

“Vivemos em uma sociedade extremamente dependente de ciência e tecnologia, onde quase ninguém sabe coisa alguma sobre ciência e tecnologia” (Carl Sagan).

Embeddings de Texto

Paulo Ricardo Lisboa de Almeida



0 Problema

“A inteligência artificial é um grande campo, e este é um grande livro. Tentamos explorar toda a extensão do assunto, que abrange lógica, probabilidade e matemática do contínuo, além de percepção, raciocínio, aprendizado, ação e, ainda, tudo o que se refere à eletrônica, ...”

0 Problema

“A inteligência artificial é um grande campo, e este é um grande livro. Tentamos explorar toda a extensão do assunto, que abrange lógica, probabilidade e matemática do contínuo, além de percepção, raciocínio, aprendizado, ação e, ainda, tudo o que se refere à eletrônica, ...”

Como segmentar o texto?

Como representar as palavras em vetores em um espaço?

Solução Ingênua

O que acontece se cortarmos o texto por espaços (e afins), e dar um número para cada palavra?

1 2 3 4 5 6 7 8 9 4 5
A inteligência artificial é um grande campo, e este é um
6 7 8 9 10 11 12 13 14
grande livro. Tentamos explorar toda a extensão do assunto,
15 16 17 18 8 19 13 20
que abrange lógica, probabilidade e matemática do contínuo

Solução Ingênua

1 2 3 4 5 6 7 8 9 4 5
A inteligência artificial é um grande campo, e este é um
6 7 8 9 10 11 12 13 14
grande livro. Tentamos explorar toda a extensão do assunto,
15 16 17 18 8 19 13 20
que abrange lógica, probabilidade e matemática do contínuo

Quebrar somente em espaços é suficiente? E vírgulas? E nomes compostos (Rio de Janeiro)? E demais caracteres de quebra?

Solução Ingênua

1 2 3 4 5 6 7 8 9 4 5
A inteligência artificial é um grande campo, e este é um
6 7 8 9 10 11 12 13 14
grande livro. Tentamos explorar toda a extensão do assunto,
15 16 17 18 8 19 13 20
que abrange lógica, probabilidade e matemática do contínuo

Quebrar somente em espaços é suficiente? E vírgulas? E nomes compostos (Rio de Janeiro)? E demais caracteres de quebra?

O que acontece se descobrirmos ou surgir uma palavra nova?

Solução Ingênua

1 2 3 4 5 6 7 8 9 4 5
A inteligência artificial é um grande campo, e este é um
6 7 8 9 10 11 12 13 14
grande livro. Tentamos explorar toda a extensão do assunto,
15 16 17 18 8 19 13 20
que abrange lógica, probabilidade e matemática do contínuo

Quebrar somente em espaços é suficiente? E vírgulas? E nomes compostos (Rio de Janeiro)? E demais caracteres de quebra?

O que acontece se descobrirmos ou surgir uma palavra nova?

Por que matemática está mais longe de lógica do que de probabilidade (pela distância L2)?

Solução Ingênua

1 2 3 4 5 6 7 8 9 4 5
A inteligência artificial é um grande campo, e este é um
6 7 8 9 10 11 12 13 14
grande livro. Tentamos explorar toda a extensão do assunto,
15 16 17 18 8 19 13 20
que abrange lógica, probabilidade e matemática do contínuo

Quebrar somente em espaços é suficiente? E vírgulas? E nomes compostos (Rio de Janeiro)? E demais caracteres de quebra?

O que acontece se descobrirmos ou surgir uma palavra nova?

Por que matemática está mais longe de lógica do que de probabilidade (pela distância L2)?

Palavras relacionadas deveriam estar próximas no espaço (cozinha, cozinhar, cozinheiro, cozido, ...).

Transformando em Tokens (*Tokenization*)

Poderíamos considerar apenas os caracteres individuais para separar as palavras.

Transformando em Tokens (*Tokenization*)

Poderíamos considerar apenas os caracteres individuais para separar as palavras.

Flexível, mas obriga a rede a descobrir como montar palavras válidas.

Token: grupo de poucos caracteres. Pode incluir:

Palavras inteiras, desde que sejam comuns.

Fragmentos de palavras.

Caracteres.

Ideia

1. Comece com uma lista de caracteres individuais.
2. Verifique a combinação de par de tokens mais comum.
3. Adicione a combinação no dicionário de tokens.
4. Enquanto não atingir critério de parada (máximo de tokens, número de ocorrências mínimo do token mais comum ...), volte para 2.

Sennrich, Rico, Barry Haddow, and Alexandra Birch. "Neural machine translation of rare words with subword units." 2015.

Neural Machine Translation of Rare Words with Subword Units

Rico Sennrich and Barry Haddow and Alexandra Birch

School of Informatics, University of Edinburgh
{f1101, sennrich}@inf.ed.ac.uk, {b.haddow, a.birch}@ed.ac.uk

Abstract

Neural machine translation (NMT) models typically operate with a fixed vocabulary, but translation is an open-vocabulary problem. Previous work addresses the handling of rare-of-vocabulary words by looking off to a dictionary. In this paper, we introduce a simple and more effective approach, making the NMT model capable of open-vocabulary translation by encoding rare and unknown words as sequences of subword units. This is based on the intuition that various word classes are translatable via similar sets of subwords. For instance, names (via character copying or transliteration), compounds (via concatenation), derivations, and regular inflected forms (via phonological and morphological regularities). We discuss the suitability of different word segmentation techniques, including simple character n-gram models and a segmentation based on the fast, open-vocabulary morpheme algorithm, and empirically show that subword models improve over look-off-dictionary baselines for the NMT V3 translation tasks English-German and English-Russian by up to 1.1 and 1.3 BLEU, respectively.

1 Introduction

Neural machine translation has recently shown impressive results (Kalchbrenner and Blunton, 2015; Sennrich et al., 2014; Bahdanau et al., 2015). However, the translation of rare words is an open problem. The vocabularies of neural models is typically limited to 30000-50000 words, but translation is an open-vocabulary prob-

lem, and especially for languages with productive word formation processes such as agglutinative and compounding, translation models require mechanisms that go below the word level. As an example, consider compounds such as the German *Abwasser/Abwasserbehandlung* “wastage water treatment plant”, for which a segmented, variable-length representation is intuitively more appealing than encoding the word as a fixed-length vector.

For word-level NMT models, the translation of out-of-vocabulary words has been addressed through a look-off-to-a-dictionary look-up (Barrat et al., 2015; Loeig et al., 2015b). We note that such techniques make assumptions that often do not hold true in practice. For instance, there is not always a 1:1 correspondence between source and target words because of variance in the degree of morphological complexity between languages, like as an increasingly compounding example, *Abwasser*, word-level models are unable to translate or generate missing words. Copying subwords instead of the target text, as done by (Barrat et al., 2015; Loeig et al., 2015b), is a reasonable strategy for names, but morphological changes and transliteration in other regions, especially of compound affixes.

We investigate NMT models that operate on the level of subword units. Our main goal is to model open-vocabulary translation in the NMT network itself, without requiring a look-off model for rare words. In addition to making the translation process simpler, we also find that the subword models achieve better accuracy in the translation of rare words than large-vocabulary models and look-off-dictionaries, and are able to productively generate new words that were not seen in training time. Our analysis shows that the neural networks are able to learn compounding and transliteration from subword representations.

This paper has two main contributions:

- We show that open-vocabulary neural ma-

¹The research presented in this publication was conducted as part of the project “Neural Machine Translation of Rare Words with Subword Units” funded by the Scottish Funding Council.

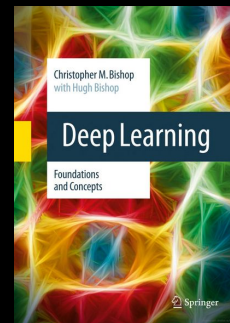
Exemplo

Tokens: a c d e f i k l o p P r s t

Peter Piper picked a peck of pickled peppers

1. Comece com uma lista de caracteres individuais.
2. Verifique a combinação de par de tokens mais comum.
3. Adicione a combinação no dicionário de tokens.
4. Enquanto não atingir critério de parada.

Bishop, C. M., Bishop, H.
Deep Learning: Foundations
and Concepts. 2023.



Exemplo

Tokens: a c d e f i k l o p P r s t **pe**

Peter Piper picked a peck of pickled peppers

Peter Pi**per** picked a **pe**ck of pickled **pe**ppers

1. Comece com uma lista de caracteres individuais.
2. Verifique a combinação de par de tokens mais comum.
3. Adicione a combinação no dicionário de tokens.
4. Enquanto não atingir critério de parada.

Exemplo

Tokens: a c d e f i k l o p P r s t **pe** **ck**

Peter Piper picked a peck of pickled peppers

Peter **Pi**per picked a **pe**ck of pickled **pe**ppers

Peter **Pi**per **pick**ed a **pe**ck of **pick**led **pe**ppers

1. Comece com uma lista de caracteres individuais.
2. **Verifique a combinação de par de tokens mais comum.**
3. **Adicione a combinação no dicionário de tokens.**
4. Enquanto não atingir critério de parada.

Exemplo

Tokens: a c d e f i k l o p P r s t **pe** **ck** **pi**

Peter Piper picked a peck of pickled peppers

Peter Pi**per** picked a **pe**ck of pickled **pe**ppers

Peter Pi**per** pick**e**d a **pe**ck of pick**l**ed **pe**ppers

Peter Pi**per** **p**ick**e**d a **pe**ck of **p**ick**l**ed **pe**ppers

1. Comece com uma lista de caracteres individuais.
2. Verifique a combinação de par de tokens mais comum.
3. Adicione a combinação no dicionário de tokens.
4. Enquanto não atingir critério de parada.

Exemplo

Tokens: a c d e f i k l o p P r s t **pe** **ck** **pi** **per**

Peter Piper picked a peck of pickled peppers

Peter Pi**per** picked a **pe**ck of pickled **pe**ppers

Peter Pi**per** pic**ke**d a **pe**ck of pic**kl**ed **pe**ppers

Peter Pi**per** pi**ck**ed a **pe**ck of pi**ck**led **pe**ppers

Peter Pi**per** pi**ck**ed a **pe**ck of pi**ck**led **pe**ppers

1. Comece com uma lista de caracteres individuais.
2. Verifique a combinação de par de tokens mais comum.
3. Adicione a combinação no dicionário de tokens.
4. Enquanto não atingir critério de parada.

Problemas

Como segmentar o texto? 

Como representar as palavras em vetores em um espaço?

Vocabulário: a c d e f i k l o p P r s t pe ck pi per

Word2Vec

Peter Piper picked a peck of pickled peppers

	a	c	d	e	f	i	k	l	o	p	P	r	s	t	pe	ck	pi	per
x_0	=	[0	0	0	0	0	0	0	0	0	1	0	0	0	0	0	0	0]
x_1	=	[0	0	0	1	0	0	0	0	0	0	0	0	0	0	0	0	0]
x_2	=	[0	0	0	0	0	0	0	0	0	0	0	0	1	0	0	0	0]
...																		

One-hot encoding

Word2Vec

Vocabulário: a c d e f i k l o p P r s t pe ck pi per

Peter Piper picked a peck of pickled peppers

	a	c	d	e	f	i	k	l	o	p	P	r	s	t	pe	ck	pi	per
$\mathbf{x}_0$	=	[0	0	0	0	0	0	0	0	0	1	0	0	0	0	0	0	0]
$\mathbf{x}_1$	=	[0	0	0	1	0	0	0	0	0	0	0	0	0	0	0	0	0]
$\mathbf{x}_2$	=	[0	0	0	0	0	0	0	0	0	0	0	0	1	0	0	0	0]
...																		

Utilizamos uma matriz de embedding $\mathbf{E}$, de tamanho $D \times K$.

K é o tamanho do Vocabulário

D é o número de dimensões do embedding.

Cada vetor é projetado no espaço de embedding através de $\mathbf{v}_i = \mathbf{E}\mathbf{x}_i$.

Exemplo

$$\mathbf{E} = \begin{bmatrix} x_{1,1} & x_{1,2} & \cdots & x_{1,K} \\ x_{2,1} & x_{2,2} & \cdots & x_{2,K} \\ \vdots & \vdots & \ddots & \vdots \\ x_{D,1} & x_{D,2} & \cdots & x_{D,K} \end{bmatrix} \quad \mathbf{x} = \begin{bmatrix} x_1 \\ x_2 \\ \vdots \\ x_K \end{bmatrix}$$

Exemplo

$$\mathbf{E} = \begin{bmatrix} x_{1,1} & x_{1,2} & \cdots & x_{1,K} \\ x_{2,1} & x_{2,2} & \cdots & x_{2,K} \\ \vdots & \vdots & \ddots & \vdots \\ x_{D,1} & x_{D,2} & \cdots & x_{D,K} \end{bmatrix} \rightarrow \mathbf{x} = \begin{bmatrix} x_1 \\ x_2 \\ \vdots \\ x_K \end{bmatrix}$$

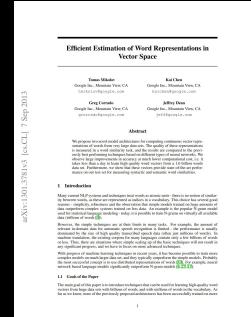
Devido ao one-hot encoding, o vetor resultante é dado pela coluna correspondente em $\mathbf{E}$, sem precisar realizar cálculos!

Word2vec

Uma técnica clássica para aprender a matriz $\mathbf{E}$ é utilizar o word2vec.

- MLP de duas camadas.
 - Entrada e saída do tamanho do vocabulário.

Mikolov, T., Chen, K., Corrado, G., & Dean, J. Efficient estimation of word representations in vector space. 2013.

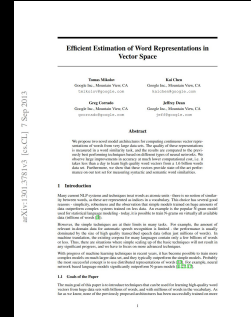


Word2vec

Uma técnica clássica para aprender a matriz $\mathbf{E}$ é utilizar o word2vec.

- MLP de duas camadas.
 - Entrada e saída do tamanho do vocabulário.
- Conjunto de treinamento construído considerando uma janela deslizante de palavras de tamanho M .
 - Comum utilizar $M=5$.

Mikolov, T., Chen, K., Corrado, G., & Dean, J. Efficient estimation of word representations in vector space. 2013.



Word2vec - Treinamento

Continuous bag of words

Considerando $M=5$

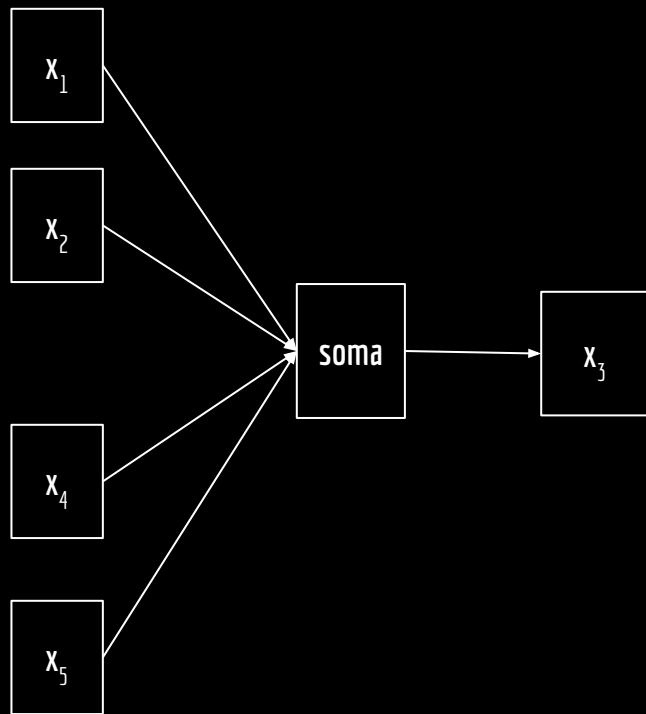
Dada uma string de treinamento, são enviados os tokens de índice 1,2,4 e 5.

O objetivo de rede é prever o token central (índice 3).

A função de loss é a soma das entradas.

A soma deve ser igual ao token central.

Word2vec - Treinamento



Self-supervised Learning

Treinamento auto-supervisionado.

Não precisamos rotular os dados, já que o treino é realizado ocultando parte do dado existente.

Mas note que o texto original precisa fazer sentido. Logo, alguém precisou escrever!

Word2vec

Depois de treinada, a transposta dos pesos da segunda camada da rede são usados como a Matriz **E**.

Word2vec

Depois de treinada, a transposta dos pesos da segunda camada da rede são usados como a Matriz **E**.

Os pesos da primeira camada são usados se usarmos o método skip-grams.

Pesquise como funciona o método baseado em skip-grams.

Semântica

Paris está para França assim como Itália está para ...

$$v(\text{Paris}) - v(\text{França}) + v(\text{Itália}) \approx v(\text{Roma})$$

Visualização

Veja um exemplo de embeddings em <https://projector.tensorflow.org>

Problemas

Como segmentar o texto? 

Como representar as palavras em vetores em um espaço? 

Uso

Os embeddings são comumente usados atualmente como um pré-processamento.

Entrada para uma rede profunda, como redes baseadas em transformers (próximas aulas).

Podemos usar de forma fixa (pré-treinado), ou como uma camada inicial treinável.

Ideias

As ideias são aplicáveis em outros tipos de dados.

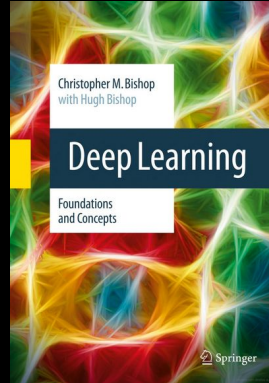
- Imagens.
- Áudio.
- Códigos fonte.
- ...

Exercícios

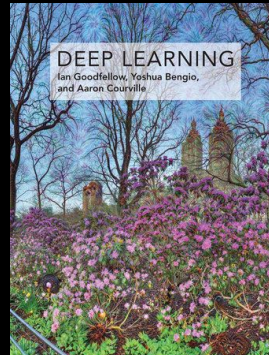
1. Execute os exemplos disponibilizados no Google Colab.

Referências

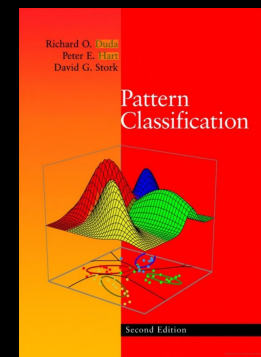
Bishop, C. M., Bishop, H. Deep Learning: Foundations and Concepts. 2023.



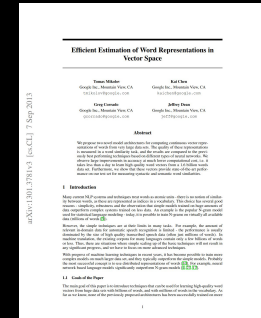
Goodfellow, I., Bengio, Y., Courville, A. Deep Learning. 2016.



Duda, R. O., Hart, P. E., Stork, D. G. Pattern Classification. 2012.



Mikolov, T., Chen, K., Corrado, G., & Dean, J. Efficient estimation of word representations in vector space. 2013.



JURAFSKY, Daniel; MARTIN, James H. Speech and Language Processing: An Introduction to Natural Language Processing, Computational Linguistics, and Speech Recognition with Language Models. 3. Disponível em: <https://web.stanford.edu/~jurafsky/slp3>.

Licença

Esta obra está licenciada com uma Licença [Creative Commons Atribuição 4.0 Internacional](https://creativecommons.org/licenses/by/4.0/).