

CSCI567 Machine Learning (Fall 2025)

Haipeng Luo

University of Southern California

Nov 7, 2025

Administration

Will discuss HW3 solutions in today's discussion session.

HW4 (last homework) will be released soon.

Outline

- 1 Review of last lecture
- 2 Recurrent Neural Network
- 3 Transformers

Outline

- 1 Review of last lecture
- 2 Recurrent Neural Network
- 3 Transformers

Hidden Markov Models

Model parameters:

- initial distribution**

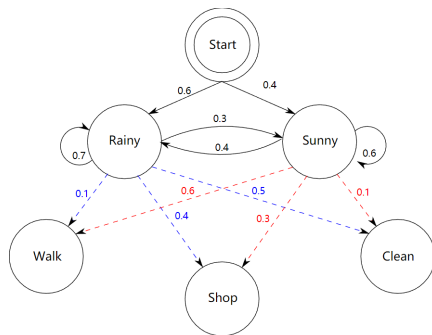
$$P(Z_1 = s) = \pi_s$$

- transition distribution**

$$P(Z_{t+1} = s' \mid Z_t = s) = a_{s,s'}$$

- emission distribution**

$$P(X_t = o \mid Z_t = s) = b_{s,o}$$



Baum–Welch algorithm

Step 0 Initialize the parameters $(\pi, \mathbf{A}, \mathbf{B})$

Step 1 (E-Step) Fixing the parameters, **compute forward and backward messages for all sample sequences**, then use these to compute $\gamma_s^{(n)}(t)$ and $\xi_{s,s'}^{(n)}(t)$ for each n, t, s, s' .

Step 2 (M-Step) Update parameters:

$$\pi_s \propto \sum_n \gamma_s^{(n)}(1), \quad a_{s,s'} \propto \sum_n \sum_{t=1}^{T-1} \xi_{s,s'}^{(n)}(t), \quad b_{s,o} \propto \sum_n \sum_{t:x_t=o} \gamma_s^{(n)}(t)$$

Step 3 Return to Step 1 if not converged

Viterbi Algorithm

Viterbi Algorithm

For each $s \in [S]$, compute $\delta_s(1) = \pi_s b_{s,x_1}$.

For each $t = 2, \dots, T$,

- for each $s \in [S]$, compute

$$\delta_s(t) = b_{s,x_t} \max_{s'} a_{s',s} \delta_{s'}(t-1)$$

$$\Delta_s(t) = \operatorname{argmax}_{s'} a_{s',s} \delta_{s'}(t-1)$$

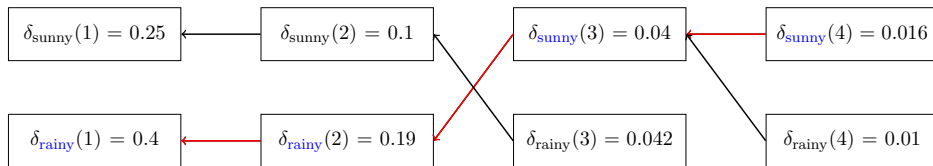
Backtracking: let $z_T^* = \operatorname{argmax}_s \delta_s(T)$.

For each $t = T, \dots, 2$: set $z_{t-1}^* = \Delta_{z_t^*}(t)$.

Output the most likely path $z_1^*, \dots, z_T^*$.

Example

Arrows represent $\Delta_s(t)$, backtracking = follow the arrows.



The most likely path is **“rainy, rainy, sunny, sunny”**.

Outline

- 1 Review of last lecture
- 2 Recurrent Neural Network
 - RNN: model
 - RNN: training and testing
- 3 Transformers

Recall: language models via HMM

- today the _____
- today the price _____
- today the price of _____
- today the price of gold _____

final result: *today the price of gold per ton, while production of shoe lasts and shoe industry, the bank intervened just after it considered and rejected an IMF demand to rebuild depleted European stocks.*

Surprisingly grammatical! but *incoherent...*

company	0.153
bank	0.153
price	0.077
italian	0.039
emirate	0.039
...	

of	0.308
for	0.050
it	0.046
to	0.046
is	0.031
...	

the	0.072
18	0.043
oil	0.043
its	0.036
gold	0.018
...	

How to improve this?

Key ideas for improvement:

How to improve this?

Key ideas for improvement:

- represent **words as vectors**, enabling **differentiable operations**

How to improve this?

Key ideas for improvement:

- represent **words as vectors**, enabling **differentiable operations**
- flexible **sequence to sequence** architecture (not just vector to vector)

How to improve this?

Key ideas for improvement:

- represent **words as vectors**, enabling **differentiable operations**
- flexible **sequence to sequence** architecture (not just vector to vector)
- **shared components** (just like **filters** in CNN)

How to improve this?

Key ideas for improvement:

- represent **words as vectors**, enabling **differentiable operations**
- flexible **sequence to sequence** architecture (not just vector to vector)
- **shared components** (just like **filters** in CNN)

Recurrent Neural Network (RNN) is one solution (popular before transformers).

Acknowledgements

Very useful resources:

- RNN cheatsheet from Stanford CS 230
 - <https://stanford.edu/~shervine/teaching/cs-230/cheatsheet-recurrent-neural-networks>
- Visualizing a tiny RNN
 - <https://joshvarty.github.io/VisualizingRNNs/>
- Character-level RNN
 - <https://karpathy.github.io/2015/05/21/rnn-effectiveness/>

Words as vectors

Simplest approach: **one-hot sparse encoding**

Words as vectors

Simplest approach: **one-hot sparse encoding**

- suppose there are d words in the vocabulary

Words as vectors

Simplest approach: **one-hot sparse encoding**

- suppose there are d words in the vocabulary
- represent the i -th of them by the d -dimensional basis vector

$$e_i = (0, \dots, 0, \underset{\substack{\uparrow \\ i\text{-th entry}}}{1}, 0, \dots, 0) \in \mathbb{R}^d$$

Words as vectors

Simplest approach: **one-hot sparse encoding**

- suppose there are d words in the vocabulary
- represent the i -th of them by the d -dimensional basis vector

$$e_i = (0, \dots, 0, \underset{\substack{\uparrow \\ i\text{-th entry}}}{1}, 0, \dots, 0) \in \mathbb{R}^d$$

Issues: does not convey any semantic meanings

Words as vectors: embedding

Word embedding: similar words are closer in their vector representation

Words as vectors: embedding

Word embedding: similar words are closer in their vector representation

- popular approaches: **word2vec**, **GloVe** (see project)

Words as vectors: embedding

Word embedding: similar words are closer in their vector representation

- popular approaches: **word2vec**, **GloVe** (see project)
- can even perform meaningful **algebraic operations**

Words as vectors: embedding

Word embedding: similar words are closer in their vector representation

- popular approaches: **word2vec**, **GloVe** (see project)
- can even perform meaningful **algebraic operations**
 - example: *man is to woman as king is to?*

Words as vectors: embedding

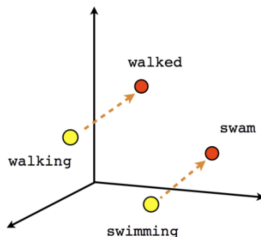
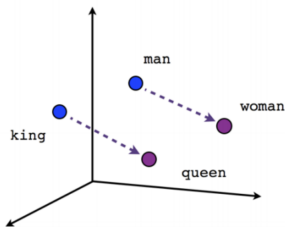
Word embedding: similar words are closer in their vector representation

- popular approaches: **word2vec**, **GloVe** (see project)
- can even perform meaningful **algebraic operations**
 - example: *man is to woman as king is to?*
 - can be answered by finding the word with the closest embedding to $\text{vec}(\text{woman}) - \text{vec}(\text{man}) + \text{vec}(\text{king})$, which happens to be “**queen**”

Words as vectors: embedding

Word embedding: similar words are closer in their vector representation

- popular approaches: **word2vec**, **GloVe** (see project)
- can even perform meaningful **algebraic operations**
 - example: *man is to woman as king is to?*
 - can be answered by finding the word with the closest embedding to $\text{vec}(\text{woman}) - \text{vec}(\text{man}) + \text{vec}(\text{king})$, which happens to be “**queen**”



Unifying two approaches

Let $\mathbf{x} \in \mathbb{R}^d$ be the one-hot encoding of a word, and matrix $\mathbf{E} \in \mathbb{R}^{d_e \times d}$ be some **embedding matrix**, then $\mathbf{E}\mathbf{x}$ is the embedding for this word

Unifying two approaches

Let $\mathbf{x} \in \mathbb{R}^d$ be the one-hot encoding of a word, and matrix $\mathbf{E} \in \mathbb{R}^{d_e \times d}$ be some **embedding matrix**, then $\mathbf{E}\mathbf{x}$ is the embedding for this word

- $\mathbf{E}$ can be fixed (i.e., from word2vec or GloVe), where the i -th column is the embedding for the i -th word

Unifying two approaches

Let $\mathbf{x} \in \mathbb{R}^d$ be the one-hot encoding of a word, and matrix $\mathbf{E} \in \mathbb{R}^{d_e \times d}$ be some **embedding matrix**, then $\mathbf{E}\mathbf{x}$ is the embedding for this word

- $\mathbf{E}$ can be fixed (i.e., from word2vec or GloVe), where the i -th column is the embedding for the i -th word
- or $\mathbf{E}$ can be *learned* (e.g., via backpropagation in DL pipeline),

Unifying two approaches

Let $\mathbf{x} \in \mathbb{R}^d$ be the one-hot encoding of a word, and matrix $\mathbf{E} \in \mathbb{R}^{d_e \times d}$ be some **embedding matrix**, then $\mathbf{E}\mathbf{x}$ is the embedding for this word

- $\mathbf{E}$ can be fixed (i.e., from word2vec or GloVe), where the i -th column is the embedding for the i -th word
- or $\mathbf{E}$ can be *learned* (e.g., via backpropagation in DL pipeline), making it application specific (common especially if data are huge)

Unifying two approaches

Let $\mathbf{x} \in \mathbb{R}^d$ be the one-hot encoding of a word, and matrix $\mathbf{E} \in \mathbb{R}^{d_e \times d}$ be some **embedding matrix**, then $\mathbf{E}\mathbf{x}$ is the embedding for this word

- $\mathbf{E}$ can be fixed (i.e., from word2vec or GloVe), where the i -th column is the embedding for the i -th word
- or $\mathbf{E}$ can be *learned* (e.g., via backpropagation in DL pipeline), making it application specific (common especially if data are huge)

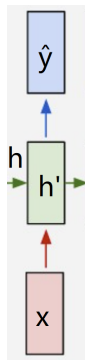
In the remaining, we simply use one-hot representation, but keep in mind it could be passed through some $\mathbf{E}$ implicitly

A recurrent layer

from $\hat{\mathbf{y}} = \mathbf{f}(\mathbf{x})$ to

A recurrent layer

from $\hat{y} = f(x)$ to $(\hat{y}, h') = f(x, h)$



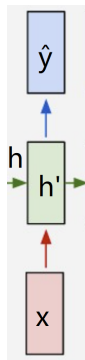
A recurrent layer

from $\hat{y} = f(x)$ to $(\hat{y}, h') = f(x, h)$

- h is “hidden state” (like HMM), updated via

$$h' = \sigma(W h + U x + b_h)$$

where σ is an activation function



A recurrent layer

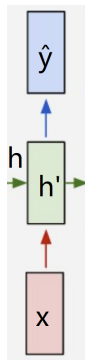
from $\hat{y} = f(x)$ to $(\hat{y}, h') = f(x, h)$

- h is “hidden state” (like HMM), updated via

$$h' = \sigma(Wh + Ux + b_h)$$

where σ is an activation function

- $\hat{y} = Vh' + b_y$ is the output



A recurrent layer

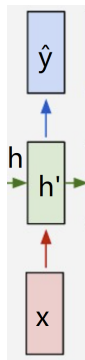
from $\hat{y} = f(x)$ to $(\hat{y}, h') = f(x, h)$

- h is “hidden state” (like HMM), updated via

$$h' = \sigma(Wh + Ux + b_h)$$

where σ is an activation function

- $\hat{y} = Vh' + b_y$ is the output
- $x, \hat{y} \in \mathbb{R}^d, h, h' \in \mathbb{R}^{d_h}$



A recurrent layer

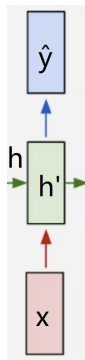
from $\hat{y} = f(x)$ to $(\hat{y}, h') = f(x, h)$

- h is “hidden state” (like HMM), updated via

$$h' = \sigma(W h + U x + b_h)$$

where σ is an activation function

- $\hat{y} = V h' + b_y$ is the output
- $x, \hat{y} \in \mathbb{R}^d, h, h' \in \mathbb{R}^{d_h}$
- weight matrices $W \in \mathbb{R}^{d_h \times d_h}, U \in \mathbb{R}^{d_h \times d}, V \in \mathbb{R}^{d \times d_h}$



A recurrent layer

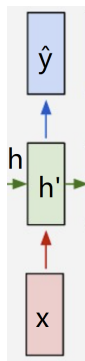
from $\hat{y} = f(x)$ to $(\hat{y}, h') = f(x, h)$

- h is “hidden state” (like HMM), updated via

$$h' = \sigma(W h + U x + b_h)$$

where σ is an activation function

- $\hat{y} = V h' + b_y$ is the output
- $x, \hat{y} \in \mathbb{R}^d, h, h' \in \mathbb{R}^{d_h}$
- weight matrices $W \in \mathbb{R}^{d_h \times d_h}, U \in \mathbb{R}^{d_h \times d}, V \in \mathbb{R}^{d \times d_h}$
- bias terms $b_h \in \mathbb{R}^{d_h}, b_y \in \mathbb{R}^d$

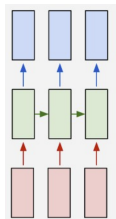


Recurrent layer applied recursively

Given a **sequence** $x_1, x_2, \dots$, can apply f **recursively**:

Recurrent layer applied recursively

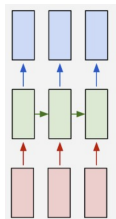
Given a **sequence** $x_1, x_2, \dots$, can apply f **recursively**:



- $h_0 = 0$
- $(\hat{y}_1, h_1) = f(x_1, h_0)$

Recurrent layer applied recursively

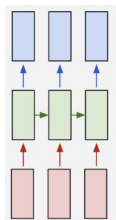
Given a **sequence** $x_1, x_2, \dots$, can apply f **recursively**:



- $h_0 = 0$
- $(\hat{y}_1, h_1) = f(x_1, h_0)$
- $(\hat{y}_2, h_2) = f(x_2, h_1)$
- $\dots$

Recurrent layer applied recursively

Given a **sequence** $x_1, x_2, \dots$, can apply f **recursively**:

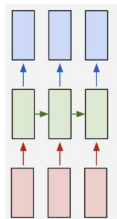


- $h_0 = 0$
- $(\hat{y}_1, h_1) = f(x_1, h_0)$
- $(\hat{y}_2, h_2) = f(x_2, h_1)$
- $\dots$

This is **one** recurrent layer unfolded (over steps), *not many different layers*.

Recurrent layer applied recursively

Given a **sequence** $x_1, x_2, \dots$, can apply f **recursively**:



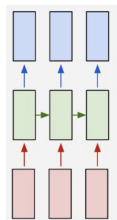
- $h_0 = 0$
- $(\hat{y}_1, h_1) = f(x_1, h_0)$
- $(\hat{y}_2, h_2) = f(x_2, h_1)$
- $\dots$

This is **one** recurrent layer unfolded (over steps), *not many different layers*.

The same f (i.e, W, U, V, b_h, b_y) is **shared** in all steps (similar to CNN's filters shared across different spatial locations).

Recurrent layer applied recursively

Given a **sequence** $x_1, x_2, \dots$, can apply f **recursively**:



- $h_0 = 0$
- $(\hat{y}_1, h_1) = f(x_1, h_0)$
- $(\hat{y}_2, h_2) = f(x_2, h_1)$
- $\dots$

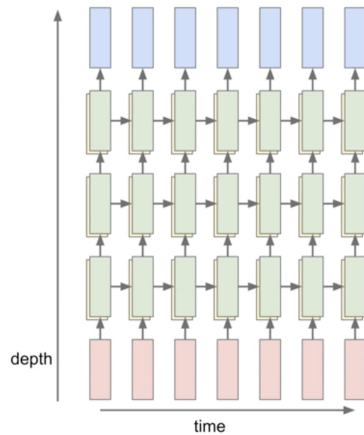
This is **one** recurrent layer unfolded (over steps), *not many different layers*.

The same f (i.e, W, U, V, b_h, b_y) is **shared** in all steps (similar to CNN's filters shared across different spatial locations).

Hidden state h_t summarizes information up to step t .

Making it “deep”

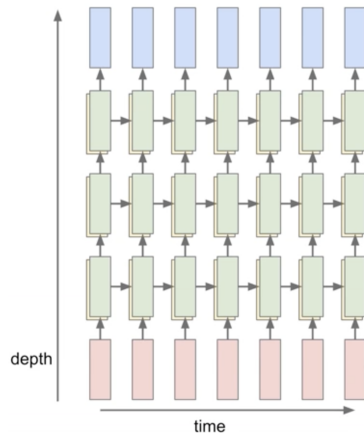
Stack multiple recurrent layers:



Making it “deep”

Stack multiple recurrent layers:

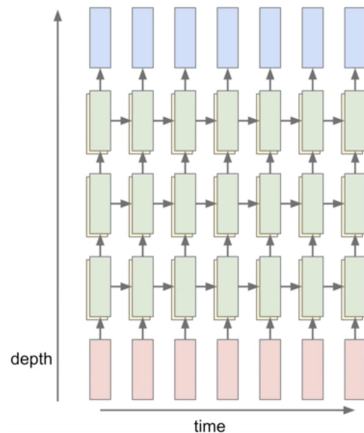
- hidden states become the inputs of the next layer



Making it “deep”

Stack multiple recurrent layers:

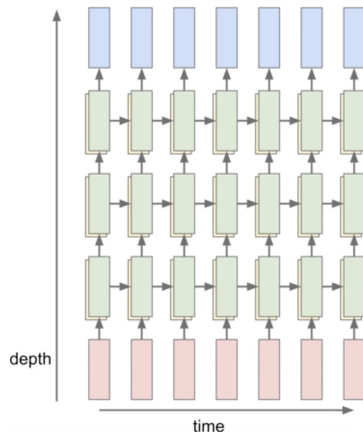
- hidden states become the inputs of the next layer
- different layers learn different W, U, b_h



Making it “deep”

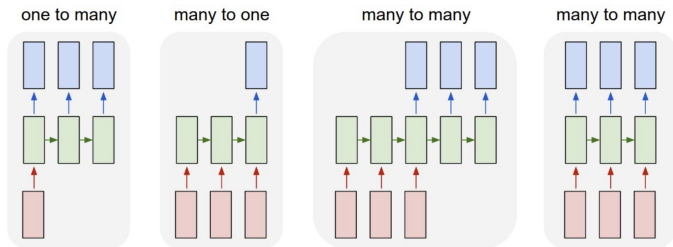
Stack multiple recurrent layers:

- hidden states become the inputs of the next layer
- different layers learn different W, U, b_h
- last layer learns V, b_y and output $\hat{y}$



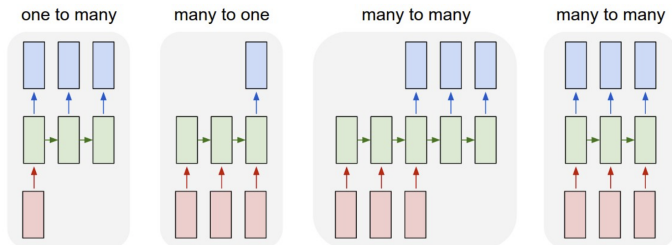
A flexible sequence-to-sequence model

Many possible structures and applications:



A flexible sequence-to-sequence model

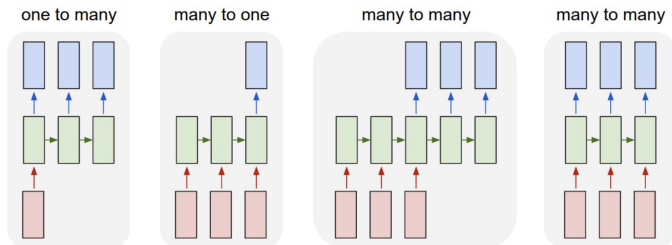
Many possible structures and applications:



- **one-to-many:** image captioning

A flexible sequence-to-sequence model

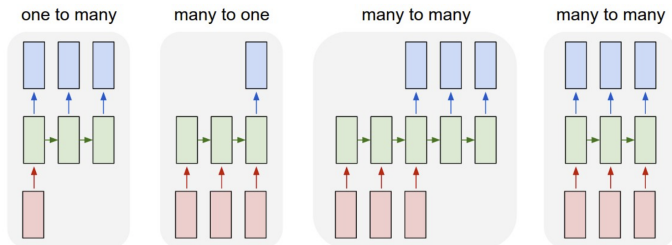
Many possible structures and applications:



- **one-to-many:** image captioning
- **many-to-one:** sentiment classification

A flexible sequence-to-sequence model

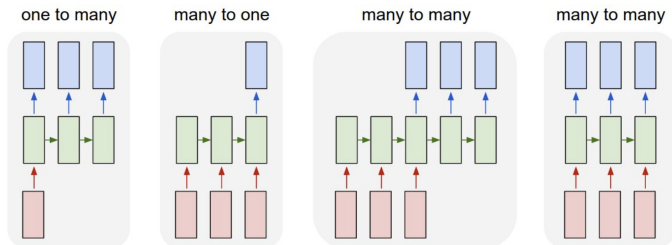
Many possible structures and applications:



- **one-to-many:** image captioning
- **many-to-one:** sentiment classification
- **many-to-many:** machine translation, question answering

A flexible sequence-to-sequence model

Many possible structures and applications:



- **one-to-many:** image captioning
- **many-to-one:** sentiment classification
- **many-to-many:** machine translation, question answering
- **(aligned) many-to-many:** POS tagging, name entity recognition

How to train an RNN

Take **text generation** (unsupervised learning) as an example:

How to train an RNN

Take **text generation** (unsupervised learning) as an example:

- given a corpus, train an RNN that learns $P(\mathbf{x}_t \mid \mathbf{x}_{1:t-1})$

How to train an RNN

Take **text generation** (unsupervised learning) as an example:

- given a corpus, train an RNN that learns $P(\mathbf{x}_t \mid \mathbf{x}_{1:t-1})$

For each sequence $\mathbf{x}_1, \dots, \mathbf{x}_T \in \mathbb{R}^d$ (one-hot representation) in the corpus

- feed $\mathbf{x}_1, \dots, \mathbf{x}_{T-1}$ into the current RNN to get $\hat{\mathbf{y}}_1, \dots, \hat{\mathbf{y}}_{T-1} \in \mathbb{R}^d$

How to train an RNN

Take **text generation** (unsupervised learning) as an example:

- given a corpus, train an RNN that learns $P(\mathbf{x}_t \mid \mathbf{x}_{1:t-1})$

For each sequence $\mathbf{x}_1, \dots, \mathbf{x}_T \in \mathbb{R}^d$ (one-hot representation) in the corpus

- feed $\mathbf{x}_1, \dots, \mathbf{x}_{T-1}$ into the current RNN to get $\hat{\mathbf{y}}_1, \dots, \hat{\mathbf{y}}_{T-1} \in \mathbb{R}^d$
- each $\hat{\mathbf{y}}_t$ defines a distribution over the next word via **softmax**:
 $P(\text{next word} = i) \propto \exp(\hat{\mathbf{y}}_{t,i})$

How to train an RNN

Take **text generation** (unsupervised learning) as an example:

- given a corpus, train an RNN that learns $P(\mathbf{x}_t \mid \mathbf{x}_{1:t-1})$

For each sequence $\mathbf{x}_1, \dots, \mathbf{x}_T \in \mathbb{R}^d$ (one-hot representation) in the corpus

- feed $\mathbf{x}_1, \dots, \mathbf{x}_{T-1}$ into the current RNN to get $\hat{\mathbf{y}}_1, \dots, \hat{\mathbf{y}}_{T-1} \in \mathbb{R}^d$
- each $\hat{\mathbf{y}}_t$ defines a distribution over the next word via **softmax**:
 $P(\text{next word} = i) \propto \exp(\hat{\mathbf{y}}_{t,i})$
- based on the **true label** $\mathbf{x}_{t+1}$, each $\hat{\mathbf{y}}_t$ incurs **cross-entropy** loss

$$-\ln \left(\frac{\exp(\hat{\mathbf{y}}_t^\top \mathbf{x}_{t+1})}{\sum_{i=1}^d \exp(\hat{\mathbf{y}}_{t,i})} \right)$$

How to train an RNN

Take **text generation** (unsupervised learning) as an example:

- given a corpus, train an RNN that learns $P(\mathbf{x}_t \mid \mathbf{x}_{1:t-1})$

For each sequence $\mathbf{x}_1, \dots, \mathbf{x}_T \in \mathbb{R}^d$ (one-hot representation) in the corpus

- feed $\mathbf{x}_1, \dots, \mathbf{x}_{T-1}$ into the current RNN to get $\hat{\mathbf{y}}_1, \dots, \hat{\mathbf{y}}_{T-1} \in \mathbb{R}^d$
- each $\hat{\mathbf{y}}_t$ defines a distribution over the next word via **softmax**:
 $P(\text{next word} = i) \propto \exp(\hat{\mathbf{y}}_{t,i})$
- based on the **true label** $\mathbf{x}_{t+1}$, each $\hat{\mathbf{y}}_t$ incurs **cross-entropy** loss

$$-\ln \left(\frac{\exp(\hat{\mathbf{y}}_t^\top \mathbf{x}_{t+1})}{\sum_{i=1}^d \exp(\hat{\mathbf{y}}_{t,i})} \right)$$

- update the RNN parameters using **backpropagation over the total loss**

Demo

Tiny RNN, predicting the **next bit of a binary sequence**

- <https://joshvarty.github.io/VisualizingRNNs/>

Demo

Tiny RNN, predicting the **next bit of a binary sequence**

- <https://joshvarty.github.io/VisualizingRNNs/>
- the entire vocabulary is just $\{0, 1\}$ ($d = 2$)

Demo

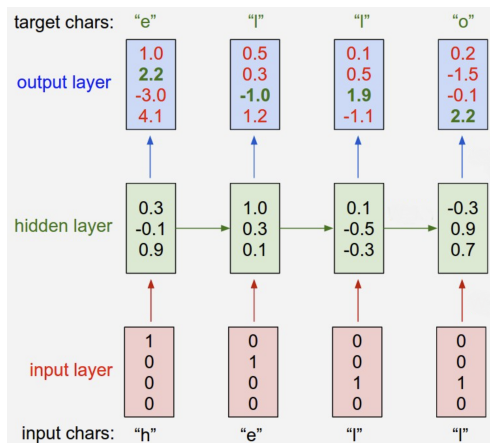
Tiny RNN, predicting the **next bit of a binary sequence**

- <https://joshvarty.github.io/VisualizingRNNs/>
- the entire vocabulary is just $\{0, 1\}$ ($d = 2$)
- one-layer RNN with $d_h = 3$, so parameters are
 $\mathbf{W} \in \mathbb{R}^{3 \times 3}, \mathbf{U} \in \mathbb{R}^{3 \times 2}, \mathbf{V} \in \mathbb{R}^{2 \times 3}, \mathbf{b}_h \in \mathbb{R}^3, \mathbf{b}_y \in \mathbb{R}^2$

Another demo

Min-Char RNN, predicting the **next character of a sequence**

- <https://karpathy.github.io/2015/05/21/rnn-effectiveness/>



Generation after training

Keep sampling from $\text{softmax}(\hat{\mathbf{y}}_t)$ as the next input $\mathbf{x}_{t+1}$ to RNN

Generation after training

Keep sampling from $\text{softmax}(\hat{\mathbf{y}}_t)$ as the next input $\mathbf{x}_{t+1}$ to RNN

Can control how “random” the generation is via $\text{softmax}(\beta \cdot \hat{\mathbf{y}}_t)$

Generation after training

Keep sampling from $\text{softmax}(\hat{\mathbf{y}}_t)$ as the next input $\mathbf{x}_{t+1}$ to RNN

Can control how “random” the generation is via $\text{softmax}(\beta \cdot \hat{\mathbf{y}}_t)$

- $1/\beta$ is called **temperature**

Generation after training

Keep sampling from $\text{softmax}(\hat{\mathbf{y}}_t)$ as the next input $\mathbf{x}_{t+1}$ to RNN

Can control how “random” the generation is via $\text{softmax}(\beta \cdot \hat{\mathbf{y}}_t)$

- $1/\beta$ is called **temperature**
- larger temperature (smaller β) leads to more random outputs

Generation after training

Keep sampling from $\text{softmax}(\hat{\mathbf{y}}_t)$ as the next input $\mathbf{x}_{t+1}$ to RNN

Can control how “random” the generation is via $\text{softmax}(\beta \cdot \hat{\mathbf{y}}_t)$

- $1/\beta$ is called **temperature**
- larger temperature (smaller β) leads to more random outputs
 - $\beta = 0$, uniform output (**maximum entropy**)

Generation after training

Keep sampling from $\text{softmax}(\hat{\mathbf{y}}_t)$ as the next input $\mathbf{x}_{t+1}$ to RNN

Can control how “random” the generation is via $\text{softmax}(\beta \cdot \hat{\mathbf{y}}_t)$

- $1/\beta$ is called **temperature**
- larger temperature (smaller β) leads to more random outputs
 - $\beta = 0$, uniform output (**maximum entropy**)
 - $\beta = \infty$, deterministically output $\text{argmax}_i \hat{\mathbf{y}}_{t,i}$ (“**hard**” **max**)

Generation after training

A few remarkable examples from Min-Char RNN:

- corpus: $\text{\LaTeX}$ source code of an algebraic geometry book (16MB)

Generation after training

A few remarkable examples from Min-Char RNN:

- corpus: $\text{\LaTeX}$ source code of an algebraic geometry book (16MB)
- generate source code that **almost complies**

For $\bigoplus_{n=1,\dots,m}$ where $\mathcal{L}_{m,\bullet} = 0$, hence we can find a closed subset $\mathcal{H}$ in $\mathcal{H}$ and any sets $\mathcal{F}$ on X , U is a closed immersion of S , then $U \rightarrow T$ is a separated algebraic space.

Proof. Proof of (1). It also start we get

$$S = \text{Spec}(R) = U \times_X U \times_X U$$

and the comparico in the fibre product covering we have to prove the lemma generated by $\coprod Z \times_U U \rightarrow V$. Consider the maps M along the set of points Sch_{fppf} and $U \rightarrow U$ is the fibre category of S in U in Section, ?? and the fact that any U affine, see Morphisms, Lemma ?? . Hence we obtain a scheme S and any open subset $W \subset U$ in $\text{Sh}(G)$ such that $\text{Spec}(R') \rightarrow S$ is smooth or an

$$U = \bigcup U_i \times_{S_i} U_i$$

which has a nonzero morphism we may assume that f_i is of finite presentation over S . We claim that $\mathcal{O}_{X,x}$ is a scheme where $x, x', s'' \in S'$ such that $\mathcal{O}_{X,x'} \rightarrow \mathcal{O}'_{X',x'}$ is separated. By Algebra, Lemma ?? we can define a map of complexes $\text{GL}_{S'}(x'/S'')$ and we win. $\square$

To prove study we see that $\mathcal{F}|_U$ is a covering of $\mathcal{X}'$, and $\mathcal{T}_i$ is an object of $\mathcal{F}_{X/S}$ for $i > 0$ and $\mathcal{F}_p$ exists and let $\mathcal{F}_i$ be a presheaf of $\mathcal{O}_X$ -modules on $\mathcal{C}$ as a $\mathcal{F}$ -module. In particular $\mathcal{F} = U/\mathcal{F}$ we have to show that

$$\widetilde{M}^\bullet = \mathcal{I}^\bullet \otimes_{\text{Spec}(k)} \mathcal{O}_{S,s} - i_X^{-1} \mathcal{F}$$

Generation after training

A few remarkable examples from Min-Char RNN:

- corpus: L^AT_EX source code of an algebraic geometry book (16MB)
- generate source code that almost complies
- the model understands complex syntactic structures

For $\bigoplus_{n=1,\dots,m}$ where $\mathcal{L}_{m,\bullet} = 0$, hence we can find a closed subset $\mathcal{H}$ in $\mathcal{H}$ and any sets $\mathcal{F}$ on X , U is a closed immersion of S , then $U \rightarrow T$ is a separated algebraic space.

Proof. Proof of (1). It also start we get

$$S = \mathrm{Spec}(R) = U \times_X U \times_X U$$

and the comparicoly in the fibre product covering we have to prove the lemma generated by $\coprod Z \times_U U \rightarrow V$. Consider the maps M along the set of points Sch_{fppf} and $U \rightarrow U$ is the fibre category of S in U in Section, ?? and the fact that any U affine, see Morphisms, Lemma ?? . Hence we obtain a scheme S and any open subset $W \subset U$ in $Sh(G)$ such that $\mathrm{Spec}(R') \rightarrow S$ is smooth or an

$$U = \bigcup U_i \times_{S_i} U_i$$

which has a nonzero morphism we may assume that f_i is of finite presentation over S . We claim that $\mathcal{O}_{X,x}$ is a scheme where $x, x', s'' \in S'$ such that $\mathcal{O}_{X,x'} \rightarrow \mathcal{O}'_{X',x'}$ is separated. By Algebra, Lemma ?? we can define a map of complexes $\mathrm{GL}_{S'}(x'/S'')$ and we win. $\square$

To prove study we see that $\mathcal{F}|_U$ is a covering of $\mathcal{X}'$, and $\mathcal{T}_i$ is an object of $\mathcal{F}_{X/S}$ for $i > 0$ and $\mathcal{F}_p$ exists and let $\mathcal{F}_i$ be a presheaf of $\mathcal{O}_{X,S}$ -modules on $\mathcal{C}$ as a $\mathcal{F}$ -module. In particular $\mathcal{F} = U/\mathcal{F}$ we have to show that

$$\widetilde{M}^\bullet = \mathcal{I}^\bullet \otimes_{\mathrm{Spec}(k)} \mathcal{O}_{S,s} - i_X^{-1} \mathcal{F}$$

Generation after training

A few remarkable examples from Min-Char RNN:

- corpus: **Linux source code** (474MB of C code); 10M parameters

Generation after training

A few remarkable examples from Min-Char RNN:

- corpus: **Linux source code** (474MB of C code); 10M parameters
- generate codes with very few syntactic errors

```
static int indicate_policy(void)
{
    int error;
    if (fd == MARN_EPT) {
        /*
         * The kernel blank will coeld it to userspace.
         */
        if (ss->segment < mem_total)
            unblock_graph_and_set_blocked();
        else
            ret = 1;
        goto bail;
    }
    segaddr = in_SB(in.addr);
    selector = seg / 16;
    setup_works = true;
    for (i = 0; i < blocks; i++) {
        seq = buf[i++];
        bpf = bd->bd.next + i * search;
        if (fd) {
            current = blocked;
        }
    }
    rw->name = "Getjbbregs";
    bprm_self_clear1(&iv->version);
    regs->new = blocks[(BPF_STATS << info->historidac)] | PFMR_CLOBATHINC_SECONDS << 12;
    return segtable;
}
```

Generation after training

A few remarkable examples from Min-Char RNN:

- corpus: **Linux source code** (474MB of C code); 10M parameters
- generate codes with very few syntactic errors
- uses strings/pointers properly, open/close brackets correctly, good indentation, even add comments

```
static int indicate_policy(void)
{
    int error;
    if (fd == MARN_EPT) {
        /*
         * The kernel blank will coeld it to userspace.
         */
        if (ss->segment < mem_total)
            unblock_graph_and_set_blocked();
        else
            ret = 1;
        goto bail;
    }
    segaddr = in_SB(in.addr);
    selector = seg / 16;
    setup_works = true;
    for (i = 0; i < blocks; i++) {
        seq = buf[i++];
        bpf = bd->bd.next + i * search;
        if (fd) {
            current = blocked;
        }
    }
    rw->name = "Getjbbregs";
    bprm_self_clearl(&iv->version);
    regs->new = blocks[(BPF_STATS << info->historidac)] | PFMR_CLOBATHINC_SECONDS << 12;
    return segtable;
}
```

A closer look at some neurons

Some neurons (entries of the hidden state h) are quite **interpretable** (even though most are not)

A closer look at some neurons

Some neurons (entries of the hidden state h) are quite **interpretable** (even though most are not)

- can visualize this by coloring the input character based on the value of this neuron (red = large value, blue = small value)

A closer look at some neurons

Some neurons (entries of the hidden state h) are quite **interpretable** (even though most are not)

- can visualize this by coloring the input character based on the value of this neuron (red = large value, blue = small value)

The sole importance of the crossing of the Berezina lies in the fact that it plainly and indubitably proved the fallacy of all the plans for cutting off the enemy's retreat and the soundness of the only possible line of action--the one Kutuzov and the general mass of the army demanded--namely, simply to follow the enemy up. The French crowd fled at a continually increasing speed and all its energy was directed to reaching its goal. It fled like a wounded animal and it was impossible to block its path. This was shown not so much by the arrangements it made for crossing as by what took place at the bridges. When the bridges broke down, unarmed soldiers, people from Moscow and women with children who were with the French transport, all--carried on by vis inertiae--pressed forward into boats and into the ice-covered water and did not, surrender.

A neuron sensitive to the **position** in line

A closer look at some neurons

Some neurons (entries of the hidden state h) are quite **interpretable** (even though most are not)

- can visualize this by coloring the input character based on the value of this neuron (**red = large value**, **blue = small value**)

"You mean to imply that I have nothing to eat out of.... on the contrary, I can supply you with everything even if you want to give dinner parties," warmly replied Chichagov, who tried by every word he spoke to prove his own rectitude and therefore imagined Kutuzov to be animated by the same desire.

Kutuzov, shrugging his shoulders, replied with his subtle penetrating smile: "I meant merely to say what I said."

A neuron that is turned on inside **quotes**

A closer look at some neurons

Some neurons (entries of the hidden state h) are quite **interpretable** (even though most are not)

- can visualize this by coloring the input character based on the value of this neuron (**red = large value**, **blue = small value**)

```
static int __dequeue_signal(struct sigpending *pending,
                           siginfo_t *info)
{
    int sig = next_signal(pending, mask);
    if (sig) {
        if (current->notifier) {
            if (sigismember(current->notifier_mask, sig)) {
                if (!current->notifier(current->notifier_data)) {
                    clear_thread_flag(TIF_SIGPENDING);
                    return 0;
                }
            }
        }
        collect_signal(sig, pending, info);
    }
    return sig;
}
```

A neuron that is turned on inside **if statements**

A closer look at some neurons

Some neurons (entries of the hidden state h) are quite **interpretable** (even though most are not)

- can visualize this by coloring the input character based on the value of this neuron (red = large value, blue = small value)

```
#ifdef CONFIG_AUDITSYSCALL
static inline int audit_match_class_bits(int class, u32 *mask)
{
    int i;
    if (classes[class]) {
        for (i = 0; i < AUDIT_BITMASK_SIZE; i++)
            if (mask[i] & classes[class][i])
                return 0;
    }
    return 1;
}
```

A neuron sensitive to the **depth** of an expression

Final notes

Instead of using a character or a word as each x , often use a **token** (word or sub-word)

Final notes

Instead of using a character or a word as each x , often use a **token** (word or sub-word)

- “apple” is a token

Final notes

Instead of using a character or a word as each x , often use a **token** (word or sub-word)

- “apple” is a token
- “unbelievable” is 3 tokens (“un”, “believ”, “able”)

Final notes

Instead of using a character or a word as each x , often use a **token** (word or sub-word)

- “apple” is a token
- “unbelievable” is 3 tokens (“un”, “believ”, “able”)
- can reduce the size of vocabulary

Final notes

Instead of using a character or a word as each x , often use a **token** (word or sub-word)

- “apple” is a token
- “unbelievable” is 3 tokens (“un”, “believ”, “able”)
- can reduce the size of vocabulary

Directly applying backpropagation to RNN leads to *vanishing/exploding gradient* issues when T is large

Final notes

Instead of using a character or a word as each x , often use a **token** (word or sub-word)

- “apple” is a token
- “unbelievable” is 3 tokens (“un”, “believ”, “able”)
- can reduce the size of vocabulary

Directly applying backpropagation to RNN leads to *vanishing/exploding gradient* issues when T is large

- W is applied T times at the end of the sequence (so roughly W^T)

Final notes

Instead of using a character or a word as each x , often use a **token** (word or sub-word)

- “apple” is a token
- “unbelievable” is 3 tokens (“un”, “believ”, “able”)
- can reduce the size of vocabulary

Directly applying backpropagation to RNN leads to *vanishing/exploding gradient* issues when T is large

- W is applied T times at the end of the sequence (so roughly W^T)
- some fixes: **Long Short-Term Memory** (LSTM) and **Gated Recurrent Units** (GRU)

Outline

- 1 Review of last lecture
- 2 Recurrent Neural Network
- 3 Transformers
 - Self-attention
 - Other components
 - Training and testing

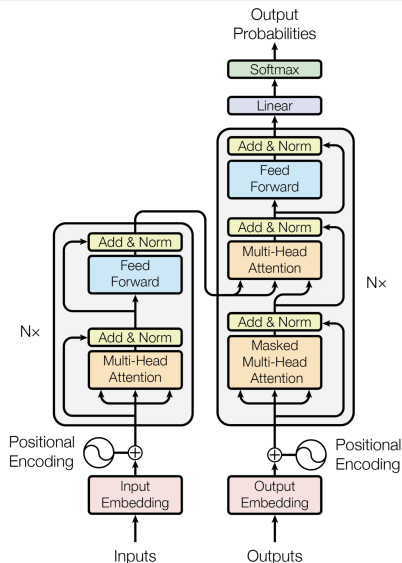
Transformers

Issues of RNN: *must compress all previous info into a single state h*

Transformers

Issues of RNN: *must compress all previous info into a single state h*

A solution that dominates all other models currently: **transformers**



Acknowledgements

Very useful resources:

- original paper: “**Attention Is All You Need**” (200K+ citation by now)
 - <https://arxiv.org/pdf/1706.03762>
- The Illustrated Transformer (most pictures are from here)
 - <https://jalammar.github.io/illustrated-transformer/>
- a super cool Nano-GPT visualization
 - <https://bbycroft.net/llm>
- A Multiscale Visualization of Attention
 - <https://arxiv.org/pdf/1906.05714>

Key idea: self-attention

Example: “*The animal didn’t cross the street because **it** was too tired*”

Key idea: self-attention

Example: “*The animal didn’t cross the street because **it** was too tired*”

- Does “**it**” refer to “animal” or “street”?

Key idea: self-attention

Example: “*The animal didn’t cross the street because **it** was too tired*”

- Does “**it**” refer to “animal” or “street”?
- trivial for human, but how to design a model that understands this?

Key idea: self-attention

Example: “*The animal didn’t cross the street because **it** was too tired*”

- Does “**it**” refer to “animal” or “street”?
- trivial for human, but how to design a model that understands this?
- intuitively, when looking at the word “it”, the model should pay **attention** to the word “animal”

Key idea: self-attention

Example: “*The animal didn’t cross the street because **it** was too tired*”

- Does “**it**” refer to “animal” or “street”?
- trivial for human, but how to design a model that understands this?
- intuitively, when looking at the word “it”, the model should pay **attention** to the word “animal”
- An **attention head** does exactly this

Attention head

An attention head

- takes a sequence of inputs $\mathbf{x}_1, \dots, \mathbf{x}_T \in \mathbb{R}^d$ and outputs another sequence $\mathbf{z}_1, \dots, \mathbf{z}_T \in \mathbb{R}^{d_v}$ (similar to **hidden states** of RNN)

Attention head

An attention head

- takes a sequence of inputs $\mathbf{x}_1, \dots, \mathbf{x}_T \in \mathbb{R}^d$ and outputs another sequence $\mathbf{z}_1, \dots, \mathbf{z}_T \in \mathbb{R}^{d_v}$ (similar to **hidden states** of RNN)
- parametrized by three matrices (and corresponding biases, omitted for simplicity): $\mathbf{W}_Q \in \mathbb{R}^{d \times d_k}$, $\mathbf{W}_K \in \mathbb{R}^{d \times d_k}$, $\mathbf{W}_V \in \mathbb{R}^{d \times d_v}$

Attention head

An attention head

- takes a sequence of inputs $\mathbf{x}_1, \dots, \mathbf{x}_T \in \mathbb{R}^d$ and outputs another sequence $\mathbf{z}_1, \dots, \mathbf{z}_T \in \mathbb{R}^{d_v}$ (similar to **hidden states** of RNN)
- parametrized by three matrices (and corresponding biases, omitted for simplicity): $\mathbf{W}_Q \in \mathbb{R}^{d \times d_k}$, $\mathbf{W}_K \in \mathbb{R}^{d \times d_k}$, $\mathbf{W}_V \in \mathbb{R}^{d \times d_v}$
 - computes a **query** vector for each input $\mathbf{x}_t$ as $\mathbf{q}_t = \mathbf{W}_Q^\top \mathbf{x}_t \in \mathbb{R}^{d_k}$

Attention head

An attention head

- takes a sequence of inputs $\mathbf{x}_1, \dots, \mathbf{x}_T \in \mathbb{R}^d$ and outputs another sequence $\mathbf{z}_1, \dots, \mathbf{z}_T \in \mathbb{R}^{d_v}$ (similar to **hidden states** of RNN)
- parametrized by three matrices (and corresponding biases, omitted for simplicity): $\mathbf{W}_Q \in \mathbb{R}^{d \times d_k}$, $\mathbf{W}_K \in \mathbb{R}^{d \times d_k}$, $\mathbf{W}_V \in \mathbb{R}^{d \times d_v}$
 - computes a **query** vector for each input $\mathbf{x}_t$ as $\mathbf{q}_t = \mathbf{W}_Q^\top \mathbf{x}_t \in \mathbb{R}^{d_k}$
 - computes a **key** vector for each input $\mathbf{x}_t$ as $\mathbf{k}_t = \mathbf{W}_K^\top \mathbf{x}_t \in \mathbb{R}^{d_k}$

Attention head

An attention head

- takes a sequence of inputs $\mathbf{x}_1, \dots, \mathbf{x}_T \in \mathbb{R}^d$ and outputs another sequence $\mathbf{z}_1, \dots, \mathbf{z}_T \in \mathbb{R}^{d_v}$ (similar to **hidden states** of RNN)
- parametrized by three matrices (and corresponding biases, omitted for simplicity): $\mathbf{W}_Q \in \mathbb{R}^{d \times d_k}$, $\mathbf{W}_K \in \mathbb{R}^{d \times d_k}$, $\mathbf{W}_V \in \mathbb{R}^{d \times d_v}$
 - computes a **query** vector for each input $\mathbf{x}_t$ as $\mathbf{q}_t = \mathbf{W}_Q^\top \mathbf{x}_t \in \mathbb{R}^{d_k}$
 - computes a **key** vector for each input $\mathbf{x}_t$ as $\mathbf{k}_t = \mathbf{W}_K^\top \mathbf{x}_t \in \mathbb{R}^{d_k}$
 - computes a **value** vector for each input $\mathbf{x}_t$ as $\mathbf{v}_t = \mathbf{W}_V^\top \mathbf{x}_t \in \mathbb{R}^{d_v}$

Attention head

An attention head

- takes a sequence of inputs $\mathbf{x}_1, \dots, \mathbf{x}_T \in \mathbb{R}^d$ and outputs another sequence $\mathbf{z}_1, \dots, \mathbf{z}_T \in \mathbb{R}^{d_v}$ (similar to **hidden states** of RNN)
- parametrized by three matrices (and corresponding biases, omitted for simplicity): $\mathbf{W}_Q \in \mathbb{R}^{d \times d_k}$, $\mathbf{W}_K \in \mathbb{R}^{d \times d_k}$, $\mathbf{W}_V \in \mathbb{R}^{d \times d_v}$
 - computes a **query** vector for each input $\mathbf{x}_t$ as $\mathbf{q}_t = \mathbf{W}_Q^\top \mathbf{x}_t \in \mathbb{R}^{d_k}$
 - computes a **key** vector for each input $\mathbf{x}_t$ as $\mathbf{k}_t = \mathbf{W}_K^\top \mathbf{x}_t \in \mathbb{R}^{d_k}$
 - computes a **value** vector for each input $\mathbf{x}_t$ as $\mathbf{v}_t = \mathbf{W}_V^\top \mathbf{x}_t \in \mathbb{R}^{d_v}$
- the output $\mathbf{z}_t$ is the “**answer**” to the query of $\mathbf{q}_t$

Attention head (cont.)

Input

Thinking

Machines

Embedding

x_1 

x_2 

Queries

q_1 

q_2 



W^Q

Keys

k_1 

k_2 



W^K

Values

v_1 

v_2 



W^V

Attention head (cont.)

The output z_t is the “**answer**” to the query of q_t . *How?*

Attention head (cont.)

The output z_t is the “**answer**” to the query of q_t . *How?*

- imagine: you make a Google query (q_t),

Attention head (cont.)

The output z_t is the “**answer**” to the query of q_t . *How?*

- imagine: you make a Google query (q_t), and it returns a list of website titles ($k_{1:T}$);

Attention head (cont.)

The output z_t is the “**answer**” to the query of q_t . *How?*

- imagine: you make a **Google query** (q_t), and it returns a list of **website titles** ($k_{1:T}$); clicking a title (k_τ) leads you to a **website** (v_τ).

Attention head (cont.)

The output z_t is the “**answer**” to the query of q_t . *How?*

- imagine: you make a **Google query** (q_t), and it returns a list of **website titles** ($k_{1:T}$); clicking a title (k_τ) leads you to a **website** (v_τ).
- You then summarize the answer using all **websites** ($v_{1:T}$),

Attention head (cont.)

The output z_t is the “**answer**” to the query of q_t . *How?*

- imagine: you make a **Google query** (q_t), and it returns a list of **website titles** ($k_{1:T}$); clicking a title (k_τ) leads you to a **website** (v_τ).
- You then summarize the answer using all **websites** ($v_{1:T}$), each with a different **weight** based on how relevant/close its **title** is to your **query**

Attention head (cont.)

The output z_t is the “**answer**” to the query of q_t . *How?*

- imagine: you make a **Google query** (q_t), and it returns a list of **website titles** ($k_{1:T}$); clicking a title (k_τ) leads you to a **website** (v_τ).
- You then summarize the answer using all **websites** ($v_{1:T}$), each with a different **weight** based on how relevant/close its **title** is to your **query**
- formally, the final answer z_t is the **weighted sum** of $v_1, \dots, v_T$,

Attention head (cont.)

The output z_t is the “**answer**” to the query of q_t . *How?*

- imagine: you make a **Google query** (q_t), and it returns a list of **website titles** ($k_{1:T}$); clicking a title (k_τ) leads you to a **website** (v_τ).
- You then summarize the answer using all **websites** ($v_{1:T}$), each with a different **weight** based on how relevant/close its **title** is to your **query**
- formally, the final answer z_t is the **weighted sum** of $v_1, \dots, v_T$, with weights computed via

$$\text{softmax} \left(\frac{q_t^\top k_1}{\sqrt{d_k}}, \dots, \frac{q_t^\top k_T}{\sqrt{d_k}} \right)$$

Attention head (cont.)

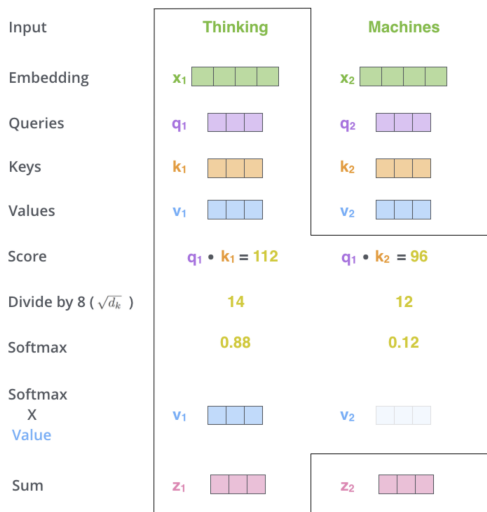
The output z_t is the “**answer**” to the query of q_t . *How?*

- imagine: you make a **Google query** (q_t), and it returns a list of **website titles** ($k_{1:T}$); clicking a title (k_τ) leads you to a **website** (v_τ).
- You then summarize the answer using all **websites** ($v_{1:T}$), each with a different **weight** based on how relevant/close its **title** is to your **query**
- formally, the final answer z_t is the **weighted sum** of $v_1, \dots, v_T$, with weights computed via

$$\text{softmax} \left(\frac{q_t^\top k_1}{\sqrt{d_k}}, \dots, \frac{q_t^\top k_T}{\sqrt{d_k}} \right)$$

where $q_t^\top k_\tau$ is the **attention score** from input x_t to input x_τ

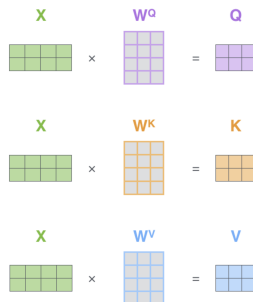
Attention head (cont.)



Attention head (cont.)

Matrix notation:

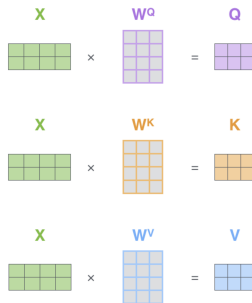
- input matrix $\mathbf{X} \in \mathbb{R}^{T \times d}$, obtained by stacking $\mathbf{x}_1^\top, \dots, \mathbf{x}_T^\top$



Attention head (cont.)

Matrix notation:

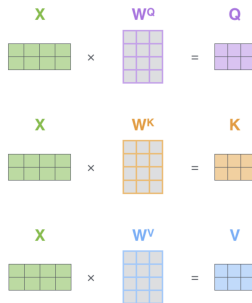
- input matrix $\mathbf{X} \in \mathbb{R}^{T \times d}$, obtained by stacking $\mathbf{x}_1^\top, \dots, \mathbf{x}_T^\top$
- query matrix $\mathbf{Q} = \mathbf{X}\mathbf{W}_Q \in \mathbb{R}^{T \times d_k}$



Attention head (cont.)

Matrix notation:

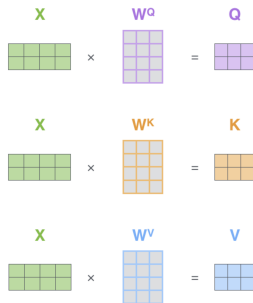
- input matrix $\mathbf{X} \in \mathbb{R}^{T \times d}$, obtained by stacking $\mathbf{x}_1^\top, \dots, \mathbf{x}_T^\top$
- query matrix $\mathbf{Q} = \mathbf{X}\mathbf{W}_Q \in \mathbb{R}^{T \times d_k}$
- key matrix $\mathbf{K} = \mathbf{X}\mathbf{W}_K \in \mathbb{R}^{T \times d_k}$



Attention head (cont.)

Matrix notation:

- input matrix $\mathbf{X} \in \mathbb{R}^{T \times d}$, obtained by stacking $\mathbf{x}_1^\top, \dots, \mathbf{x}_T^\top$
- query matrix $\mathbf{Q} = \mathbf{X}\mathbf{W}_Q \in \mathbb{R}^{T \times d_k}$
- key matrix $\mathbf{K} = \mathbf{X}\mathbf{W}_K \in \mathbb{R}^{T \times d_k}$
- value matrix $\mathbf{V} = \mathbf{X}\mathbf{W}_V \in \mathbb{R}^{T \times d_v}$



Attention head (cont.)

Matrix notation:

- input matrix $\mathbf{X} \in \mathbb{R}^{T \times d}$, obtained by stacking $\mathbf{x}_1^\top, \dots, \mathbf{x}_T^\top$
- query matrix $\mathbf{Q} = \mathbf{X}\mathbf{W}_Q \in \mathbb{R}^{T \times d_k}$
- key matrix $\mathbf{K} = \mathbf{X}\mathbf{W}_K \in \mathbb{R}^{T \times d_k}$
- value matrix $\mathbf{V} = \mathbf{X}\mathbf{W}_V \in \mathbb{R}^{T \times d_v}$
- attention score matrix $\mathbf{Q}\mathbf{K}^\top \in \mathbb{R}^{T \times T}$

$$\begin{array}{c} \mathbf{X} \\ \text{3x3 grid} \end{array} \times \begin{array}{c} \mathbf{W}_Q \\ \text{3x4 grid} \end{array} = \begin{array}{c} \mathbf{Q} \\ \text{3x4 grid} \end{array}$$

$$\begin{array}{c} \mathbf{X} \\ \text{3x3 grid} \end{array} \times \begin{array}{c} \mathbf{W}_K \\ \text{3x4 grid} \end{array} = \begin{array}{c} \mathbf{K} \\ \text{3x4 grid} \end{array}$$

$$\begin{array}{c} \mathbf{X} \\ \text{3x3 grid} \end{array} \times \begin{array}{c} \mathbf{W}_V \\ \text{3x4 grid} \end{array} = \begin{array}{c} \mathbf{V} \\ \text{3x4 grid} \end{array}$$

$$\text{softmax}\left(\frac{\mathbf{Q} \times \mathbf{K}^\top}{\sqrt{d_k}}\right) \mathbf{V}$$

$$= \mathbf{Z}$$

(3x4 grid)

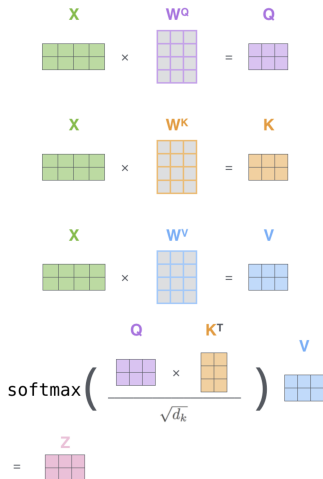
Attention head (cont.)

Matrix notation:

- input matrix $\mathbf{X} \in \mathbb{R}^{T \times d}$, obtained by stacking $\mathbf{x}_1^\top, \dots, \mathbf{x}_T^\top$
- query matrix $\mathbf{Q} = \mathbf{X}\mathbf{W}_Q \in \mathbb{R}^{T \times d_k}$
- key matrix $\mathbf{K} = \mathbf{X}\mathbf{W}_K \in \mathbb{R}^{T \times d_k}$
- value matrix $\mathbf{V} = \mathbf{X}\mathbf{W}_V \in \mathbb{R}^{T \times d_v}$
- attention score matrix $\mathbf{Q}\mathbf{K}^\top \in \mathbb{R}^{T \times T}$
- output matrix $\mathbf{Z} \in \mathbb{R}^{T \times d_v}$ is

$$\text{softmax}\left(\frac{\mathbf{Q}\mathbf{K}^\top}{\sqrt{d_k}}\right)\mathbf{V}$$

where softmax is *applied row-wise*



Attention head (cont.)

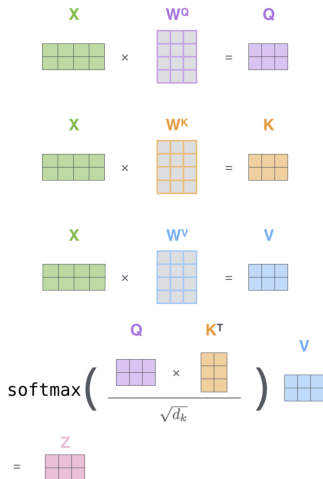
Matrix notation:

- input matrix $\mathbf{X} \in \mathbb{R}^{T \times d}$, obtained by stacking $\mathbf{x}_1^\top, \dots, \mathbf{x}_T^\top$
- query matrix $\mathbf{Q} = \mathbf{X}\mathbf{W}_Q \in \mathbb{R}^{T \times d_k}$
- key matrix $\mathbf{K} = \mathbf{X}\mathbf{W}_K \in \mathbb{R}^{T \times d_k}$
- value matrix $\mathbf{V} = \mathbf{X}\mathbf{W}_V \in \mathbb{R}^{T \times d_v}$
- attention score matrix $\mathbf{Q}\mathbf{K}^\top \in \mathbb{R}^{T \times T}$
- output matrix $\mathbf{Z} \in \mathbb{R}^{T \times d_v}$ is

$$\text{softmax}\left(\frac{\mathbf{Q}\mathbf{K}^\top}{\sqrt{d_k}}\right)\mathbf{V}$$

where softmax is *applied row-wise*

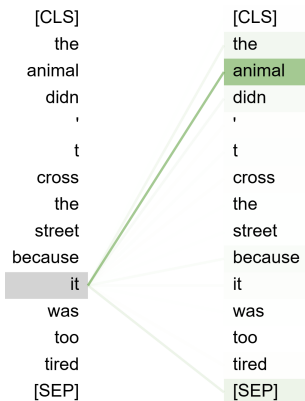
$O(T^2)$ complexity (ignoring d, d_k, d_v)



Visualization of an attention head

[link](#)

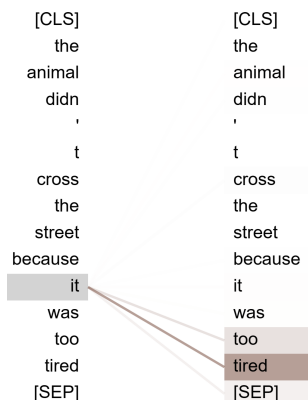
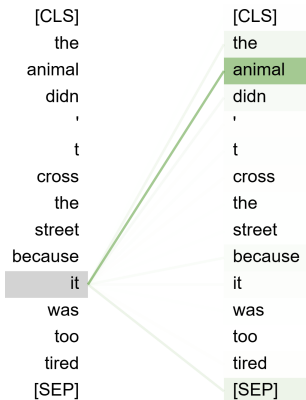
- the darker the color, the larger the attention score
- “it” attends to “animal” in one head,



Visualization of an attention head

[link](#)

- the darker the color, the larger the attention score
- “it” attends to “animal” in one head, and “tired” in another head

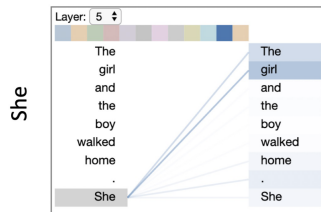


Visualization of an attention head

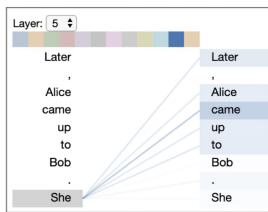
[link](#)

More examples

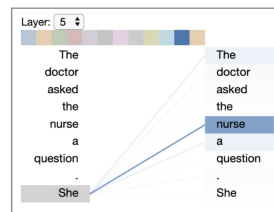
Gender-specific term



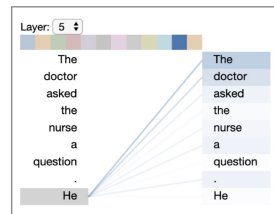
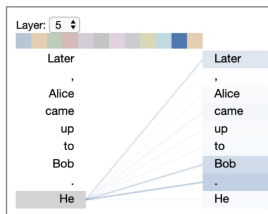
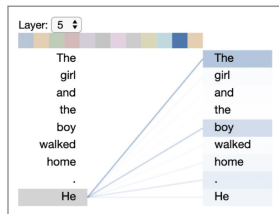
Name



Occupation



He

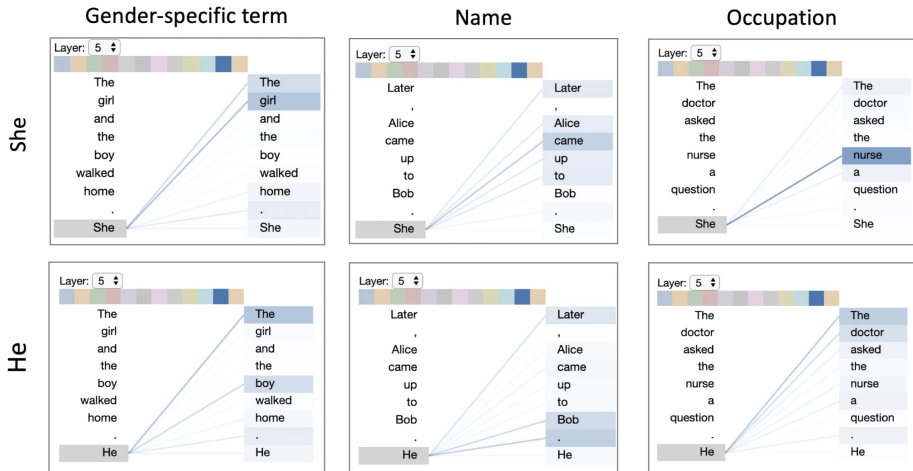


Visualization of an attention head

[link](#)

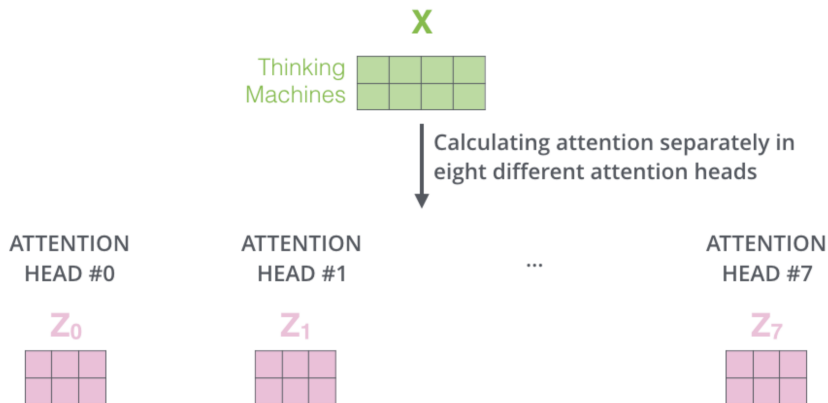
More examples

- all from unsupervised learning; *no one tells the model to learn these!*



Multi-head attention

Pass X to multiple attention-heads, each with **different parameters**



Multi-head attention (cont.)

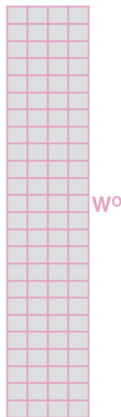
Concatenate outputs of different heads and then **project** again

1) Concatenate all the attention heads



2) Multiply with a weight matrix W^O that was trained jointly with the model

$\times$



3) The result would be the Z matrix that captures information from all the attention heads. We can send this forward to the FFNN



Multi-head attention (cont.)

Concatenate outputs of different heads and then **project** again

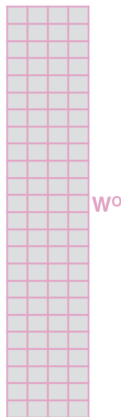
- final output dimension is $\mathbb{R}^{T \times d}$, same as inputs X

1) Concatenate all the attention heads



2) Multiply with a weight matrix W^O that was trained jointly with the model

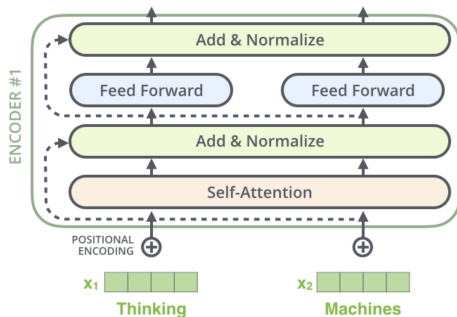
$\times$



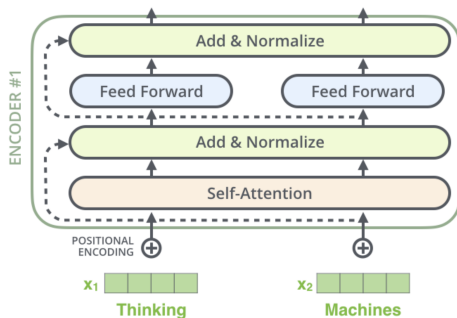
3) The result would be the Z matrix that captures information from all the attention heads. We can send this forward to the FFNN



Zooming out: a complete encoder



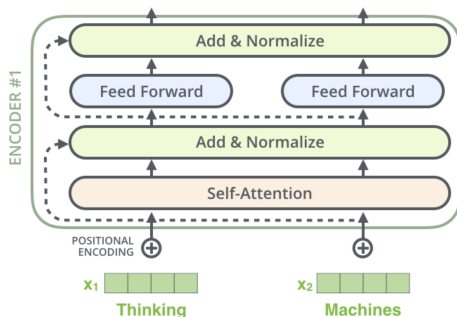
Zooming out: a complete encoder



Positional encoding

- fix the issue that attention-head *does not have positional info*

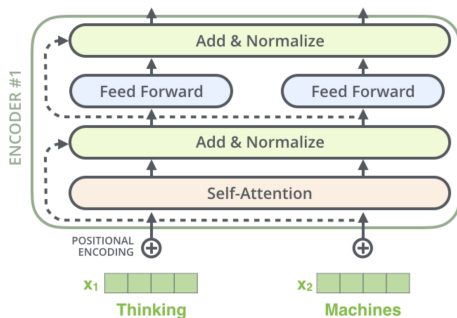
Zooming out: a complete encoder



Positional encoding

- fix the issue that attention-head *does not have positional info*
- via a **positional embedding matrix** $E_P \in \mathbb{R}^{d \times T}$, fixed or learned,

Zooming out: a complete encoder



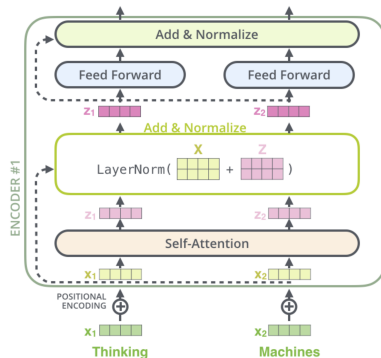
Positional encoding

- fix the issue that attention-head *does not have positional info*
- via a **positional embedding matrix** $E_P \in \mathbb{R}^{d \times T}$, fixed or learned,
- $x_t \leftarrow x_t + E_P e_t$, i.e., add the t -th column of E_P to x_t

Zooming out: a complete encoder (cont.)

Two more components to stabilize and speed up training:

1. Residual pathway

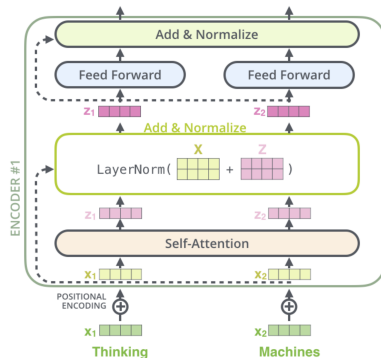


Zooming out: a complete encoder (cont.)

Two more components to stabilize and speed up training:

1. Residual pathway

- add input to output, $Z \leftarrow Z + X$

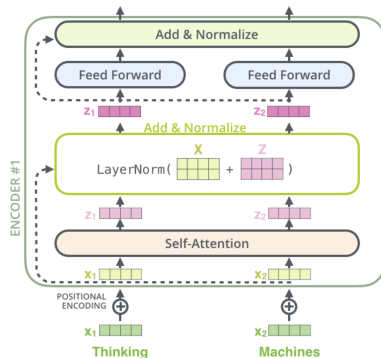


Zooming out: a complete encoder (cont.)

Two more components to stabilize and speed up training:

1. Residual pathway

- add input to output, $Z \leftarrow Z + X$
- an idea from **Residual Networks** to deal with **vanishing gradients**



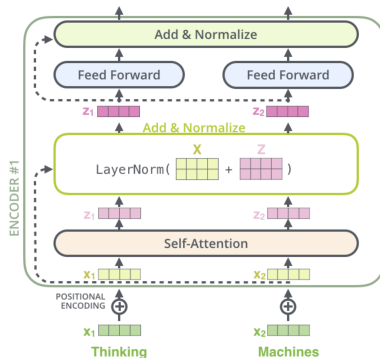
Zooming out: a complete encoder (cont.)

Two more components to stabilize and speed up training:

1. Residual pathway

- add input to output, $Z \leftarrow Z + X$
- an idea from **Residual Networks** to deal with **vanishing gradients**

2. Layer normalization:



Zooming out: a complete encoder (cont.)

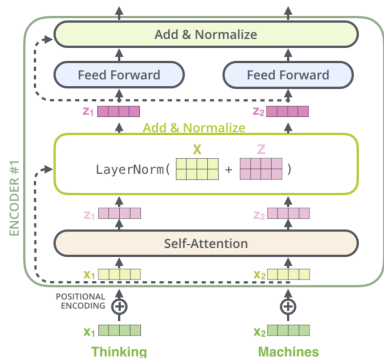
Two more components to stabilize and speed up training:

1. Residual pathway

- add input to output, $Z \leftarrow Z + X$
- an idea from **Residual Networks** to deal with **vanishing gradients**

2. Layer normalization:

- for each $z_t \in \mathbb{R}^d$, normalize it to **zero-mean** and **unit-variance** (**across features**)



Zooming out: a complete encoder (cont.)

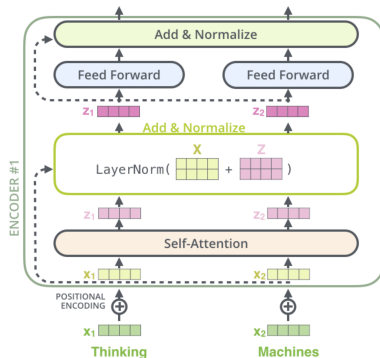
Two more components to stabilize and speed up training:

1. Residual pathway

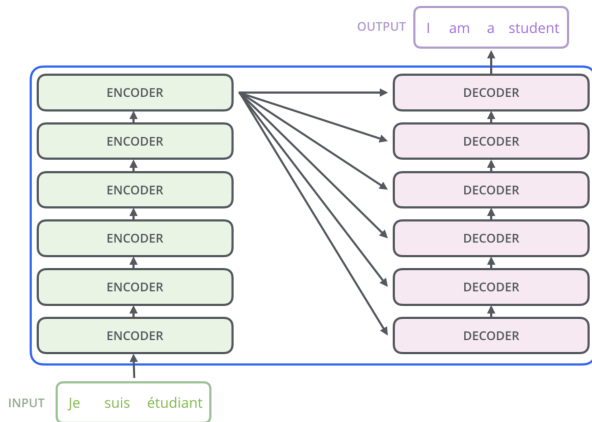
- add input to output, $Z \leftarrow Z + X$
- an idea from **Residual Networks** to deal with **vanishing gradients**

2. Layer normalization:

- for each $z_t \in \mathbb{R}^d$, normalize it to **zero-mean** and **unit-variance** (**across features**)
- similar but different from batch normalization (where you normalize each feature to zero-mean and unit-variance **across samples**)

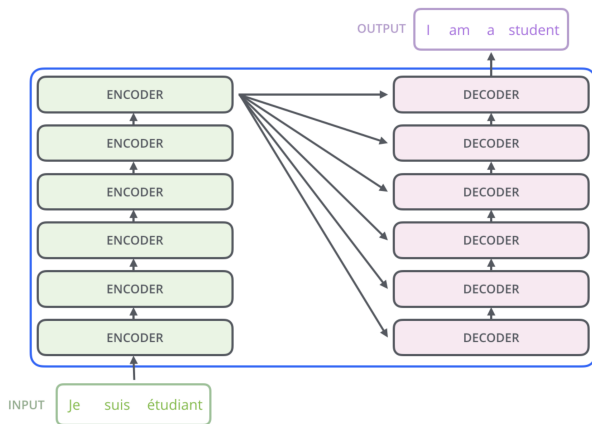


Zooming out: stacking encoders and decoders



Encoder: summarizes the input into a useful representation

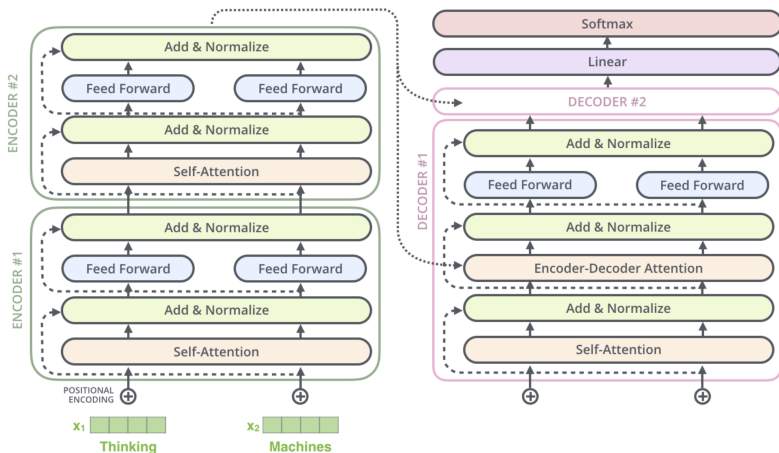
Zooming out: stacking encoders and decoders



Encoder: summarizes the input into a useful representation

Decoder: generates outputs

A closer look at decoders



Extra component: **encoder-decoder attention**

Encoder-decoder attention

Same idea as self-attention, except key and value matrices are computed using output $\mathbf{Z}_{\text{enc}}$ of the final encoder:

Encoder-decoder attention

Same idea as self-attention, except key and value matrices are computed using output $\mathbf{Z}_{\text{enc}}$ of the final encoder:

- query matrix $\mathbf{Q} = \mathbf{X}\mathbf{W}_Q \in \mathbb{R}^{T_{\text{dec}} \times d_k}$ (as usual)

Encoder-decoder attention

Same idea as self-attention, except key and value matrices are computed using output $\mathbf{Z}_{\text{enc}}$ of the final encoder:

- query matrix $\mathbf{Q} = \mathbf{X}\mathbf{W}_Q \in \mathbb{R}^{T_{\text{dec}} \times d_k}$ (as usual)
- key matrix $\mathbf{K}_{\text{enc}} = \mathbf{Z}_{\text{enc}}\mathbf{W}_K \in \mathbb{R}^{T_{\text{enc}} \times d_k}$

Encoder-decoder attention

Same idea as self-attention, except key and value matrices are computed using output $\mathbf{Z}_{\text{enc}}$ of the final encoder:

- query matrix $\mathbf{Q} = \mathbf{X}\mathbf{W}_Q \in \mathbb{R}^{T_{\text{dec}} \times d_k}$ (as usual)
- key matrix $\mathbf{K}_{\text{enc}} = \mathbf{Z}_{\text{enc}}\mathbf{W}_K \in \mathbb{R}^{T_{\text{enc}} \times d_k}$
- value matrix $\mathbf{V}_{\text{enc}} = \mathbf{Z}_{\text{enc}}\mathbf{W}_V \in \mathbb{R}^{T_{\text{enc}} \times d_v}$

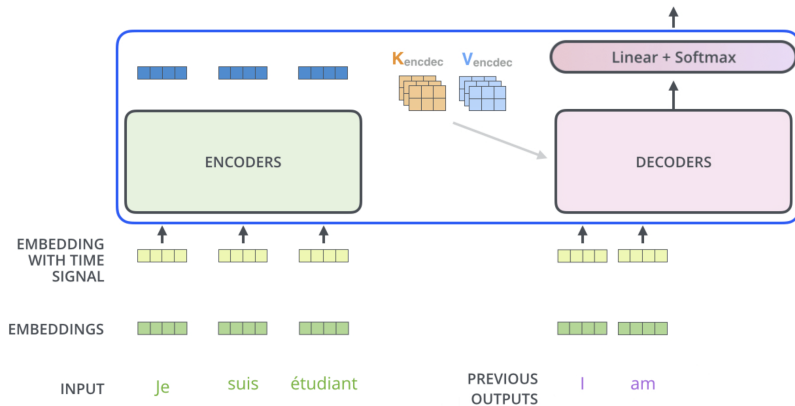
Encoder-decoder attention

Same idea as self-attention, except key and value matrices are computed using output $\mathbf{Z}_{\text{enc}}$ of the final encoder:

- query matrix $\mathbf{Q} = \mathbf{X}\mathbf{W}_Q \in \mathbb{R}^{T_{\text{dec}} \times d_k}$ (as usual)
- key matrix $\mathbf{K}_{\text{enc}} = \mathbf{Z}_{\text{enc}}\mathbf{W}_K \in \mathbb{R}^{T_{\text{enc}} \times d_k}$
- value matrix $\mathbf{V}_{\text{enc}} = \mathbf{Z}_{\text{enc}}\mathbf{W}_V \in \mathbb{R}^{T_{\text{enc}} \times d_v}$

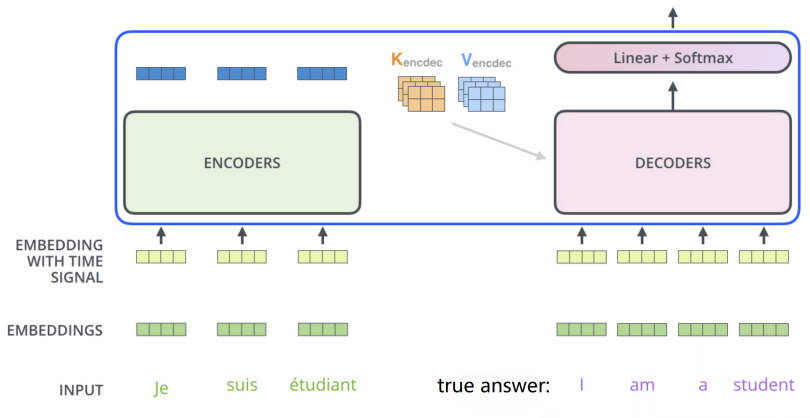
Intuition: find answer from the encoded representation of original inputs

Generating answers



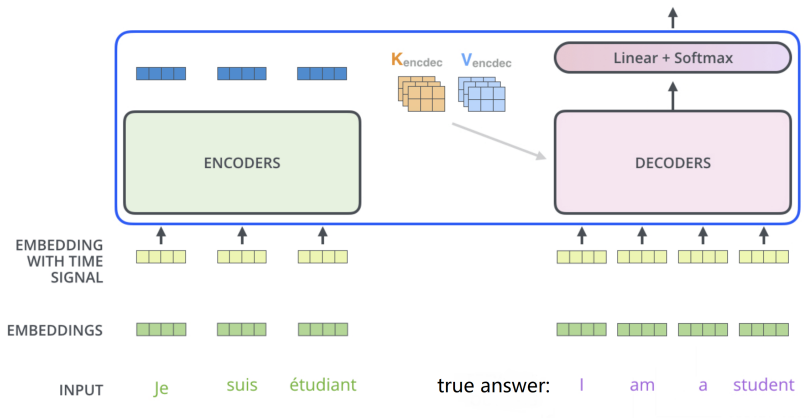
Use previously generated text as inputs of the decoder

Training



Use cross-entropy loss again,

Training



Use cross-entropy loss again, and apply

- **teacher forcing:** use the true answer as inputs of the decoder

Teaching forcing and causal mask

In teaching forcing, to avoid generating a word by peeking at future words, need to use a **causal mask** in the decoder's **self-attention heads**:

Teaching forcing and causal mask

In teaching forcing, to avoid generating a word by peeking at future words, need to use a **causal mask** in the decoder's **self-attention heads**:

$$QK^{\top} \leftarrow QK^{\top} + \mathbf{M}$$

Teaching forcing and causal mask

In teaching forcing, to avoid generating a word by peeking at future words, need to use a **causal mask** in the decoder's **self-attention heads**:

$$QK^{\top} \leftarrow QK^{\top} + \mathbf{M}$$

where

$$\mathbf{M} = \begin{bmatrix} 0 & -\infty & -\infty & \cdots & -\infty \\ 0 & 0 & -\infty & \cdots & -\infty \\ 0 & 0 & 0 & \cdots & -\infty \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & 0 & \cdots & 0 \end{bmatrix} \in \mathbb{R}^{T_{\text{dec}} \times T_{\text{dec}}}$$

Teaching forcing and causal mask

In teaching forcing, to avoid generating a word by peeking at future words, need to use a **causal mask** in the decoder's **self-attention heads**:

$$QK^{\top} \leftarrow QK^{\top} + \mathbf{M}$$

where

$$\mathbf{M} = \begin{bmatrix} 0 & -\infty & -\infty & \cdots & -\infty \\ 0 & 0 & -\infty & \cdots & -\infty \\ 0 & 0 & 0 & \cdots & -\infty \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & 0 & \cdots & 0 \end{bmatrix} \in \mathbb{R}^{T_{\text{dec}} \times T_{\text{dec}}}$$

so a word at position t never attends to words at positions $> t$

Teaching forcing and causal mask

In teaching forcing, to avoid generating a word by peeking at future words, need to use a **causal mask** in the decoder's **self-attention heads**:

$$QK^{\top} \leftarrow QK^{\top} + \mathbf{M}$$

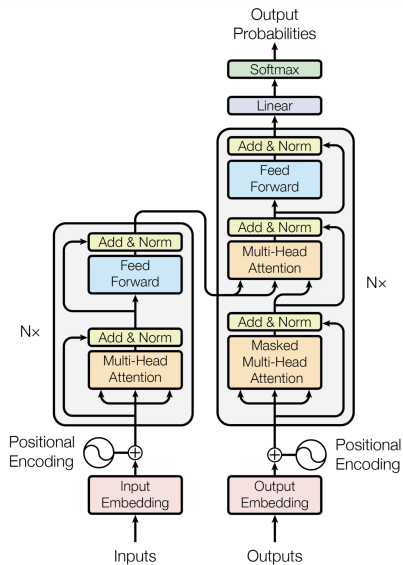
where

$$\mathbf{M} = \begin{bmatrix} 0 & -\infty & -\infty & \cdots & -\infty \\ 0 & 0 & -\infty & \cdots & -\infty \\ 0 & 0 & 0 & \cdots & -\infty \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & 0 & \cdots & 0 \end{bmatrix} \in \mathbb{R}^{T_{\text{dec}} \times T_{\text{dec}}}$$

so a word at position t never attends to words at positions $> t$

Q: should we use causal mask for encoder-decoder attention heads?

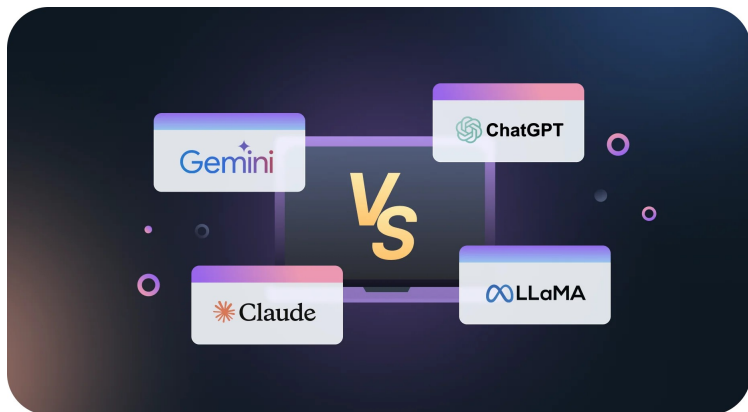
That's it!



Large language models

Large language models (LLMs) are all based on transformers

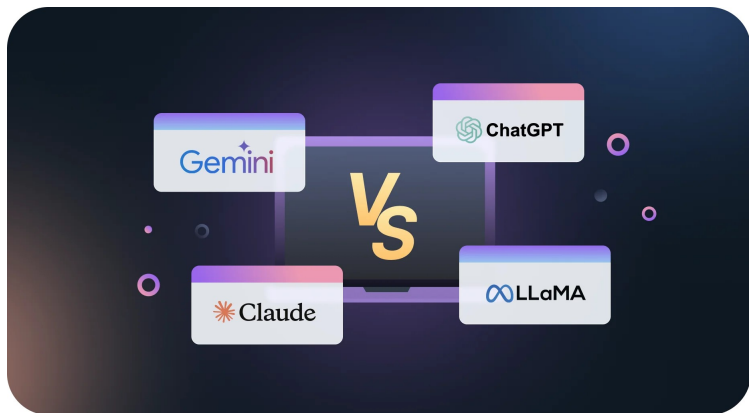
- estimated #parameters for GPT5: **trillions**



Large language models

Large language models (LLMs) are all based on transformers

- estimated #parameters for GPT5: **trillions**



A cool 3D visualization of a **nano-GPT**: <https://bbycroft.net/llm>

Training Large language models

Unsupervised pre-training

- via next word prediction using a *huge* training set (e.g., the entire internet)

Training Large language models

Unsupervised pre-training

- via next word prediction using a *huge* training set (e.g., the entire internet)

Fine-tuning

- using a *labeled* dataset for a specific task (translation, question answering, etc.)

Training Large language models

Unsupervised pre-training

- via next word prediction using a *huge* training set (e.g., the entire internet)

Fine-tuning

- using a *labeled* dataset for a specific task (translation, question answering, etc.)

Reinforcement Learning with Human Feedback (RLHF)

- get *preference feedback* from human: *which answer is better?*
- more on this in the next two weeks