

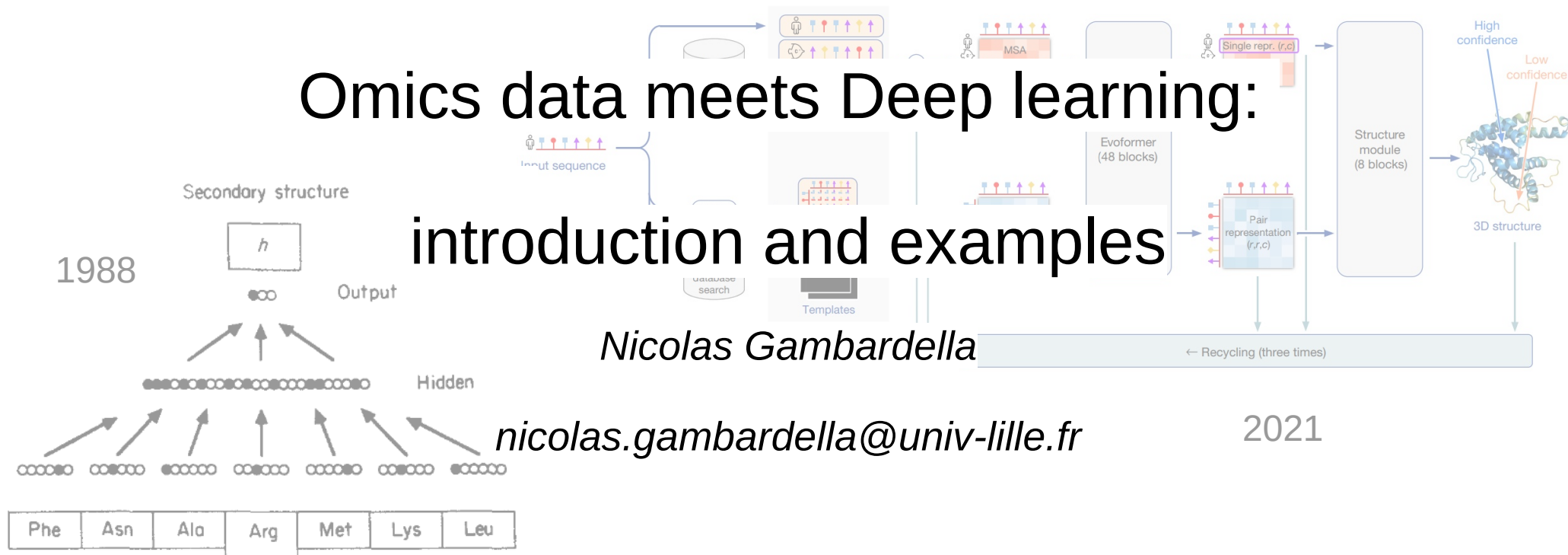
Omics data meets Deep learning:

introduction and examples

Nicolas Gambardella

nicolas.gambardella@univ-lille.fr

2021





ChatGPT ▾

NI



Create an image for my presentation



Suggest a recipe based on a photo of my fridge



Create a workout plan



Write a report based on my data



Message ChatGPT



ChatGPT can make mistakes. Check important info.



What are we going to (attempt to) cover?

- 1) What is artificial intelligence?
- 2) Multi-Layer Perceptrons (MLP)
- 3) MLPs in action
- 4) Convolutional Neural Networks (CNN)
- 5) Embeddings and latent space
- 6) Encoder-Decoders,
(variational) AutoEncoders (VAEs)
- 7) Attention and the Transformer
- 8) Graph Neural Networks (GNNs)
- 9) Alphafold2
- 10) Recurrent Neural Networks (RNNs)
and Long Short-Term Memory (LSTM)
- 11) RNNs are back: Rise of the Mamba
- 12) Generative Adversarial Networks (GANs)

1

What is artificial intelligence?

Some terminology

Assessments, evaluations, decisions, predictions
made by software tools

Artificial
intelligence

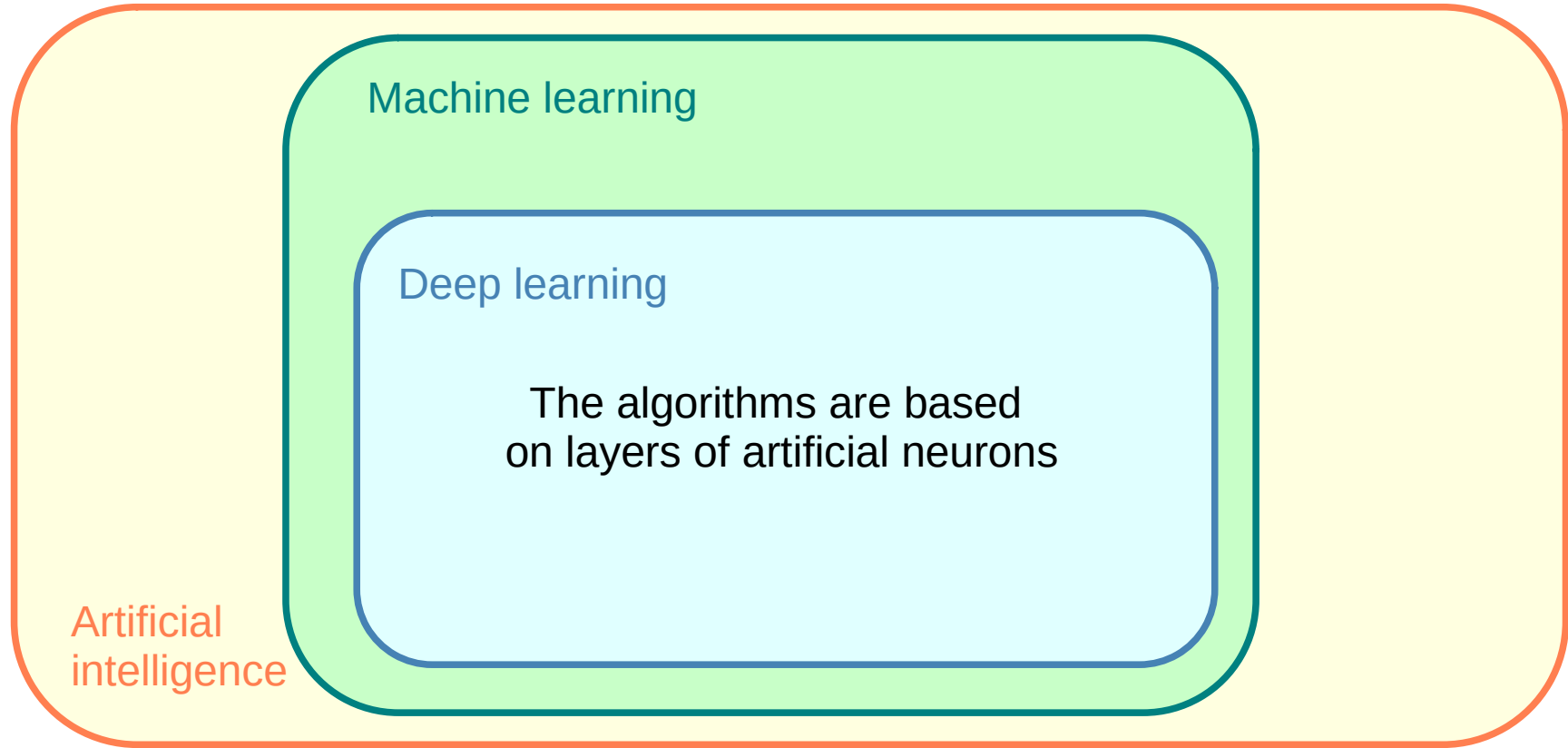
Some terminology

Machine learning

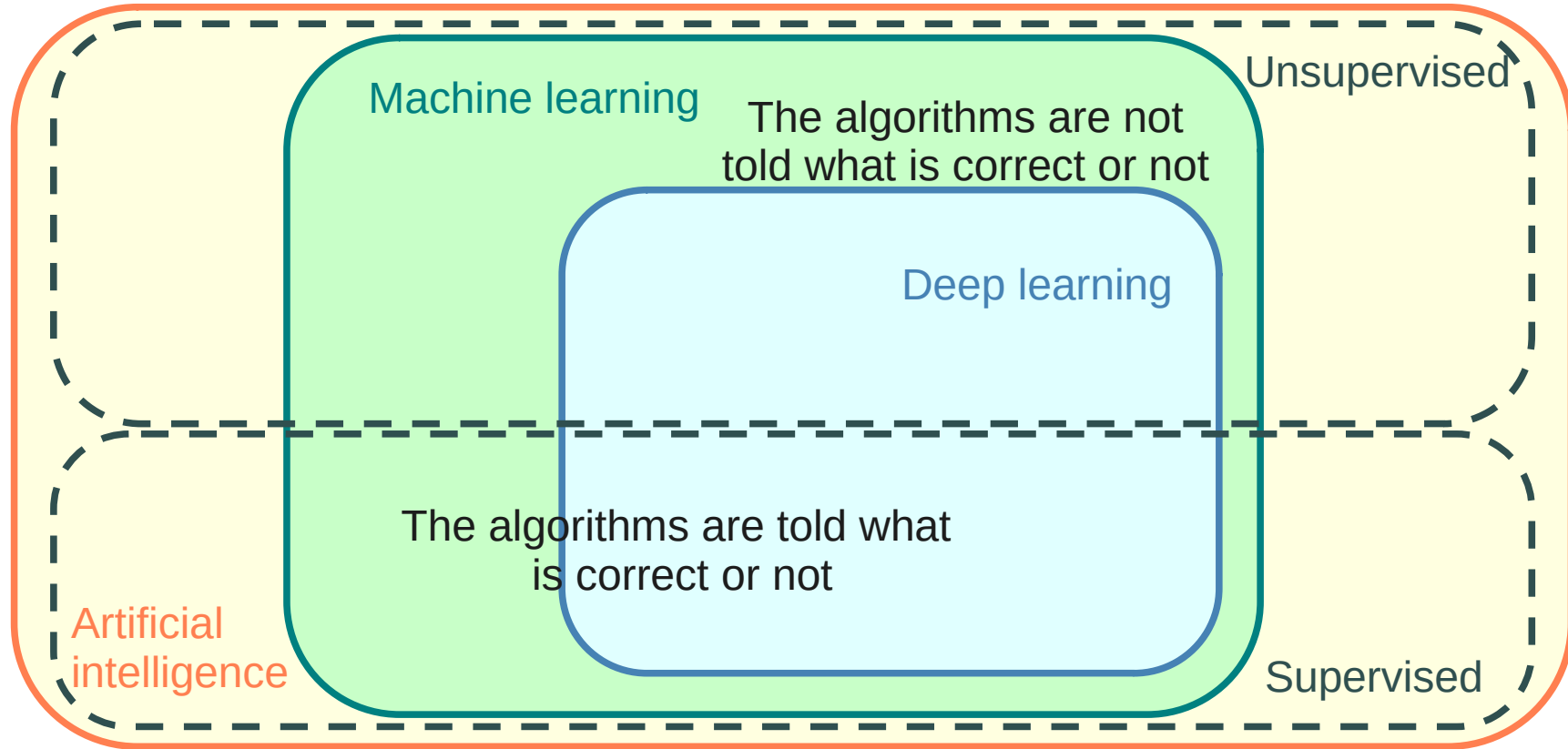
The rules are learned from the data

Artificial
intelligence

Some terminology

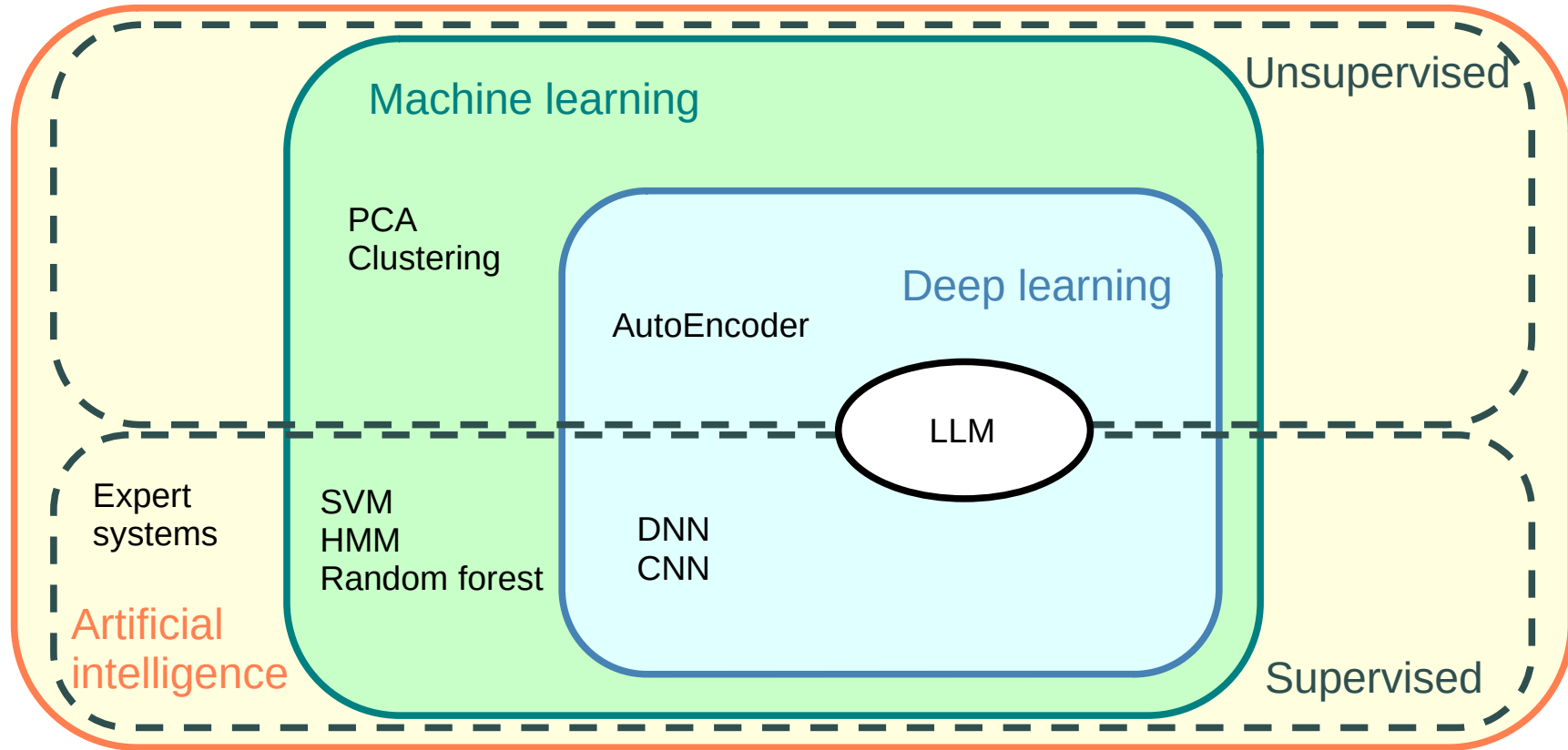


Some terminology

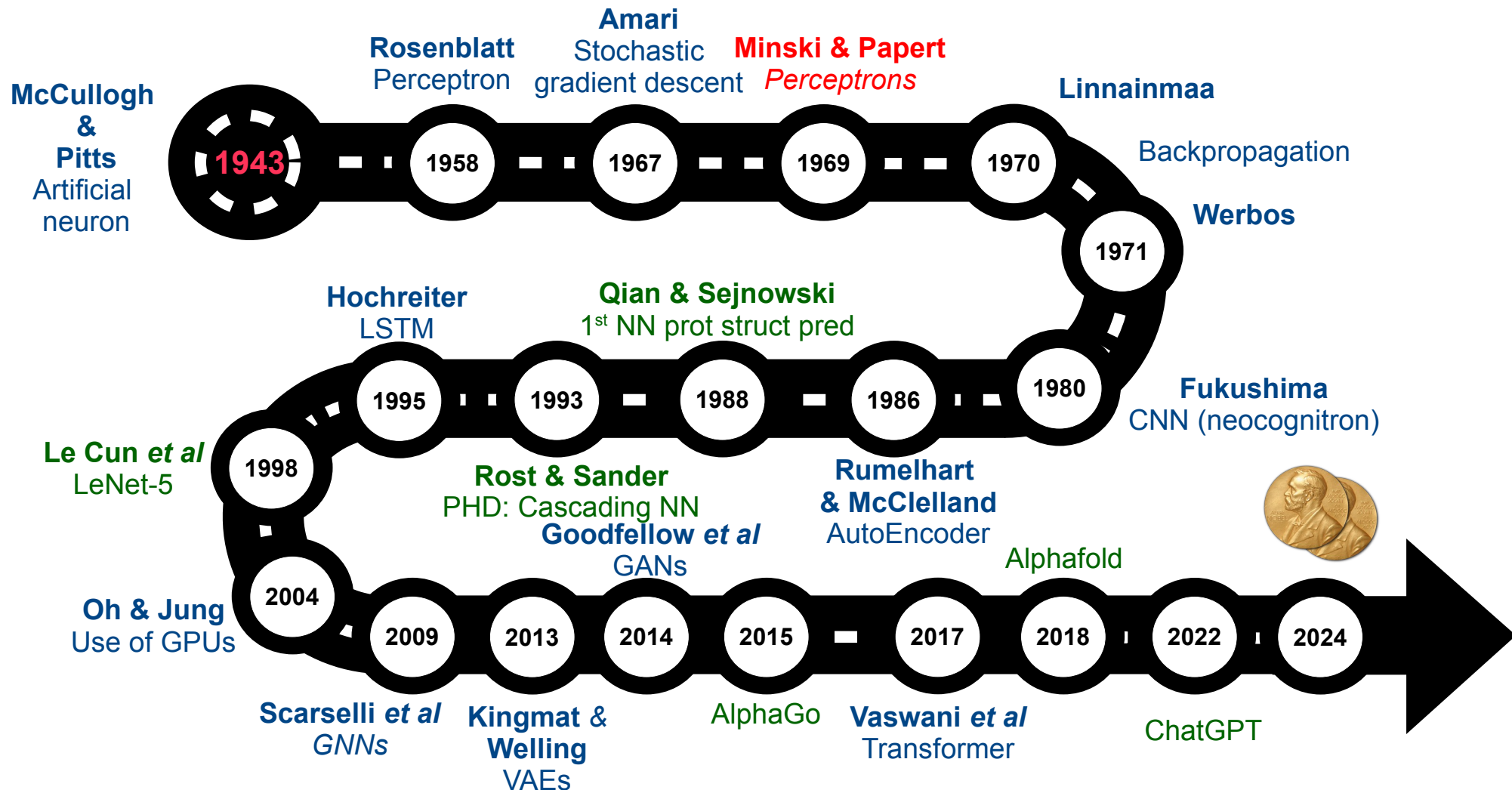


(NB: I consider reinforcement learning as part of supervised, but this is controversial)

Some terminology



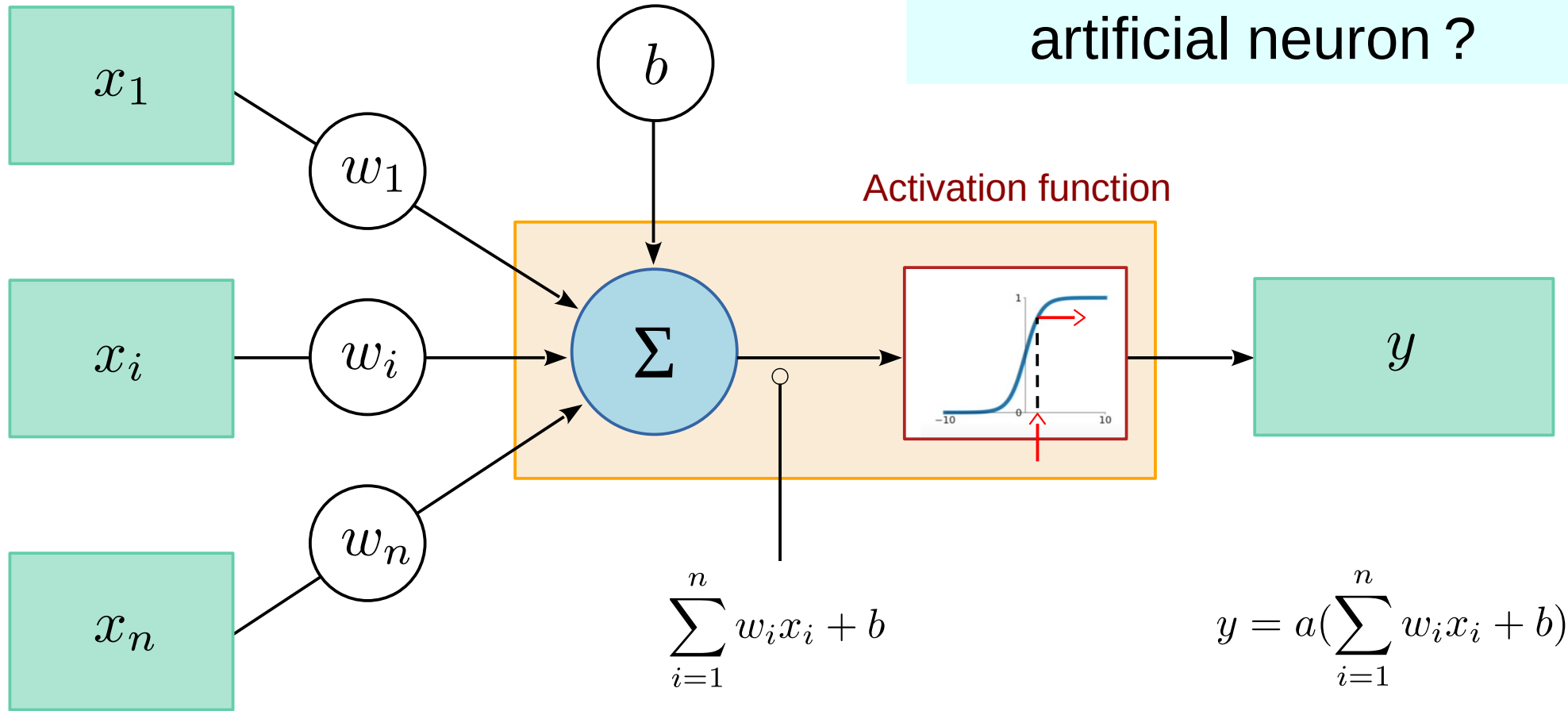
(NB: I consider reinforcement learning as part of supervised, but this is controversial)



2

Multi-Layer Perceptrons (MLP)
a.k.a Fully Connected Networks (FCN)

What is an artificial neuron ?

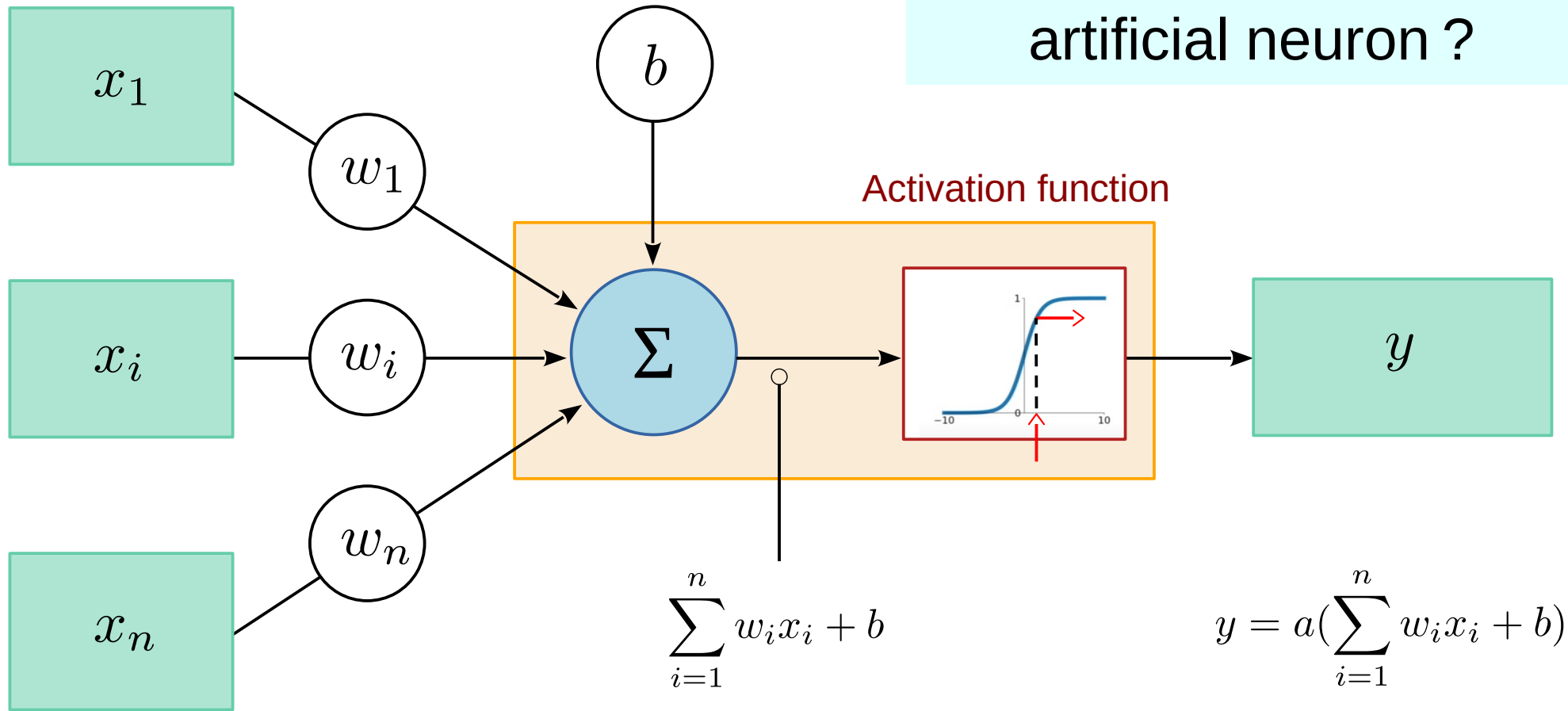


McCulloch and Pitts (1943) A logical calculus of the ideas immanent in nervous activity. *Bull Math Biophys* 5:115-133

Rosenblatt (1958) The perceptron: a probabilistic model for information storage and organization in the brain. *Psychol Rev* 65(6):386-408

Widrow and Hoff (1960) Adaptive Switching circuits. *WESCON Convention record* part IV: 96-104

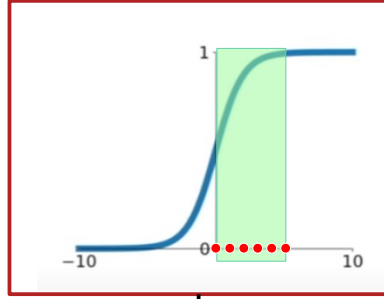
What is an artificial neuron ?



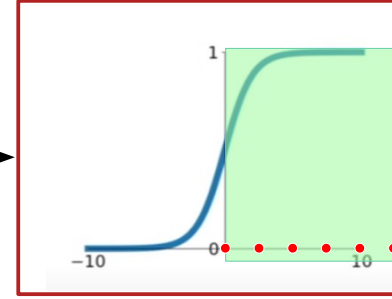
NB: when the activation function is logistic (sigmoid), this is actually a logistic regression...

Impact of the weights and the bias

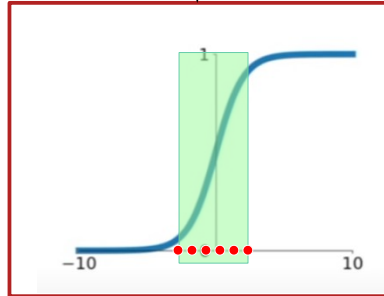
inputs = [0,1,2,3,4,5]



$$w = 3$$

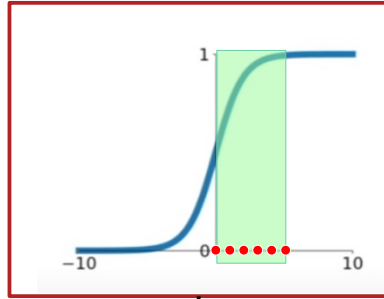


$$b = -2.5$$

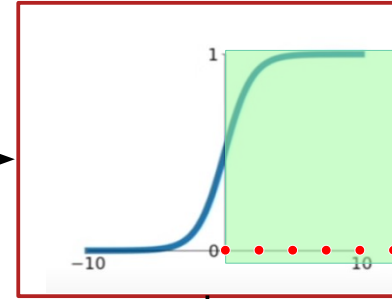


Impact of the weights and the bias

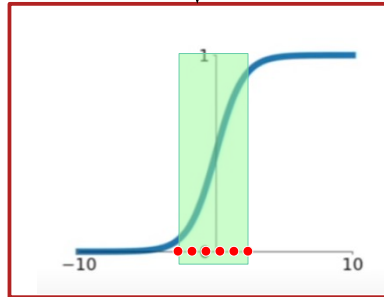
inputs = [0,1,2,3,4,5]



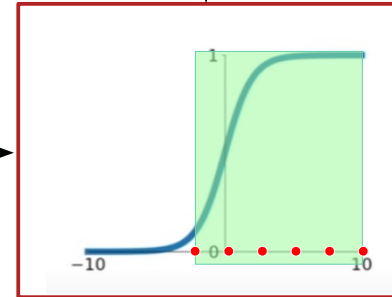
$$w = 3$$



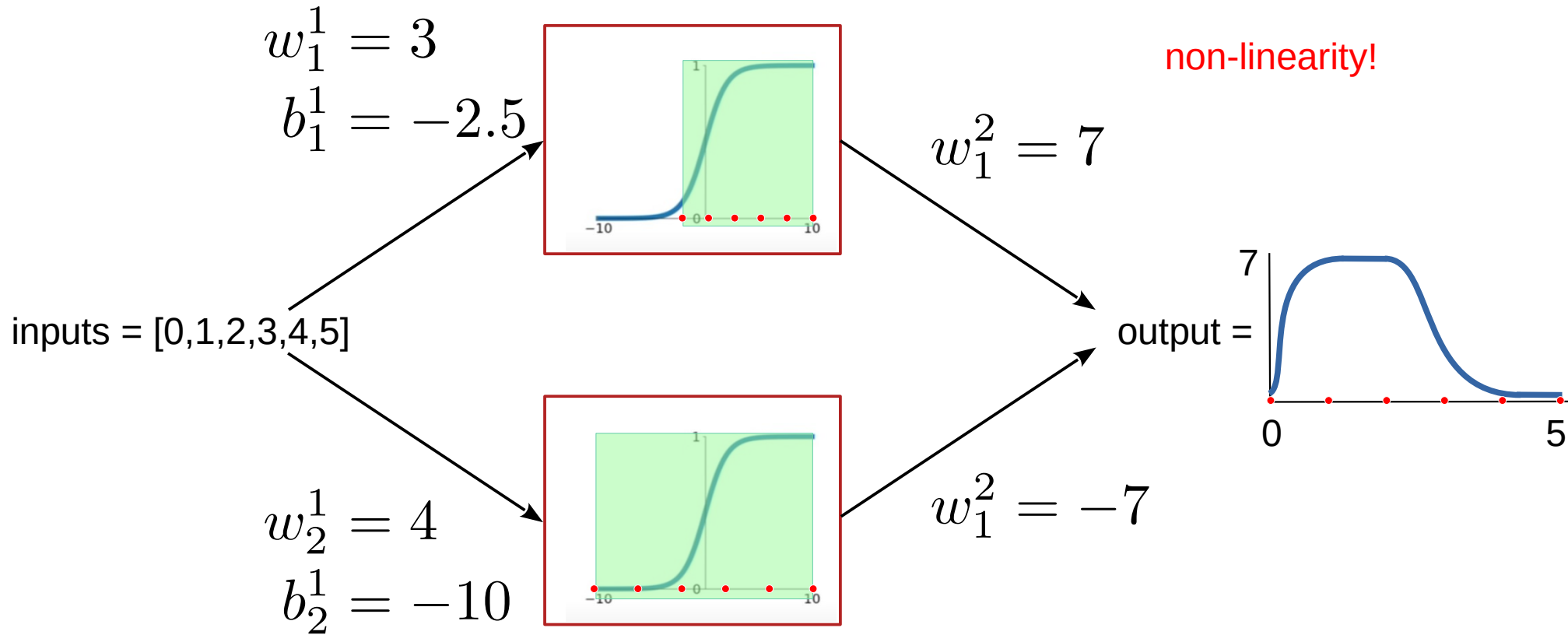
$$b = -2.5$$



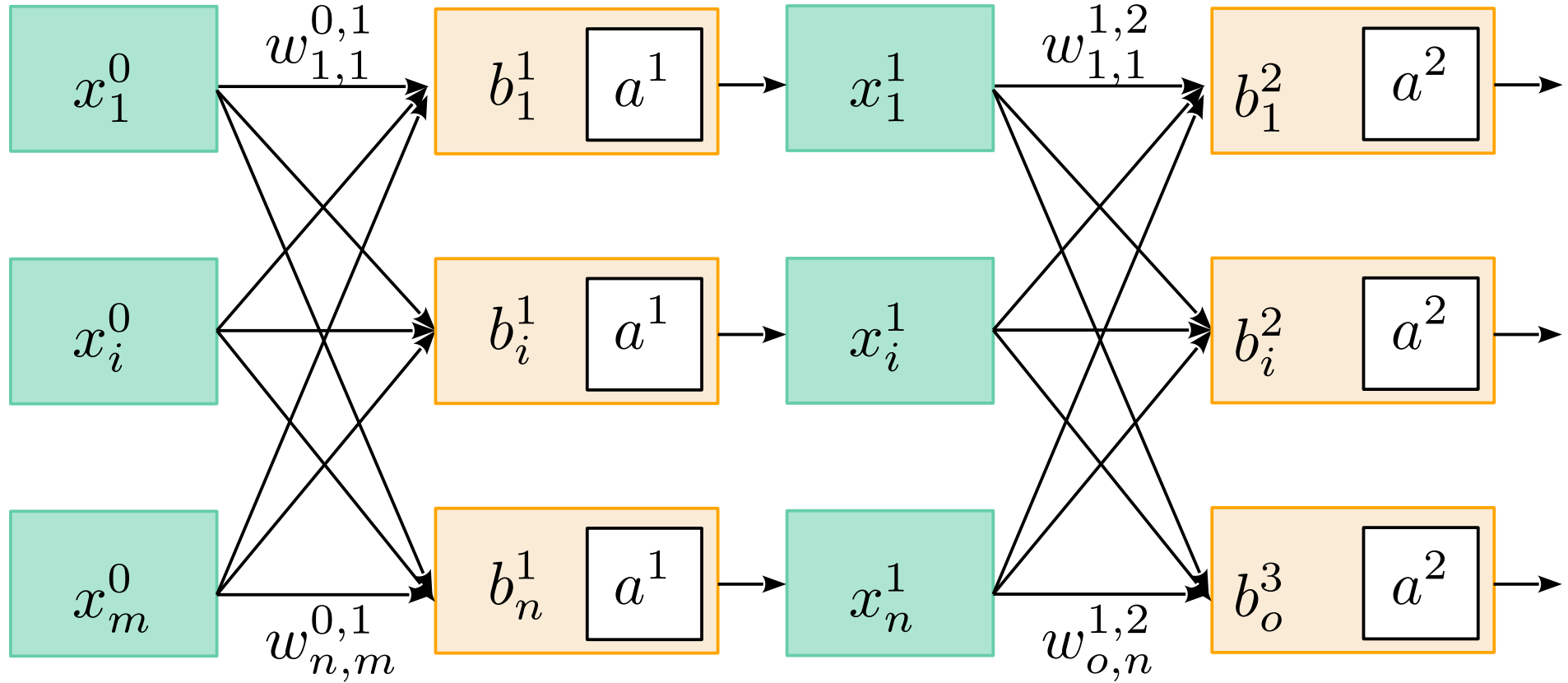
$$w = 3$$



The magic happens with several neurons



Then we add layers (the “Deep”)



And then we add layers (the “Deep”)

f_1 and f_2 can be different

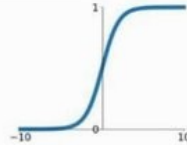
x_1^0

x_i^0

x_m^0

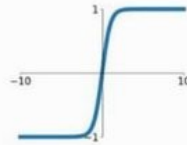
Sigmoid

$$\sigma(x) = \frac{1}{1+e^{-x}}$$



tanh

$$\tanh(x)$$



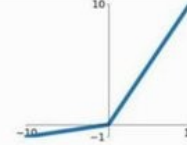
ReLU

$$\max(0, x)$$



Leaky ReLU

$$\max(0.1x, x)$$



Maxout

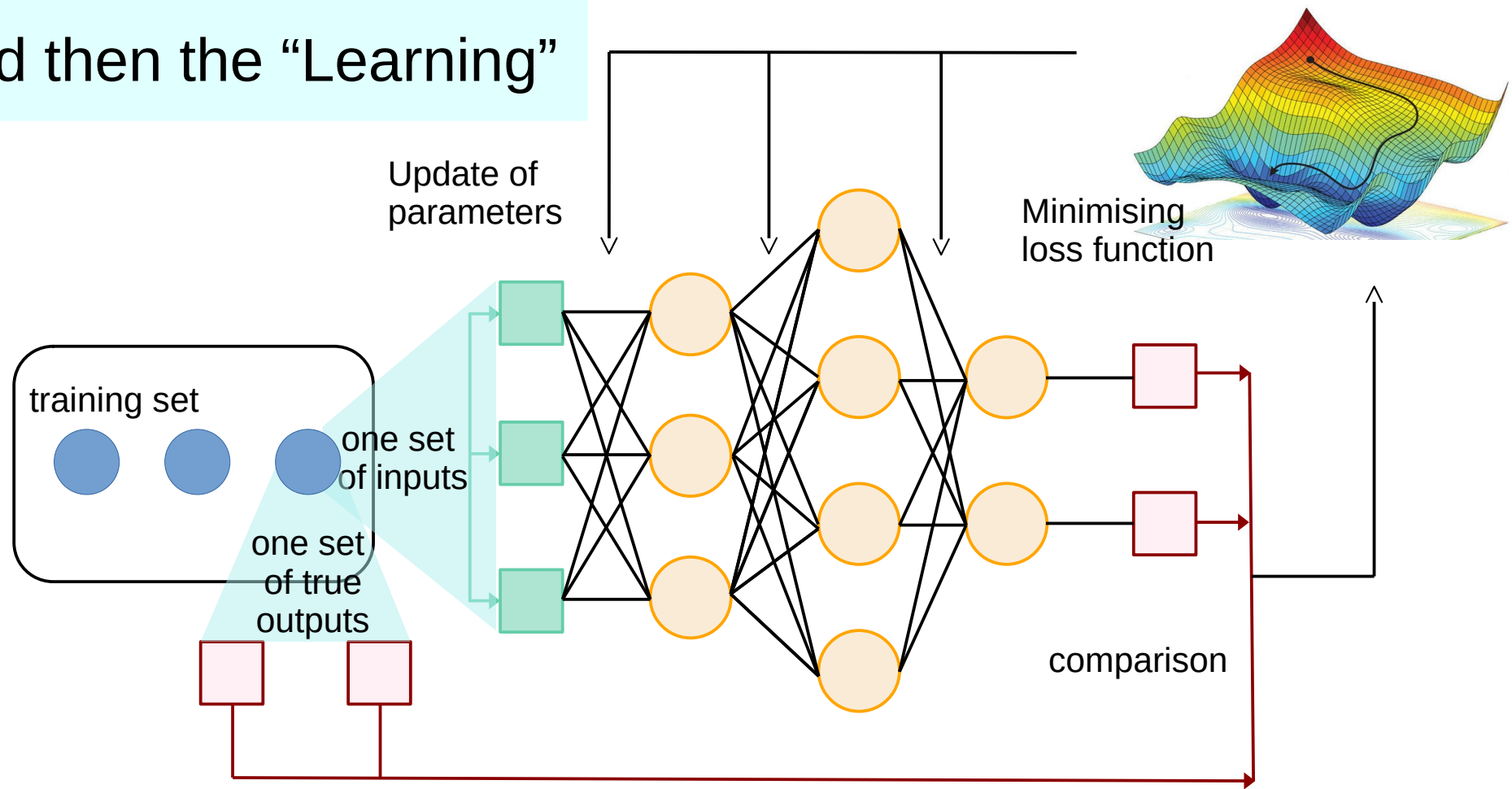
$$\max(w_1^T x + b_1, w_2^T x + b_2)$$

ELU

$$\begin{cases} x & x \geq 0 \\ \alpha(e^x - 1) & x < 0 \end{cases}$$



And then the “Learning”



Widrow and Hoff (1960) Adaptive Switching circuits. *WESCON Convention record* part IV: 96-104

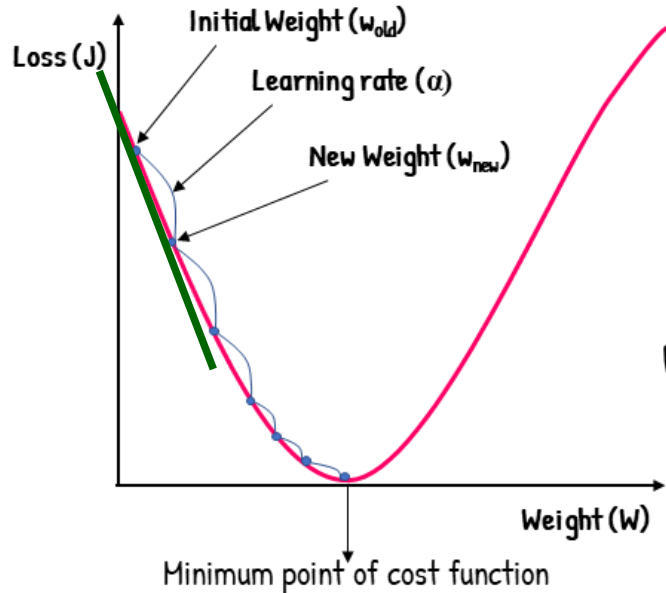
S Amari (1967). A theory of adaptive pattern classifier. *IEEE Transactions*. EC (16): 279–307

S Linnainmaa (1970-1976). The representation of the cumulative rounding error of an algorithm as a Taylor expansion of the local rounding errors (Masters). University of Helsinki. p. 6–7.

P Werbos (1971-1982) Applications of advances in non-linear sensitivity analysis. *LNCIS* 38: 762-770

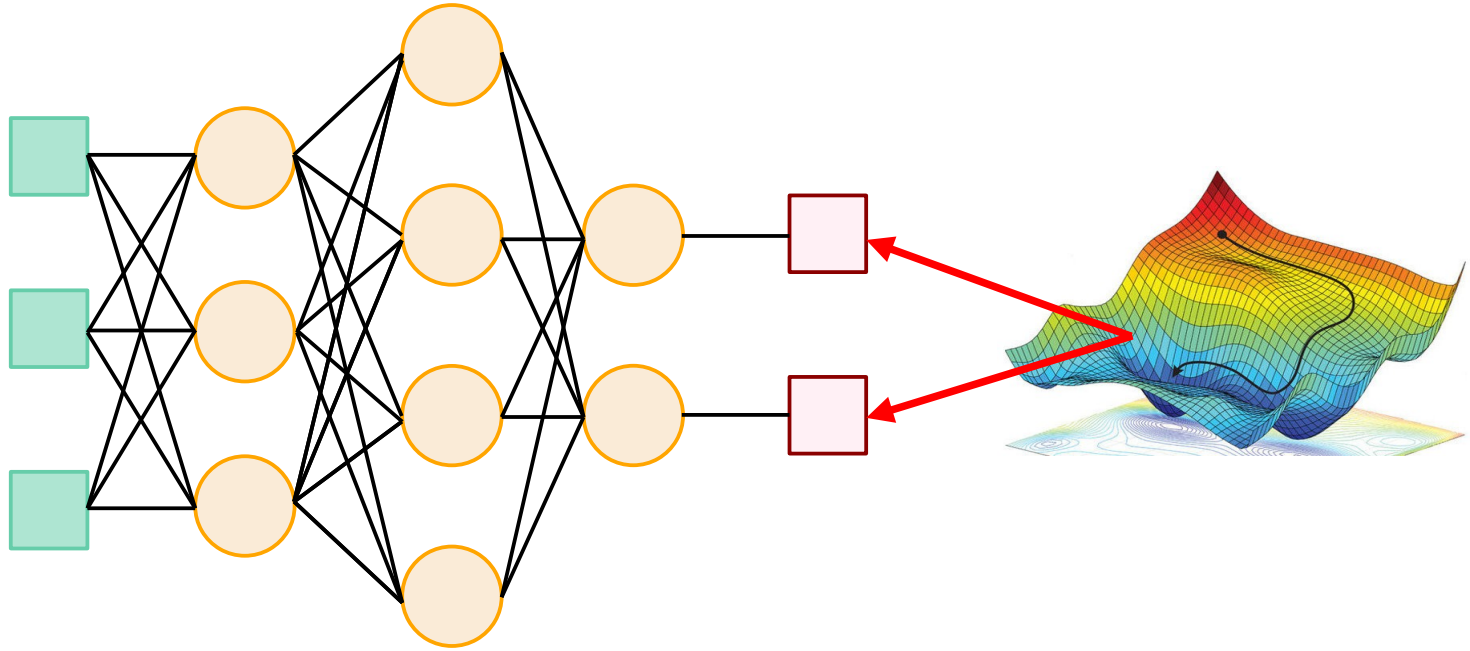
Optimization, e.g., gradient descent

To minimise the loss function (called L , C , or most often J), we will calculate the “gradient” – the derivative – of the loss function with a set of parameters and calculate a new set using this gradient and the *learning rate*.

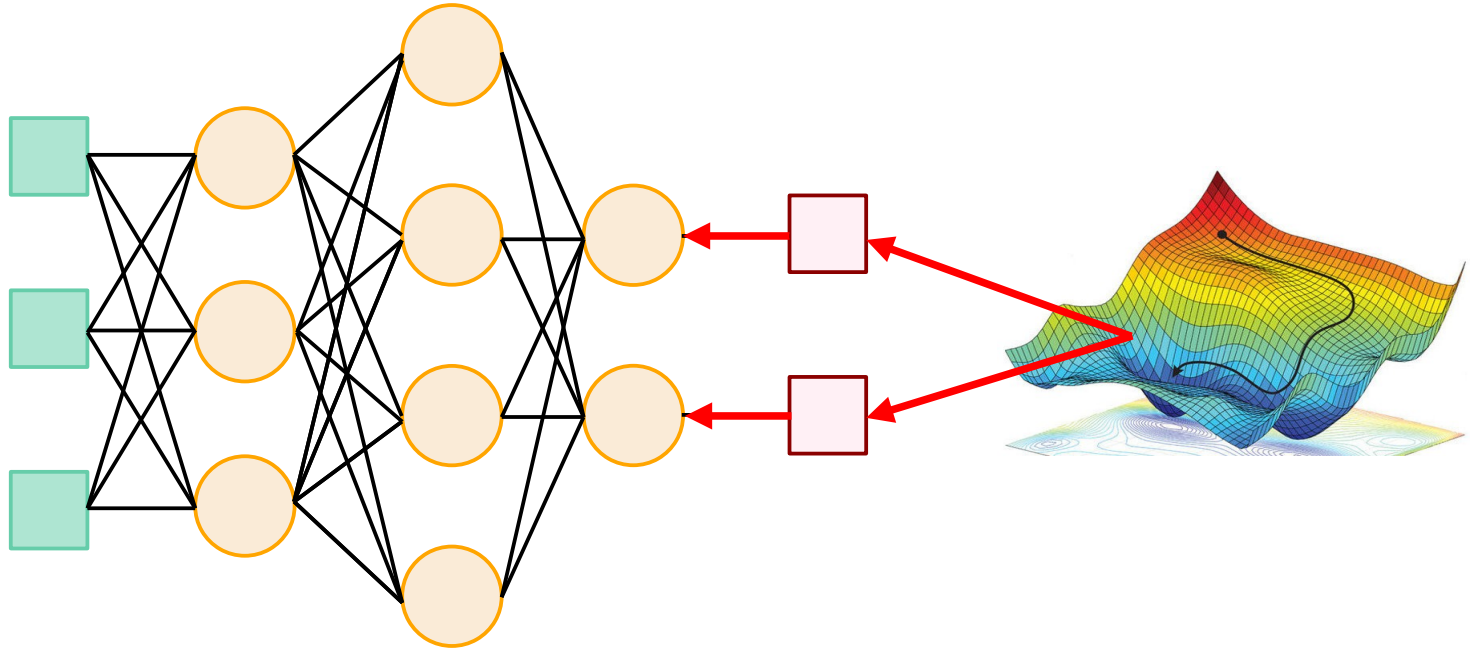


$$w_{new} = w_{old} - \alpha \frac{\delta J}{\delta w}$$

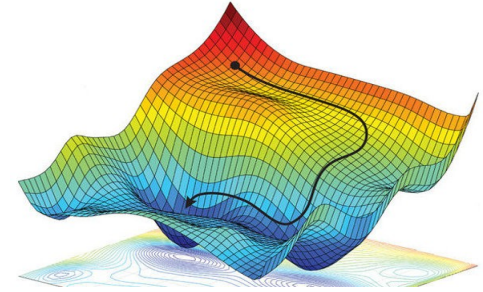
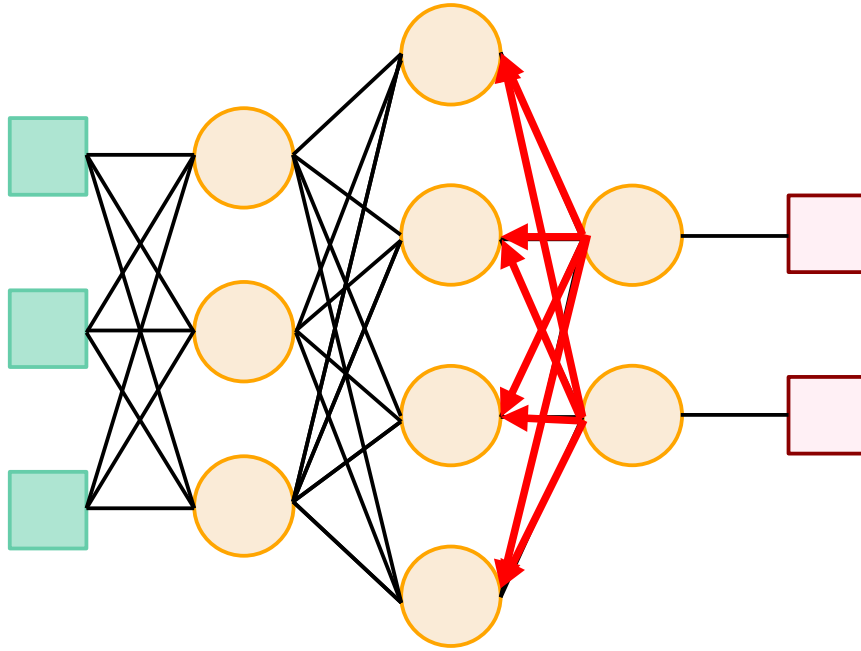
Error backpropagation one layer at a time



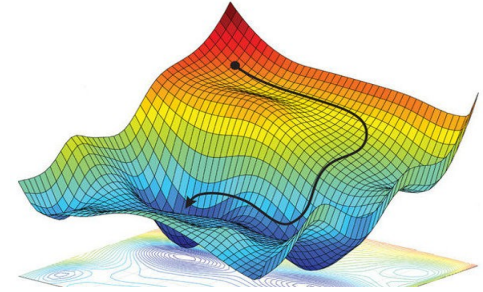
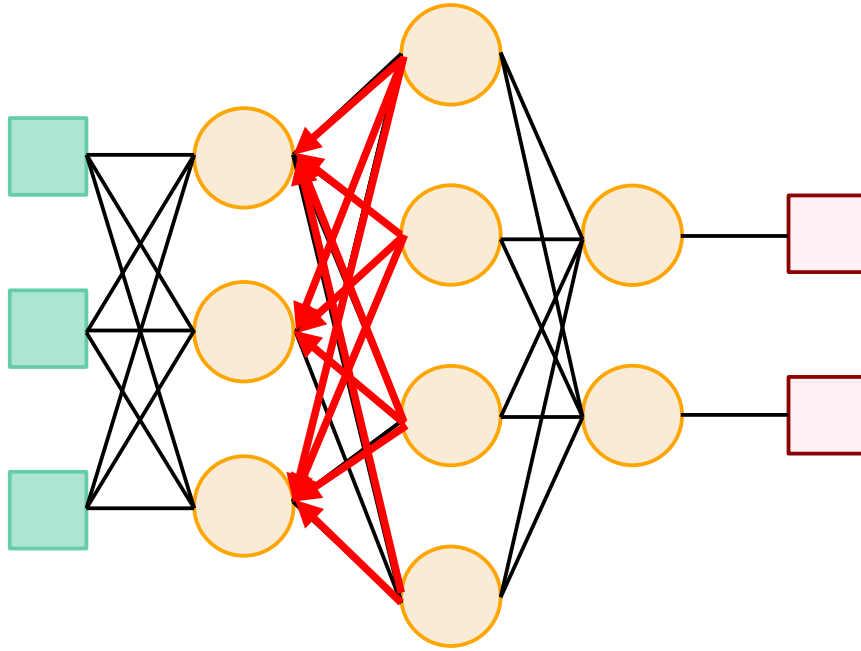
Error backpropagation one layer at a time



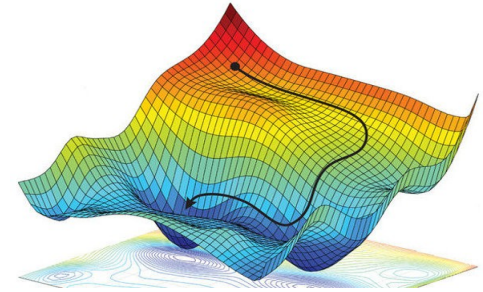
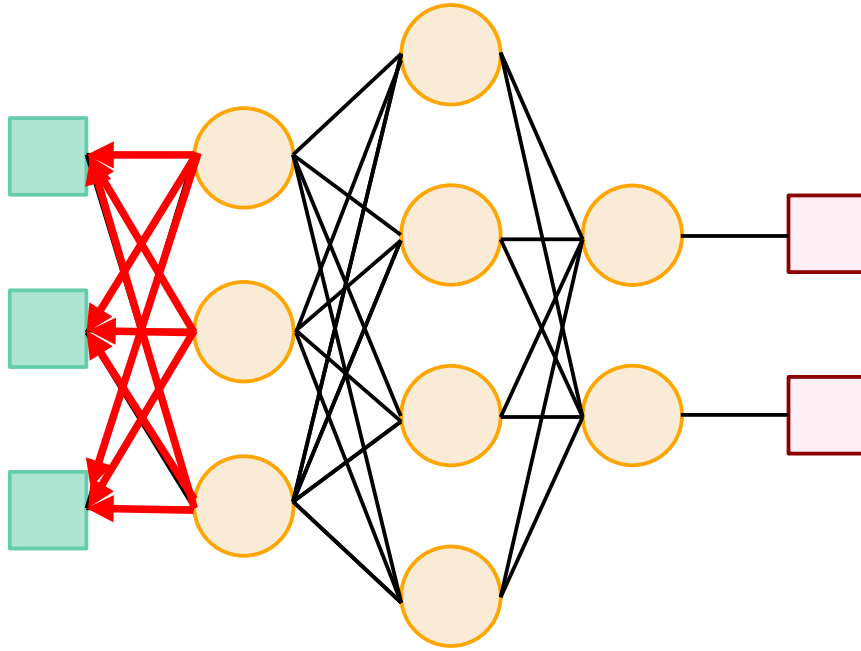
Error backpropagation one layer at a time



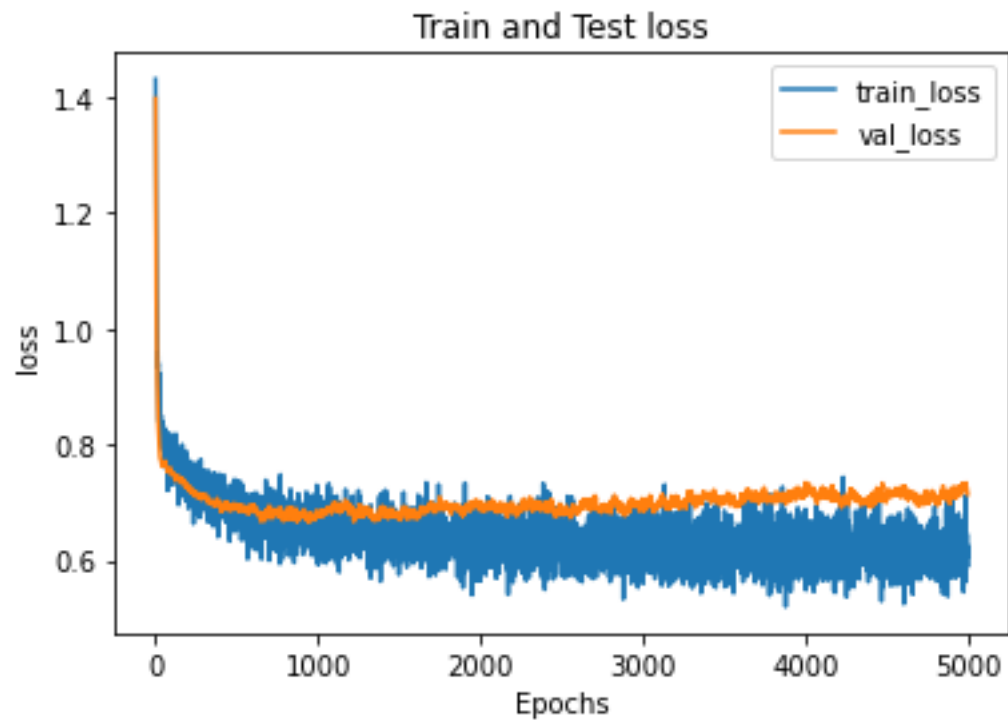
Error backpropagation one layer at a time



Error backpropagation one layer at a time



Learning



Training, testing, and validation sets

“validation” (never seen)
Same for all model instances
Used to assess the model at the end

Training set: used to learn

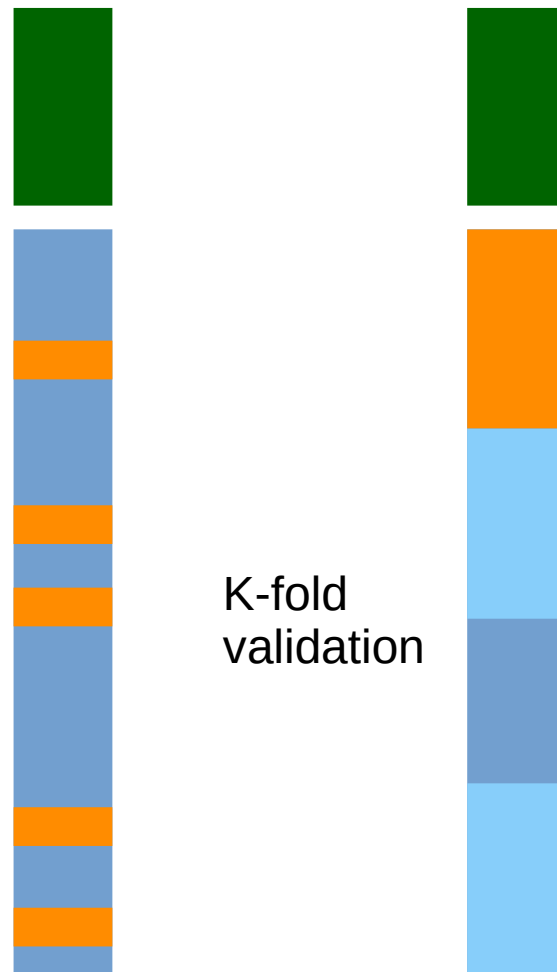
(training + test set = learning set)

“test” set: used to assess the model
during the learning phase
Different for each model instance

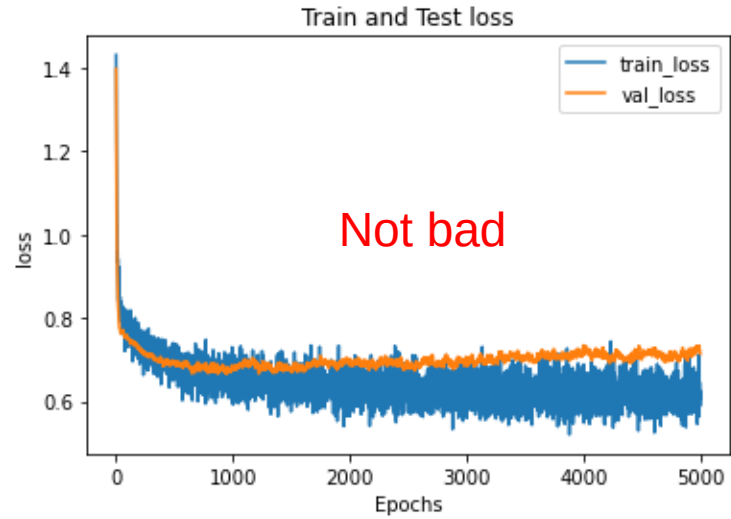
Beware: “validation” and “test” are used the other way
around a lot in deep learning, at the opposite of all other
fields of machine learning, or even life science in general

Random
Test samples

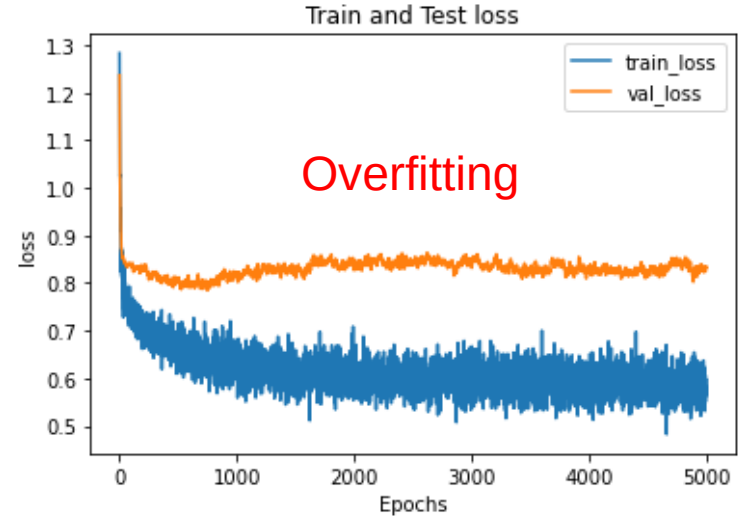
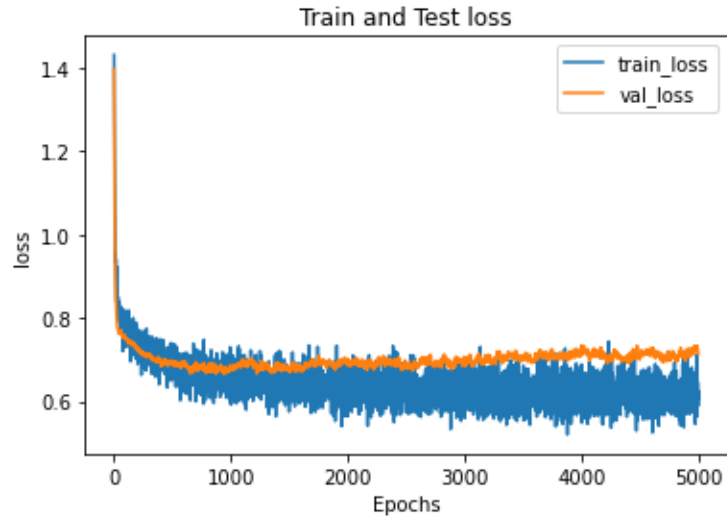
K-fold
validation



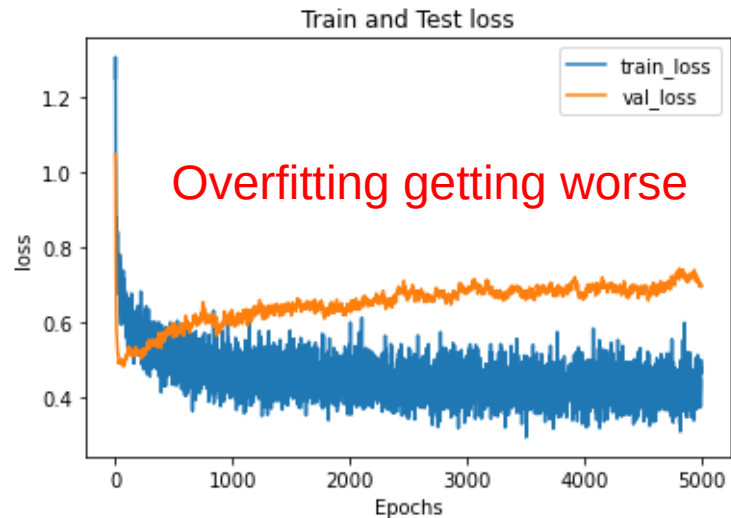
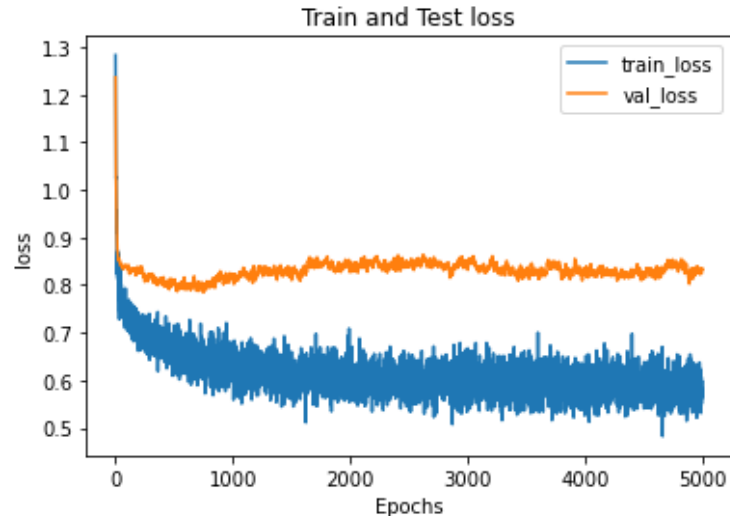
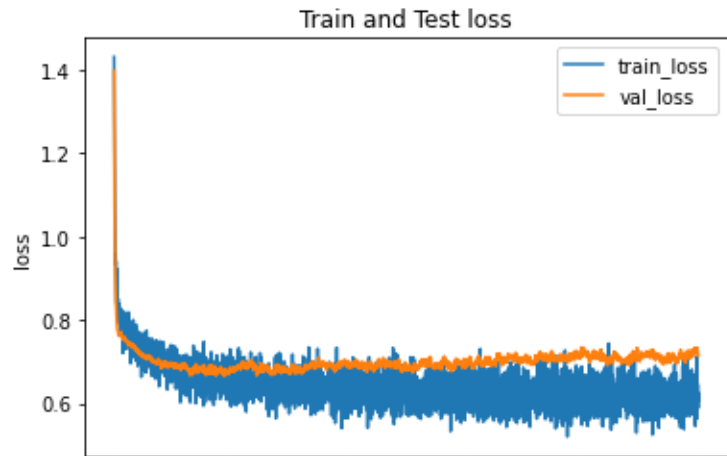
Learning and overfitting



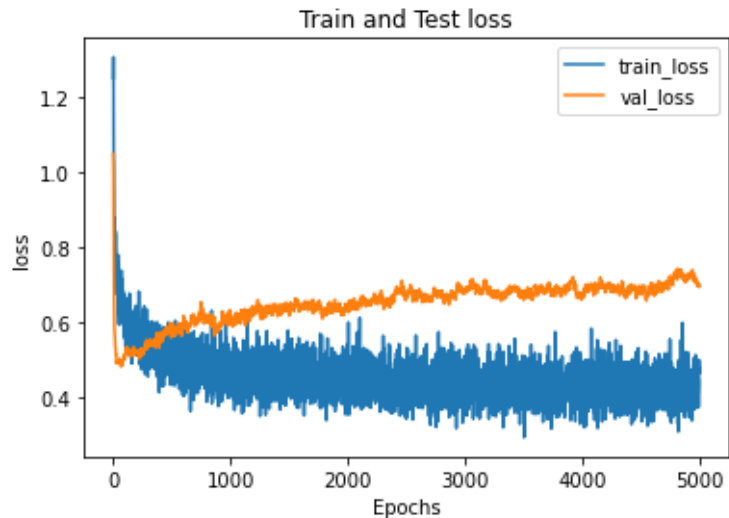
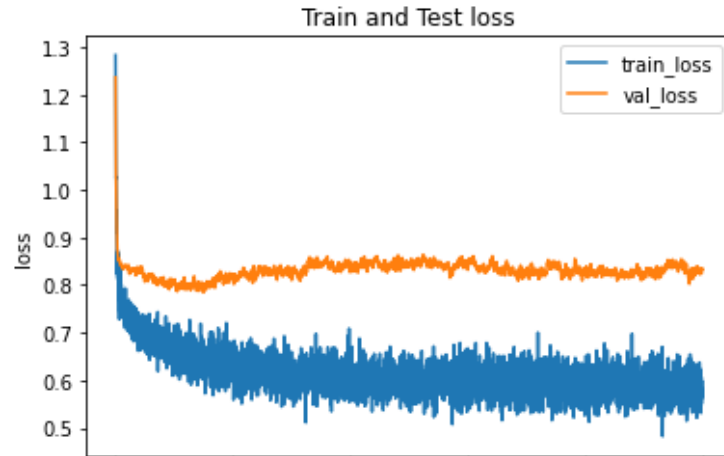
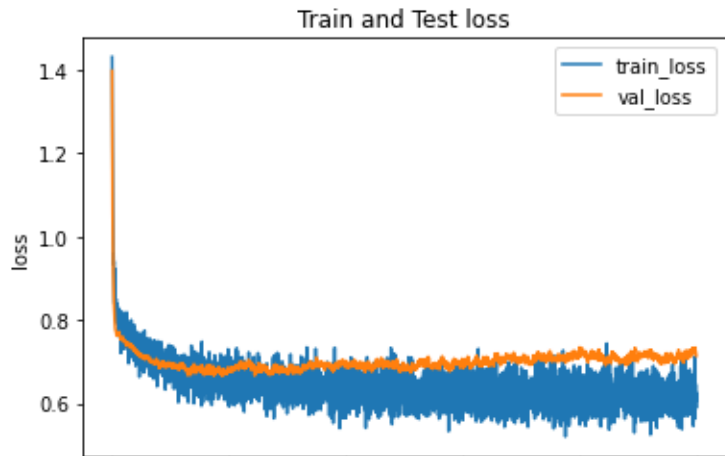
Learning and overfitting



Learning and overfitting



Learning and overfitting



Why a test AND a validation set?

Underperforming on the test set means the model overfitted the training set, the parameters are too specific of the training samples. This overfitting is learned.

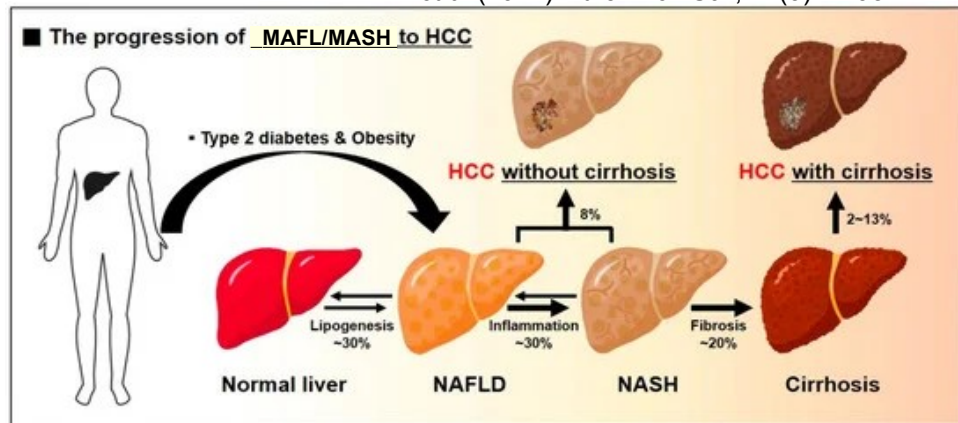
Underperforming on the validation set means the hyperparameters are too specific of the test set as well! When you modifies the structure of the model to avoid overfitting, you actually made the model overfit the entire learning set. YOU biased the model. This overfitting is built-in.

3

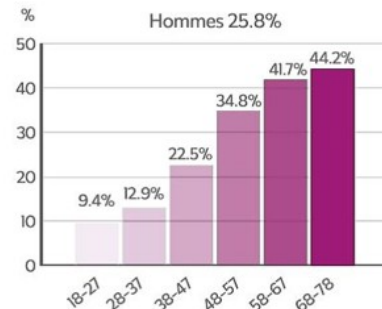
MLPs in action
Multi-omics real-world example

Let's try to recognise a disease severity

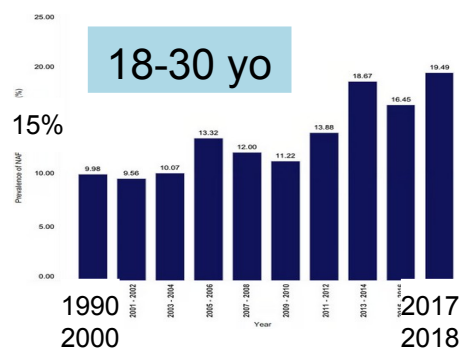
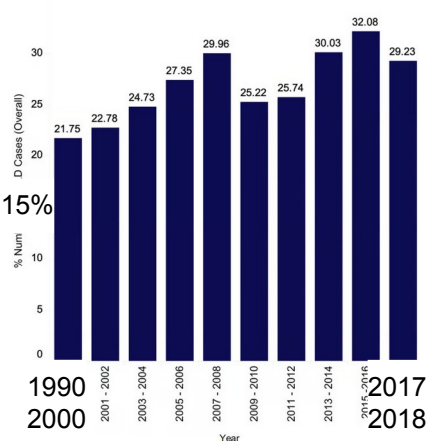
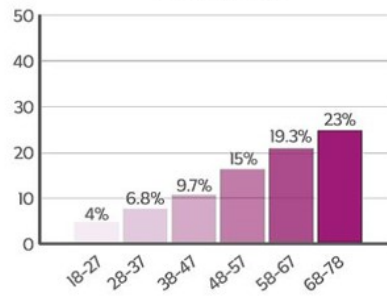
Kim et al (2021) *Int. J. Mol. Sci.*, 22(9): 4495



Prévalence en fonction de l'âge et du sexe

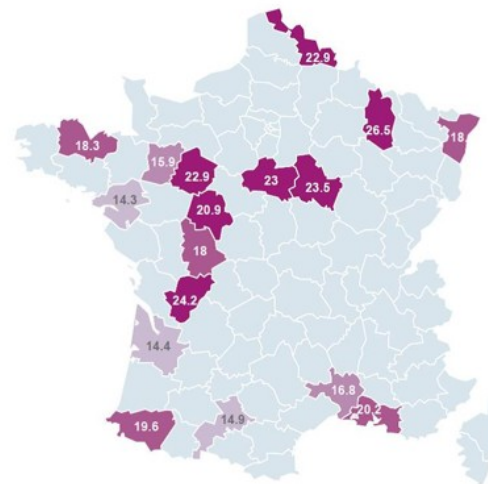


Femmes 11.4%



Kim et al (2022) *Met. Target Organ Damage*, 2: 19

NALFD (now MASLD)

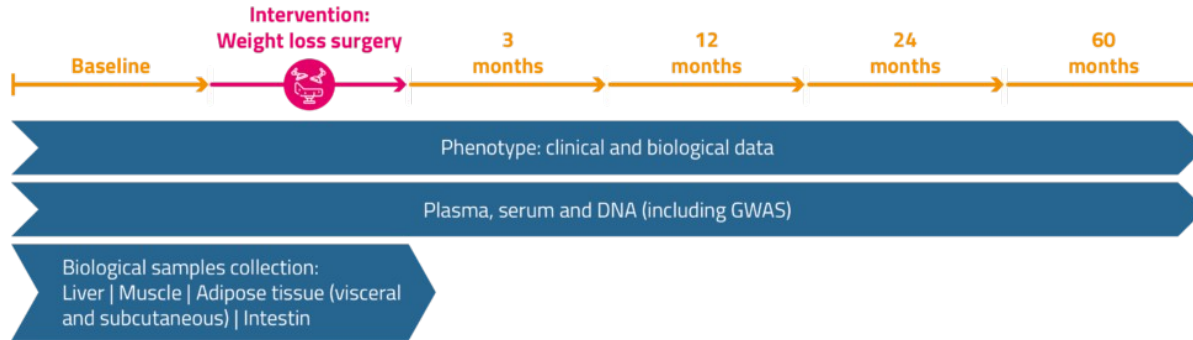


Répartition par régions

NASH (now MASH)

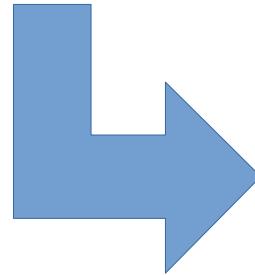
Paris MASH Meeting (11-12 juillet 2019)

Cohort



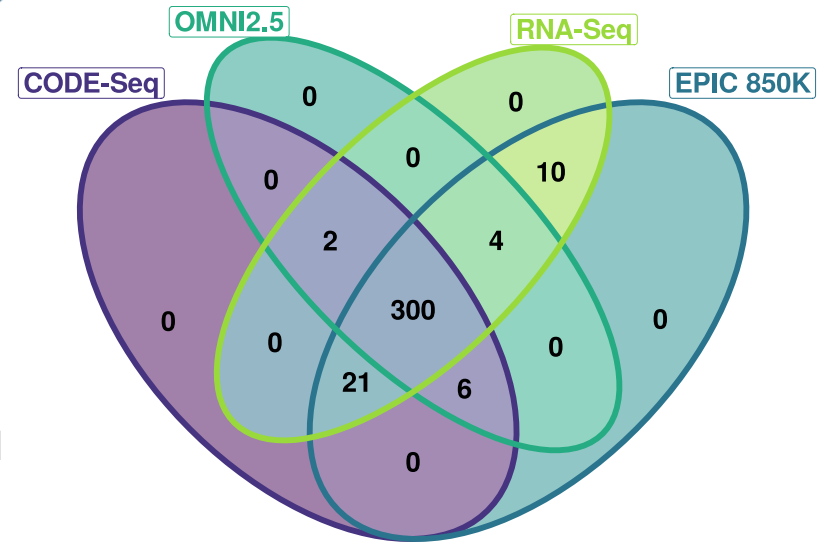
ABOS (Biological Atlas of Severe Obesity)

All subjects had bariatric surgery



PreciNASH project

ABOS subset: Only European ancestry and unrelated individuals



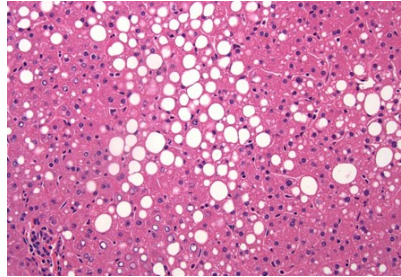
(+66 clinical and personal data
+1076 identified metabolites
in blood and liver)

Subject grouping

Scoring on liver biopsy with the method from Kleiner and Brunt 2005

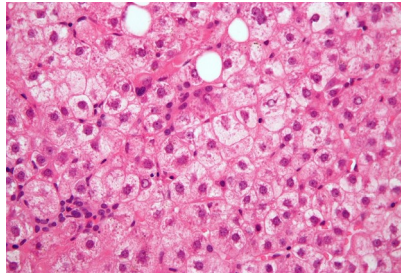
Steatosis

Categorical [0-3] from
quantitative measurement



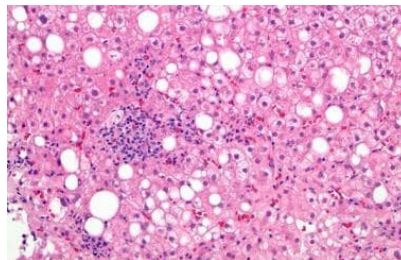
Ballooning

Categorical [0-2]
= {none, some, much}



Inflammation

Categorical [0-3] from
number of foci



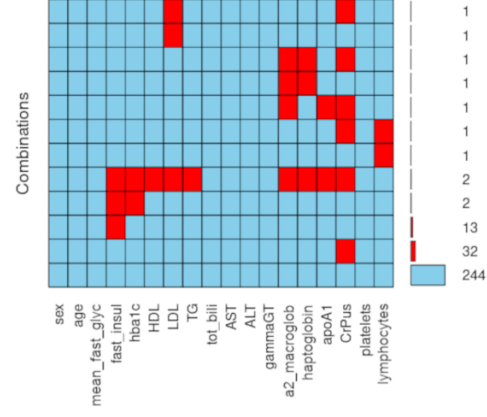
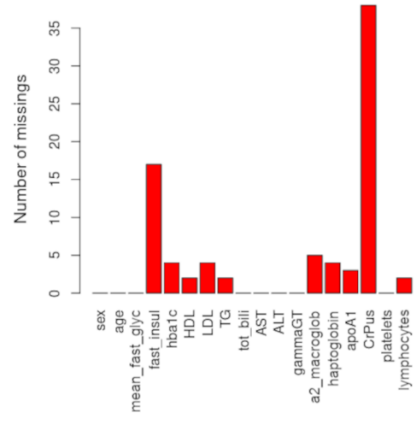
Final score:

Healthy: $S = 0, B = 0, I = 0$ $n = 80$

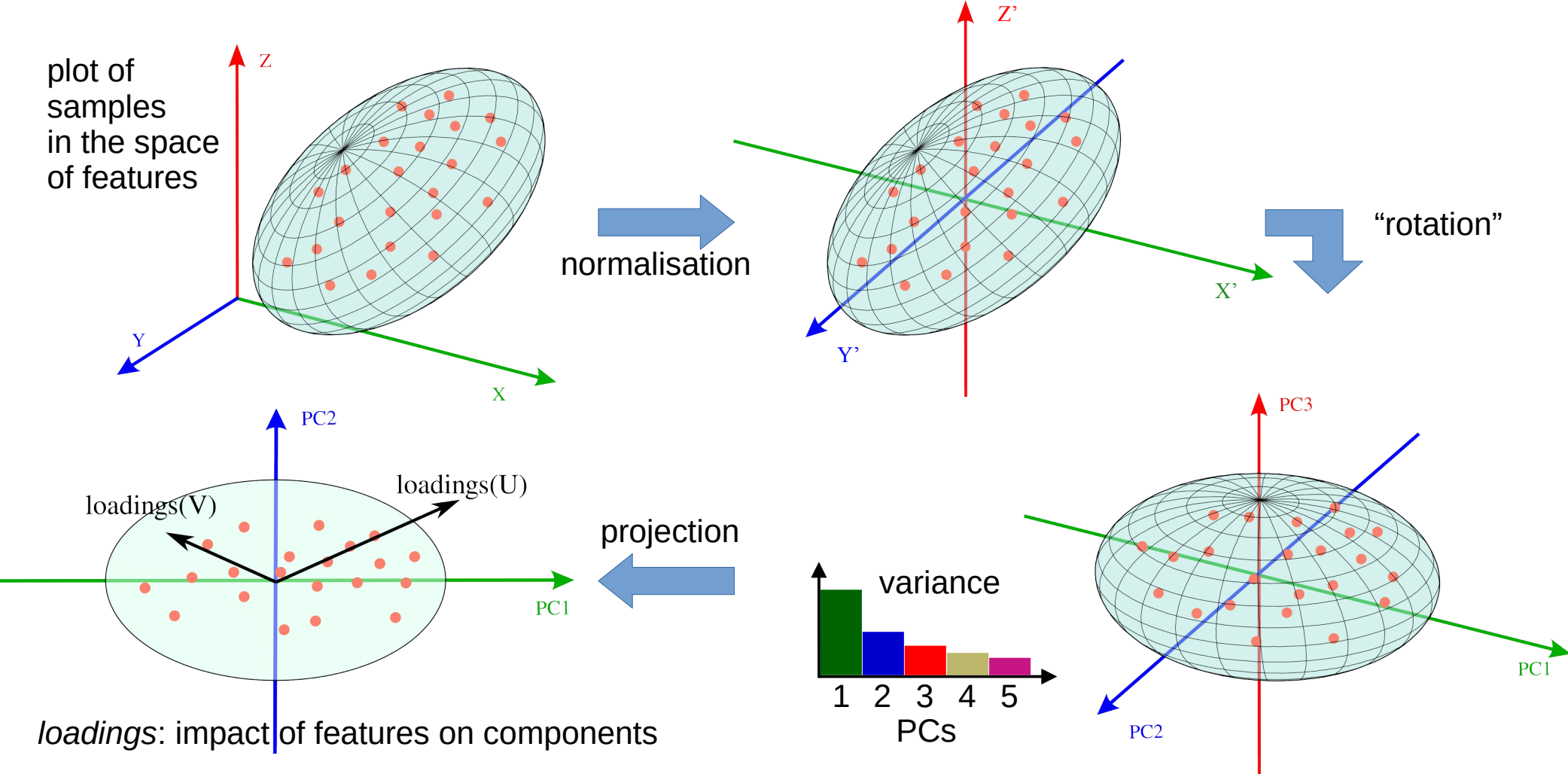
NAFL: $S > 1, B = 0, I \geq 1$ $n = 137$
 $S > 1, B > 1, I = 0$

NASH: $S > 0, B > 0, I > 0$ $n = 83$

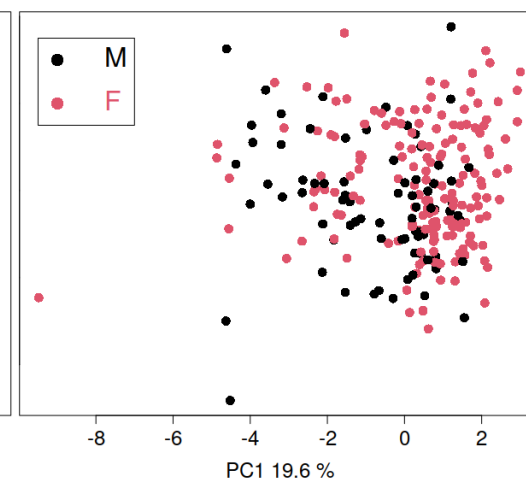
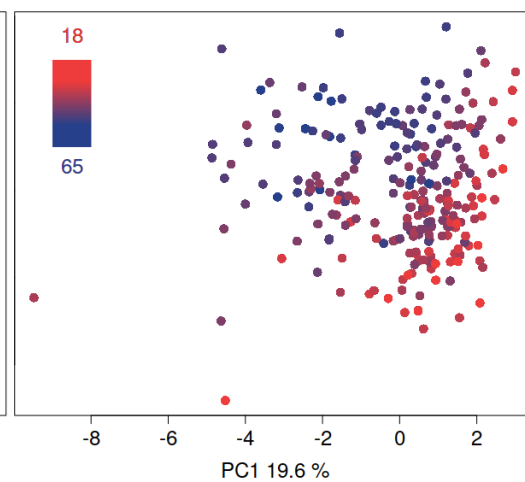
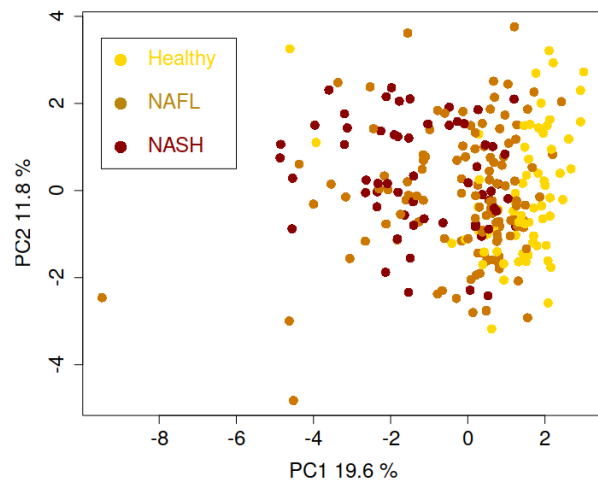
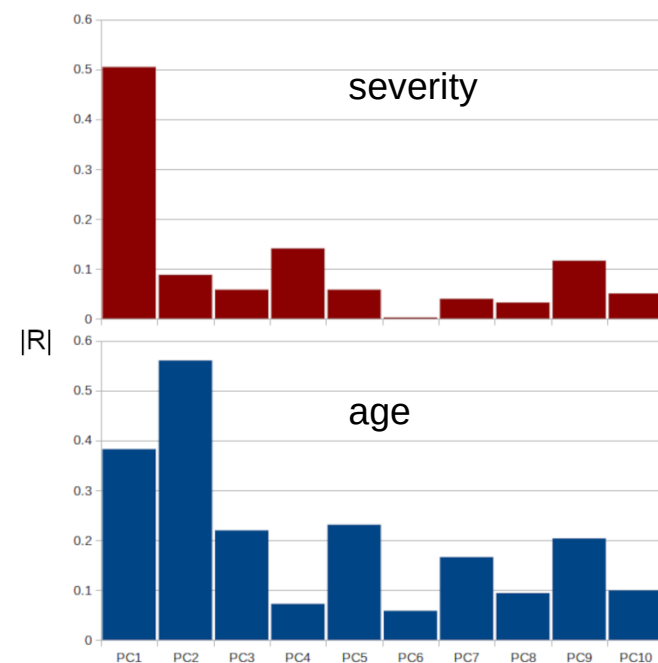
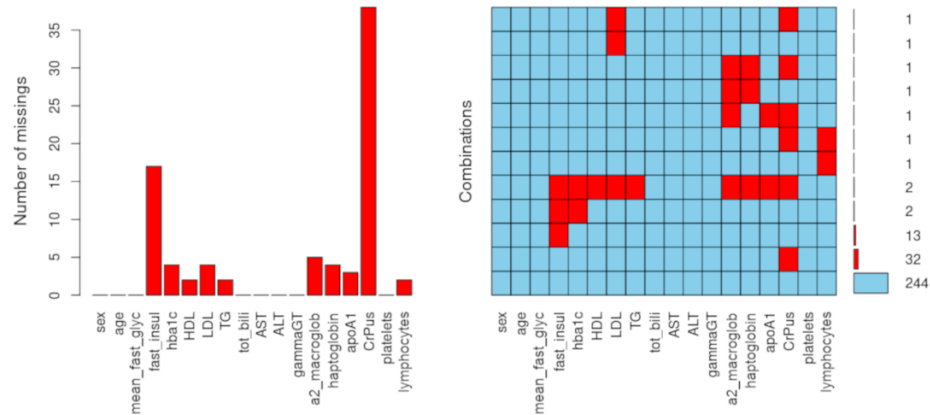
Clinical data



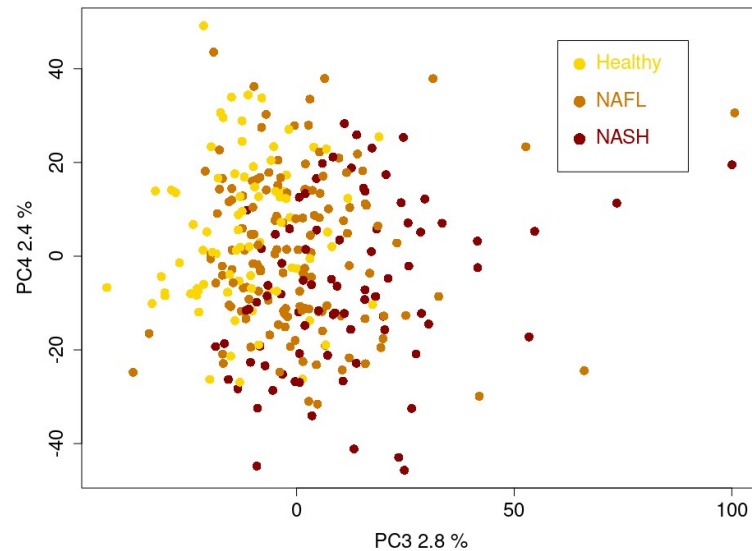
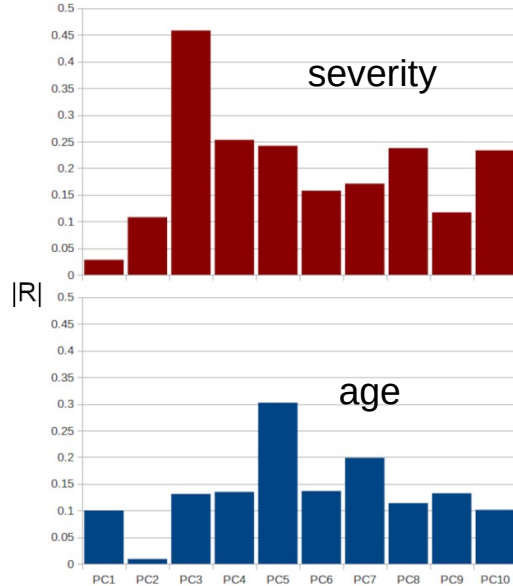
Principal component analysis (PCA)



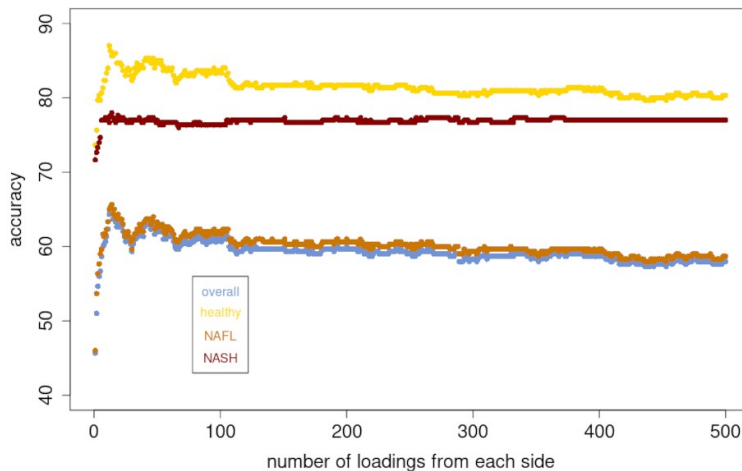
Clinical data



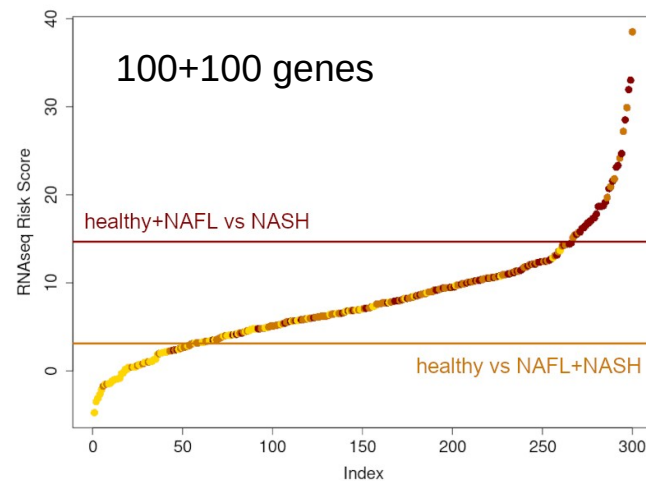
RNAseq

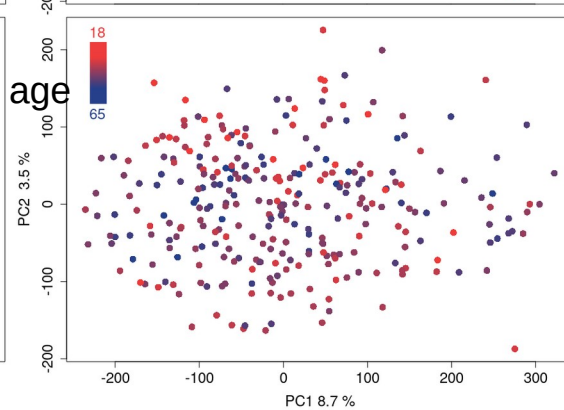
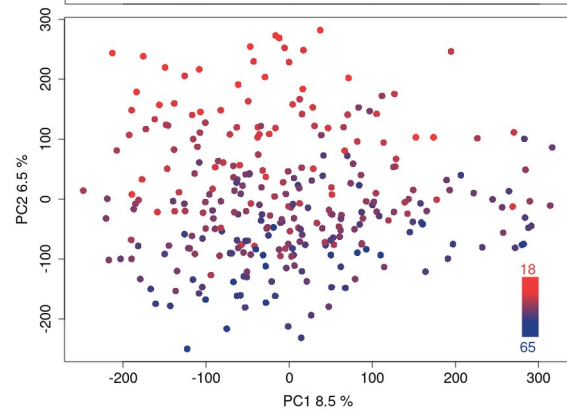
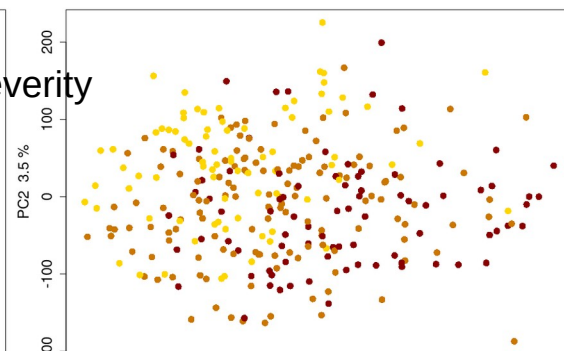
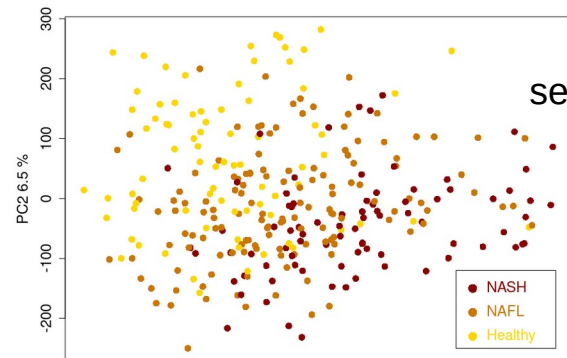
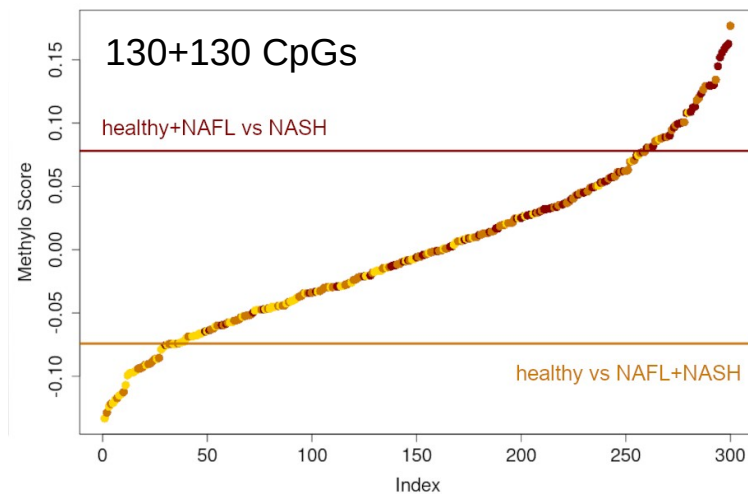
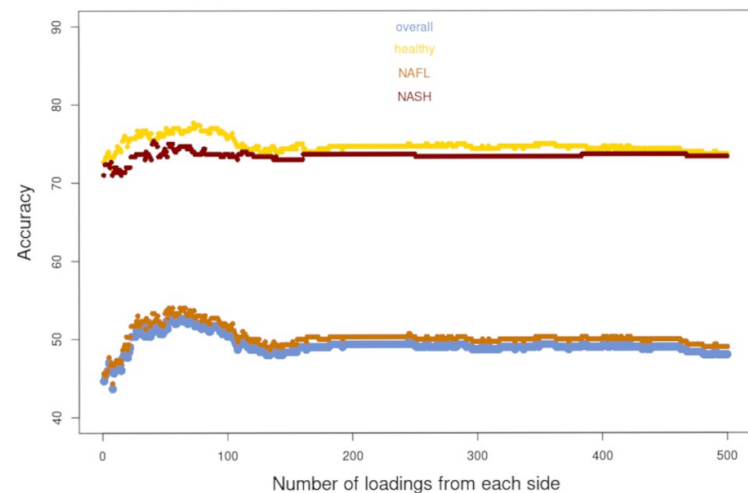
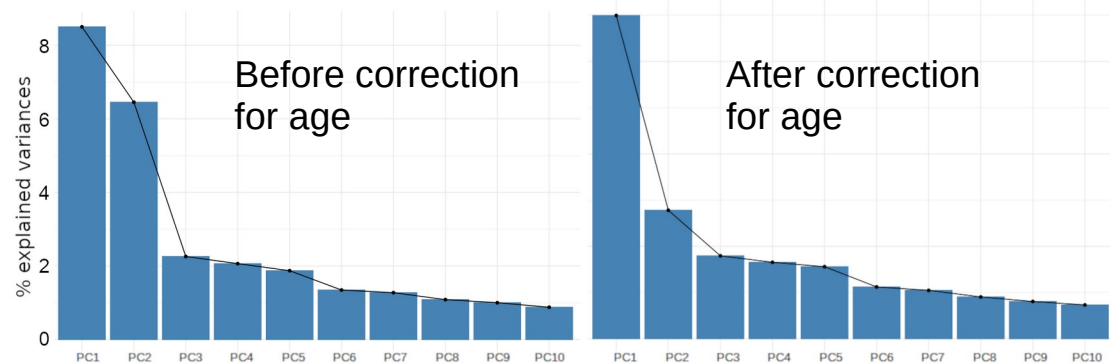


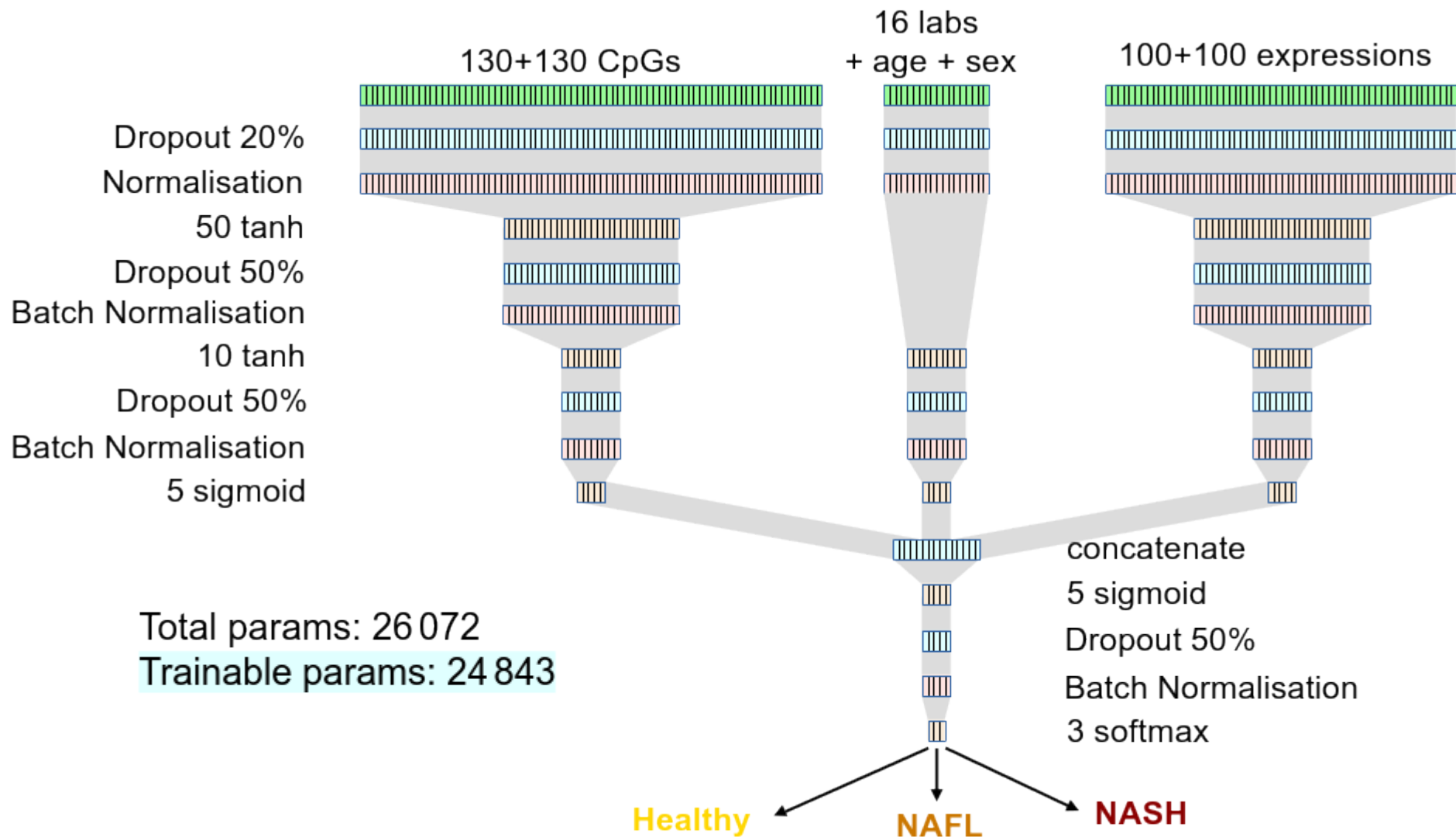
Score based on
gene expression
and gene “loadings”
(impact of a gene
on a given principal
component)



Logistic regression
to find the thresholds
best separating the
severity groups

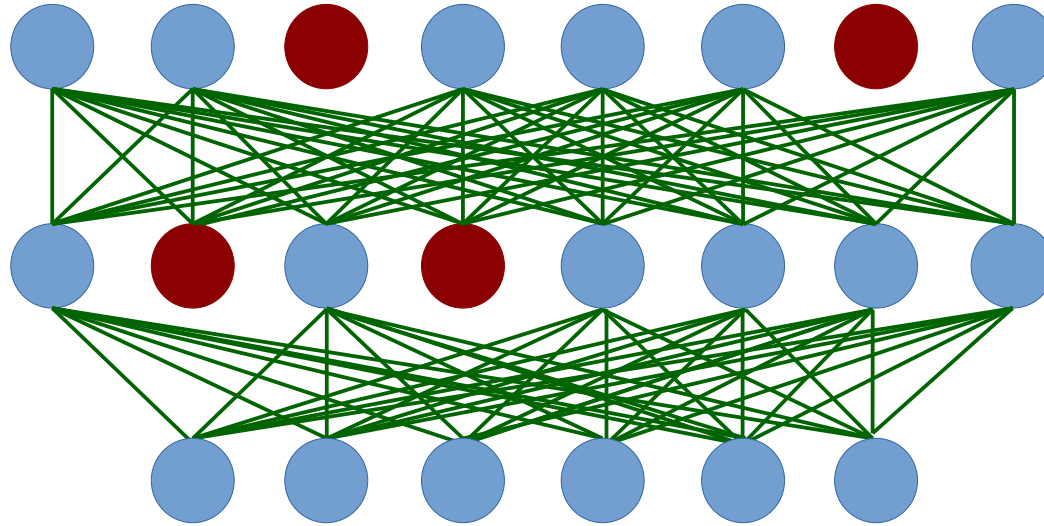






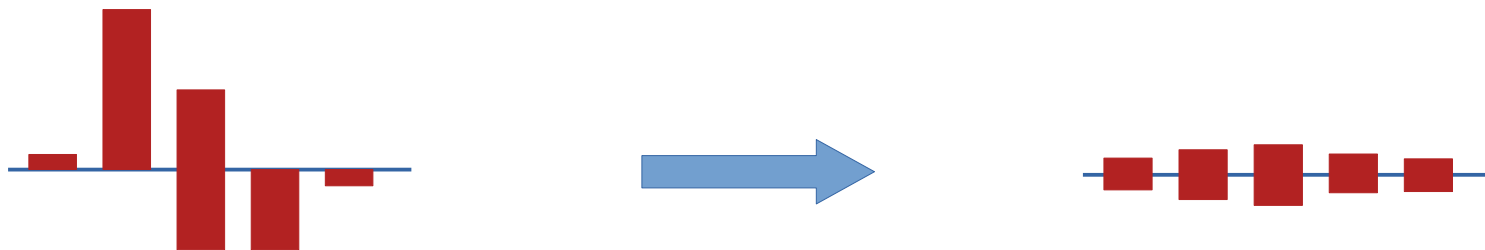
Dropout

- 1) Purpose: avoiding overfitting
- 2) Disable some connections at random (set the weights at 0)



Srivastava et al (2014) Dropout: A Simple Way to Prevent Neural Networks from Overfitting. JMLR 15(56):1929–1958

Normalisation



- 1) Normalisation during data processing
- 2) Normalisation before training: on the whole dataset
- 3) Normalisation during training: after dropout
- 4) Batch normalisation: normalisation on the current batch (not on the entire dataset)
- 5) Layer normalisation: normalisation of the input of a layer

Ba, Kiros, Hinton (2016) Layer Normalization. arXiv:1607.06450v1

Ioffe and Szegedy (2015) Batch Normalization: Accelerating Deep Network Training by Reducing Internal Covariate Shift. Proc 32nd Intl Conf Machine Learning, Lille, France volume 37

Evaluating a model's performance

		Predicted	
		Positive	Negative
Actual	Positive	True positive	False negative
	Negative	False positive	True negative

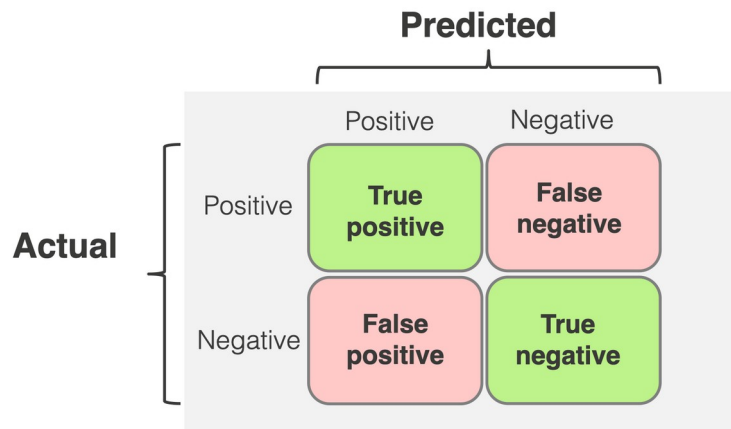
$$\text{Accuracy} = (TP+TN)/(TP+FN+TN+FP)$$

$$\text{Precision} = TP/(TP+FP)$$

$$\text{Sensitivity (true positive rate)} = TP/(TP+FN)$$

$$\text{Specificity (true negative rate)} = TN/(TN+FP)$$

Evaluating a model's performance



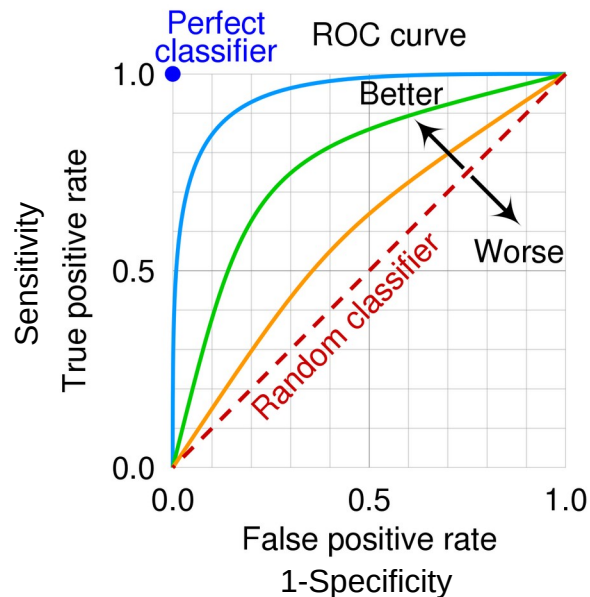
$$\text{Accuracy} = (TP+TN)/(TP+FN+TN+FP)$$

$$\text{Precision} = TP/(TP+FP)$$

$$\text{Sensitivity (true positive rate)} = TP/(TP+FN)$$

$$\text{Specificity (true negative rate)} = TN/(TN+FP)$$

Receiver operating characteristic (ROC) curve



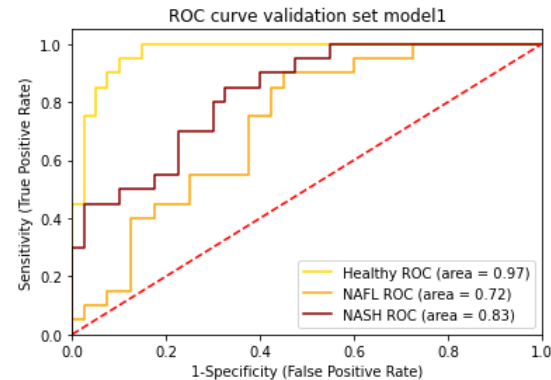
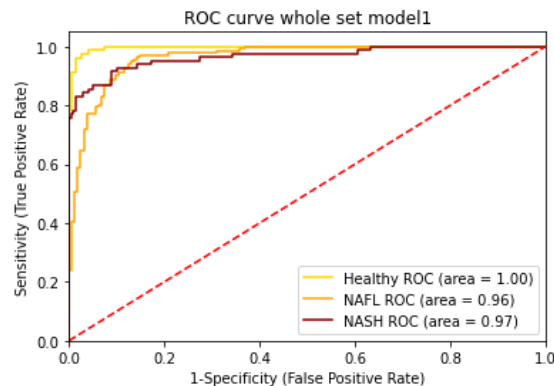
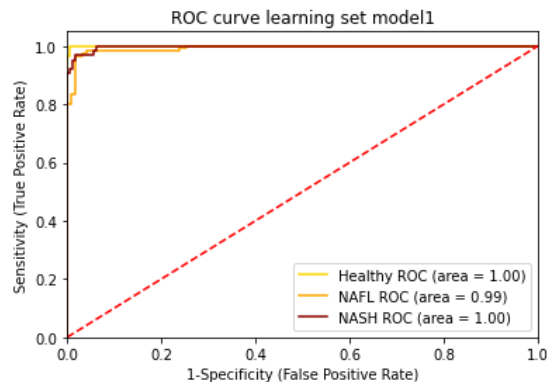
Different accuracies on different datasets

Learning set

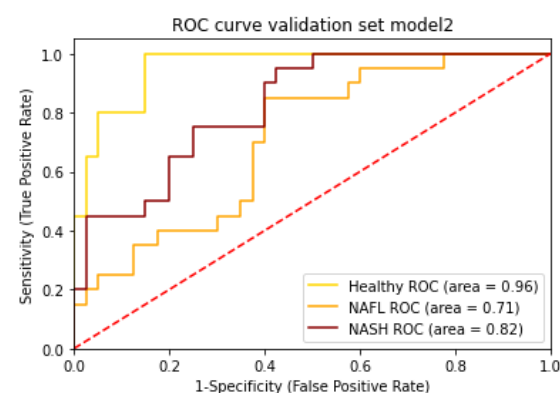
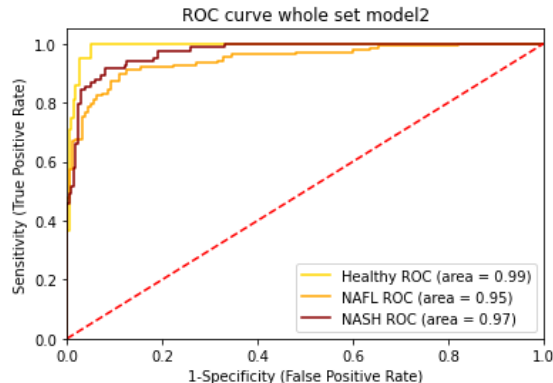
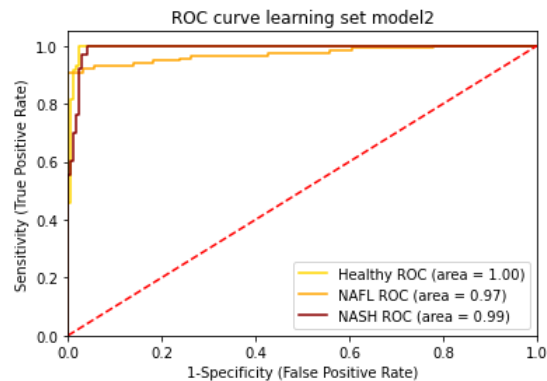
Whole set

Validation set

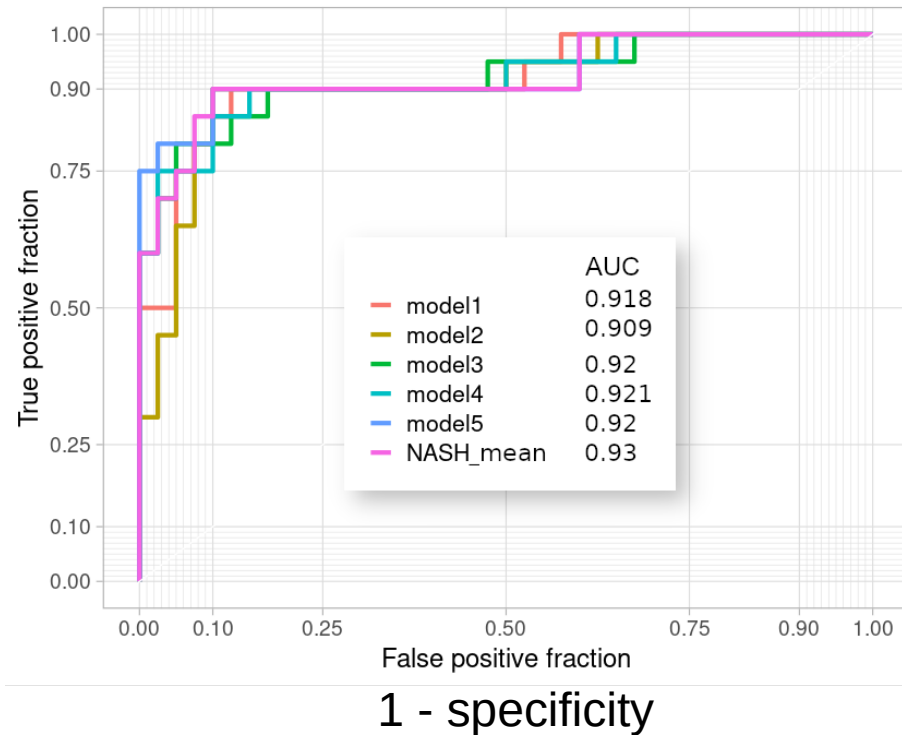
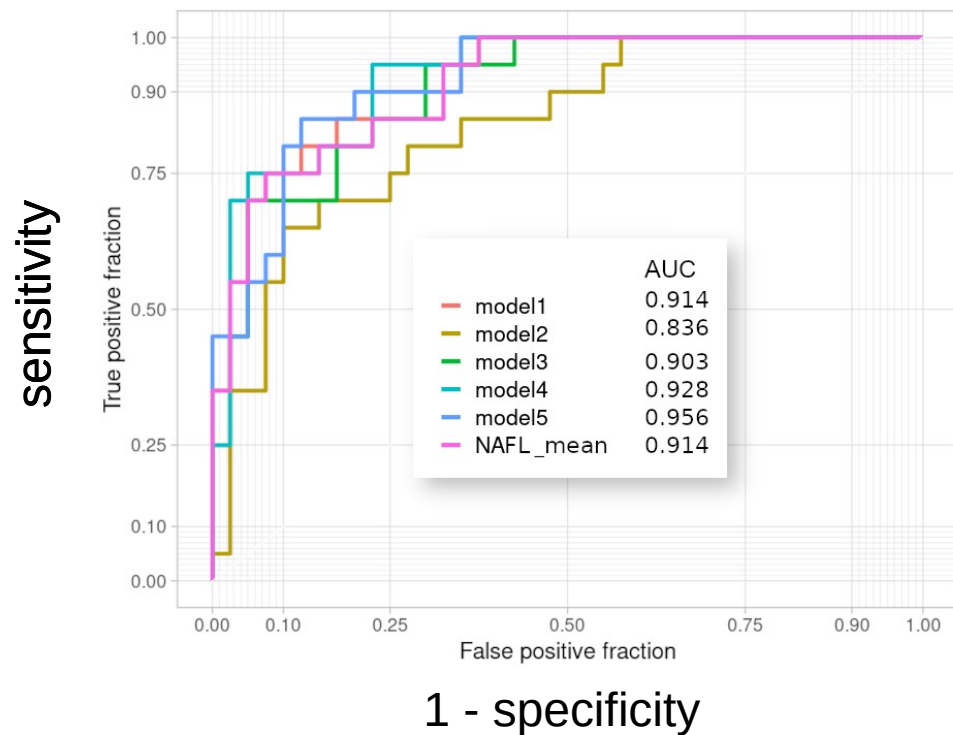
Model 1

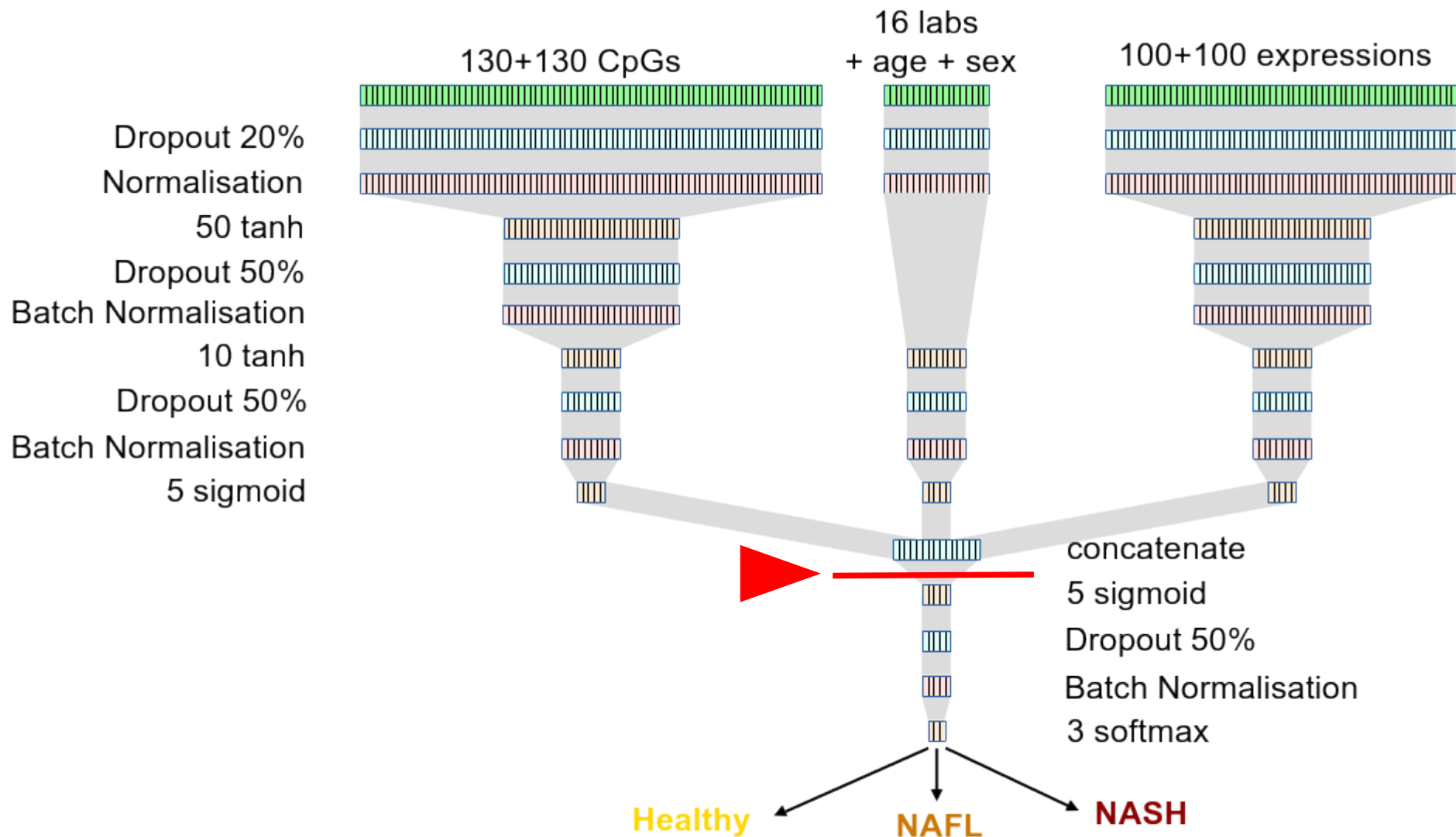


Model 2



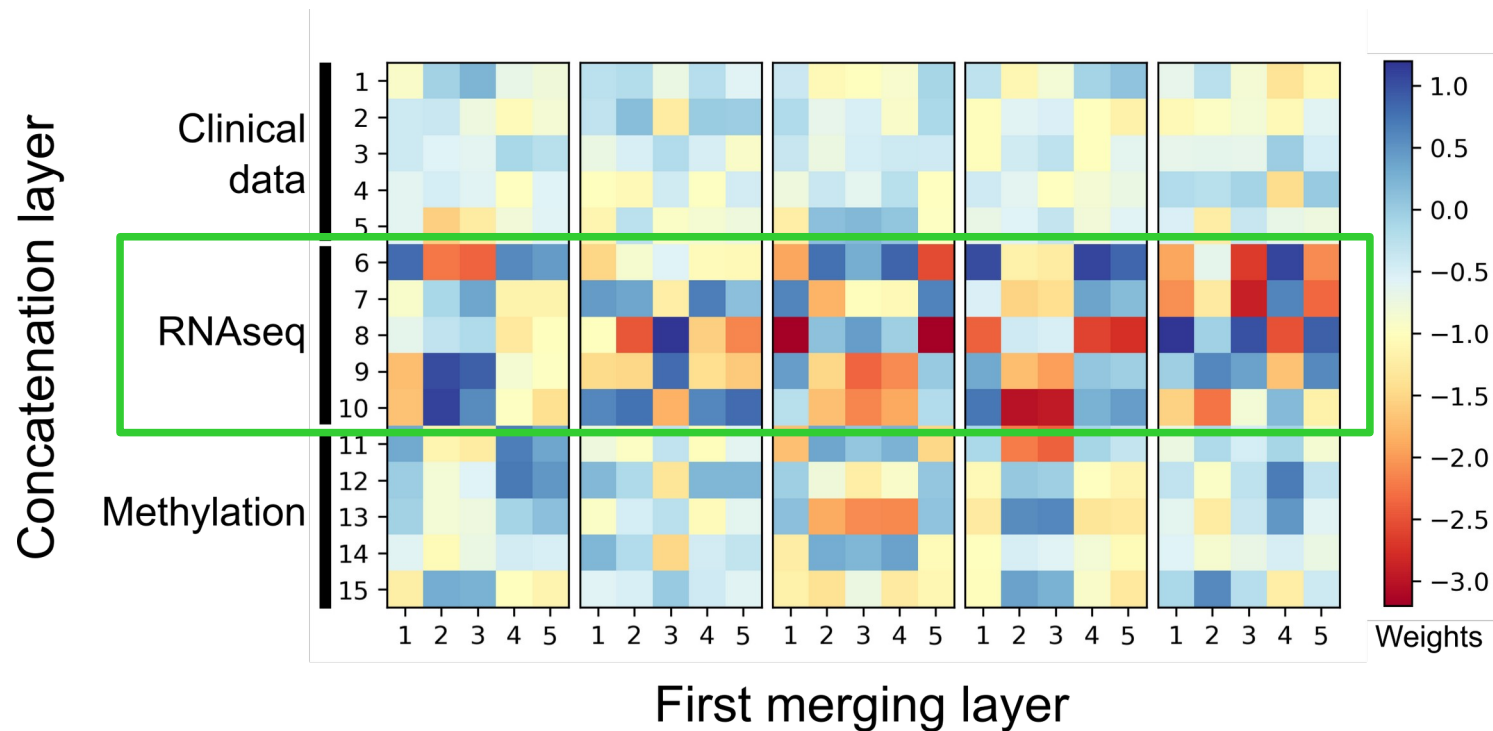
How good is the model to distinguish NAFL and NASH?

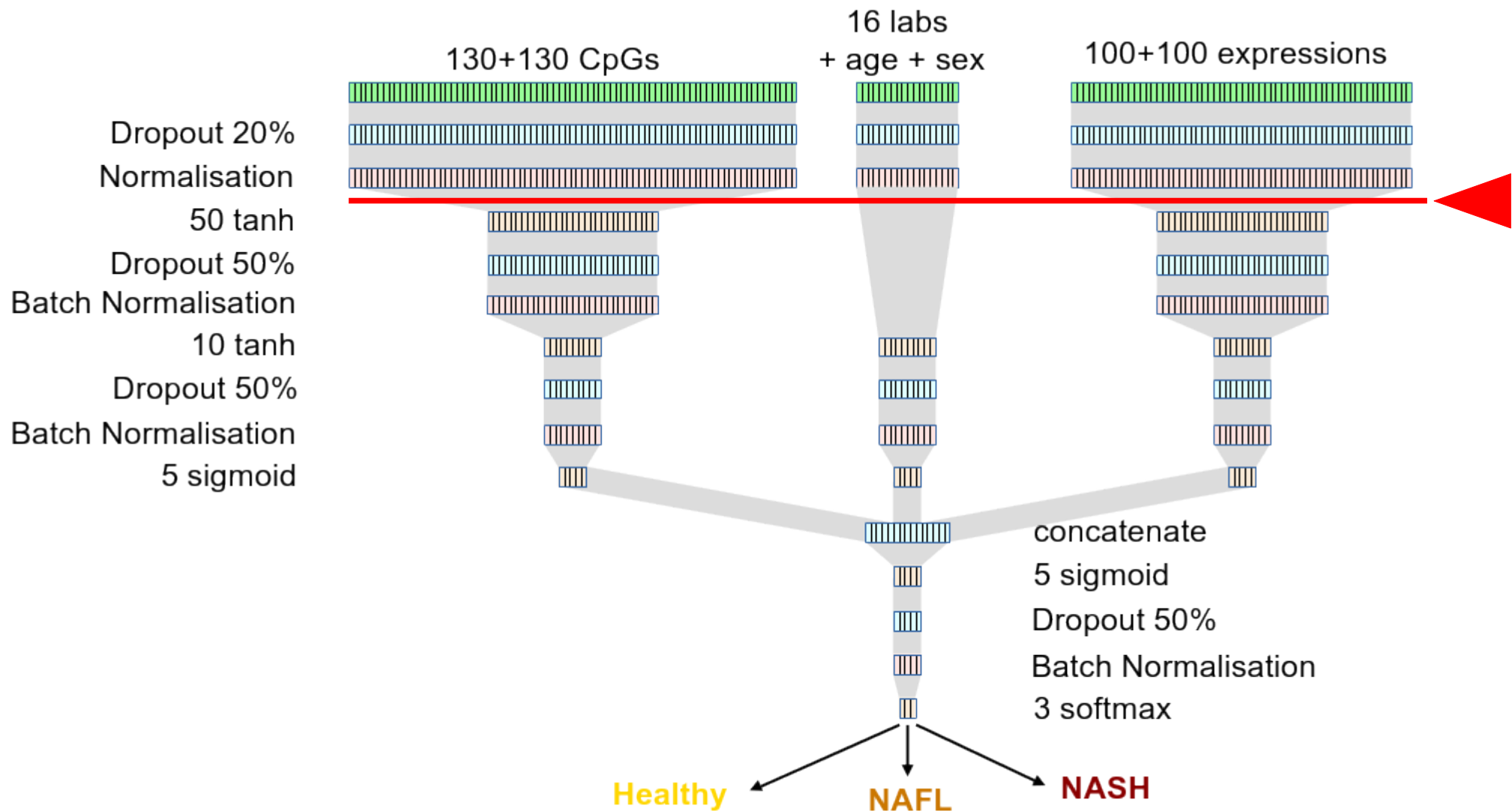




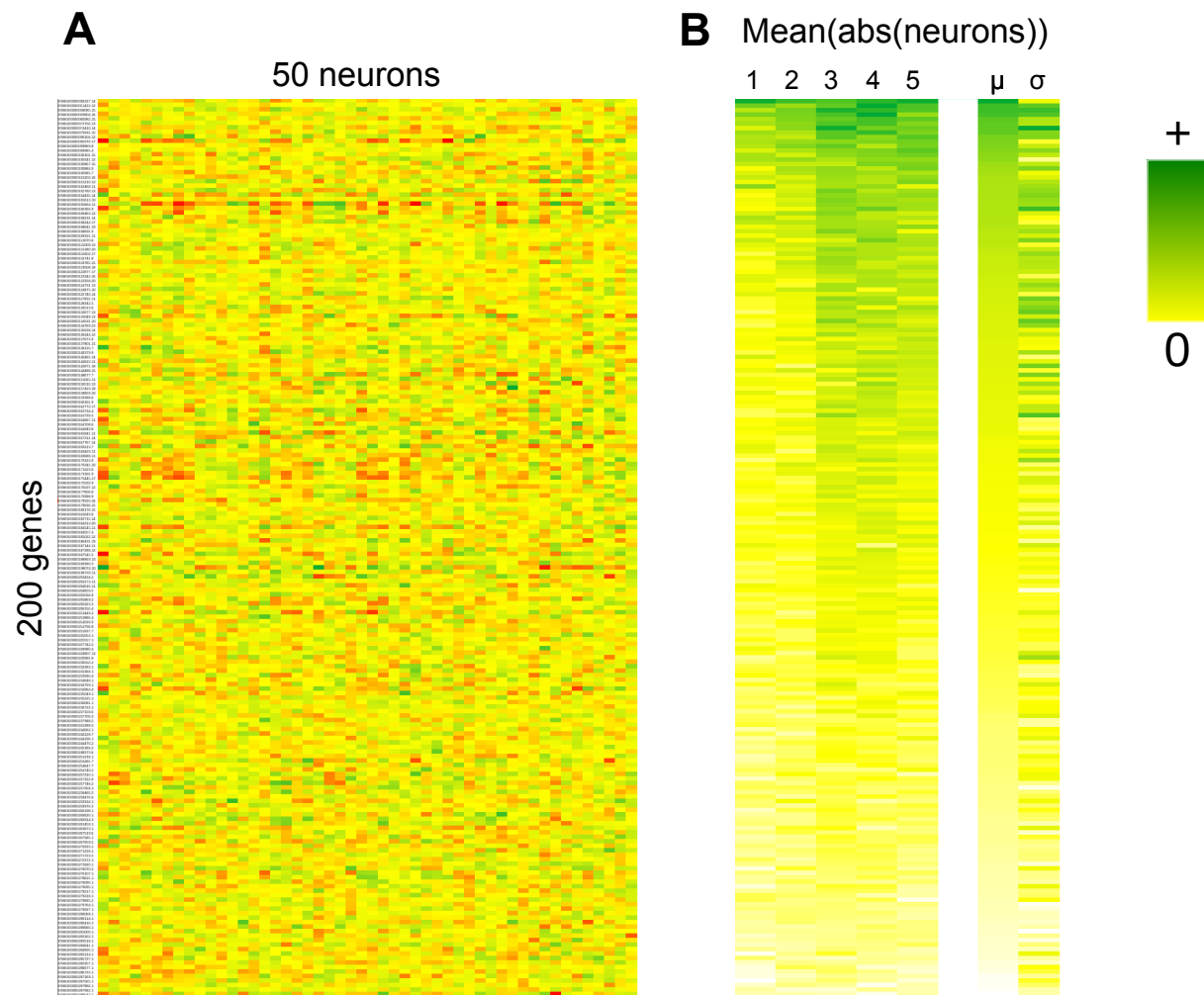
Peeping into the black box

The RNAseq module has the most impact on output

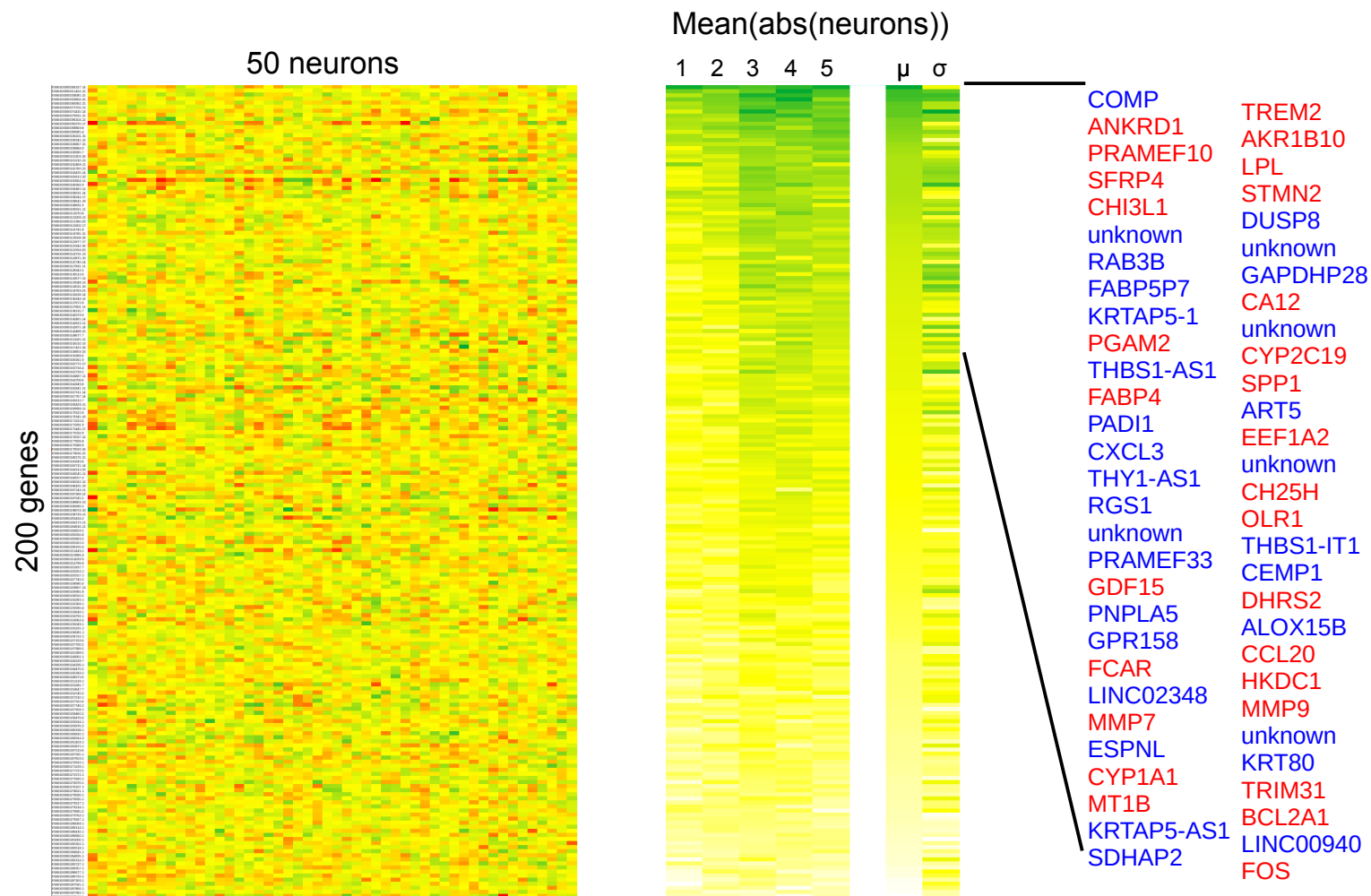




Independently trained models learn from the same genes



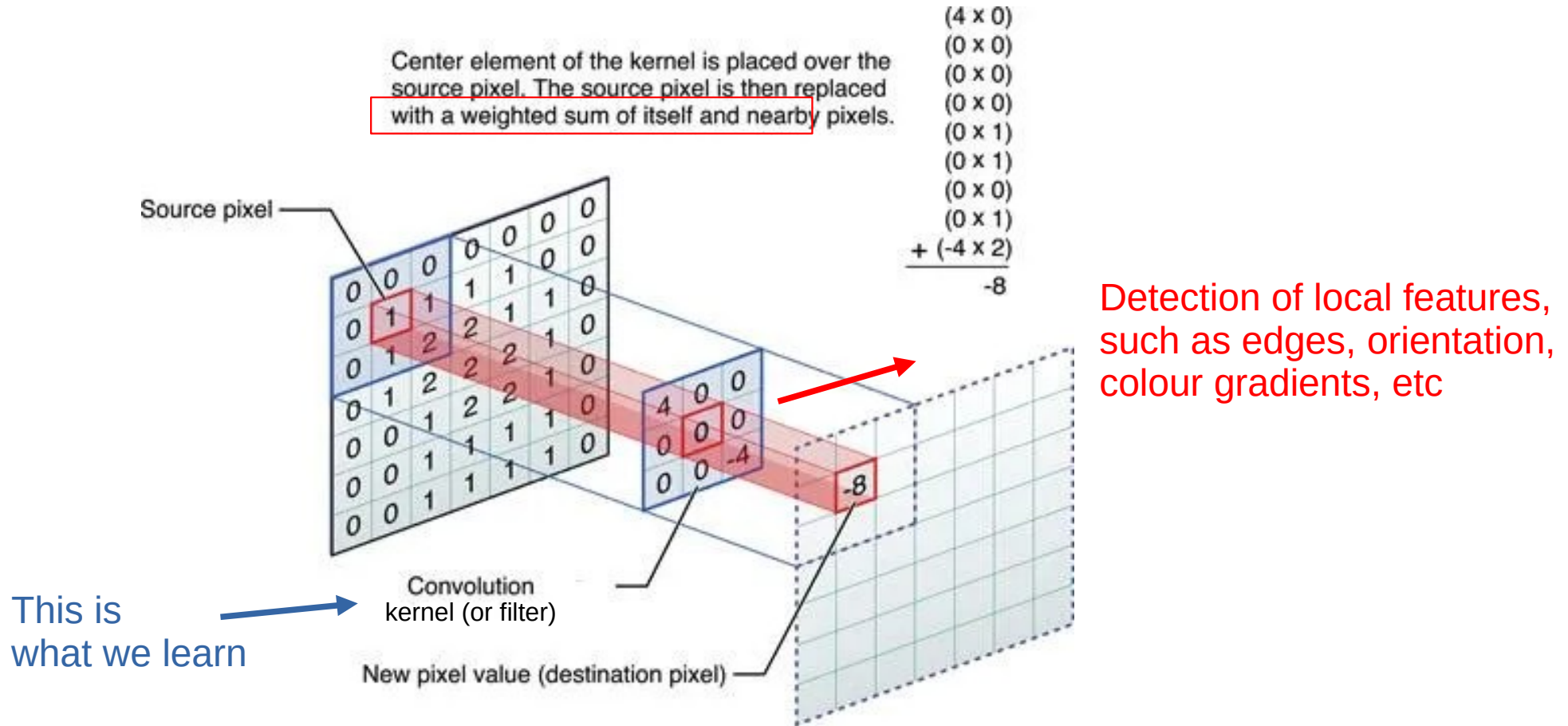
Known and new genes in MASLD severity



4

Convolutional Neural Networks (CNN)

We can detect local features by linking neighbouring inputs



Example of feature detection: vertical edges

Sobel Kernel
(3 x 3)

1	0	-1
2	0	-2
1	0	-1

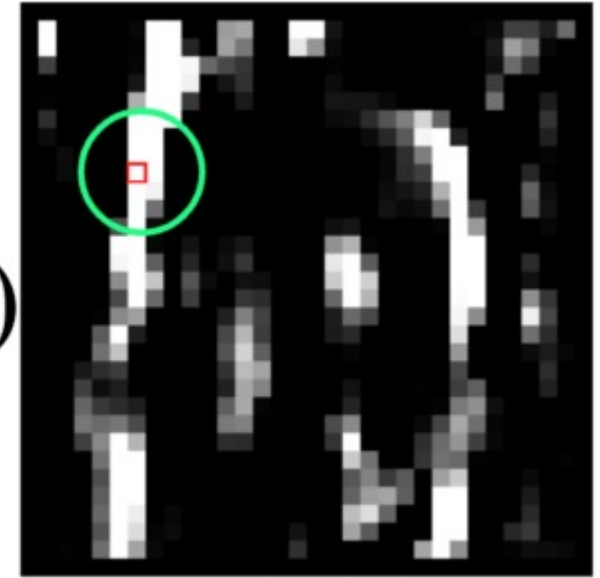


input image

$$\begin{pmatrix} 245 \times 1 + 171 \times 0 + 32 \times -1 \\ + 250 \times 2 + 158 \times 0 + 36 \times -2 \\ + 237 \times 1 + 181 \times 0 + 53 \times -1 \end{pmatrix} = 825$$

kernel:

left sobel

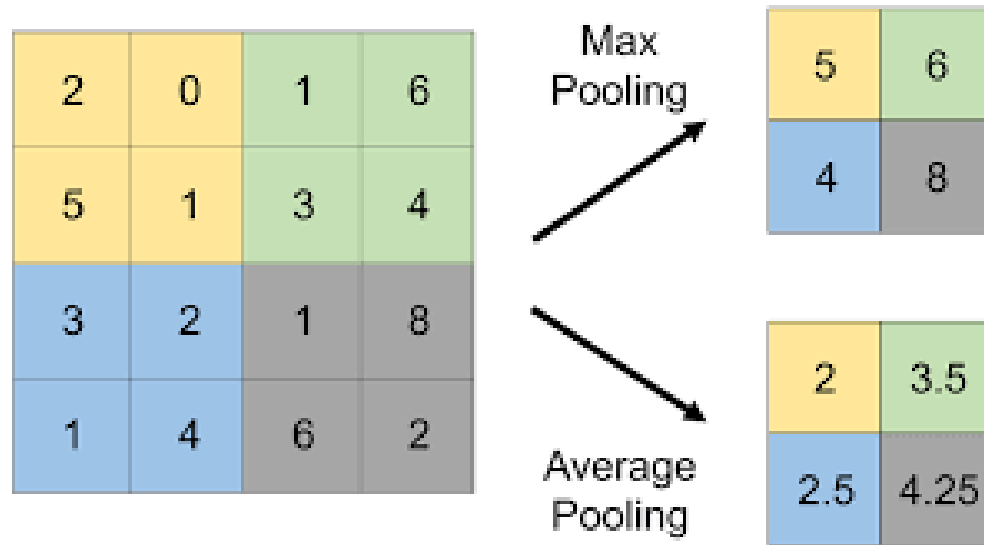


output image

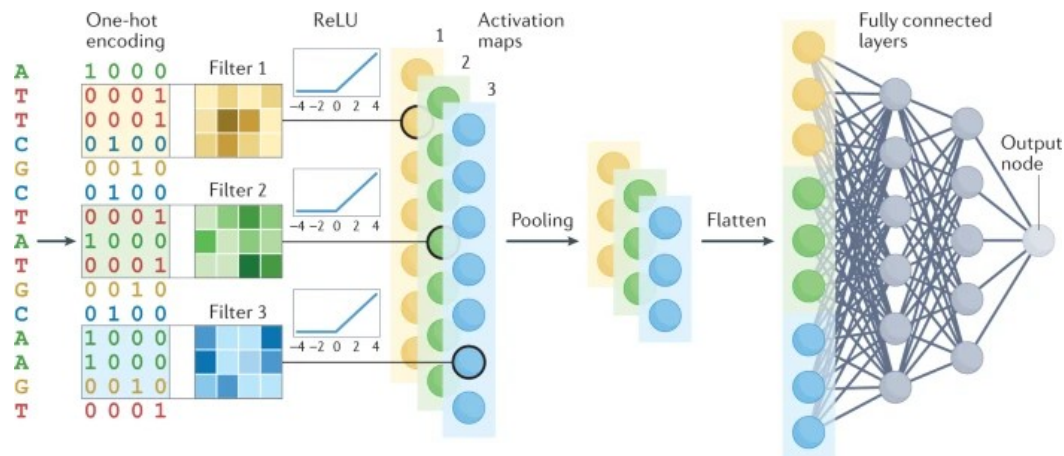
<https://setosa.io/ev/image-kernels/>

NB: here, we provide the kernel.
In CNNs, the kernel is learned

Downsampling: Max/average pooling



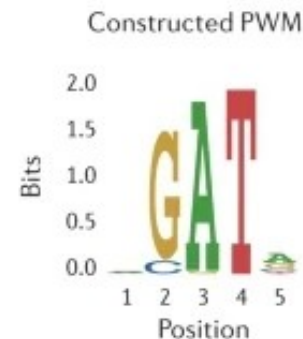
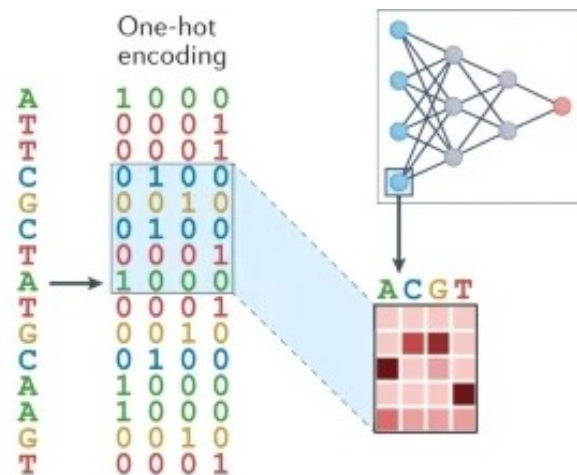
Everything is an image: DNA sequences to images



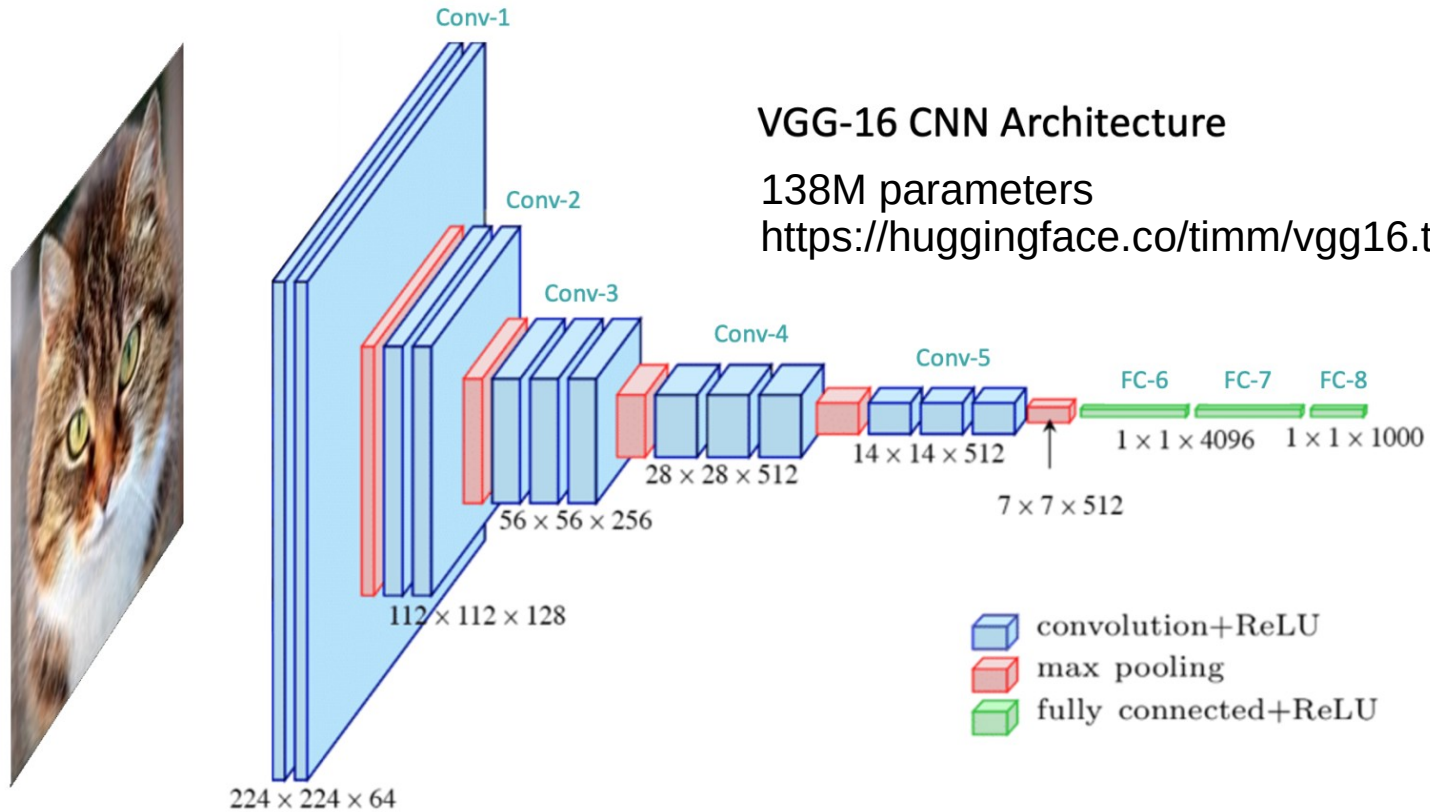
Obtaining genetics insights from deep learning via explainable artificial intelligence

[Gherman Novakovsky](#), [Nick Dexter](#), [Maxwell W. Libbrecht](#) , [Wyeth W. Wasserman](#)  & [Sara Mostafavi](#) 

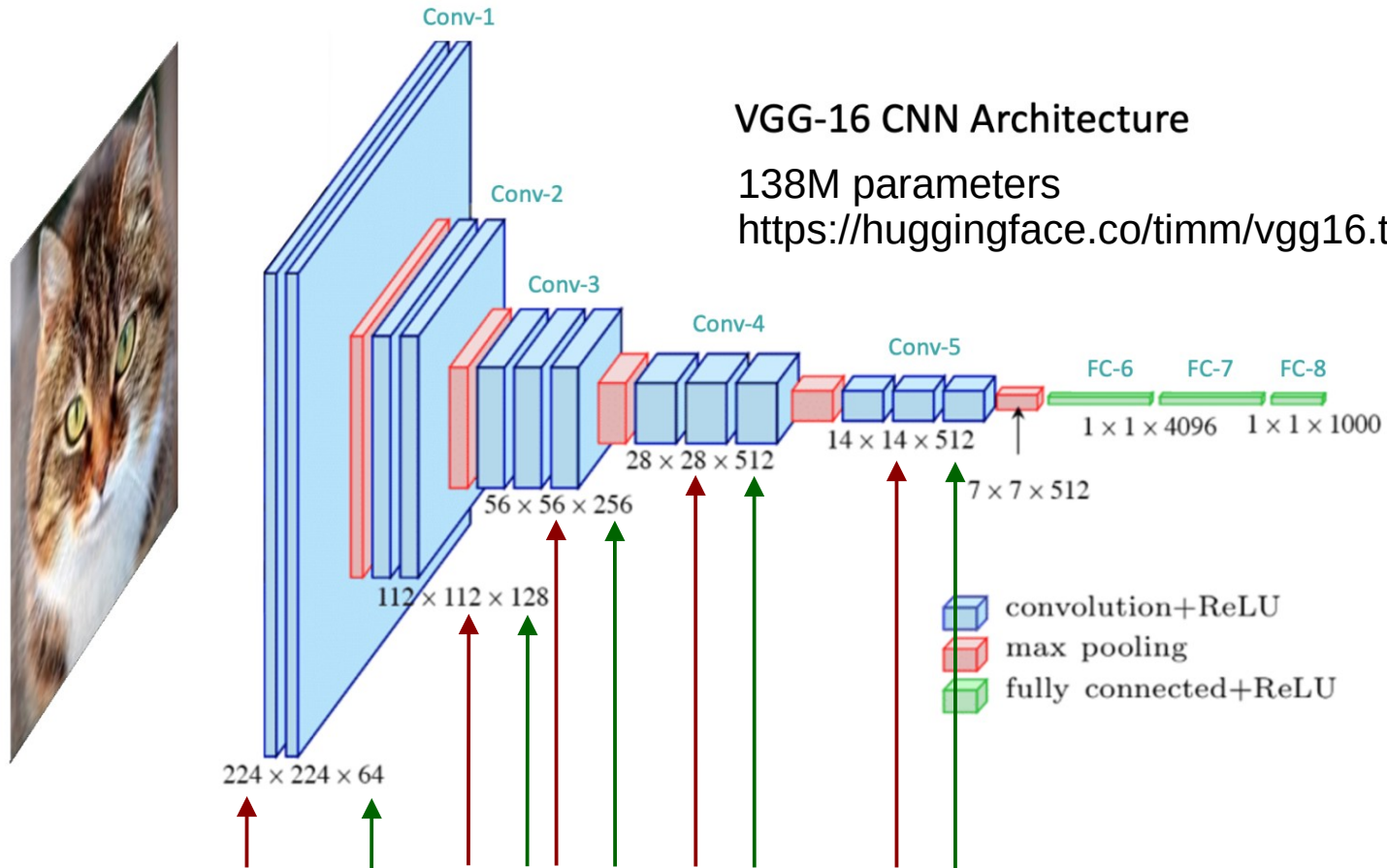
Nature Reviews Genetics 24, 125–137 (2023) | [Cite this article](#)



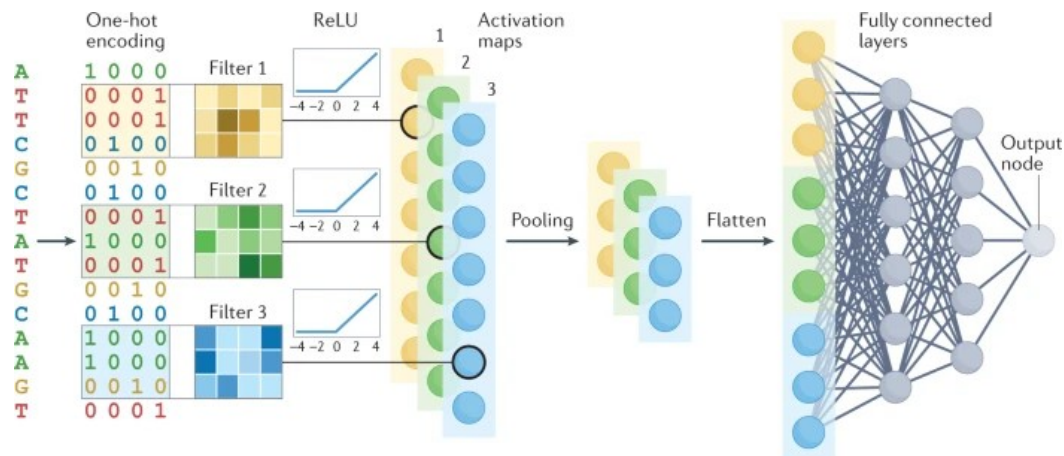
Real example: VGG



Decreasing size, increasing feature number



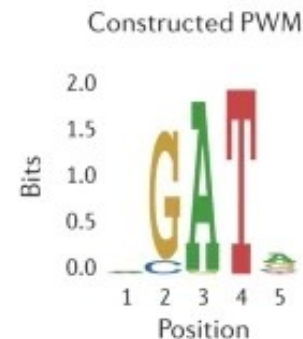
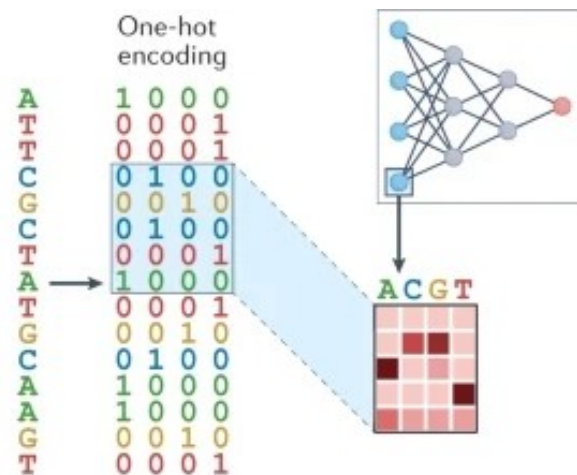
Everything is an image: DNA sequences to images



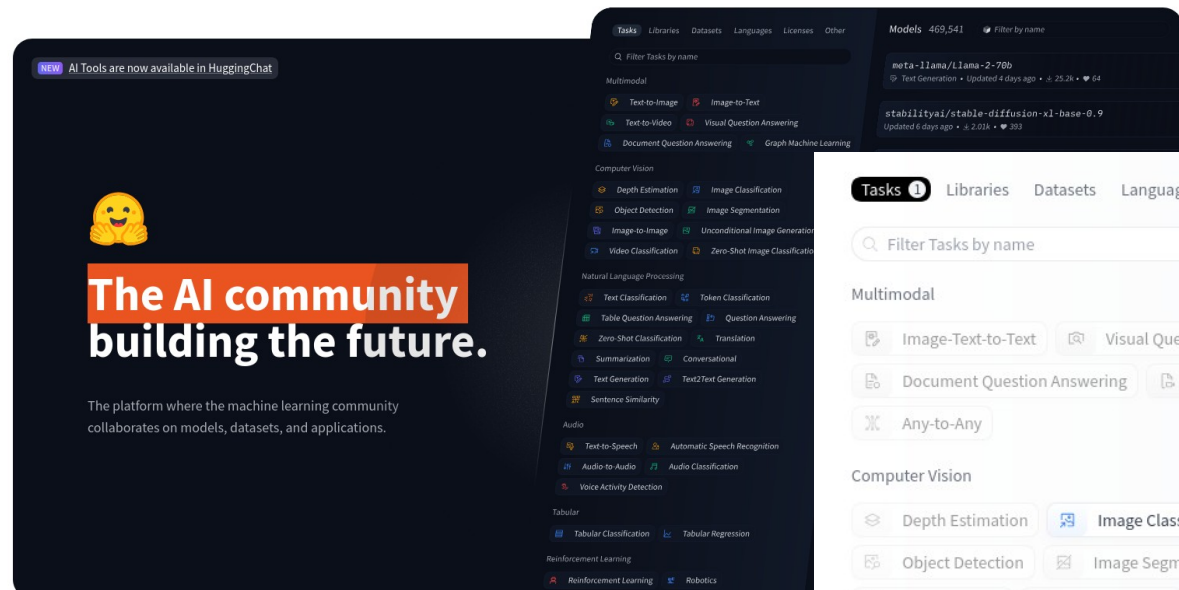
Obtaining genetics insights from deep learning via explainable artificial intelligence

[Gherman Novakovsky](#), [Nick Dexter](#), [Maxwell W. Libbrecht](#) , [Wyeth W. Wasserman](#)  & [Sara Mostafavi](#) 

Nature Reviews Genetics 24, 125–137 (2023) | [Cite this article](#)



Where to find models: Hugging face



Trending on 🤖 this week

Models Spaces

openai/whisper-large-v3-turbo
Updated 4 days ago • 62.3k • 698

nvidia/NVLM-D-72B
Updated 4 days ago • 7.69k • 491

black-forest-labs/FLUX.1-dev
Updated Aug 16 • 1.14M • 5,22k

ostris/OpenFLUX.1
Updated 4 days ago • 7.97k • 359

meta-llama/Llama-3.2-11B-Vision-Instruct
Updated 8 days ago • 427k • 563

Open NotebookLM 640

Whisper Turbo 394

Kolors Virtual Try-On 3,69k

FacePoke 304

Expression Editor 353

Browse 400k+ models Browse 150k+ applications Browse 100k+ datasets

Tasks Libraries Datasets Languages Licenses Other

Filter Tasks by name Reset Tasks

Multimodal

- Image-Text-to-Text
- Visual Question Answering
- Document Question Answering
- Video-Text-to-Text
- Any-to-Any

Computer Vision

- Depth Estimation
- Image Classification
- Object Detection
- Image Segmentation
- Text-to-Image
- Image-to-Text
- Image-to-Image
- Image-to-Video
- Unconditional Image Generation
- Video Classification
- Text-to-Video
- Zero-Shot Image Classification
- Mask Generation
- Zero-Shot Object Detection
- Text-to-3D
- Image-to-3D
- Image Feature Extraction
- Keypoint Detection

Natural Language Processing

- Text Classification
- Token Classification
- Table Question Answering
- Question Answering
- Zero-Shot Classification
- Translation

Models 136 vgg

amehta633/cifar-10-vgg-pretrained
Image Classification • Updated Jun 8, 2022 • 12

edadaltocg/vgg16_bn_cifar100
Image Classification • Updated Feb 20, 2023 • 43

timm/vgg11.tv_in1k
Image Classification • Updated Apr 25, 2023 • 716

timm/vgg13.tv_in1k
Image Classification • Updated Apr 25, 2023 • 523

timm/vgg16.tv_in1k
Image Classification • Updated Apr 25, 2023 • 28.3k • 5

timm/vgg19.tv_in1k
Image Classification • Updated Apr 25, 2023 • 9.14k • 2

schibfab/landscape_classification_vgg16_fine_tuned-...
Image Classification • Updated May 28, 2023

Grecko/AQUA-vgg16-imagenet
Image Classification • Updated Jun 19, 2023 • 2

Reusing models without full retraining

Skin Cancer Classification using VGG-16 and Googlenet CNN Models

January 2023 · [International Journal of Computer Applications](#) 184(42):5-9

Anju, T.E., Vimala, S. (2023). Finetuned-VGG16 CNN Model for Tissue Classification of Colorectal Cancer. In: Raj, J.S., Perikos, I., Balas, V.E. (eds) Intelligent Sustainable Systems. ICoISS 2023. Lecture Notes in Networks and Systems, vol 665. Springer, Singapore.

VGG 16 Pre-Trained Model for Early Detection of Retinal Diseases

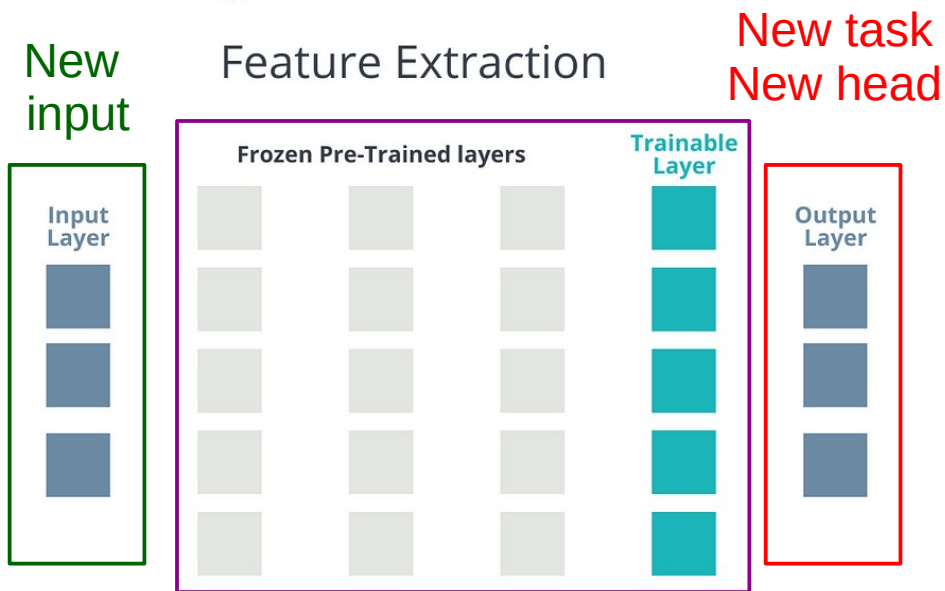
December 2023

DOI:[10.1109/SMARTGENCON60755.2023.10442010](#)

Conference: 2023 3rd International Conference on Smart Generation Computing,

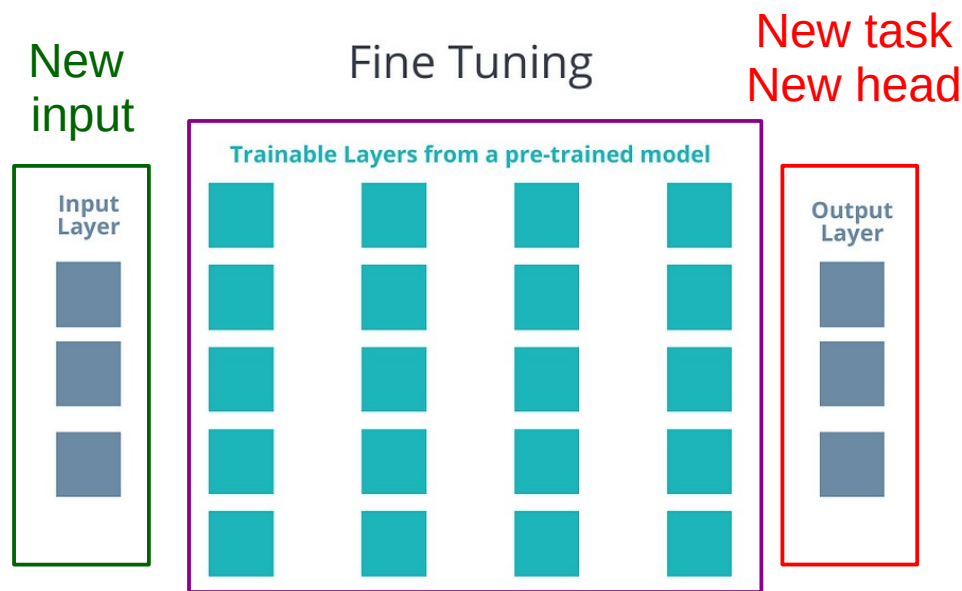
Raghuvanshi S, Dhariwal S. The VGG16 Method Is a Powerful Tool for Detecting Brain Tumors Using Deep Learning Techniques. *Engineering Proceedings*. 2023; 59(1):46. <https://doi.org/10.3390/engproc2023059046>

Transfer Learning: How Feature Extraction & Fine-Tuning work?



Existing model

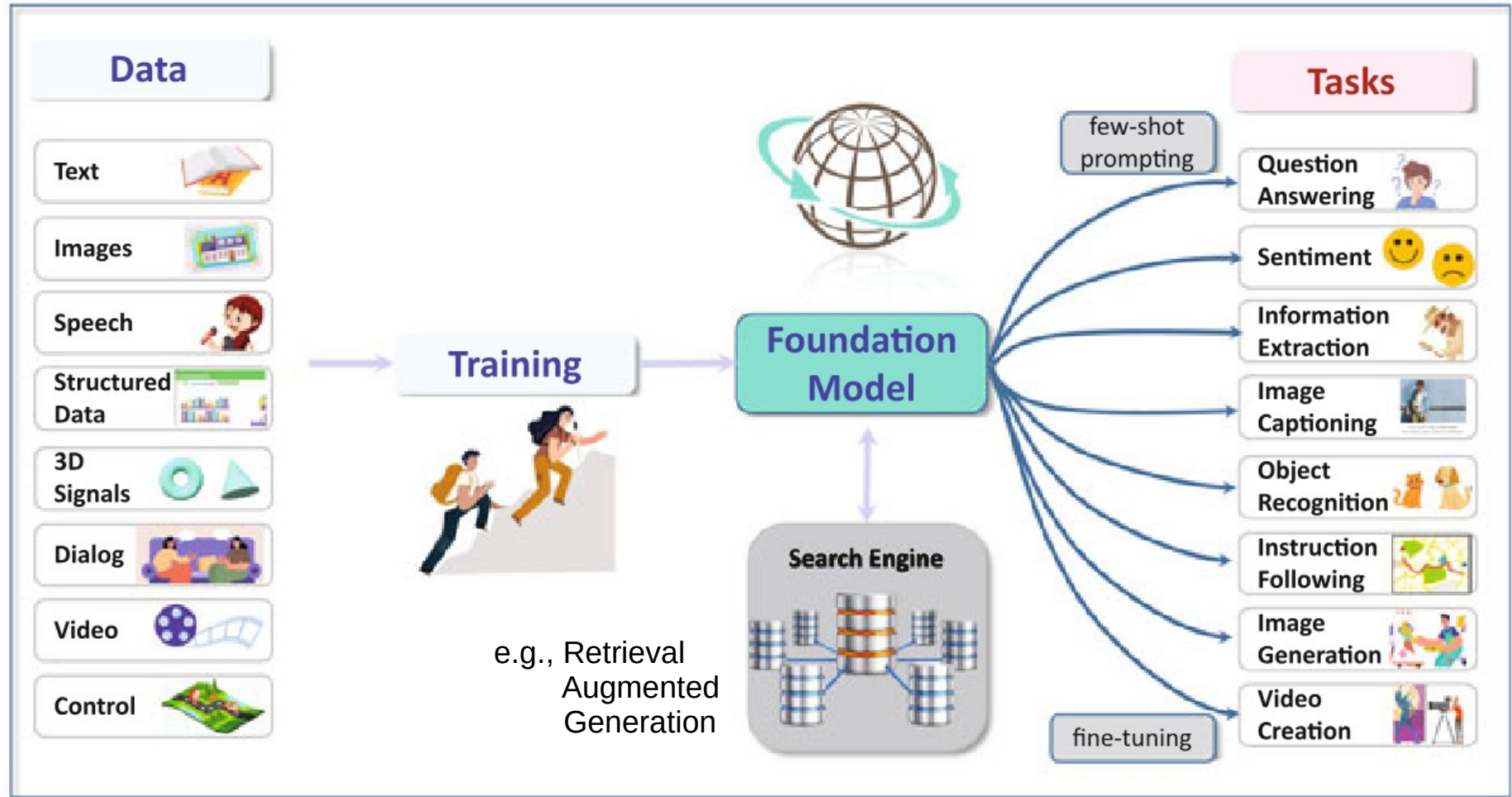
In feature extraction, you **freeze the pre-trained model** layers to preserve existing learning and **add new layers** to learn additional information.



Existing model

In fine-tuning, you **unfreeze the entire model** and **train it with a lower learning rate** to adapt to new challenges.

Extreme transfer learning: Foundation models



5

Embeddings and latent spaces

Embeddings (“plongement”)

Values in reference frame A

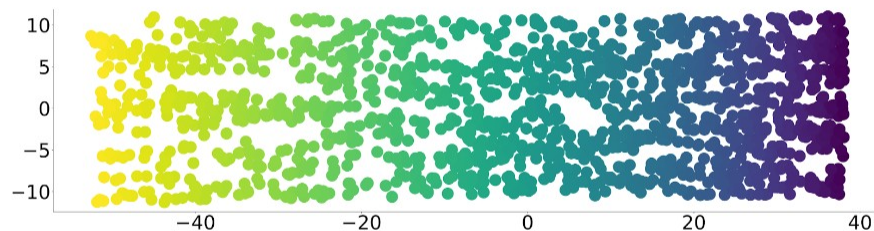
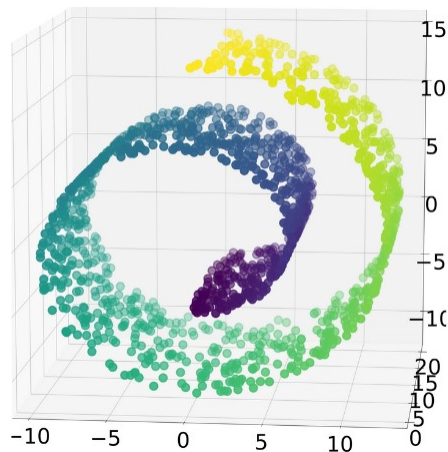
m vectors of size n
from a dictionary of o
(i.e. o coordinates)



Values in reference frame B

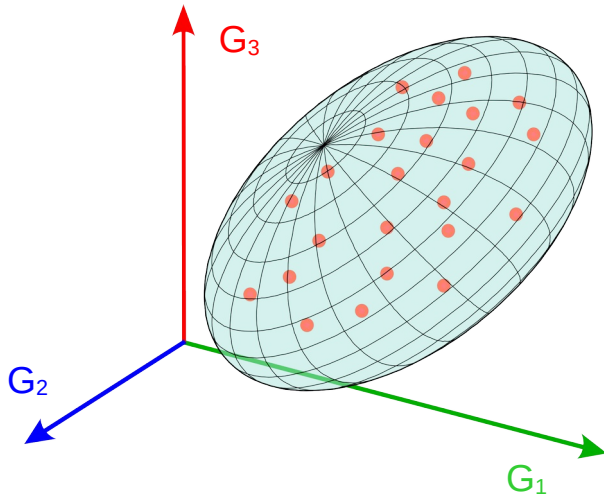
m vectors of size n'
from a dictionary of o'
(i.e. o' coordinates)

Embedding from a space with n dimensions
into a space of o dimensions



A PCA is an embedding

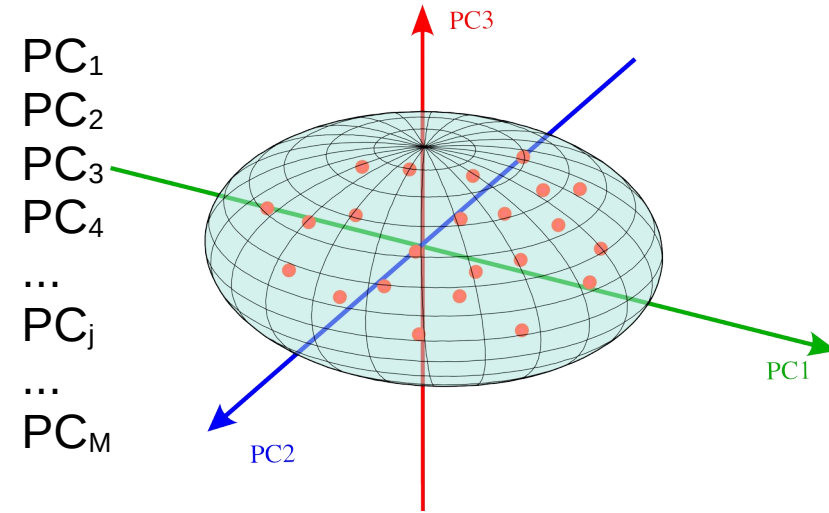
Axes = N genes
Coordinate = expressions



G_1
 G_2
 G_3
 G_4
 $\dots$
 G_i
 $\dots$
 G_N



Axes = M principal components
Coordinates = Rotations x expressions



PC_1
 PC_2
 PC_3
 PC_4
 $\dots$
 PC_j
 $\dots$
 PC_M

“Similar” objects are neighbours in the latent space

Dictionary (initial space)

	bunny	father	hamster	hutches	man	mother	rabbit	tractor	woman
bunny	1	0	0	0	0	0	0	0	0
rabbit	0	0	0	0	0	0	1	0	0
hamster	0	0	1	0	0	0	0	0	0
hutches	0	0	0	1	0	0	0	0	0

	bunny	father	hamster	hutches	man	mother	rabbit	tractor	woman
mother	0	0	0	0	0	1	0	0	0
father	0	1	0	0	0	0	0	0	0
woman	0	0	0	0	0	0	0	0	1
man	0	0	0	0	1	0	0	0	0

Semantics (Embedding space)

	alive	mammal	rodent	bipedal	gender	plural
bunny	0.9	0.7	-0.1	-0.1	0.2	-0.3
rabbit	0.9	0.8	-0.1	-0.1	0.4	-0.1
hamster	0.9	0.6	0.7	-0.3	-0.4	-0.4
hutches	-0.9	-0.3	-0.1	-0.9	0.0	0.9

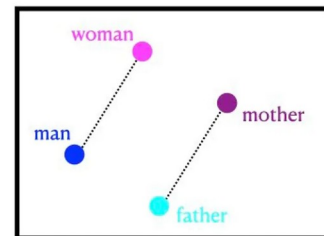
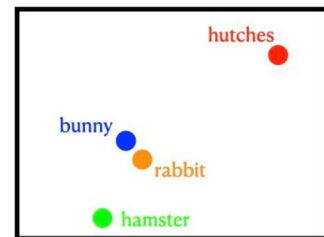
	alive	mammal	rodent	bipedal	gender	plural
mother	0.9	0.1	0.1	0.2	0.8	-0.7
father	0.9	0.2	0.1	0.2	-0.8	-0.7
woman	0.9	0.8	-0.9	0.9	0.8	-0.8
man	0.9	0.8	-0.7	0.9	-0.8	-0.8

Word

Vector Embedding

Dimensionality
Reduction

Dimensionality



2D Visualization

Arithmetic operation in latent space = semantic statement

Dictionary (initial space)

	bunny	father	hamster	hutches	man	mother	rabbit	tractor	woman
bunny	1	0	0	0	0	0	0	0	0
rabbit	0	0	0	0	0	0	1	0	0
hamster	0	0	1	0	0	0	0	0	0
hutches	0	0	0	1	0	0	0	0	0

	bunny	father	hamster	hutches	man	mother	rabbit	tractor	woman
mother	0	0	0	0	0	1	0	0	0
father	0	1	0	0	0	0	0	0	0
woman	0	0	0	0	0	0	0	0	1
man	0	0	0	0	1	0	0	0	0

Semantics (Embedding space)

	alive	mammal	rodent	bipedal	gender	plural
bunny	0.9	0.7	-0.1	-0.1	0.2	-0.3
rabbit	0.9	0.8	-0.1	-0.1	0.4	-0.1
hamster	0.9	0.6	0.7	-0.3	-0.4	-0.4
hutches	-0.9	-0.3	-0.1	-0.9	0.0	0.9

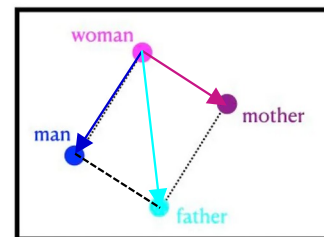
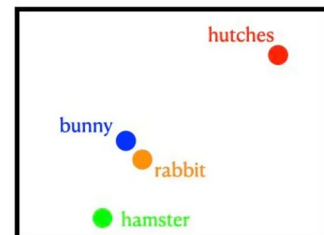
	alive	mammal	rodent	bipedal	gender	plural
mother	0.9	0.1	0.1	0.2	0.8	-0.7
father	0.9	0.2	0.1	0.2	-0.8	-0.7
woman	0.9	0.8	-0.9	0.9	0.8	-0.8
man	0.9	0.8	-0.7	0.9	-0.8	-0.8

Word

Vector Embedding

Dimensionality
Reduction

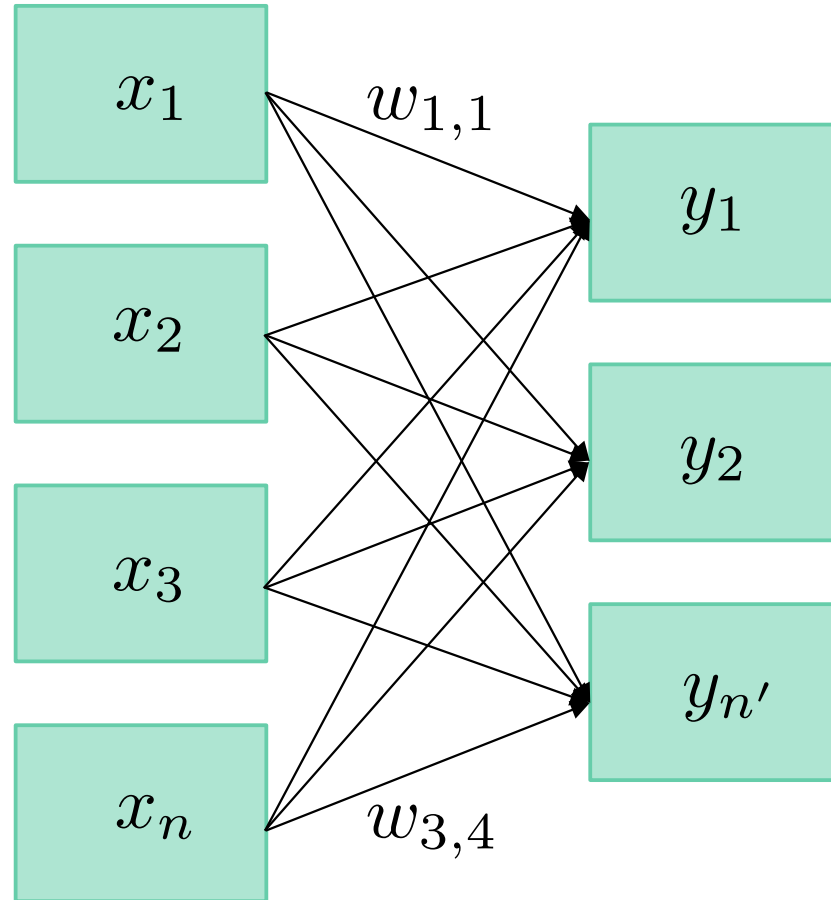
Dimensionality



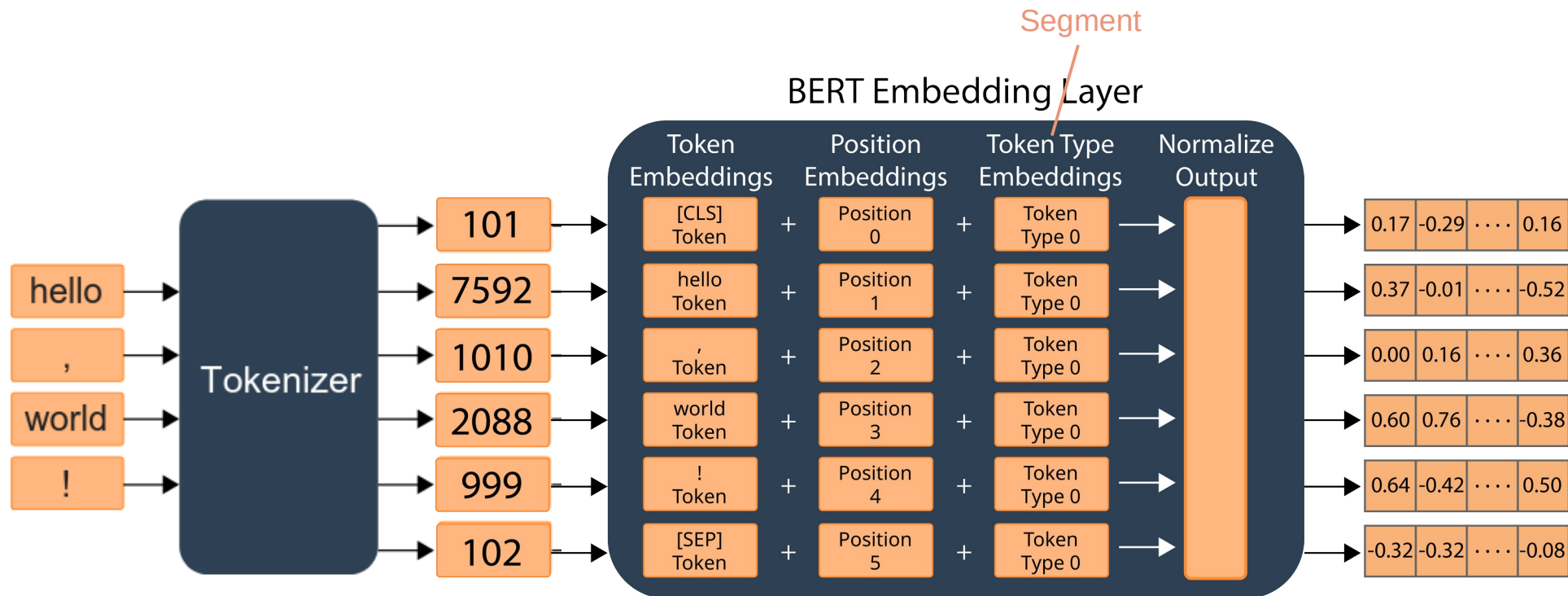
2D Visualization

In the latent space, the vector going from **woman** to **father** is equal to the vector going from **woman** to **man** plus the vector going from **woman** to **mother**

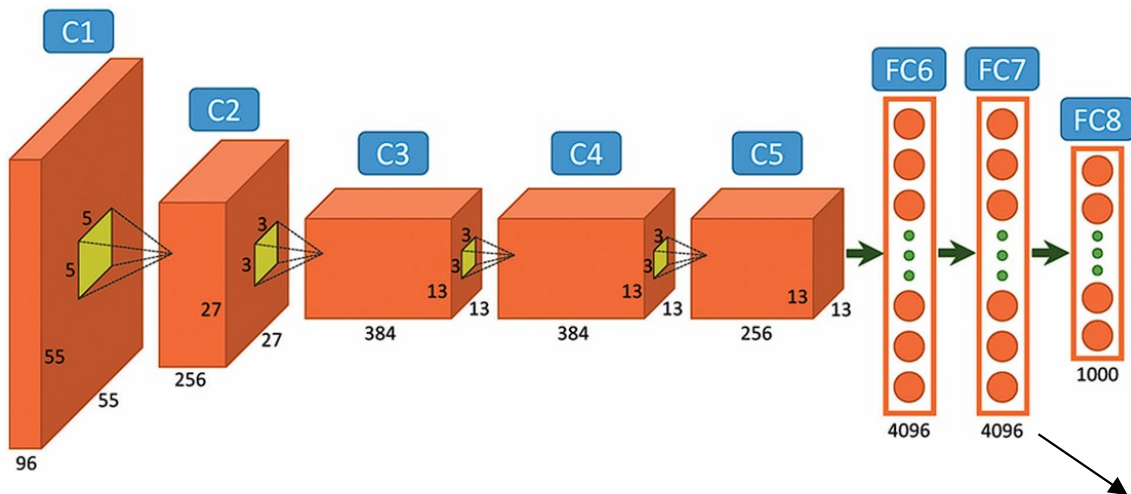
DL: Fully connected layers to learn the embedding



There can be several embeddings



Deep CNNs are “embedding” the images in a “latent space”



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>

(presenting AlexNet, the first Deep Convolutional Network)

<https://huggingface.co/debashd/AlexNet>

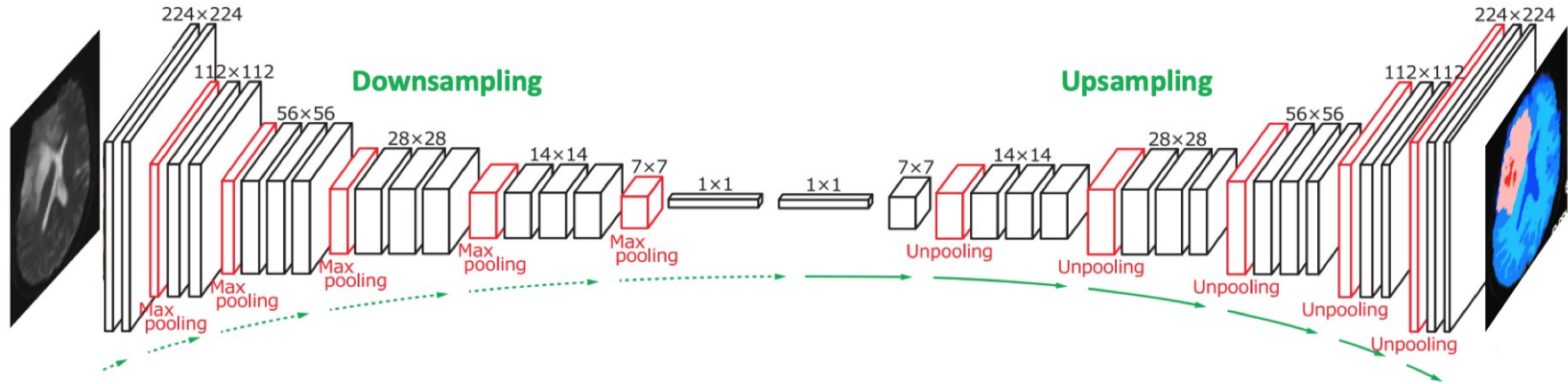
input
image



6

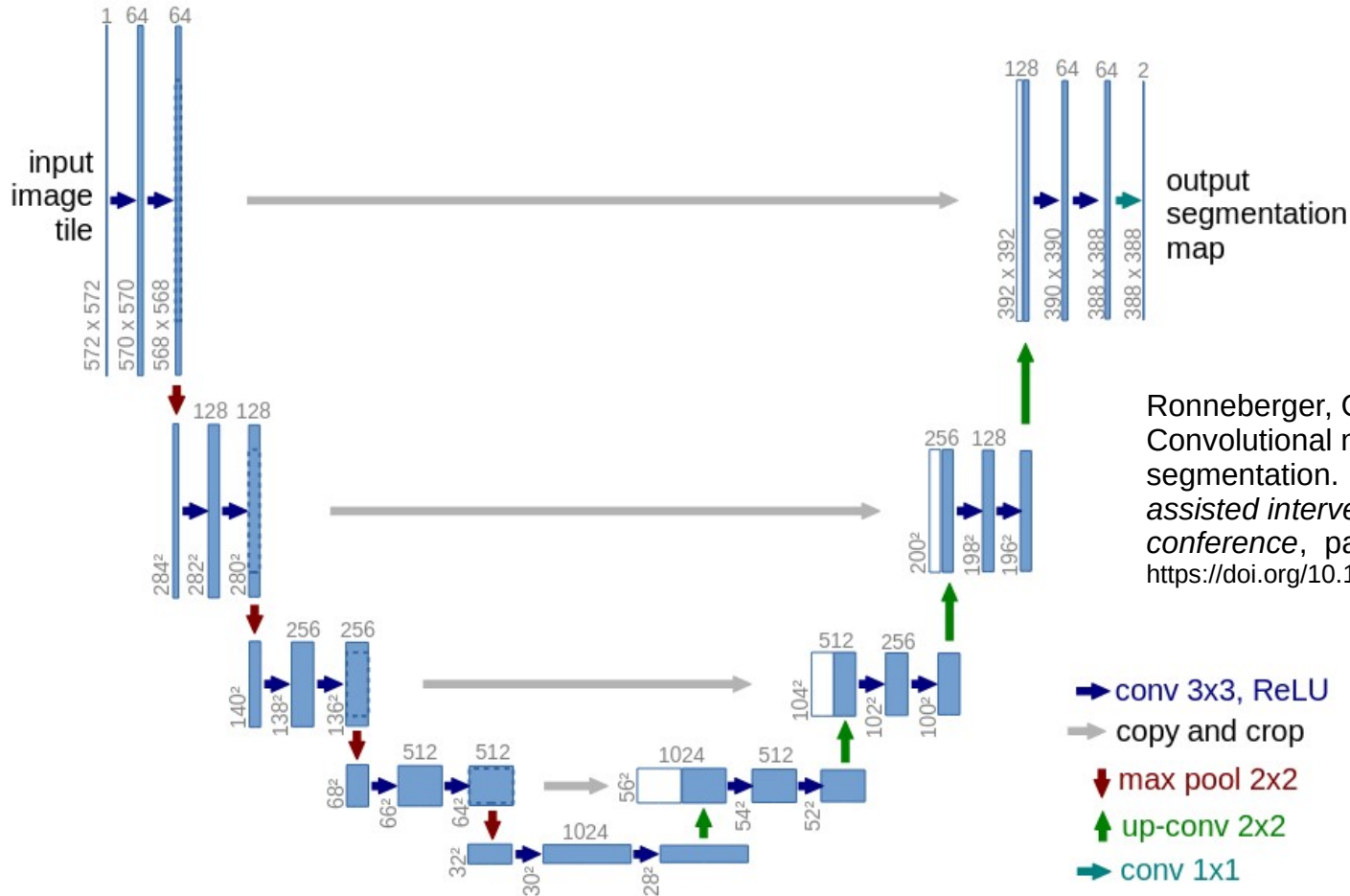
Encoder-Decoders (variational) AutoEncoders (VAEs)

Encoder-decoder



Example of segmentation to identify brain tumours

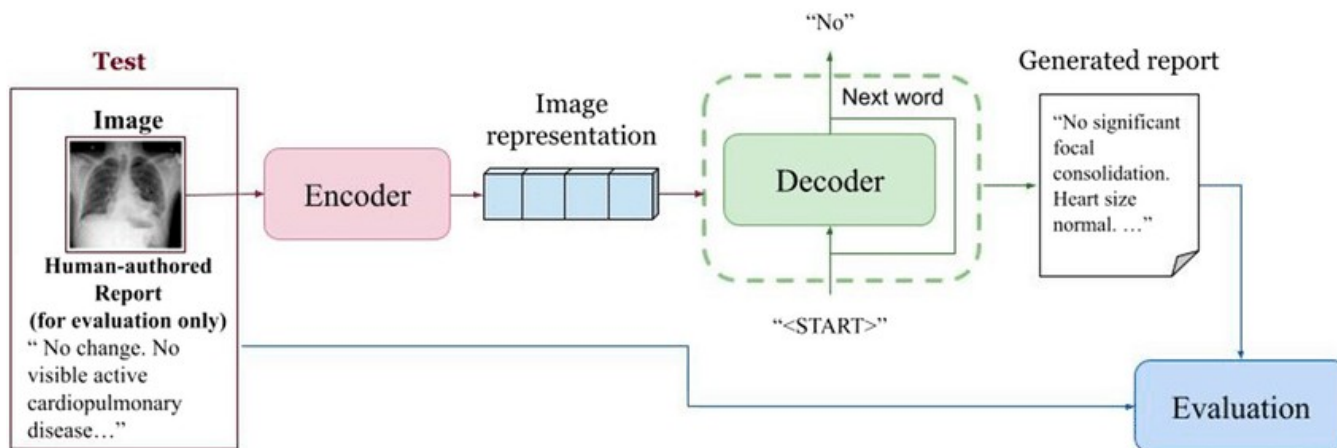
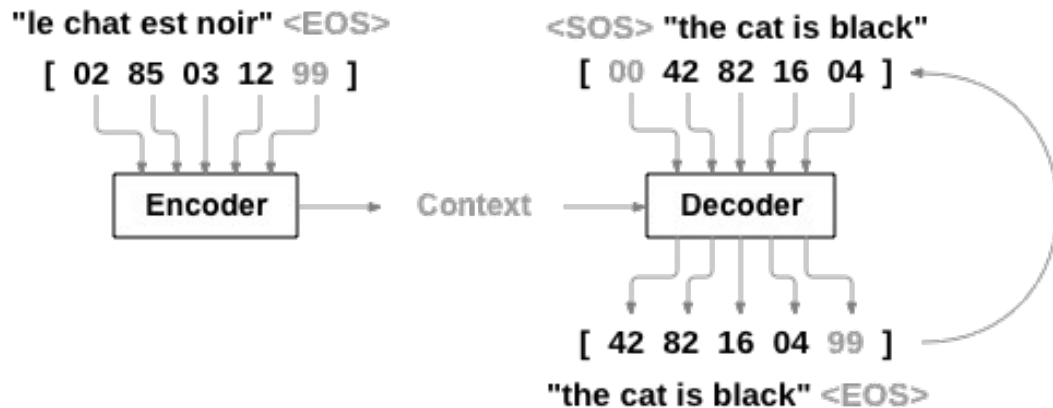
Disclaimer: it is more complicated



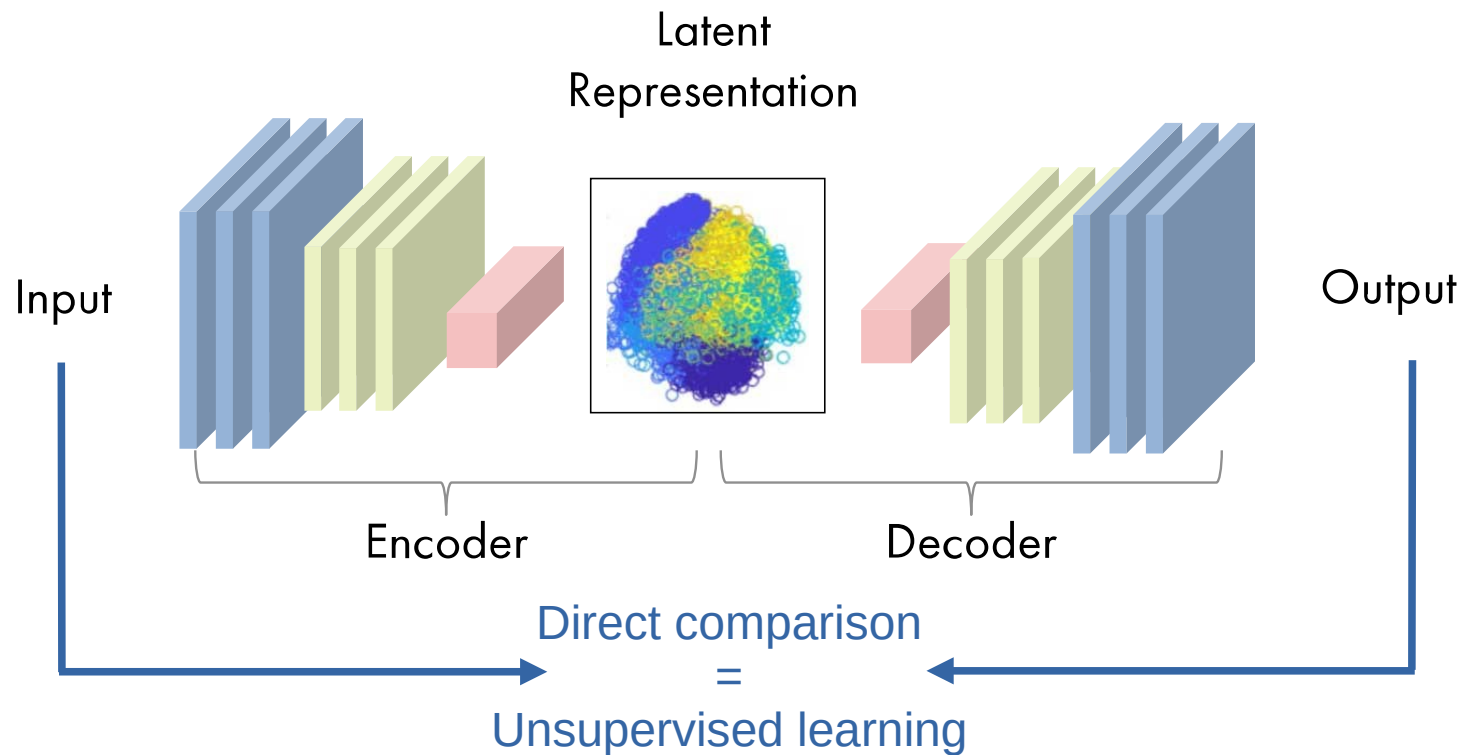
Ronneberger, O., Fischer, P., & Brox, T. (2015). U-net: Convolutional networks for biomedical image segmentation. In *Medical image computing and computer-assisted intervention—MICCAI 2015: 18th international conference, part III 18* (pp. 234-241) https://doi.org/10.1007/978-3-319-24574-4_28

Encoders and decoders can be anything

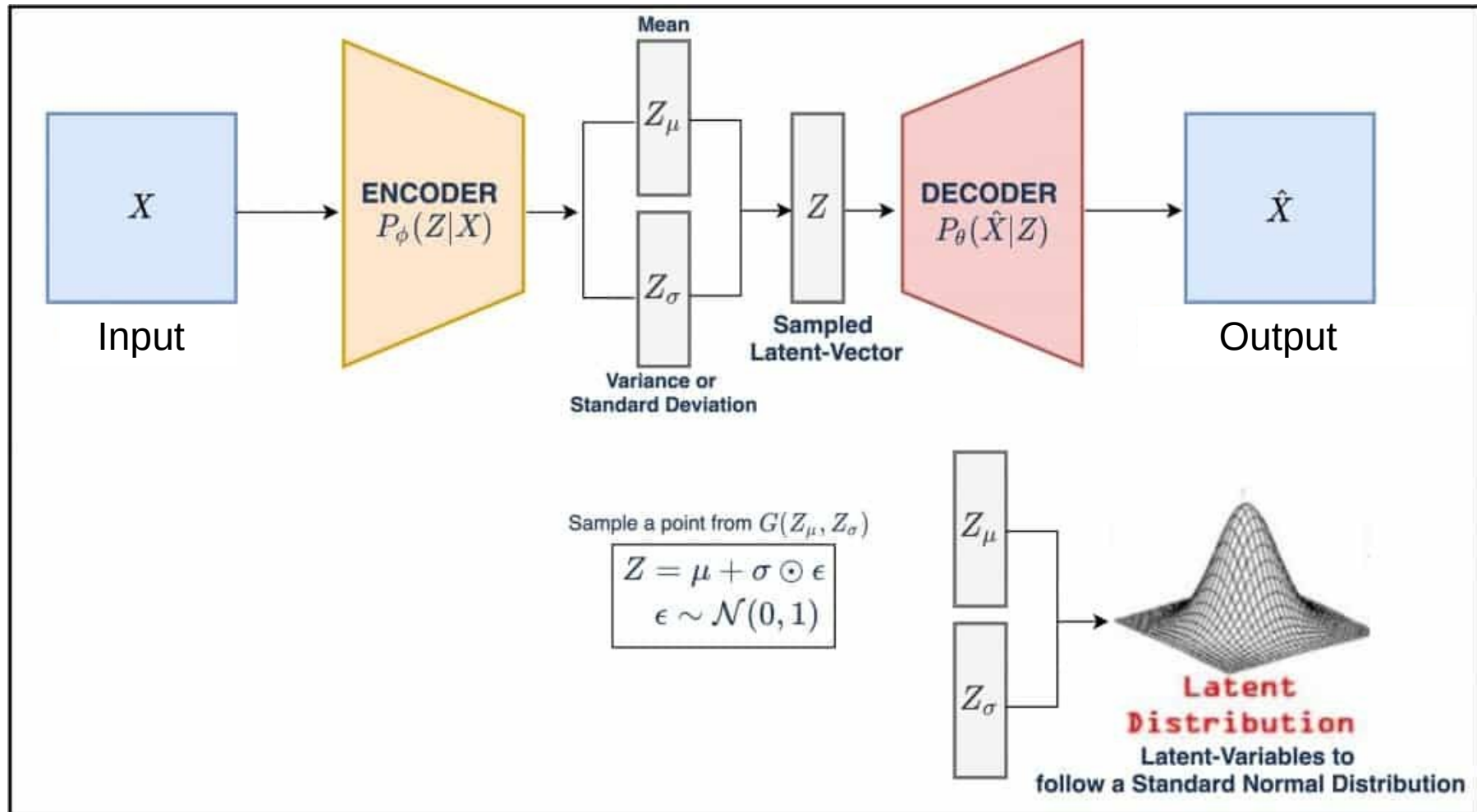
Pavlopoulos *et al* (2022)
doi:10.1007/s10115-022-01684-7



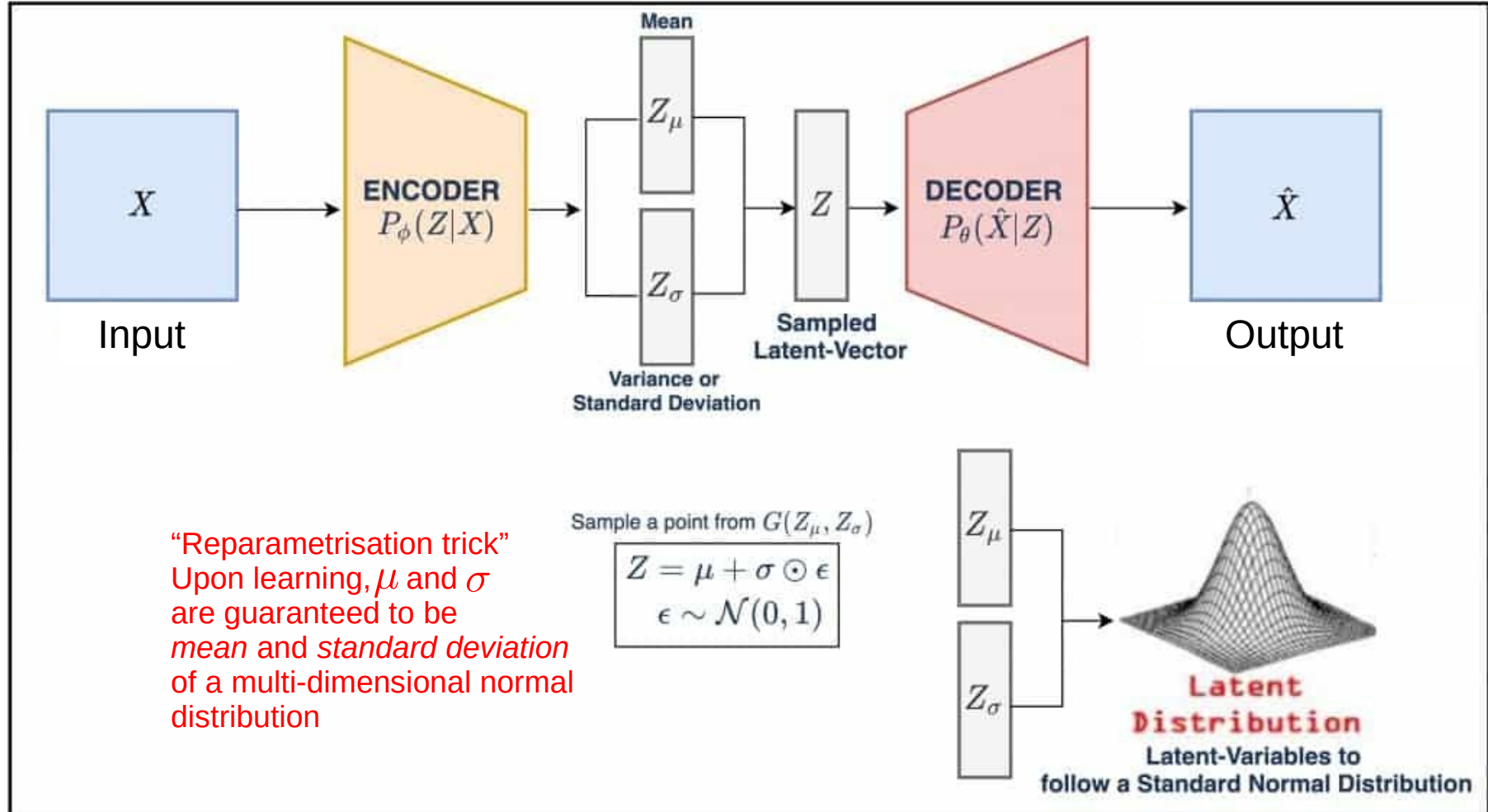
Learn by itself: AutoEncoder



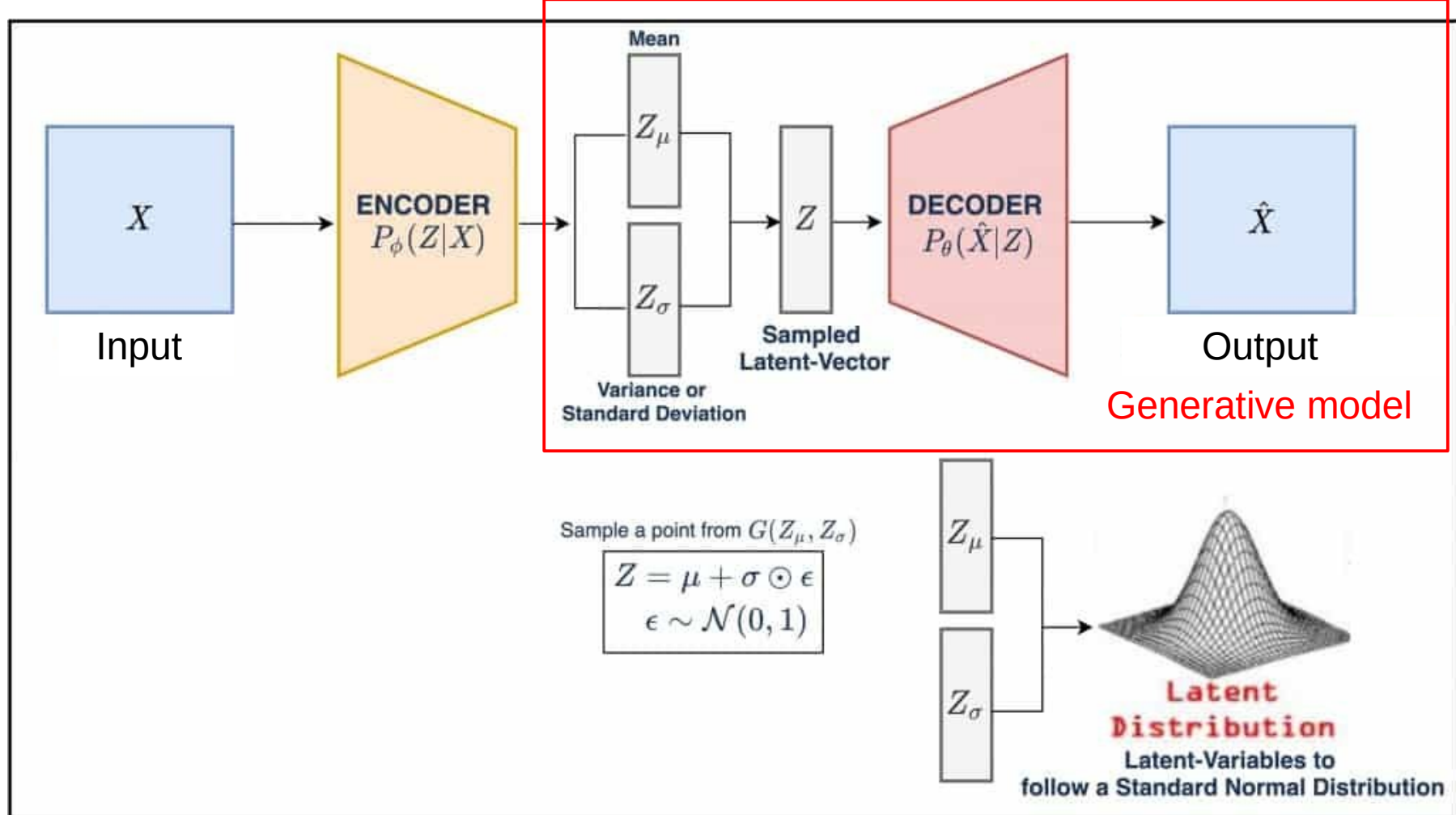
Variational Auto Encoder (VAE)

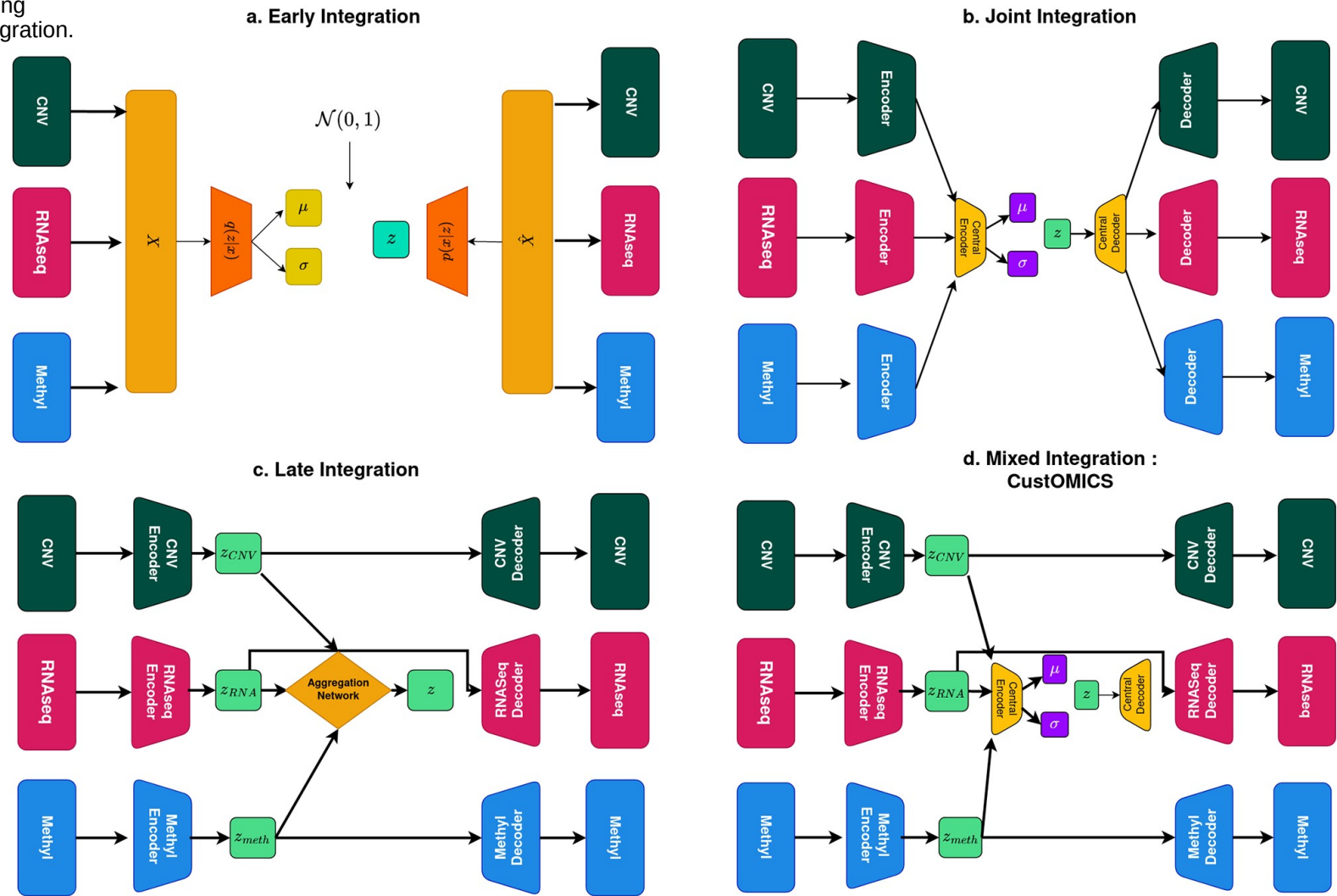


Variational Auto Encoder (VAE)



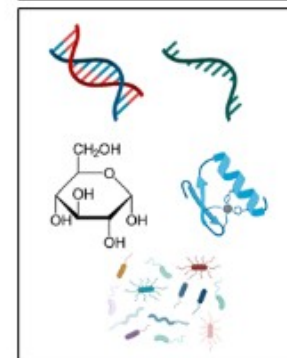
Variational Auto Encoder (VAE)





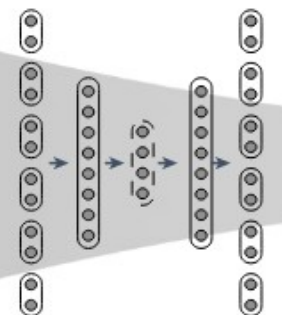
Discovery of drug–omics associations in type 2 diabetes with generative deep-learning models

Non-omics data

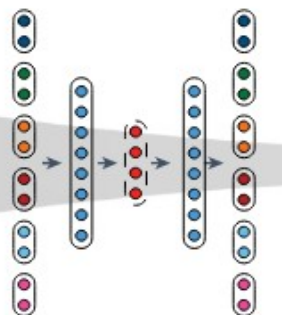


Multi-omics data

Naïve VAE model



Trained VAE model



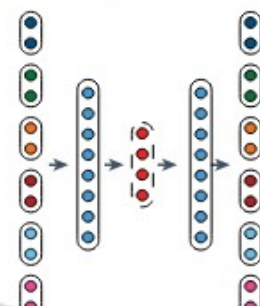
VAE hyperparameter selection

Test likelihood

Model accuracy

Model stability

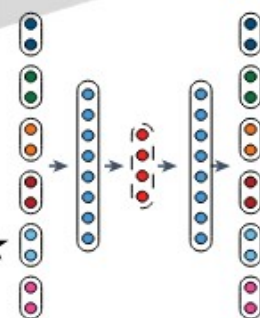
Baseline model



Significant drug–omics associations

Drug 0→1★

Drug-perturbed model



7

Attention and the Transformer

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... 140833 times as of 13 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].

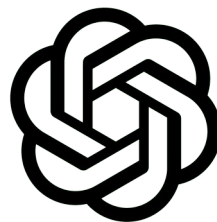
*Equal contribution. Listing order is random. Jakob proposed replacing RNNs with self-attention and started the effort to evaluate this idea. Ashish, with Illia, designed and implemented the first Transformer models and has been crucially involved in every aspect of this work. Noam proposed scaled dot-product 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.

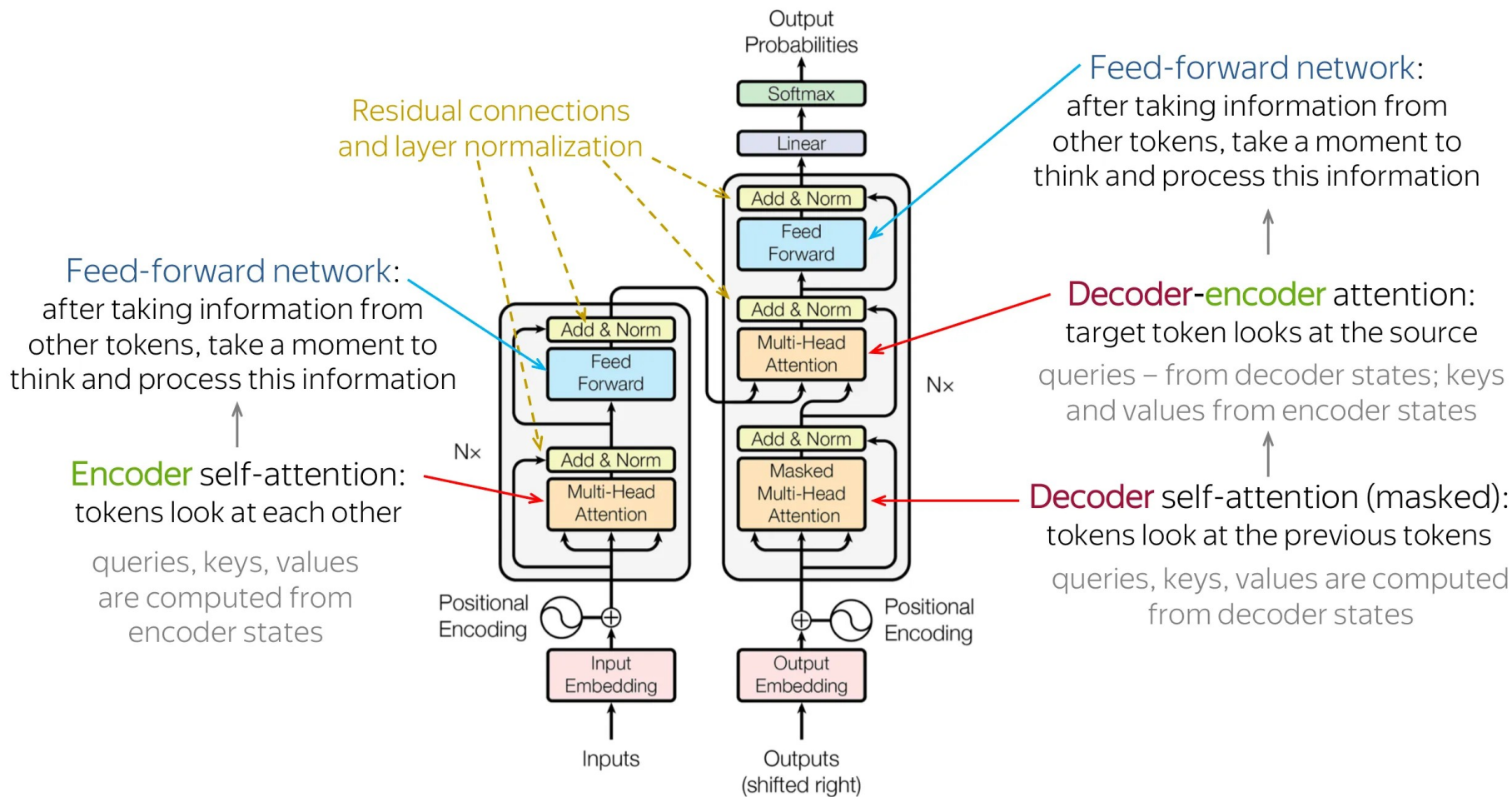
[‡]Work performed while at Google Research.

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

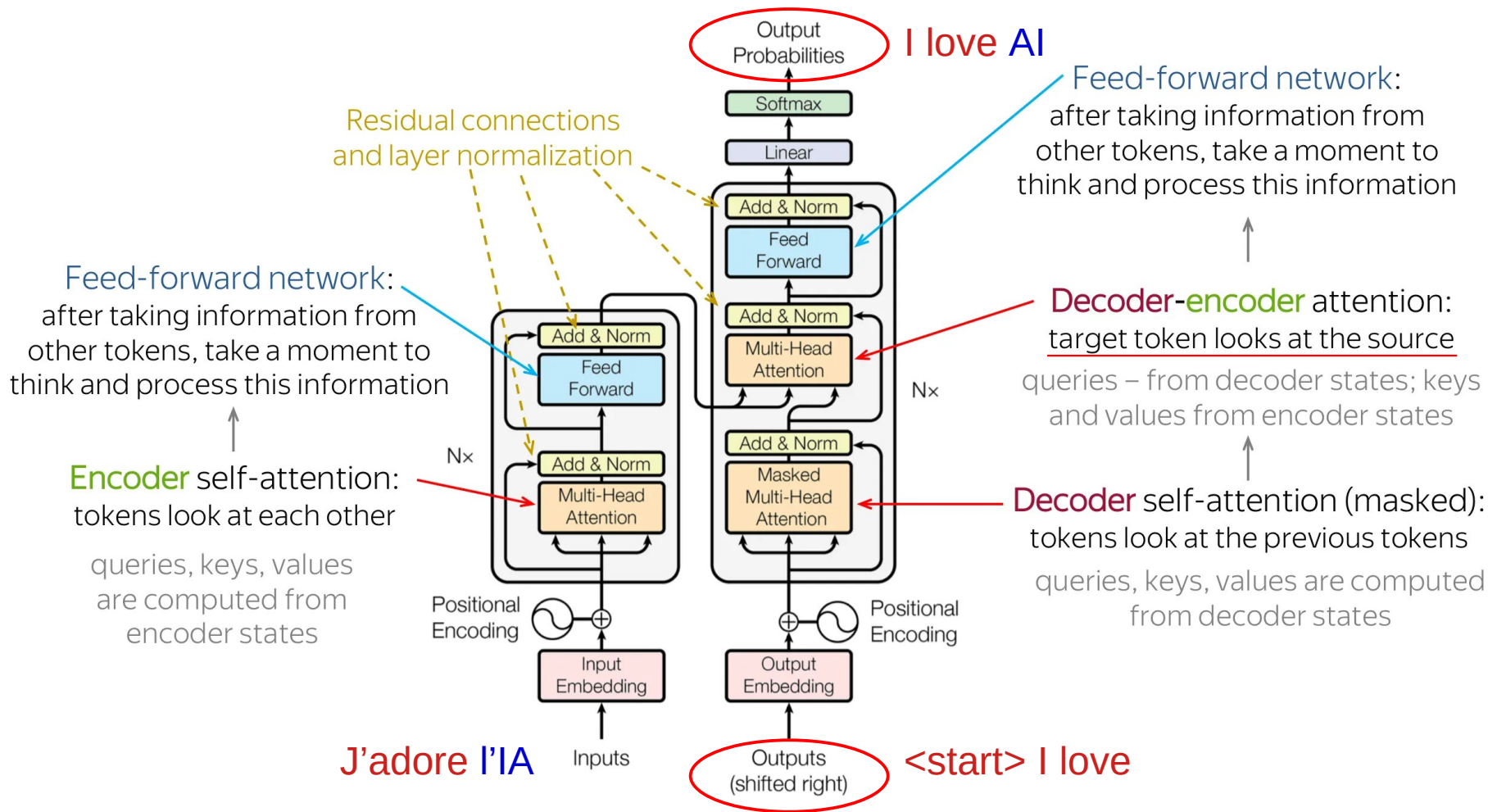
The paper that changed everything: the Transformer



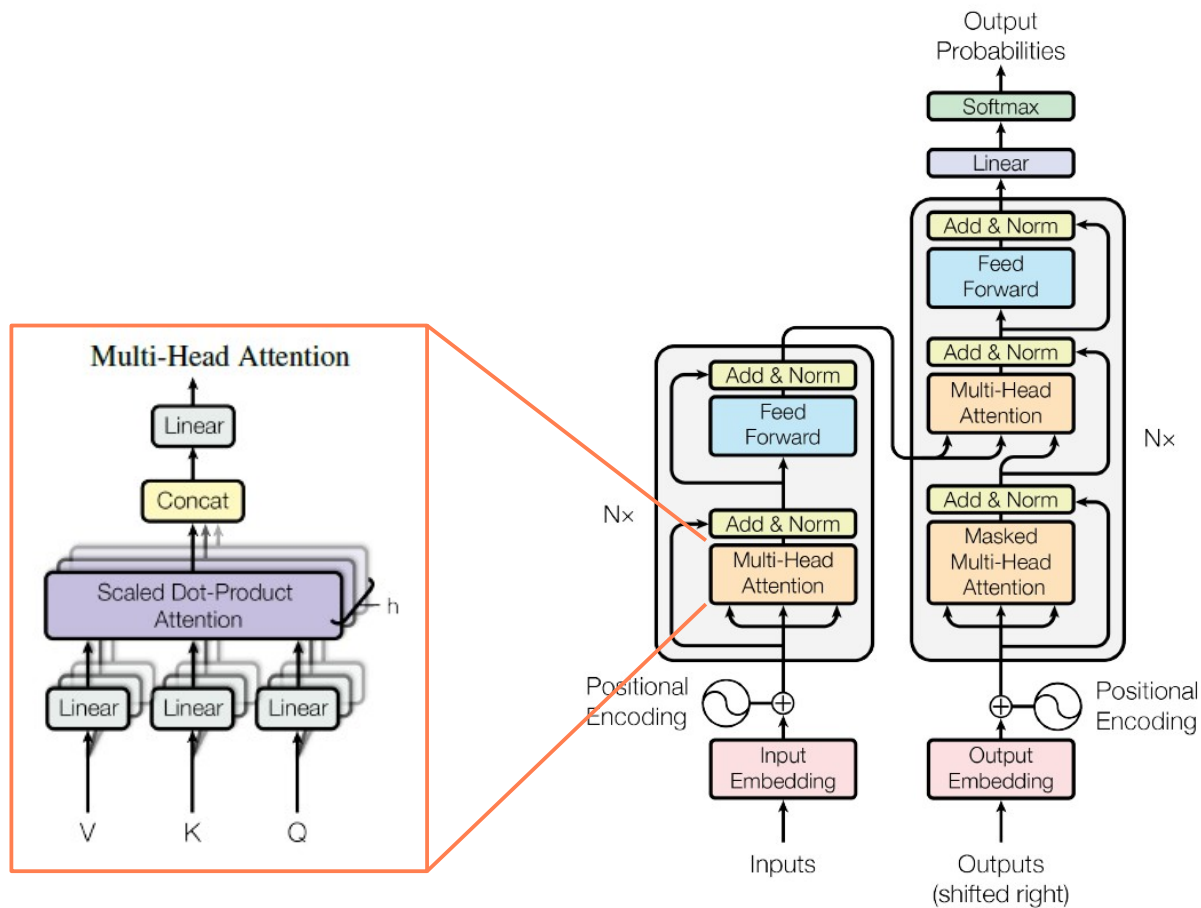
The Transformer: Memory + context = attention



The Transformer: Memory + context = attention



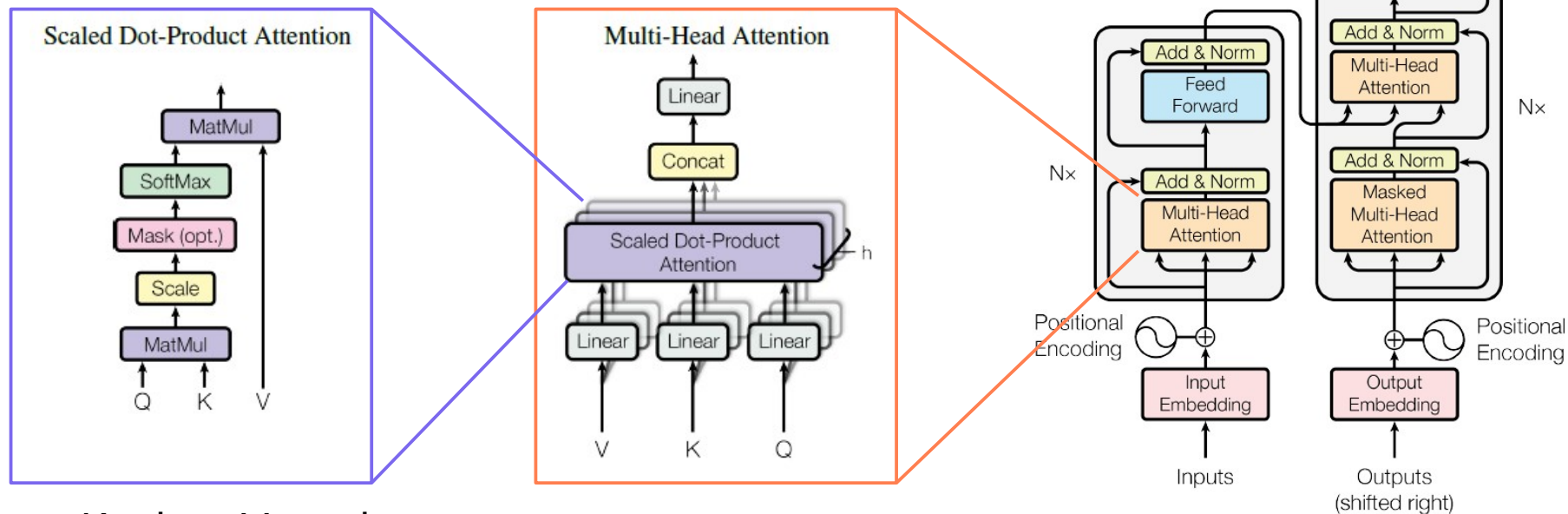
Attention in the Transformer



Q = query, K = key, V = value

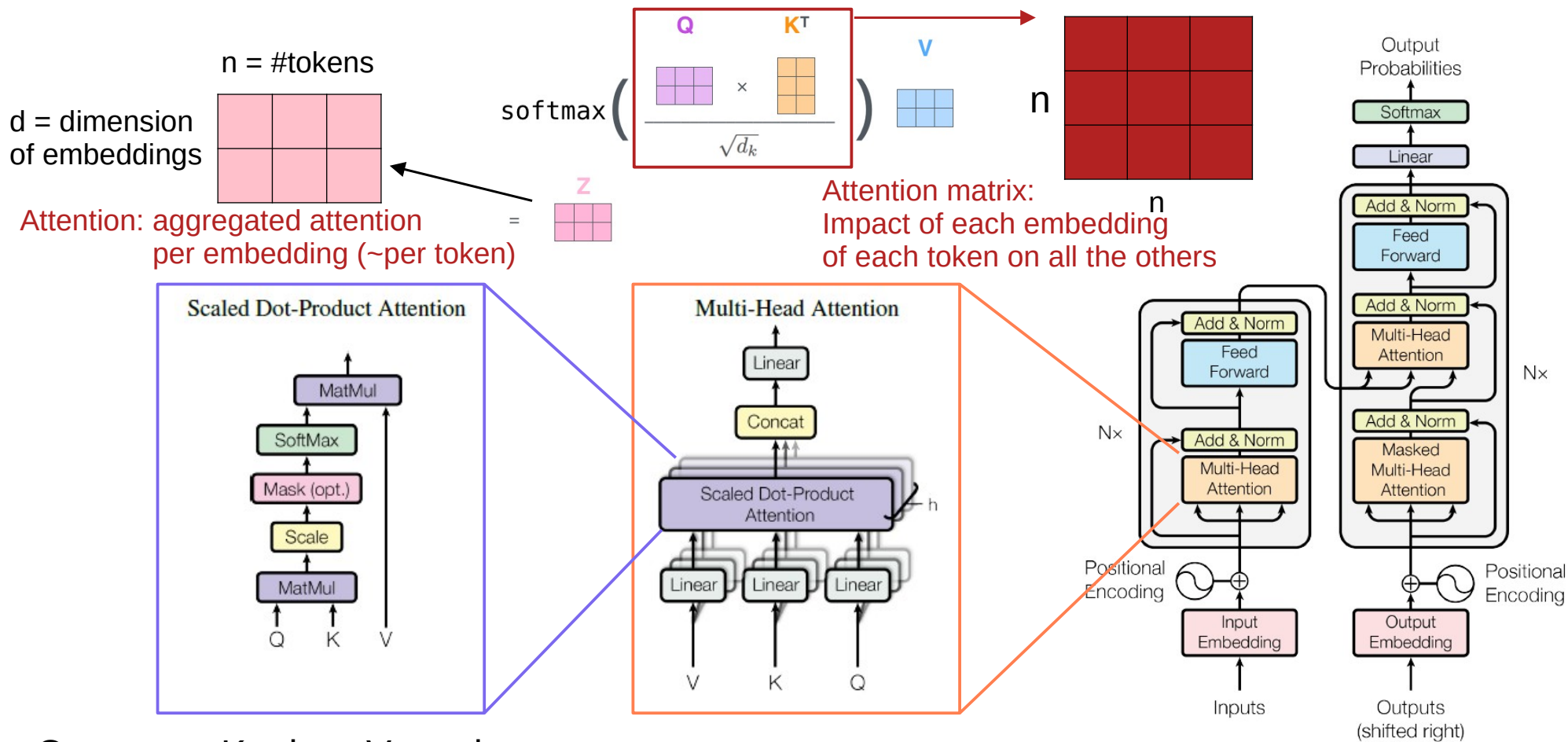
Attention in the Transformer

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



Q = query, K = key, V = value

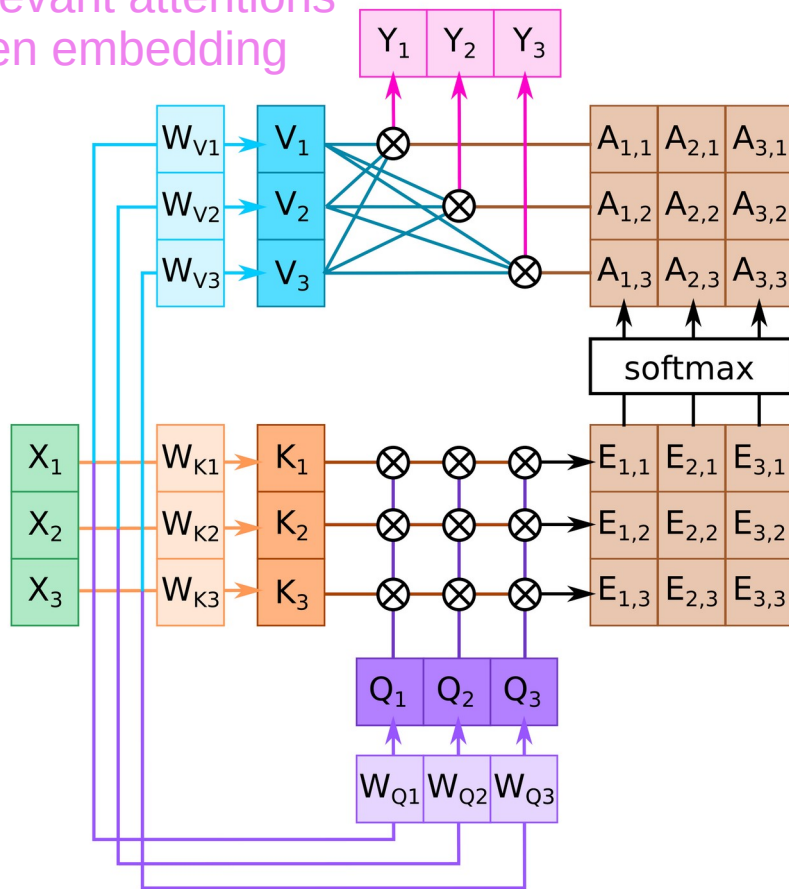
Attention in the Transformer



Q = query, K = key, V = value

Attention in the Transformer

Actual relevant attentions
1 per token embedding



Query: What are the things I am looking for?

Key: What are the thing I have?

Value: What are the things that I will communicate?

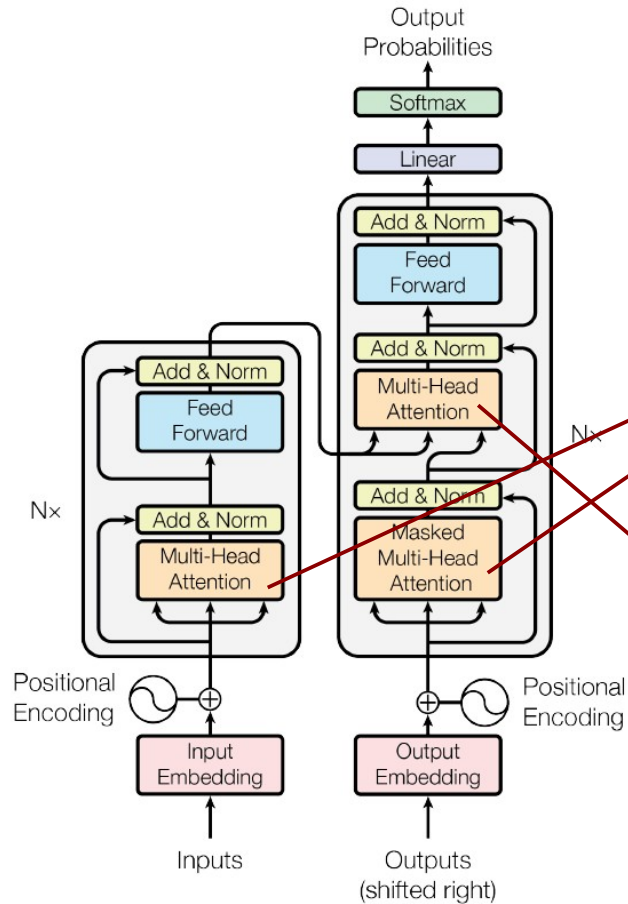
The attention of all token embeddings
on all token embeddings

X is entered

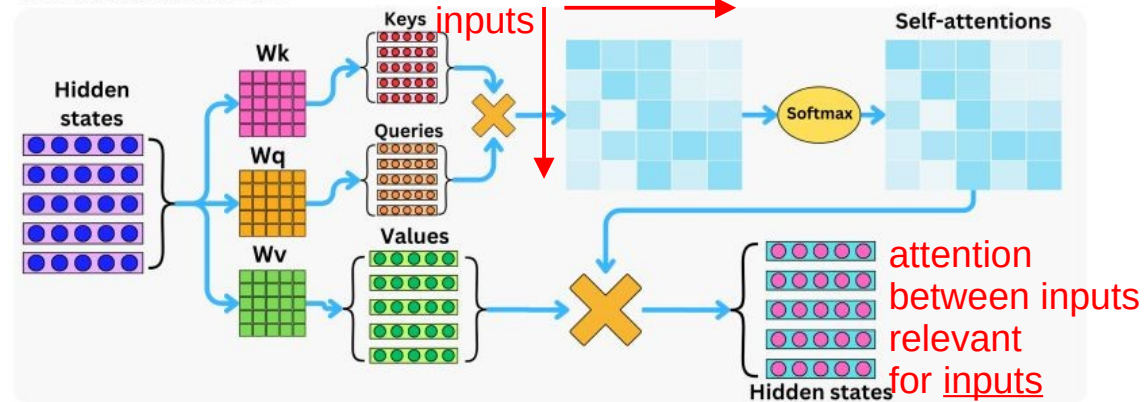
W_Q, W_K , and W_V are learned
Everything else is computed

Only one sequence (X) as input,
transformed into different sequences
of embeddings when multiplied
by the learned weights W_Q, W_K , and W_V

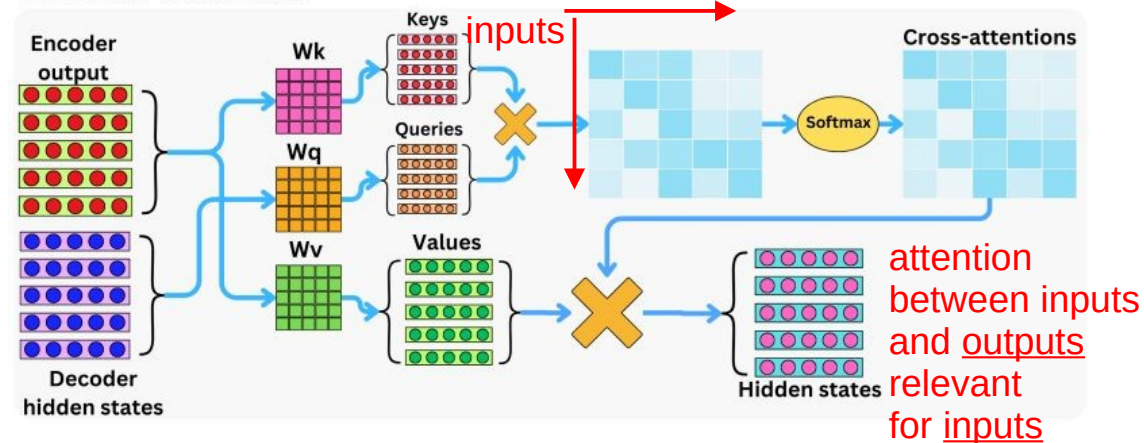
Self versus Cross-attention



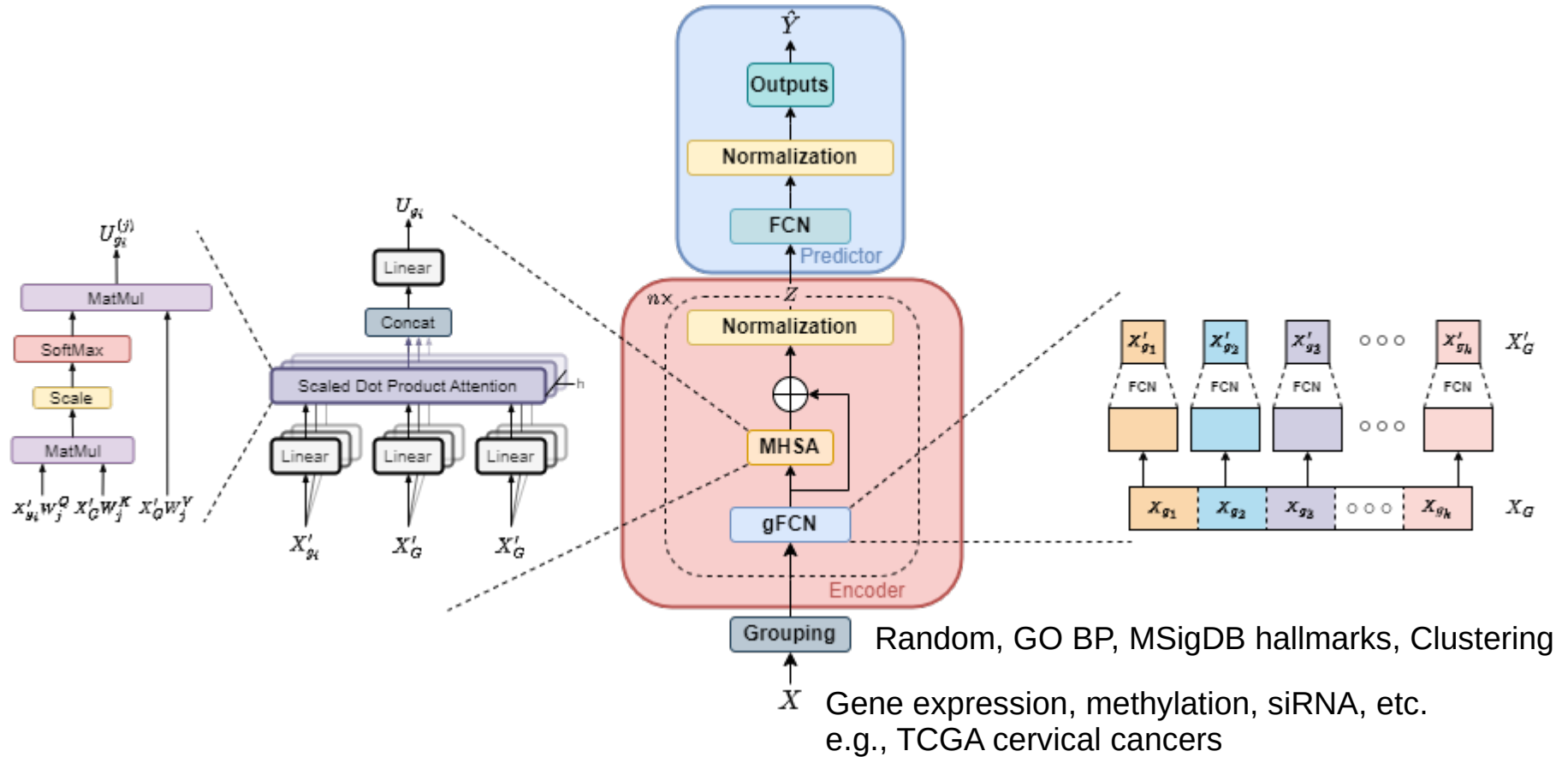
The Self-Attentions



The Cross-Attentions



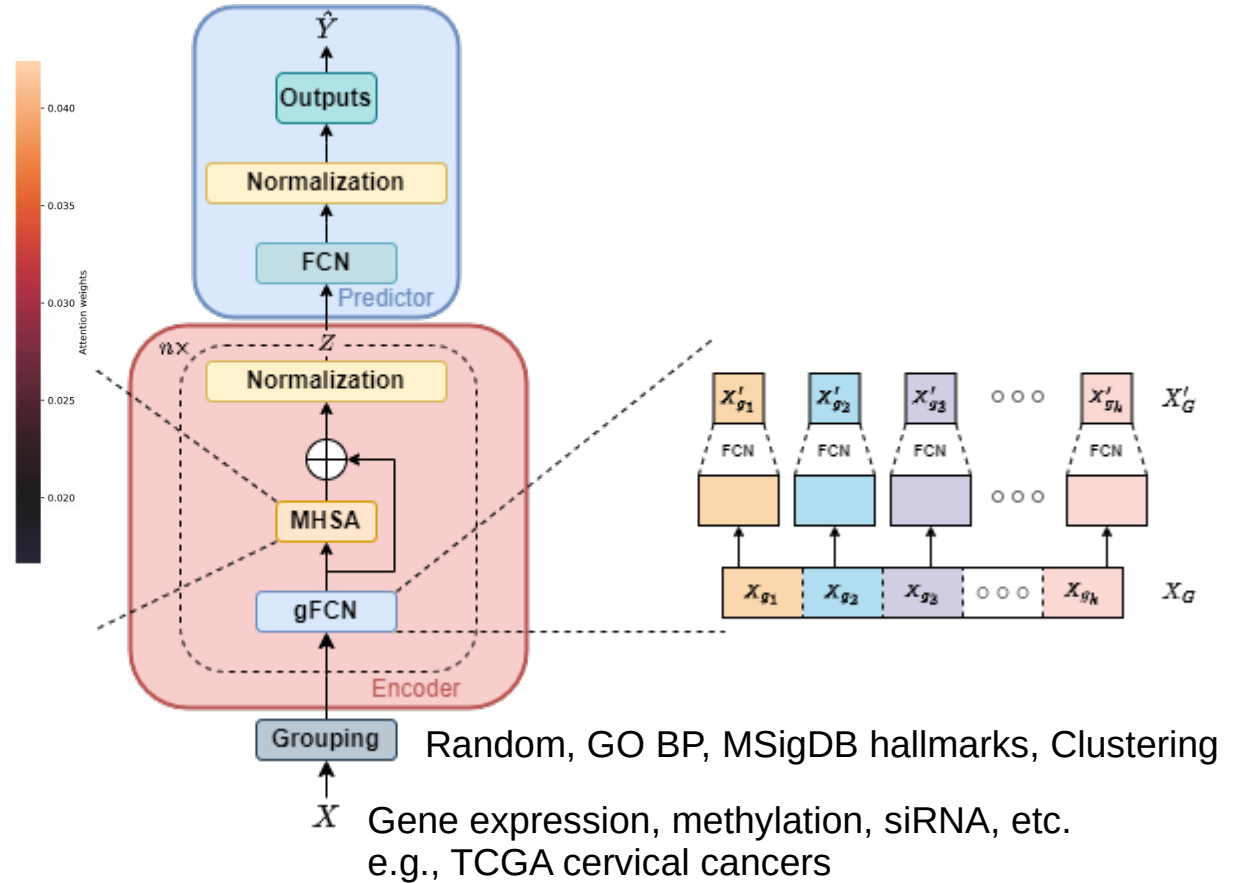
AttOmics



AttOmics



Pathways affected in cervical cancer

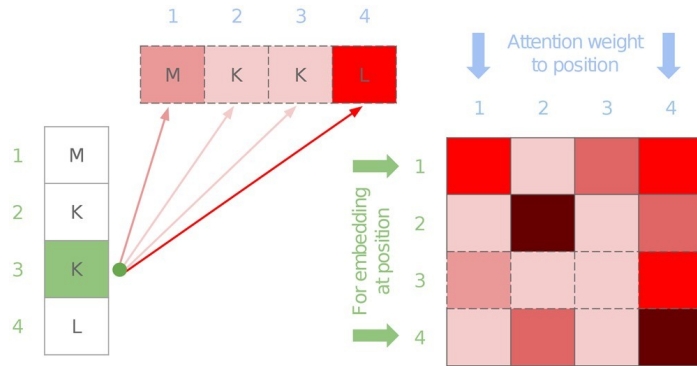


EnzBERT

Predicting enzymatic function of protein sequences with attention 🧠

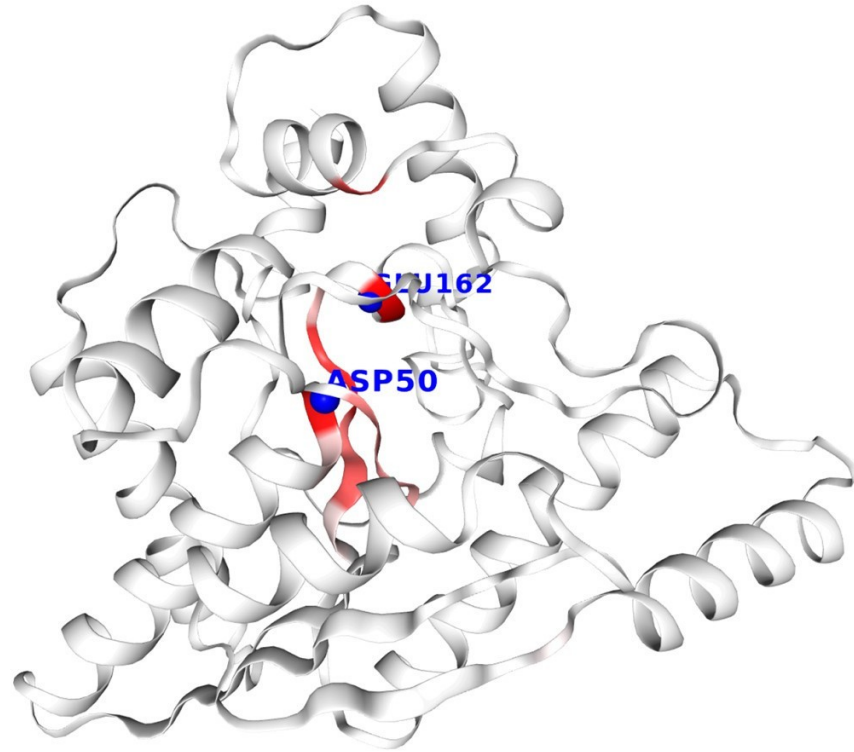
Nicolas Buton ✉, François Coste, Yann Le Cunff

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



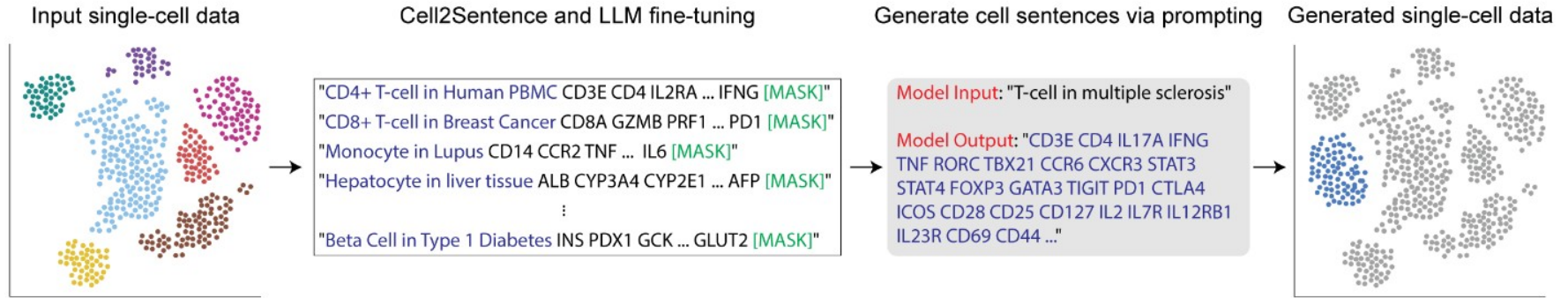
Aggregated attention
for each token (amino acid)

Nh(3)-dependent nad(+) synthetase

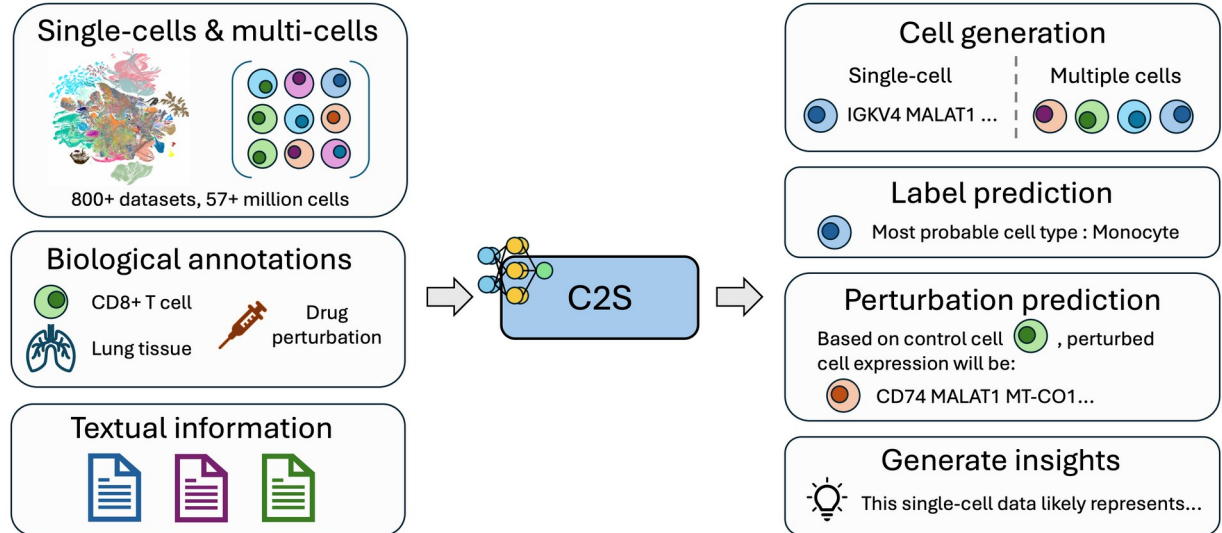


0 MSMQEKIMRE LHVKPSIDPK QEIEDRVNLF KQYVKKTGAK GFVLGISGQ DSTLAGRLAQ LAVESIREEG GDAQFIIVRL PHGTQQDEDD AQLALKFIKP
1 DKSWKFDIKS TVSAFSDQYQ QETGDQLTDF NKGNVKARTR MIAQYAIGGQ EGLLVLSIDH AAEAVTGFFT KYGDGGADLL PLTGLTKRQG RTLLKELGAP
2 ERLYLKEPTA DLLDEKPQQS DETELGISD EIDDYLEGKE VSAKVSEALE KRYSMTEHKR QVPASMFDDW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

Cell2Sentence

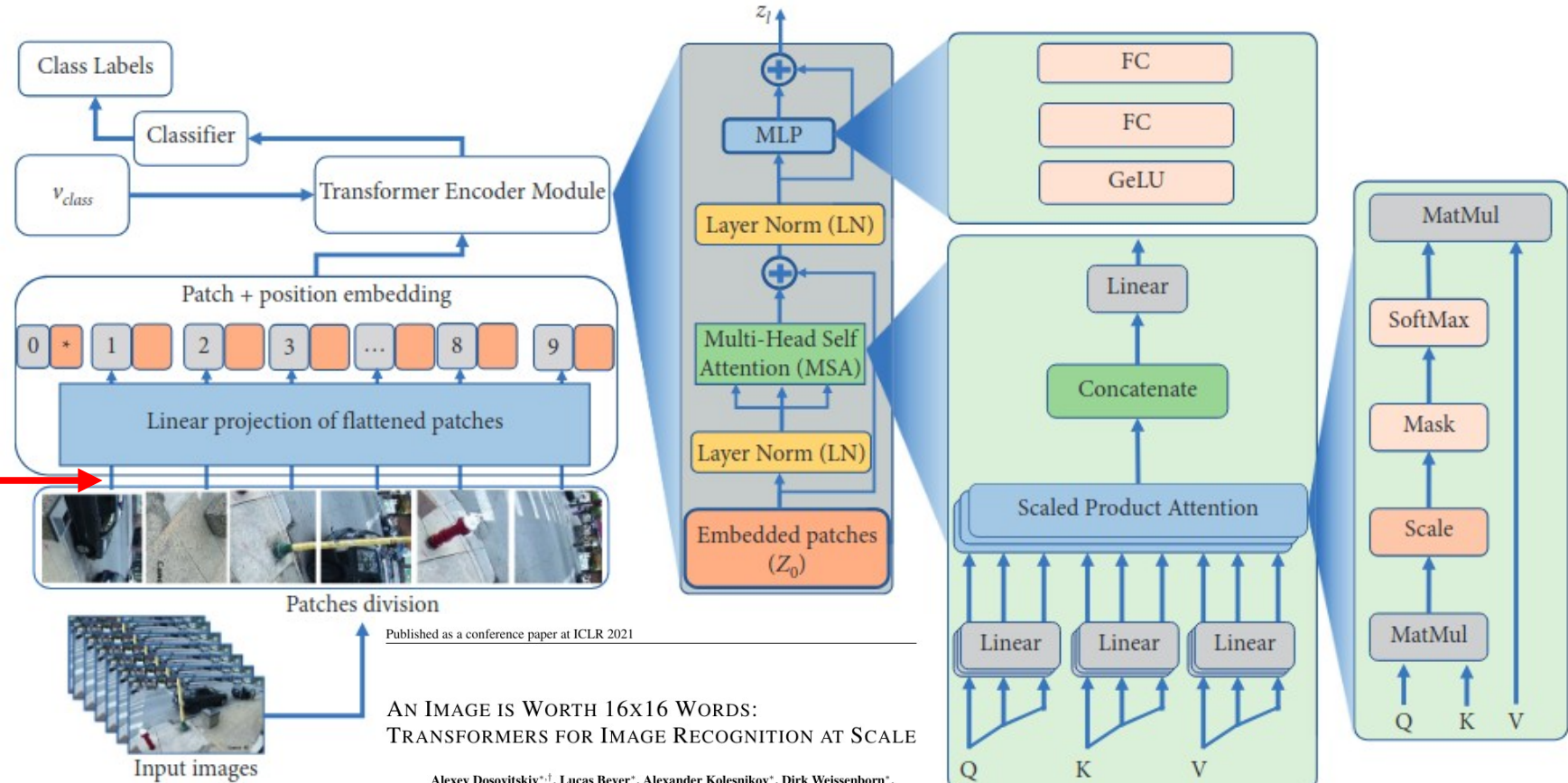


Levine *et al* (2024). Cell2Sentence:
Teaching Large Language Models
the Language of Biology. *BioRxiv*
<https://doi.org/10.1101/2023.09.11.557287>



Vision Transformer

Patch embedding by a CNN



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

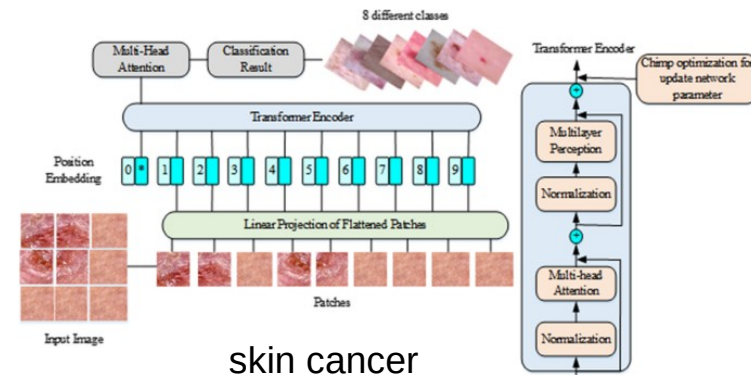
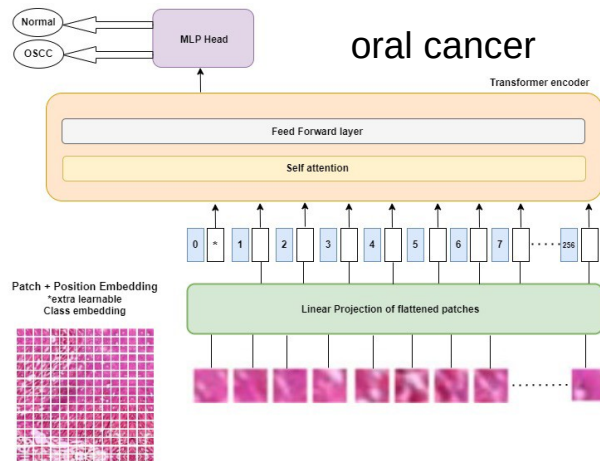
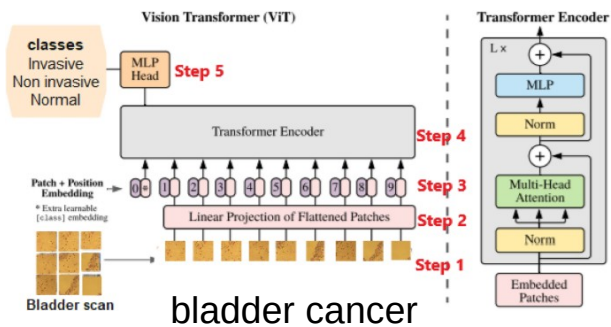
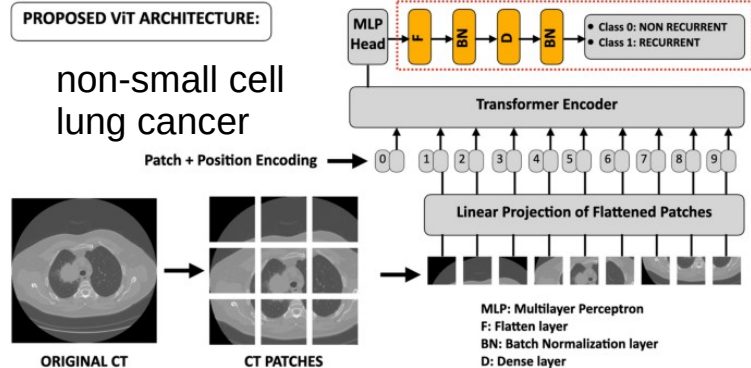
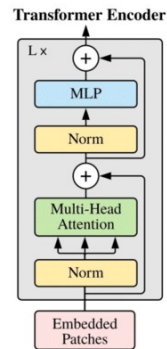
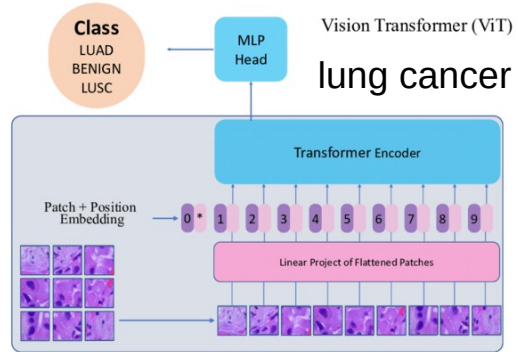
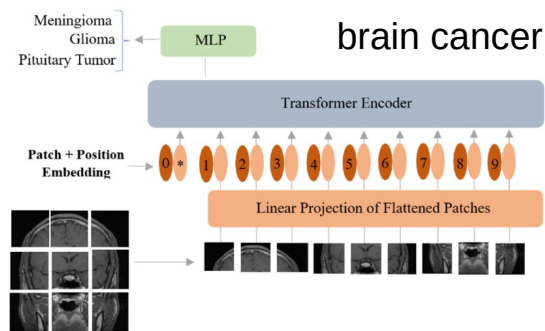
{adosovitskiy, neilhoulby}@google.com

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

CNN = local features

ViT = relations between distant features

ViTs are replacing vanilla CNNs

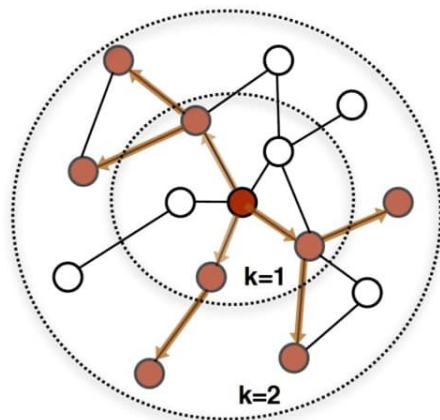


8

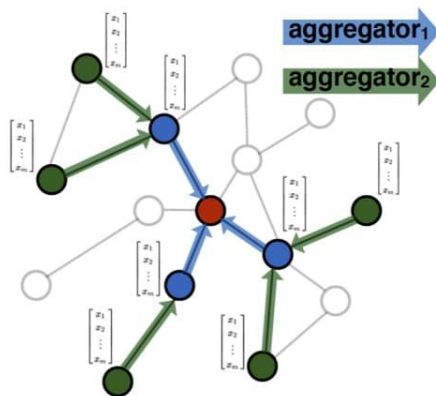
Graph Neural Networks (GNNs)

Graph Neural Networks (GNNs)

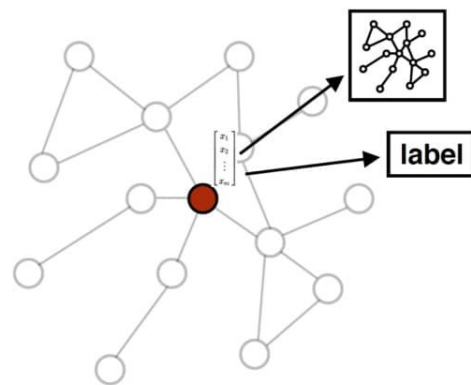
source: <https://blogs.nvidia.com/blog/what-are-graph-neural-networks/>



1. Sample neighborhood



2. Aggregate feature information from neighbors



3. Predict graph context and label using aggregated information

IEEE TRANSACTIONS ON NEURAL NETWORKS, VOL. 20, NO. 1, JANUARY 2009

61

The Graph Neural Network Model

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

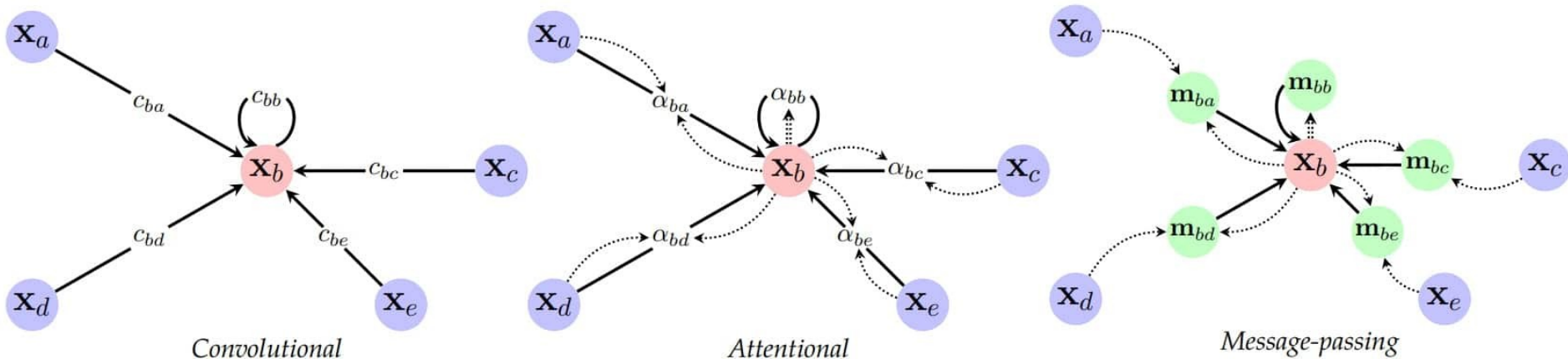
Convolutional Neural Networks on Graphs with Fast Localized Spectral Filtering

Michaël Defferrard Xavier Bresson Pierre Vandergheynst
EPFL, Lausanne, Switzerland
{michael.defferrard,xavier.bresson,pierre.vandergheynst}@epfl.ch

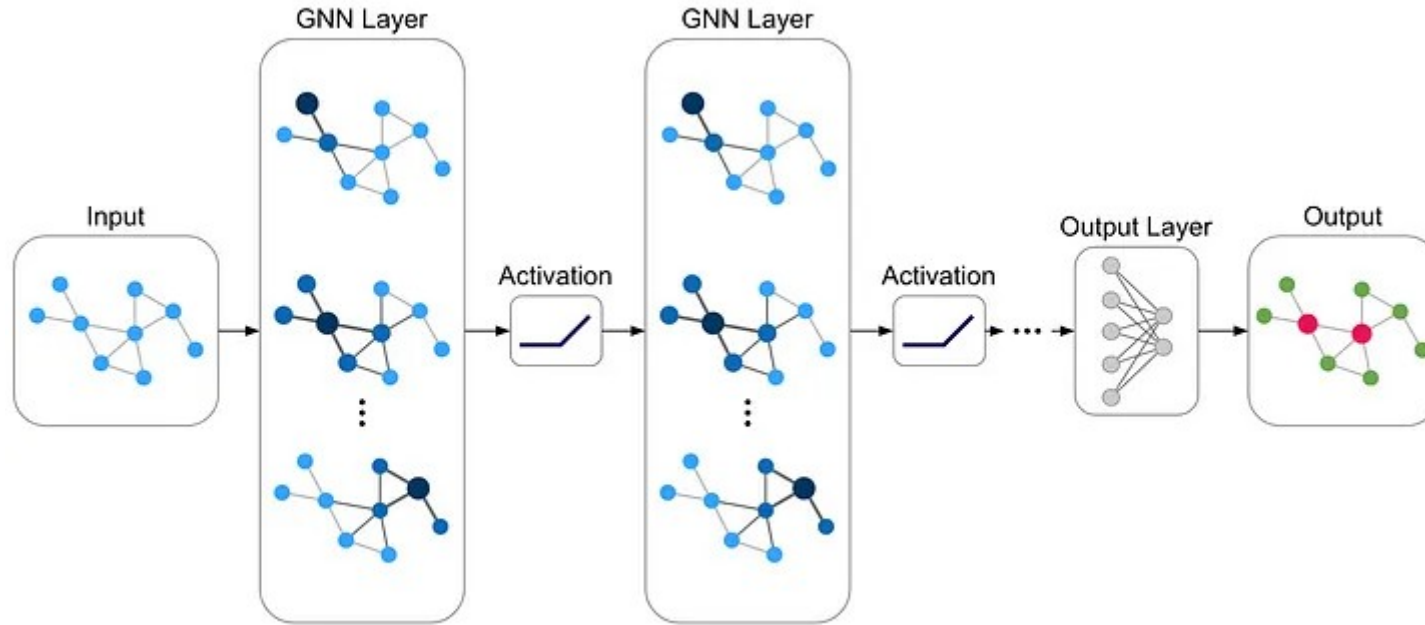
30th Conference on Neural Information Processing Systems (NIPS 2016), Barcelona, Spain.

Many different ways to update GNNs

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

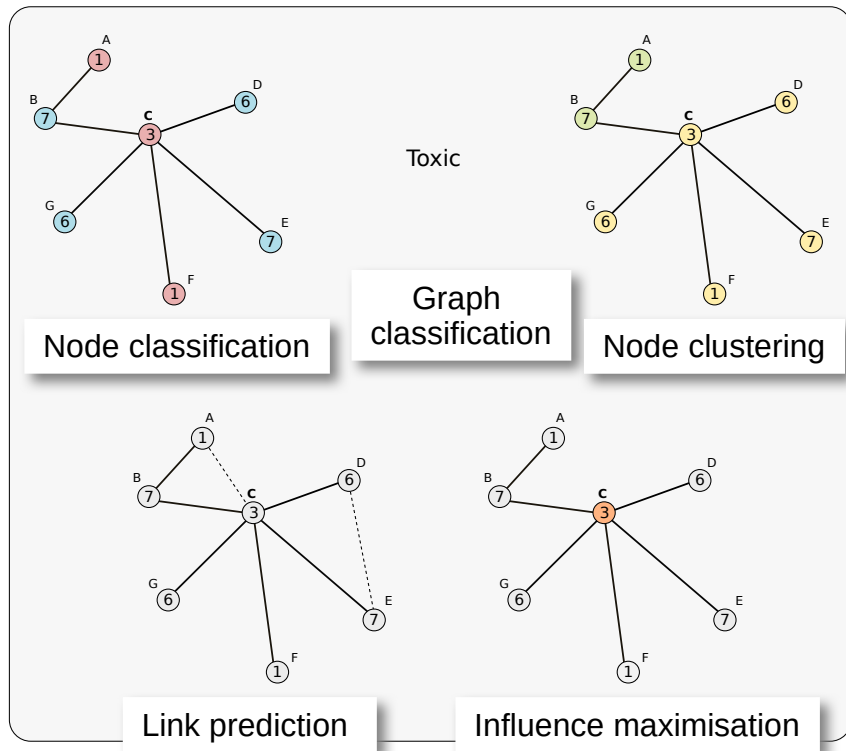
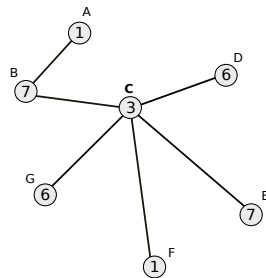


Graph Neural Networks (GNNs)



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

GNNs: What can we do?

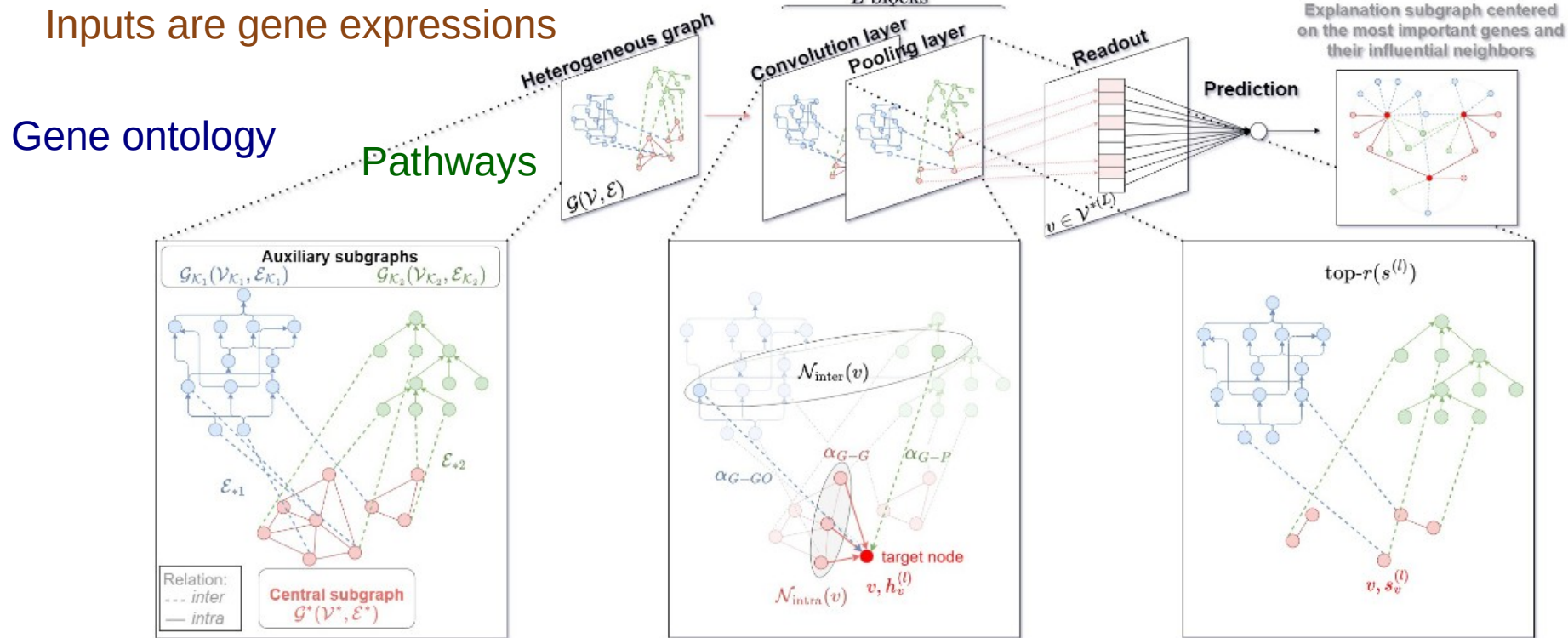


Source: Understanding Convolutions on Graphs
<https://distill.pub/2021/understanding-gnns/>

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

Both by Google Research teams

GNN can be heterogeneous



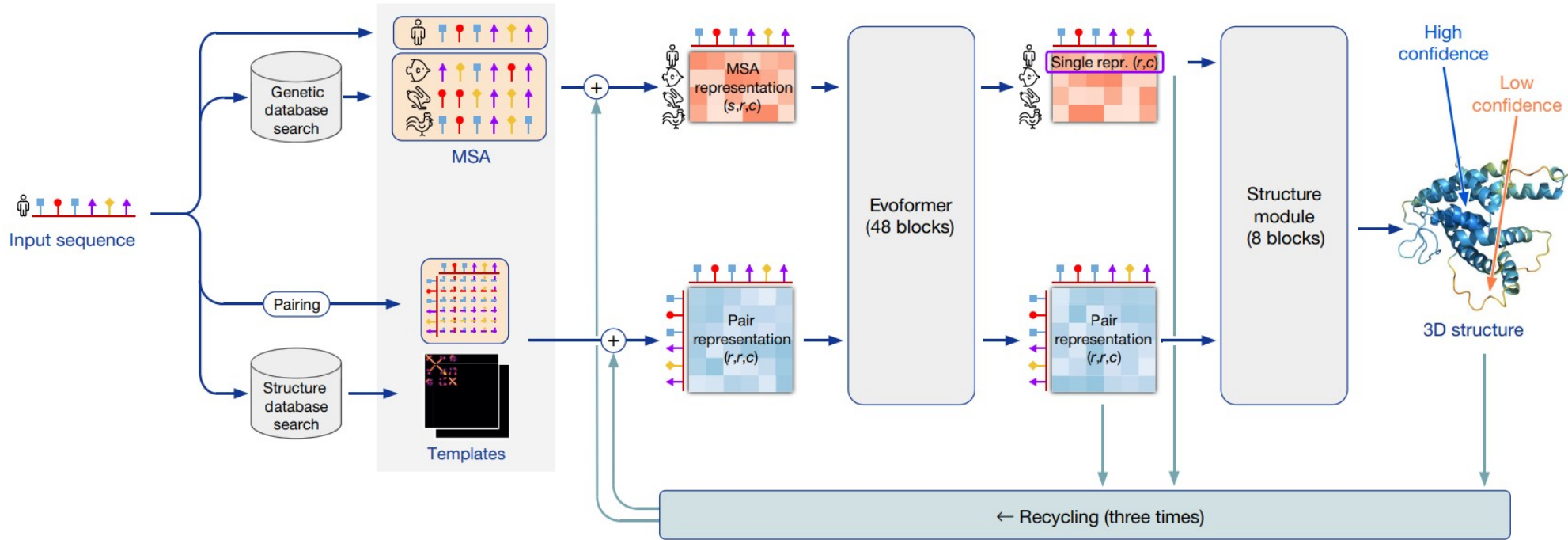
**BioHAN: a Knowledge-based Heterogeneous Graph
Neural Network for precision medicine on
transcriptomic data**

9

AlphaFold2



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

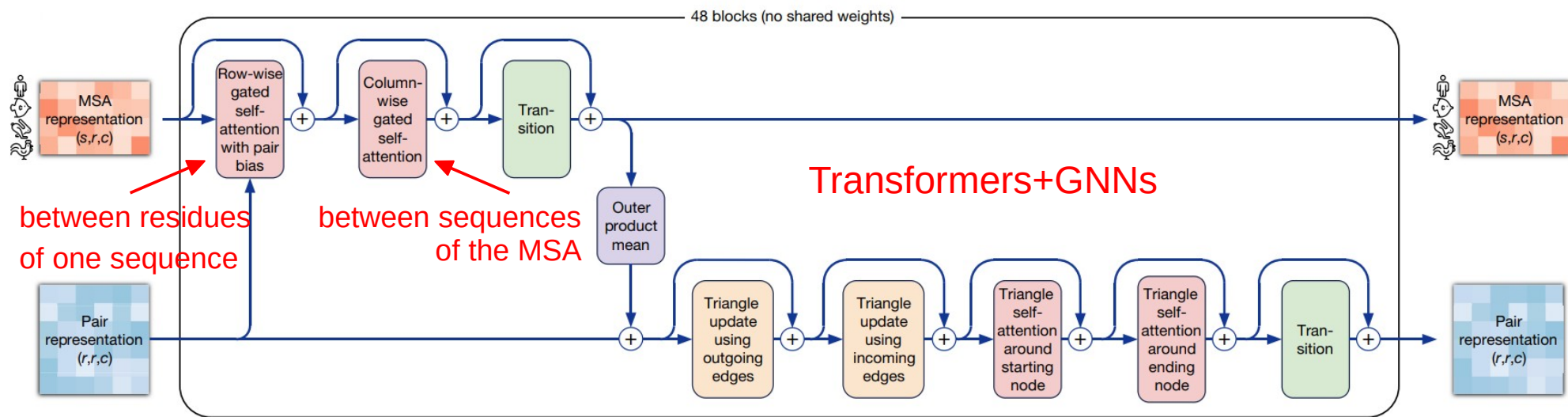
Check for updates

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}

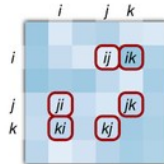
~93 million parameters (weights+biases)

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

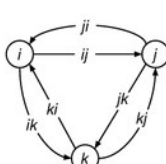
AlphaFold2: evoformer



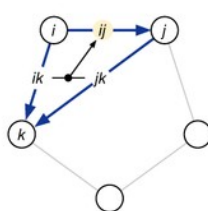
Pair representation (r, r, c)



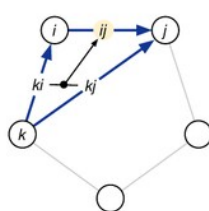
Corresponding edges in a graph



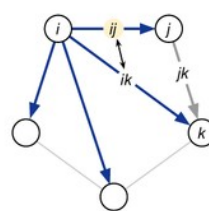
Triangle multiplicative update using 'outgoing' edges



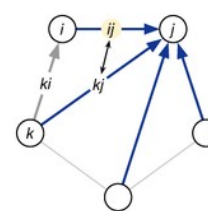
Triangle multiplicative update using 'incoming' edges



Triangle self-attention around starting node

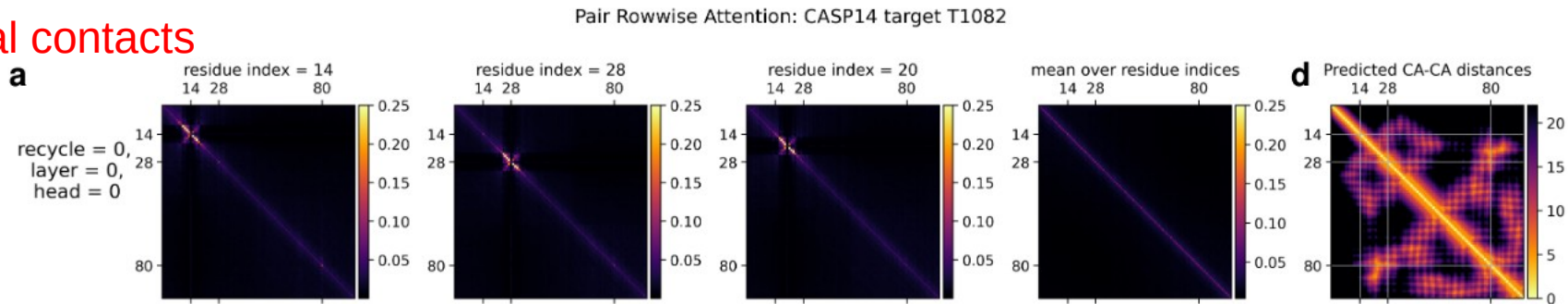


Triangle self-attention around ending node

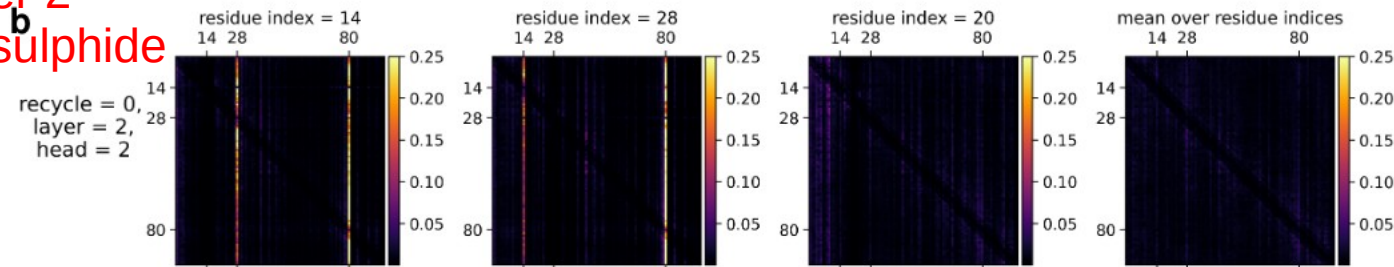


Row-wise attention: between residues of a sequence

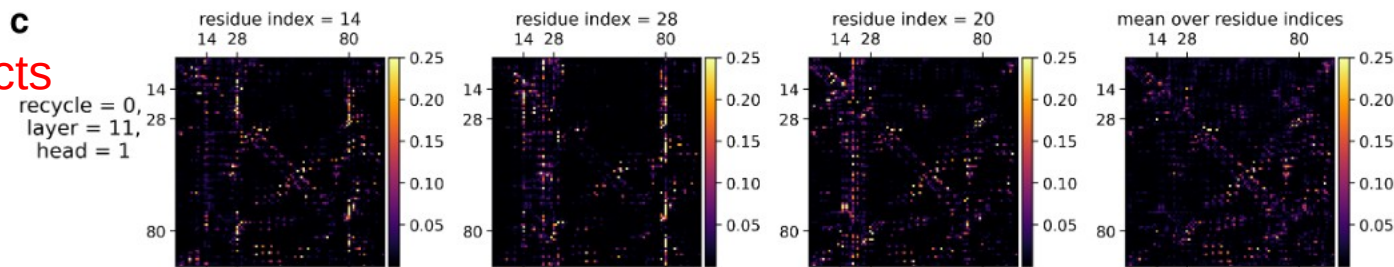
Layer 0 = local contacts



Head 2 of layer 2 recognises disulphide bonds



At layer 11, heads recognise distant contacts

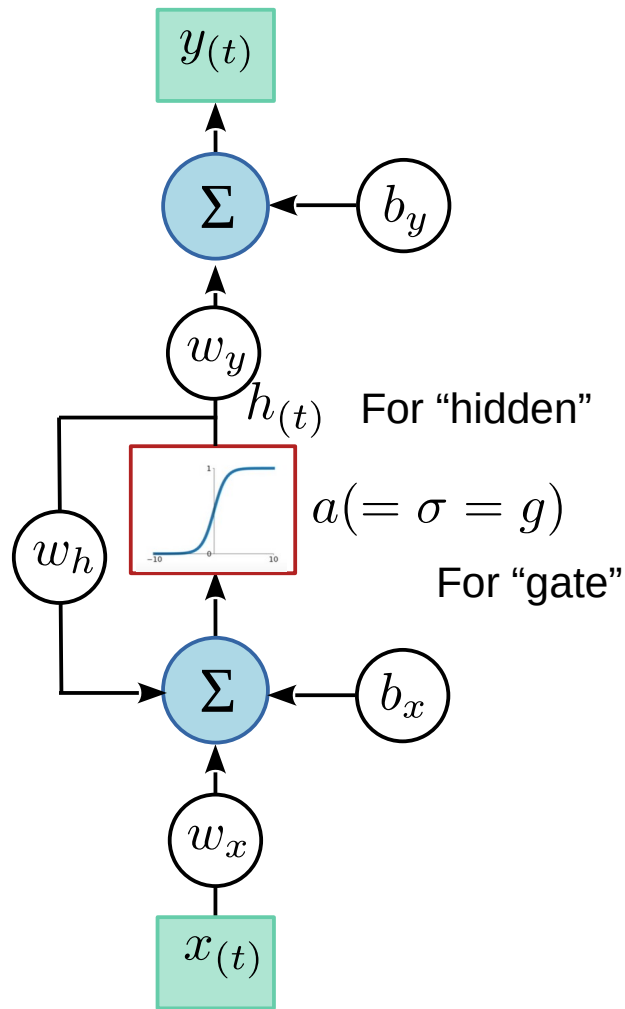


Source:
AlphaFold2 paper
supplementary info

10

Analysing and predicting series:
Recurrent Neural Networks (RNNs)
Long Short-Term Memory (LSTM)

RNN: 1 cell (here, 1 neuron)



NB: implicit "identity" activation function

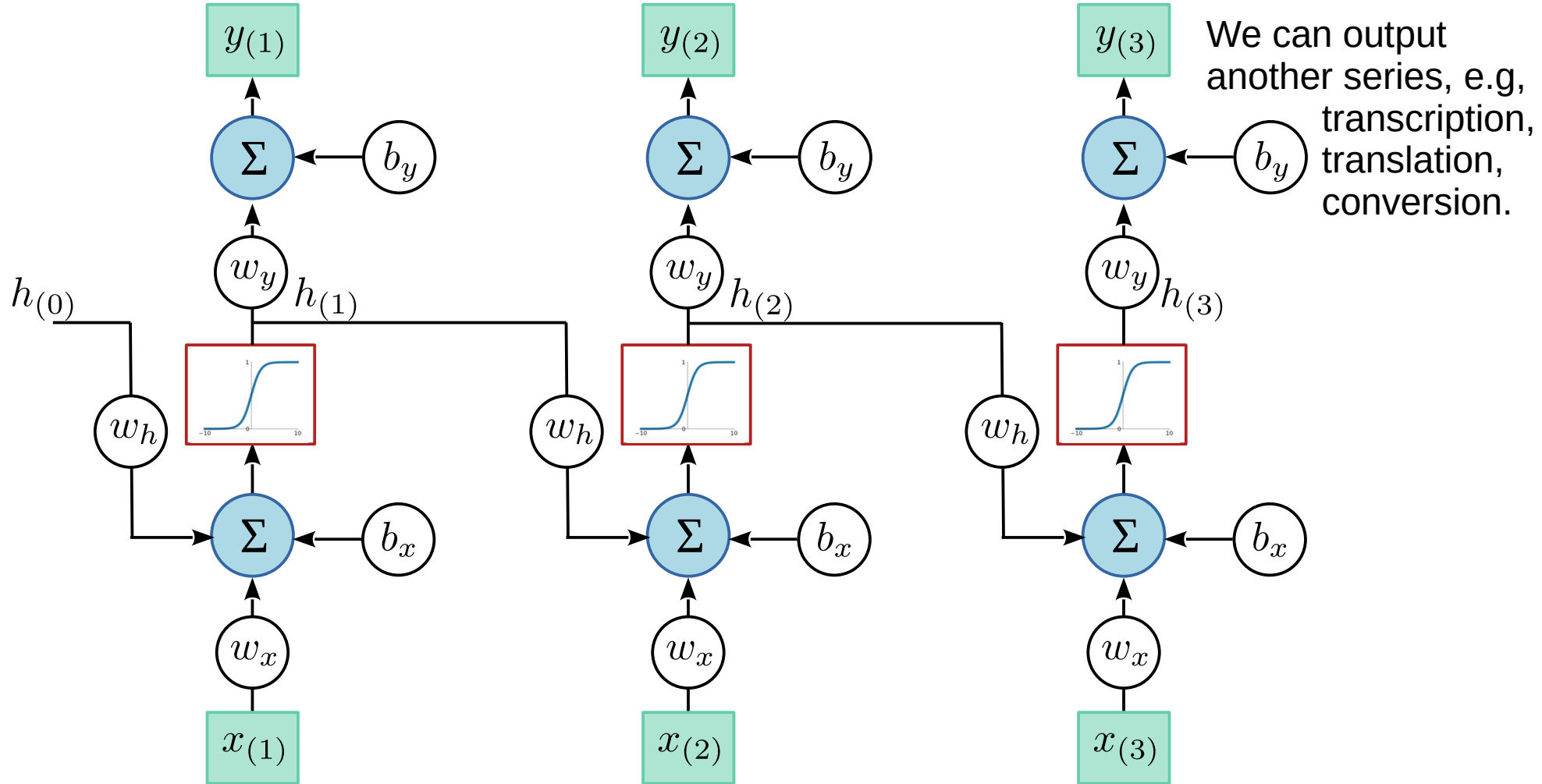
$$y(t) = w_y \times h(t) + b_y$$

not the same timepoint

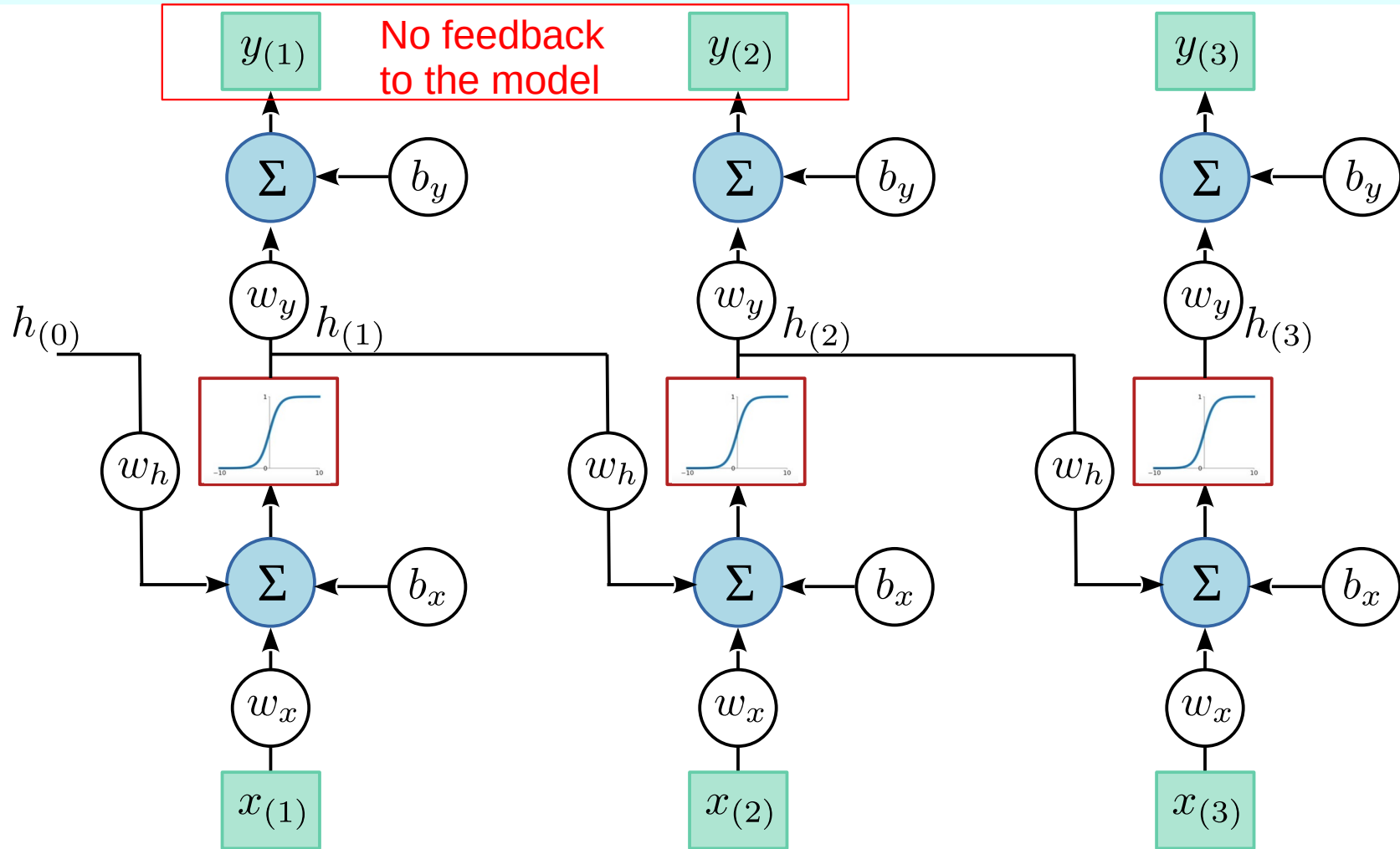
$$h(t) = a(w_x \times x(t) + b_x + w_h \times h_{(t-1)})$$

"memory"

RNN: 1 cell - unfolding

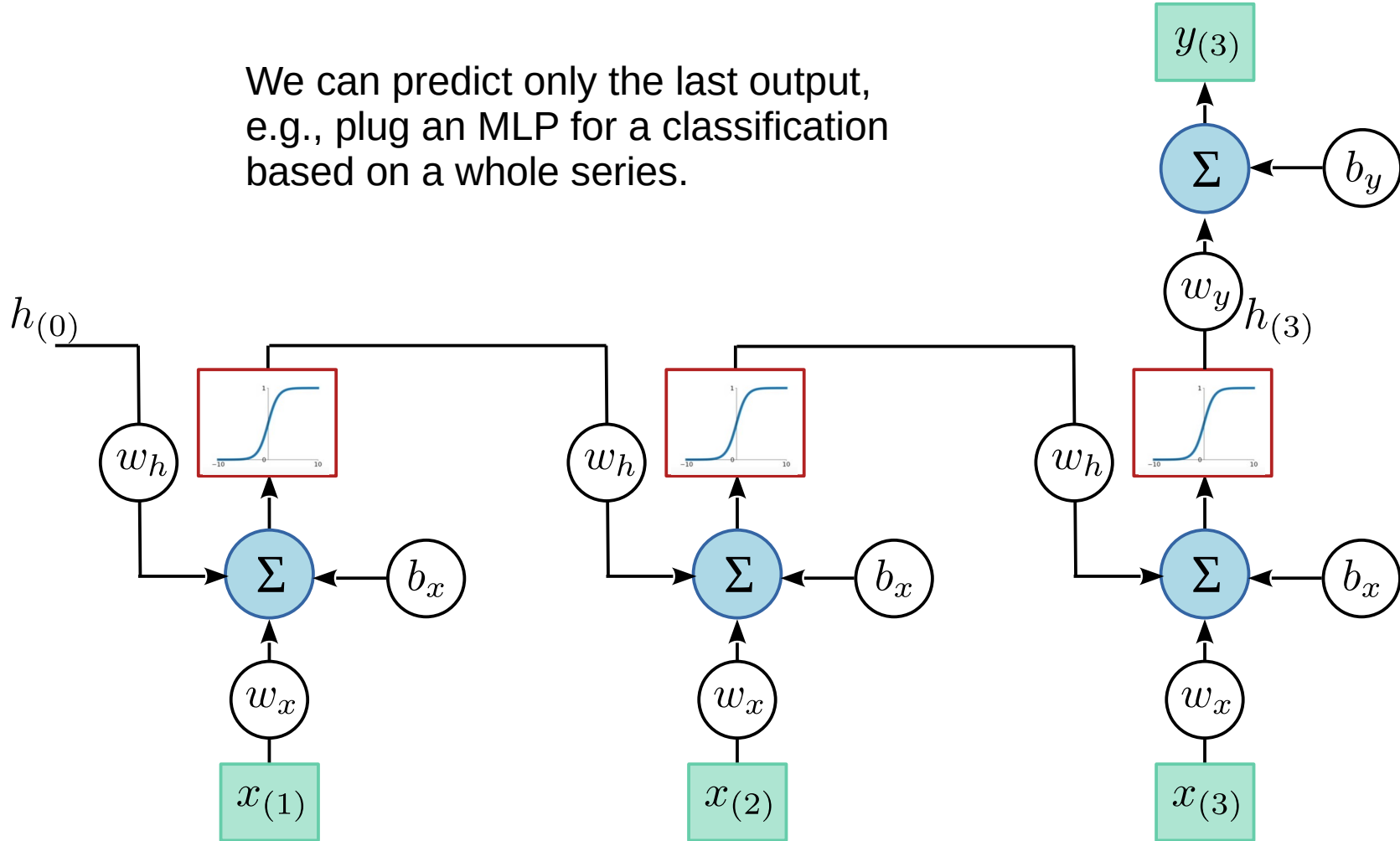


RNN: 1 cell - unfolding

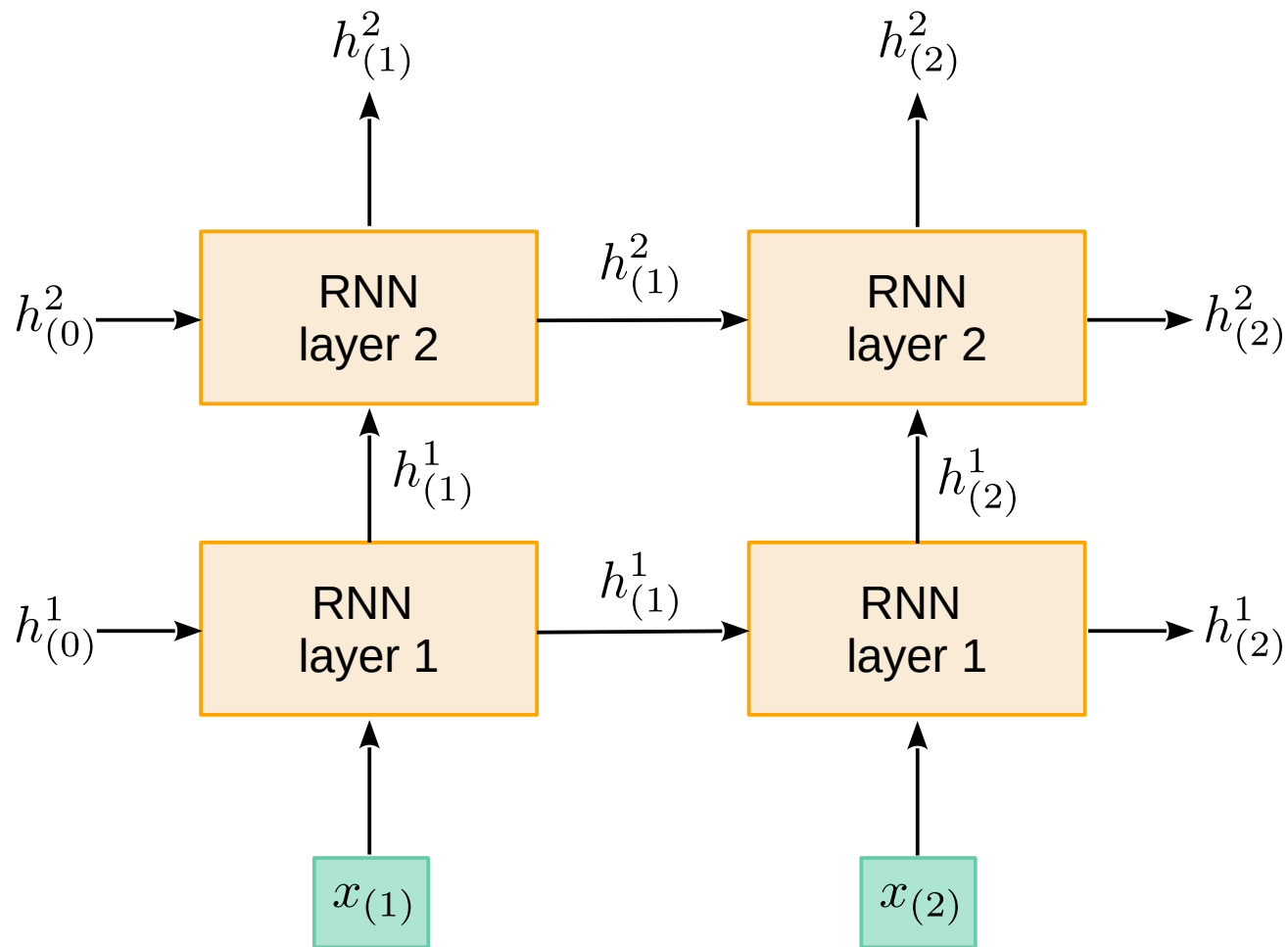


RNN: 1 cell - unfolding

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

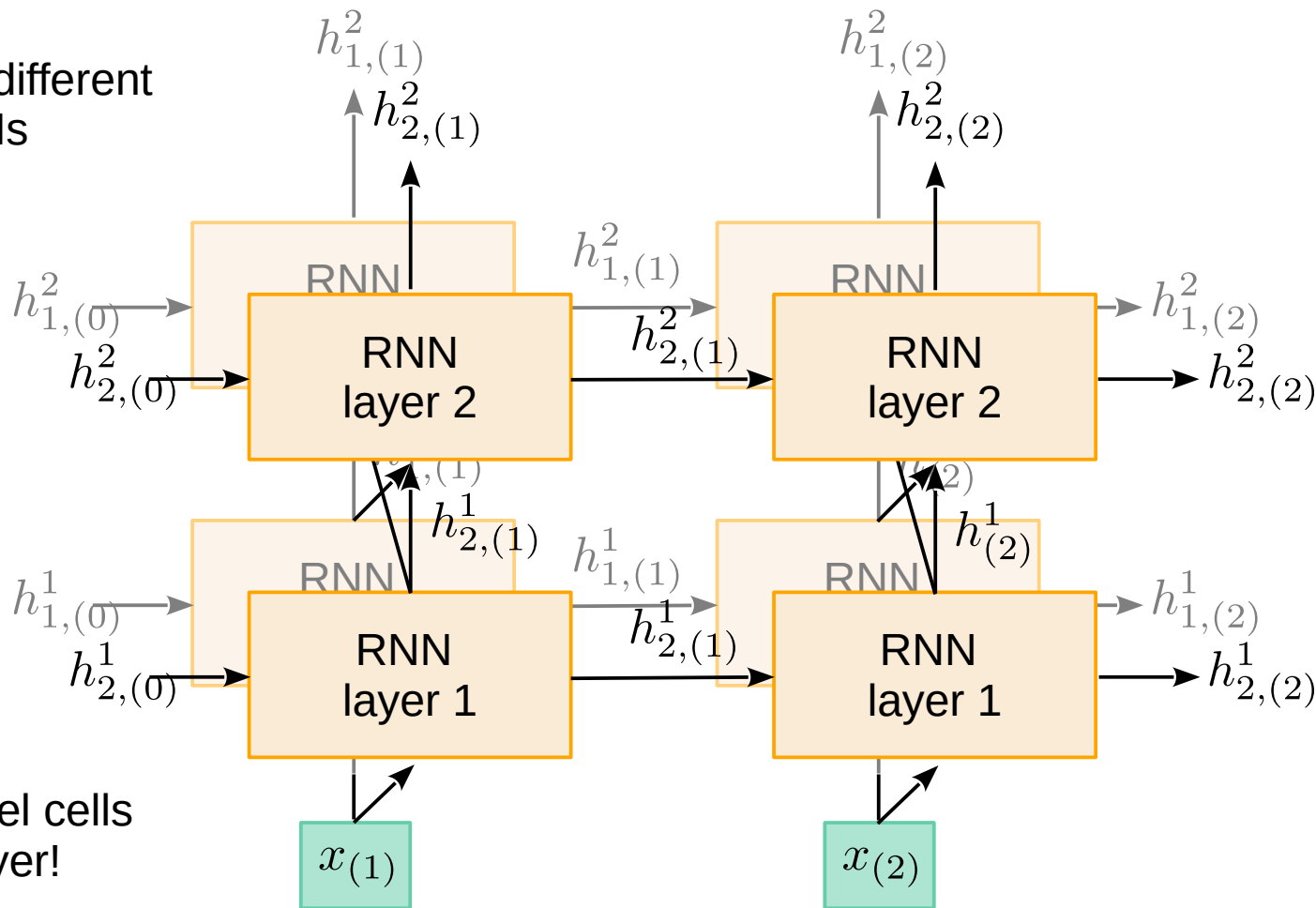


Stacked RNNs



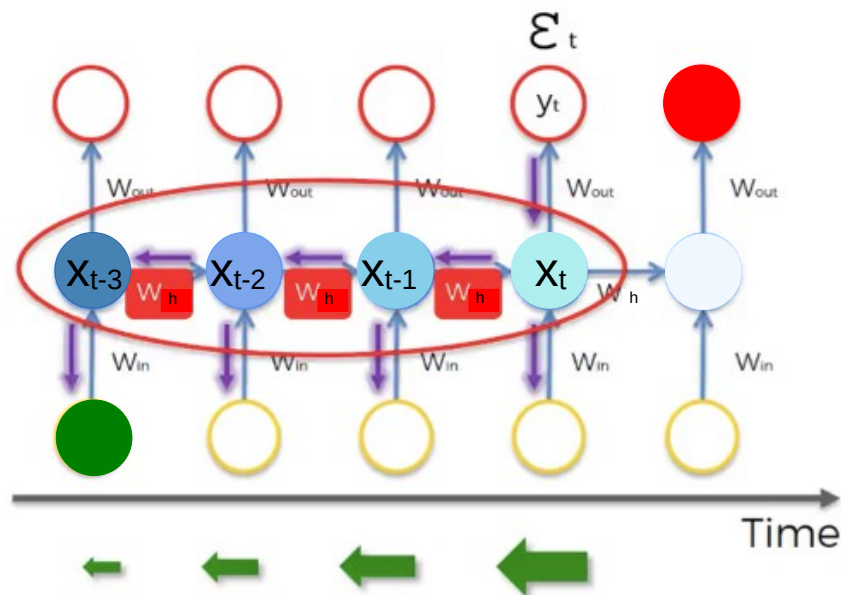
Several RNNs may learn different patterns in parallel

Same idea as different
kernels in CNNs



No connexions
between parallel cells
of the same layer!

Vanishing and exploding gradients: the network forgets



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

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

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

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

Adapted from:
SuperDataScience

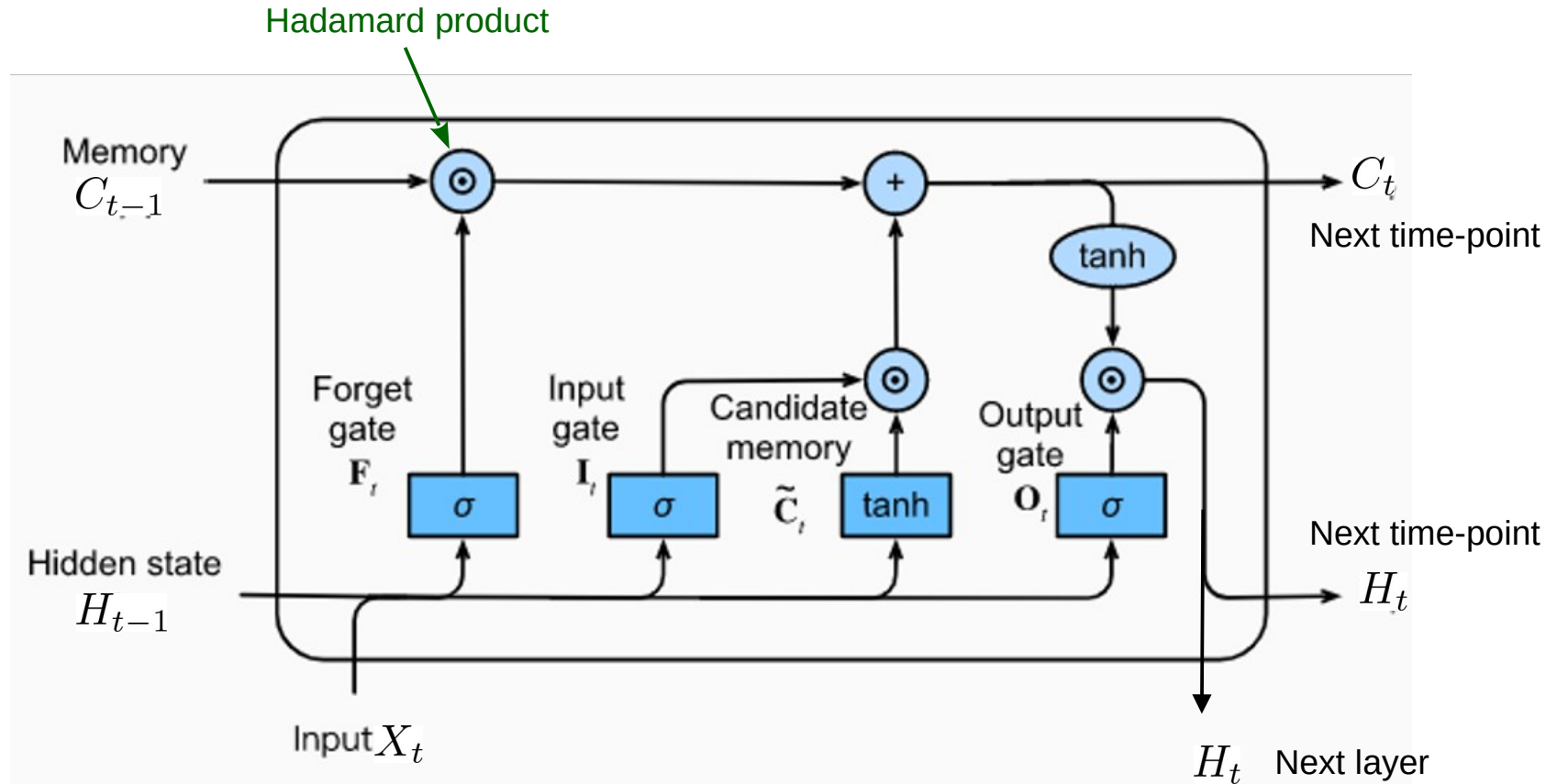
E.g.: activation function = ReLU

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

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

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

Solution: Long Short-Term Memory (LSTM)

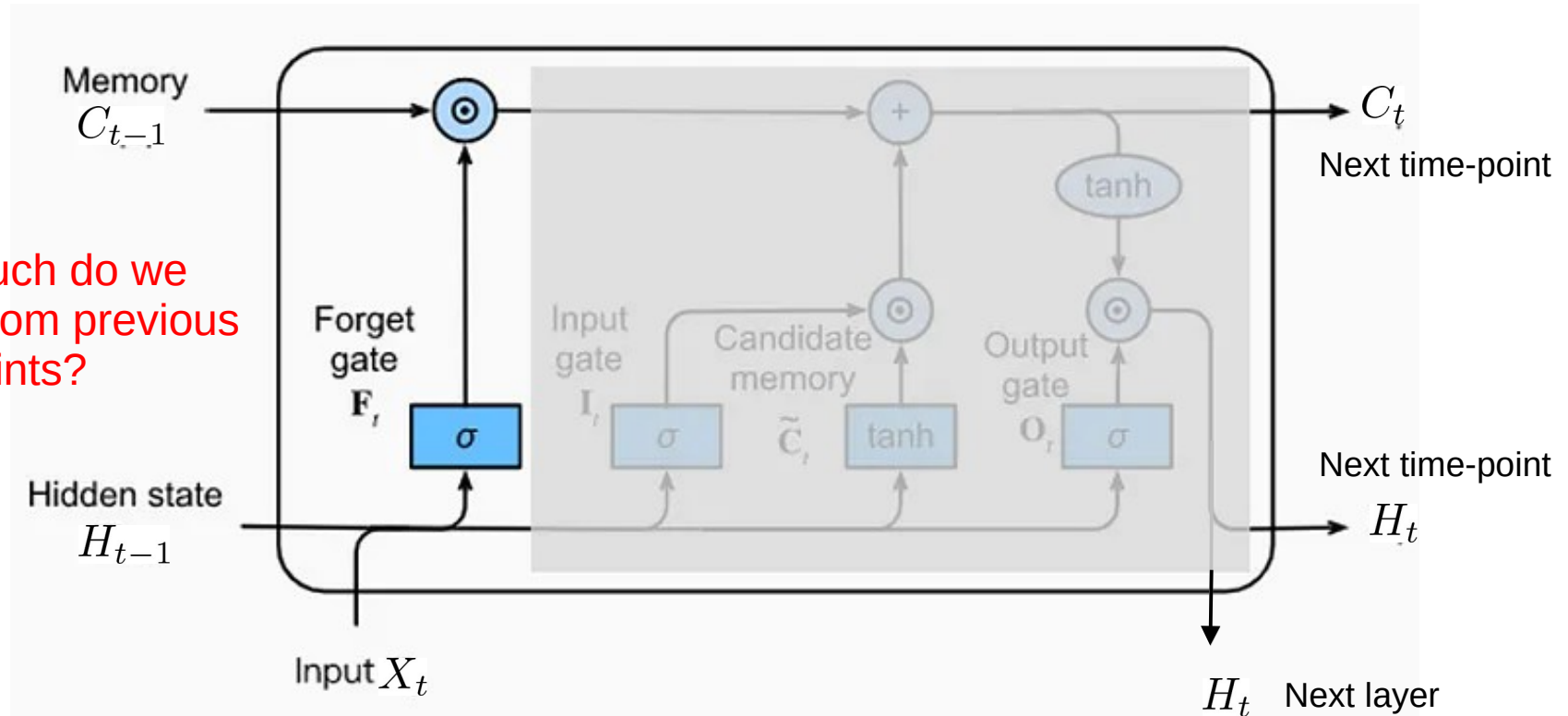


Hochreiter and Schmidhuber (1997) Long short-term memory. *Neur Comput*, 9(8):1735-1780

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

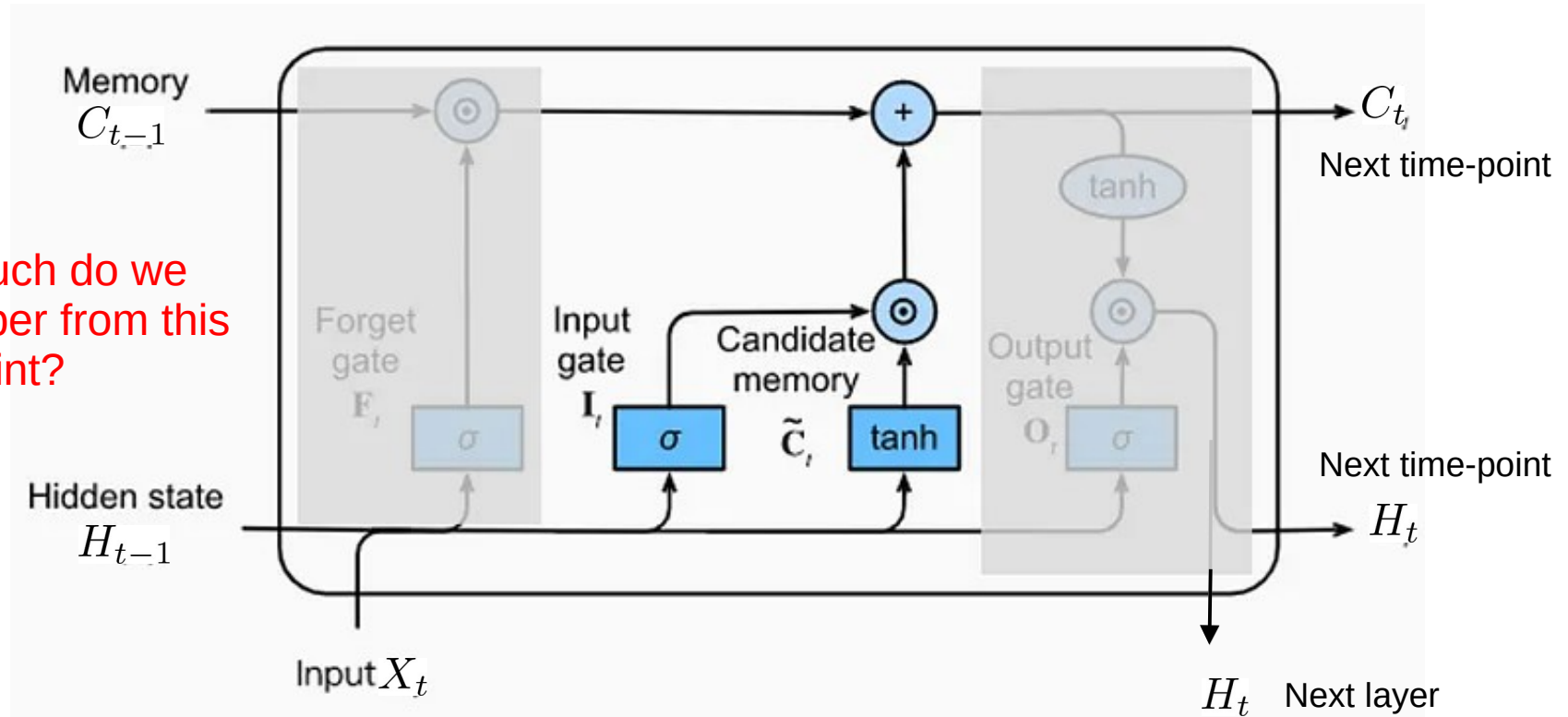
Solution: Long Short-Term Memory (LSTM)

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



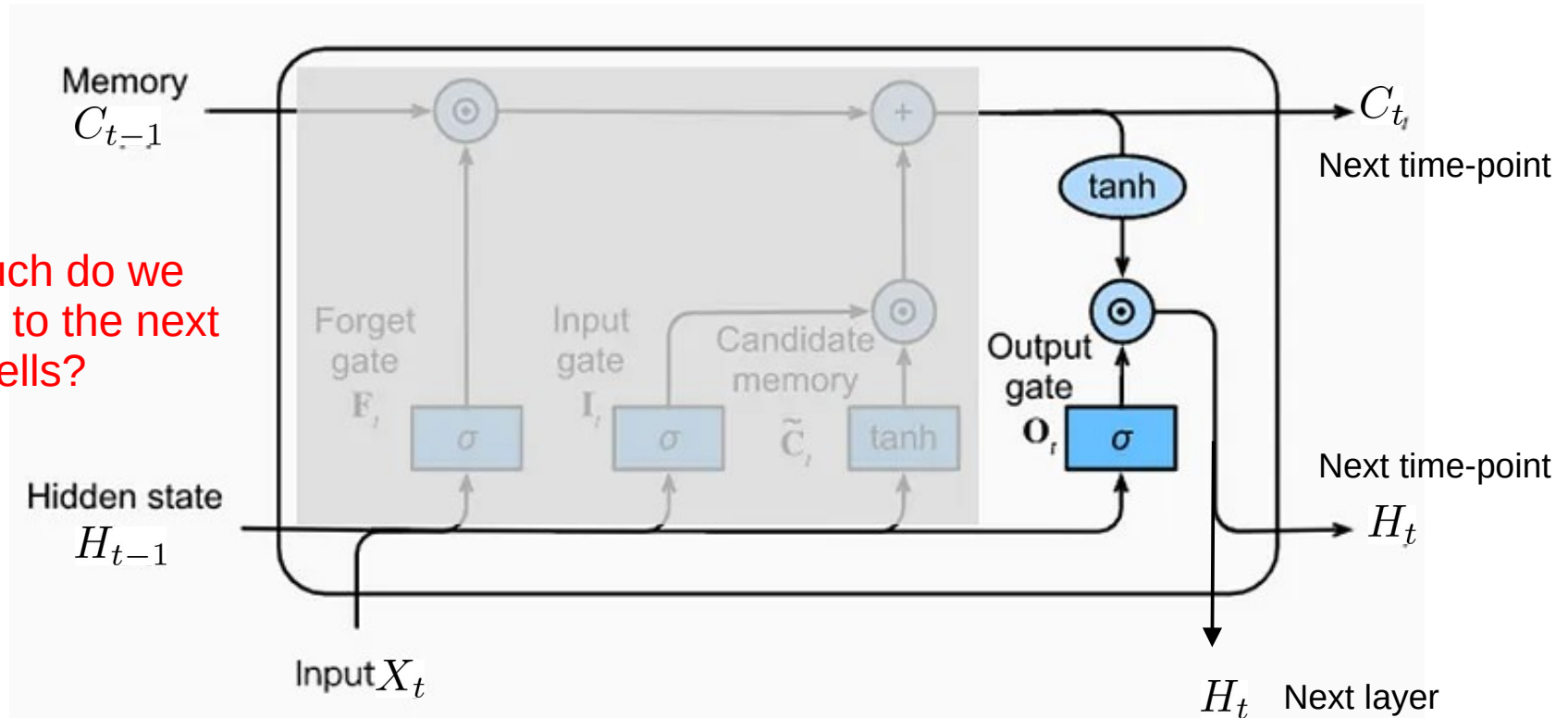
Solution: Long Short-Term Memory (LSTM)

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



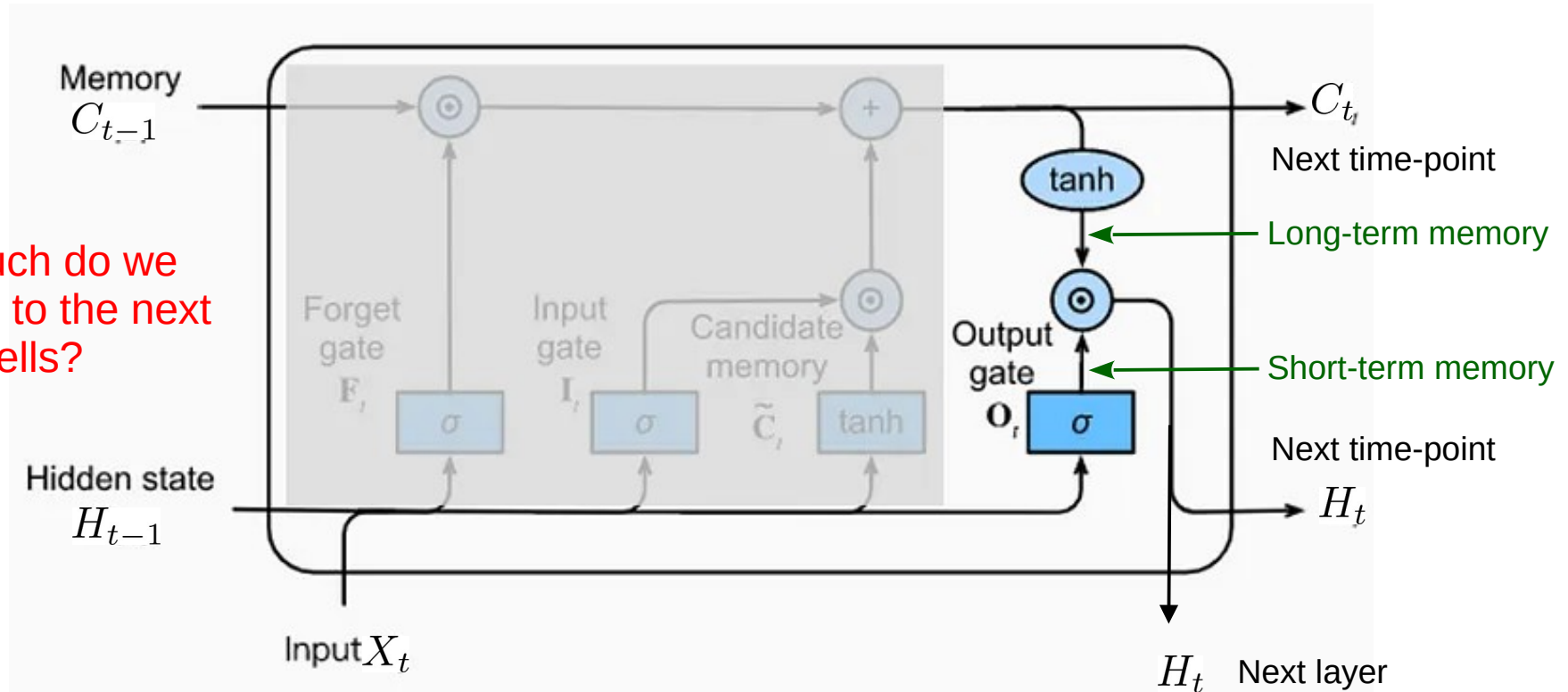
Solution: Long Short-Term Memory (LSTM)

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

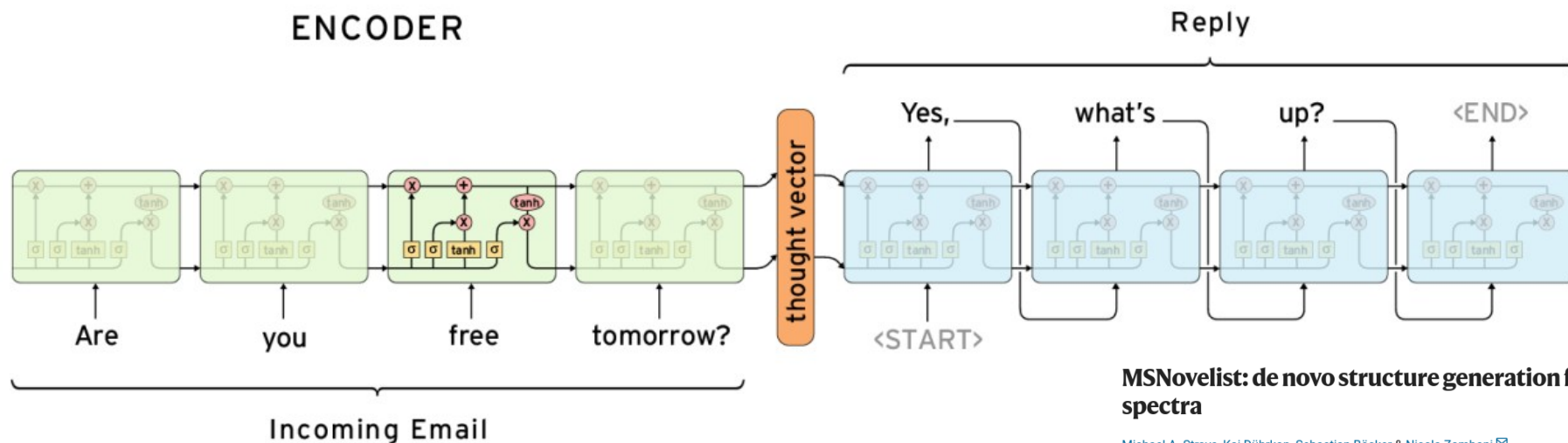


Solution: Long Short-Term Memory (LSTM)

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



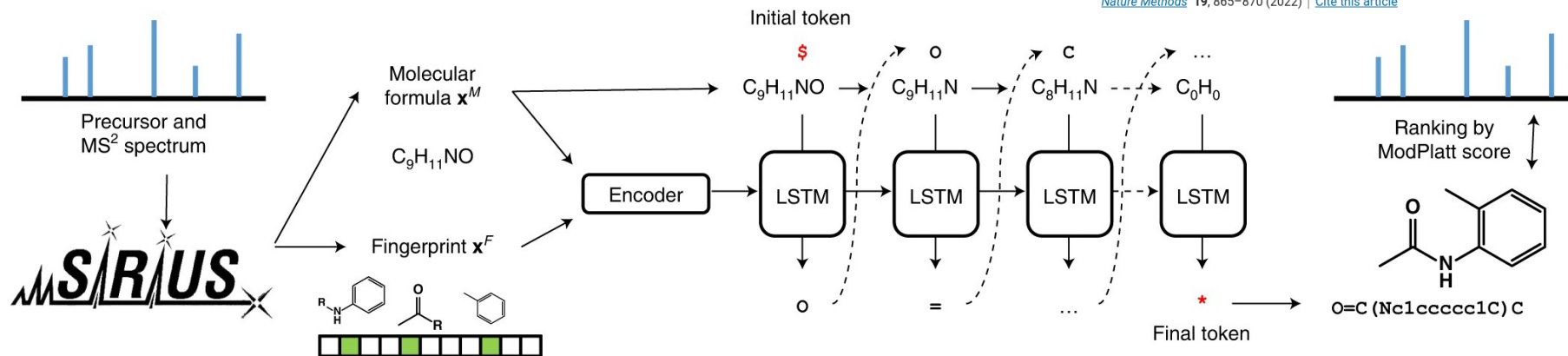
LSTMs for Encoder-Decoder



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

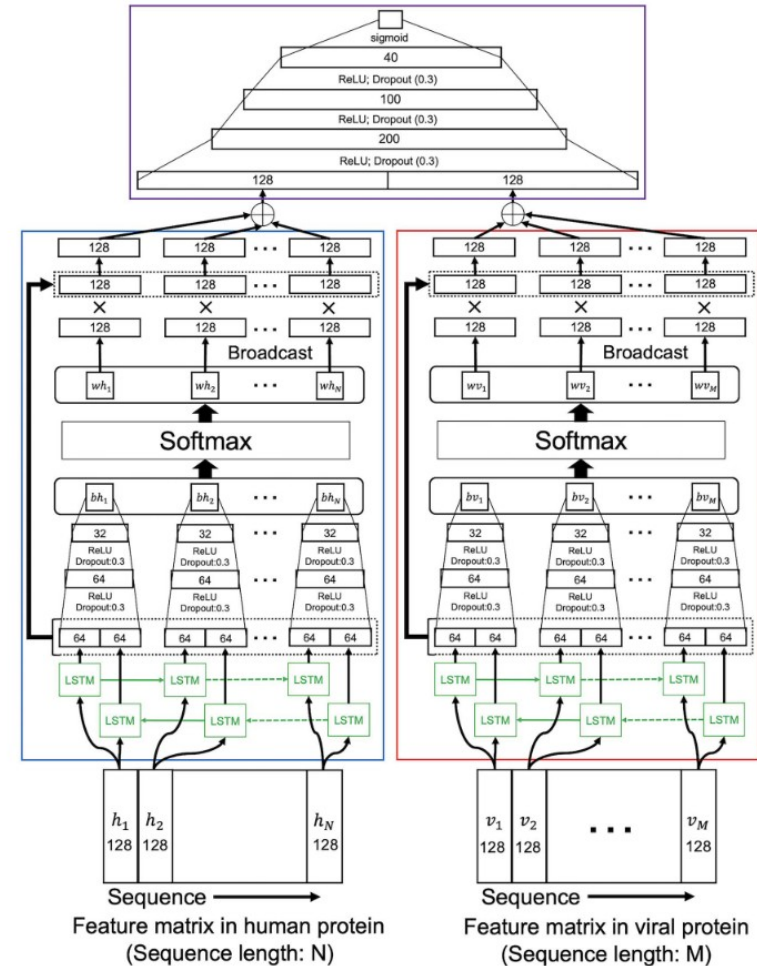
OXFORD

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

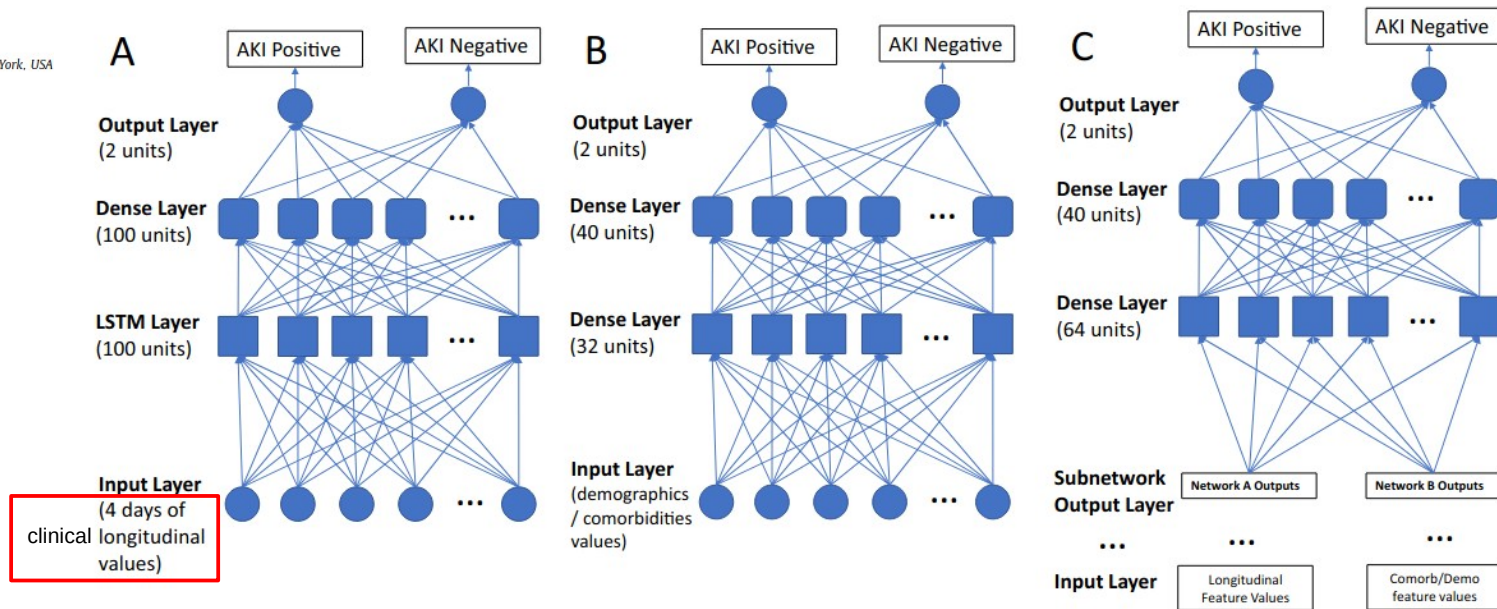
journal homepage: www.elsevier.com/locate/ijid



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



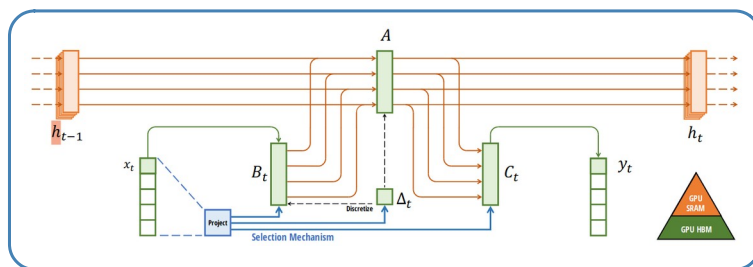
11

RNNs are back
Rise of the Mamba



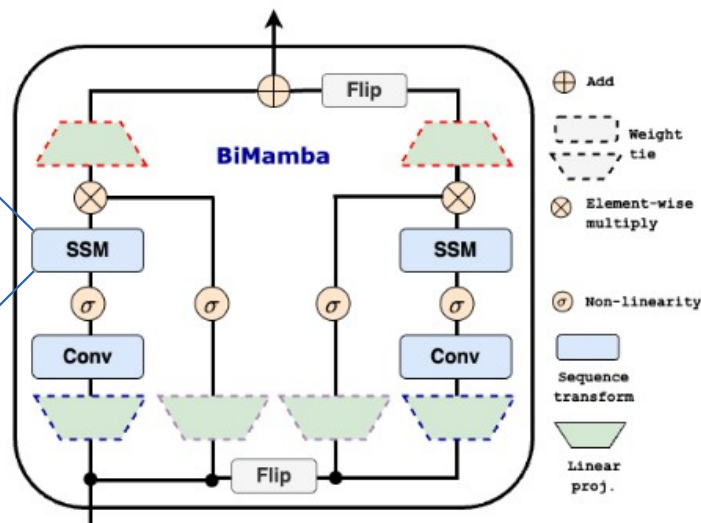
Selective State Space Model

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



SSM = State Space Model

Prediction/classification



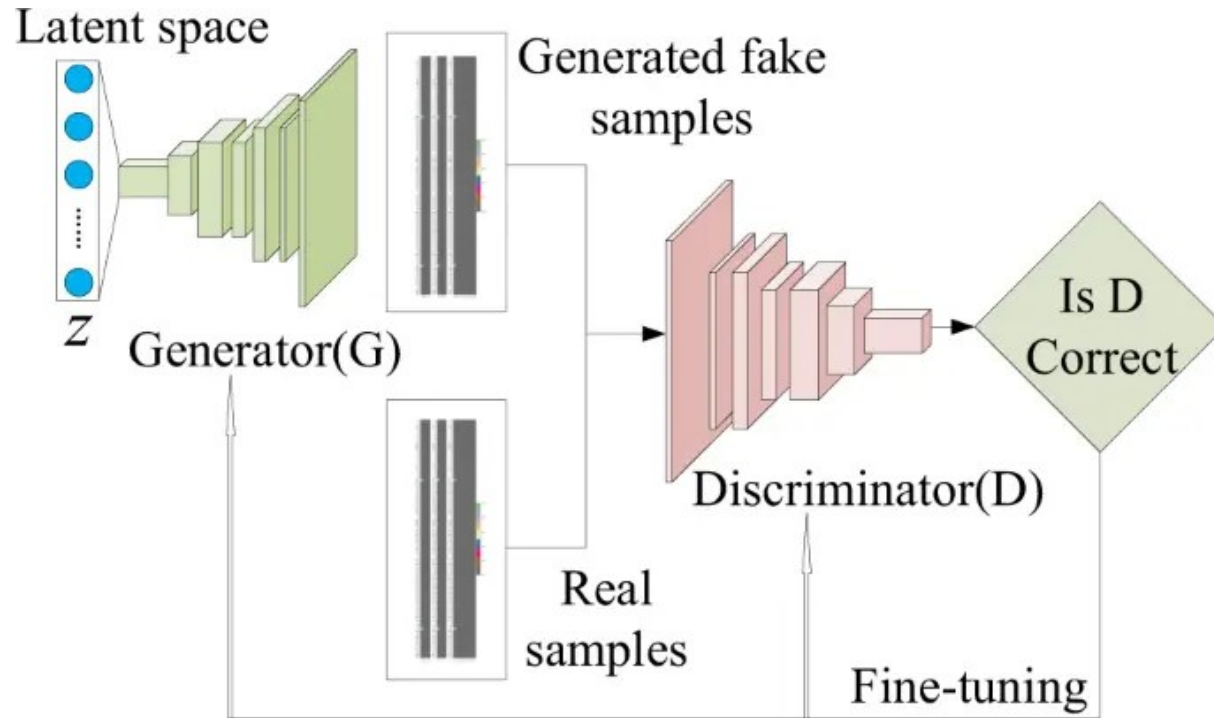
...AGGCTAGGATATCGATAAGCTGACTGAT...

The model
 learns about
 variants'
 Interactions
 → Reusable
 foundation
 model

12

Generative Adversarial Networks (GANs)

Learning by trying to trick itself: Generative Adversarial Network (GAN)



The Generator gets ever better at generating good fakes
The Discriminator gets ever better at detecting them

GAN for synthetic data



Neurocomputing

Volume 321, 10 December 2018, Pages 321-331



GAN-based synthetic medical image augmentation for increased CNN performance in liver lesion classification

[Maayan Frid-Adar^a](#), [Idit Diamant^a](#), [Eyal Klang^b](#), [Michal Amitai^b](#), [Jacob Goldberger^c](#), [Hayit Greenspan^a](#) ✉

Brain tumor image generation using an aggregation of GAN models with style transfer

[Debaditya Mukherjee](#), [Pritam Saha](#), [Dmitry Kaplun](#) ✉, [Aleksandr Sinitca](#) & [Ram Sarkar](#)

[Scientific Reports](#) 12, Article number: 9141 (2022) | [Cite this article](#)



Computer Methods and Programs in
Biomedicine

Volume 195, October 2020, 105568



A GAN-based image synthesis method for skin lesion classification

[Zhiwei Qin^a](#), [Zhao Liu^b](#) ✉, [Ping Zhu^a](#) ✉, [Yongbo Xue^a](#)

[Home](#) > [Proceedings of International Conference on Artificial Intelligence and Applications](#) >
Conference paper

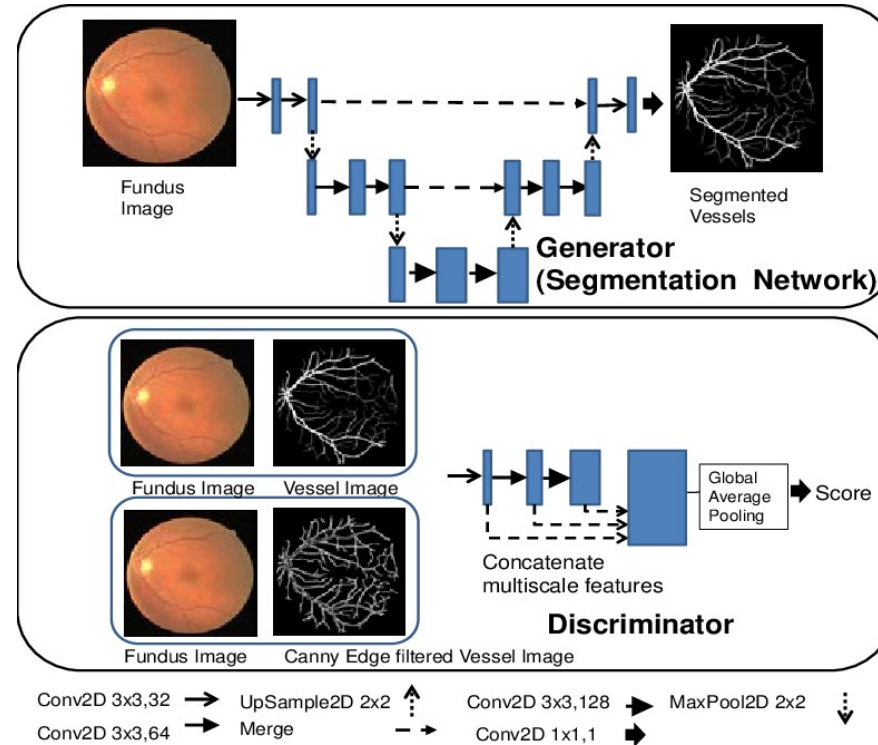
MR Image Synthesis Using Generative Adversarial Networks for Parkinson's Disease Classification

Conference paper | First Online: 02 July 2020

pp 317–327 | [Cite this conference paper](#)

[Sukhpal Kaur](#) ✉, [Himanshu Aggarwal](#) & [Rinkle Rani](#)

GAN for segmentation



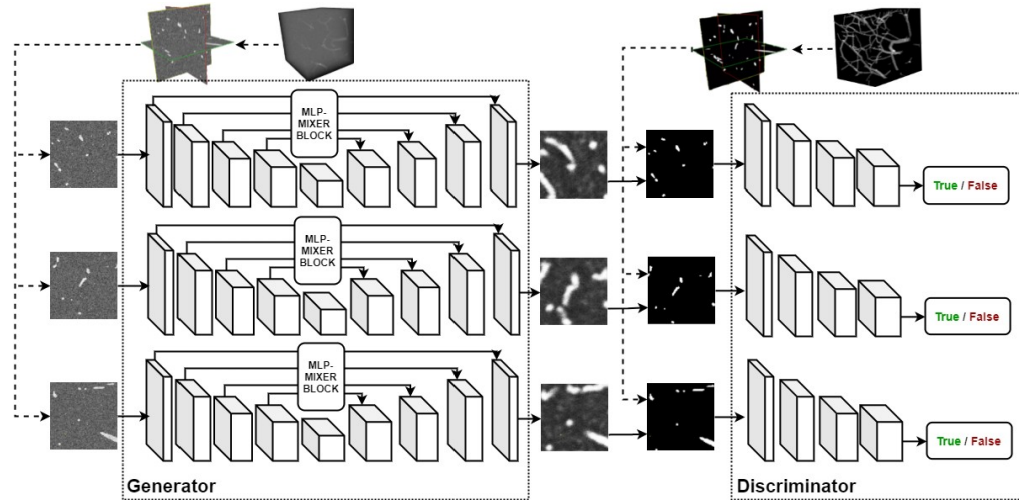
Tjio, G., Li, S., Xu, X., Ting, D.S.W., Liu, Y., Goh, R.S.M. (2019).

Multi-discriminator Generative Adversarial Networks for Improved Thin Retinal Vessel Segmentation.

In: Fu, H., Garvin, M., MacGillivray, T., Xu, Y., Zheng, Y. (eds) Ophthalmic Medical Image Analysis. OMIA 2019.

Lecture Notes in Computer Science(), vol 11855. Springer, Cham. https://doi.org/10.1007/978-3-030-32956-3_18

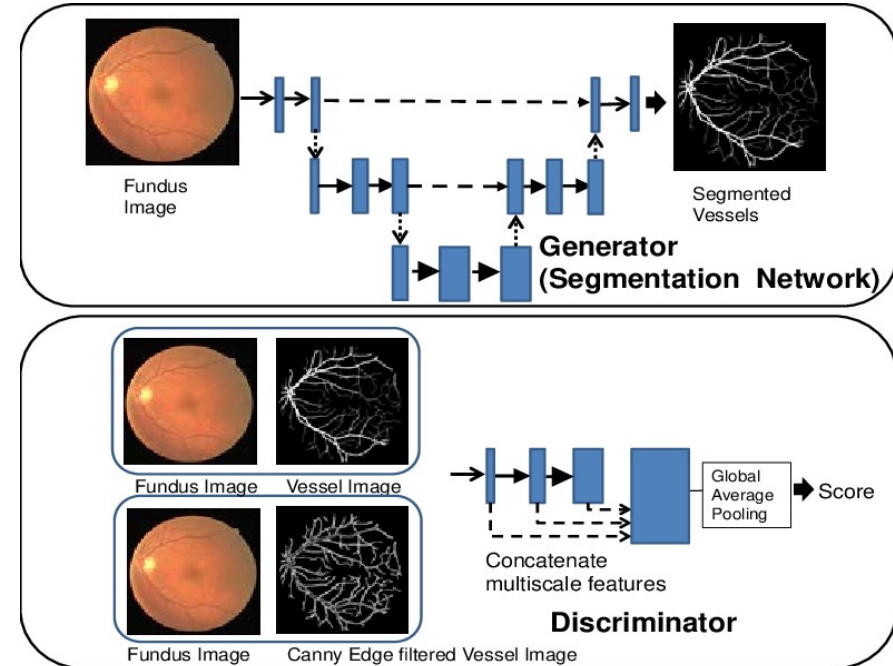
GAN can be built out of any DL architecture



DOI: 10.1109/ICASSP49357.2023.10096997 • Corpus ID: 250626500

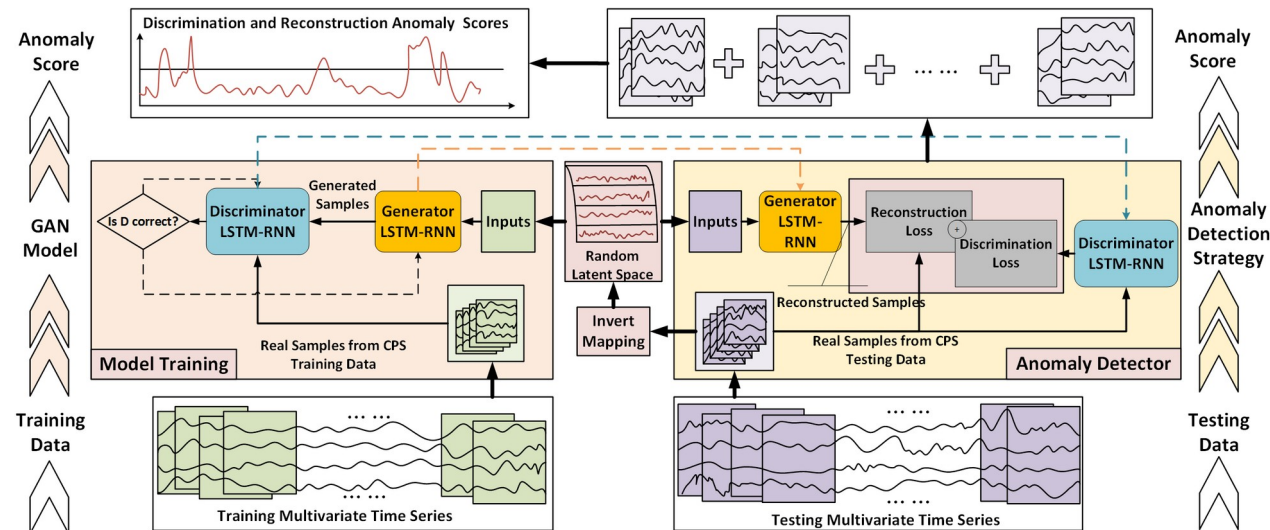
MLP-GAN for Brain Vessel Image Segmentation

B. Xie, Hao Tang, +2 authors Yan Yan • Published in *IEEE International Conference...* 17 July 2022 • Computer Science, Medicine



Conv2D 3x3,32 → UpSample2D 2x2
 Conv2D 3x3,64 → Merge
 Conv2D 3x3,128 → MaxPool2D 2x2
 Conv2D 1x1,1 →

GAN can be built out of any DL architecture

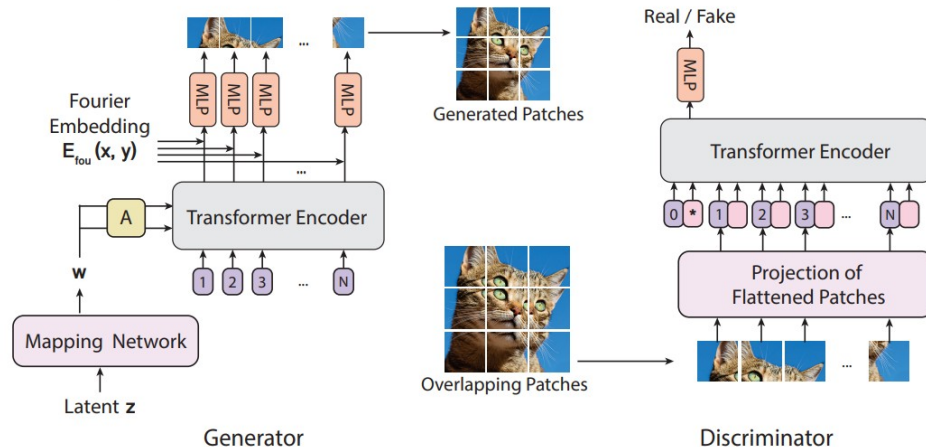


ViTGAN: Training GANs with Vision Transformers

Kwonjoon Lee¹ Huiwen Chang² Lu Jiang² Han Zhang²
 Zhuowen Tu¹ Ce Liu²
¹UC San Diego ²Google Research
 {kw1042,ztu}@ucsd.edu {huiwenchang,lujiang,zhanghan,celiu}@google.com

MAD-GAN: Multivariate Anomaly Detection for Time Series Data with Generative Adversarial Networks

Dan Li¹, Dacheng Chen¹, Lei Shi¹, Baihong Jin², Jonathan Goh³, and See-Kiong Ng¹



Acknowledgements

Philippe Froguel
Amélie Bonnefond

Smaïn Fettes

Mathilde Boissel
Lijiao Ning

Emma Henriques
ShuangShuang Geng

Anne-Sophie Ledoux
Vincent Massy

Amna Khamis
Stefan Gaget
Mehdi Derhourhi

Hélène de Gavre
Mélanie Hocquet

François Pattou

R package developers
(*VIM, MICE, DESeq2, ChAMP*)

Python package developers
(*Tensorflow/Keras, Numpy, Pandas, Scikit-learn*)