

Machine Learning, Artificial Intelligence, and Natural Language Processing

Marcelo C. Medeiros

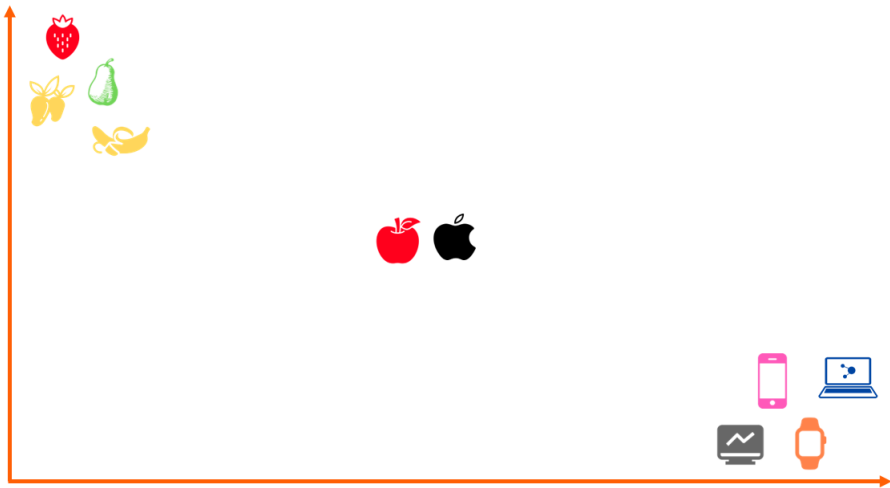
Department of Economics
The University of Illinois at Urbana-Champaign

Attention

Transformers

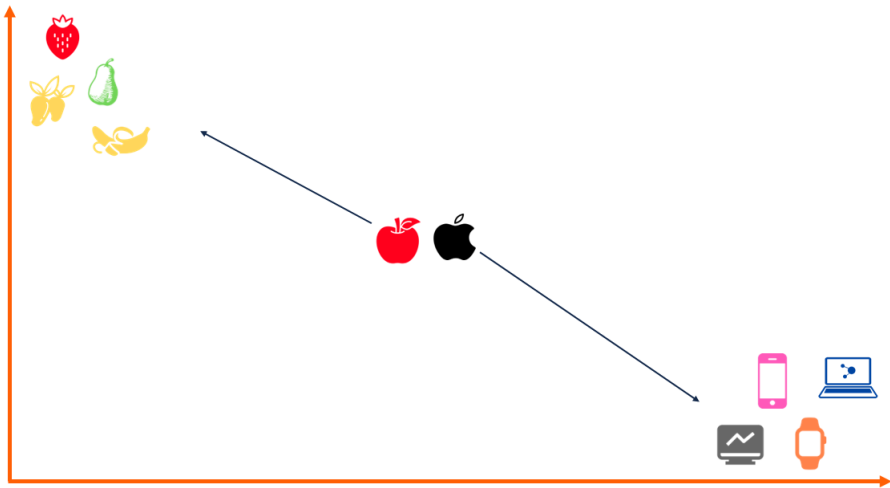
Attention Mechanism

Searching for a better word embedding



Attention Mechanism

Searching for a better word embedding



Attention Mechanism

Searching for a better word embedding



Attention Is All You Need

Ashish Vaswani*
Google Brain
avaswani@google.com

Noam Shazeer*
Google Brain
noam@google.com

Niki Parmar*
Google Research
nikip@google.com

Jakob Uszkoreit*
Google Research
usz@google.com

Llion Jones*
Google Research
llion@google.com

Aidan N. Gomez*[†]
University of Toronto
aidan@cs.toronto.edu

Łukasz Kaiser*
Google Brain
lukaszkaizer@google.com

Illia Polosukhin*[‡]
illia.polosukhin@gmail.com

Attention Mechanism

The problem of static embeddings - like Word2Vec

- ▶ They are static – The embedding for a word does not reflect how its meaning changes in context

Attention Mechanism

The problem of static embeddings - like Word2Vec

- ▶ They are static – The embedding for a word does not reflect how its meaning changes in context
- ▶ For example,

The chicken didn't cross the road because
it was too tired

Attention Mechanism

The problem of static embeddings - like Word2Vec

- ▶ They are static – The embedding for a word does not reflect how its meaning changes in context
- ▶ For example,

The chicken didn't cross the road because
it was too tired

- * What is the meaning represented in the static embedding for “it”?

Attention Mechanism

The problem of static embeddings - like Word2Vec

- ▶ They are static – The embedding for a word does not reflect how its meaning changes in context
- ▶ For example,

The chicken didn't cross the road because
it was too tired

* What is the meaning represented in the static embedding for “*it*”?

- ▶ Consider now

The chicken didn't cross the road because
it was too wide

Attention Mechanism

The problem of static embeddings - like Word2Vec

- ▶ They are static – The embedding for a word does not reflect how its meaning changes in context
- ▶ For example,

The chicken didn't cross the road because
it was too tired

* What is the meaning represented in the static embedding for “*it*”?

- ▶ Consider now

The chicken didn't cross the road because
it was too wide

* Is the meaning the same?

Attention Mechanism

Contextual embeddings

- ▶ **Problem:** The representation of the meaning of a word should be different in different contexts
- ▶ **Contextual Embedding:** each word has a different vector that expresses different meanings depending on the surrounding words
- ▶ **Solution:** Attention mechanism.

The chicken didn't cross the road because it ...

- ▶ What should be the properties of “it”?

The chicken didn't cross the road because it was too **tired**

The chicken didn't cross the road because it was too **wide**

- ▶ At this point in the sentence, it's probably referring to either the chicken or the street

Attention Mechanism

Basic idea

- ▶ Build up the contextual embedding from a word by selectively integrating information from all the **neighboring words**
- ▶ A word “attends to” some neighboring words more than others

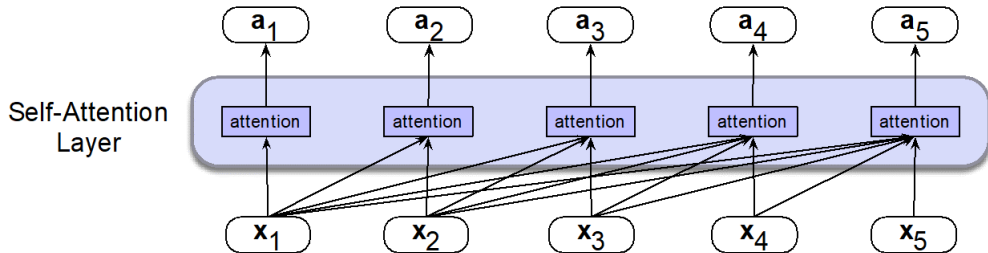
Attention Mechanism

Definition

- ▶ Attention is a mechanism for helping compute the embedding for a token by selectively attending to and integrating information from surrounding tokens
- ▶ Mathematically: a method for computing a (“smart”) weighted sum of vectors

Attention Mechanism

Left-to-right



Attention Mechanism

Simplified version

- ▶ Sequence of token embeddings: $\mathbf{x}_1, \mathbf{x}_2, \mathbf{x}_3, \mathbf{x}_4, \mathbf{x}_5, \mathbf{x}_6, \mathbf{x}_7, \mathbf{x}_i$
- ▶ Measure distance between two tokens: **inner product**

$$\text{score}(\mathbf{x}_i, \mathbf{x}_j) = \mathbf{x}_i' \mathbf{x}_j$$

- ▶ Transform the score into weights between 0 and 1:

$$\alpha_{ij} = \text{softmax}(\text{score}(\mathbf{x}_i, \mathbf{x}_j)) = \frac{e^{\text{score}(\mathbf{x}_i, \mathbf{x}_j)}}{\sum_{k=1}^8 e^{\text{score}(\mathbf{x}_i, \mathbf{x}_k)}}, \quad j \leq i$$

- ▶ Attention:

$$\mathbf{a}_i = \sum_{j \leq i} \alpha_{ij} \mathbf{x}_j$$

- ▶ **High-level idea:** instead of using embedding vectors directly, we will represent three separate roles each vector x_i plays:
 1. **query:** element that is being compared to the preceding inputs
 2. **key:** preceding input that is being compared to the current element to determine a similarity
 3. **value:** a value of a preceding element that gets weighted and summed

Attention Mechanism

Attention head

- ▶ Matrices to project each vector $\mathbf{x}_i$ into a representation of its role as query, key, or value

- * Query: Q

- * Key: K

- * Value: V

- ▶ Therefore,

$$\mathbf{q}_i = Q\mathbf{x}_i \quad \mathbf{k}_i = K\mathbf{x}_i \quad \mathbf{v}_i = V\mathbf{x}_i$$

- ▶ The three matrices are learned from data

Attention Mechanism

Attention head

- ▶ Take the representation in terms of query, key, and value matrices

$$\mathbf{q}_i = \mathbf{Q}\mathbf{x}_i \quad \mathbf{k}_i = \mathbf{K}\mathbf{x}_i \quad \mathbf{v}_i = \mathbf{V}\mathbf{x}_i$$

- ▶ Similarity (score) between $\mathbf{x}_i$ and prior token $\mathbf{x}_j$: $\mathbf{q}_i'\mathbf{k}_j/\sqrt{d}$

- ▶ d is the dimension of the embeddings

- ▶ Therefore,

$$\alpha_{ij} = \text{softmax}(\mathbf{q}_i'\mathbf{k}_j), \quad j \leq i$$

- ▶ Attention:

$$\mathbf{a}_i = \sum_{j \leq i} \alpha_{ij} \mathbf{v}_j$$

Attention Mechanism

Multi-head attention

- ▶ Instead of one attention head, we'll have lots of them
- ▶ **Intuition:** each head might be attending to the context for different purposes

Attention Mechanism

Multi-head attention

- ▶ For each head h :

$$\mathbf{q}_i^h = \mathbf{Q}^h \mathbf{x}_i \quad \mathbf{k}_i^h = \mathbf{K}^h \mathbf{x}_i \quad \mathbf{v}_i^h = \mathbf{V}^h \mathbf{x}_i$$

- ▶ Similarity (score) between $\mathbf{x}_i^h$ and prior token $\mathbf{x}_j$: $\mathbf{q}_i^{h'} \mathbf{k}_j^h / \sqrt{d_h}$

- ▶ $d_h = d/H$, H is the number of heads

- ▶ Therefore,

$$\alpha_{ij}^h = \text{softmax}(\mathbf{q}_i^{h'} \mathbf{k}_j^h), \quad \text{and} \quad \mathbf{a}_i^h = \sum_{j \leq i} \alpha_{ij}^h \mathbf{v}_j^h$$

- ▶ Hence,

$$\mathbf{a}_i = \left(\mathbf{a}_i^1 \cdots \mathbf{a}_i^H \right) \mathbf{O}$$

Attention = contextual embedding

Attention

Transformers

- ▶ **Goal:** given a sequence of words, predict which word will come next
- ▶ It is based on NN architectures plus attention mechanism

Transformers

The big picture

