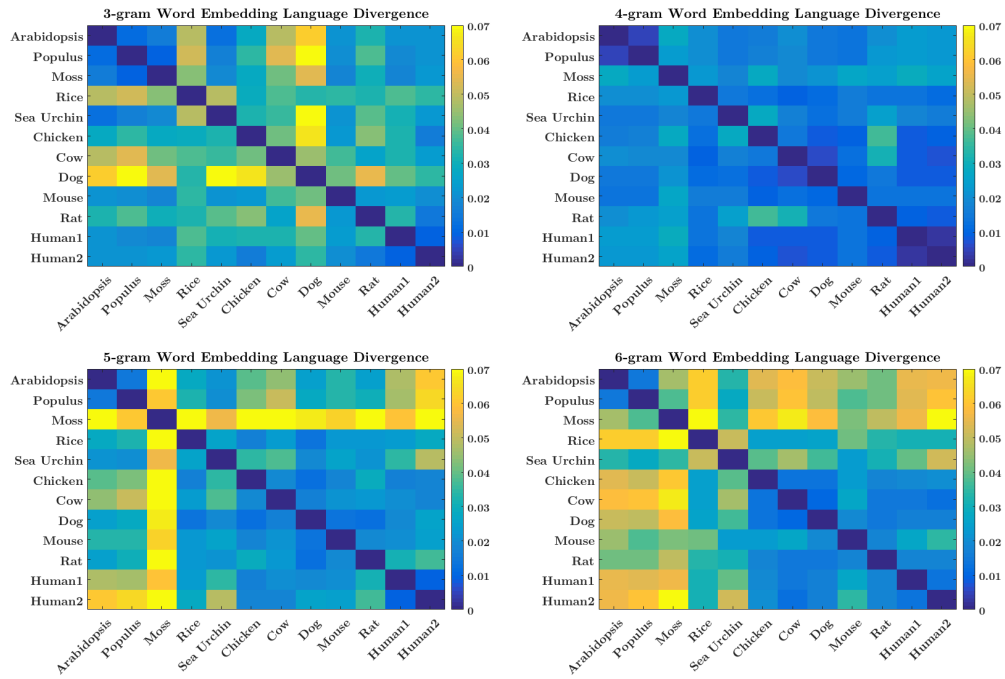


Figure 2: Visualization of word embedding language divergence in twelve different genomes belonging to 12 organisms for various n-gram segments. Our results indicate that evolutionarily closer species have higher proximity in the syntax and semantics of their genomes.

Acknowledgments

Fruitful discussions with David Bamman, Meshkat Ahmadi, and Mohsen Mahdavi are gratefully acknowledged.

References

- [Asgari and Mofrad2015] Ehsaneddin Asgari and Mohammad RK Mofrad. 2015. Continuous distributed representation of biological sequences for deep proteomics and genomics. *PloS one*, 10(11):e0141287.
- [Berg-Kirkpatrick and Klein2010] Taylor Berg-Kirkpatrick and Dan Klein. 2010. Phylogenetic grammar induction. In *Proceedings of the 48th Annual Meeting of the Association for Computational Linguistics*, pages 1288–1297. Association for Computational Linguistics.
- [Blei2012] David M Blei. 2012. Probabilistic topic models. *Communications of the ACM*, 55(4):77–84.
- [Bouchard-Côté et al.2013] Alexandre Bouchard-Côté, David Hall, Thomas L Griffiths, and Dan Klein. 2013. Automated reconstruction of ancient languages using probabilistic models of sound change. *Proceedings of the National Academy of Sciences*, 110(11):4224–4229.
- [Christodouloupoulos and Steedman2015] Christos Christodouloupoulos and Mark Steedman. 2015. A massively parallel corpus: the bible in 100 languages. *Language resources and evaluation*, 49(2):375–395.
- [Ciobanu and Dinu2014] Alina Maria Ciobanu and Liviu P. Dinu. 2014. An etymological approach to cross-language orthographic similarity. application on romanian. In *Proceedings of the 2014 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, pages 1047–1058, Doha, Qatar, October. Association for Computational Linguistics.
- [Collobert and Weston2008] Ronan Collobert and Jason Weston. 2008. A unified architecture for natural language processing: Deep neural networks with multitask learning. In *Proceedings of the 25th international conference on Machine learning*, pages 160–167. ACM.
- [Collobert et al.2011] Ronan Collobert, Jason Weston, Léon Bottou, Michael Karlen, Koray Kavukcuoglu, and Pavel Kuksa. 2011. Natural language processing (almost) from scratch. *The Journal of Machine Learning Research*, 12:2493–2537.
- [Cong and Liu2014] Jin Cong and Haitao Liu. 2014. Approaching human language with complex networks. *Physics of life reviews*, 11(4):598–618.
- [Costa et al.2011] Luciano da Fontoura Costa, Osvaldo N Oliveira Jr, Gonzalo Travieso, Francisco Aparecido Rodrigues, Paulino Ribeiro Villas Boas, Lucas Antikueira, Matheus Palhares Viana, and Luis Enrique Correa Rocha. 2011. Analyzing and modeling real-world phenomena with complex networks: a survey of applications. *Advances in Physics*, 60(3):329–412.
- [Firth1975] John Rupert Firth. 1975. *Modes of meaning*. College Division of Bobbs-Merrill Company.
- [Fraser and Marcu2007] Alexander Fraser and Daniel Marcu. 2007. Measuring word alignment quality for statistical machine translation. *Computational Linguistics*, 33(3):293–303.
- [Ganapathiraju et al.2002] Madhavi Ganapathiraju, Deborah Weissner, Roni Rosenfeld, Jaime Carbonell, Raj Reddy, and Judith Klein-Seetharaman. 2002. Comparative n-gram analysis of whole-genome protein sequences. In *Proceedings of the second international conference on Human Language Technology Research*, pages 76–81. Morgan Kaufmann Publishers Inc.
- [Gao et al.2014] Yuyang Gao, Wei Liang, Yuming Shi, and Qiuling Huang. 2014. Comparison of directed and weighted co-occurrence networks of six languages. *Physica A: Statistical Mechanics and its Applications*, 393:579–589.
- [Glorot et al.2011] Xavier Glorot, Antoine Bordes, and Yoshua Bengio. 2011. Domain adaptation for large-scale sentiment classification: A deep learning approach. In *Proceedings of the 28th International Conference on Machine Learning (ICML-11)*, pages 513–520.
- [Hall and Klein2010] David Hall and Dan Klein. 2010. Finding cognate groups using phylogenies. In *Proceedings of the 48th Annual Meeting of the Association for Computational Linguistics*, pages 1030–1039. Association for Computational Linguistics.
- [Han et al.2013] Lushan Han, Tim Finin, Paul McNamee, Akanksha Joshi, and Yelena Yesha. 2013. Improving word similarity by augmenting pmi with estimates of word polysemy. *Knowledge and Data Engineering, IEEE Transactions on*, 25(6):1307–1322.
- [Hinton1984] Geoffrey E Hinton. 1984. Distributed representations. *Computer Science Department, Carnegie Mellon University*.
- [i Cancho and Solé2001] Ramon Ferrer i Cancho and Richard V Solé. 2001. The small world of human language. *Proceedings of the Royal Society of London B: Biological Sciences*, 268(1482):2261–2265.
- [Johnson1967] Stephen C Johnson. 1967. Hierarchical clustering schemes. *Psychometrika*, 32(3):241–254.
- [Krámský1959] Jiří Krámský. 1959. A quantitative typology of languages. *Language and speech*, 2(2):72–85.

- [Kroeber and Chrétien1937] Alfred L Kroeber and C Douglas Chrétien. 1937. Quantitative classification of indo-european languages. *Language*, 13(2):83–103.
- [Kulkarni et al.2015] Vivek Kulkarni, Rami Al-Rfou, Bryan Perozzi, and Steven Skiena. 2015. Statistically significant detection of linguistic change. In *Proceedings of the 24th International Conference on World Wide Web*, pages 625–635. International World Wide Web Conferences Steering Committee.
- [Lardilleux and Lepage2009] Adrien Lardilleux and Yves Lepage. 2009. Sampling-based multilingual alignment. In *Recent Advances in Natural Language Processing*, pages 214–218.
- [Levy and Goldberg2014] Omer Levy and Yoav Goldberg. 2014. Neural word embedding as implicit matrix factorization. In *Advances in Neural Information Processing Systems*, pages 2177–2185.
- [Liu and Cong2013] HaiTao Liu and Jin Cong. 2013. Language clustering with word co-occurrence networks based on parallel texts. *Chinese Science Bulletin*, 58(10):1139–1144.
- [Mantegna et al.1995] RN Mantegna, SV Buldyrev, AL Goldberger, S Havlin, C-K Peng, M Simons, and HE Stanley. 1995. Systematic analysis of coding and noncoding dna sequences using methods of statistical linguistics. *Physical Review E*, 52(3):2939.
- [Marjorie and Rees-Miller2001] M Marjorie and Janie Rees-Miller. 2001. Language in social contexts. *Contemporary Linguistics*, pages 537–590.
- [McMahon and McMahon2003] April McMahon and Robert McMahon. 2003. Finding families: quantitative methods in language classification. *Transactions of the Philological Society*, 101(1):7–55.
- [Mikolov et al.2013a] Tomas Mikolov, Kai Chen, Greg Corrado, and Jeffrey Dean. 2013a. Efficient estimation of word representations in vector space. *arXiv preprint arXiv:1301.3781*.
- [Mikolov et al.2013b] Tomas Mikolov, Ilya Sutskever, Kai Chen, Greg S Corrado, and Jeff Dean. 2013b. Distributed representations of words and phrases and their compositionality. In *Advances in neural information processing systems*, pages 3111–3119.
- [Mikolov et al.2013c] Tomas Mikolov, Wen-tau Yih, and Geoffrey Zweig. 2013c. Linguistic regularities in continuous space word representations. In *HLT-NAACL*, pages 746–751.
- [Miller1995] George A Miller. 1995. Wordnet: a lexical database for english. *Communications of the ACM*, 38(11):39–41.
- [Mohammad and Hirst2006] Saif Mohammad and Graeme Hirst. 2006. Distributional measures of concept-distance: A task-oriented evaluation. In *Proceedings of the 2006 Conference on Empirical Methods in Natural Language Processing*, pages 35–43. Association for Computational Linguistics.
- [Och and Ney2003] Franz Josef Och and Hermann Ney. 2003. A systematic comparison of various statistical alignment models. *Computational linguistics*, 29(1):19–51.
- [Och2003] FJ Och. 2003. Giza++ software.
- [Resnik1999] Philip Resnik. 1999. Semantic similarity in a taxonomy: An information-based measure and its application to problems of ambiguity in natural language. *J. Artif. Intell. Res.(JAIR)*, 11:95–130.
- [Sankaran et al.1950] CR Sankaran, AD Taskar, and PC Ganeshsundaram. 1950. Quantitative classification of languages. *Bulletin of the Deccan College Research Institute*, pages 85–111.
- [Saxonov et al.2000] Serge Saxonov, Iraj Daizadeh, Alexei Fedorov, and Walter Gilbert. 2000. Eid: the exon-intron databasean exhaustive database of protein-coding intron-containing genes. *Nucleic acids research*, 28(1):185–190.
- [Schwenk2007] Holger Schwenk. 2007. Continuous space language models. *Computer Speech & Language*, 21(3):492–518.
- [Searls2002] David B Searls. 2002. The language of genes. *Nature*, 420(6912):211–217.
- [Shepelev and Fedorov2006] Valery Shepelev and Alexei Fedorov. 2006. Advances in the exon-intron database (eid). *Briefings in bioinformatics*, 7(2):178–185.
- [Smith2016] Andrew DM Smith. 2016. Dynamic models of language evolution: The linguistic perspective.
- [Stamatatos et al.2000] Efstathios Stamatatos, Nikos Fakotakis, and George Kokkinakis. 2000. Text genre detection using common word frequencies. In *Proceedings of the 18th conference on Computational linguistics-Volume 2*, pages 808–814. Association for Computational Linguistics.
- [Turney et al.2010] Peter D Turney, Patrick Pantel, et al. 2010. From frequency to meaning: Vector space models of semantics. *Journal of artificial intelligence research*, 37(1):141–188.
- [Xu et al.2003] Wei Xu, Xin Liu, and Yihong Gong. 2003. Document clustering based on non-negative matrix factorization. In *Proceedings of the 26th annual international ACM SIGIR conference on Research and development in informaion retrieval*, pages 267–273. ACM.
- [Yandell and Majoros2002] Mark D Yandell and William H Majoros. 2002. Genomics and natural language processing. *Nature Reviews Genetics*, 3(8):601–610.