



Deep Learning for Genomic Prediction



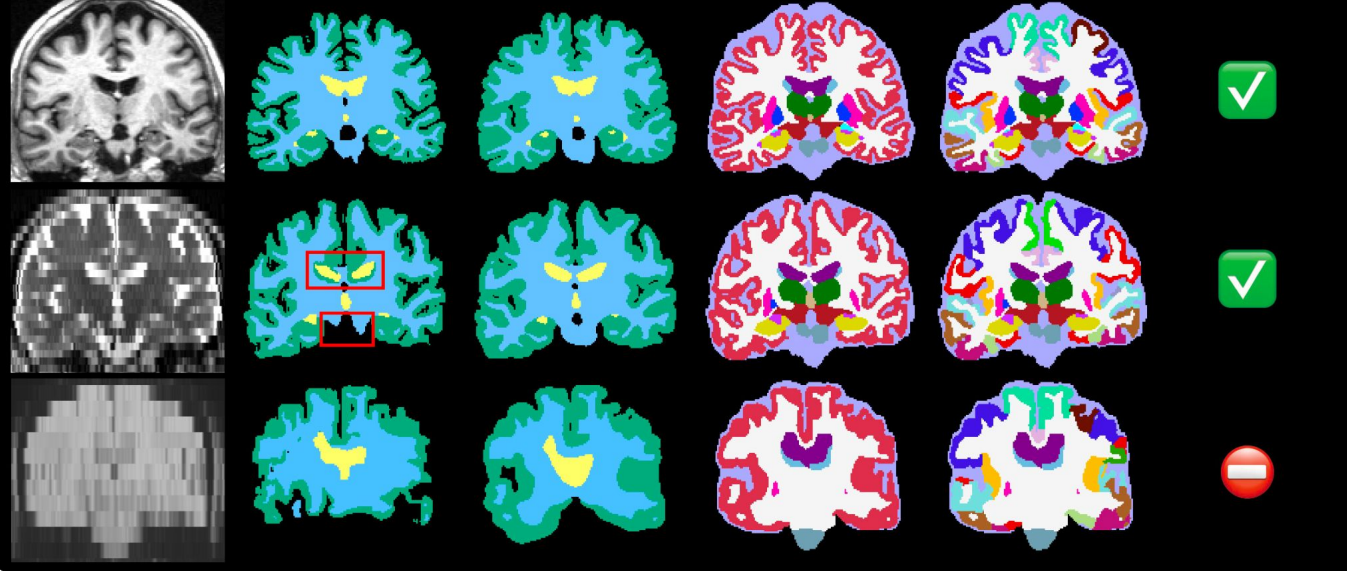


Auto-Encoders (AE) Generative Adversarial Networks (GAN)

Redes neuronales convolucionales

Billot, Benjamin, et al. *Proceedings of the National Academy of Sciences* 120.9 (2023)

Segmentación y clasificación de imágenes médicas



Detección de ánimo facial



Detección de objetos en video

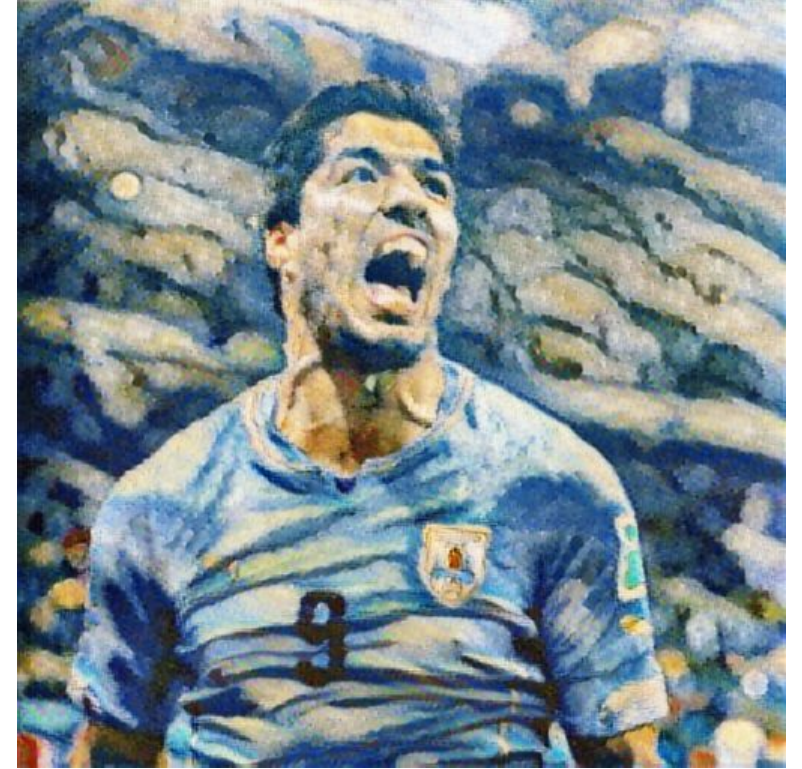
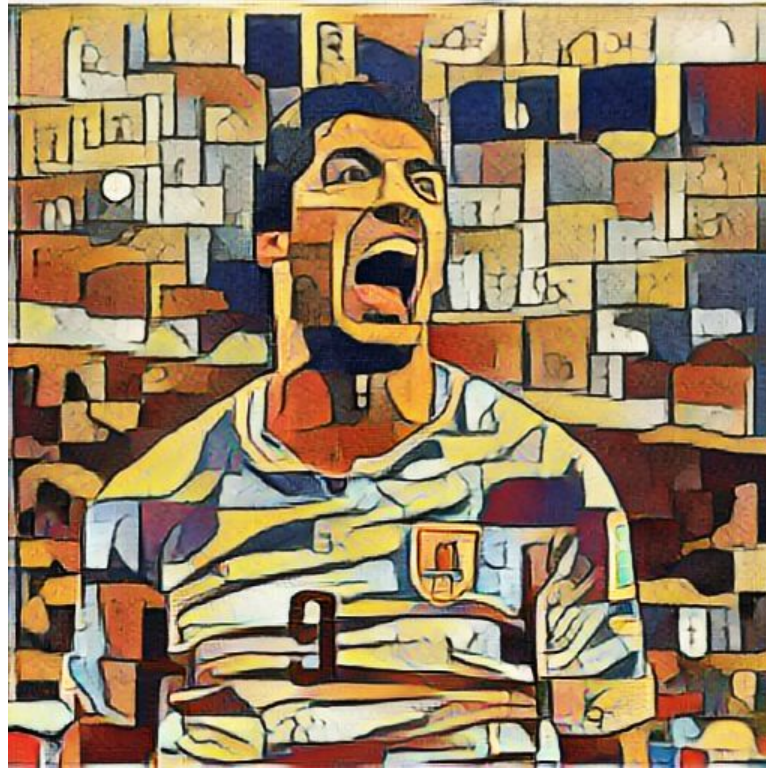
Istanbul Traffic - YOLO COCO - Object Detection (youtube.com)



Transferencia de estilo

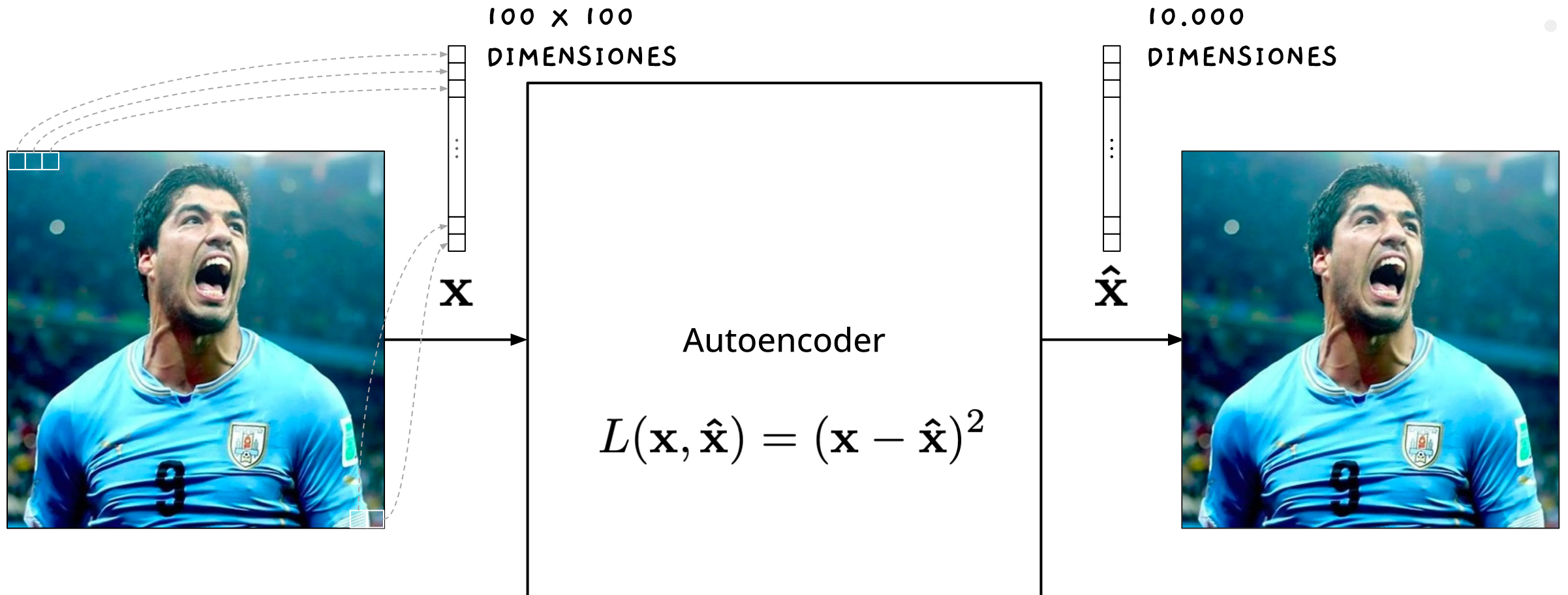
Artistic Style Transfer with Convolutional Neural Network, Manjeet Singh (medium.com)

Redes neuronales convolucionales

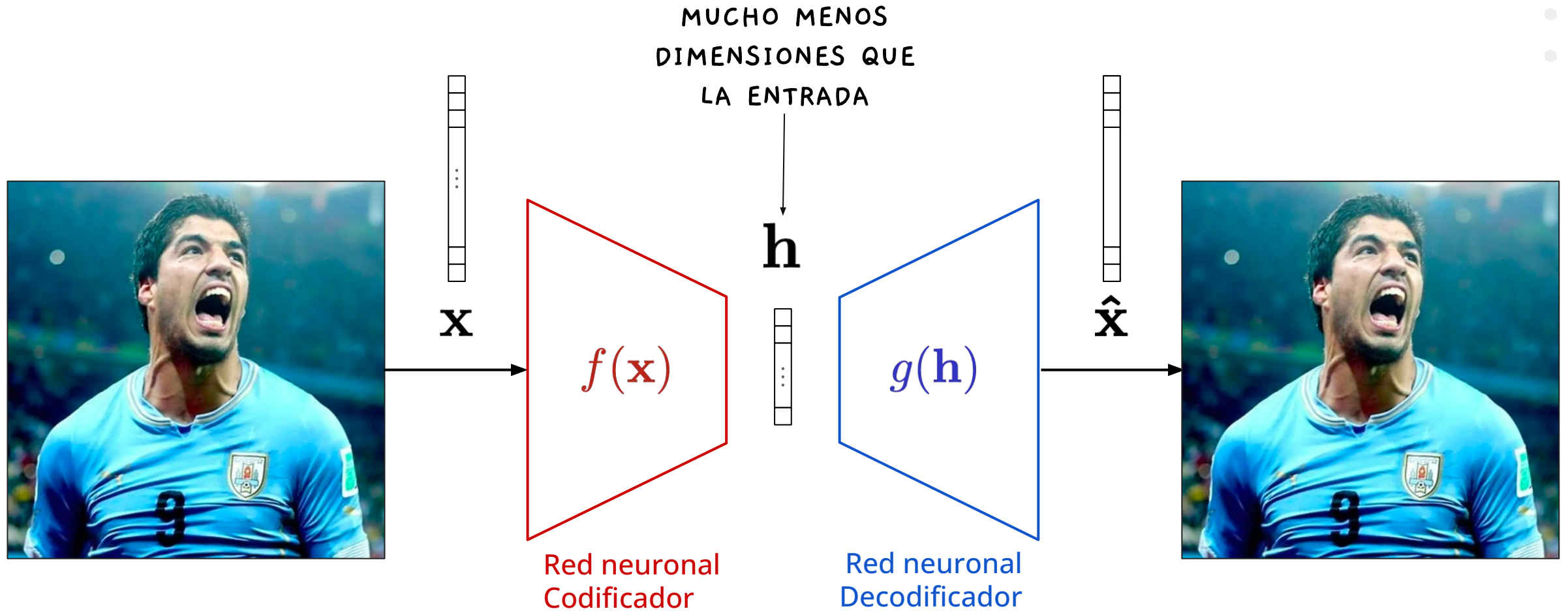


Autoencoders (AE)

'propio', 'de uno
misma' codificación: nueva
representación
(autobiografía) eficiente

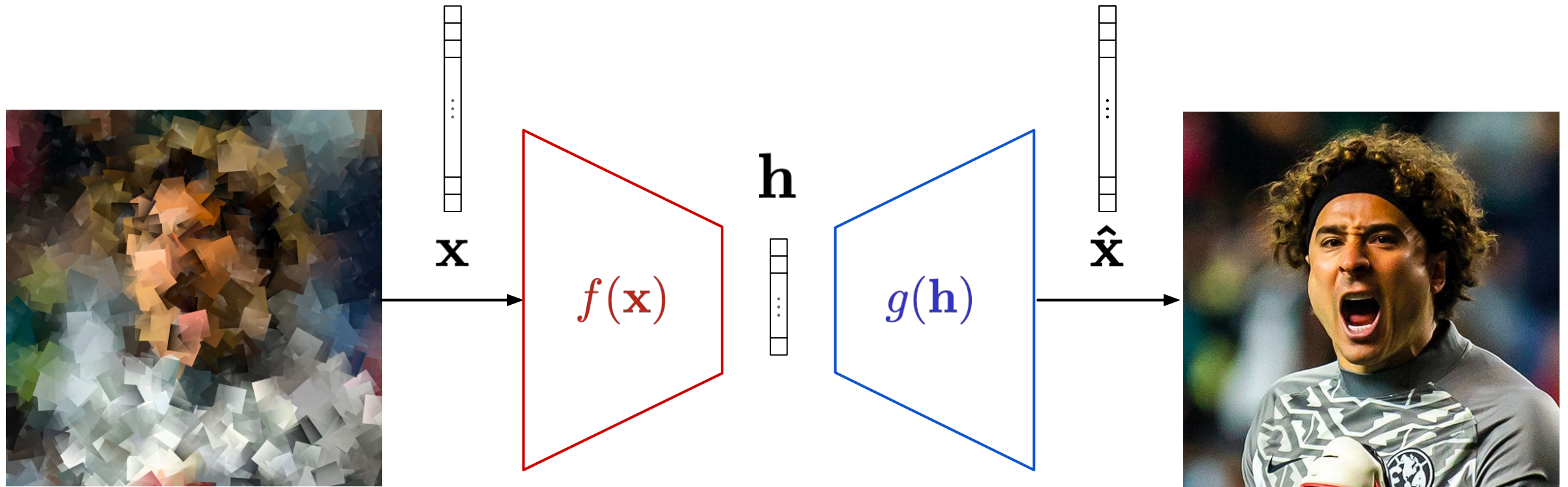


Autoencoders (AE)



$\mathbf{h}$ es un *resumen* de $\mathbf{x}$, describe lo *esencial* de la entrada dejando de lado *pequeñas perturbaciones*

Autoencoders (AE)



Autoencoders (AE)

Colorear imágenes



becominghuman.ai

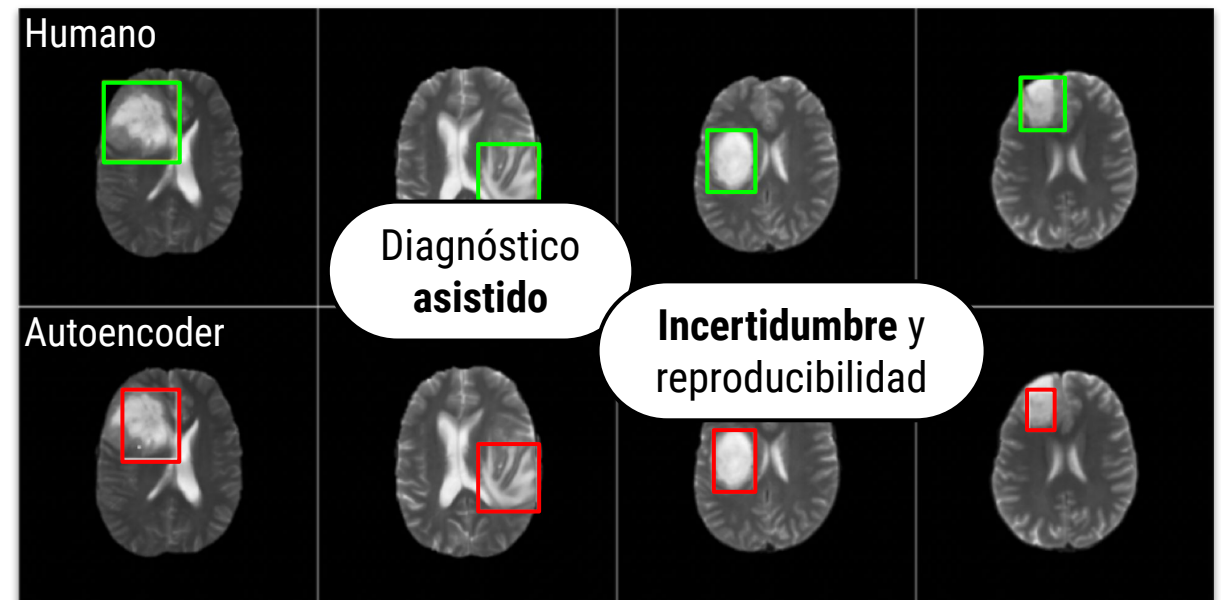
- Compresión de imágenes
- Eliminación de ruido en imágenes
- Sistema de recomendación
- Imputación de valores perdidos
- Generación de datos

Corregir iluminación



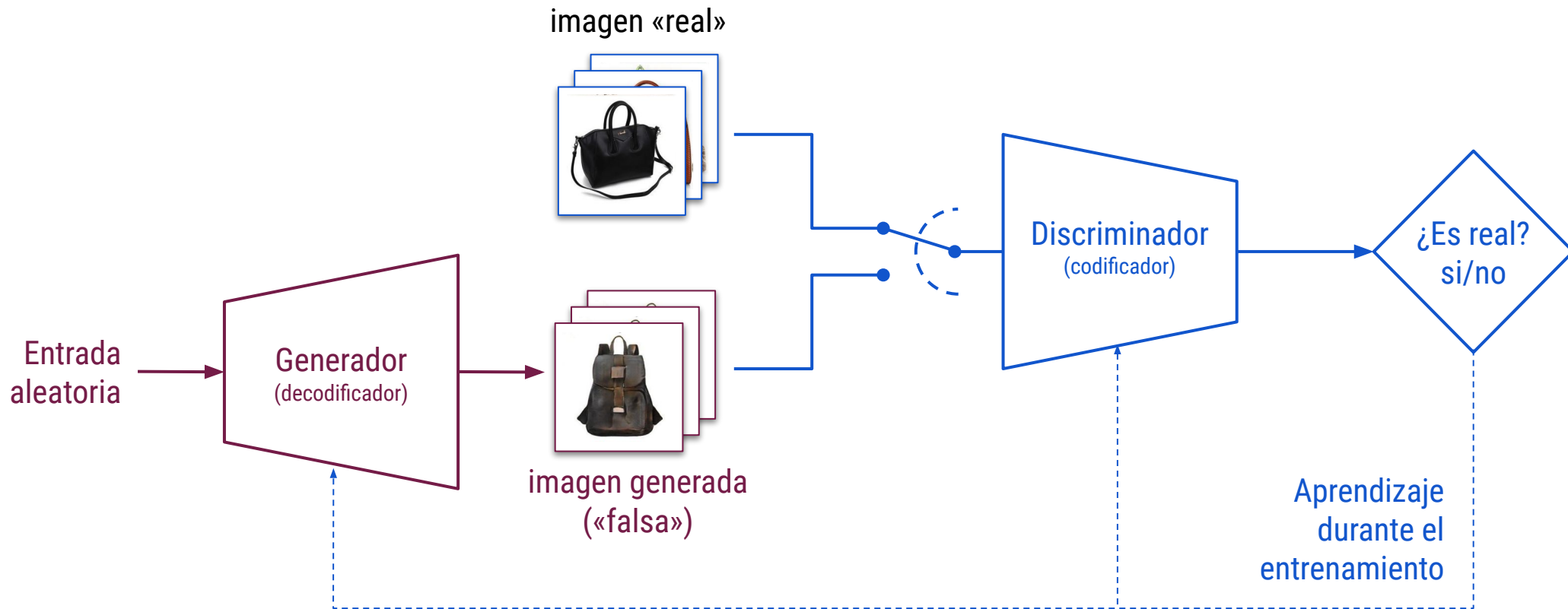
Dias Da Cruz, Steve, et al. *Sensors* 22.13 (2022)

Detectar anomalías en imágenes médicas



Chatterjee, Soumick, et al. *Computers in Biology and Medicine* 149 (2022)

Redes generativas adversarias (GAN)



- La idea clave: la competencia mejora el desempeño.

Redes generativas adversarias (GAN)

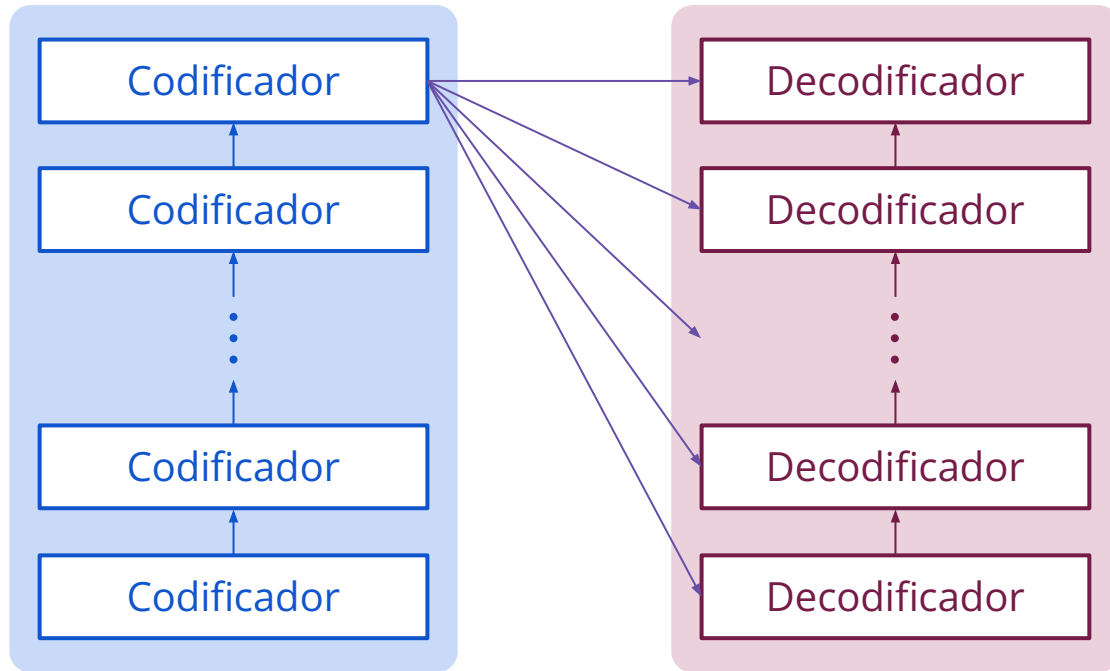


Estas personas no existen
<https://thispersondoesnotexist.com/>

Transformers

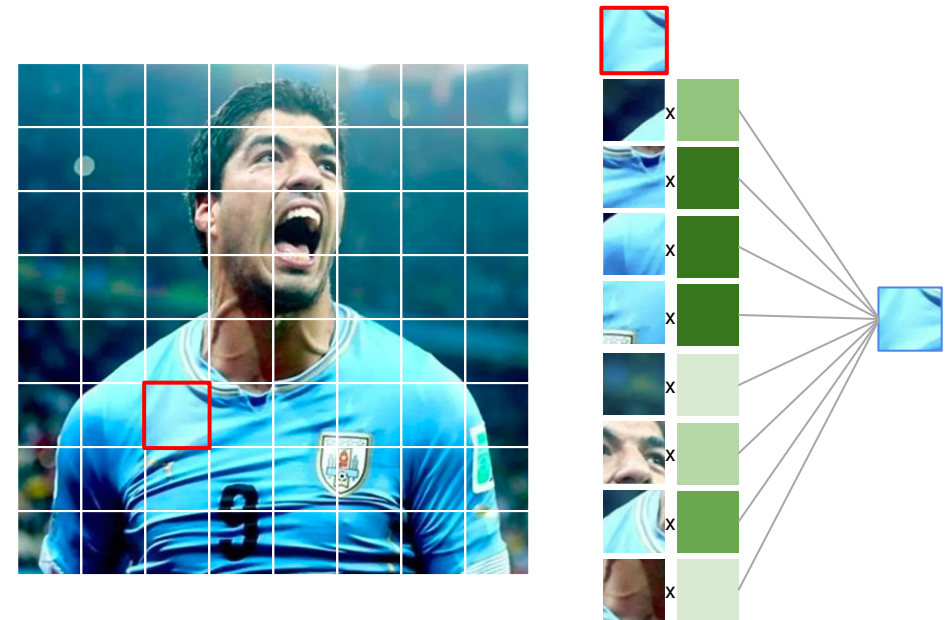
 **ChatGPT** ⇒ Generative Pre-trained Transformer

Arquitectura codificador-decodificador



Auto-atención

Representar una parte de *algo* como una combinación de sus otras partes dada la *relevancia* entre ellas.



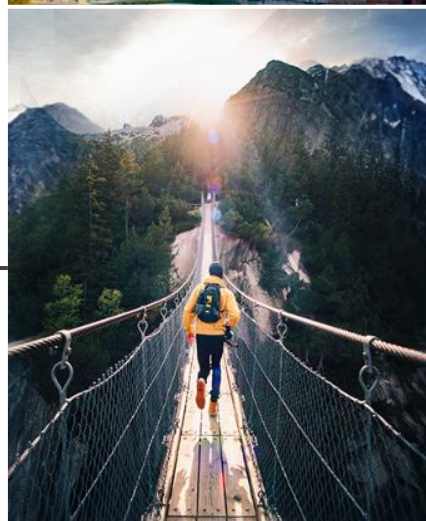
?



?

Queremos generar (sintetizar) la imagen de entrada (consulta) hacia los lados.

?



?

?



?

Transformers



Chang, Huiwen, et al. "Maskgit: Masked generative image transformer." Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition. 2022.



Transformers

Transformer model

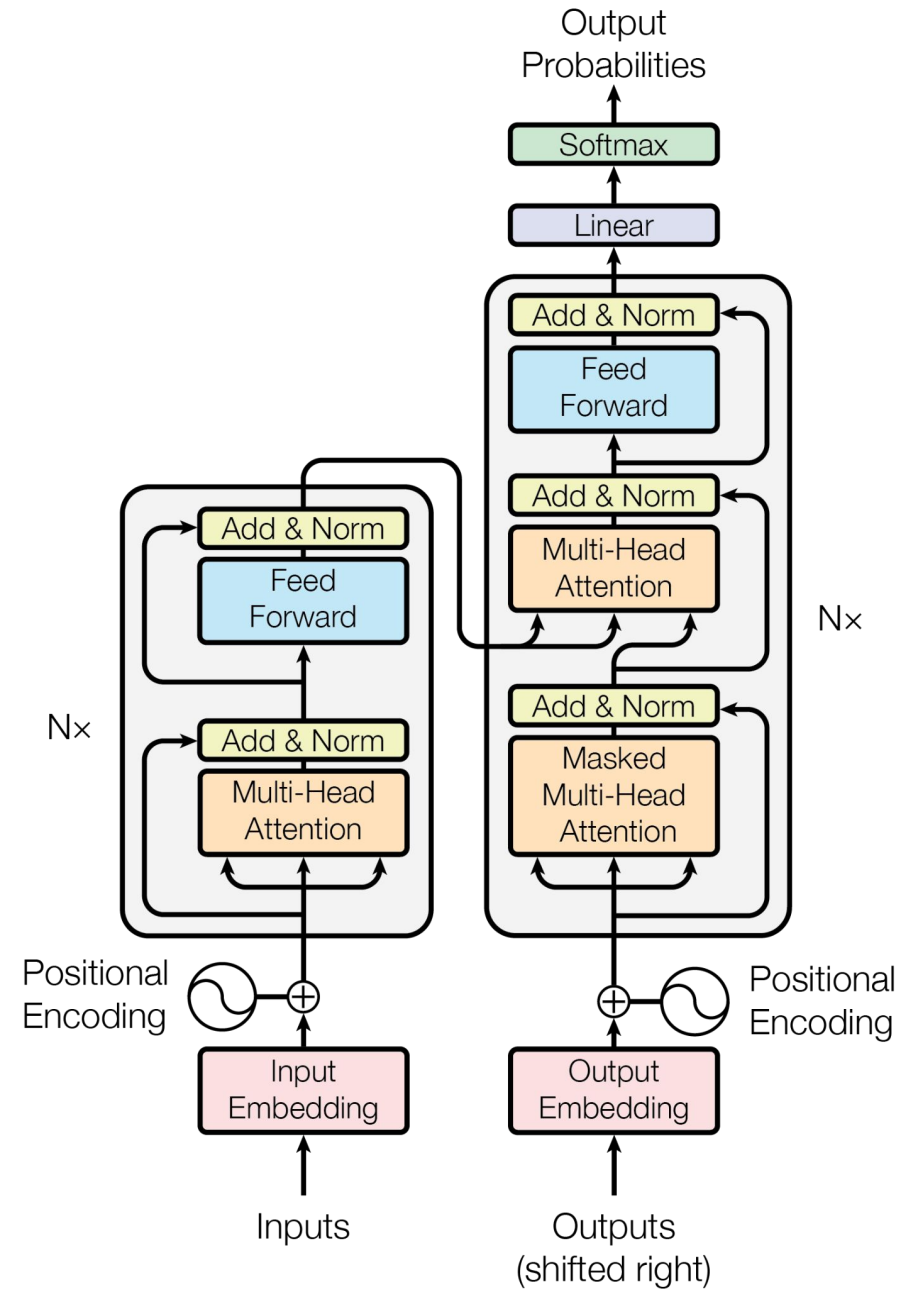
- Can draw global dependencies between sequences.
- Originally created for Natural Language Processing (NLP).
- Bidirectional structure: Encoder - Decoder.
- Core block: Attention module.
- Some known applications:



175 billion
parameters

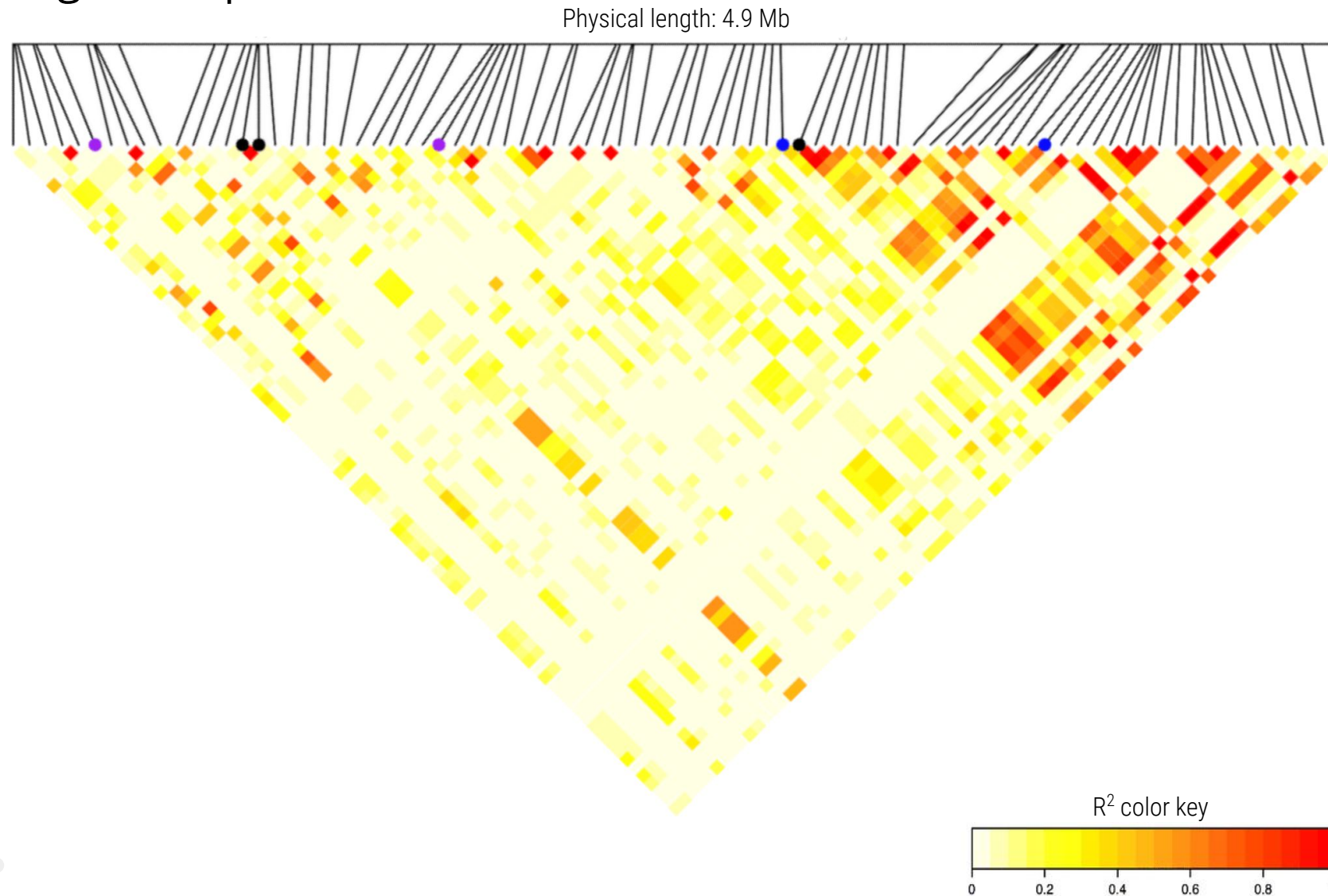


3.5 billion
parameters



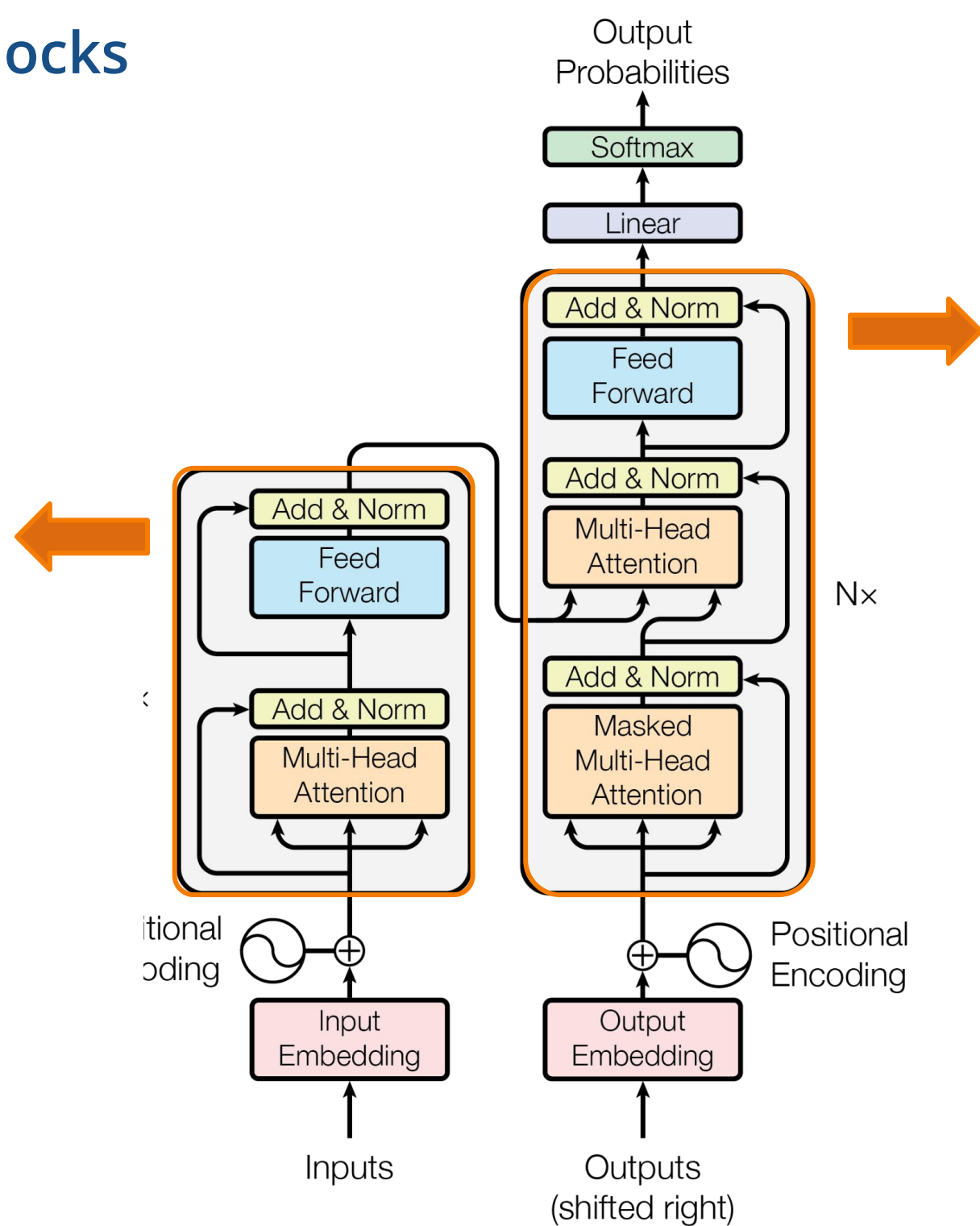
Why using Transformers in Genomic Prediction?

Pairwise Linkage Disequilibrium:



Transformer's blocks

Encoder: extract semantics or representations from the data.

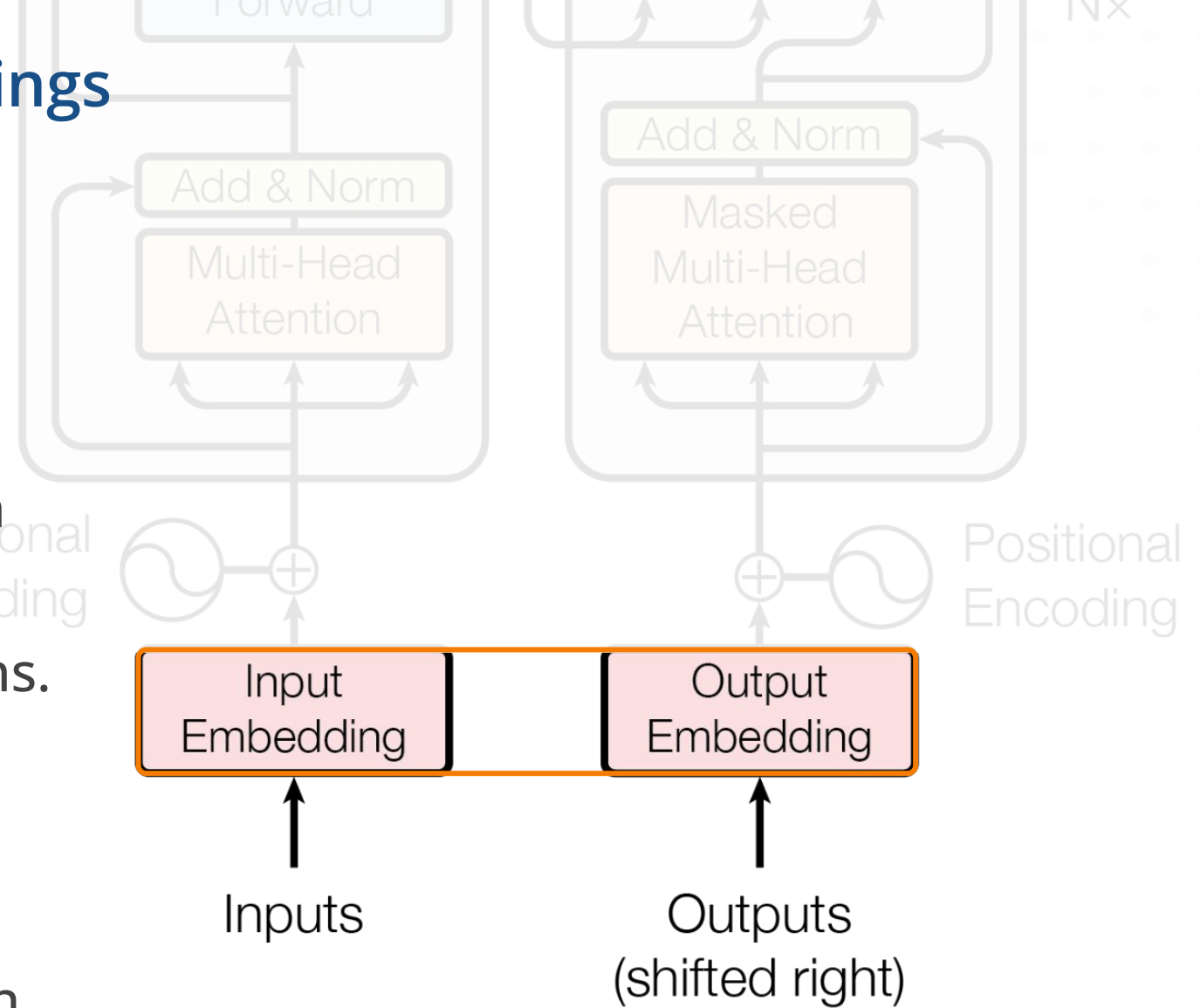


Decoder: used for generating new data (i.e.: create new sentences, translating the input, etc.)

Transformer's blocks: Embeddings

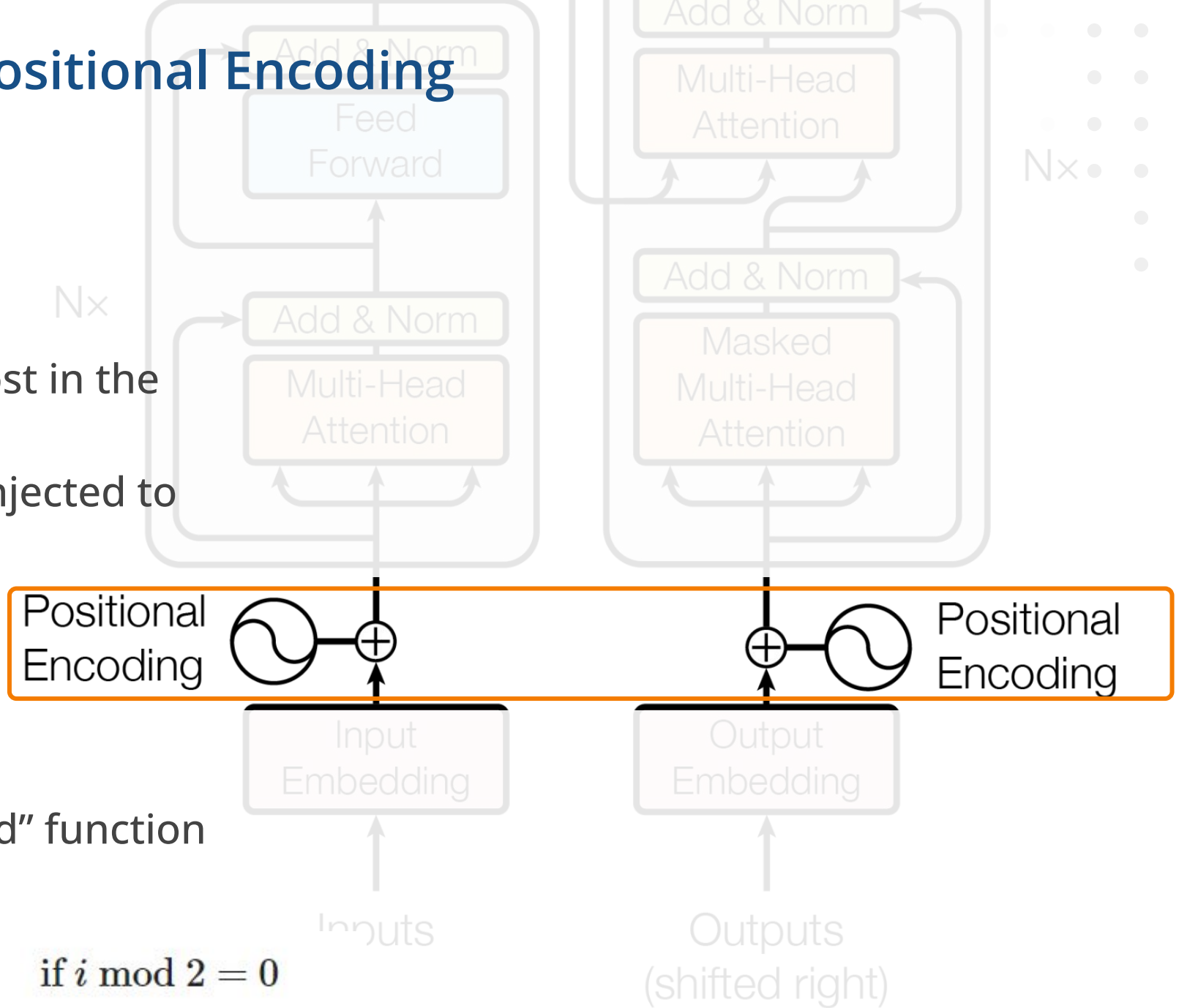
Transforms the input data into an acceptable format for the model.

- Each word is turned into tokens.
- Each token is mapped to an index of a look up table.
- Each entry maps to a learned embedding of d_{model} dimension.



Transformer's blocks: Positional Encoding

- The order of tokens is lost in the calculations.
- The sequence order is injected to each token.

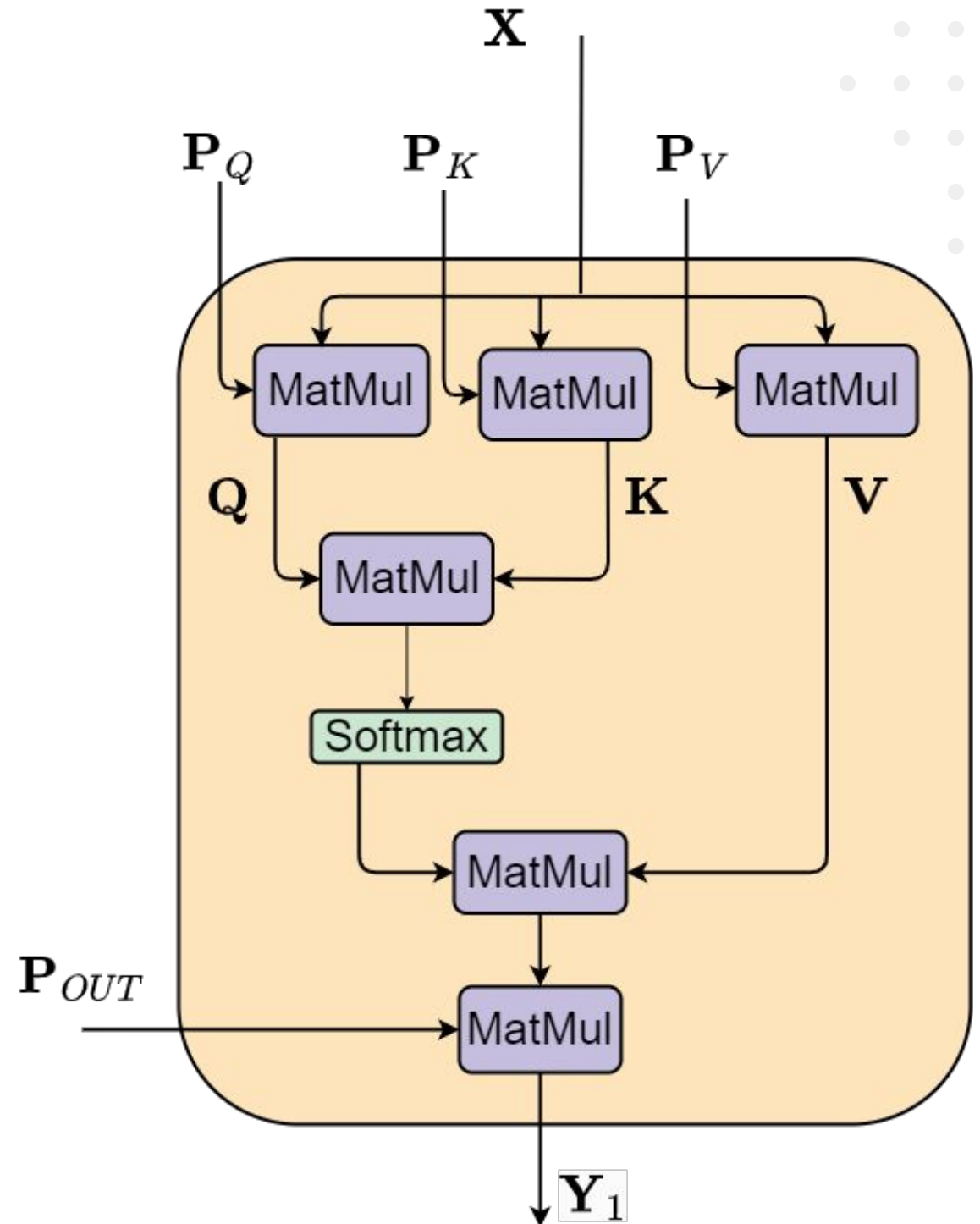


- “Attention is all you need” function for Positional Encoding:

$$P_D^i = \begin{cases} \sin\left(\frac{D}{10000^{i/d_m}}\right), & \text{if } i \bmod 2 = 0 \\ \cos\left(\frac{D}{10000^{((i-1)/d_m)}}\right), & \text{otherwise} \end{cases}$$

Self-Attention

- Attention allows models to "focus" on different parts of an input; while considering the entire context.
- Self-attention is an attention mechanism where each vector of a given input sequence attends to the entire sequence.
- Three important matrices:
 - Query: asking for information.
 - Key: saying we have the information.
 - Value: getting the information.



Self-Attention

Compute Query, Key and Value:

$$\mathbf{P}_Q =$$

2	0	1
0	1	1

$$\mathbf{P}_K =$$

0	1	1
2	0	0

$$\mathbf{P}_V =$$

1	1	3
0	1	0

$$\mathbf{Q} = \mathbf{P}_Q \mathbf{X}$$

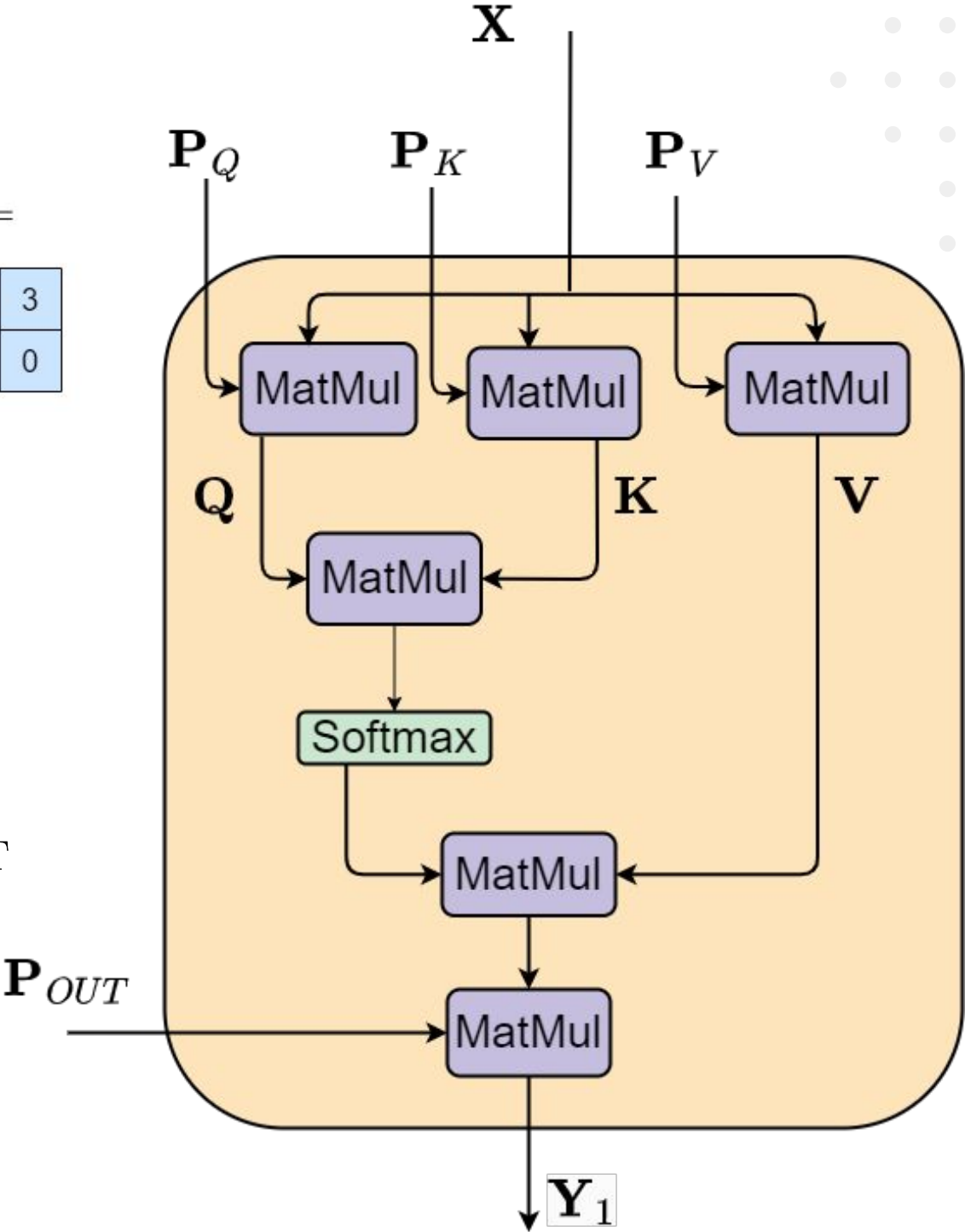
$$\mathbf{K} = \mathbf{P}_K \mathbf{X}$$

$$\mathbf{V} = \mathbf{P}_V \mathbf{X}$$

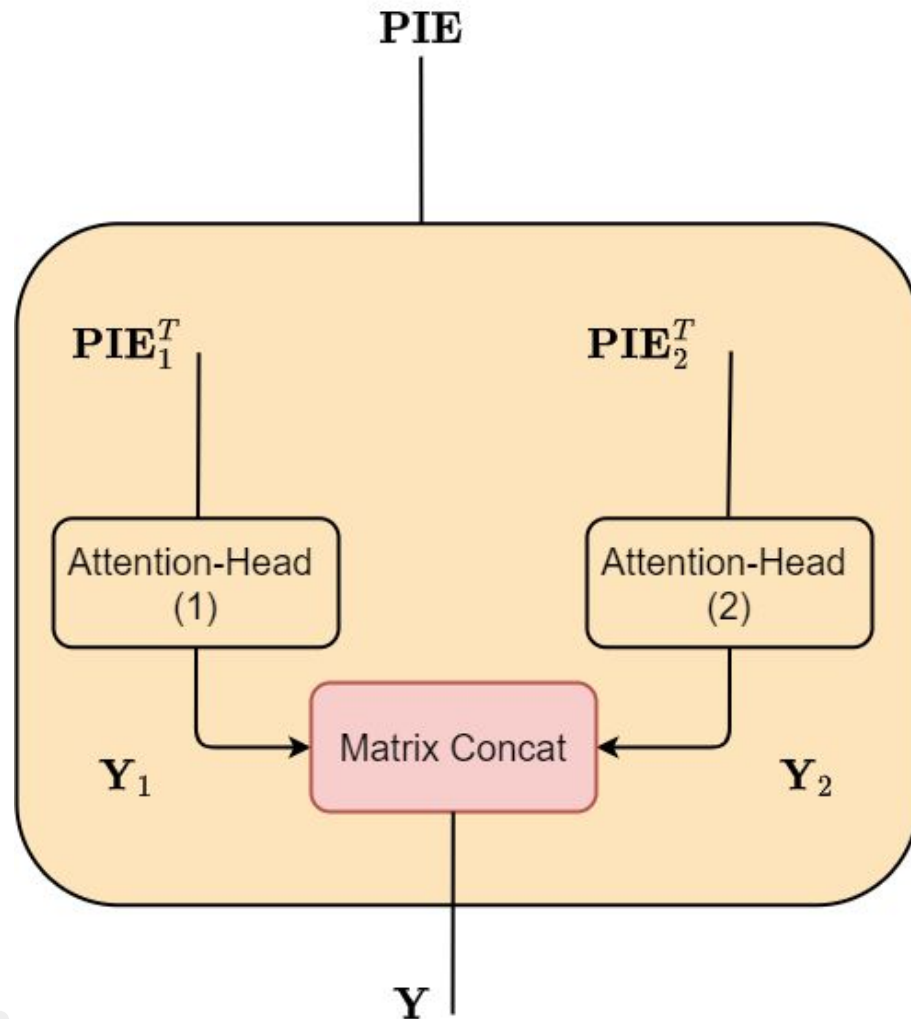
$$\mathbf{Y} = \text{Softmax} \left(\frac{\mathbf{Q} \times \mathbf{K}}{\sqrt{k}} \right) \times \mathbf{V} \times \mathbf{P}_{OUT}$$

Attention Score

The information we keep



Multi-Head Self-Attention



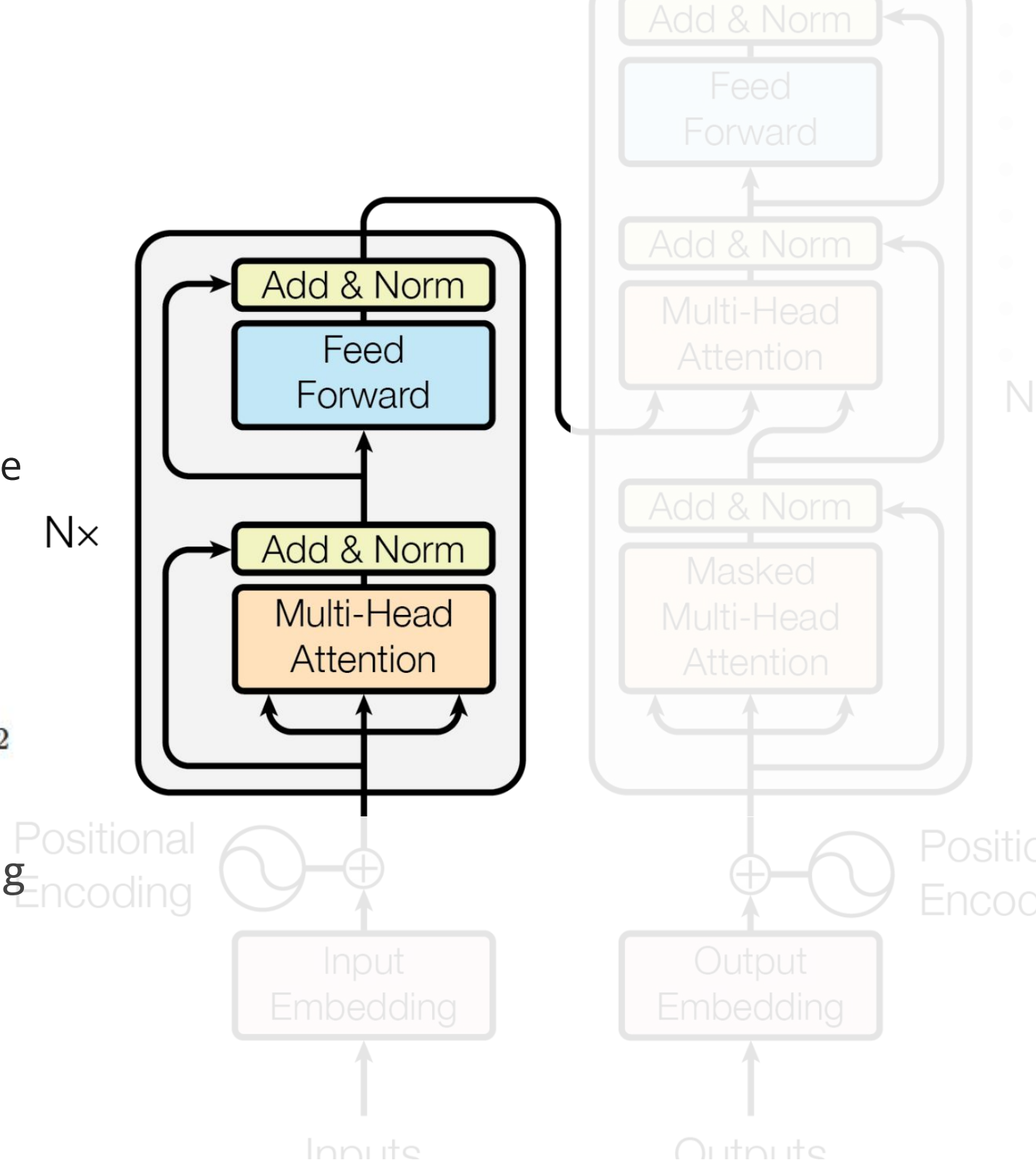
- Parallelization to optimize computation
- Separates each vector in k smaller vectors with dimension d_{model}/k
 k : number of heads
- Outputs the concatenation of the attention heads' outputs.

Transformer's blocks: Encoder

- Multi-Head Self-Attention to capture the dependencies of the input sequence with itself.
- Feed-Forward: then projecting these new token representations to a space that the next block can use more optimally

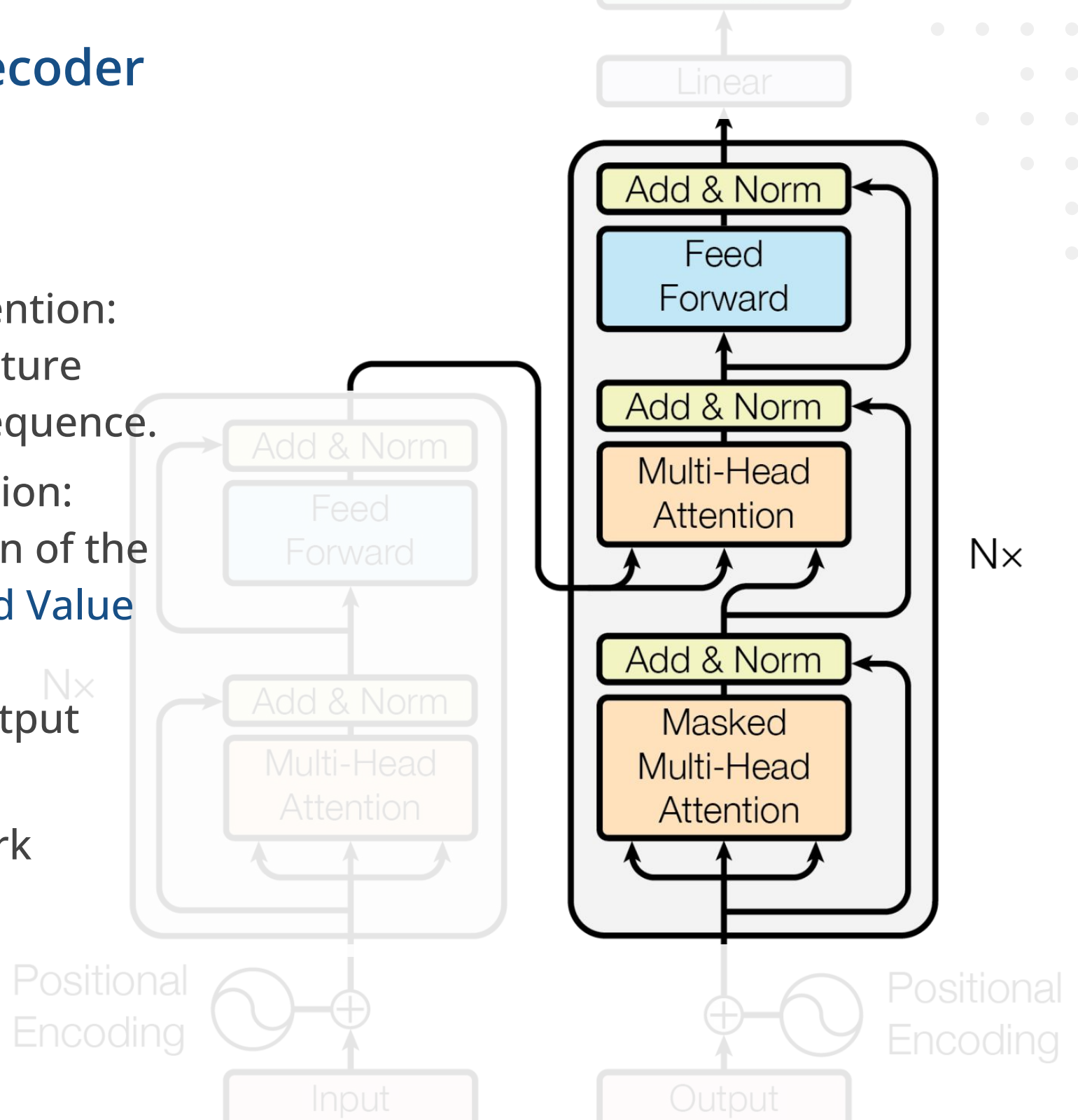
$$\text{FFN}(x) = \max(0, xW_1 + b_1)W_2 + b_2$$

- The Add&Norm layer is a residual connection to prevent the vanishing gradients problem.



Transformer's blocks: Decoder

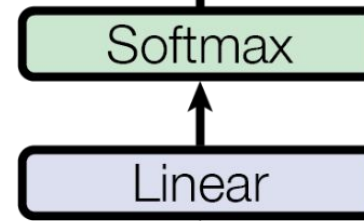
- Masked Multi-Head Attention: masks information of future positions of the input sequence.
- Cross Multi-Head Attention: incorporates information of the encoder output (**Key and Value matrix**) and the masked multi-head attention output (**Query matrix**).
- FFN and Add&Norm work as in the Encoder.



Transformer's blocks: Output Layer

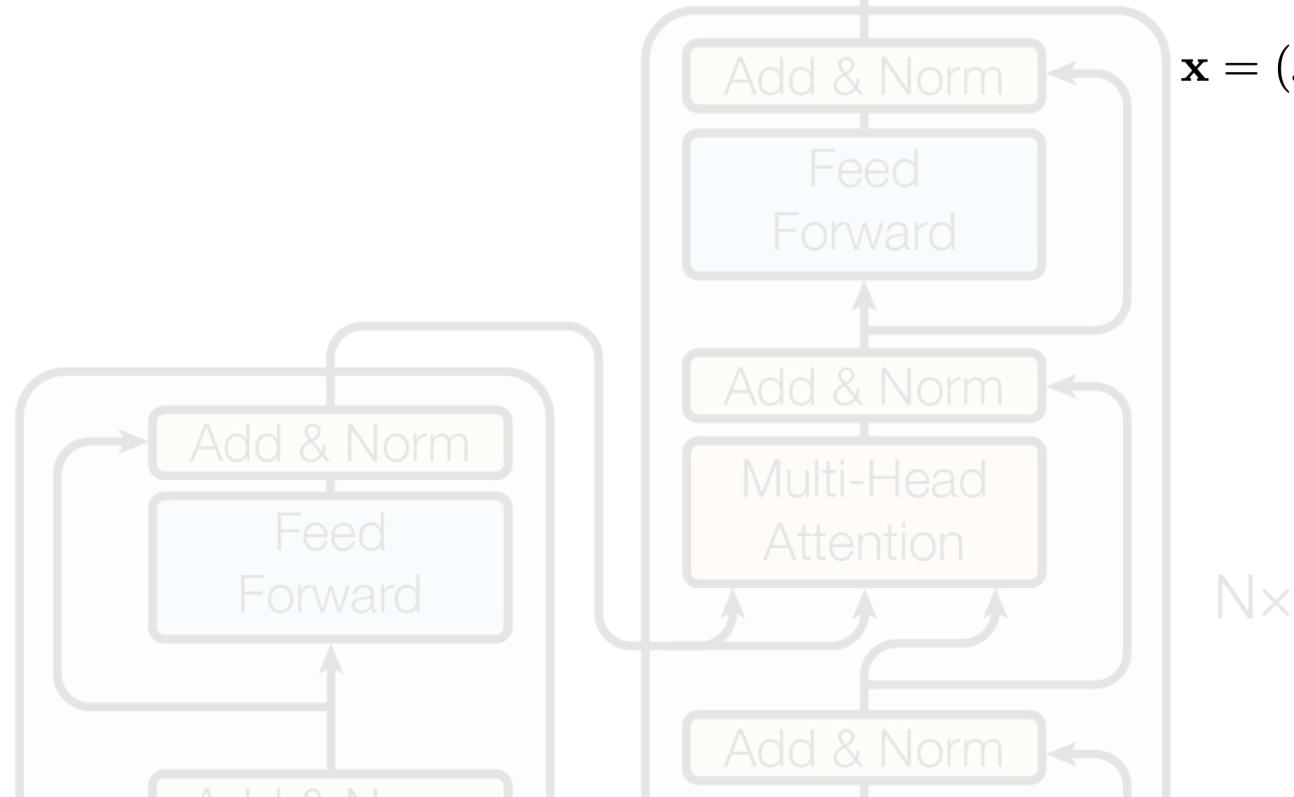
- Linear layer that takes the output dimension to the number of possible outputs.
- Softmax activation
- MAP estimation

Output
Probabilities



$$\sigma(\mathbf{x})_i = \frac{e^{x_i}}{\sum_{j=1}^K e^{x_j}}$$
$$i = 1, \dots, K$$

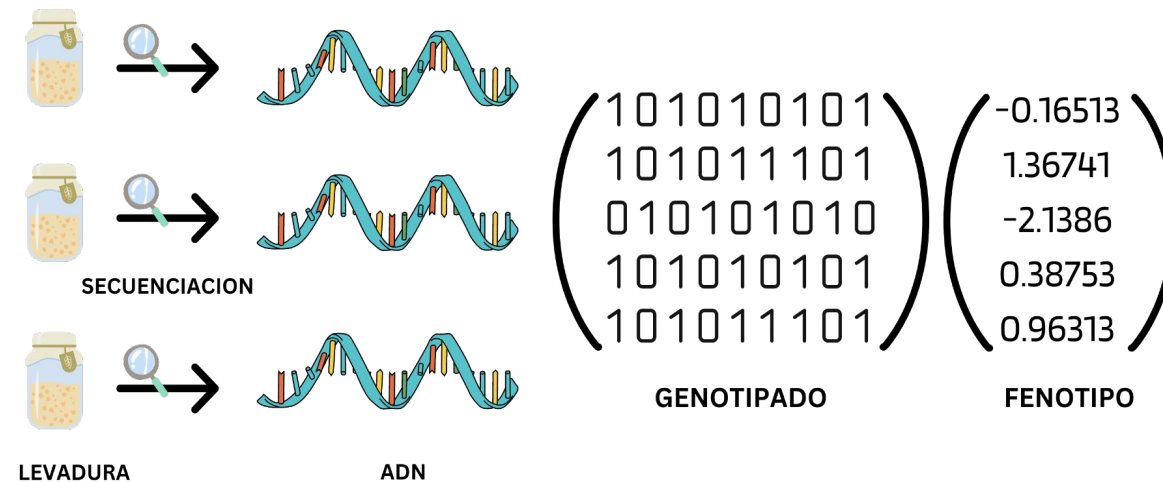
$$\mathbf{x} = (x_1, \dots, x_K) \in \mathbb{R}^K$$



Practical

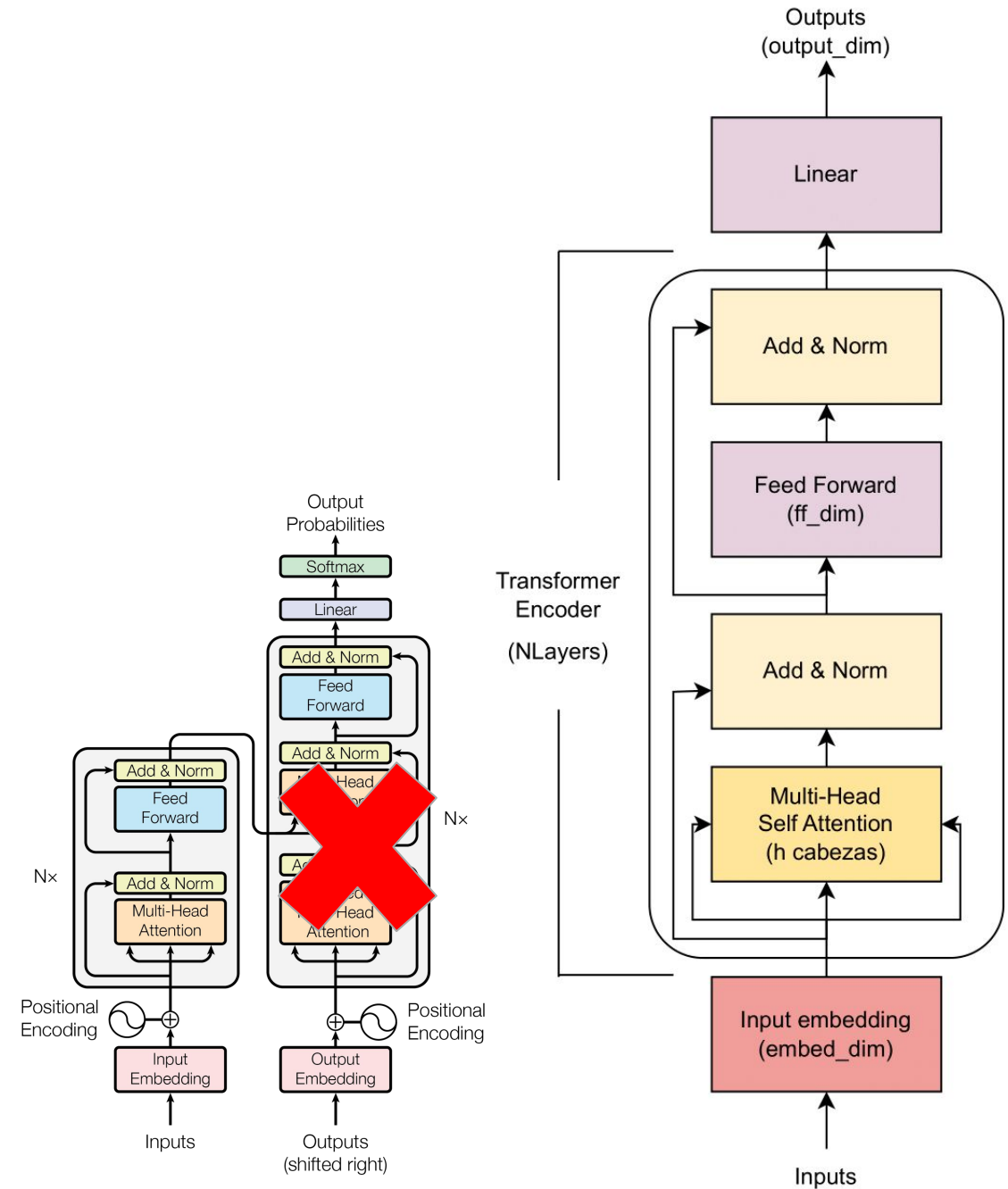
Yeast database:

- Phenotype: yeast growth in 48 different environments
 - 1008 samples with 11623 SNPs
- Working with two environments: Lactate and Lactose.
 - Single-trait: predict one phenotype per model
 - Multi-trait: predict both phenotypes training one single model



Practical

- GPTTransformer (Jubair et al., 2021)
 - Barley
- No Decoder, we want to learn the semantics of the data.
- Linear layer substitutes the Embedding and Positional Encoding layers
- Linear output



Practical

- Input: batches of 8 individuals.
- Metrics:
 - MSE
 - Pearson Correlation Coefficient
- Output dimension 1 or 2 depending on single-trait or multi-trait prediction.

