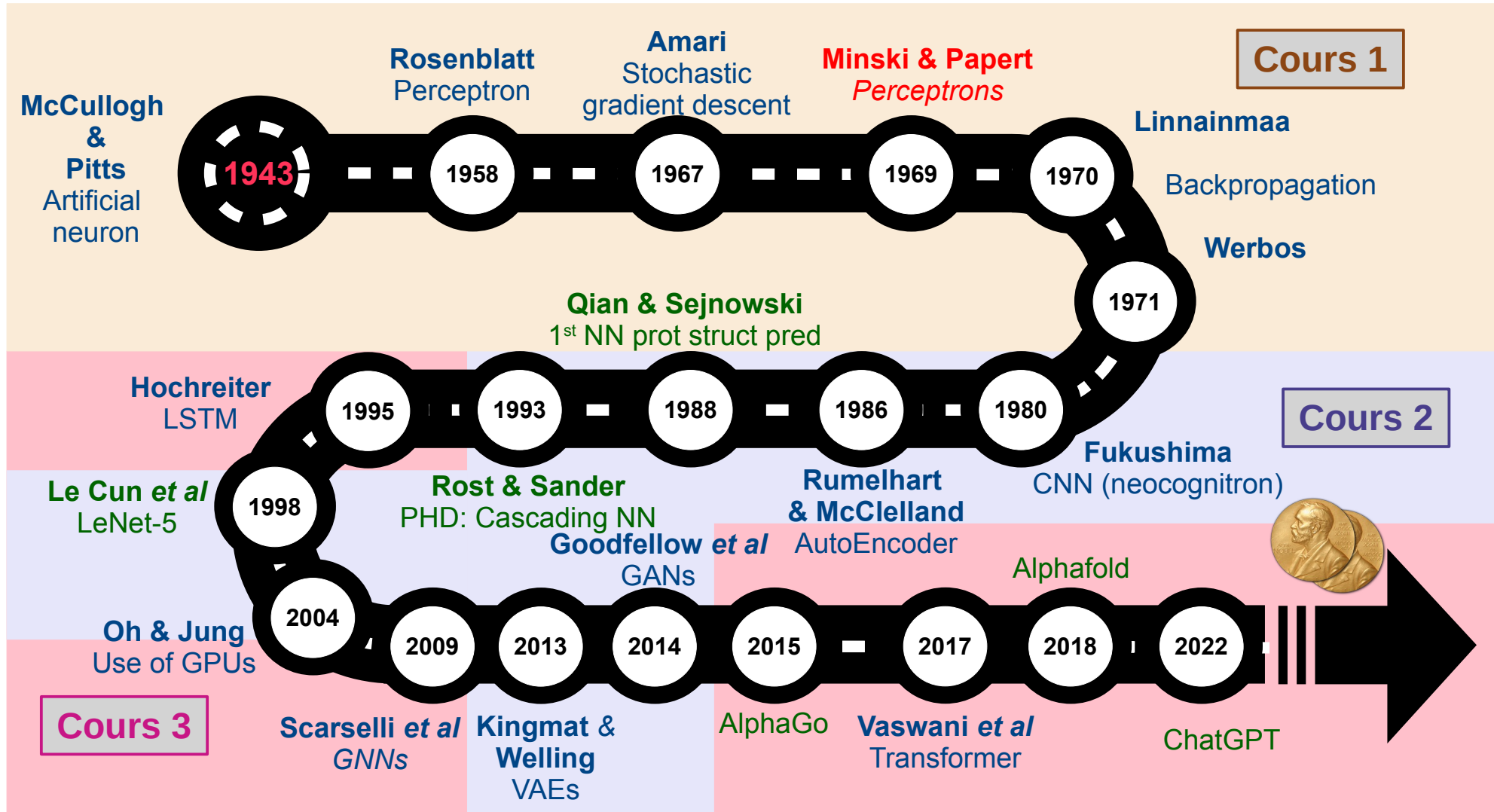


Beyond MLPs

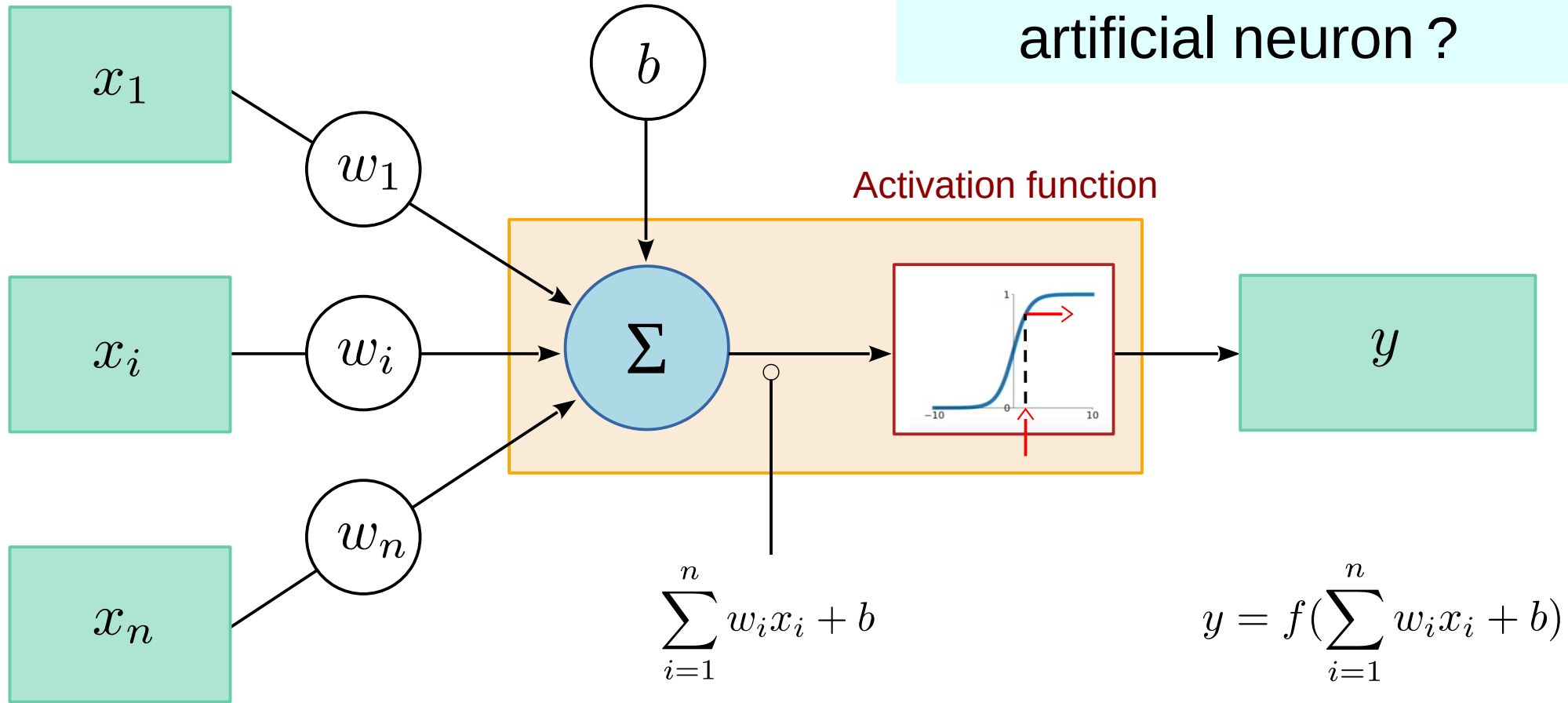
Part 2/2: RNNs, Attention and GNNs

Nicolas Gambardella

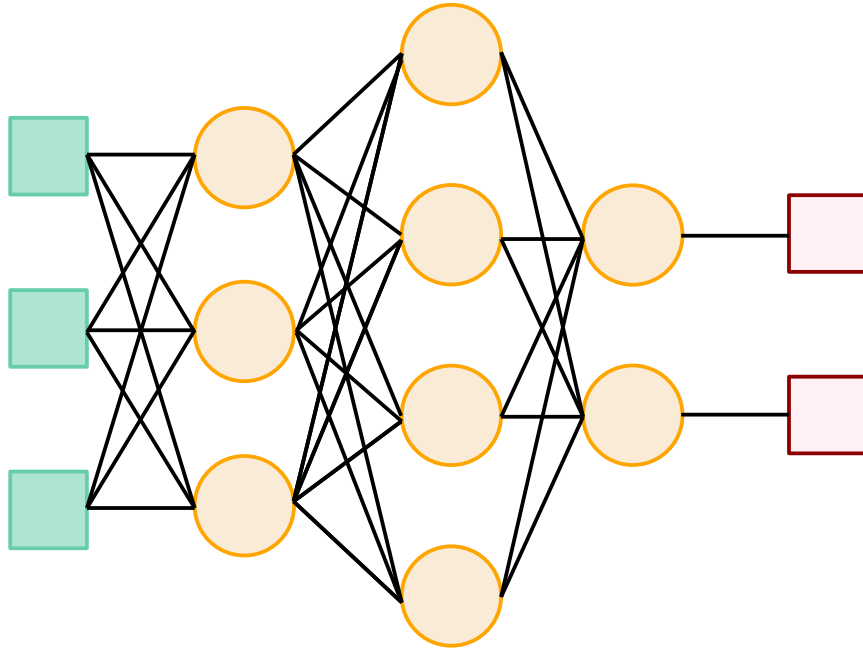
nicolas.gambardella@univ-lille.fr



What is an artificial neuron ?

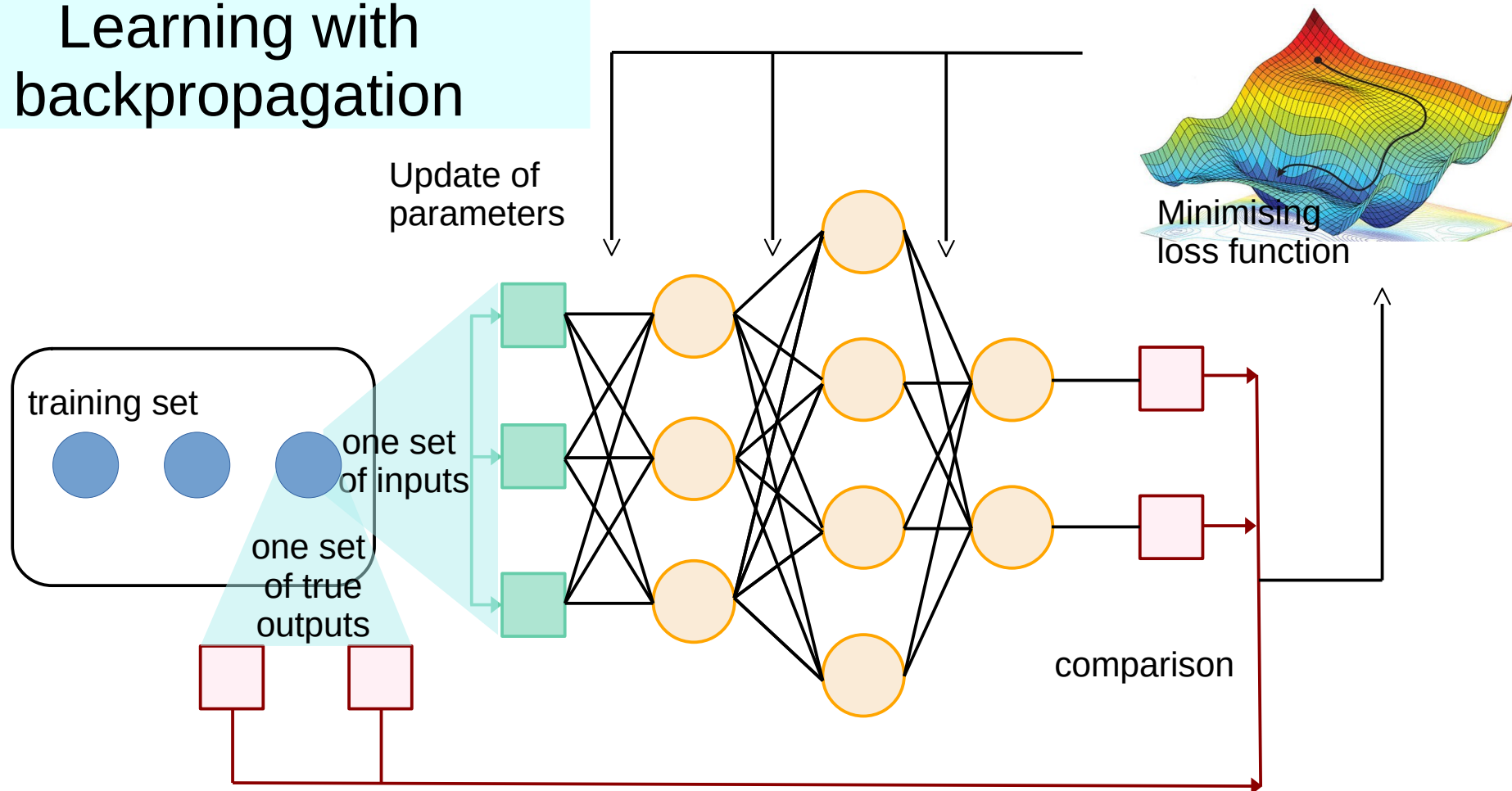


All inputs can be independent and everything connected to everything

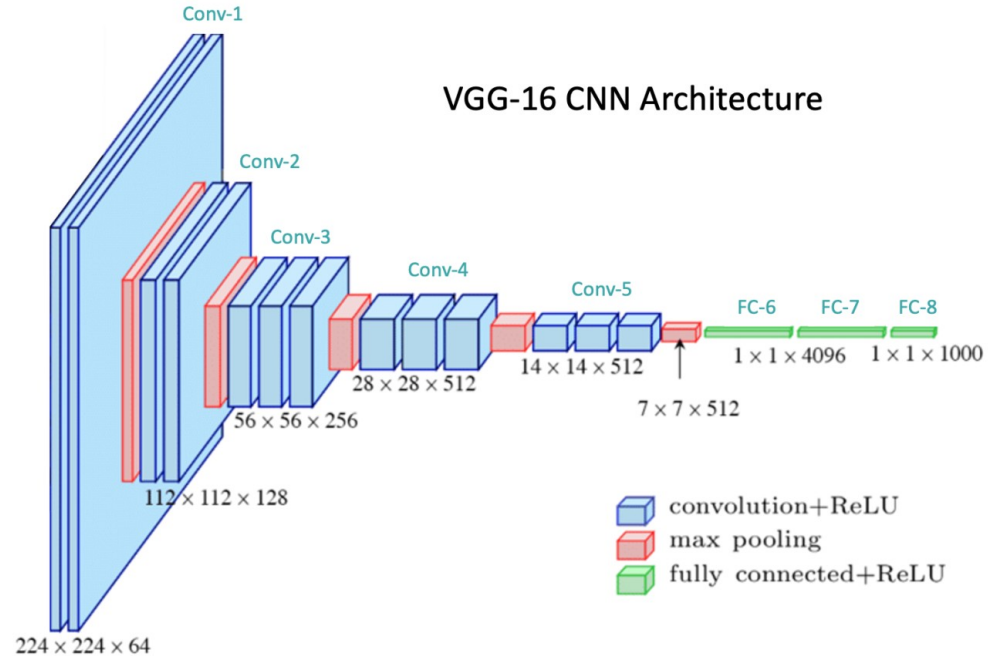
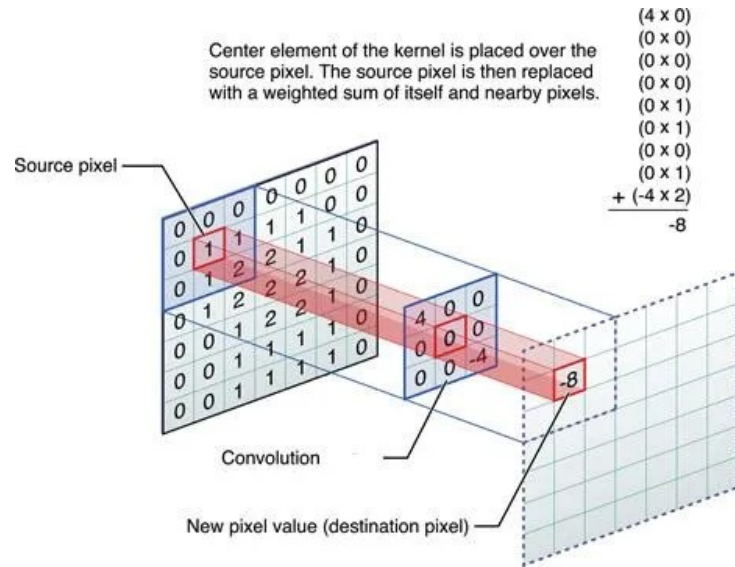


Multi-Layer Perceptrons (MLP)
or Dense neural networks (DNN)
made of Fully Connected layers (FC)

Learning with backpropagation



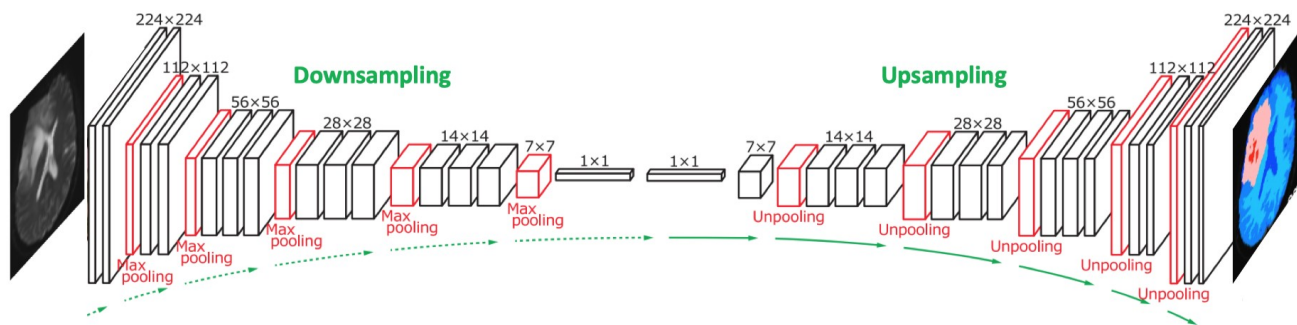
We can detect local features by linking neighbouring inputs



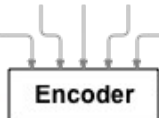
Deep Convolutional Neural Networks (CNN)

Encoder networks can embed information in a latent space

Decoder networks can reconstruct the information from it

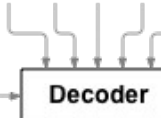


"le chat est noir" <EOS>
[02 85 03 12 99]

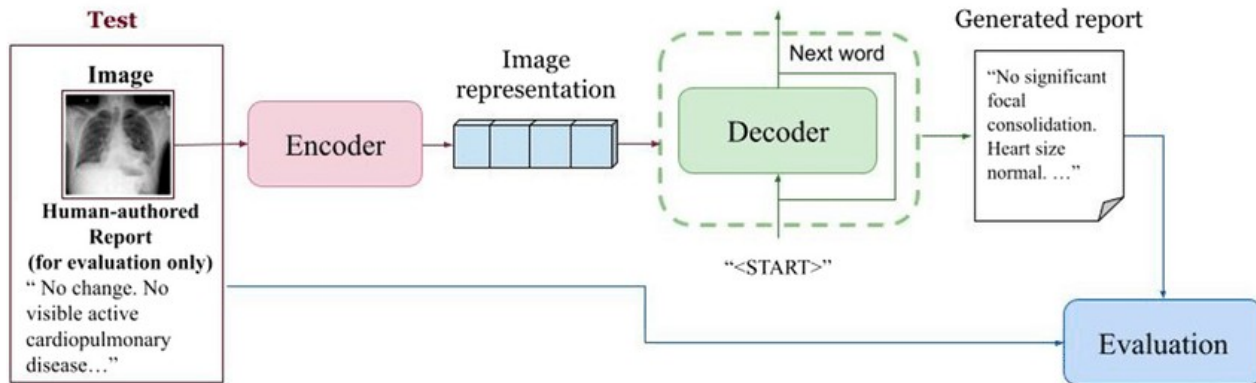


Context

<SOS> "the cat is black"
[00 42 82 16 04]

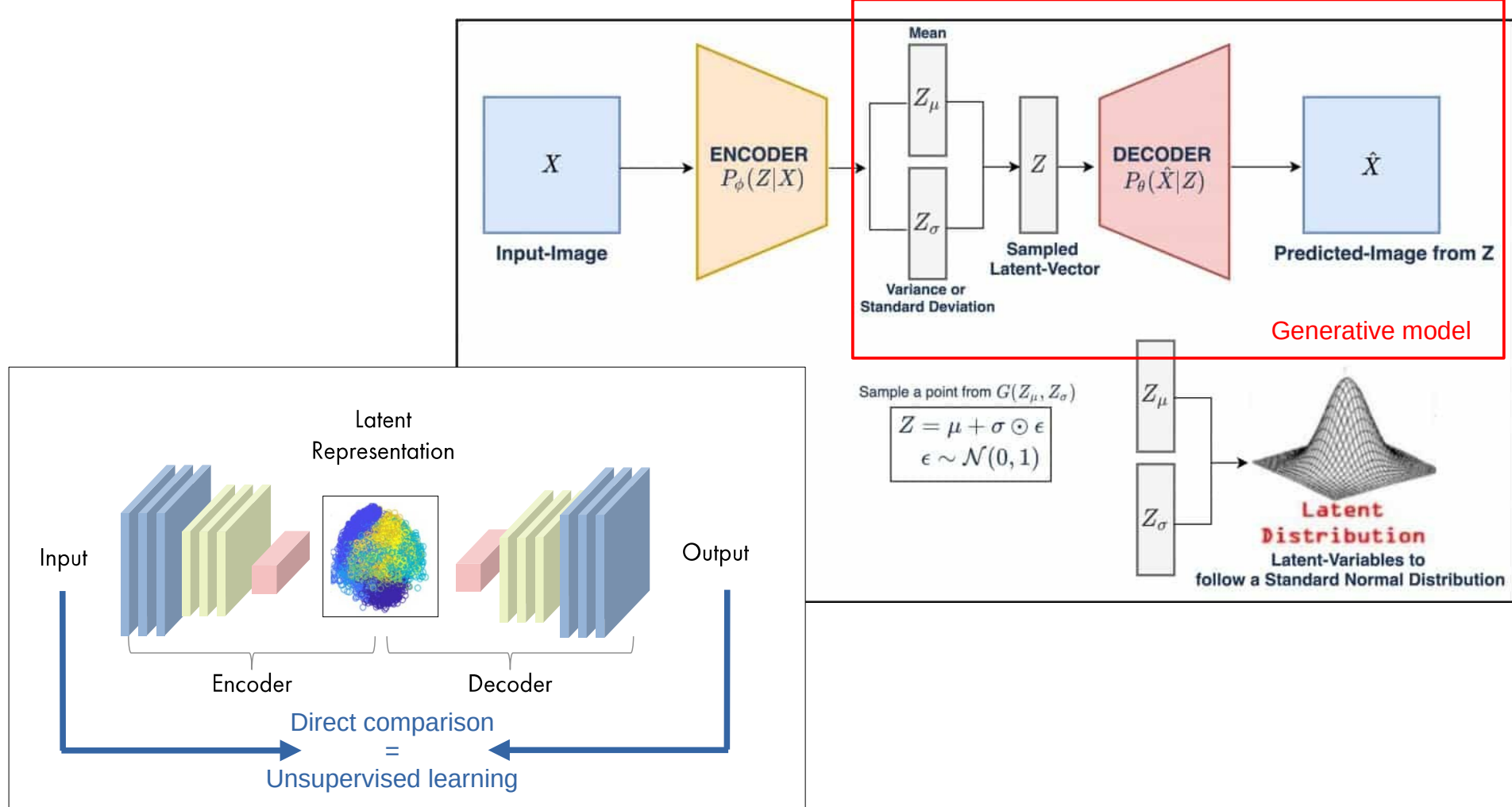


[42 82 16 04 99]
"the cat is black" <EOS>

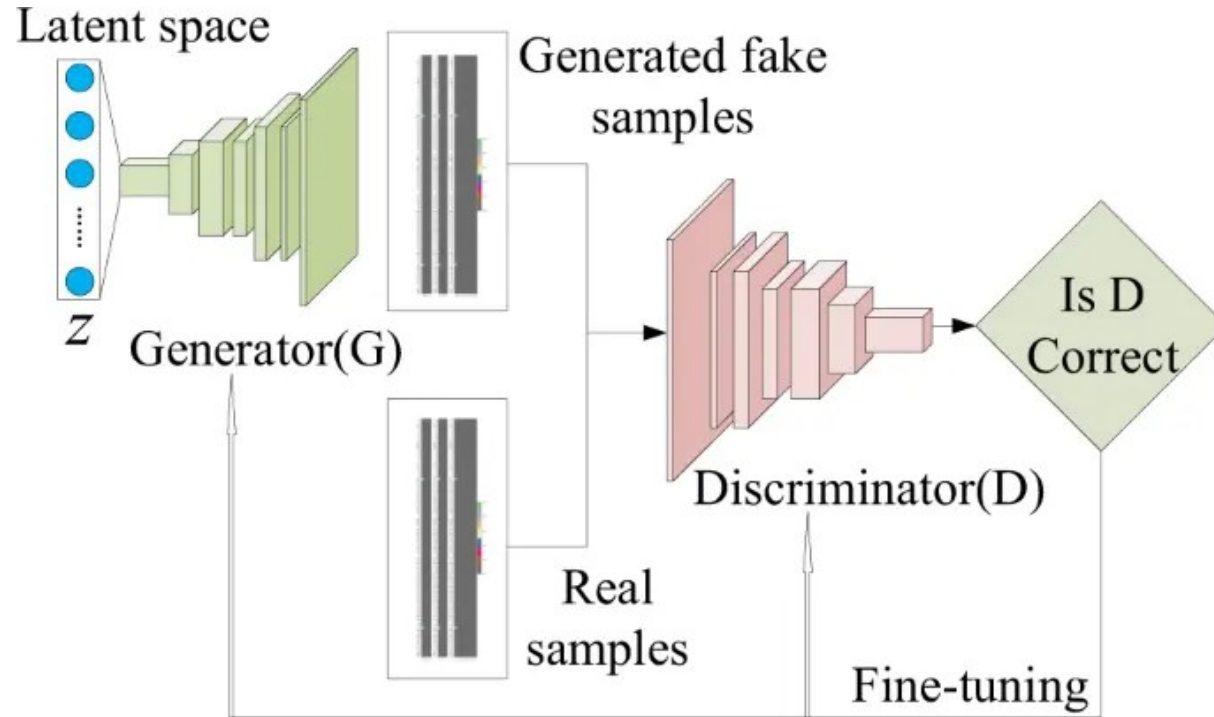


AutoEncoders can train themselves unsupervised

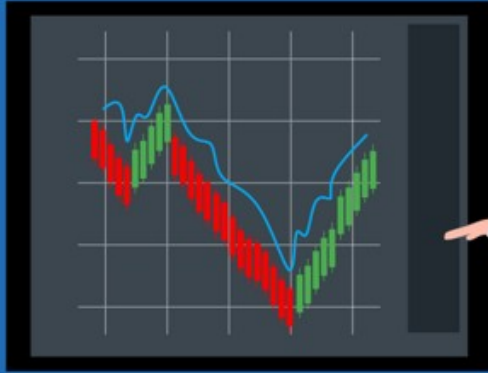
Variational AutoEncoders learn mean and std distributions



Generative Adversarial Networks learn by trying to deceive themselves



Time series and sequences of variable lengths



Time series and sequences of variable lengths

Speech recognition



"The quick brown fox jumped over the lazy dog."

Music generation



Sentiment classification

"There is nothing to like in this movie."



DNA sequence analysis

AGCCCCTGTGAGGAACTAG



AGCCCCTGTGAGGAACTAG

Machine translation

Voulez-vous chanter avec moi?



Do you want to sing with me?

Video activity recognition



Running

Name entity recognition

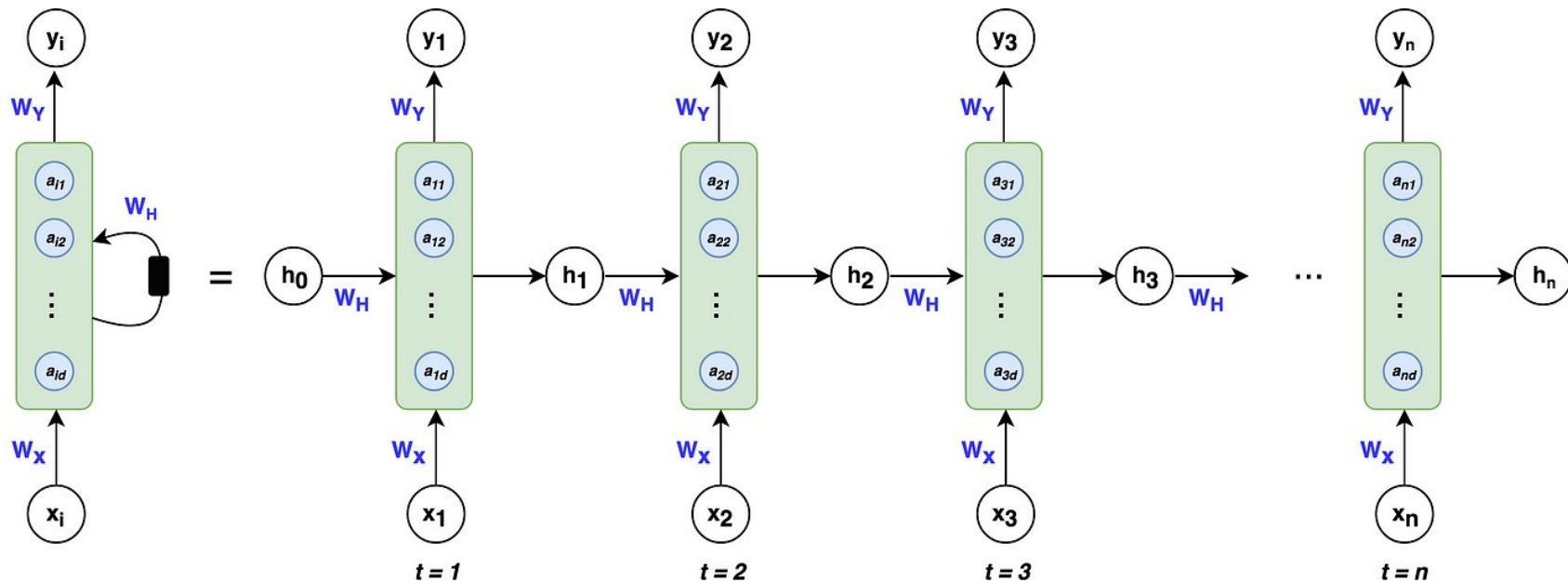
Yesterday, Harry Potter met Hermione Granger.



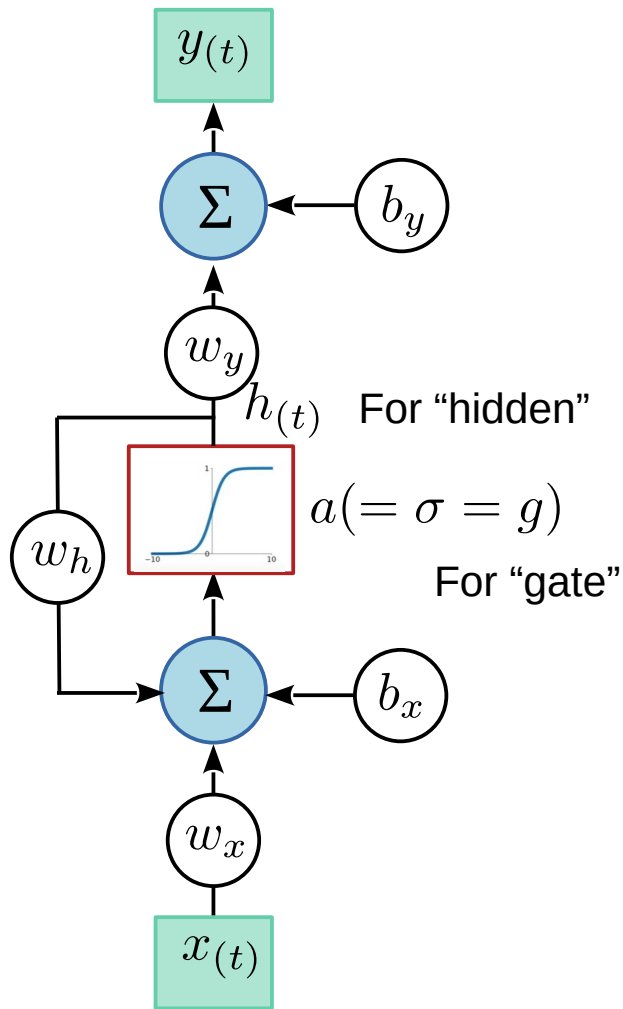
Yesterday, **Harry Potter** met **Hermione Granger**.

Andrew Ng

Recurrent Neural Networks: successive inputs are not independent



RNN: 1 cell (here, 1 neuron)



NB: implicit “identity” activation function

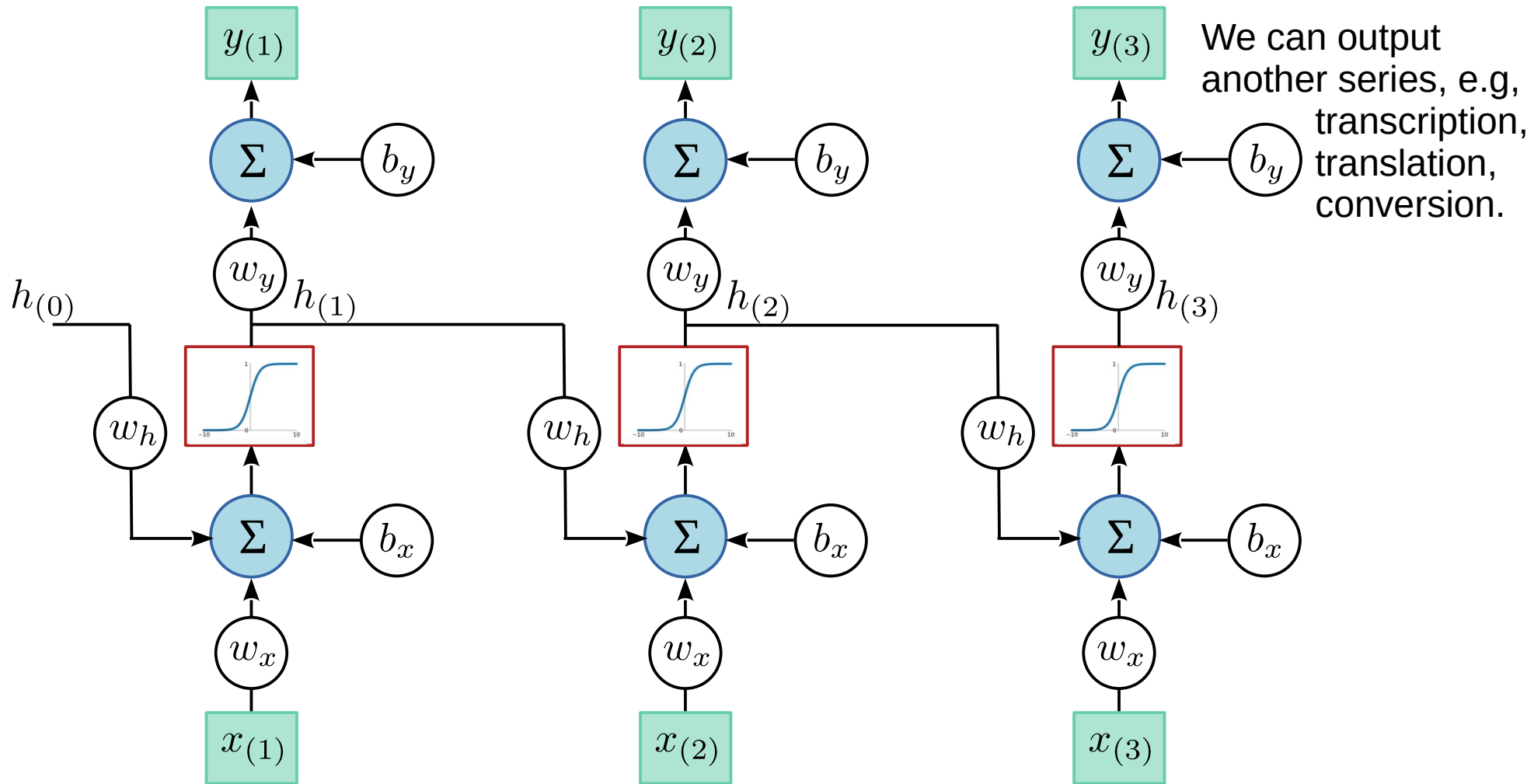
$$y(t) = w_y \times h(t) + b_y$$

not the same timepoint

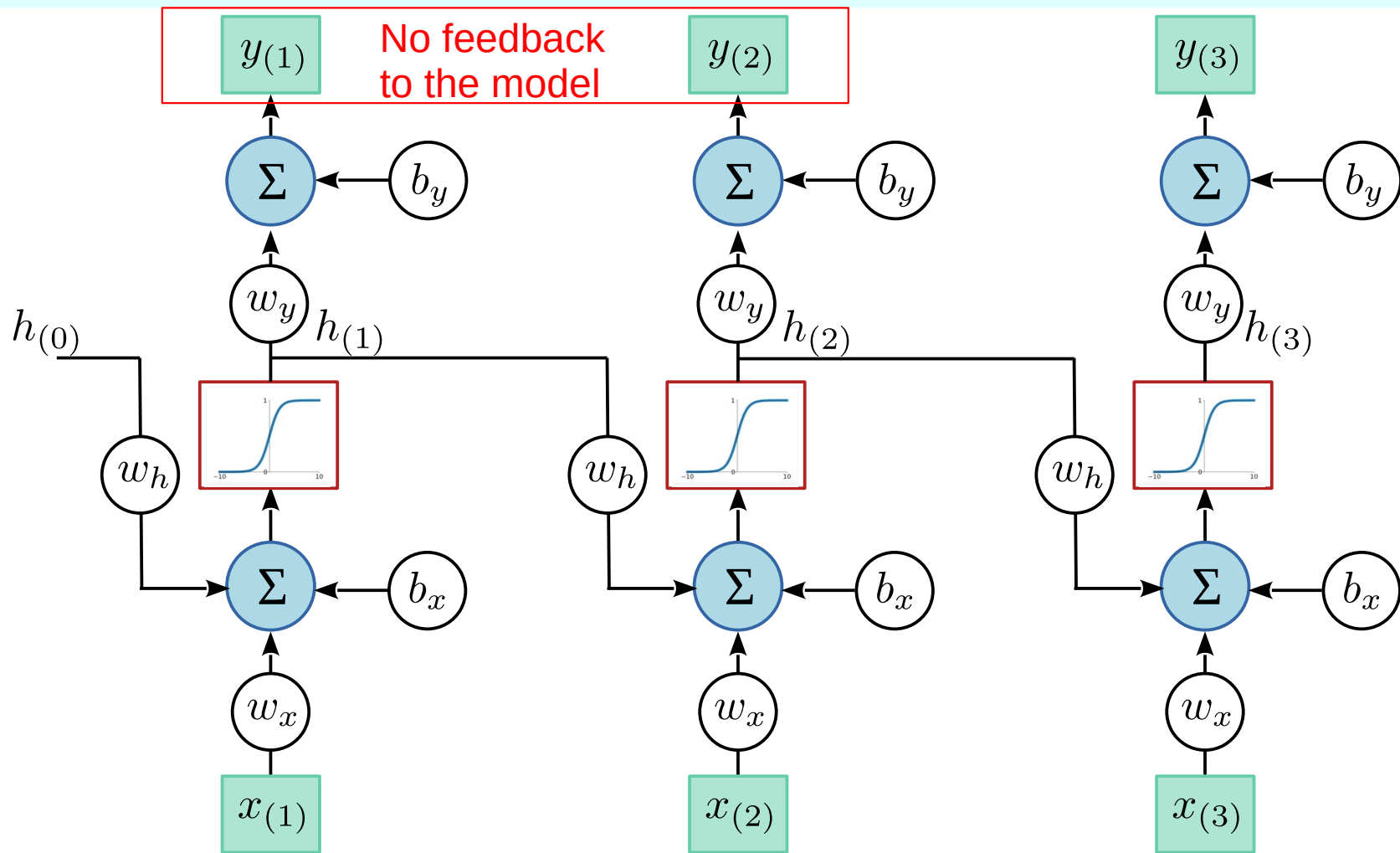
$$h(t) = a(w_x \times x(t) + b_x + w_h \times h_{(t-1)})$$

“memory”

RNN: 1 cell - unfolding

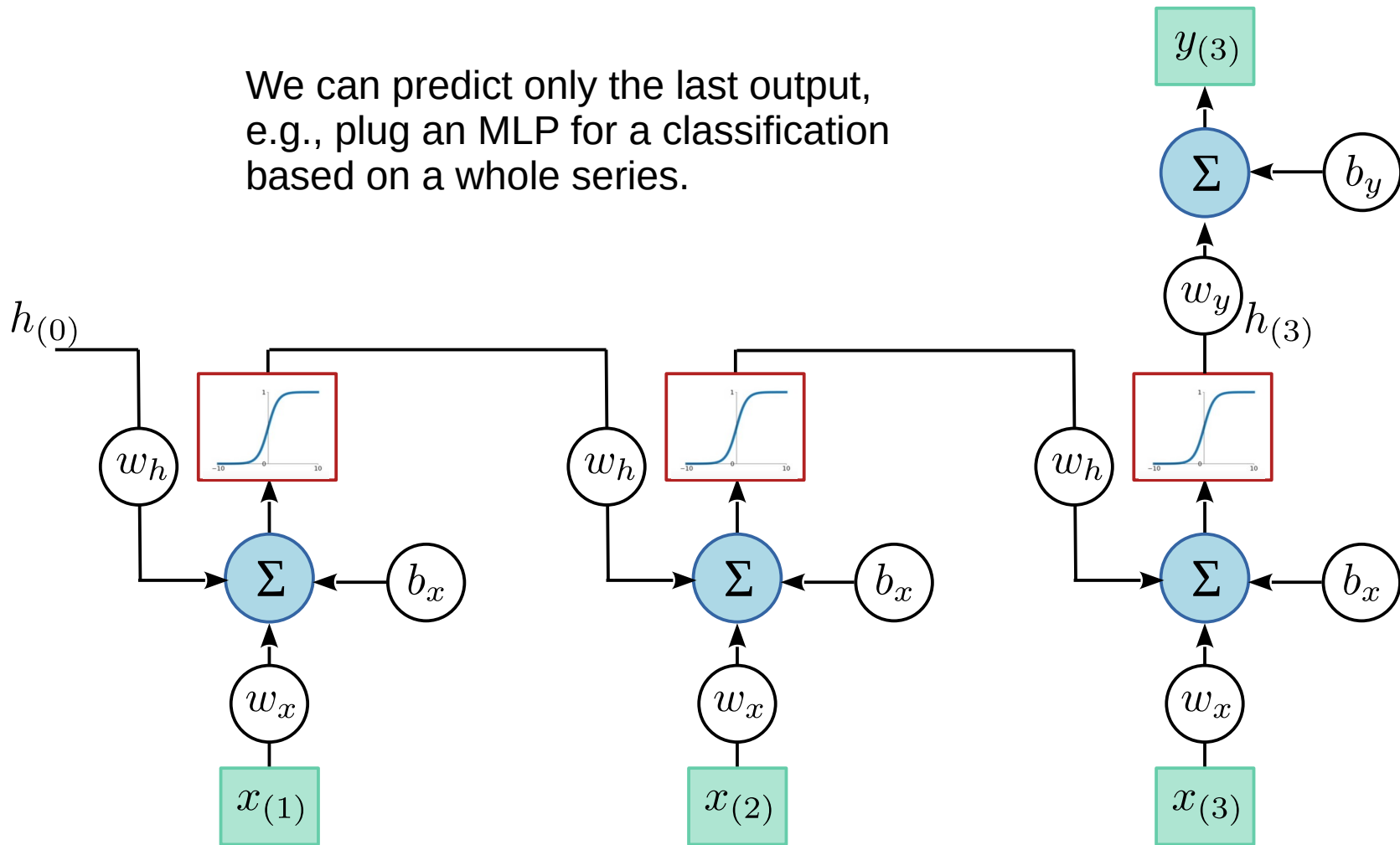


RNN: 1 cell - unfolding

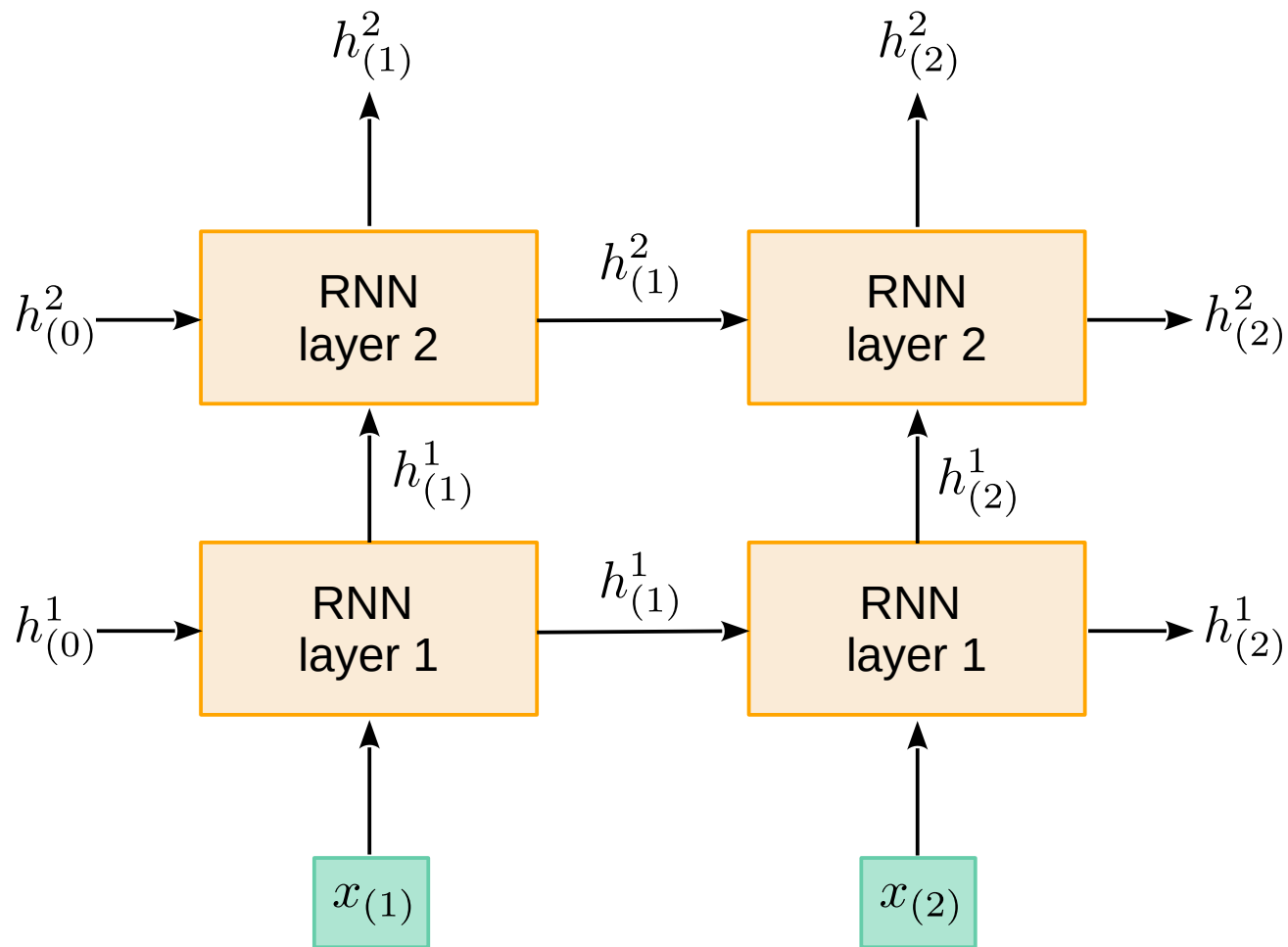


RNN: 1 cell - unfolding

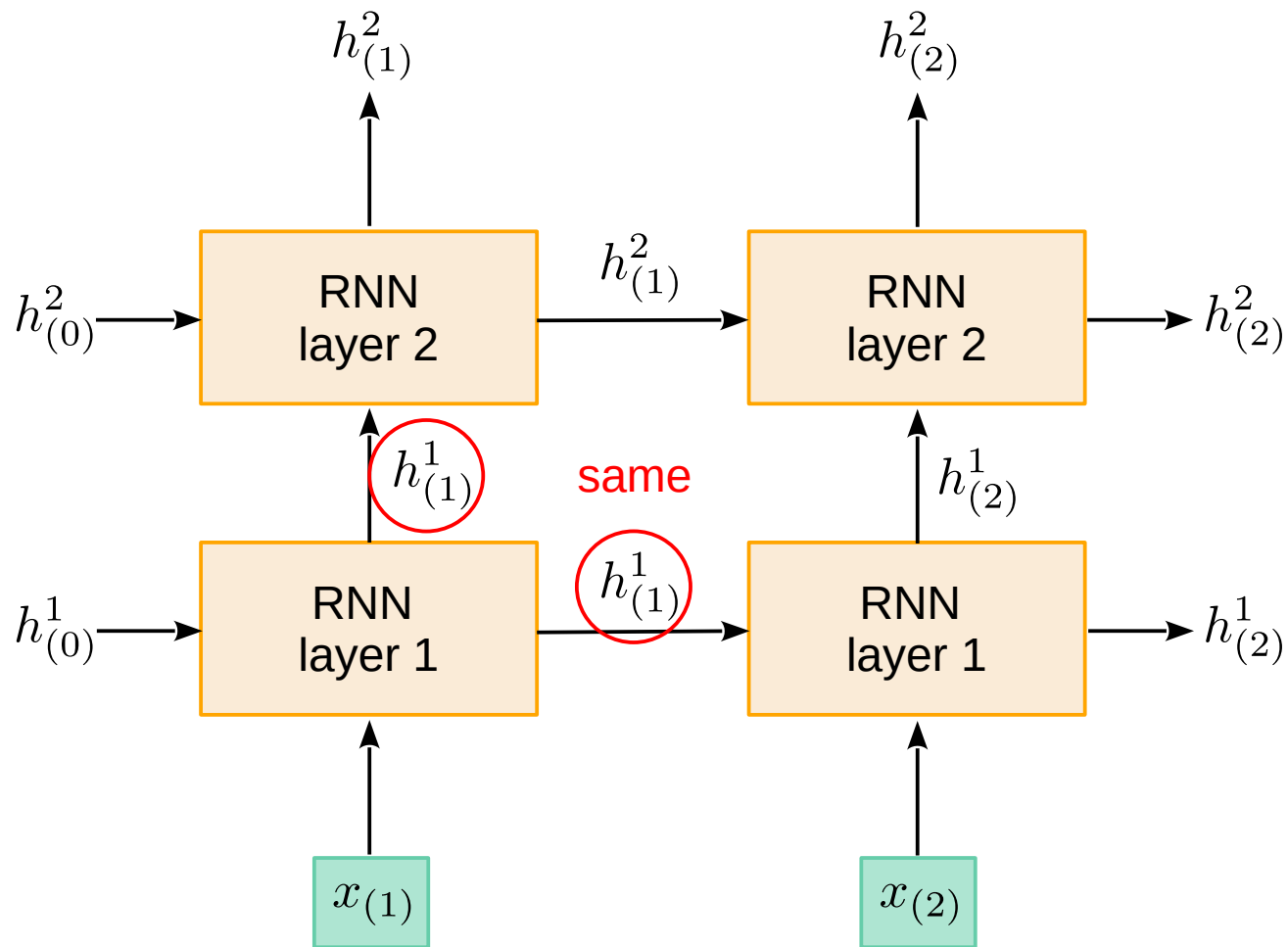
We can predict only the last output, e.g., plug an MLP for a classification based on a whole series.



Stacked RNNs

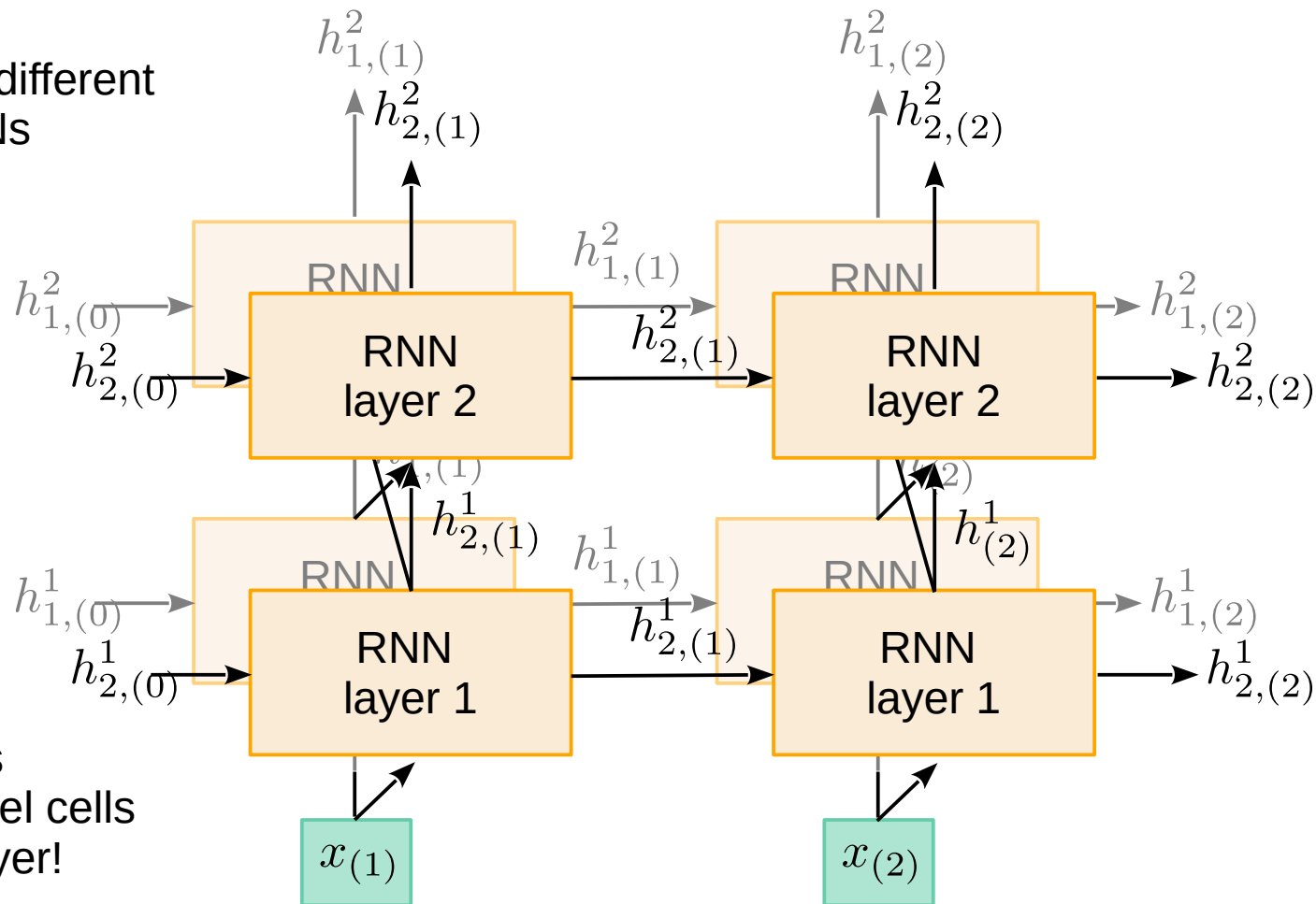


Stacked RNN



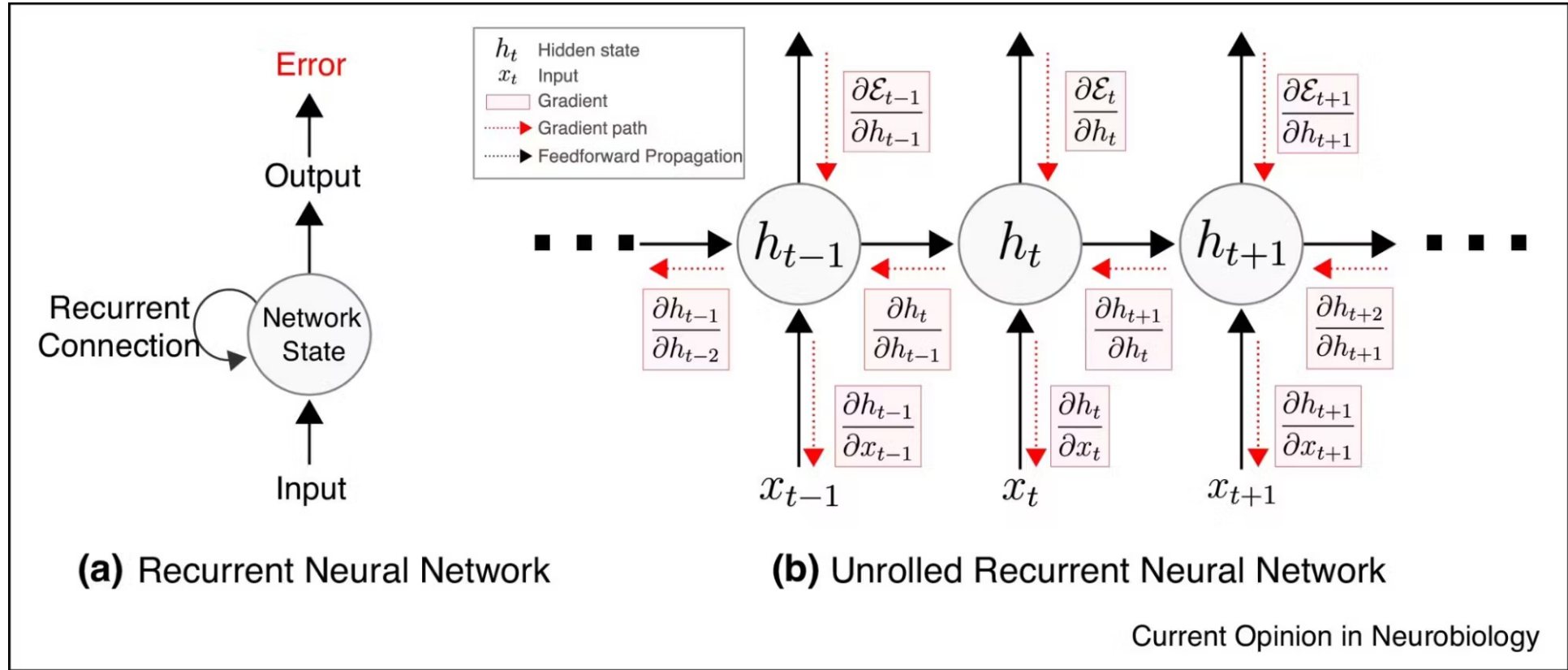
Several RNNs may learn different patterns in parallel

Same idea as different kernels in CNNs

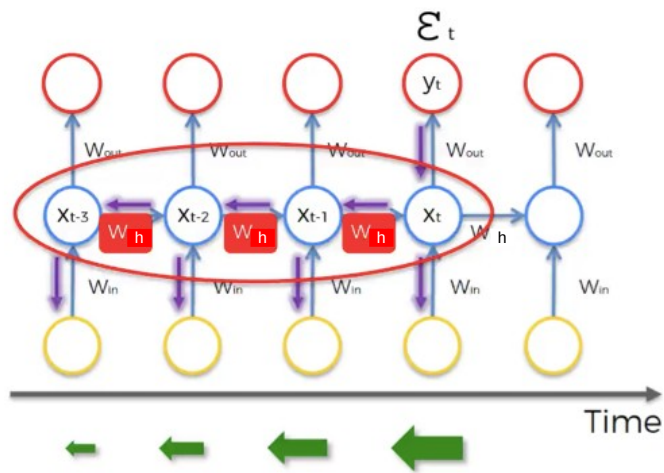


No connexions
between parallel cells
of the same layer!

Learning by backpropagation in time



Exploding and vanishing gradients



Source:
SuperDataScience

$$\frac{\partial \mathcal{E}}{\partial \theta} = \sum_{1 \leq t \leq T} \frac{\partial \mathcal{E}_t}{\partial \theta}$$

$$\frac{\partial \mathcal{E}_t}{\partial \theta} = \sum_{1 \leq k \leq t} \left(\frac{\partial \mathcal{E}_t}{\partial \mathbf{x}_t} \frac{\partial \mathbf{x}_t}{\partial \mathbf{x}_k} \frac{\partial^+ \mathbf{x}_k}{\partial \theta} \right)$$

$$\frac{\partial \mathbf{x}_t}{\partial \mathbf{x}_k} = \prod_{t \geq i > k} \frac{\partial \mathbf{x}_i}{\partial \mathbf{x}_{i-1}} = \prod_{t \geq i > k} \mathbf{W}_h^T \text{sc} \text{diag}(\sigma'(\mathbf{x}_{i-1}))$$

$W_h \sim \text{small} \rightarrow \text{Vanishing}$
 $W_h \sim \text{large} \rightarrow \text{Exploding}$

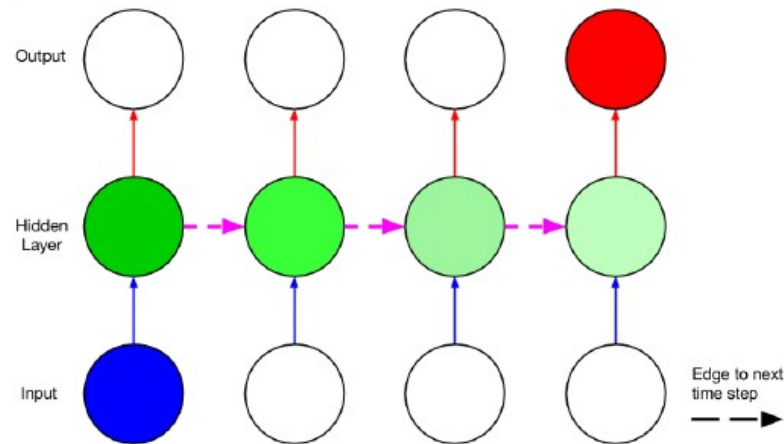
E.g.: activation function = ReLU

$$\frac{\partial x_i}{\partial x_{i-1}} = w_h$$

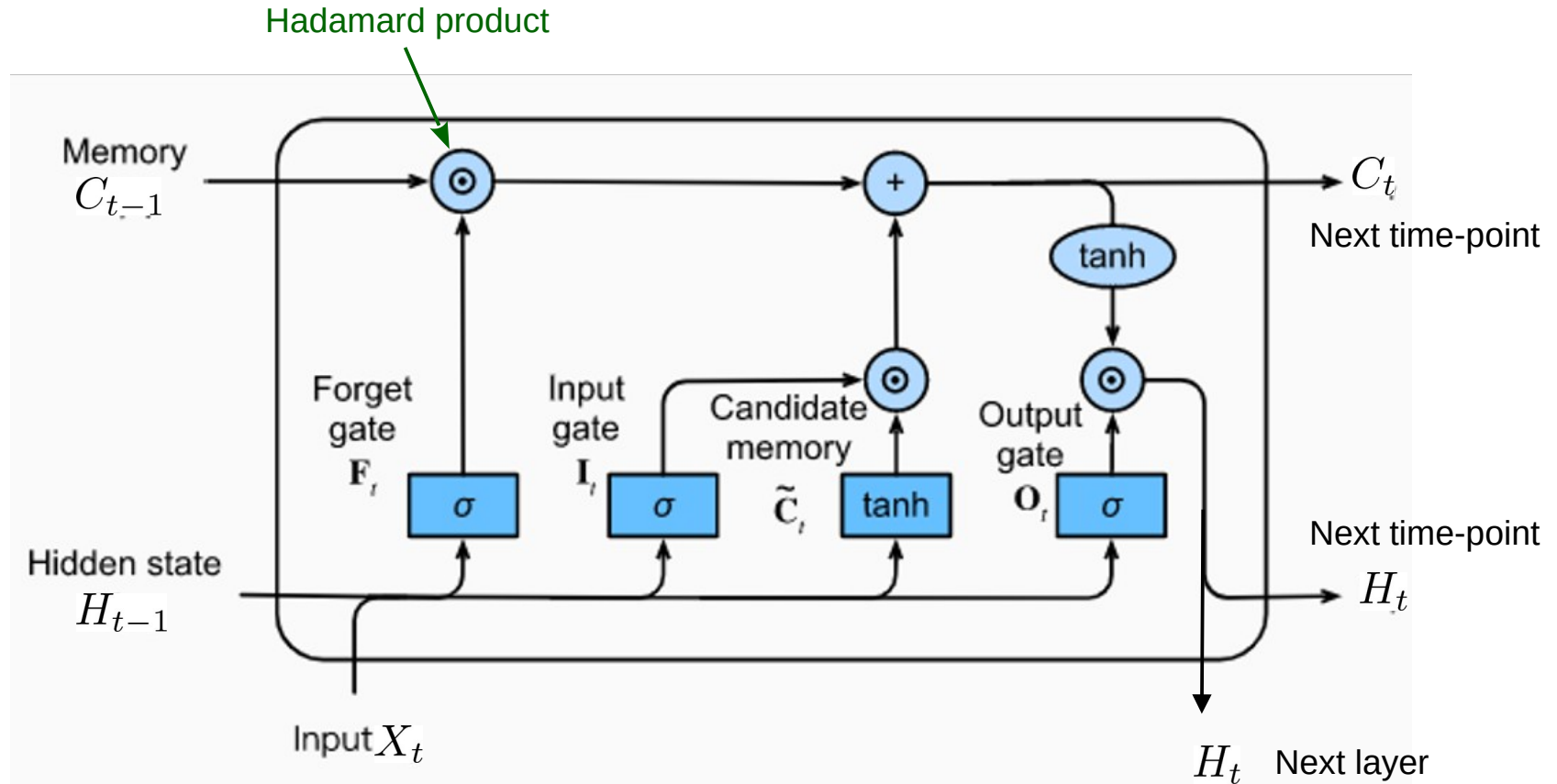
$$w_h = 0.1; \frac{\partial x_{10}}{\partial x_1} = w_h^{10} = 0.0000000001$$

$$w_h = 10; \frac{\partial x_{10}}{\partial x_1} = w_h^{10} = 10000000000$$

Green = sensitivity of output on input



Solution: Long Short-Term Memory (LSTM)

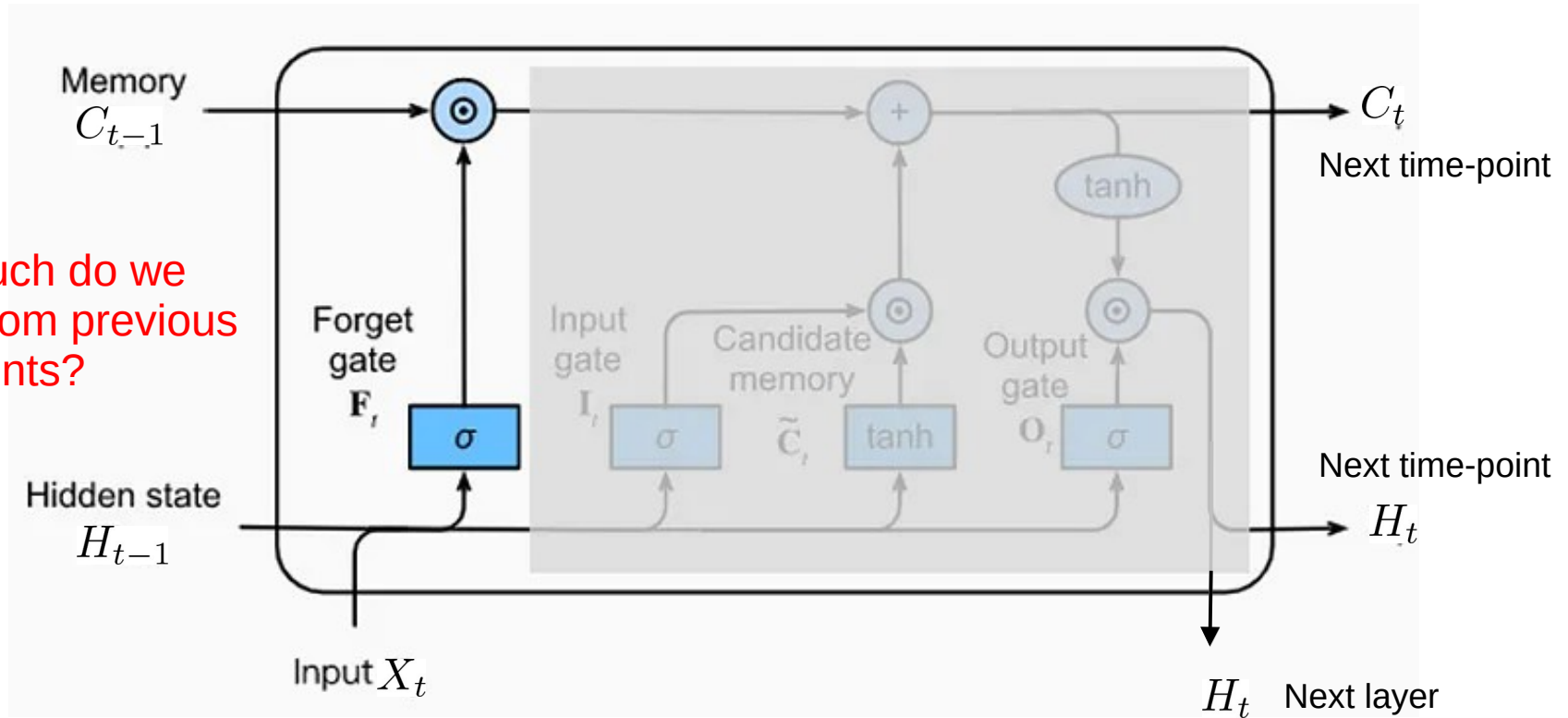


Hochreiter and Schmidhuber (1997) Long short-term memory. *Neur Comput*, 9(8):1735-1780

Source: Ottavio Calzone (2002) An Intuitive Explanation of LSTM. <https://medium.com/@ottaviocalzone>

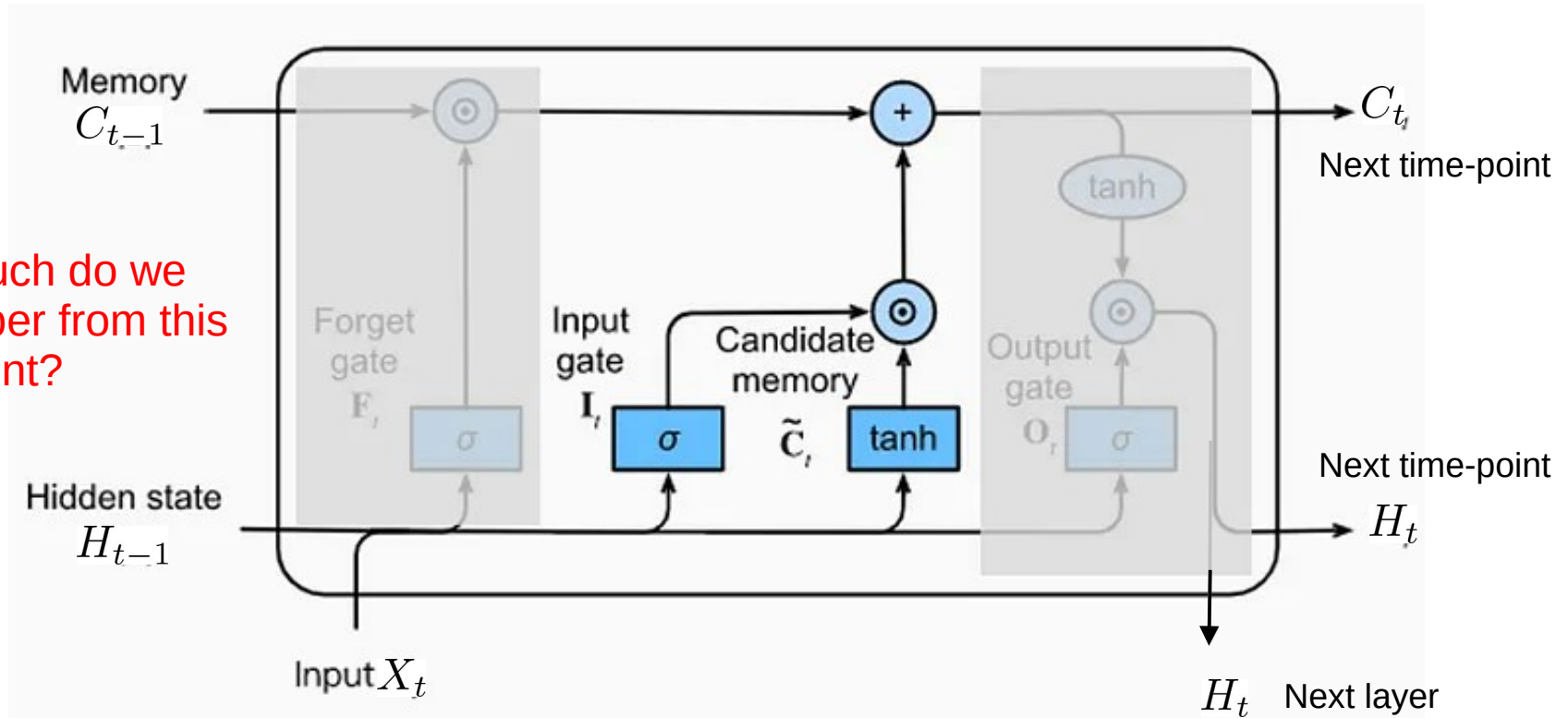
Solution: Long Short-Term Memory (LSTM)

1-How much do we forget from previous time-points?



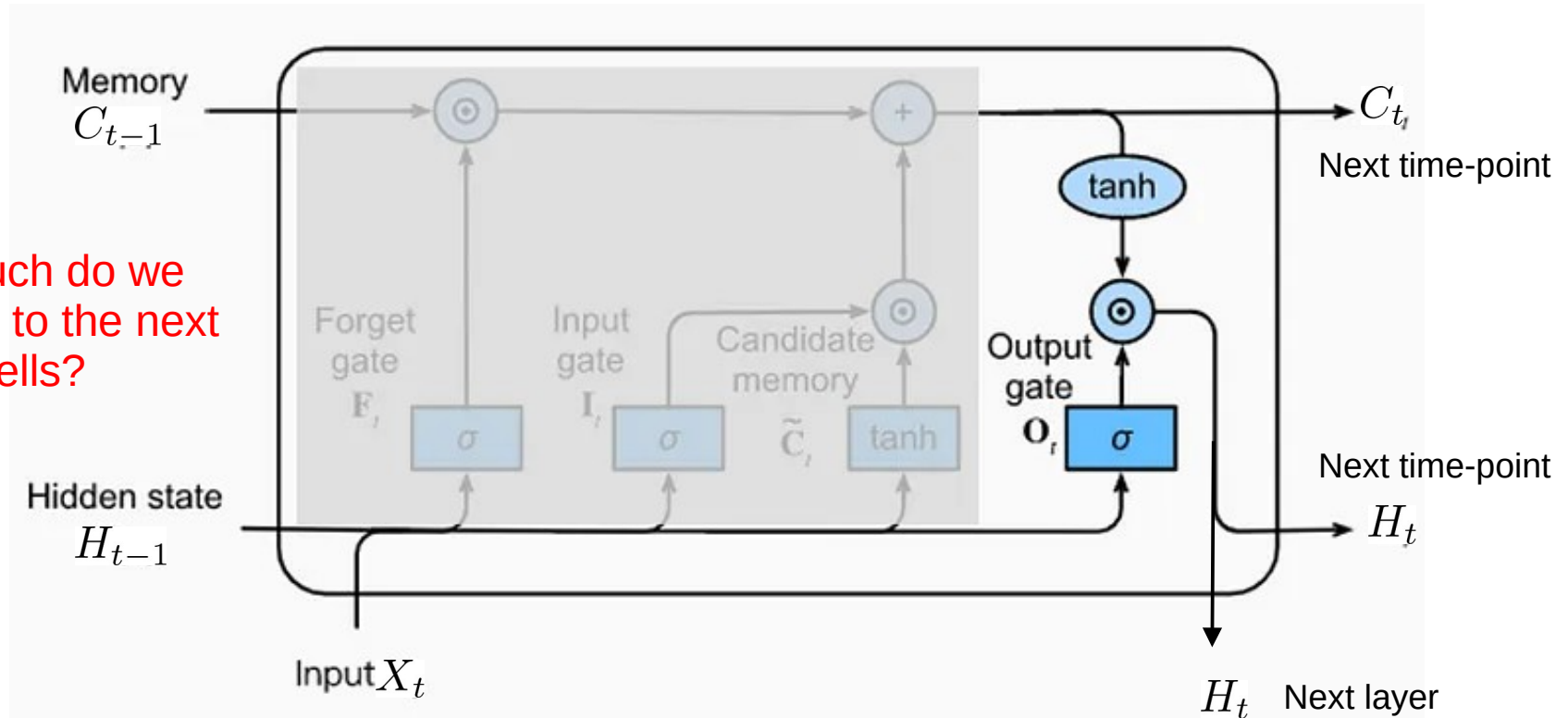
Solution: Long Short-Term Memory (LSTM)

2-How much do we remember from this time-point?



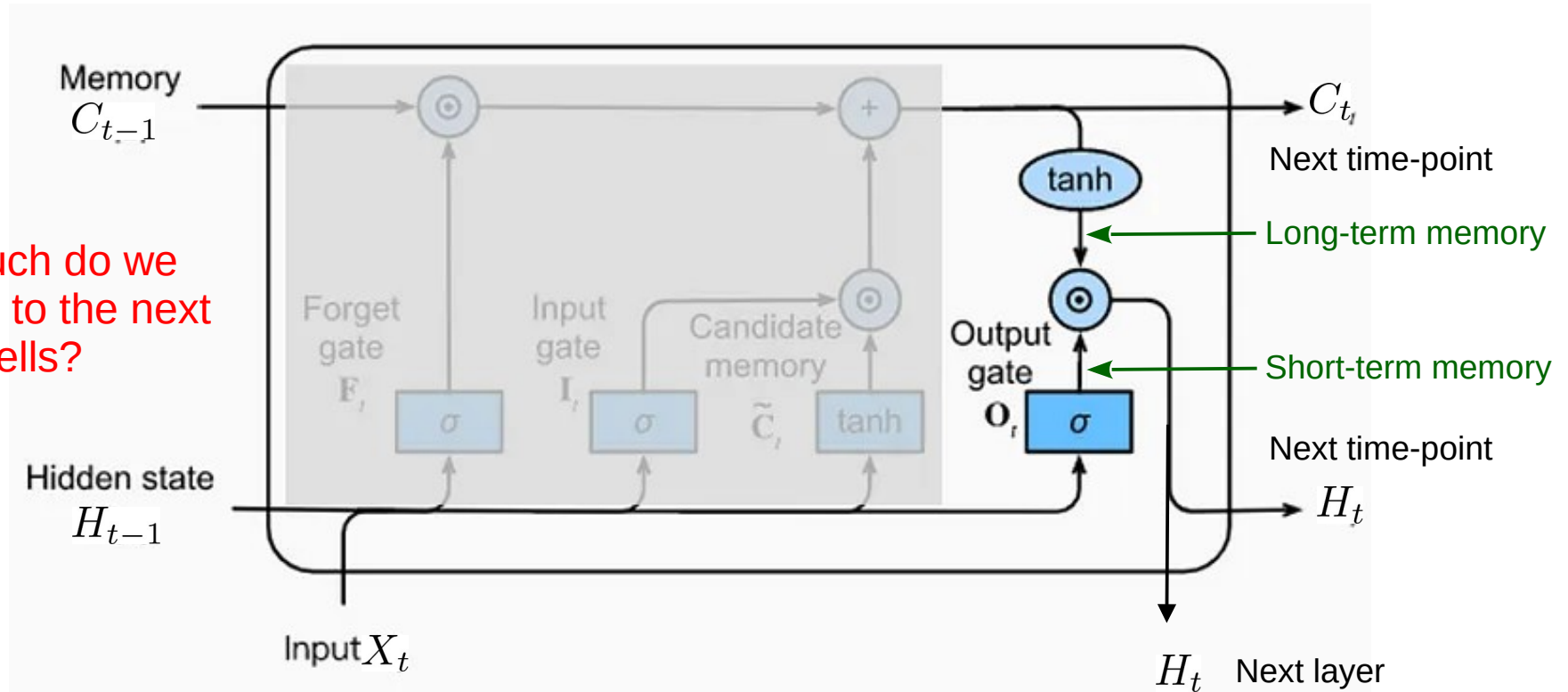
Solution: Long Short-Term Memory (LSTM)

3-How much do we transmit to the next LSTM cells?

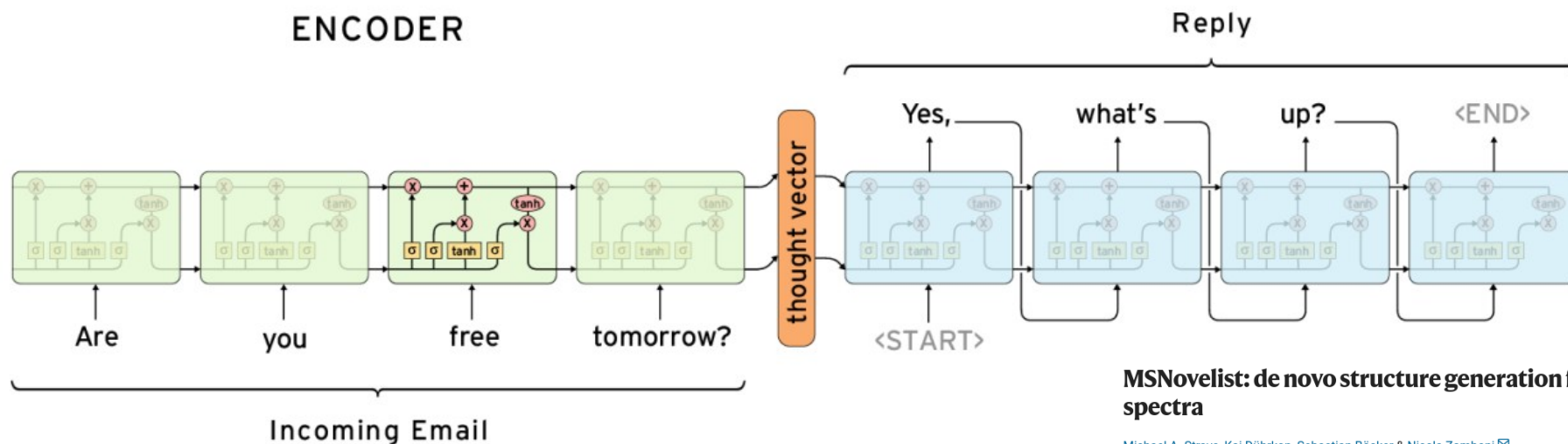


Solution: Long Short-Term Memory (LSTM)

3-How much do we transmit to the next LSTM cells?



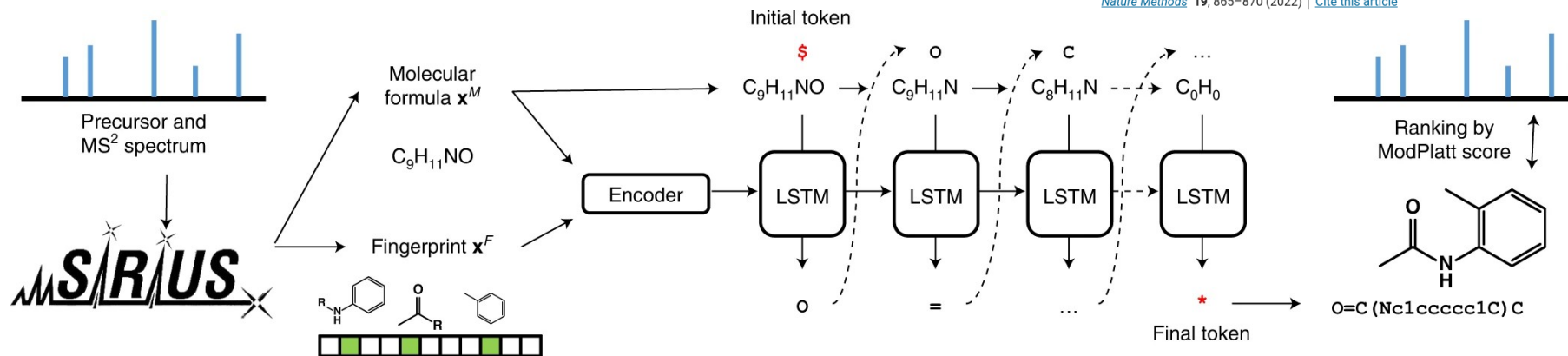
LSTMs for Encoder-Decoder



MSNovelist: de novo structure generation from mass spectra

Michael A. Stravs, Kai Dührkop, Sebastian Böcker & Nicola Zamboni

Nature Methods 19, 865–870 (2022) | [Cite this article](#)



Example in bioinformatics

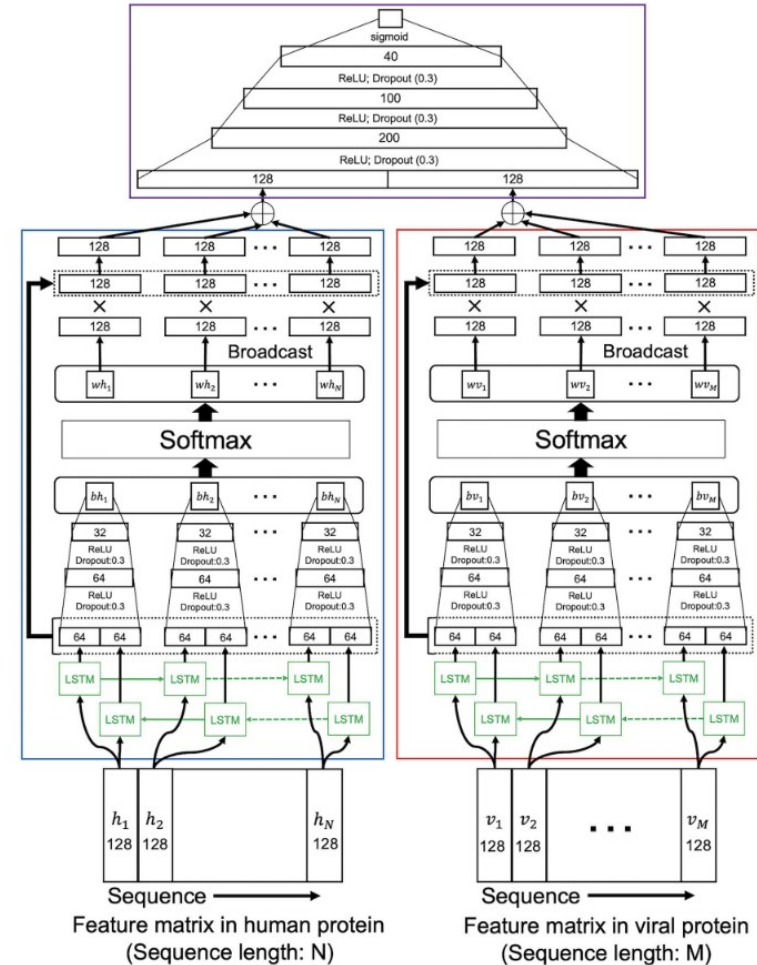
OXFORD

Briefings in Bioinformatics, 22(6), 2021, 1–9

<https://doi.org/10.1093/bib/bbab228>
Problem Solving Protocol

LSTM-PHV: prediction of human-virus protein–protein interactions by LSTM with word2vec

Sho Tsukiyama, Md Mehedi Hasan, Satoshi Fujii and Hiroyuki Kurata



Example in clinical setting



Contents lists available at [ScienceDirect](https://www.sciencedirect.com)

International Journal of Infectious Diseases

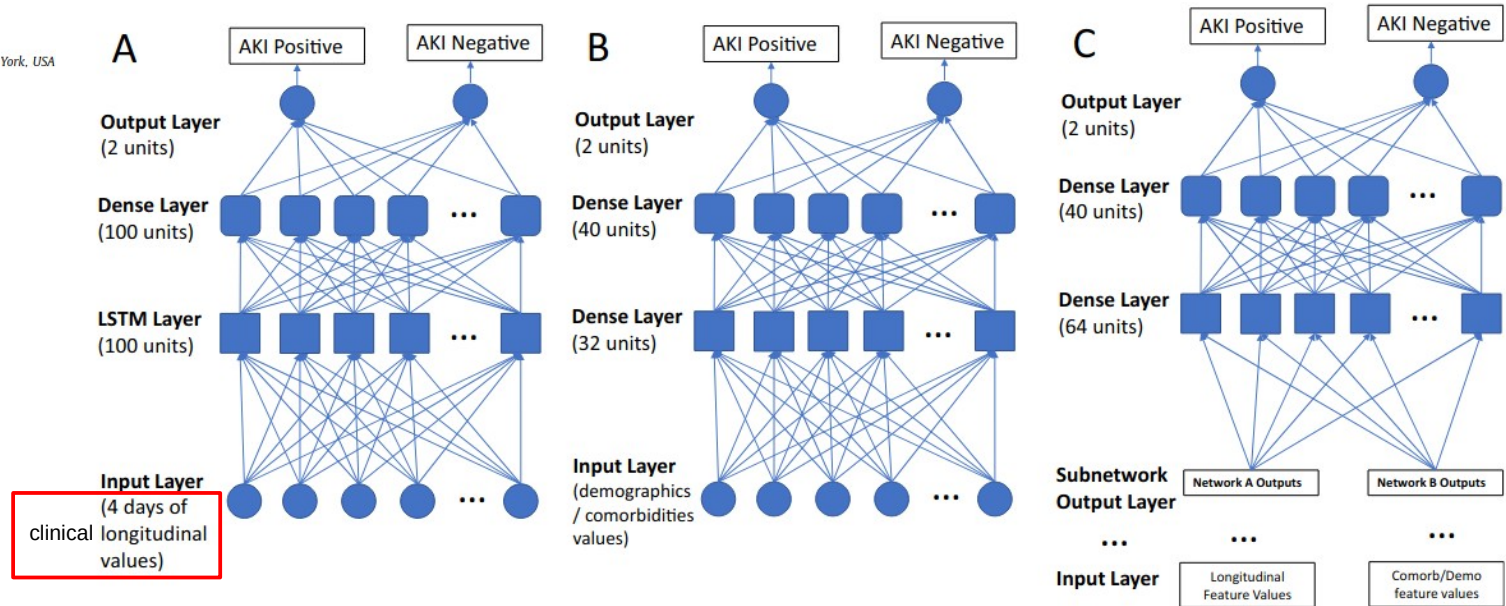
journal homepage: www.elsevier.com/locate/ijid



Long-short-term memory machine learning of longitudinal clinical data accurately predicts acute kidney injury onset in COVID-19: a two-center study

Justin Y. Lu, Joanna Zhu, Jocelyn Zhu, Tim Q Duong*

Department of Radiology, Montefiore Medical Center, Albert Einstein College of Medicine, New York, USA



The paper that changed everything: the Transformer

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

Lukasz Kaiser*
Google Brain
lukaszkaiser@google.com

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

Abstract

The dominant sequence transduction models are based on complex recurrent or convolutional neural networks that include an encoder and a decoder. The best performing models also connect the encoder and decoder through an attention mechanism. We propose a new simple network architecture, the Transformer, based solely on attention mechanisms, dispensing with recurrence and convolutions entirely. Experiments on two machine translation tasks show these models to be superior in quality while being more parallelizable and requiring significantly less time to train. Our model achieves 28.4 BLEU on the WMT 2014 English-to-German translation task, improving over the existing best results, including ensembles, by over 2 BLEU. On the WMT 2014 English-to-French translation task, our model establishes a new single-model state-of-the-art BLEU score of 41.0 after training for 3.5 days on eight GPUs, a small fraction of the training costs of the best models from the literature.

1 Introduction

Recurrent neural networks, long short-term memory [12] and gated recurrent [7] neural networks in particular, have been firmly established as state of the art approaches in sequence modeling and transduction problems such as language modeling and machine translation [29, 2, 5]. Numerous efforts have since continued to push the boundaries of recurrent language models and encoder-decoder architectures [31, 21, 13].

*Equal contribution. Listing order is random. Jakob proposed replacing RNNs with self-attention and started the effort to evaluate this idea. Ashish, with Illia, designed and implemented the first Transformer models and has been crucially involved in every aspect of this work. Noam proposed scaled dot-product attention, multi-head attention and the parameter-free position representation and became the other person involved in nearly every detail. Niki designed, implemented, tuned and evaluated countless model variants in our original codebase and tensor2tensor. Llion also experimented with novel model variants, was responsible for our initial codebase, and efficient inference and visualizations. Lukasz and Aidan spent countless long days designing various parts of and implementing tensor2tensor, replacing our earlier codebase, greatly improving results and massively accelerating our research.

[†]Work performed while at Google Brain.

[‡]Work performed while at Google Research.

The paper that changed everything: the Transformer

Attention Is All You Need

Cool title

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

Lukasz Kaiser*
Google Brain
lukaszkaier@google.com

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

Abstract

The dominant sequence transduction models are based on complex recurrent or convolutional neural networks that include an encoder and a decoder. The best performing models also connect the encoder and decoder through an attention mechanism. We propose a new simple network architecture, the Transformer, based solely on attention mechanisms, dispensing with recurrence and convolutions entirely. Experiments on two machine translation tasks show these models to be superior in quality while being more parallelizable and requiring significantly less time to train. Our model achieves 28.4 BLEU on the WMT 2014 English-to-German translation task, improving over the existing best results, including ensembles, by over 2 BLEU. On the WMT 2014 English-to-French translation task, our model establishes a new single-model state-of-the-art BLEU score of 41.0 after training for 3.5 days on eight GPUs, a small fraction of the training costs of the best models from the literature.

1 Introduction

Recurrent neural networks, long short-term memory [12] and gated recurrent [7] neural networks in particular, have been firmly established as state-of-the-art approaches in sequence modeling and

*Equal contribution. Listing order is random.

[†]Equal contribution. Listing order is random. Jakob proposed replacing RNNs with self-attention and started the effort to evaluate this idea. Ashish, with Illia, designed and implemented the first Transformer models and has been crucially involved in every aspect of this work. Noam proposed scaled dot-product attention, multi-head attention and the parameter-free position representation and became the other person involved in nearly every detail. Niki designed, implemented, tuned and evaluated countless model variants in our original codebase and tensor2tensor. Llion also experimented with novel model variants, was responsible for our initial codebase, and efficient inference and visualizations. Lukasz and Aidan spent countless long days designing various parts of and implementing tensor2tensor, replacing our earlier codebase, greatly improving results and massively accelerating our research.

[‡]Work performed while at Google Brain.

[†]Work performed while at Google Research.

31st Conference on Neural Information Processing Systems (NIPS 2017), Long Beach, CA, USA.

All authors equal

Cited... 138640 times as of 27 October 2024!

Never published in a journal

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

Lukasz Kaiser*
Google Brain
lukaszkaier@google.com

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

Abstract

The dominant sequence transduction models are based on complex recurrent or convolutional neural networks that include an encoder and a decoder. The best performing models also connect the encoder and decoder through an attention mechanism. We propose a new simple network architecture, the Transformer, based solely on attention mechanisms, dispensing with recurrence and convolutions entirely. Experiments on two machine translation tasks show these models to be superior in quality while being more parallelizable and requiring significantly less time to train. Our model achieves 28.4 BLEU on the WMT 2014 English-to-German translation task, improving over the existing best results, including ensembles, by over 2 BLEU. On the WMT 2014 English-to-French translation task, our model establishes a new single-model state-of-the-art BLEU score of 41.0 after training for 3.5 days on eight GPUs, a small fraction of the training costs of the best models from the literature.

1 Introduction

Recurrent neural networks, long short-term memory [12] and gated recurrent [7] neural networks in particular, have been firmly established as state of the art approaches in sequence modeling and transduction problems such as language modeling and machine translation [29, 2, 5]. Numerous efforts have since continued to push the boundaries of recurrent language models and encoder-decoder architectures [31, 21, 13].

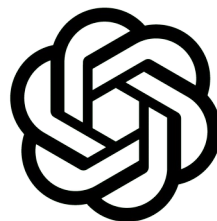
*Equal contribution. Listing order is random. Jakob proposed replacing RNNs with self-attention and started the effort to evaluate this idea. Ashish, with Illia, designed and implemented the first Transformer models and has been crucially involved in every aspect of this work. Noam proposed scaled dot-product attention, multi-head attention and the parameter-free position representation and became the other person involved in nearly every detail. Niki designed, implemented, tuned and evaluated countless model variants in our original codebase and tensor2tensor. Llion also experimented with novel model variants, was responsible for our initial codebase, and efficient inference and visualizations. Lukasz and Aidan spent countless long days designing various parts of and implementing tensor2tensor, replacing our earlier codebase, greatly improving results and massively accelerating our research.

[†]Work performed while at Google Brain.

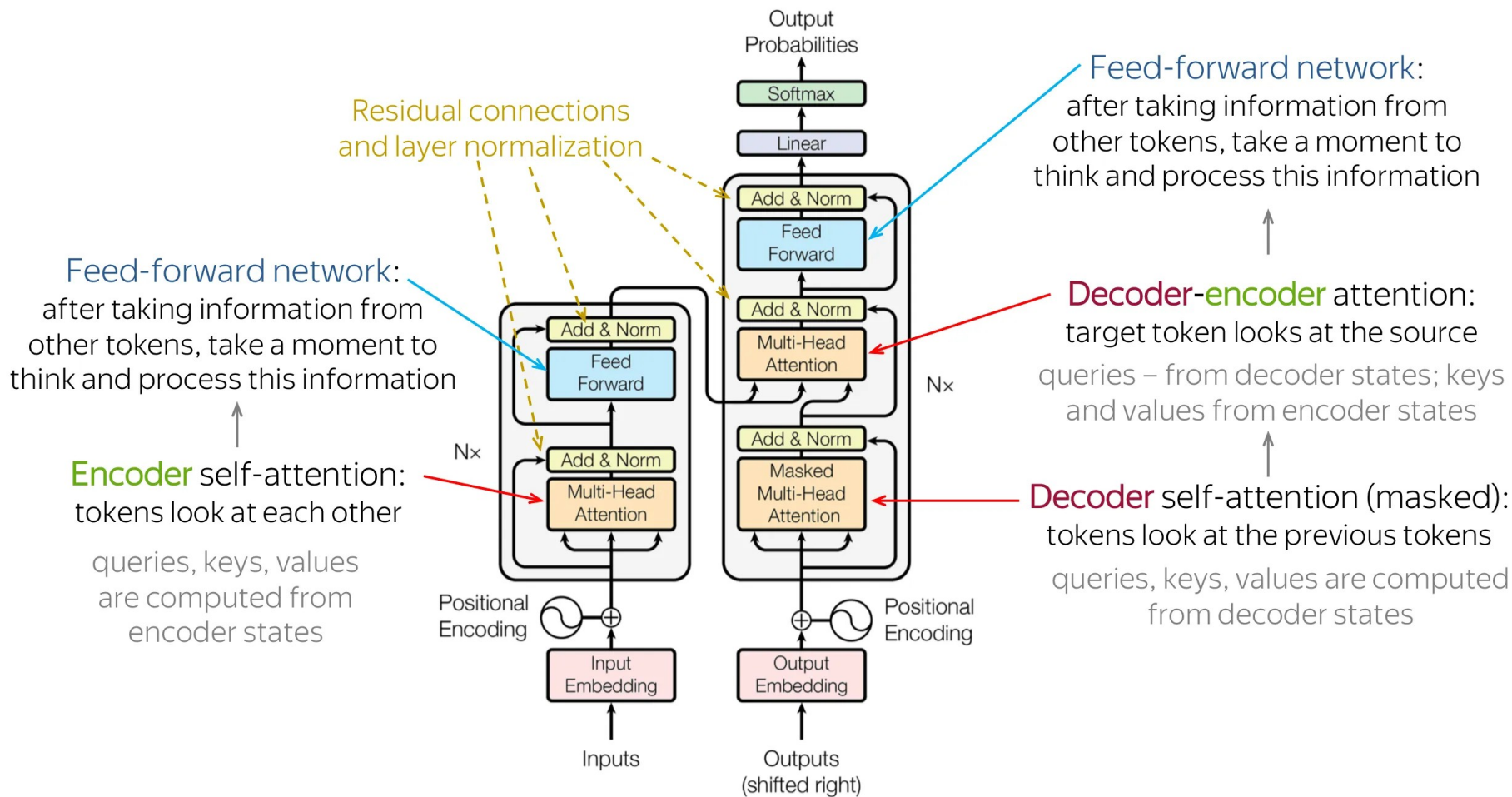
[‡]Work performed while at Google Research.

31st Conference on Neural Information Processing Systems (NIPS 2017), Long Beach, CA, USA.

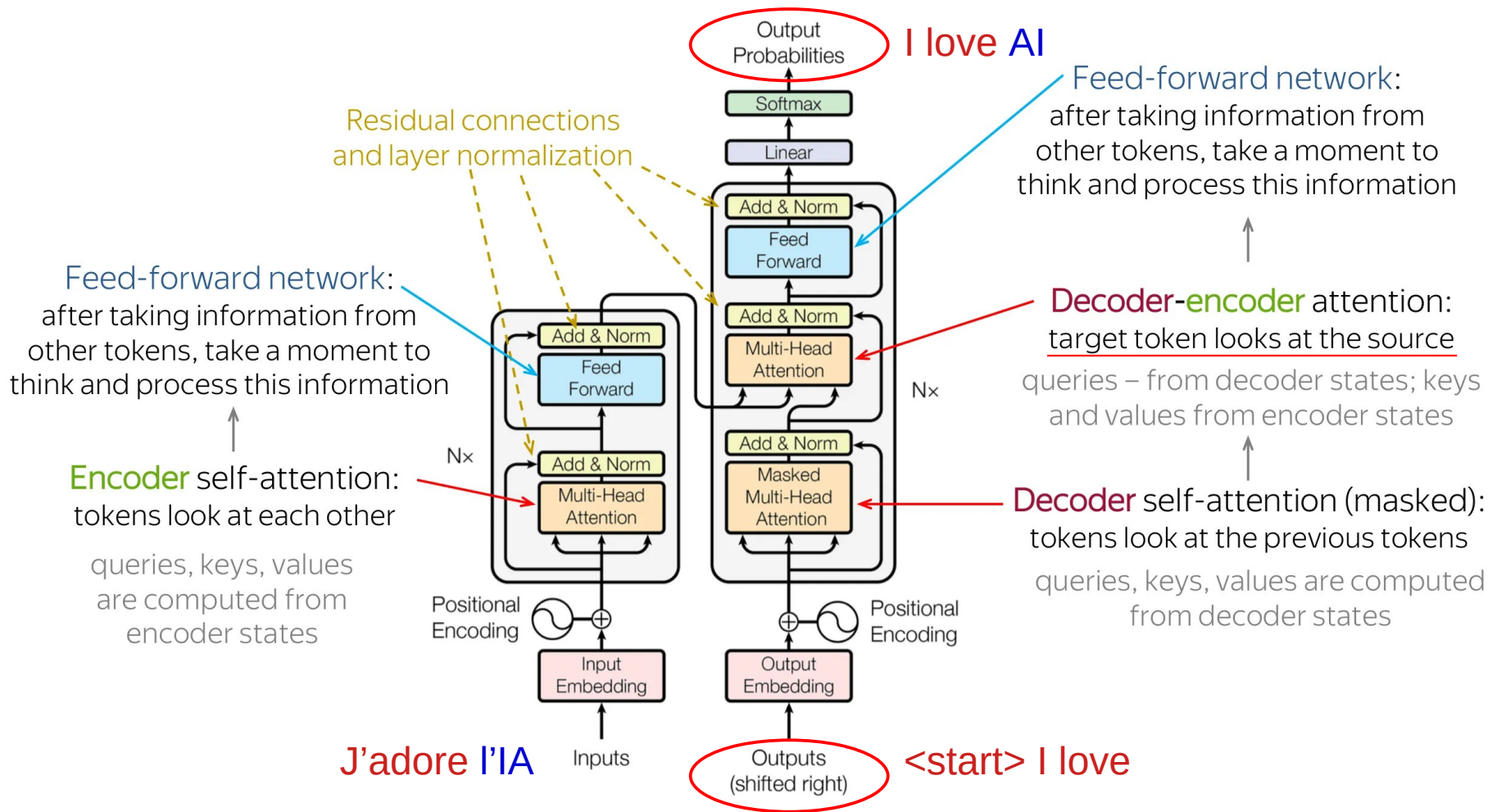
The paper that changed everything: the Transformer



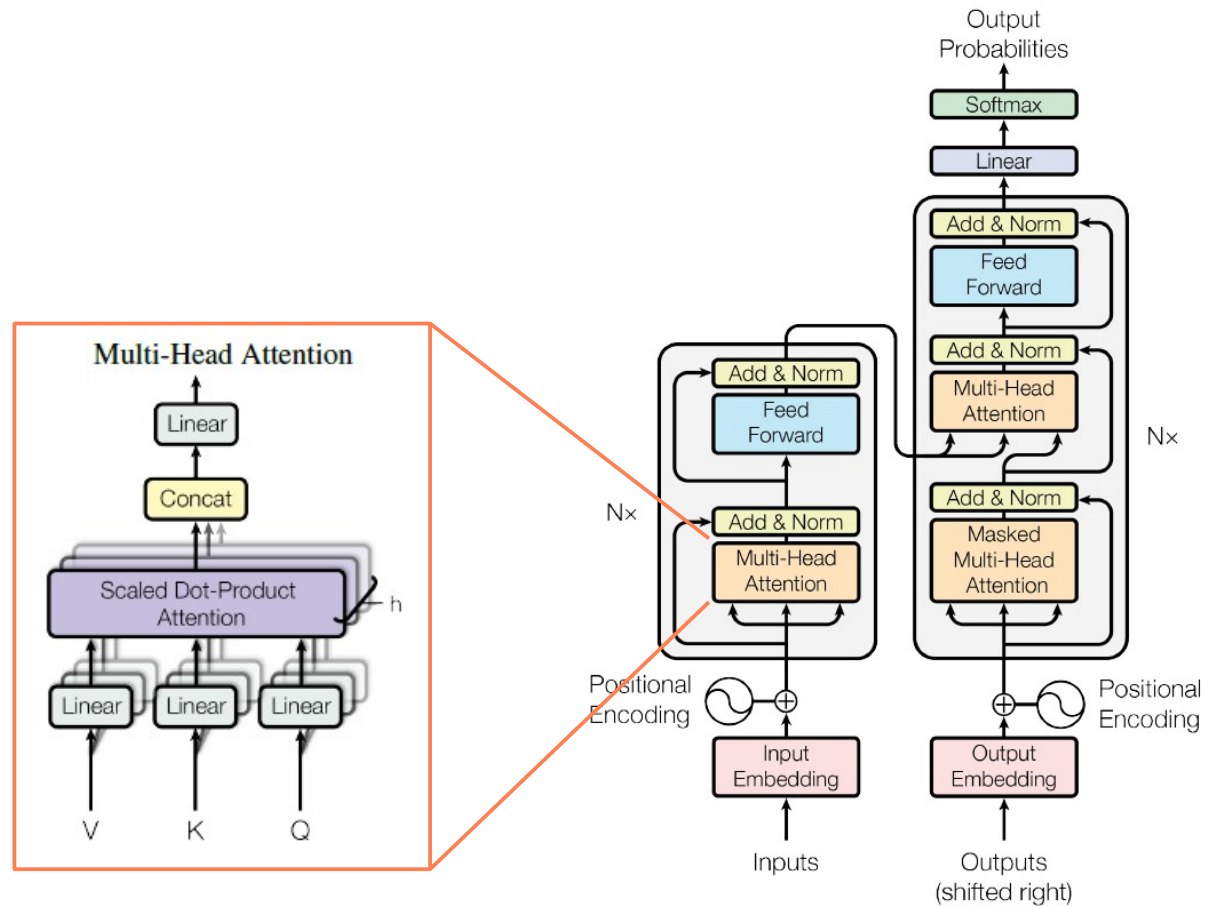
The Transformer: Memory + context = attention



The Transformer: Memory + context = attention

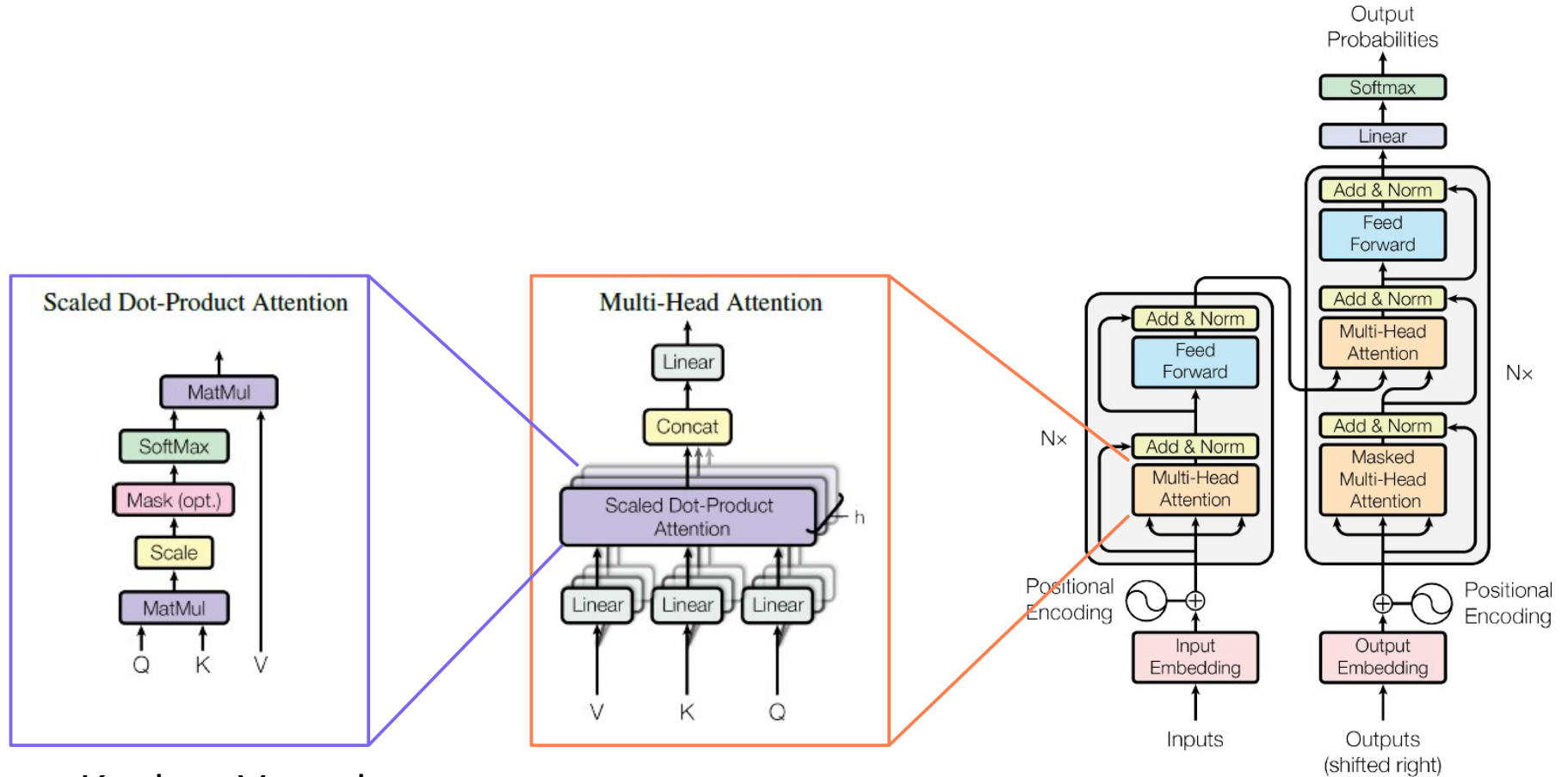


Attention in the Transformer



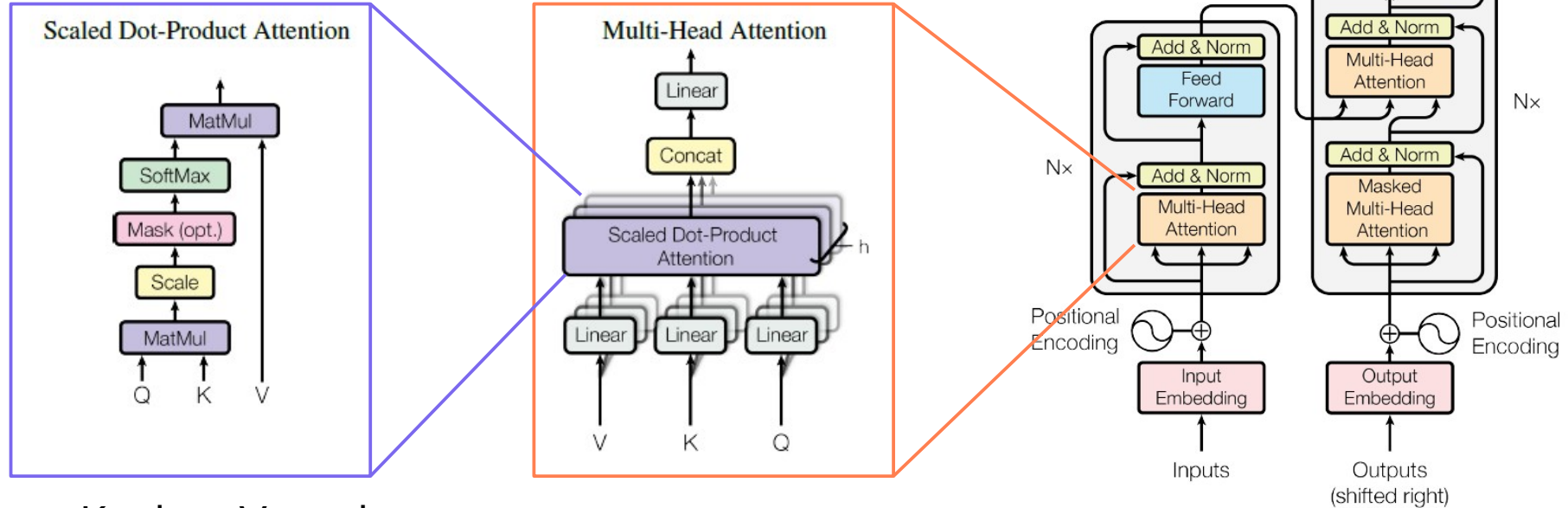
Q = query, K = key, V = value

Attention in the Transformer



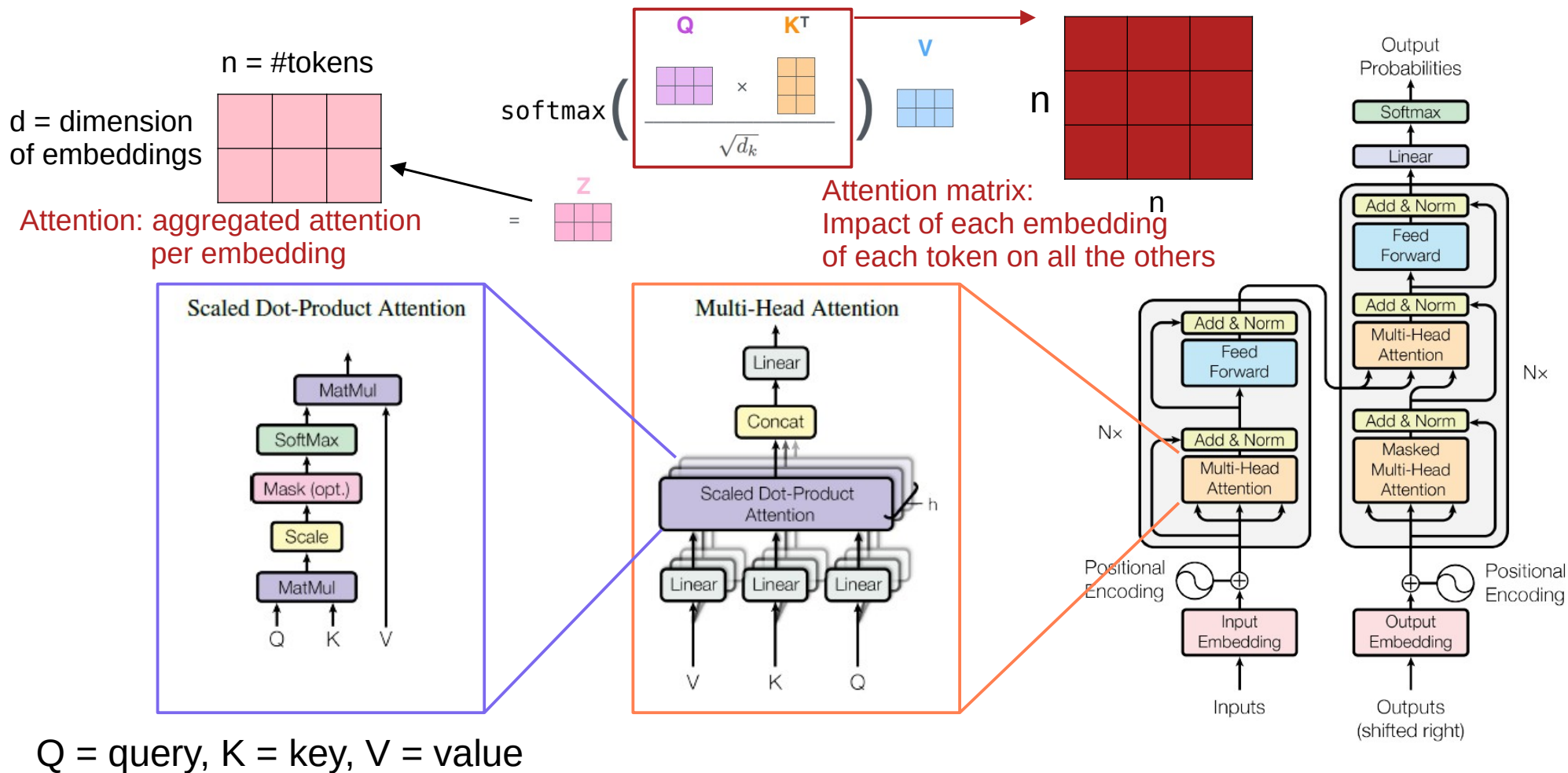
Attention in the Transformer

$$\text{softmax}\left(\frac{\begin{matrix} \text{Q} \\ \text{K}^T \end{matrix}}{\sqrt{d_k}}\right) \begin{matrix} \text{V} \end{matrix} = \begin{matrix} \text{Z} \end{matrix}$$



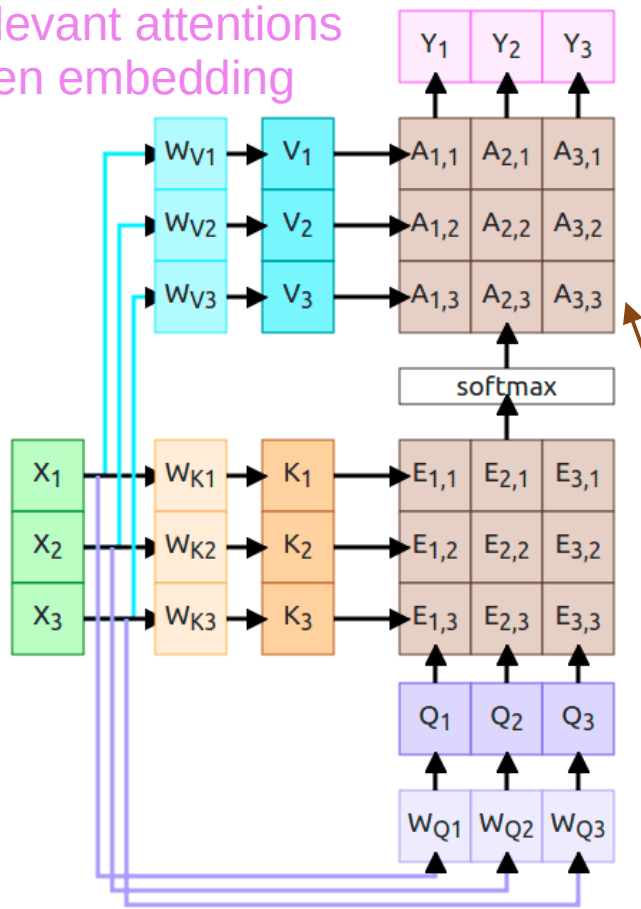
Q = query, K = key, V = value

Attention in the Transformer



Attention in the Transformer

Actual relevant attentions
1 per token embedding



Self-Attention Layer

Input vectors: $\mathbf{X}$ (shape $N_X \times D_X$)

Key matrix: $\mathbf{W}_K$ (Shape $D_X \times D_Q$)

Value matrix: $\mathbf{W}_V$ (Shape $D_X \times D_V$)

Query matrix: $\mathbf{W}_Q$ (Shape $D_X \times D_Q$)

Query vectors: $\mathbf{Q} = \mathbf{X}\mathbf{W}_Q$ (Shape $N_X \times D_Q$)

Key vectors: $\mathbf{K} = \mathbf{X}\mathbf{W}_K$ (Shape $N_X \times D_Q$)

Value vectors: $\mathbf{V} = \mathbf{X}\mathbf{W}_V$ (Shape $N_X \times D_V$)

Similarity function: *scaled dot product*

Similarities: $\mathbf{E} = \mathbf{Q}\mathbf{K}^T$ (shape $N_X \times N_X$), $E_{i,j} = \mathbf{Q}_i \cdot \mathbf{K}_j / \sqrt{D_Q}$

Attention weights: $\mathbf{A} = \text{softmax}(\mathbf{E}, \text{dim} = 1)$ (shape $N_X \times N_X$)

Output: $\mathbf{Y} = \mathbf{A}\mathbf{V}$ (shape $N_X \times D_V$) where $\mathbf{Y}_i = \sum_j (\mathbf{A}_{i,j}, \mathbf{V}_j)$

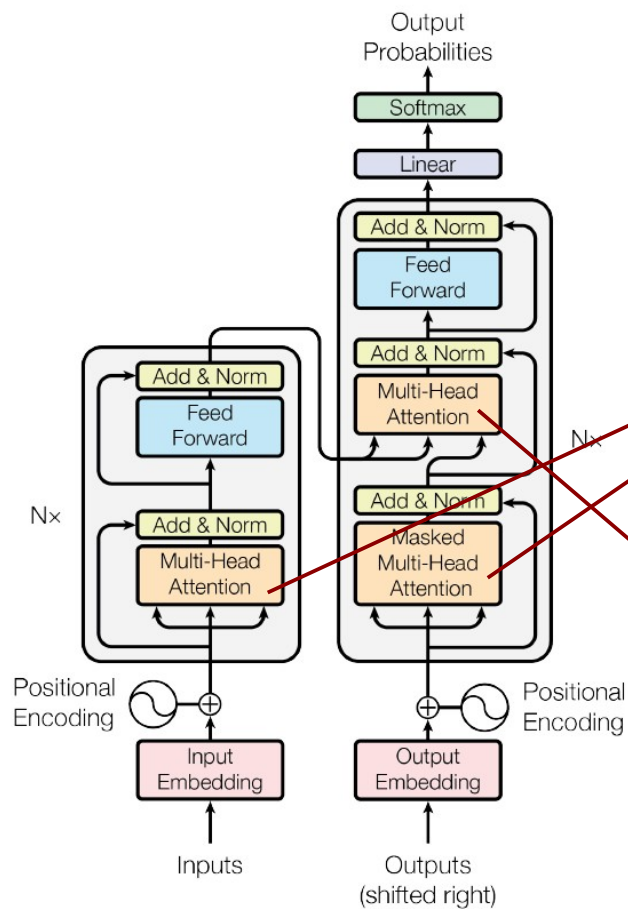
The attention of all token embeddings
on all token embeddings

$\mathbf{X}$ is entered

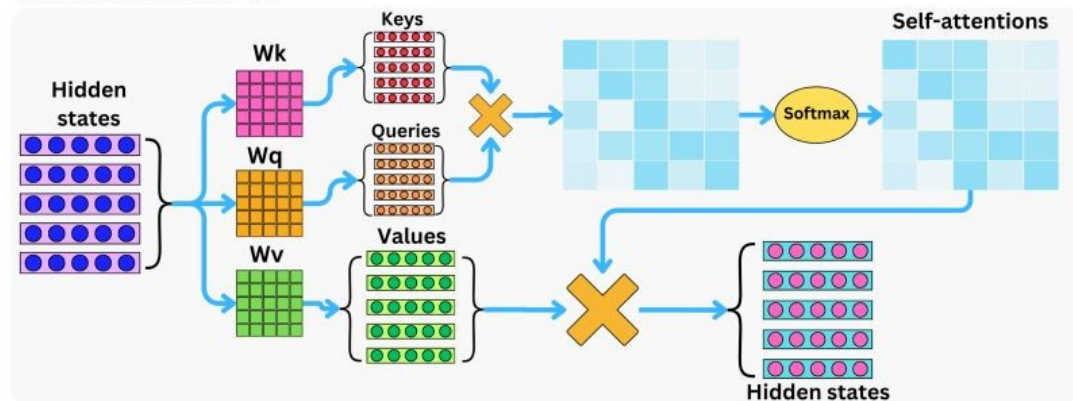
$\mathbf{W}_Q$, $\mathbf{W}_K$, and $\mathbf{W}_V$ are learned

Everything else is computed

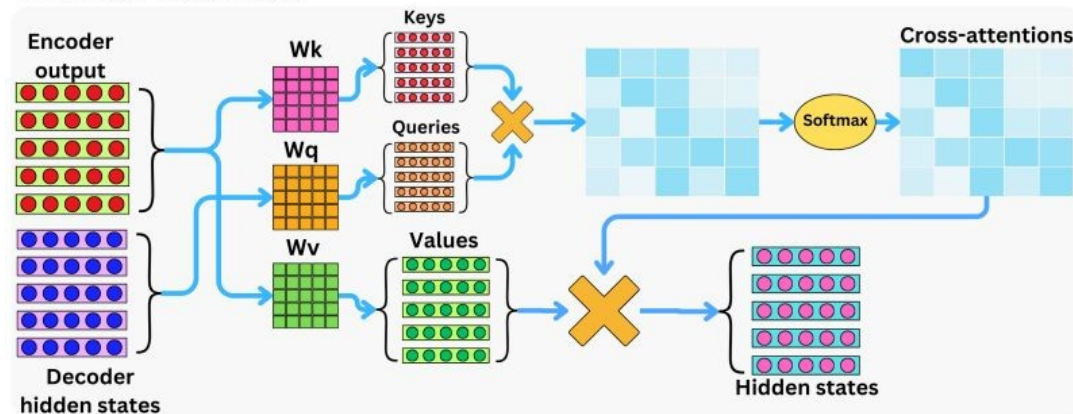
Self versus cross-attention



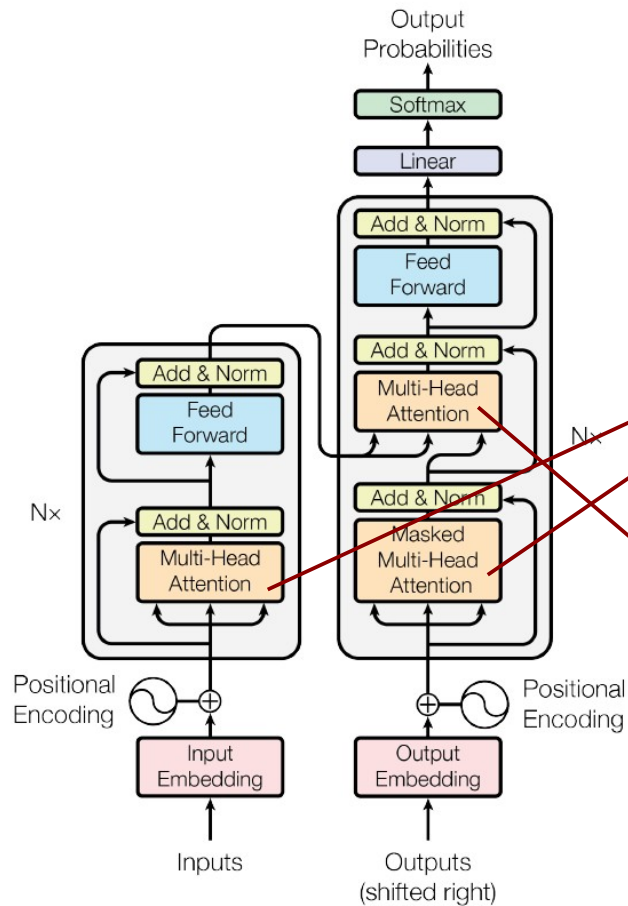
The Self-Attentions



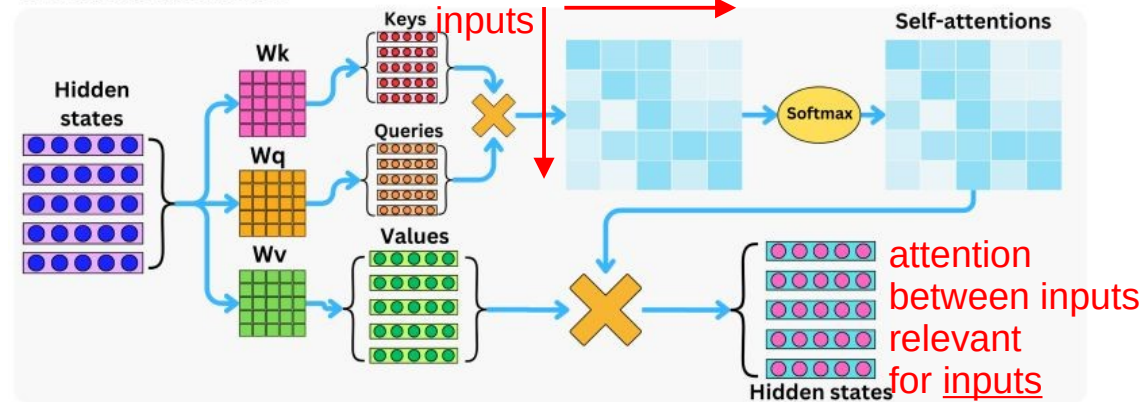
The Cross-Attentions



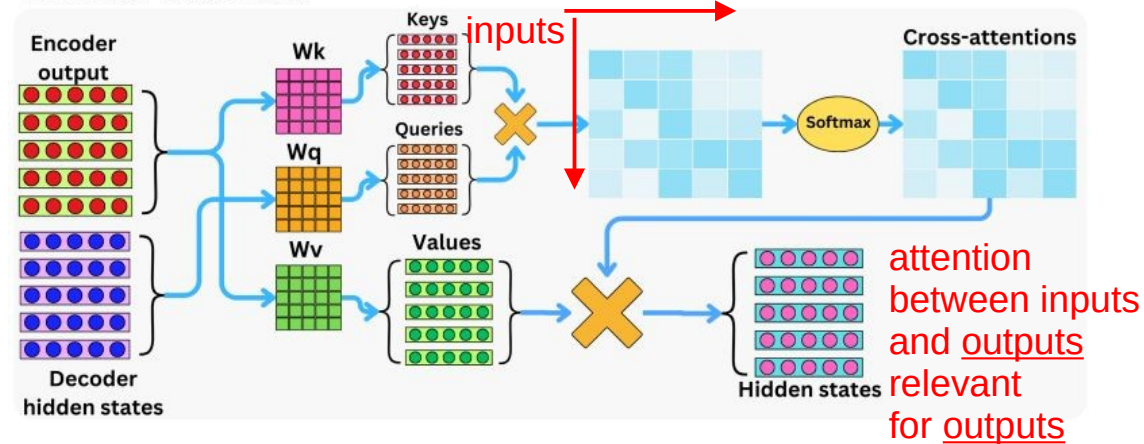
Self versus cross-attention



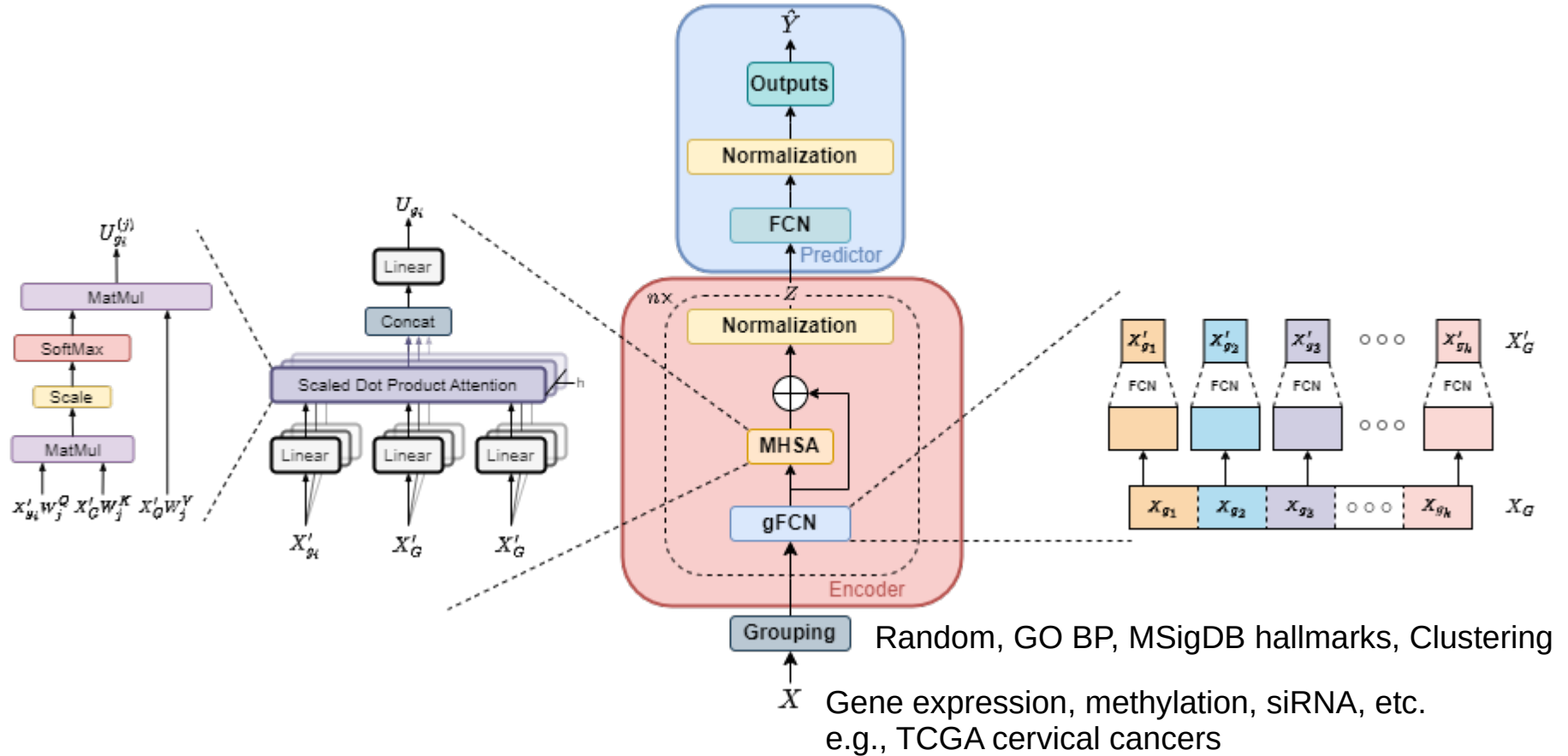
The Self-Attentions



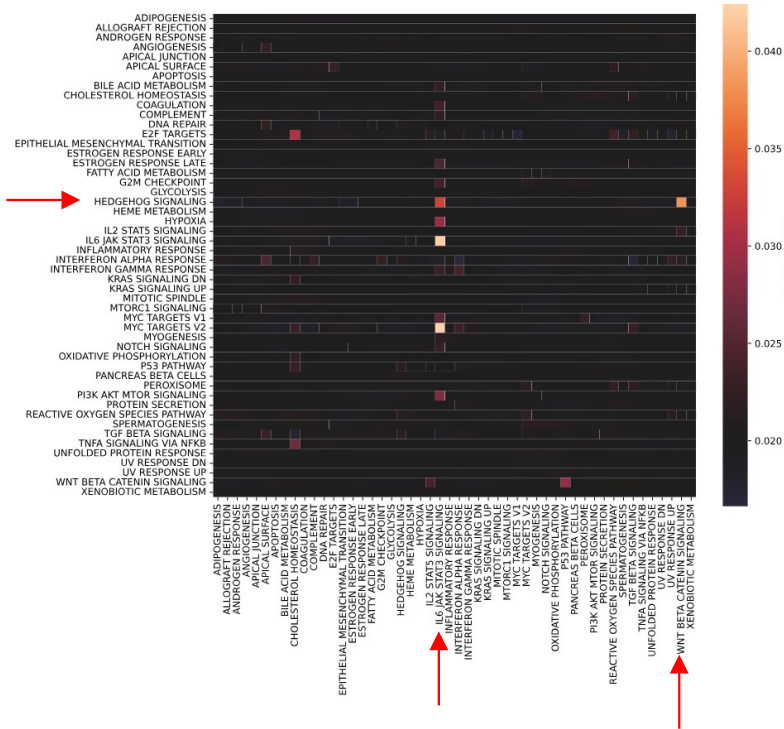
The Cross-Attentions



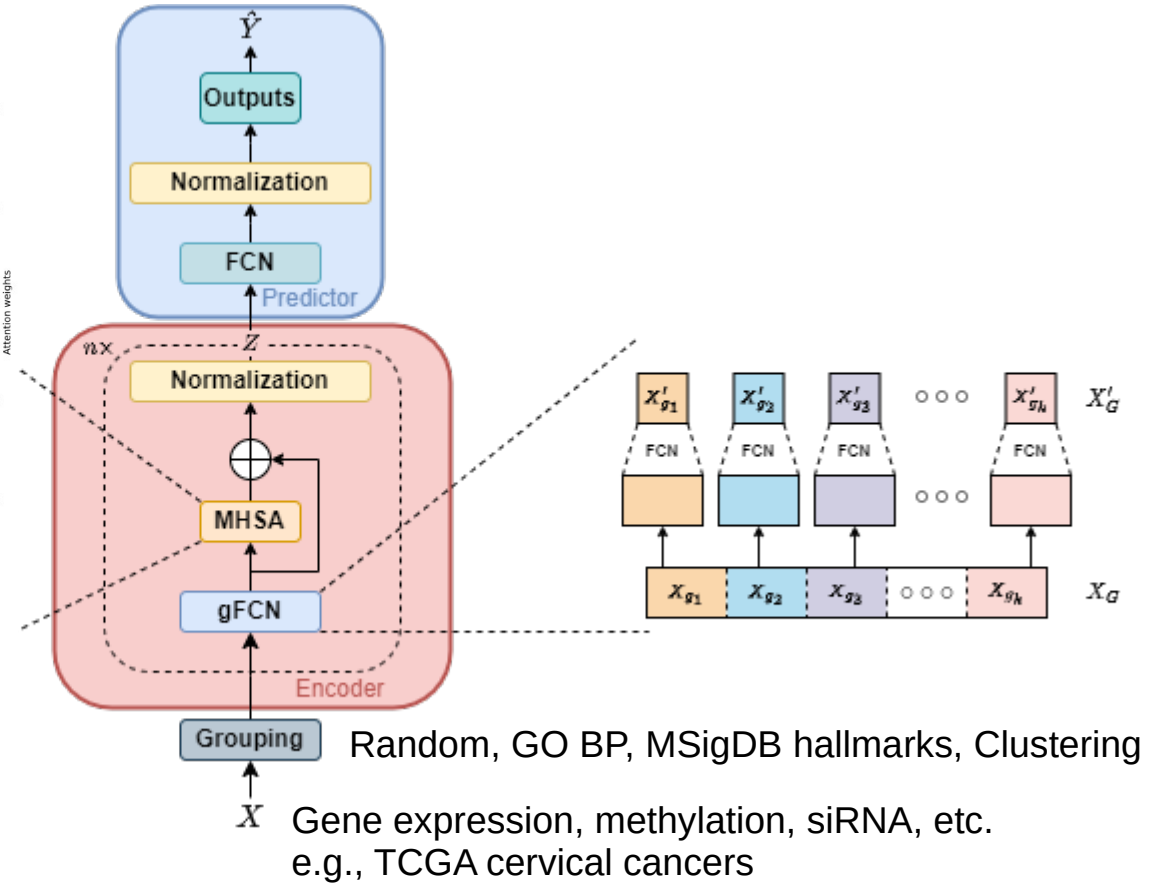
AttOmics



AttOmics



Pathways affected in cervical cancer

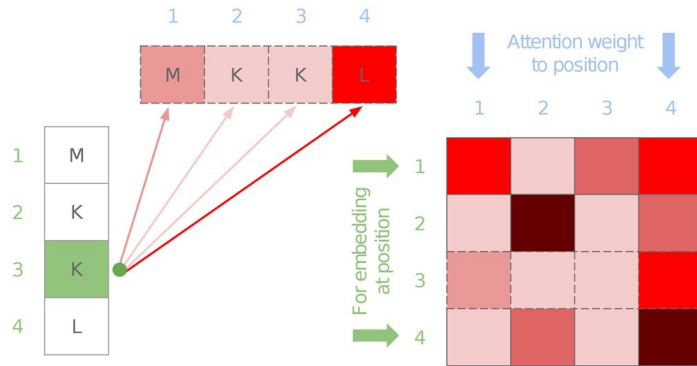


EnzBERT

Predicting enzymatic function of protein sequences with attention 🧠

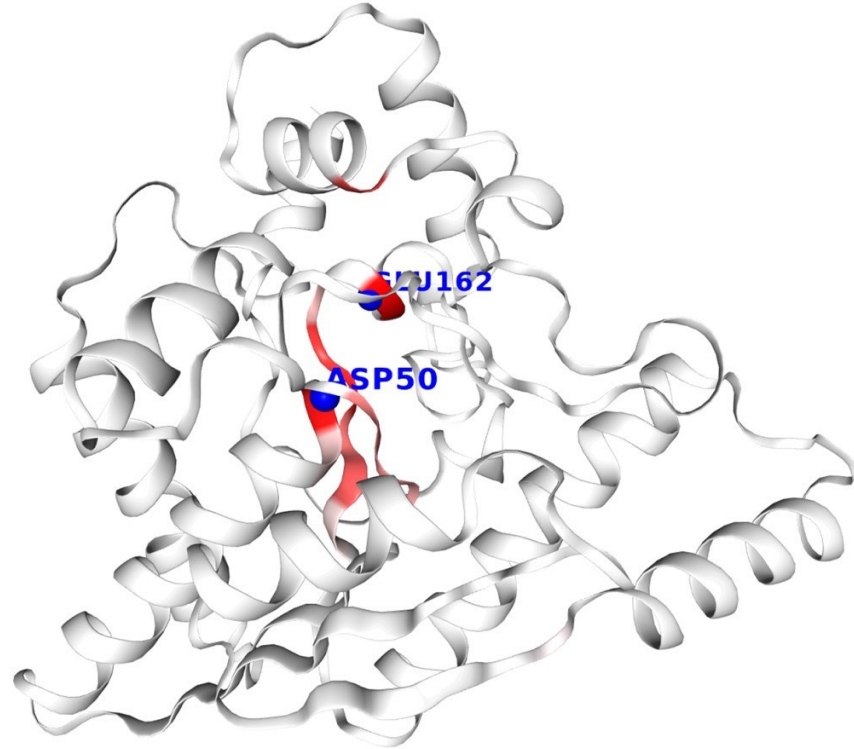
Nicolas Buton ✉, François Coste, Yann Le Cunff

Bioinformatics, Volume 39, Issue 10, October 2023, btad620, <https://doi.org/10.1093/bioinformatics/btad620>



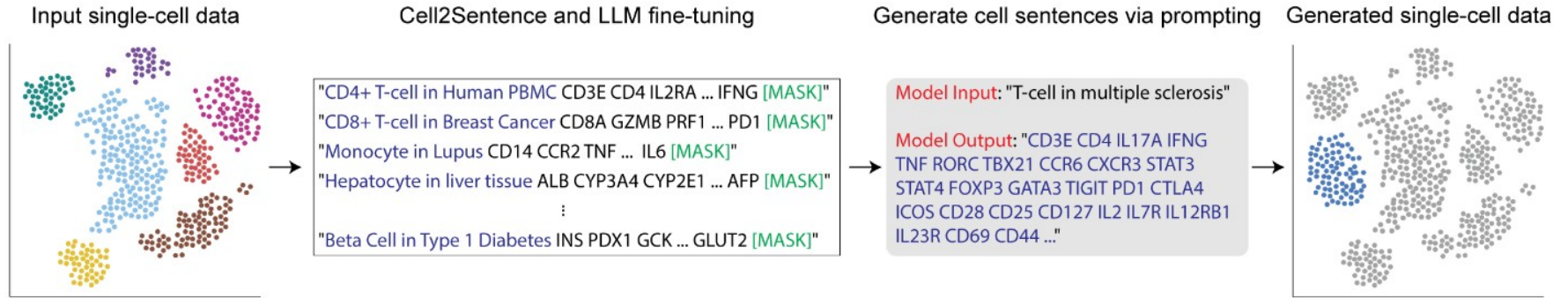
Aggregated attention
for each token (amino acid)

Nh(3)-dependent nad(+) synthetase

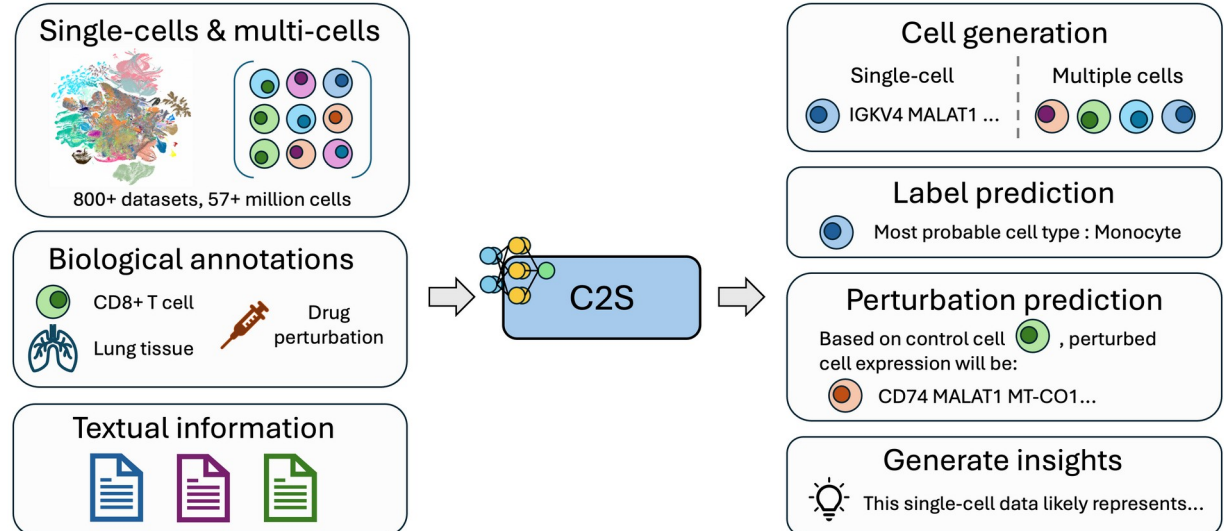


0 MSMQEKIMRE LHVKPSIDPK QEIEDRVNLF KQYVKKTGAK GFVLGISGQ DSTLAGRLAQ LAVESIREEG GDAQFIIVRL PHGTQQDEDD AQLALKFIKP
1 DKSWKFDIKS TVSAFSDQYQ QETGDQLTDF NKGNVKARTR MIAQYAIGGQ EGLLVLSIDH AAEAVTGFFT KYGDGGADLL PLTGLTKRQG RTLLKELGAP
2 ERLYLKEPTA DLLDEKPQQS DETELGISD EIDDYLEGKE VSAKVSEALE KRYSMTEHKR QVPASMFDDW WK

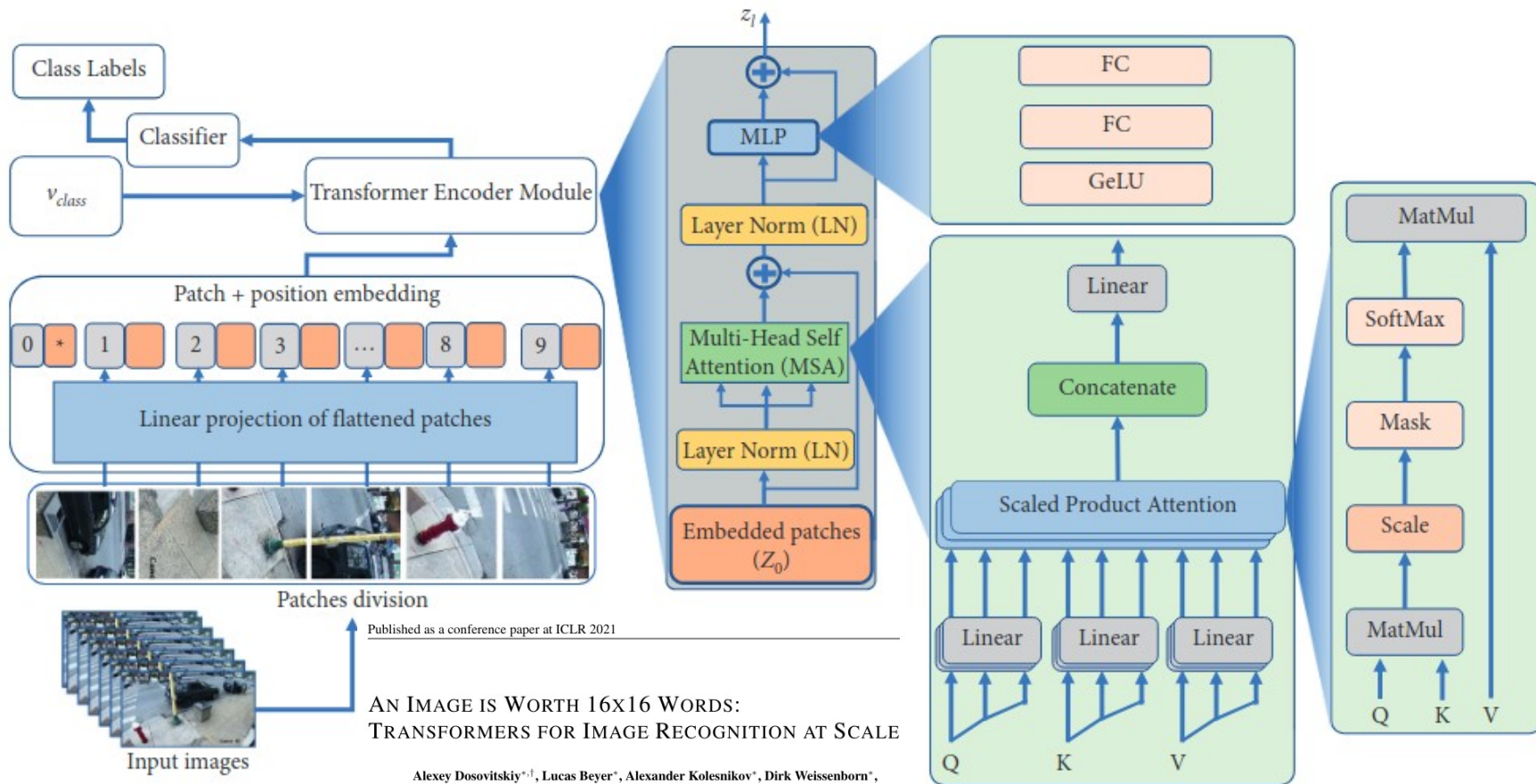
Cell2Sentence



Levine *et al* (2024). Cell2Sentence:
Teaching Large Language Models
the Language of Biology. *BioRxiv*
<https://doi.org/10.1101/2023.09.11.557287>



Vision Transformer

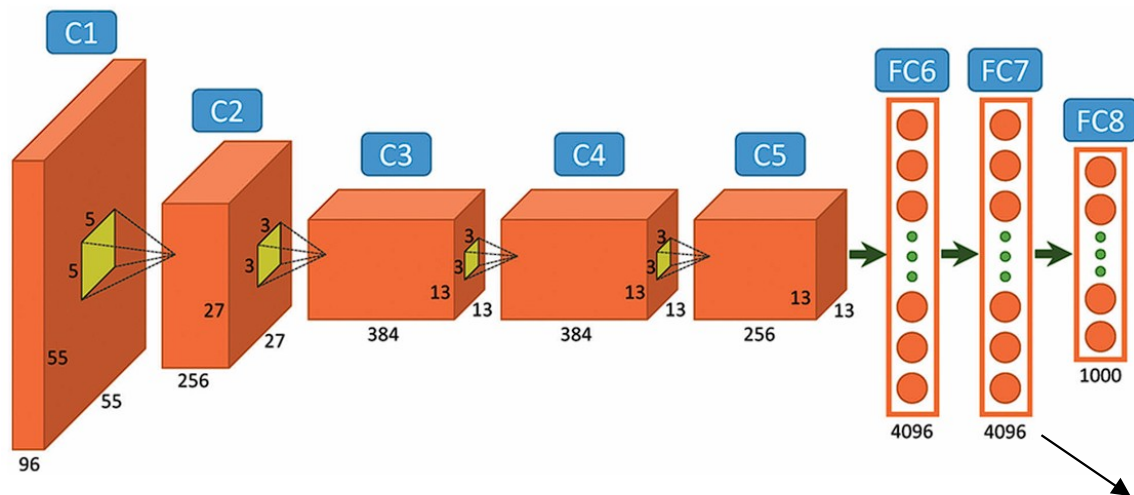


Alexey Dosovitskiy^{*,†}, Lucas Beyer^{*}, Alexander Kolesnikov^{*}, Dirk Weissenborn^{*}, Xiaohua Zhai^{*}, Thomas Unterthiner, Mostafa Dehghani, Matthias Minderer, Georg Heigold, Sylvain Gelly, Jakob Uszkoreit, Neil Houlsby^{*,†}
^{*}equal technical contribution, [†]equal advising
 Google Research, Brain Team
 {adosovitskiy, neilhoulby}@google.com

CNN = local features
 ViT = relations between distant features

source: <https://doi.org/10.1155/2022/3454167>

Patches are embedded by CNNs

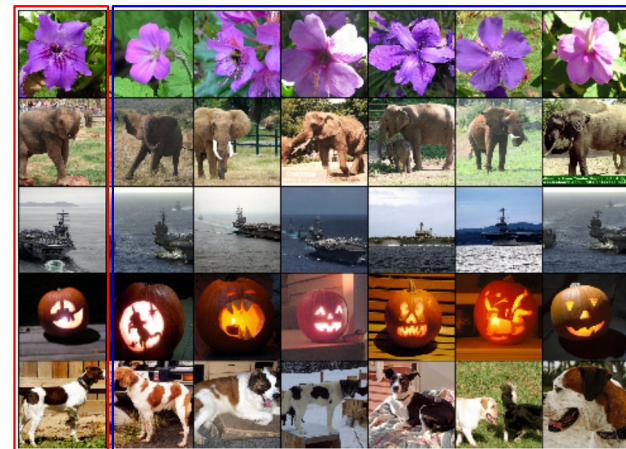


6 nearest neighbours in the 4096 dimension space

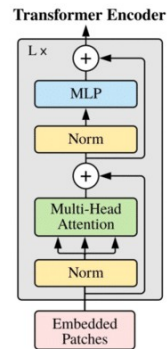
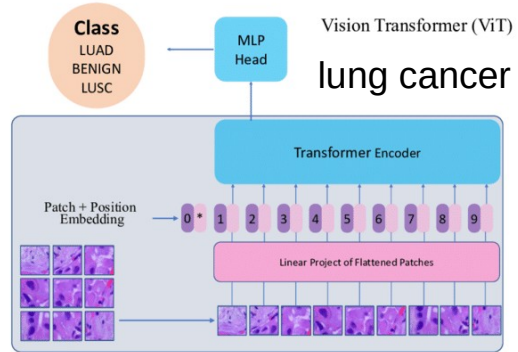
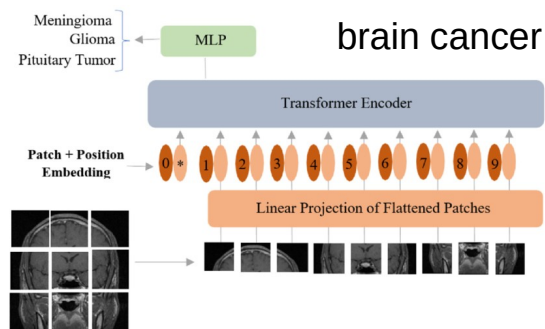
Krizhevsky A, Sutskever I, Hinton GE (2012)
ImageNet Classification with Deep
Convolutional Neural Networks
[https://proceedings.neurips.cc/paper/2012/file/
c399862d3b9d6b76c8436e924a68c45b-Paper.pdf](https://proceedings.neurips.cc/paper/2012/file/c399862d3b9d6b76c8436e924a68c45b-Paper.pdf)

(presenting AlexNet, the first Deep Convolutional Network)

input
image

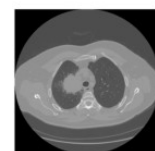


ViTs are replacing vanilla CNNs

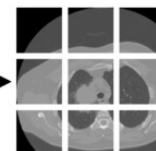


PROPOSED ViT ARCHITECTURE:

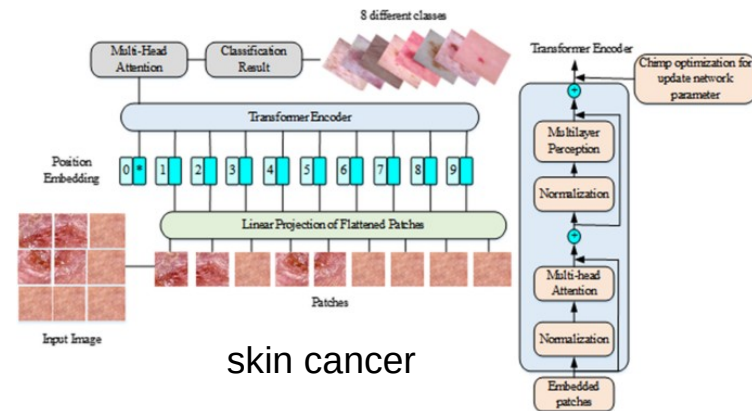
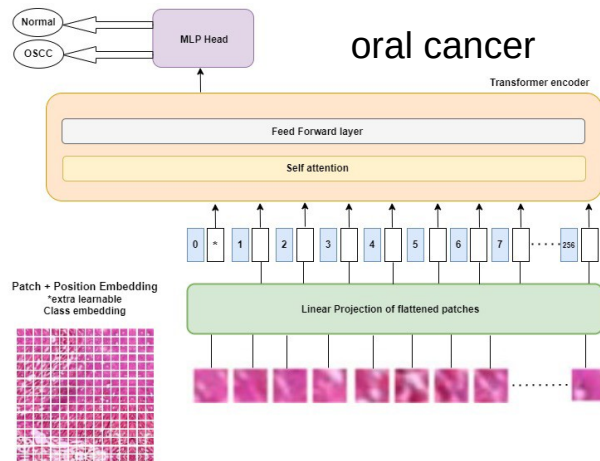
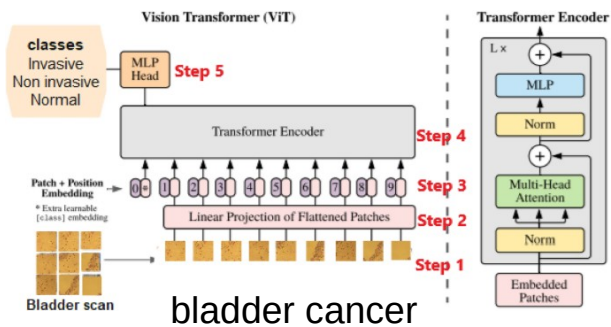
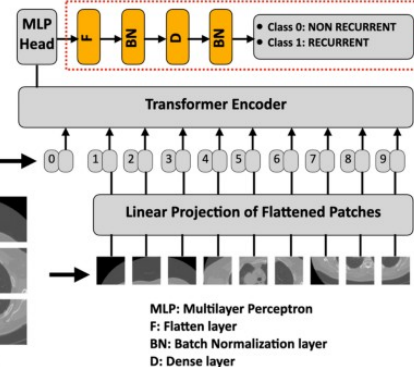
non-small cell lung cancer



ORIGINAL CT

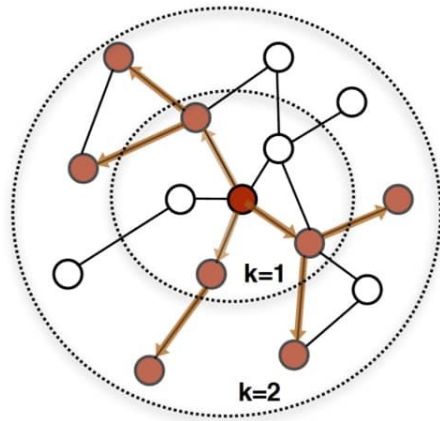


CT PATCHES

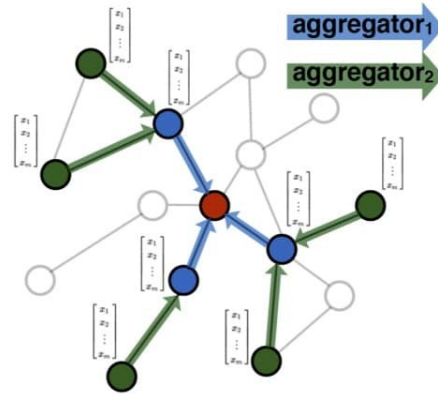


Graph Neural Networks (GNNs)

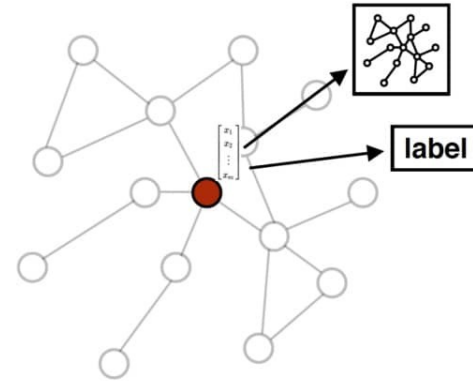
source: <https://blogs.nvidia.com/blog/what-are-graph-neural-networks/>



1. Sample neighborhood



2. Aggregate feature information from neighbors

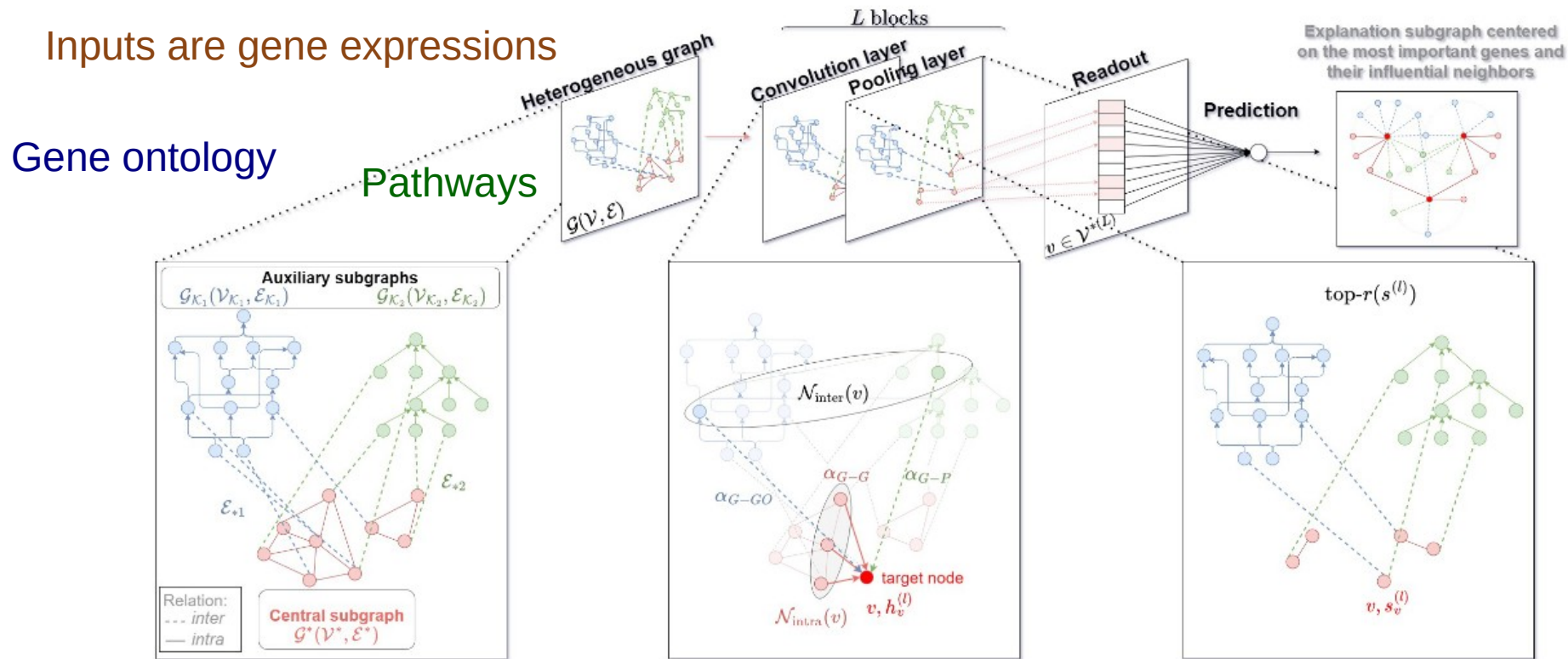


3. Predict graph context and label using aggregated information

The Graph Neural Network Model

Franco Scarselli, Marco Gori, *Fellow, IEEE*, Ah Chung Tsoi, Markus Hagenbuchner, *Member, IEEE*, and Gabriele Monfardini

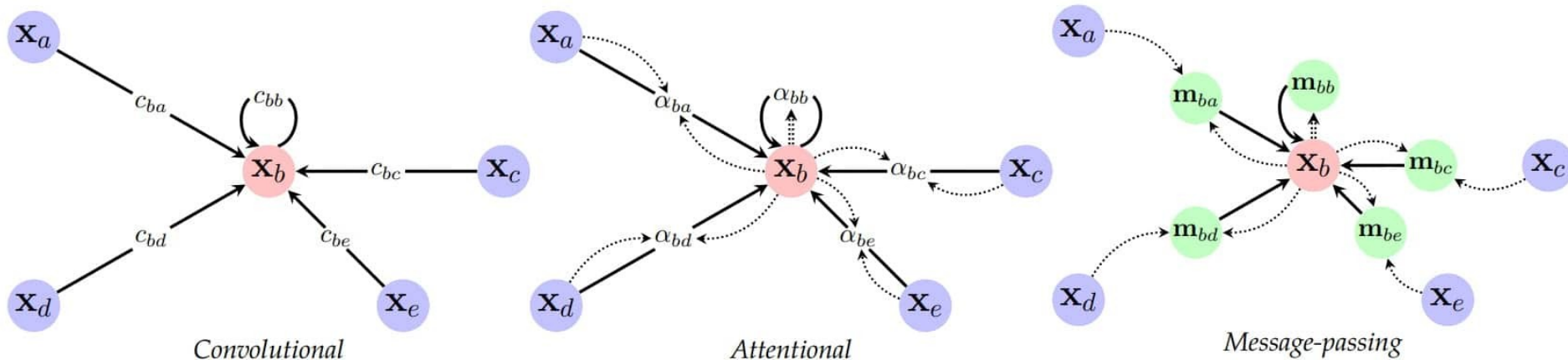
GNN can be heterogeneous



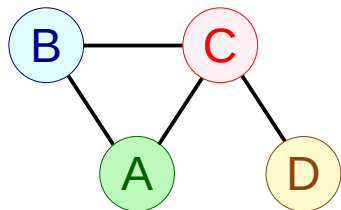
**BioHAN: a Knowledge-based Heterogeneous Graph
Neural Network for precision medicine on
transcriptomic data**

Many different ways to update GNNs

source: <https://blogs.nvidia.com/blog/what-are-graph-neural-networks/>



Example of GNN convolution



$$D = \begin{matrix} & \begin{matrix} A & B & C & D \end{matrix} \\ \begin{matrix} A \\ B \\ C \\ D \end{matrix} & \begin{bmatrix} 2 & 0 & 0 & 0 \\ 0 & 2 & 0 & 0 \\ 0 & 0 & 3 & 0 \\ 0 & 0 & 0 & 1 \end{bmatrix} \end{matrix}$$

Degree matrix

$$A = \begin{bmatrix} 0 & 1 & 1 & 0 \\ 1 & 0 & 1 & 0 \\ 1 & 1 & 0 & 1 \\ 0 & 0 & 1 & 0 \end{bmatrix}$$

Adjacency matrix

$$L = D - A = \begin{bmatrix} 2 & -1 & -1 & 0 \\ -1 & 2 & -1 & 0 \\ -1 & -1 & 3 & -1 \\ 0 & 0 & -1 & 1 \end{bmatrix}$$

Laplacian matrix

NB: undirected graph $\rightarrow$ all matrices are symmetric
This could be different for a directed graph

**Convolutional Neural Networks on Graphs
with Fast Localized Spectral Filtering**

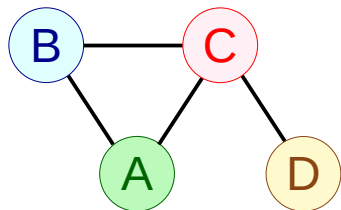
Michaël Defferrard Xavier Bresson Pierre Vandergheynst

EPFL, Lausanne, Switzerland

{michael.defferrard,xavier.bresson,pierre.vandergheynst}@epfl.ch

30th Conference on Neural Information Processing Systems (NIPS 2016), Barcelona, Spain.

Example of GNN convolution



$$D = \begin{matrix} & \begin{matrix} A & B & C & D \end{matrix} \\ \begin{matrix} A \\ B \\ C \\ D \end{matrix} & \begin{bmatrix} 2 & 0 & 0 & 0 \\ 0 & 2 & 0 & 0 \\ 0 & 0 & 2 & 0 \\ 0 & 0 & 0 & 1 \end{bmatrix} \end{matrix}$$

Degree matrix

$$A = \begin{bmatrix} 0 & 1 & 1 & 0 \\ 1 & 0 & 1 & 0 \\ 1 & 1 & 0 & 1 \\ 0 & 0 & 1 & 0 \end{bmatrix}$$

Adjacency matrix

$$L = D - A = \begin{bmatrix} 2 & -1 & -1 & 0 \\ -1 & 2 & -1 & 0 \\ -1 & -1 & 3 & -1 \\ 0 & 0 & -1 & 1 \end{bmatrix}$$

Laplacian matrix

identity matrix
no influence from neighbours

influence from
immediate neighbours

influence from neighbours and
neighbours of neighbours

$$p_w(L) = w_0 I + w_1 L + w_2 L^2 + \dots + w_k L^k$$

$$x' = p_w(L)x$$

$$w = [w_0, w_1, w_2, \dots, w_k] \quad \text{kernel de convolution}$$

**Convolutional Neural Networks on Graphs
with Fast Localized Spectral Filtering**

Michaël Defferrard

Xavier Bresson

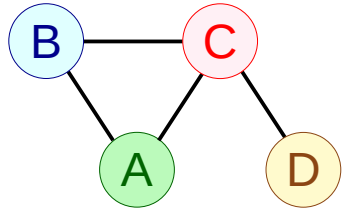
Pierre Vandergheynst

EPFL, Lausanne, Switzerland

{michael.defferrard,xavier.bresson,pierre.vandergheynst}@epfl.ch

30th Conference on Neural Information Processing Systems (NIPS 2016), Barcelona, Spain.

Example of GNN convolution



D only influences C

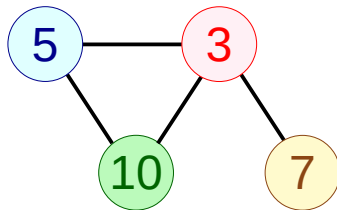
$$L = \begin{bmatrix} 2 & -1 & -1 & 0 \\ -1 & 2 & -1 & 0 \\ -1 & -1 & 3 & -1 \\ 0 & 0 & -1 & 1 \end{bmatrix}$$

D influences A, B, C

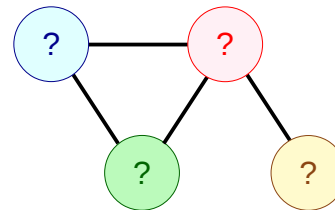
$$L^2 = \begin{bmatrix} 6 & -3 & -4 & 1 \\ -3 & 6 & -4 & 1 \\ -4 & -4 & 12 & -4 \\ 1 & 1 & -4 & 2 \end{bmatrix}$$

$$w = [1, 0.1, 0.01]$$

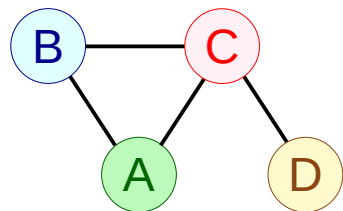
layer 1



layer 2



Example of GNN convolution



D only influences C



$$L = \begin{bmatrix} 2 & -1 & -1 & 0 \\ -1 & 2 & -1 & 0 \\ -1 & -1 & 3 & -1 \\ 0 & 0 & -1 & 1 \end{bmatrix}$$

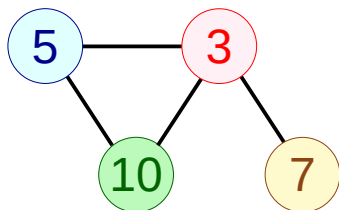
D influences A, B, C



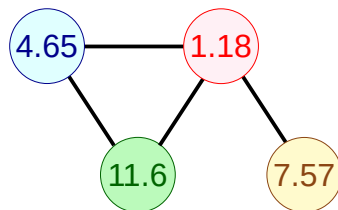
$$L^2 = \begin{bmatrix} 6 & -3 & -4 & 1 \\ -3 & 6 & -4 & 1 \\ -4 & -4 & 12 & -4 \\ 1 & 1 & -4 & 2 \end{bmatrix}$$

$$w = [1, 0.1, 0.01] \quad 1 \begin{bmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \end{bmatrix} \begin{bmatrix} 10 \\ 5 \\ 3 \\ 7 \end{bmatrix} + 0.1 \begin{bmatrix} 2 & -1 & -1 & 0 \\ -1 & 2 & -1 & 0 \\ -1 & -1 & 3 & -1 \\ 0 & 0 & -1 & 1 \end{bmatrix} \begin{bmatrix} 10 \\ 5 \\ 3 \\ 7 \end{bmatrix} + 0.01 \begin{bmatrix} 6 & -3 & -4 & 1 \\ -3 & 6 & -4 & 1 \\ -4 & -4 & 12 & -4 \\ 1 & 1 & -4 & 2 \end{bmatrix} \begin{bmatrix} 10 \\ 5 \\ 3 \\ 7 \end{bmatrix}$$

layer 1

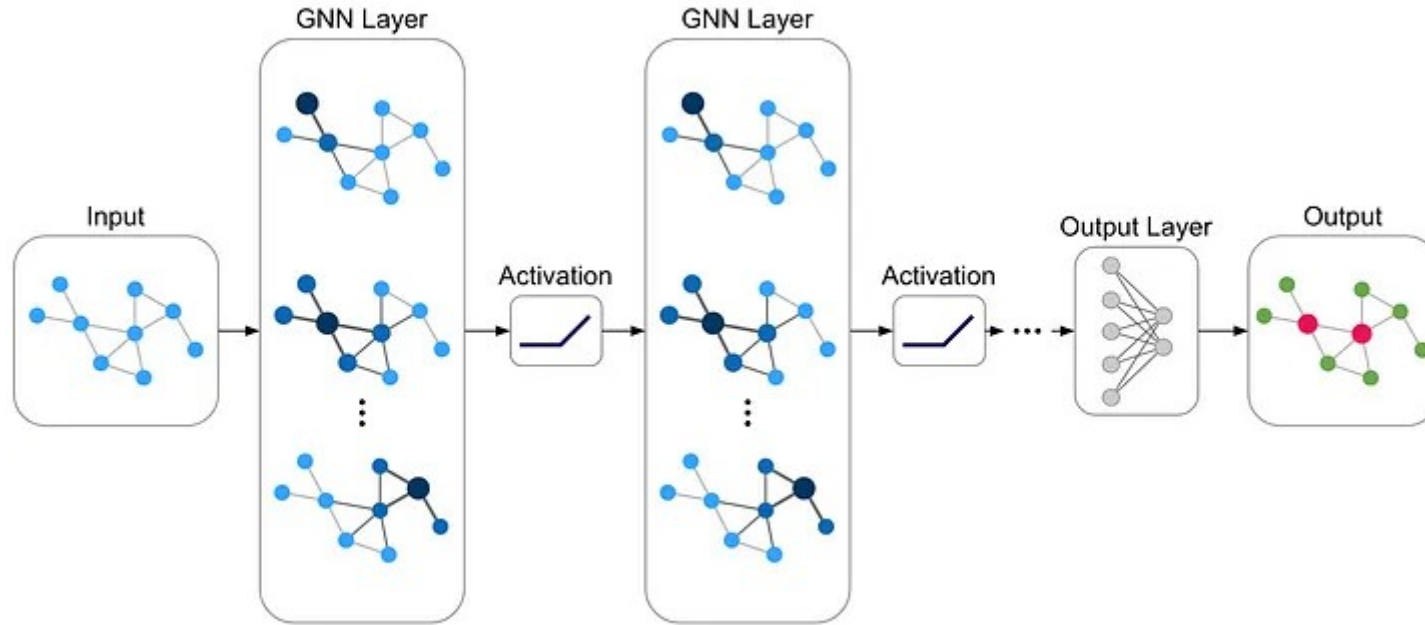


layer 2



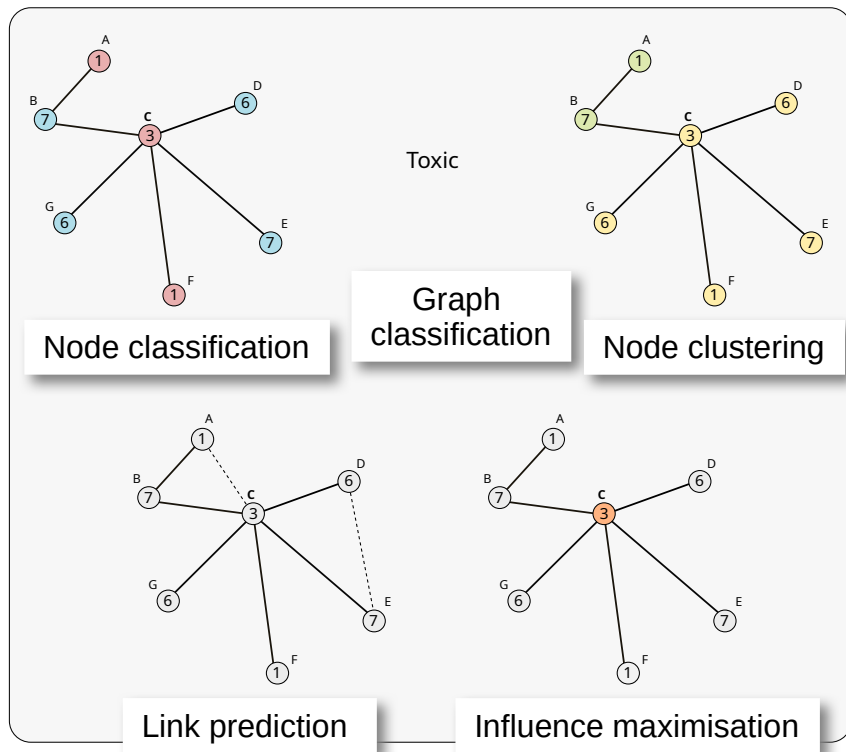
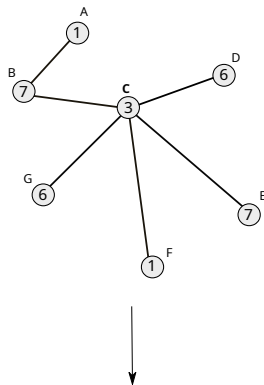
$$\sum \text{layer1} = \sum \text{layer2}$$

Graph Neural Networks (GNNs)



NB: GNNs generally comprise 3 embeddings that are updated at each iteration, i.e nodes (vertices), edges, and graph

GNNs: What can we do?



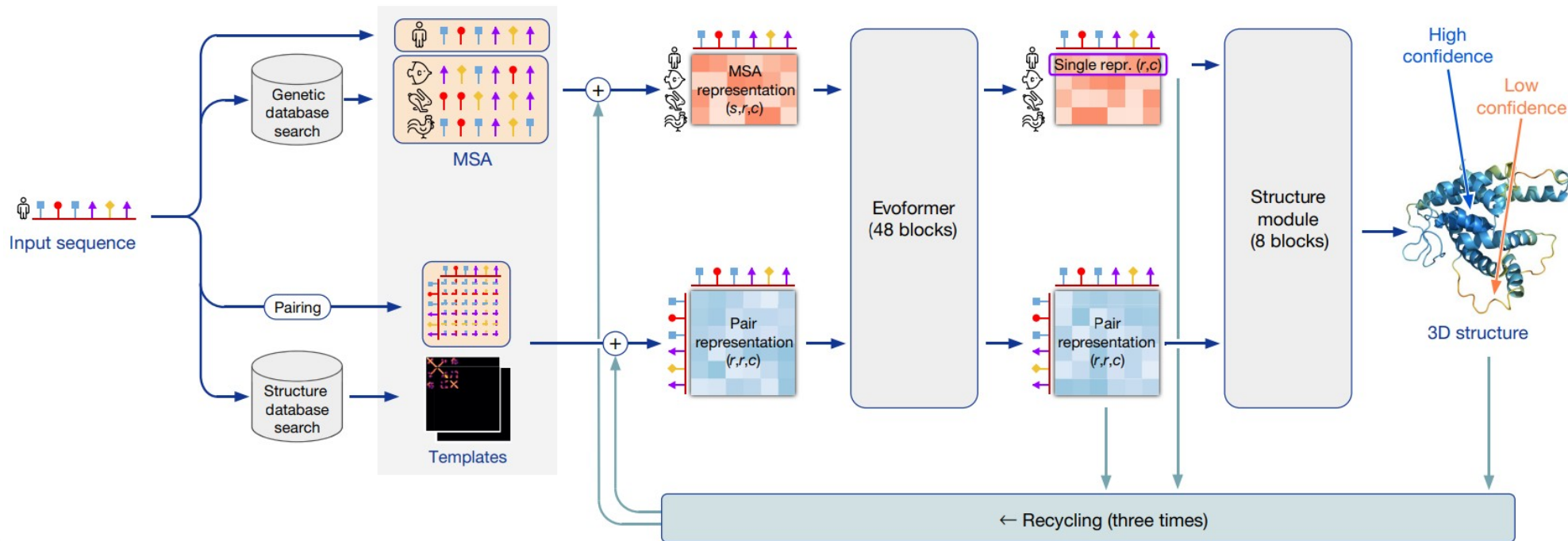
Source: Understanding Convolutions on Graphs
<https://distill.pub/2021/understanding-gnns/>

See also: A Gentle Introduction to Graph Neural Networks
<https://distill.pub/2021/gnn-intro/>

Both by Google Research teams



AlphaFold2



Highly accurate protein structure prediction with AlphaFold

<https://doi.org/10.1038/s41586-021-03819-2>

Received: 11 May 2021

Accepted: 12 July 2021

Published online: 15 July 2021

Open access

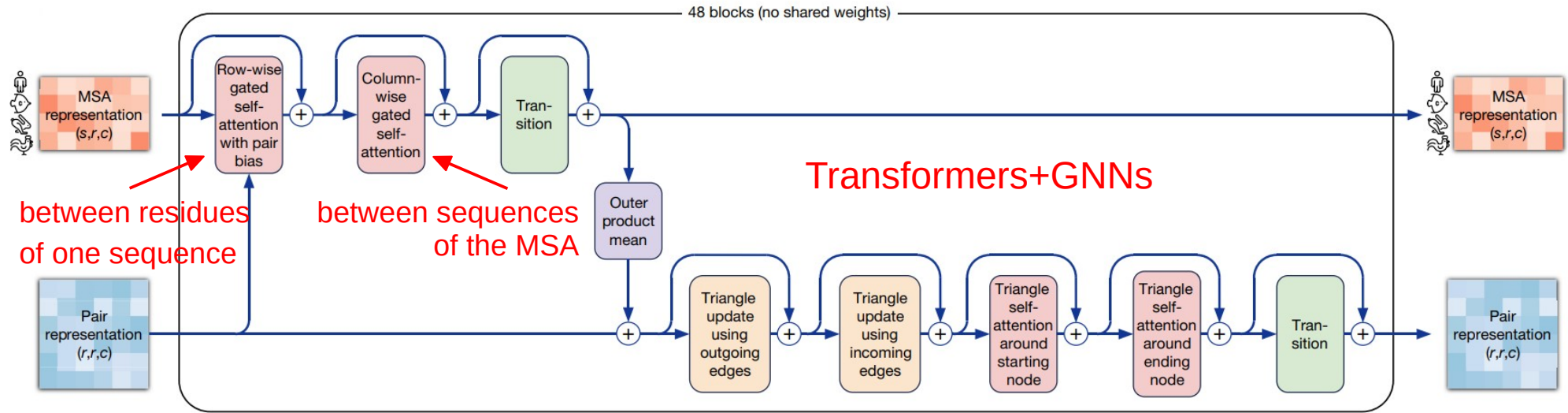
Check for updates

John Jumper^{1,4,5}, Richard Evans^{1,4}, Alexander Pritzel^{1,4}, Tim Green^{1,4}, Michael Figurnov^{1,4}, Olaf Ronneberger^{1,4}, Kathryn Tunyasuvunakool^{1,4}, Russ Bates^{1,4}, Augustin Židek^{1,4}, Anna Potapenko^{1,4}, Alex Bridgland^{1,4}, Clemens Meyer^{1,4}, Simon A. A. Kohl^{1,4}, Andrew J. Ballard^{1,4}, Andrew Cowie^{1,4}, Bernardino Romera-Paredes^{1,4}, Stanislav Nikolov^{1,4}, Rishub Jain^{1,4}, Jonas Adler¹, Trevor Back¹, Stig Petersen¹, David Reiman¹, Ellen Clancy¹, Michal Zielinski¹, Martin Steinegger^{2,3}, Michalina Pacholska¹, Tamas Berghammer¹, Sebastian Bodenstein¹, David Silver¹, Oriol Vinyals¹, Andrew W. Senior¹, Koray Kavukcuoglu¹, Pushmeet Kohli¹ & Demis Hassabis^{1,4,5}

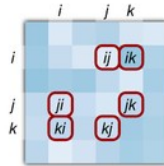
~93 million parameters (weights+biases)

<https://github.com/google-deepmind/alphafold>

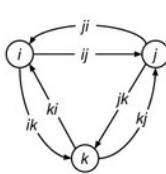
AlphaFold2: evoformer



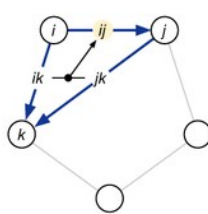
Pair representation (r, r, c)



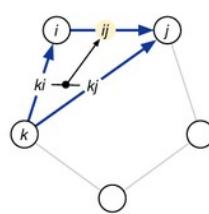
Corresponding edges in a graph



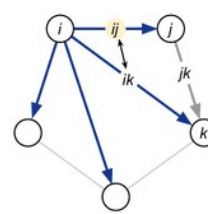
Triangle multiplicative update using 'outgoing' edges



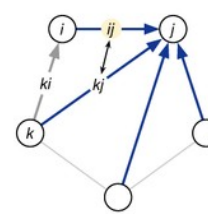
Triangle multiplicative update using 'incoming' edges



Triangle self-attention around starting node

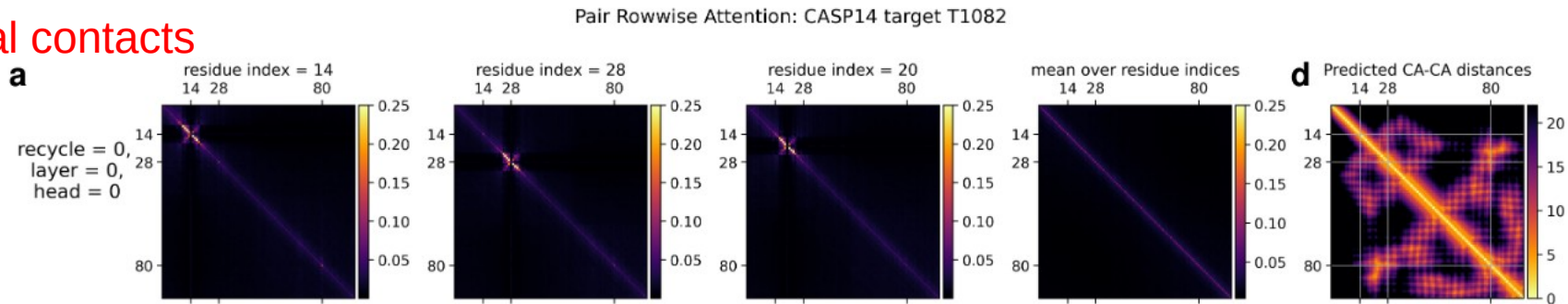


Triangle self-attention around ending node

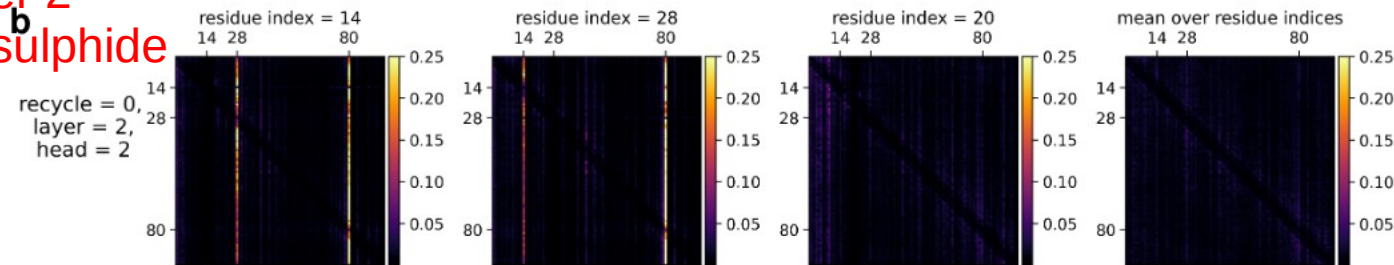


Row-wise attention: between residues of a sequence

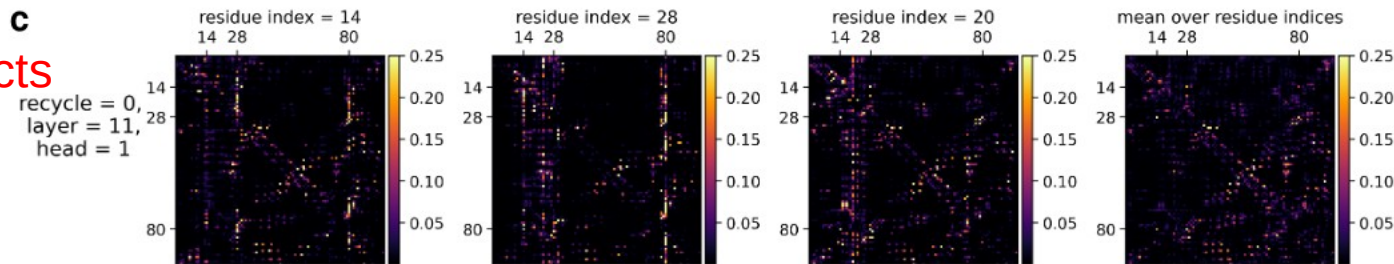
Layer 0 = local contacts



Head 2 of layer 2 recognises disulphide bonds



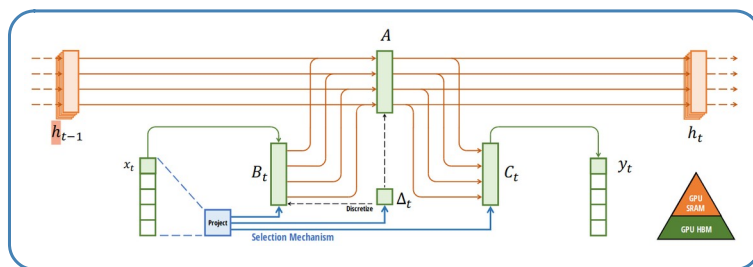
At layer 11, heads recognise distant contacts



Source:
AlphaFold2 paper
supplementary info

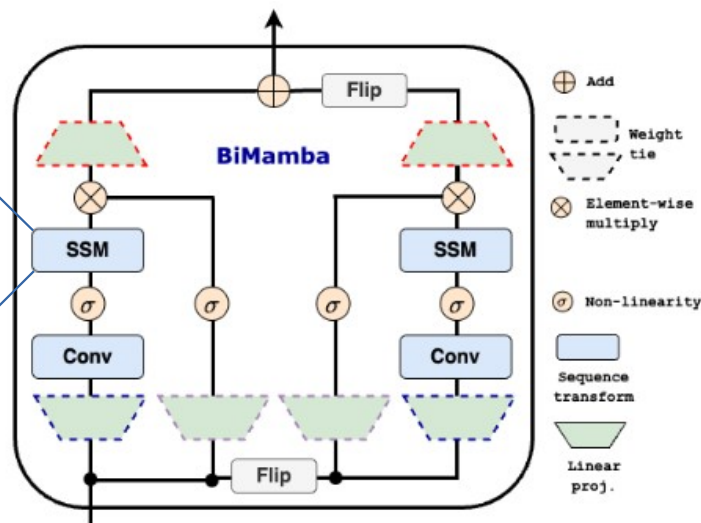
RNNs are back. Rise of the Mamba

“Attention” = embedding size x input length
→ linear growth with input length
(not quadratic like transformers)



SSM = State Space Model

Prediction/classification



...AGGCTAGGATATCGATAAGCTGACTGAT...

The model learns about variants' Interactions
→ Reusable foundation model

