

Substitute Based SCODE Word Embeddings in Supervised NLP Tasks

First Author

Affiliation / Address line 1
Affiliation / Address line 2
Affiliation / Address line 3
email@domain

Second Author

Affiliation / Address line 1
Affiliation / Address line 2
Affiliation / Address line 3
email@domain

Abstract

We analyze a word embedding method in supervised tasks. It maps words on a sphere such that words co-occurring in similar contexts lie closely. The similarity of contexts is measured by the distribution of substitutes that can fill them. We compared word embeddings, including more recent representations (Huang et al., 2012; Mikolov et al., 2013), in Named Entity Recognition (NER), Chunking, and Dependency Parsing. We examine our framework in multilingual dependency parsing as well. The results show that the proposed method achieves as good as or better results compared to the other word embeddings in the tasks we investigate. It achieves state-of-the-art results in multilingual dependency parsing. Word embeddings in 7 languages from 8 corpora are available for public use.

1 Introduction

Word embeddings represent each word with a dense, real valued vector. The dimension of word embeddings are generally small compared to the vocabulary size. They do not suffer from sparsity unlike one-hot representations which have the dimensionality of the vocabulary and a single non-zero entry. They capture semantic and syntactic similarities (Mikolov et al., 2013). They may help reduce the dependence on hand-designed features which are task and language dependent.

We analyze a word embedding method proposed in (Yatbaz et al., 2012), in supervised Natural Language Processing (NLP) tasks. The method represents the context of a word by its probable substitutes. Words with their probable substitutes are fed to a co-occurrence modeling framework (SCODE) (Maron et al., 2010). Words co-

occurring in similar context are closely embedded on a sphere. These word embeddings achieve state-of-the-art results in inducing part-of-speech (POS) tags for several languages (Yatbaz et al., 2014). However, their use in supervised tasks has not been well studied so far. This study aims to fill this gap.

Turian et al. (2010) compared word embeddings in Named Entity Recognition (NER) and Chunking. They use word embeddings as auxiliary features in existing systems. They improved results in both tasks compared to the baseline systems. Following this study, we report results in Chunking and NER benchmarks for SCODE embeddings. In addition, we examine word embeddings in dependency parsing. We report multilingual dependency parsing results for SCODE embeddings as well.

SCODE embeddings achieve comparable or better results compared to the other word embeddings. Multilingual results in dependency parsing also suggest that SCODE embeddings are consistent in achieving good results across different languages.

2 Related Work

In this section, we introduce word embeddings we mentioned in this work.

- **C&W:** Collobert and Weston (2008) introduce a convolutional neural network architecture that is capable of learning a language model and generating word embeddings from unlabeled data. The model can be fine-tuned for supervised NLP tasks.
- **HLBL:** Mnih and Hinton (2007) introduce the log-bilinear language model. It is a feed-forward neural network with one linear hidden layer and a softmax output layer. The model utilizes linear combination of word type representations of preceding words to predict the next word. Mnih and Hinton

(2009) modify this model to reduce computational cost by introducing a hierarchical structure. The architecture is then named the hierarchical log-bilinear language model.

- **GCA NLM:** Huang et al. (2012) introduce an architecture using both local and global context via a joint training objective. The training is very similar to (Collobert and Weston, 2008). They represent a word context by taking the weighted average of the representations of word types in a fixed size window around the target word token. Following (Reisinger and Mooney, 2010), they cluster word context representations for each word type to form word prototypes. These prototypes capture homonymy and polysemy relations.
- **LR-MVL:** Dhillon et al. (2011) present a spectral method to induce word embeddings. They perform the Canonical Correlation Analysis on the context of a token. They provide an algorithm to represent a target word with different vectors depending on its context. The objective function they define is convex. Thus, the method is guaranteed to converge to the optimal solution.
- **Skip-Gram NLM:** Mikolov et al. (2010) propose a two neural models to induce word embeddings. The first architecture is Continuous Bag-of-Words where the words in a window of target words is used to classify the target word. The second one is continuous Skip-Gram model in which the target word is used to classify its surrounding words. Mikolov et al. (2013) show that these representations reflect syntactic and semantic regularities.
- **SCODE Word Embeddings:** Maron et al. (2010) introduce the SCODE framework, an extension of the CODE (Globerson et al., 2007) framework. Maron et al. (2010) obtains word type representations from co-occurrence data generated by using neighbors of words. Yatbaz et al. (2012) extend this work by generating co-occurrence data using probable substitutes of words. In Section 3, we explain this framework in detail. Here, we review studies extending that work.

Baskaya et al. (2013) used SCODE word embeddings for Word Sense Induction. They achieved the best results in Semeval 2013 Shared Task (Jurgens and Klapaftis, 2013). (Cirik and Sensoy, 2013) is the first study exploiting SCODE embeddings in a supervised setup by using them as word features.

3 Substitute Based SCODE Word Embeddings

In this section, we summarize our framework based on (Yatbaz et al., 2012). In Section 3.1, we explain substitute word distributions. In Section 3.2, we explain how substitute word distributions are discretized. In Section 3.3 we introduce Spherical Co-Occurrence Data Embedding framework (Maron et al., 2010).

3.1 Substitute Word Distributions

Substitute word distributions are defined as the probability of observing a word in the context of the target word. We define the context of a target word as the sequence of words in the window of size $2n - 1$ centered at the position of the target word token. The target word is excluded in the context.

(1)“Steve Martin has already **laid** his claim to that .”

For example, in the sentence (1), the context of the word token ‘*laid*’, for $n = 4$, is ‘*Martin has already — his claim to*’ where — specifies the position of the target word token.

Table 1: Substitute word distribution for “laid” in sentence (1).

Probability	Substitute Word
0.191	staked
0.161	established
0.125	made
0.096	proved
0.094	rejected

Table 1 illustrates the substitute distribution of “laid” in (1). There is a row for each word in the vocabulary. For instance, probability of “established” occurring in the position of “laid” is 0.161 in this context.

Let target word token be in the position 0, the context spans from positions $-n + 1$ to $n - 1$. The

probability of observing each word w in vocabulary in the context of the target word token is calculated as follows:

$$\begin{aligned}
P(w_0 = w|c_w) &\propto P(w_{-n+1} \dots w_0 \dots w_{n-1}) \quad (1) \\
&= P(w_{-n+1})P(w_{-n+2}|w_{-n+1}) \\
&\quad \dots P(w_{n-1}|w_{-n+1}^{n-2}) \quad (2) \\
&\approx P(w_0|w_{-n+1}^{-1})P(w_1|w_{-n+2}^0) \\
&\quad \dots P(w_{n-1}|w_0^{n-2}) \quad (3)
\end{aligned}$$

In the Equation 1, the right-hand side is proportional to the left-hand side because $P(c_{w_0})$ is independent of any word w for w_0 . After using the chain rule, Equation 2 is obtained from the right-hand side of Equation 1. By applying n^{th} -order Markov assumption, only the closest $n - 1$ words in each term of the Equation 2 are needed which equals to the Equation 3. The Equation 3 is proportional to the Equation 2 because the context of the target word is fixed, thus, any term that does not depend on w_0 is fixed. Equation 3 are truncated or dropped near the boundaries of the sentence. (e.g. if 0 is the first word of a sentence, $P(w_0|w_{-n+1}^{-1})$ becomes $P(w_0)$). An n-gram language model provides the probabilities required for Equation 3.

3.2 Discretization of Substitute Word Distributions

The co-occurrence embedding algorithm we describe in Section 3.3, requires its input as categorical variables co-occurring together. We aim to associate words co-occurring in the same context. Although substitute word distributions represent the context of a word, they are categorical probability distributions. Thus, they should be transformed into a discrete setting.

We sample word types from substitute word distributions. The number of samples should be chosen carefully, if the number of the samples are too small, it may fail to capture the characteristics of the distribution.

Figure 1 is an example of a discretization with sampling. Substitute words are sampled from substitute word distributions of sentence (1).

3.3 Spherical Co-Occurrence Data Embedding

This section shortly reviews the Symmetric Interaction Model of the Co-occurrence Data Embedding (CODE) (Globerson et al., 2007) and its

Word Token	Substitute Word
Steve	Mr
Steve	Chris
Martin	Coppell
Martin	Wilson
has	had
has	has
already	finally
already	already
laid	made
laid	shown
his	no
his	no
claim	response
claim	testimony
to	to
to	to
that	fame
that	succeed
.	.
.	.

Figure 1: Sampling twice from the substitute word distributions of sentence (1).

extension Spherical Co-Occurrence Data Embedding (SCODE) (Maron et al., 2010).

We map co-occurrence data generated from the word types and substitute word distributions described in Section 3.2 to d dimensional Euclidean space.

Let X and Y have a joint distribution such that X and Y be are categorical variables with finite cardinality $|X|$ and $|Y|$. However we only observe a set of pairs $\{x_i, y_i\}_{i=1}^n$ drawn IID from the joint distribution of X and Y . These pairs are summarized by the empirical distributions $\bar{p}(x, y)$, $\bar{p}(x)$ and $\bar{p}(y)$. Embeddings $\phi(x)$ and $\psi(y)$ can capture the statistical relationship between the variables x and y in terms of square of Euclidean distance $d_{x,y}^2 = \|\phi(x) - \psi(y)\|^2$. In other words, pairs frequently co-occur are embedded closely in d dimensional space.

We used the following extended model Maron et al. (2010) proposed among others in (Globerson et al., 2007) :

$$p(x, y) = \frac{1}{Z} \bar{p}(x) \bar{p}(y) e^{-d_{x,y}^2} \quad (4)$$

where $Z = \sum_{x,y} \bar{p}(x) \bar{p}(y) e^{-d_{x,y}^2}$ is the normalization term. The log-likelihood of the joint distribution over all embeddings ϕ and ψ can be described as the following:

$$\ell(\phi, \psi) = \sum_{x,y} \bar{p}(x, y) \log p(x, y) \quad (5)$$

$$= \sum_{x,y} \bar{p}(x, y) (-\log Z + \log \bar{p}(x) \bar{p}(y) - d_{x,y}^2) \quad (6)$$

$$= -\log Z + \text{const} - \sum_{x,y} \bar{p}(x, y) d_{x,y}^2 \quad (7)$$

The gradient of the log-likelihood depends on the sum of embeddings $\phi(x)$ and $\psi(y)$, for $x \in X$ and $y \in Y$, and to maximize the log-likelihood, (Maron et al., 2010) use a gradient-ascent approach. The gradient is :

$$\frac{\partial \ell(\phi, \psi)}{\partial \phi(x)} = \sum_y 2\bar{p}(x, y)[\psi(y) - \phi(x)] + \frac{1}{Z} \sum_y \bar{p}(x)\bar{p}(y)[\phi(x) - \psi(y)]e^{-d_{x,y}^2} \quad (8)$$

$$\frac{\partial \ell(\phi, \psi)}{\partial \psi(y)} = \sum_x 2\bar{p}(x, y)[\phi(x) - \psi(y)] + \frac{1}{Z} \sum_x \bar{p}(x)\bar{p}(y)[\psi(y) - \phi(x)]e^{-d_{x,y}^2} \quad (9)$$

The first sum in (8) and (9), the gradient of the part with $d_{x,y}^2$ of (5) acts as an attraction force between the $\phi(x)$ and $\psi(y)$. The second sum in (8) and (9), the gradient of $-\log Z$ in (5) acts a repulsion force between the $\phi(x)$ and $\psi(y)$.

Maron et al. (2010) constrain all embeddings ϕ and ψ to lie on the d dimensional unit sphere, hence the name SCODE. A coarse approximation in which all ϕ and ψ distributed uniformly and independently on the sphere, enables Z to be approximated by a constant value. Thus, it does not require the computation of Z during training.

For the experiments in the work, we use SCODE with sampling based stochastic gradient ascent a constant approximation of Z and randomly initialized ϕ and ψ vectors.

4 Induction of Word Embeddings

This section explains how we induced Substitute Based SCODE Word Embeddings and obtain other embeddings. We report the details of unlabeled data used to induce word embeddings. We present the parameters chosen for induction. We explain how we obtain other word embeddings.

Unlabeled Data

Word embeddings require large amount of unlabeled data to efficiently capture syntactic and semantic regularities. The source of the data also may have an impact on the success of the word embedding on the labeled data. Thus, we induce word embeddings using a large unlabeled corpora.

Following (Turian et al., 2010), we used RCV1 corpus containing 190M word tokens (Rose et al., 2002) corpus. After following the preprocessing technique described in (Turian et al., 2010), the corpus has 80M word tokens.

We induce word embeddings for multilingual experiments explained in Section 5. We generate embeddings using subsamples of corresponding Tenten Corpora (Jakubíček et al., 2013) for Czech, German, Spanish and Swedish and Wikipedia dump files for Bulgarian, Hungarian. For Turkish, we used a web corpus (Sak et al., 2008). Table 2 shows the statistics of unlabeled corpora for languages.

Table 2: Unlabeled Corpora of Different Languages for Word Embeddings

Language	Corpus	Number Of Words
Bulgarian	Wikipedia	101M
Czech	Tenten	140M
English	RCV1	80M
German	Tenten	180M
Spanish	Tenten	106M
Swedish	Tenten	113M
Turkish	Web Corpus	180M

Parameters

To generate substitute word distributions, we trained a 4-gram statistical language model (LM) using SRILM (Stolcke, 2002). We used interpolated Kneser-Ney discounting. We replaced words observed less than 2 times with an unknown tag. Table 4 shows the statistics of language model corpora¹ for each language. We used FASTSUBS algorithm (Yuret, 2012) to generate top 100 substitutes words and their substitute probabilities.

We keep each word with its original capitalization. We sampled 100 substitutes per instance. The SCODE normalization constant was set to 0.166. For multilingual experiments we used 25 dimension word embeddings. We observe no significant improvements in scores when we change the number of dimensions for SCODE embeddings.

Other Word Embeddings

We downloaded word embeddings from corresponding studies^{2,3,4}(Turian et al., 2010; Dhillon et al., 2011; Huang et al., 2012). We should note that we do not use the context-aware word embeddings of (Dhillon et al., 2011). These word

¹We should note that LM corpora differ from the word embedding corpora. The first one is used to learn a LM which is then used for generating substitute words on the word embedding corpora.

²<http://metaoptimize.com/projects/wordreprs/>

³<http://www.cis.upenn.edu/~ungar/eigenwords/>

⁴<http://goo.gl/ZXv0Ot>

Table 3: Word token coverage for word embeddings.

Word Embeddings	Chunking			NER			Dependency Parsing	
	Training	Development	Test	Training & Development	Test	OOD	Training	Test
C&W	0.9800	0.9832	0.9764	0.9402	0.9359	0.9631	0.9835	0.9856
HLBL	0.9654	0.9675	0.9621	0.9549	0.9503	0.9777	0.9691	0.9674
GCA NLM	0.8230	0.8271	0.8139	0.6971	0.6760	0.8208	0.8322	0.8270
LR-MVL	0.9806	0.9839	0.9778	0.9422	0.9380	0.9637	0.9841	0.9862
Skip-Gram NLM	0.9848	0.9877	0.9827	0.9117	0.9075	0.9614	0.9833	0.9852
SCODE	0.9848	0.9877	0.9827	0.9117	0.9075	0.9614	0.9833	0.9852

Table 4: Unlabeled Corpora for Language Modeling

Language	Corpus	Number Of Words
Bulgarian	Wikipedia	850M
Czech	Tenten	1.79B
English	ukWac	2B
German	Tenten	1.8B
Spanish	Tenten	2.4B
Swedish	Tenten	113M
Turkish	Web Corpus	1.8B

embeddings are scaled with parameter $\sigma = 0.1$, since Turian et al. (2010) have shown that word embeddings achieve their optima at this value. We use 50-dimension of each word embeddings in all comparisons.

To induce Skip-Gram NLM embeddings, we ran the code provided on the website⁵ of (Mikolov et al., 2010; Mikolov et al., 2013) on the RCV1 corpus. We used Skip-Gram model with default parameters. We changed words occurring less than 2 times with an unknown tag. The performance of Skip-Gram NLM and SCODE word embeddings do not improve with scaling, thus, we use them without scaling.

We report word token coverage for word embeddings in Table 3. For each task, an unknown word in the training or test phase is replaced with the word embedding of unknown tag. Thus, the word embedding method with high coverage suffers less from unknown words, which in turn affects its success. Table 3 shows the word token coverage for each task and their corresponding datasets. GCA NLM has the lowest coverage in all tasks, which may explain its level of performance.

5 Experiments

In this section, we detail the experiments. We introduce tasks in which we compared word embeddings, the data used, and parameter choices made.

⁵<https://code.google.com/p/word2vec/>

We report results for each task.

Chunking

We used CoNLL-2000 Shared task Chunking as the first benchmark (Tjong Kim Sang and Buchholz, 2000). The data is from Penn Treebank which is a newswire text from Wall Street Journal (Marcus et al., 1999). The training set contains 8.9K sentences. The development set contains 1K sentences and the test set has 2K.

Table 5: Features Used In CRF Chunker

- Word features: w_i for i in $\{-2, -1, 0, +1, +2\}$, $w_i \wedge w_{i+1}$ for i in $\{-1, 0\}$
- Tag features: w_i for i in $\{-2, -1, 0, +1, +2\}$, $t_i \wedge t_{i+1}$ for i in $\{-2, -1, 0, +1\}$, $t_i \wedge t_{i+1} \wedge t_{i+2}$ for i in $\{-2, -1, 0\}$.
- Embedding features: $e_i[d]$ for i in $\{-2, -1, 0, +1, +2\}$, where d ranges over the dimensions of the embedding e_i .

We used publicly available implementation of (Turian et al., 2010). It is a CRF based chunker using features described in Table 5. The only hyperparameters of the model was L2-regularization σ which is optimal at 2. After successfully replicating results in that work⁶, we ran experiments for new word embeddings.

In Table 6, we report F1-score of word embeddings and the score of the baseline chunker that is not using word embeddings. They all improve baseline chunker, however, improvement is marginal for all of them. The best score is achieved by SCODE embeddings trained on RCV1 corpus.

Named Entity Recognition

The second benchmark is CoNLL-2003 shared task Named Entity Recognition (Tjong Kim Sang and De Meulder, 2003). The data is extracted from RCV1 Corpus. Training, development, and

⁶We report our replication of results for word embeddings which differs from (Dhillon et al., 2011).

Table 6: Chunking Results for Word Embeddings. The ones in bold font are the highest scores in their columns.

Word Embeddings	Development Score	Test Score
Baseline	0.9416	0.9379
C&W	0.9466	0.9410
HLBL	0.9463	0.9400
GCA NLM	0.9425	0.9402
LR-MVL	0.9458	0.9416
Skip-Gram NLM	0.9400	0.9402
SCODE	0.9430	0.9429

test set contains 14K, 3.3K and 3.5 sentences. Annotated named entities are location, organization and miscellaneous names. (Tjong Kim Sang and De Meulder, 2003) details the number of named entities and data preprocessing. In addition, (Turian et al., 2010) evaluated word embeddings on an out-of-domain (OOD) data containing 2.4K sentences (Chinchor, 1997).

Table 7: Features Used In Regularized Averaged Perceptron. Word embeddings are used the same way as in Table 5.

- Previous two predictions y_{i-1} and y_{i-2}
- Current word x_i
- x_i word type information : all-capitalized, is-capitalized, all-digits, alphanumeric etc.
- Prefixes and suffixes of x_i , if the word contains hyphens, then the tokens between the hyphens
- Tokens in the window $c = (x_{i-2}, x_{i-1}, x_i, x_{i+1}, x_{i+2})$
- Capitalization pattern in the window c
- Conjunction of c and y_{i-1}

We used publicly available implementation of (Turian et al., 2010). It is a regularized averaged perceptron model using features described in Table 7. After we replicated results of that work, we ran the same experiments for new word embeddings. It is important to note that, unlike (Turian et al., 2010), we did not use any non-local features or gazetteers because we wanted to measure the performance gain of word embeddings alone. The only hyperparameter is the number of epochs for the perceptron. The perceptron stops when there is no improvement for 10 epochs on the development set. The best epoch on development set is used for the final model.

Table 8 summarizes the result of NER exper-

iments. The first three rows from (Turian et al., 2010), report the baseline and the best results for C&W and HLBL embeddings. The baseline system does not use word embeddings as features. All of the word embeddings significantly improve the baseline system. SCODE embeddings trained on RCV1 corpus achieves the best score on test set and Out of Domain Test (OOD) set. Note that RCV1 corpus is the superset of NER training and test data. Thus, C&W, HLBL and SCODE on RCV1 embeddings are from the same data source.

Table 8: NER Results for Word Embeddings. The ones in bold fonts are the highest scores in their columns.

Word Embeddings	Development	Test	OOD
Baseline	0.9003	0.8439	0.6748
C&W 200-dim	0.9246	0.8796	0.7551
HLBL 100-dim	0.9200	0.8813	0.7525
C&W	0.9227	0.8793	0.7574
HLBL	0.9146	0.8705	0.7293
GCA NLM	0.9	0.8467	0.6752
LR-MVL	0.9171	0.8683	0.7323
Skip-Gram NLM	0.9095	0.8647	0.7194
SCODE	0.9207	0.8835	0.7739

Dependency Parsing

We chose CoNLL-2008 data (Surdeanu et al., 2008) as the benchmark to compare word embeddings in English Dependency Parsing. For computational reasons, we fixed the training set to the first 5K sentences of CoNLL 2008 English dataset. However, we conducted experiments using full training set with SCODE embeddings. For multilingual experiments, we chose CoNLL-2006 Shared Task languages Bulgarian, Spanish, Czech, German, Swedish, and Turkish (Buchholz and Marsi, 2006).

We used a framework (Lei et al., 2014) that is capable of incorporating word embeddings in dependency parsing. It reduces the dimensionality of head-modifier feature vectors by learning a tensor of low rank. The model is able to combine features from state-of-the-art parsers MST Parser (McDonald et al., 2005) and Turbo Parser (Martins et al., 2013) as well as low-rank tensor features which includes word embeddings. Features used in the model is listed in Table 9.

There are two hyperparameters γ and r . The first one balances tensor features and traditional MST/Turbo features. The second one is the rank of the tensor. We set the hyperparameters $\gamma = 0.3$ and $r = 50$ and ran third-order model to get

Table 9: Features Used In Low-Rank Tensor based Dependency Parser

- Unigram Features: for current word x_i form, lemma and POS tag of $x_{i,i-1,i+2}$, morphology of x_i , bias
- Bigram Features : previous and current POS tag, the current and next POS tag, current POS and lemma, current lemma and morphology
- Trigram Features: POS tag of the previous, current, and next word.
- Embedding features: $e_i[d]$ for i in $\{-1,0,+1\}$, where d ranges over the dimensions of the embedding e_i .

comparable result in that work.

Table 10 shows the Unlabeled Accuracy Scores for word embeddings and the baseline parser which is not using word embeddings. Each word embedding shows improvements over baseline parser. However, improvements are marginal, similar to Chunking results. SCODE embeddings trained on RCV1 corpus achieve the best scores among others.

Table 10: Dependency Parsing Results for ConLL 2008 English Data for Word Embeddings. The ones in bold font are the highest scores in their columns.

Word Embeddings	Training Score	Test Score
Baseline	0.9447	0.8976
C&W	0.9332	0.9007
HLBL	0.9459	0.9013
GCA NLM	0.9140	0.8985
LR-MVL	0.9308	0.9016
Skip-Gram NLM	0.9397	0.9014
SCODE	0.9444	0.9028

We report Multilingual Dependency Parsing scores in Table 11. In the first column, the results reported in (Lei et al., 2014) is listed. In the second column, the state-of-the-art results before (Lei et al., 2014). In the third column, the parser using the SCODE embeddings are listed. SCODE embeddings improve parsers for 6 out of 7 languages and achieve the best results for 5 out of 7 of them.

6 Conclusion

We analyzed SCODE word embeddings in supervised NLP tasks. SCODE word embeddings are previously used in unsupervised part of speech tagging (Yatbaz et al., 2012; Cirik, 2013; Yatbaz et al., 2014) and word sense induction (Baskaya et al., 2013). Their first use in a supervised set-

Table 11: Dependency Parsing Results for ConLL 2006 Languages for SCODE Embeddings. English results are from ConLL 2008. The ones in bold font are the highest scores in their rows.

Language	Baseline	State-of-The-Art	SCODE Embeddings
Bulgarian	0.9350	0.9402	0.9413
Czech	0.9050	0.9032	0.9038
English	0.9302	0.9322	0.9344
German	0.9197	0.9241	0.9233
Spanish	0.8800	0.8796	0.8823
Swedish	0.9100	0.9162	0.9165
Turkish	0.7684	0.7755	0.7783

ting was in dependency parsing (Cirik and Sensoy, 2013), however, results were inconclusive. Lei et al. (2014) successfully make use of SCODE embeddings as additional features in dependency parsing.

We compared SCODE word embeddings with existing word embeddings in Chunking, NER, and Dependency Parsing. For all these benchmarks, we used publicly available implementations. They all are near state-of-the-art solutions in these tasks. SCODE word embeddings are at least good as other word embeddings or achieved better results.

We analyzed SCODE embeddings in multilingual Dependency Parsing. SCODE embeddings are consistent in improving the baseline systems. Note that other word embeddings are not studied in multilingual settings yet. SCODE word embeddings and the code used in generating embeddings in this work is publicly available⁷.

References

- Osman Baskaya, Enis Sert, Volkan Cirik, and Deniz Yuret. 2013. Ai-ku: Using substitute vectors and co-occurrence modeling for word sense induction and disambiguation. *Atlanta, Georgia, USA*, page 300.
- Sabine Buchholz and Erwin Marsi. 2006. Conll-x shared task on multilingual dependency parsing. In *Proceedings of the Tenth Conference on Computational Natural Language Learning*, pages 149–164. Association for Computational Linguistics.
- Nancy Chinchor. 1997. Muc-7 named entity task definition.
- Volkan Cirik and Hüsnü Sensoy. 2013. The ai-ku system at the spmrl 2013 shared task: Unsupervised features for dependency parsing. In *Proceedings of the Fourth Workshop on Statistical Parsing of Morphologically-Rich Languages*, pages 68–75.

⁷link

- Volkan Cirik. 2013. Addressing ambiguity in unsupervised part-of-speech induction with substitute vectors. *ACL 2013*, page 117.
- R. Collobert and J. Weston. 2008. A Unified Architecture for Natural Language Processing: Deep Neural Networks with Multitask Learning. In *International Conference on Machine Learning, ICML*.
- Paramveer Dhillon, Dean P Foster, and Lyle H Ungar. 2011. Multi-View Learning of Word Embeddings via CCA. In *Advances in Neural Information Processing Systems*, pages 199–207.
- Amir Globerson, Gal Chechik, Fernando Pereira, and Naftali Tishby. 2007. Euclidean Embedding of Co-occurrence Data. *J. Mach. Learn. Res.*, 8:2265–2295, December.
- Eric H. Huang, Richard Socher, Christopher D. Manning, and Andrew Y. Ng. 2012. Improving Word Representations via Global Context and Multiple Word Prototypes. In *Proceedings of the 50th Annual Meeting of the Association for Computational Linguistics: Long Papers - Volume 1*, ACL '12, pages 873–882, Stroudsburg, PA, USA. Association for Computational Linguistics.
- Miloš Jakubíček, Adam Kilgariff, Vojtěch Kovář, Pavel Rychlý, and Vít Suchomel. 2013. The ten-ten corpus family. In *International Conference on Corpus Linguistics, Lancaster*.
- David Jurgens and Ioannis Klapaftis. 2013. Semeval-2013 task 13: Word sense induction for graded and non-graded senses. In *Second Joint Conference on Lexical and Computational Semantics (*SEM)*, volume 2, pages 290–299.
- Tao Lei, Yu Xin, Yuan Zhang, Regina Barzilay, and Tommi Jaakkola. 2014. Low-Rank Tensors for Scoring Dependency Structures. In *Proceedings of the 52nd Annual Meeting of the Association for Computational Linguistics*. Association for Computational Linguistics.
- Mitchell P. Marcus, Beatrice Santorini, Mary Ann Marcinkiewicz, and Ann Taylor. 1999. Treebank-3. Linguistic Data Consortium, Philadelphia.
- Yariv Maron, Michael Lamar, and Elie Bienenstock. 2010. Sphere Embedding: An Application to Part-of-Speech Induction. In J. Lafferty, C. K. I. Williams, J. Shawe-Taylor, R.S. Zemel, and A. Culotta, editors, *Advances in Neural Information Processing Systems 23*, pages 1567–1575.
- André FT Martins, Miguel B Almeida, and Noah A Smith. 2013. Turning on the turbo: Fast third-order non-projective turbo parsers. *Proc. of ACL*.
- Ryan McDonald, Koby Crammer, and Fernando Pereira. 2005. Online large-margin training of dependency parsers. In *Proceedings of the 43rd Annual Meeting on Association for Computational Linguistics*, pages 91–98. Association for Computational Linguistics.
- Tomas Mikolov, Martin Karafiát, Lukáš Burget, Jan Cernocky, and Sanjeev Khudanpur. 2010. Recurrent Neural Network Based Language Model. *Proceedings of Interspeech*.
- Tomas Mikolov, Wen-tau Yih, and Geoffrey Zweig. 2013. Linguistic Regularities in Continuous Space Word Representations. *Proceedings of NAACL-HLT*, pages 746–751.
- Andriy Mnih and Geoffrey Hinton. 2007. Three New Graphical Models for Statistical Language Modelling. In *Proceedings of the 24th International Conference on Machine learning*, pages 641–648. ACM.
- Andriy Mnih and Geoffrey E Hinton. 2009. A Scalable Hierarchical Distributed Language Model. In *Advances in Neural Information Processing Systems*, pages 1081–1088.
- Joseph Reisinger and Raymond J Mooney. 2010. Multi-Prototype Vector-Space Models of Word Meaning. In *Human Language Technologies: The 2010 Annual Conference of the North American Chapter of the Association for Computational Linguistics*, pages 109–117. Association for Computational Linguistics.
- Tony Rose, Mark Stevenson, and Miles Whitehead. 2002. The reuters corpus volume 1-from yesterday's news to tomorrow's language resources. In *LREC*, volume 2, pages 827–832.
- H. Sak, T. Güngör, and M. Saraçlar. 2008. Turkish language resources: Morphological parser, morphological disambiguator and web corpus. *Advances in natural language processing*, pages 417–427.
- Andreas Stolcke. 2002. SRILM – An extensible language modeling toolkit. In *Proceedings International Conference on Spoken Language Processing*, pages 257–286, November.
- Mihai Surdeanu, Richard Johansson, Adam Meyers, Lluís Màrquez, and Joakim Nivre. 2008. The conll-2008 shared task on joint parsing of syntactic and semantic dependencies. In *Proceedings of the Twelfth Conference on Computational Natural Language Learning*, pages 159–177. Association for Computational Linguistics.
- Erik F Tjong Kim Sang and Sabine Buchholz. 2000. Introduction to the conll-2000 shared task: Chunking. In *Proceedings of the 2nd workshop on Learning language in logic and the 4th conference on Computational natural language learning-Volume 7*, pages 127–132. Association for Computational Linguistics.
- Erik F Tjong Kim Sang and Fien De Meulder. 2003. Introduction to the conll-2003 shared task: Language-independent named entity recognition. In *Proceedings of the seventh conference on Natural language learning at HLT-NAACL 2003-Volume 4*, pages 142–147. Association for Computational Linguistics.

Joseph Turian, Lev Ratinov, and Yoshua Bengio. 2010. Word representations: A simple and general method for semi-supervised learning. In *Proceedings of the 48th Annual Meeting of the Association for Computational Linguistics*, pages 384–394. Association for Computational Linguistics.

Mehmet Ali Yatbaz, Enis Sert, and Deniz Yuret. 2012. Learning Syntactic Categories Using Paradigmatic Representations of Word Context. In *Proceedings of the 2012 Joint Conference on Empirical Methods in Natural Language Processing and Computational Natural Language Learning*, pages 940–951. Association for Computational Linguistics.

Mehmet Ali Yatbaz, Enis Sert, and Deniz Yuret. 2014. Unsupervised instance-based part of speech induction using probable substitutes. In *Proceedings of COLING 2014*. The COLING 2014 Organizing Committee.

Deniz Yuret. 2012. FASTSUBS: An Efficient and Exact Procedure for Finding the Most Likely Lexical Substitutes Based on an N-Gram Language Model. *Signal Processing Letters, IEEE*, 19(11):725–728, Nov.