

Exploring Transformer-Based Models for Cascading Extremes

Clemente Ferrer

`clemente.ferrer@usm.cl`

Department of Mathematics
Universidad Técnica Federico Santa María

Joint work with **Miguel de Carvalho & Ronny Vallejos**

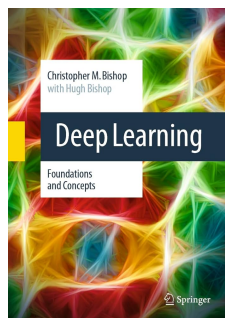
WORKSHOP ON GENERATIVE AI FOR EXTREME EVENTS

Edinburgh, UK
June 13, 2025

During this talk I will cover the following topics:

- ▶ Why use Transformers?
- ▶ A Literature Review on Transformers.
- ▶ Transformer in Statistics.
- ▶ Large Language Models and Generative Pre-trained Transformers (GPT).
- ▶ Ideas on Transformers for Chained Extremes.

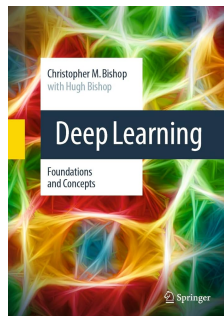
Most of the figures and ideas in this talk were taken from the book by [Bishop and Bishop \(2024\)](#).



During this talk I will cover the following topics:

- ▶ Why use Transformers?
- ▶ A Literature Review on Transformers.
- ▶ Transformer in Statistics.
- ▶ Large Language Models and Generative Pre-trained Transformers (GPT).
- ▶ Ideas on Transformers for Chained Extremes.

Most of the figures and ideas in this talk were taken from the book by [Bishop and Bishop \(2024\)](#).

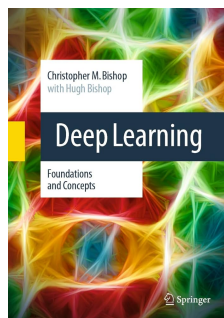


Literature Review

During this talk I will cover the following topics:

- ▶ Why use Transformers?
- ▶ A Literature Review on Transformers.
- ▶ Transformer in Statistics.
- ▶ Large Language Models and Generative Pre-trained Transformers (GPT).
- ▶ Ideas on Transformers for Chained Extremes.

Most of the figures and ideas in this talk were taken from the book by [Bishop and Bishop \(2024\)](#).

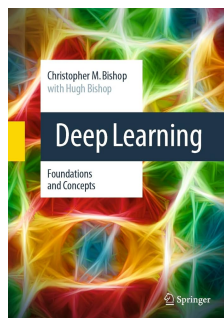


Literature Review

During this talk I will cover the following topics:

- ▶ Why use Transformers?
- ▶ A Literature Review on Transformers.
- ▶ Transformer in Statistics.
- ▶ Large Language Models and Generative Pre-trained Transformers (GPT).
- ▶ Ideas on Transformers for Chained Extremes.

Most of the figures and ideas in this talk were taken from the book by [Bishop and Bishop \(2024\)](#).

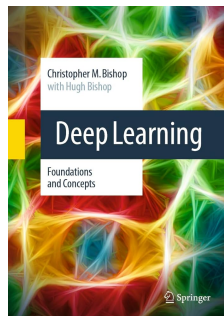


Literature Review

During this talk I will cover the following topics:

- ▶ Why use Transformers?
- ▶ A Literature Review on Transformers.
- ▶ Transformer in Statistics.
- ▶ Large Language Models and Generative Pre-trained Transformers (GPT).
- ▶ Ideas on Transformers for Chained Extremes.

Most of the figures and ideas in this talk were taken from the book by [Bishop and Bishop \(2024\)](#).

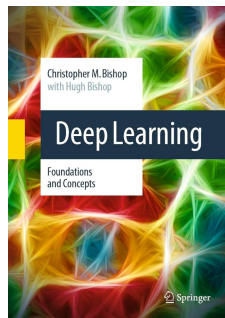


Literature Review

During this talk I will cover the following topics:

- ▶ Why use Transformers?
- ▶ A Literature Review on Transformers.
- ▶ Transformer in Statistics.
- ▶ Large Language Models and Generative Pre-trained Transformers (GPT).
- ▶ Ideas on Transformers for Chained Extremes.

Most of the figures and ideas in this talk were taken from the book by [Bishop and Bishop \(2024\)](#).



Why use Transformers?

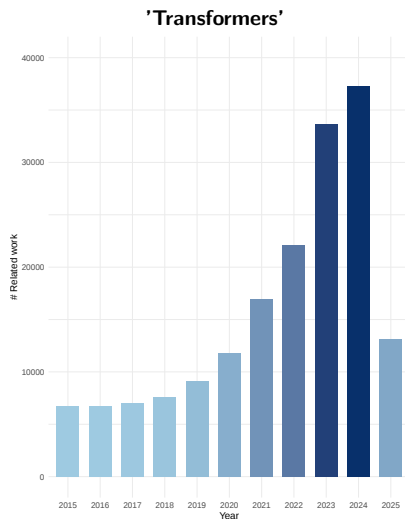


Figure: Number of works (articles, preprints, book chapter) referring to 'Transformers' since 2015 (Source: OpenAlex.org).

Literature Review

Suppose N tokens

$$\mathbf{x}_1, \dots, \mathbf{x}_N \leftarrow \text{Input data.}$$

where $\mathbf{x}_n = (x_{n1}, \dots, x_{nD})^T \in \mathbb{R}^D$, $n = 1, \dots, N$. The elements of the tokens are called features.

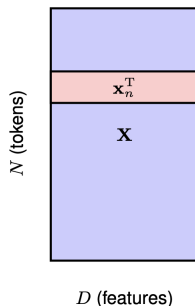


Figure: The structure of the data matrix $\mathbf{X}$, of dimension $N \times D$, in which row n represents the transposed data vector $\mathbf{x}_n^T$.

Literature Review

Suppose N tokens

$$\mathbf{x}_1, \dots, \mathbf{x}_N \leftarrow \text{Input data.}$$

where $\mathbf{x}_n = (x_{n1}, \dots, x_{nD})^T \in \mathbb{R}^D$, $n = 1, \dots, N$. The elements of the tokens are called **features**.

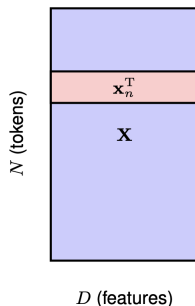


Figure: The structure of the data matrix $\mathbf{X}$, of dimension $N \times D$, in which row n represents the transposed data vector $\mathbf{x}_n^T$.

Literature Review

GOAL

Define a **transformation** from tokens $\mathbf{x}_1, \dots, \mathbf{x}_N$ to output embeddings $\mathbf{y}_1, \dots, \mathbf{y}_N$, where each $\mathbf{y}_n$ is a weighted combination of **ALL** $\mathbf{x}_m$, with weights **reflecting the relevance of $\mathbf{x}_m$ to $\mathbf{y}_n$** .

PROPOSAL

$$\mathbf{y}_n = \sum_{m=1}^N \alpha_{nm} \mathbf{x}_m,$$

where $\alpha_{nm} \geq 0$ are **attention weights** such that $\sum_{m=1}^N \alpha_{nm} = 1$.

FIRST APPROACH

$$a_{nm} = \frac{\exp(\mathbf{x}_n^T \mathbf{x}_m)}{\sum_{m'=1}^N \exp(\mathbf{x}_n^T \mathbf{x}_{m'})} \quad \leftarrow \quad \text{Softmax.}$$

Using matrix notation

$$\mathbf{Y} = \text{Softmax}(\mathbf{X}\mathbf{X}^T)\mathbf{X} \quad \leftarrow \quad \text{Self-Attention.}$$

Literature Review

GOAL

Define a **transformation** from tokens $\mathbf{x}_1, \dots, \mathbf{x}_N$ to output embeddings $\mathbf{y}_1, \dots, \mathbf{y}_N$, where each $\mathbf{y}_n$ is a weighted combination of **ALL** $\mathbf{x}_m$, with weights **reflecting the relevance of $\mathbf{x}_m$ to $\mathbf{y}_n$** .

PROPOSAL

$$\mathbf{y}_n = \sum_{m=1}^N \alpha_{nm} \mathbf{x}_m,$$

where $\alpha_{nm} \geq 0$ are **attention weights** such that $\sum_{m=1}^N \alpha_{nm} = 1$.

FIRST APPROACH

$$a_{nm} = \frac{\exp(\mathbf{x}_n^T \mathbf{x}_m)}{\sum_{m'=1}^N \exp(\mathbf{x}_n^T \mathbf{x}_{m'})} \quad \leftarrow \quad \text{Softmax.}$$

Using matrix notation

$$\mathbf{Y} = \text{Softmax}(\mathbf{X}\mathbf{X}^T)\mathbf{X} \quad \leftarrow \quad \text{Self-Attention.}$$

Literature Review

GOAL

Define a **transformation** from tokens $\mathbf{x}_1, \dots, \mathbf{x}_N$ to output embeddings $\mathbf{y}_1, \dots, \mathbf{y}_N$, where each $\mathbf{y}_n$ is a weighted combination of **ALL** $\mathbf{x}_m$, with weights **reflecting the relevance of $\mathbf{x}_m$ to $\mathbf{y}_n$** .

PROPOSAL

$$\mathbf{y}_n = \sum_{m=1}^N \alpha_{nm} \mathbf{x}_m,$$

where $\alpha_{nm} \geq 0$ are **attention weights** such that $\sum_{m=1}^N \alpha_{nm} = 1$.

FIRST APPROACH

$$a_{nm} = \frac{\exp(\mathbf{x}_n^T \mathbf{x}_m)}{\sum_{m'=1}^N \exp(\mathbf{x}_n^T \mathbf{x}_{m'})} \quad \leftarrow \quad \text{Softmax.}$$

Using matrix notation

$$\mathbf{Y} = \text{Softmax}(\mathbf{X}\mathbf{X}^T)\mathbf{X} \quad \leftarrow \quad \text{Self-Attention.}$$

Literature Review

GOAL

Define a **transformation** from tokens $\mathbf{x}_1, \dots, \mathbf{x}_N$ to output embeddings $\mathbf{y}_1, \dots, \mathbf{y}_N$, where each $\mathbf{y}_n$ is a weighted combination of **ALL** $\mathbf{x}_m$, with weights **reflecting the relevance of $\mathbf{x}_m$ to $\mathbf{y}_n$** .

PROPOSAL

$$\mathbf{y}_n = \sum_{m=1}^N \alpha_{nm} \mathbf{x}_m,$$

where $\alpha_{nm} \geq 0$ are **attention weights** such that $\sum_{m=1}^N \alpha_{nm} = 1$.

FIRST APPROACH

$$a_{nm} = \frac{\exp(\mathbf{x}_n^T \mathbf{x}_m)}{\sum_{m'=1}^N \exp(\mathbf{x}_n^T \mathbf{x}_{m'})} \quad \leftarrow \quad \text{Softmax.}$$

Using matrix notation

$$\mathbf{Y} = \text{Softmax}(\mathbf{X}\mathbf{X}^T)\mathbf{X} \quad \leftarrow \quad \text{Self-Attention.}$$

Literature Review

GOAL

Define a **transformation** from tokens $\mathbf{x}_1, \dots, \mathbf{x}_N$ to output embeddings $\mathbf{y}_1, \dots, \mathbf{y}_N$, where each $\mathbf{y}_n$ is a weighted combination of **ALL** $\mathbf{x}_m$, with weights **reflecting the relevance of $\mathbf{x}_m$ to $\mathbf{y}_n$** .

PROPOSAL

$$\mathbf{y}_n = \sum_{m=1}^N \alpha_{nm} \mathbf{x}_m,$$

where $\alpha_{nm} \geq 0$ are **attention weights** such that $\sum_{m=1}^N \alpha_{nm} = 1$.

FIRST APPROACH

$$a_{nm} = \frac{\exp(\mathbf{x}_n^T \mathbf{x}_m)}{\sum_{m'=1}^N \exp(\mathbf{x}_n^T \mathbf{x}_{m'})} \quad \leftarrow \quad \text{Softmax.}$$

Using matrix notation

$$\mathbf{Y} = \text{Softmax}(\mathbf{X}\mathbf{X}^T)\mathbf{X} \quad \leftarrow \quad \text{Self-Attention.}$$

Literature Review

LIMITATION

The mapping $\{\mathbf{x}_n\} \mapsto \{\mathbf{y}_n\}$ lacks learnable parameters and thus cannot adapt to data. Additionally, all features in each $\mathbf{x}_n$ contribute equally to attention weights, limiting the model's ability to focus selectively on informative features.

PARTIAL SOLUTION

Define

$$\tilde{\mathbf{X}} = \mathbf{X}\mathbf{U}, \quad \mathbf{U} \in \mathbb{R}^{D \times D},$$

where $\mathbf{U}$ is learnable, analogous to a 'layer' in standard neural networks. Then

$$\mathbf{Y} = \text{Softmax}(\mathbf{X}\mathbf{U}\mathbf{U}^T\mathbf{X}^T)\mathbf{X}\mathbf{U}.$$

STILL A PROBLEM

Although this has much more flexibility, the matrix

$$\mathbf{X}\mathbf{U}\mathbf{U}^T\mathbf{X}^T$$

is symmetric, whereas we would like the attention mechanism to support asymmetry.

Literature Review

LIMITATION

The mapping $\{\mathbf{x}_n\} \mapsto \{\mathbf{y}_n\}$ **lacks learnable parameters** and thus **cannot** adapt to data. Additionally, **all features** in each $\mathbf{x}_n$ **contribute equally** to attention weights, limiting the model's ability to focus selectively on informative features.

PARTIAL SOLUTION

Define

$$\tilde{\mathbf{X}} = \mathbf{X}\mathbf{U}, \quad \mathbf{U} \in \mathbb{R}^{D \times D},$$

where $\mathbf{U}$ is **learnable**, analogous to a 'layer' in standard neural networks. Then

$$\mathbf{Y} = \text{Softmax}(\mathbf{X}\mathbf{U}\mathbf{U}^T\mathbf{X}^T)\mathbf{X}\mathbf{U}.$$

STILL A PROBLEM

Although this has much more flexibility, the matrix

$$\mathbf{X}\mathbf{U}\mathbf{U}^T\mathbf{X}^T$$

is **symmetric**, whereas we would like the **attention mechanism** to support **asymmetry**.

Literature Review

LIMITATION

The mapping $\{\mathbf{x}_n\} \mapsto \{\mathbf{y}_n\}$ **lacks learnable parameters** and thus **cannot** adapt to data. Additionally, **all features** in each $\mathbf{x}_n$ **contribute equally** to attention weights, limiting the model's ability to focus selectively on informative features.

PARTIAL SOLUTION

Define

$$\tilde{\mathbf{X}} = \mathbf{X}\mathbf{U}, \quad \mathbf{U} \in \mathbb{R}^{D \times D},$$

where $\mathbf{U}$ is **learnable**, analogous to a 'layer' in standard neural networks. Then

$$\mathbf{Y} = \text{Softmax}(\mathbf{X}\mathbf{U}\mathbf{U}^T\mathbf{X}^T)\mathbf{X}\mathbf{U}.$$

STILL A PROBLEM

Although this has much more flexibility, the matrix

$$\mathbf{X}\mathbf{U}\mathbf{U}^T\mathbf{X}^T$$

is **symmetric**, whereas we would like the **attention mechanism** to support **asymmetry**.

Literature Review

LIMITATION

The mapping $\{\mathbf{x}_n\} \mapsto \{\mathbf{y}_n\}$ **lacks learnable parameters** and thus **cannot** adapt to data. Additionally, **all features** in each $\mathbf{x}_n$ **contribute equally** to attention weights, limiting the model's ability to focus selectively on informative features.

PARTIAL SOLUTION

Define

$$\tilde{\mathbf{X}} = \mathbf{X}\mathbf{U}, \quad \mathbf{U} \in \mathbb{R}^{D \times D},$$

where $\mathbf{U}$ is **learnable**, analogous to a 'layer' in standard neural networks. Then

$$\mathbf{Y} = \text{Softmax}(\mathbf{X}\mathbf{U}\mathbf{U}^T\mathbf{X}^T)\mathbf{X}\mathbf{U}.$$

STILL A PROBLEM

Although this has much more flexibility, the matrix

$$\mathbf{X}\mathbf{U}\mathbf{U}^T\mathbf{X}^T$$

is **symmetric**, whereas we would like the **attention mechanism** to support **asymmetry**.

Literature Review

SOLUTION

Vaswani et al. (2017) proposed defining separate **query**, **key**, and **value** matrices each having their own independent linear transformations:

$$\mathbf{Q} = \mathbf{X}\mathbf{W}^{(q)}, \quad \mathbf{K} = \mathbf{X}\mathbf{W}^{(k)}, \quad \mathbf{V} = \mathbf{X}\mathbf{W}^{(v)},$$

where $\mathbf{W}^{(q)}, \mathbf{W}^{(k)} \in \mathbb{R}^{D \times D_k}$ and $\mathbf{W}^{(v)} \in \mathbb{R}^{D \times D_v}$, where D_v governs the dimensionality of the output. Therefore

$$\mathbf{Y} = \text{Softmax}(\mathbf{Q}\mathbf{K}^T)\mathbf{V} \quad \leftarrow \quad \mathbf{Q}\mathbf{K}^T \text{ are the attention weights.}$$

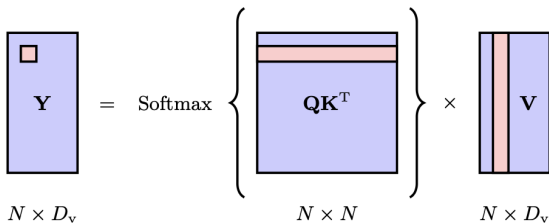


Figure: Illustration of the evaluation of the output from an attention layer given the query, key, and value matrices $\mathbf{Q}$, $\mathbf{K}$, and $\mathbf{V}$, respectively.

Literature Review

SOLUTION

Vaswani et al. (2017) proposed defining separate **query**, **key**, and **value** matrices each having their own independent linear transformations:

$$\mathbf{Q} = \mathbf{X}\mathbf{W}^{(q)}, \quad \mathbf{K} = \mathbf{X}\mathbf{W}^{(k)}, \quad \mathbf{V} = \mathbf{X}\mathbf{W}^{(v)},$$

where $\mathbf{W}^{(q)}, \mathbf{W}^{(k)} \in \mathbb{R}^{D \times D_k}$ and $\mathbf{W}^{(v)} \in \mathbb{R}^{D \times D_v}$, where D_v governs the dimensionality of the output. Therefore

$$\mathbf{Y} = \text{Softmax}(\mathbf{Q}\mathbf{K}^T)\mathbf{V} \quad \leftarrow \quad \mathbf{Q}\mathbf{K}^T \text{ are the } \text{attention weights}.$$

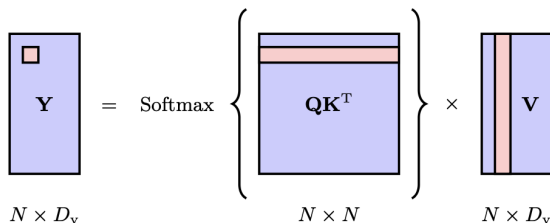


Figure: Illustration of the evaluation of the output from an attention layer given the query, key, and value matrices $\mathbf{Q}$, $\mathbf{K}$, and $\mathbf{V}$, respectively.

Literature Review

FINAL ADJUSTMENT

To **avoid vanishing gradients** in softmax, scale the dot product of query and key vectors by $\sqrt{D_k}$, yielding the output of the attention layer as:

$$\mathbf{Y} = \text{Attention}(\mathbf{Q}, \mathbf{K}, \mathbf{V}) \equiv \text{Softmax}\left(\frac{\mathbf{Q}\mathbf{K}^T}{\sqrt{D_k}}\right) \mathbf{V}$$

Scaled dot-product self-attention

This structure constitutes a **single attention head**.

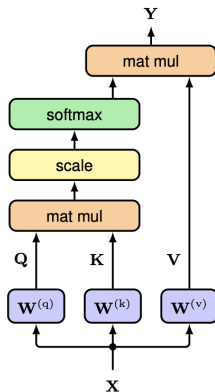


Figure: Information flow. Here 'mat mul' denotes matrix multiplication, and 'scale' refers to the normalization of the argument to the softmax using $\sqrt{D_k}$.

Literature review

MODELING MULTIPLE DEPENDENCY PATTERNS

To capture different types of relationships in the data simultaneously we replicate the attention mechanism into **multiple heads**.

Suppose we have H attention heads indexed by $h = 1, \dots, H$ of the form

$$\mathbf{H}_h = \text{Attention}(\mathbf{Q}_h, \mathbf{K}_h, \mathbf{V}_h), \quad \mathbf{Q}_h = \mathbf{X}\mathbf{W}_h^{(q)}, \quad \mathbf{K}_h = \mathbf{X}\mathbf{W}_h^{(k)}, \quad \mathbf{V}_h = \mathbf{X}\mathbf{W}_h^{(v)}.$$

The **heads are first concatenated** into a single matrix, and the result is then linearly transformed using a matrix $\mathbf{W}^{(o)}$ to give a combined output in the form

$$\mathbf{Y}(\mathbf{X}) = \text{Concat}(\mathbf{H}_1, \dots, \mathbf{H}_H) \mathbf{W}^{(o)}, \in \mathbb{R}^{N \times D_v}.$$

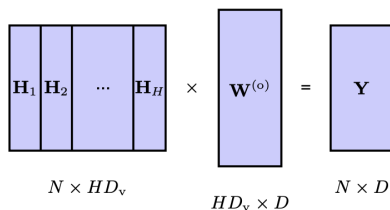


Figure: Network architecture for multi-head attention.

Literature review

MODELING MULTIPLE DEPENDENCY PATTERNS

To capture different types of relationships in the data simultaneously we replicate the attention mechanism into **multiple heads**.

Suppose we have H attention heads indexed by $h = 1, \dots, H$ of the form

$$\mathbf{H}_h = \text{Attention}(\mathbf{Q}_h, \mathbf{K}_h, \mathbf{V}_h), \quad \mathbf{Q}_h = \mathbf{X}\mathbf{W}_h^{(q)}, \quad \mathbf{K}_h = \mathbf{X}\mathbf{W}_h^{(k)}, \quad \mathbf{V}_h = \mathbf{X}\mathbf{W}_h^{(v)}.$$

The **heads are first concatenated** into a single matrix, and the result is then linearly transformed using a matrix $\mathbf{W}^{(o)}$ to give a combined output in the form

$$\mathbf{Y}(\mathbf{X}) = \text{Concat}(\mathbf{H}_1, \dots, \mathbf{H}_H) \mathbf{W}^{(o)}, \in \mathbb{R}^{N \times D_v}.$$

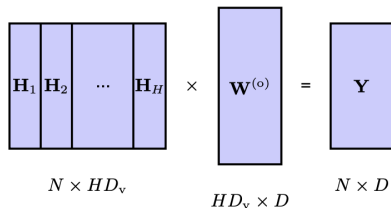


Figure: Network architecture for multi-head attention.

Literature review

MODELING MULTIPLE DEPENDENCY PATTERNS

To capture different types of relationships in the data simultaneously we replicate the attention mechanism into **multiple heads**.

Suppose we have H attention heads indexed by $h = 1, \dots, H$ of the form

$$\mathbf{H}_h = \text{Attention}(\mathbf{Q}_h, \mathbf{K}_h, \mathbf{V}_h), \quad \mathbf{Q}_h = \mathbf{X}\mathbf{W}_h^{(q)}, \quad \mathbf{K}_h = \mathbf{X}\mathbf{W}_h^{(k)}, \quad \mathbf{V}_h = \mathbf{X}\mathbf{W}_h^{(v)}.$$

The **heads are first concatenated** into a single matrix, and the result is then linearly transformed using a matrix $\mathbf{W}^{(o)}$ to give a combined output in the form

$$\mathbf{Y}(\mathbf{X}) = \text{Concat}(\mathbf{H}_1, \dots, \mathbf{H}_H) \mathbf{W}^{(o)}, \in \mathbb{R}^{N \times D_v}.$$

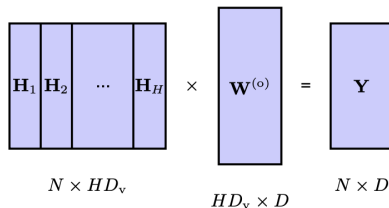


Figure: Network architecture for multi-head attention.

Literature review

MODELING MULTIPLE DEPENDENCY PATTERNS

To capture different types of relationships in the data simultaneously we replicate the attention mechanism into **multiple heads**.

Suppose we have H attention heads indexed by $h = 1, \dots, H$ of the form

$$\mathbf{H}_h = \text{Attention}(\mathbf{Q}_h, \mathbf{K}_h, \mathbf{V}_h), \quad \mathbf{Q}_h = \mathbf{X}\mathbf{W}_h^{(q)}, \quad \mathbf{K}_h = \mathbf{X}\mathbf{W}_h^{(k)}, \quad \mathbf{V}_h = \mathbf{X}\mathbf{W}_h^{(v)}.$$

The **heads are first concatenated** into a single matrix, and the result is then linearly transformed using a matrix $\mathbf{W}^{(o)}$ to give a combined output in the form

$$\mathbf{Y}(\mathbf{X}) = \text{Concat}(\mathbf{H}_1, \dots, \mathbf{H}_H) \mathbf{W}^{(o)}, \in \mathbb{R}^{N \times D_v}.$$

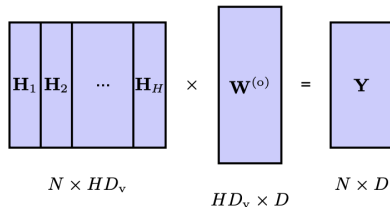


Figure: Network architecture for multi-head attention.

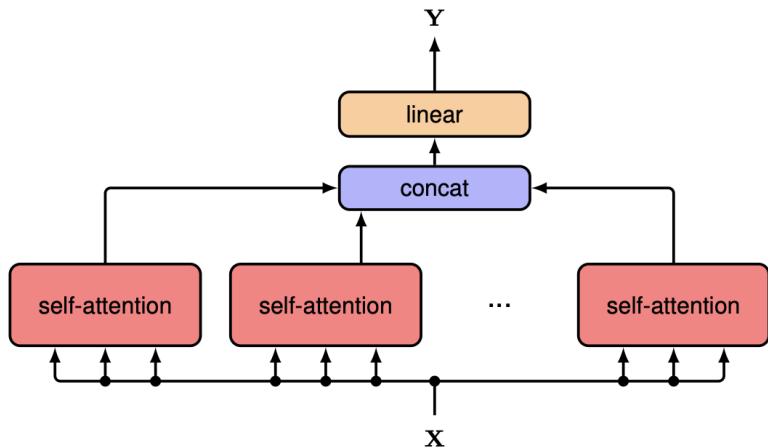


Figure: Information flow in a multi-head attention layer.

Literature Review

TRAINING IMPROVEMENTS

To improve training efficiency, the output is followed by a *residual connection* and a *layer normalization* (Ba, Kiros, and Hinton, 2016)

$$\mathbf{Z} = \text{LayerNorm}(\mathbf{Y}(\mathbf{X}) + \mathbf{X}).$$

STILL WORK TO DO

An MLP is added after the attention layer to break the linearity of the output and increase the expressive capabilities of the attention layer:

$$\tilde{\mathbf{X}} = \text{LayerNorm}(\text{MLP}(\mathbf{Z}) + \mathbf{Z}).$$

Again, this neural network layer can be improved by using a *residual connection* and *layer normalization*. A short notation:

$$\tilde{\mathbf{X}} = \text{TF}_{\theta}(\mathbf{X})$$

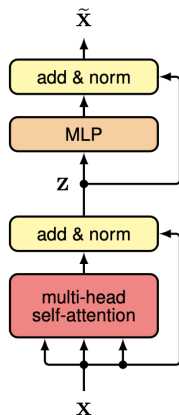


Figure: One layer of the transformer architecture. Here, 'add and norm' denotes a residual connection followed by layer normalization.

Literature Review

TRAINING IMPROVEMENTS

To improve training efficiency, the output is followed by a *residual connection* and a *layer normalization* (Ba, Kiros, and Hinton, 2016)

$$\mathbf{Z} = \text{LayerNorm}(\mathbf{Y}(\mathbf{X}) + \mathbf{X}).$$

STILL WORK TO DO

An MLP is added after the attention layer to break the linearity of the output and increase the expressive capabilities of the attention layer:

$$\tilde{\mathbf{X}} = \text{LayerNorm}(\text{MLP}(\mathbf{Z}) + \mathbf{Z}).$$

Again, this neural network layer can be improved by using a *residual connection* and *layer normalization*. A short notation:

$$\tilde{\mathbf{X}} = \text{TF}_{\theta}(\mathbf{X})$$

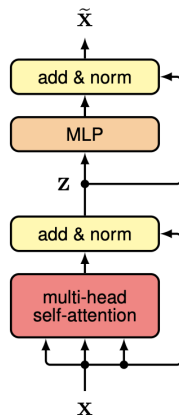


Figure: One layer of the transformer architecture. Here, 'add and norm' denotes a residual connection followed by layer normalization.

Literature Review

TRAINING IMPROVEMENTS

To **improve training efficiency**, the output is followed by a *residual connection* and a *layer normalization* (Ba, Kiros, and Hinton, 2016)

$$\mathbf{Z} = \text{LayerNorm}(\mathbf{Y}(\mathbf{X}) + \mathbf{X}).$$

STILL WORK TO DO

An MLP is added after the attention layer to **break the linearity of the output** and increase the expressive capabilities of the attention layer:

$$\tilde{\mathbf{X}} = \text{LayerNorm}(\text{MLP}(\mathbf{Z}) + \mathbf{Z}).$$

Again, this neural network layer can be improved by using a *residual connection* and *layer normalization*. A short notation:

$$\tilde{\mathbf{X}} = \text{TF}_{\theta}(\mathbf{X})$$

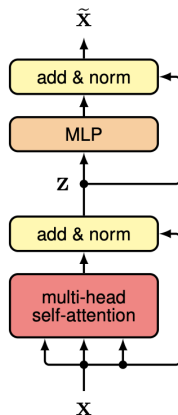


Figure: One layer of the transformer architecture. Here, 'add and norm' denotes a residual connection followed by layer normalization.

Literature Review

TRAINING IMPROVEMENTS

To **improve training efficiency**, the output is followed by a *residual connection* and a *layer normalization* (Ba, Kiros, and Hinton, 2016)

$$\mathbf{Z} = \text{LayerNorm}(\mathbf{Y}(\mathbf{X}) + \mathbf{X}).$$

STILL WORK TO DO

An MLP is added after the attention layer to **break the linearity of the output** and increase the expressive capabilities of the attention layer:

$$\tilde{\mathbf{X}} = \text{LayerNorm}(\text{MLP}(\mathbf{Z}) + \mathbf{Z}).$$

Again, this neural network layer can be improved by using a *residual connection* and *layer normalization*. A short notation:

$$\tilde{\mathbf{X}} = \text{TF}_{\theta}(\mathbf{X})$$

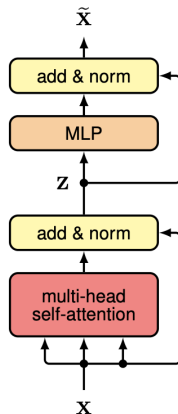


Figure: One layer of the transformer architecture. Here, 'add and norm' denotes a residual connection followed by layer normalization.

Literature Review

TRAINING IMPROVEMENTS

To **improve training efficiency**, the output is followed by a *residual connection* and a *layer normalization* (Ba, Kiros, and Hinton, 2016)

$$\mathbf{Z} = \text{LayerNorm}(\mathbf{Y}(\mathbf{X}) + \mathbf{X}).$$

STILL WORK TO DO

An MLP is added after the attention layer to **break the linearity of the output** and increase the expressive capabilities of the attention layer:

$$\tilde{\mathbf{X}} = \text{LayerNorm}(\text{MLP}(\mathbf{Z}) + \mathbf{Z}).$$

Again, this neural network layer can be improved by using a *residual connection* and *layer normalization*. A short notation:

$$\tilde{\mathbf{X}} = \text{TF}_{\theta}(\mathbf{X})$$

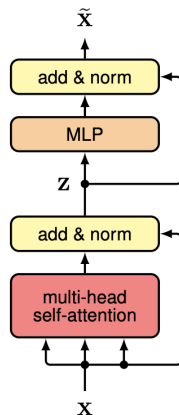


Figure: One layer of the transformer architecture. Here, 'add and norm' denotes a residual connection followed by layer normalization.

Transformers in Statistics

Garg et al. (2022) show empirically that **standard Transformers** can be trained from scratch to perform in-context learning of linear functions.

NOTATION

Let $\mathcal{X}$ be the input feature space, and $\mathcal{Y}$ be the output/label space, such that $\mathcal{F} = \{f : \mathcal{X} \rightarrow \mathcal{Y}\}$ is **class of functions** between this spaces.

IN-CONTEXT LEARNING SETTING

- 1) Sample $f \sim P_{\mathcal{F}}$ and N i.i.d. inputs tokens $\mathbf{x}_1, \dots, \mathbf{x}_N \sim P_{\mathcal{X}}$.
- 2) Generate labels $\mathbf{y}_i = f(\mathbf{x}_i) \in \mathcal{Y}$ for $i = 1, \dots, N$, obtaining $\{(\mathbf{x}_i, \mathbf{y}_i)\}_{i=1}^N$.
- 3) Define a length- i prompt as

$$\mathbf{X}_{\text{prompt}}^{(i)} = (\mathbf{x}_1, \mathbf{y}_1, \dots, \mathbf{x}_{i-1}, \mathbf{y}_{i-1}, \mathbf{x}_i)^T,$$

such that the **model prediction** is $\hat{\mathbf{y}}_i = \mathbf{W} \text{TF}_{\boldsymbol{\theta}}(\mathbf{X}_{\text{prompt}}^{(i)}) \in \mathcal{Y}$, $i = 1, \dots, N$.

- 4) Training by minimizing the expected loss

$$\min_{\boldsymbol{\theta}} \mathbb{E}_{\substack{\mathbf{x}_1, \dots, \mathbf{x}_N \sim P_{\mathcal{X}} \\ f \sim P_{\mathcal{F}}}} \left[\frac{1}{N} \sum_{i=1}^N \ell(\hat{\mathbf{y}}_i, \mathbf{y}_i) \right],$$

where ℓ is an appropriate loss function.

Transformers in Statistics

Garg et al. (2022) show empirically that **standard Transformers** can be trained from scratch to perform in-context learning of linear functions.

NOTATION

Let $\mathcal{X}$ be the input feature space, and $\mathcal{Y}$ be the output/label space, such that $\mathcal{F} = \{f : \mathcal{X} \rightarrow \mathcal{Y}\}$ is **class of functions** between this spaces.

IN-CONTEXT LEARNING SETTING

- 1) Sample $f \sim P_{\mathcal{F}}$ and N i.i.d. inputs tokens $\mathbf{x}_1, \dots, \mathbf{x}_N \sim P_{\mathcal{X}}$.
- 2) Generate labels $\mathbf{y}_i = f(\mathbf{x}_i) \in \mathcal{Y}$ for $i = 1, \dots, N$, obtaining $\{(\mathbf{x}_i, \mathbf{y}_i)\}_{i=1}^N$.
- 3) Define a length- i prompt as

$$\mathbf{X}_{\text{prompt}}^{(i)} = (\mathbf{x}_1, \mathbf{y}_1, \dots, \mathbf{x}_{i-1}, \mathbf{y}_{i-1}, \mathbf{x}_i)^T,$$

such that the **model prediction** is $\hat{\mathbf{y}}_i = \mathbf{W} \text{TF}_{\boldsymbol{\theta}}(\mathbf{X}_{\text{prompt}}^{(i)}) \in \mathcal{Y}$, $i = 1, \dots, N$.

- 4) Training by minimizing the expected loss

$$\min_{\boldsymbol{\theta}} \mathbb{E}_{\substack{\mathbf{x}_1, \dots, \mathbf{x}_N \sim P_{\mathcal{X}} \\ f \sim P_{\mathcal{F}}}} \left[\frac{1}{N} \sum_{i=1}^N \ell(\hat{\mathbf{y}}_i, \mathbf{y}_i) \right],$$

where ℓ is an appropriate loss function.

Transformers in Statistics

Garg et al. (2022) show empirically that **standard Transformers** can be trained from scratch to perform in-context learning of linear functions.

NOTATION

Let $\mathcal{X}$ be the input feature space, and $\mathcal{Y}$ be the output/label space, such that $\mathcal{F} = \{f : \mathcal{X} \rightarrow \mathcal{Y}\}$ is **class of functions** between this spaces.

IN-CONTEXT LEARNING SETTING

- 1) Sample $f \sim P_{\mathcal{F}}$ and N i.i.d. inputs tokens $\mathbf{x}_1, \dots, \mathbf{x}_N \sim P_{\mathcal{X}}$.
- 2) Generate labels $\mathbf{y}_i = f(\mathbf{x}_i) \in \mathcal{Y}$ for $i = 1, \dots, N$, obtaining $\{(\mathbf{x}_i, \mathbf{y}_i)\}_{i=1}^N$.
- 3) Define a length- i prompt as

$$\mathbf{X}_{\text{prompt}}^{(i)} = (\mathbf{x}_1, \mathbf{y}_1, \dots, \mathbf{x}_{i-1}, \mathbf{y}_{i-1}, \mathbf{x}_i)^T,$$

such that the **model prediction** is $\hat{\mathbf{y}}_i = \mathbf{W} \text{TF}_{\boldsymbol{\theta}}(\mathbf{X}_{\text{prompt}}^{(i)}) \in \mathcal{Y}$, $i = 1, \dots, N$.

- 4) Training by minimizing the expected loss

$$\min_{\boldsymbol{\theta}} \mathbb{E}_{\substack{\mathbf{x}_1, \dots, \mathbf{x}_N \sim P_{\mathcal{X}} \\ f \sim P_{\mathcal{F}}}} \left[\frac{1}{N} \sum_{i=1}^N \ell(\hat{\mathbf{y}}_i, \mathbf{y}_i) \right],$$

where ℓ is an appropriate loss function.

Transformers in Statistics

Garg et al. (2022) show empirically that **standard Transformers** can be trained from scratch to perform in-context learning of linear functions.

NOTATION

Let $\mathcal{X}$ be the input feature space, and $\mathcal{Y}$ be the output/label space, such that $\mathcal{F} = \{f : \mathcal{X} \rightarrow \mathcal{Y}\}$ is **class of functions** between this spaces.

IN-CONTEXT LEARNING SETTING

- 1) Sample $f \sim P_{\mathcal{F}}$ and N i.i.d. inputs tokens $\mathbf{x}_1, \dots, \mathbf{x}_N \sim P_{\mathcal{X}}$.
- 2) Generate labels $\mathbf{y}_i = f(\mathbf{x}_i) \in \mathcal{Y}$ for $i = 1, \dots, N$, obtaining $\{(\mathbf{x}_i, \mathbf{y}_i)\}_{i=1}^N$.
- 3) Define a length- i prompt as

$$\mathbf{X}_{\text{prompt}}^{(i)} = (\mathbf{x}_1, \mathbf{y}_1, \dots, \mathbf{x}_{i-1}, \mathbf{y}_{i-1}, \mathbf{x}_i)^T,$$

such that the **model prediction** is $\hat{\mathbf{y}}_i = \mathbf{W} \text{TF}_{\boldsymbol{\theta}}(\mathbf{X}_{\text{prompt}}^{(i)}) \in \mathcal{Y}$, $i = 1, \dots, N$.

- 4) Training by minimizing the expected loss

$$\min_{\boldsymbol{\theta}} \mathbb{E}_{\substack{\mathbf{x}_1, \dots, \mathbf{x}_N \sim P_{\mathcal{X}} \\ f \sim P_{\mathcal{F}}}} \left[\frac{1}{N} \sum_{i=1}^N \ell(\hat{\mathbf{y}}_i, \mathbf{y}_i) \right],$$

where ℓ is an appropriate loss function.

Transformers in Statistics

Garg et al. (2022) show empirically that **standard Transformers** can be trained from scratch to perform in-context learning of linear functions.

NOTATION

Let $\mathcal{X}$ be the input feature space, and $\mathcal{Y}$ be the output/label space, such that $\mathcal{F} = \{f : \mathcal{X} \rightarrow \mathcal{Y}\}$ is **class of functions** between this spaces.

IN-CONTEXT LEARNING SETTING

- 1) Sample $f \sim P_{\mathcal{F}}$ and N i.i.d. inputs tokens $\mathbf{x}_1, \dots, \mathbf{x}_N \sim P_{\mathcal{X}}$.
- 2) Generate labels $\mathbf{y}_i = f(\mathbf{x}_i) \in \mathcal{Y}$ for $i = 1, \dots, N$, obtaining $\{(\mathbf{x}_i, \mathbf{y}_i)\}_{i=1}^N$.
- 3) Define a length- i prompt as

$$\mathbf{X}_{\text{prompt}}^{(i)} = (\mathbf{x}_1, \mathbf{y}_1, \dots, \mathbf{x}_{i-1}, \mathbf{y}_{i-1}, \mathbf{x}_i)^{\top},$$

such that the **model prediction** is $\hat{\mathbf{y}}_i = \mathbf{W} \text{TF}_{\theta}(\mathbf{X}_{\text{prompt}}^{(i)}) \in \mathcal{Y}$, $i = 1, \dots, N$.

- 4) Training by minimizing the expected loss

$$\min_{\theta} \mathbb{E}_{\substack{\mathbf{x}_1, \dots, \mathbf{x}_n \sim P_{\mathcal{X}} \\ f \sim P_{\mathcal{F}}}} \left[\frac{1}{N} \sum_{i=1}^N \ell(\hat{\mathbf{y}}_i, \mathbf{y}_i) \right],$$

where ℓ is an appropriate loss function.

Transformers in Statistics

Garg et al. (2022) show empirically that **standard Transformers** can be trained from scratch to perform in-context learning of linear functions.

NOTATION

Let $\mathcal{X}$ be the input feature space, and $\mathcal{Y}$ be the output/label space, such that $\mathcal{F} = \{f : \mathcal{X} \rightarrow \mathcal{Y}\}$ is **class of functions** between this spaces.

IN-CONTEXT LEARNING SETTING

- 1) Sample $f \sim P_{\mathcal{F}}$ and N i.i.d. inputs tokens $\mathbf{x}_1, \dots, \mathbf{x}_N \sim P_{\mathcal{X}}$.
- 2) Generate labels $\mathbf{y}_i = f(\mathbf{x}_i) \in \mathcal{Y}$ for $i = 1, \dots, N$, obtaining $\{(\mathbf{x}_i, \mathbf{y}_i)\}_{i=1}^N$.
- 3) Define a length- i prompt as

$$\mathbf{X}_{\text{prompt}}^{(i)} = (\mathbf{x}_1, \mathbf{y}_1, \dots, \mathbf{x}_{i-1}, \mathbf{y}_{i-1}, \mathbf{x}_i)^T,$$

such that the **model prediction** is $\hat{\mathbf{y}}_i = \mathbf{W}^{\text{TF}}_{\theta}(\mathbf{X}_{\text{prompt}}^{(i)}) \in \mathcal{Y}$, $i = 1, \dots, N$.

- 4) Training by minimizing the expected loss

$$\min_{\theta} \mathbb{E}_{\substack{\mathbf{x}_1, \dots, \mathbf{x}_n \sim P_{\mathcal{X}} \\ f \sim P_{\mathcal{F}}}} \left[\frac{1}{N} \sum_{i=1}^N \ell(\hat{\mathbf{y}}_i, \mathbf{y}_i) \right],$$

where ℓ is an appropriate loss function.

Transformers in Statistics

Garg et al. (2022) show empirically that **standard Transformers** can be trained from scratch to perform in-context learning of linear functions.

NOTATION

Let $\mathcal{X}$ be the input feature space, and $\mathcal{Y}$ be the output/label space, such that $\mathcal{F} = \{f : \mathcal{X} \rightarrow \mathcal{Y}\}$ is **class of functions** between this spaces.

IN-CONTEXT LEARNING SETTING

- 1) Sample $f \sim P_{\mathcal{F}}$ and N i.i.d. inputs tokens $\mathbf{x}_1, \dots, \mathbf{x}_N \sim P_{\mathcal{X}}$.
- 2) Generate labels $\mathbf{y}_i = f(\mathbf{x}_i) \in \mathcal{Y}$ for $i = 1, \dots, N$, obtaining $\{(\mathbf{x}_i, \mathbf{y}_i)\}_{i=1}^N$.
- 3) Define a length- i prompt as

$$\mathbf{X}_{\text{prompt}}^{(i)} = (\mathbf{x}_1, \mathbf{y}_1, \dots, \mathbf{x}_{i-1}, \mathbf{y}_{i-1}, \mathbf{x}_i)^T,$$

such that the **model prediction** is $\hat{\mathbf{y}}_i = \mathbf{W}^{\text{TF}}_{\boldsymbol{\theta}}(\mathbf{X}_{\text{prompt}}^{(i)}) \in \mathcal{Y}$, $i = 1, \dots, N$.

- 4) Training by minimizing the expected loss

$$\min_{\boldsymbol{\theta}} \mathbb{E}_{\substack{\mathbf{x}_1, \dots, \mathbf{x}_n \sim P_{\mathcal{X}} \\ f \sim P_{\mathcal{F}}}} \left[\frac{1}{N} \sum_{i=1}^N \ell(\hat{\mathbf{y}}_i, \mathbf{y}_i) \right],$$

where ℓ is an appropriate loss function.

Transformers in Statistics

EXAMPLE (LINEAR FUNCTIONS)

Consider the class of linear functions

$$\mathcal{F} = \{f : \mathbb{R}^d \rightarrow \mathbb{R} \mid f(\mathbf{x}) = \mathbf{w}^T \mathbf{x}, \mathbf{w} \in \mathbb{R}^d\}.$$

Sample

$$\mathbf{w} \sim P_{\mathcal{F}} := N_d(\mathbf{0}, \mathbf{I}_d), \quad \mathbf{x}_i \stackrel{\text{i.i.d.}}{\sim} P_{\mathcal{X}} := N_d(\mathbf{0}, \mathbf{I}_d), \quad i = 1, \dots, N,$$

independently. For each f determined by $\mathbf{w}$, generate training pairs

$$\mathbf{y}_i = f(\mathbf{x}_i) = \mathbf{w}^T \mathbf{x}_i.$$

Then, train the Transformer to predict $\hat{\mathbf{y}}_i = \mathbf{W} \text{TF}_{\theta}(\mathbf{X}_{\text{prompt}}^{(i)})$ where

$$\mathbf{X}_{\text{prompt}}^{(i)} = (\mathbf{x}_1, \mathbf{y}_1, \dots, \mathbf{x}_{i-1}, \mathbf{y}_{i-1}, \mathbf{x}_i),$$

by minimizing the expected mean squared error loss

$$\min_{\theta} \mathbb{E}_{\substack{\mathbf{x}_1, \dots, \mathbf{x}_N \sim N_d(\mathbf{0}, \mathbf{I}_d) \\ f \sim N_d(\mathbf{0}, \mathbf{I}_d)}} \left[\frac{1}{N} \sum_{i=1}^N (\hat{\mathbf{y}}_i - \mathbf{y}_i)^2 \right].$$

Transformers in Statistics

EXAMPLE (LINEAR FUNCTIONS)

Consider the **class of linear functions**

$$\mathcal{F} = \{f : \mathbb{R}^d \rightarrow \mathbb{R} \mid f(\mathbf{x}) = \mathbf{w}^T \mathbf{x}, \mathbf{w} \in \mathbb{R}^d\}.$$

Sample

$$\mathbf{w} \sim P_{\mathcal{F}} := N_d(\mathbf{0}, \mathbf{I}_d), \quad \mathbf{x}_i \stackrel{\text{i.i.d.}}{\sim} P_{\mathcal{X}} := N_d(\mathbf{0}, \mathbf{I}_d), \quad i = 1, \dots, N,$$

independently. For **each f determined by $\mathbf{w}$** , generate training pairs

$$\mathbf{y}_i = f(\mathbf{x}_i) = \mathbf{w}^T \mathbf{x}_i.$$

Then, train the Transformer to predict $\hat{\mathbf{y}}_i = \mathbf{W} \text{TF}_{\theta}(\mathbf{X}_{\text{prompt}}^{(i)})$ where

$$\mathbf{X}_{\text{prompt}}^{(i)} = (\mathbf{x}_1, \mathbf{y}_1, \dots, \mathbf{x}_{i-1}, \mathbf{y}_{i-1}, \mathbf{x}_i),$$

by minimizing the expected **mean squared error loss**

$$\min_{\theta} \mathbb{E}_{\substack{\mathbf{x}_1, \dots, \mathbf{x}_N \sim N_d(\mathbf{0}, \mathbf{I}_d) \\ f \sim N_d(\mathbf{0}, \mathbf{I}_d)}} \left[\frac{1}{N} \sum_{i=1}^N (\hat{\mathbf{y}}_i - \mathbf{y}_i)^2 \right].$$

Transformers in Statistics

EXAMPLE (LINEAR FUNCTIONS)

Consider the **class of linear functions**

$$\mathcal{F} = \{f : \mathbb{R}^d \rightarrow \mathbb{R} \mid f(\mathbf{x}) = \mathbf{w}^T \mathbf{x}, \mathbf{w} \in \mathbb{R}^d\}.$$

Sample

$$\mathbf{w} \sim \mathbf{P}_{\mathcal{F}} := \mathcal{N}_d(\mathbf{0}, \mathbf{I}_d), \quad \mathbf{x}_i \stackrel{\text{i.i.d.}}{\sim} \mathbf{P}_{\mathcal{X}} := \mathcal{N}_d(\mathbf{0}, \mathbf{I}_d), \quad i = 1, \dots, N,$$

independently. For each f determined by $\mathbf{w}$, generate training pairs

$$\mathbf{y}_i = f(\mathbf{x}_i) = \mathbf{w}^T \mathbf{x}_i.$$

Then, train the Transformer to predict $\hat{\mathbf{y}}_i = \mathbf{W} \text{TF}_{\theta}(\mathbf{X}_{\text{prompt}}^{(i)})$ where

$$\mathbf{X}_{\text{prompt}}^{(i)} = (\mathbf{x}_1, \mathbf{y}_1, \dots, \mathbf{x}_{i-1}, \mathbf{y}_{i-1}, \mathbf{x}_i),$$

by minimizing the expected **mean squared error loss**

$$\min_{\theta} \mathbb{E}_{\substack{\mathbf{x}_1, \dots, \mathbf{x}_N \sim \mathcal{N}_d(\mathbf{0}, \mathbf{I}_d) \\ f \sim \mathcal{N}_d(\mathbf{0}, \mathbf{I}_d)}} \left[\frac{1}{N} \sum_{i=1}^N (\hat{\mathbf{y}}_i - \mathbf{y}_i)^2 \right].$$

Transformers in Statistics

EXAMPLE (LINEAR FUNCTIONS)

Consider the **class of linear functions**

$$\mathcal{F} = \{f : \mathbb{R}^d \rightarrow \mathbb{R} \mid f(\mathbf{x}) = \mathbf{w}^T \mathbf{x}, \mathbf{w} \in \mathbb{R}^d\}.$$

Sample

$$\mathbf{w} \sim P_{\mathcal{F}} := N_d(\mathbf{0}, \mathbf{I}_d), \quad \mathbf{x}_i \stackrel{\text{i.i.d.}}{\sim} P_{\mathcal{X}} := N_d(\mathbf{0}, \mathbf{I}_d), \quad i = 1, \dots, N,$$

independently. For **each f determined by $\mathbf{w}$** , generate training pairs

$$\mathbf{y}_i = f(\mathbf{x}_i) = \mathbf{w}^T \mathbf{x}_i.$$

Then, train the Transformer to predict $\hat{\mathbf{y}}_i = \mathbf{W} \text{TF}_{\theta}(\mathbf{X}_{\text{prompt}}^{(i)})$ where

$$\mathbf{X}_{\text{prompt}}^{(i)} = (\mathbf{x}_1, \mathbf{y}_1, \dots, \mathbf{x}_{i-1}, \mathbf{y}_{i-1}, \mathbf{x}_i),$$

by minimizing the expected **mean squared error loss**

$$\min_{\theta} \mathbb{E}_{\substack{\mathbf{x}_1, \dots, \mathbf{x}_N \sim N_d(\mathbf{0}, \mathbf{I}_d) \\ f \sim N_d(\mathbf{0}, \mathbf{I}_d)}} \left[\frac{1}{N} \sum_{i=1}^N (\hat{\mathbf{y}}_i - \mathbf{y}_i)^2 \right].$$

Transformers in Statistics

EXAMPLE (LINEAR FUNCTIONS)

Consider the [class of linear functions](#)

$$\mathcal{F} = \{f : \mathbb{R}^d \rightarrow \mathbb{R} \mid f(\mathbf{x}) = \mathbf{w}^T \mathbf{x}, \mathbf{w} \in \mathbb{R}^d\}.$$

Sample

$$\mathbf{w} \sim P_{\mathcal{F}} := N_d(\mathbf{0}, \mathbf{I}_d), \quad \mathbf{x}_i \stackrel{\text{i.i.d.}}{\sim} P_{\mathcal{X}} := N_d(\mathbf{0}, \mathbf{I}_d), \quad i = 1, \dots, N,$$

independently. For [each \$f\$ determined by \$\mathbf{w}\$](#) , generate training pairs

$$\mathbf{y}_i = f(\mathbf{x}_i) = \mathbf{w}^T \mathbf{x}_i.$$

Then, train the Transformer to predict $\hat{\mathbf{y}}_i = \mathbf{W} \text{TF}_{\theta}(\mathbf{X}_{\text{prompt}}^{(i)})$ where

$$\mathbf{X}_{\text{prompt}}^{(i)} = (\mathbf{x}_1, \mathbf{y}_1, \dots, \mathbf{x}_{i-1}, \mathbf{y}_{i-1}, \mathbf{x}_i),$$

by minimizing the expected [mean squared error loss](#)

$$\min_{\theta} \mathbb{E}_{\substack{\mathbf{x}_1, \dots, \mathbf{x}_N \sim N_d(\mathbf{0}, \mathbf{I}_d) \\ f \sim N_d(\mathbf{0}, \mathbf{I}_d)}} \left[\frac{1}{N} \sum_{i=1}^N (\hat{\mathbf{y}}_i - \mathbf{y}_i)^2 \right].$$

Transformers in Statistics

In this setting, the Transformer learns to approximate the least squares estimator.

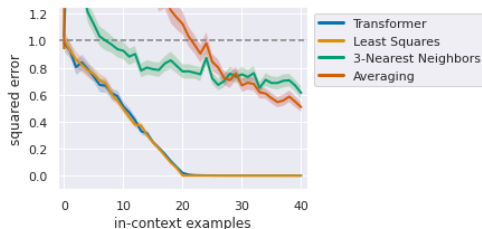


Figure: Normalized squared error of the Transformer as a function of the number of in-context examples ([Garg et al., 2022](#)).

[Bai et al. \(2023\)](#) show¹ that Transformers can implement a broad class of standard ML algorithms in context, such as least squares, Ridge regression, Lasso, convex risk minimization for GLMs.

¹Possibly the most mathematically formal paper on this topic.

Transformers in Statistics

In this setting, the Transformer learns to approximate the least squares estimator.

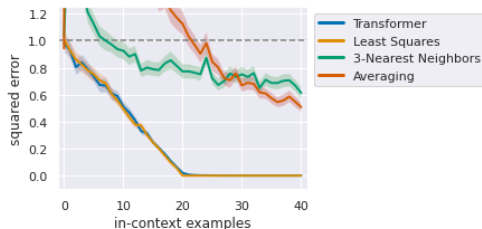


Figure: Normalized squared error of the Transformer as a function of the number of in-context examples ([Garg et al., 2022](#)).

[Bai et al. \(2023\)](#) show¹ that Transformers can implement a broad class of standard ML algorithms in context, such as least squares, Ridge regression, Lasso, convex risk minimization for GLMs.

¹Possibly the most mathematically formal paper on this topic.

Large Language Models

DECODER TRANSFORMERS

These can be used as **generative models** that create output sequences of tokens. we will focus on a class of models called GPT which stands for **generative pre-trained transformer**.

GOAL

Construct an **autoregressive model** of the form defined by

$$p(\mathbf{x}_1, \dots, \mathbf{x}_N) = \prod_{n=1}^N p(\mathbf{x}_n \mid \mathbf{x}_{1:n-1}), \quad \mathbf{x}_{1:n-1} = (\mathbf{x}_1, \dots, \mathbf{x}_{n-1}),$$

where conditional distributions $p(\mathbf{x}_n \mid \mathbf{x}_{1:n-1})$ are expressed using a Transformer neural network that is **learned from data**.

IDEA

$$\mathbf{x}_1, \dots, \mathbf{x}_N \xrightarrow{\text{Transformer layers}} \tilde{\mathbf{x}}_1, \dots, \tilde{\mathbf{x}}_N,$$

Each output needs to represent a **probability distribution** over the class of tokens:

$$\mathbf{Y} = \text{Softmax}(\tilde{\mathbf{X}}\mathbf{W}^{(p)}), \quad \mathbf{W}^{(p)} \in \mathbb{R}^{D \times K}.$$

Large Language Models

DECODER TRANSFORMERS

These can be used as **generative models** that create output sequences of tokens. we will focus on a class of models called GPT which stands for **generative pre-trained transformer**.

GOAL

Construct an **autoregressive model** of the form defined by

$$p(\mathbf{x}_1, \dots, \mathbf{x}_N) = \prod_{n=1}^N p(\mathbf{x}_n \mid \mathbf{x}_{1:n-1}), \quad \mathbf{x}_{1:n-1} = (\mathbf{x}_1, \dots, \mathbf{x}_{n-1}),$$

where conditional distributions $p(\mathbf{x}_n \mid \mathbf{x}_{1:n-1})$ are expressed using a Transformer neural network that is **learned from data**.

IDEA

$$\mathbf{x}_1, \dots, \mathbf{x}_N \xrightarrow{\text{Transformer layers}} \tilde{\mathbf{x}}_1, \dots, \tilde{\mathbf{x}}_N,$$

Each output needs to represent a **probability distribution** over the class of tokens:

$$\mathbf{Y} = \text{Softmax}(\tilde{\mathbf{X}}\mathbf{W}^{(p)}), \quad \mathbf{W}^{(p)} \in \mathbb{R}^{D \times K}.$$

Large Language Models

DECODER TRANSFORMERS

These can be used as **generative models** that create output sequences of tokens. we will focus on a class of models called GPT which stands for **generative pre-trained transformer**.

GOAL

Construct an **autoregressive model** of the form defined by

$$p(\mathbf{x}_1, \dots, \mathbf{x}_N) = \prod_{n=1}^N p(\mathbf{x}_n \mid \mathbf{x}_{1:n-1}), \quad \mathbf{x}_{1:n-1} = (\mathbf{x}_1, \dots, \mathbf{x}_{n-1}),$$

where conditional distributions $p(\mathbf{x}_n \mid \mathbf{x}_{1:n-1})$ are expressed using a Transformer neural network that is **learned from data**.

IDEA

$$\mathbf{x}_1, \dots, \mathbf{x}_N \xrightarrow{\text{Transformer layers}} \tilde{\mathbf{x}}_1, \dots, \tilde{\mathbf{x}}_N,$$

Each output needs to represent a **probability distribution** over the class of tokens:

$$\mathbf{Y} = \text{Softmax}(\tilde{\mathbf{X}}\mathbf{W}^{(p)}), \quad \mathbf{W}^{(p)} \in \mathbb{R}^{D \times K}.$$

EMBEDDING AND POSITIONAL ENCODING

We begin by defining a **dictionary**

$$\mathcal{V} = \{w_1, w_2, \dots, w_K\},$$

which is a fixed, finite set of words indexed from 1 to K . A word $\mathbf{v}_i$ is mapped to an index in the dictionary using a tokenization function $\tau: \mathcal{V} \rightarrow \{1, \dots, K\}$, such that

$$\tau(\mathbf{v}_i) = k_i \quad \text{with} \quad \mathbf{v}_i = w_{k_i}.$$

This index is then used to construct a **one-hot encoded vector** $\mathbf{o}_i \in \{0, 1\}^K$, where the k_i -th component is set to 1 and all others are 0.

This one-hot vector is then projected into a dense, learnable embedding space using an embedding matrix $\mathbf{E} \in \mathbb{R}^{D \times K}$, and positional information (Vaswani et al. 2017) is added

$$\mathbf{x}_i = \mathbf{E}\mathbf{o}_i + \mathbf{r}_i.$$

EMBEDDING AND POSITIONAL ENCODING

We begin by defining a **dictionary**

$$\mathcal{V} = \{w_1, w_2, \dots, w_K\},$$

which is a fixed, finite set of words indexed from 1 to K . A word $\mathbf{v}_i$ is **mapped to an index in the dictionary** using a tokenization function $\tau : \mathcal{V} \rightarrow \{1, \dots, K\}$, such that

$$\tau(\mathbf{v}_i) = k_i \quad \text{with} \quad \mathbf{v}_i = w_{k_i}.$$

This index is then used to construct a **one-hot encoded vector** $\mathbf{o}_i \in \{0, 1\}^K$, where the k_i -th component is set to 1 and all others are 0.

This one-hot vector is then **projected into a dense, learnable embedding space** using an embedding matrix $\mathbf{E} \in \mathbb{R}^{D \times K}$, and positional information (Vaswani et al. 2017) is added

$$\mathbf{x}_i = \mathbf{E}\mathbf{o}_i + \mathbf{r}_i.$$

EMBEDDING AND POSITIONAL ENCODING

We begin by defining a **dictionary**

$$\mathcal{V} = \{w_1, w_2, \dots, w_K\},$$

which is a fixed, finite set of words indexed from 1 to K . A word $\mathbf{v}_i$ is **mapped to an index in the dictionary** using a tokenization function $\tau : \mathcal{V} \rightarrow \{1, \dots, K\}$, such that

$$\tau(\mathbf{v}_i) = k_i \quad \text{with} \quad \mathbf{v}_i = w_{k_i}.$$

This index is then used to construct a **one-hot encoded vector** $\mathbf{o}_i \in \{0, 1\}^K$, where the k_i -th component is set to 1 and all others are 0.

This one-hot vector is then **projected into a dense, learnable embedding space** using an embedding matrix $\mathbf{E} \in \mathbb{R}^{D \times K}$, and positional information (Vaswani et al. 2017) is added

$$\mathbf{x}_i = \mathbf{E}\mathbf{o}_i + \mathbf{r}_i.$$

EMBEDDING AND POSITIONAL ENCODING

We begin by defining a [dictionary](#)

$$\mathcal{V} = \{w_1, w_2, \dots, w_K\},$$

which is a fixed, finite set of words indexed from 1 to K . A word $\mathbf{v}_i$ is [mapped to an index in the dictionary](#) using a tokenization function $\tau : \mathcal{V} \rightarrow \{1, \dots, K\}$, such that

$$\tau(\mathbf{v}_i) = k_i \quad \text{with} \quad \mathbf{v}_i = w_{k_i}.$$

This index is then used to construct a **one-hot encoded vector** $\mathbf{o}_i \in \{0, 1\}^K$, where the k_i -th component is set to 1 and all others are 0.

This one-hot vector is then [projected into a dense, learnable embedding space](#) using an embedding matrix $\mathbf{E} \in \mathbb{R}^{D \times K}$, and positional information ([Vaswani et al. 2017](#)) is added

$$\mathbf{x}_i = \mathbf{E}\mathbf{o}_i + \mathbf{r}_i.$$

GPT Architecture

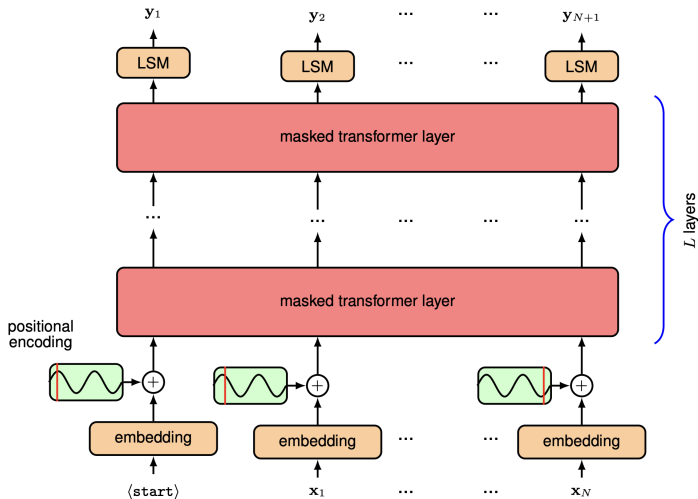


Figure: Architecture of a GPT decoder transformer network. Here 'LSM' stands for linear-softmax and denotes a linear transformation whose learnable parameters are shared across the token positions, followed by a softmax activation function.

HOW DO WE APPLY GPT TO GENERATE SEQUENCES OF EXTREME EVENTS?²

- ▶ Let $(X_1, \dots, X_d) \in \mathcal{X}^d$ be a **chain of extreme events**, e.g. $\mathcal{X} = \mathbb{R}_+$.
- ▶ Define a **finite partition** $\{\mathcal{C}_k\}_{k=1}^K$ of $\mathcal{X}$, and a **one-hot tokenization function** $\tau : \mathcal{X} \rightarrow \{0, 1\}^K$ defined as

$$\tau(X_i) = (I(X_i \in \mathcal{C}_1), \dots, I(X_i \in \mathcal{C}_K)), \quad i = 1, \dots, d.$$

- ▶ Each token is then projected into an **embedding space** using a **learnable matrix** $\mathbf{E} \in \mathbb{R}^{D \times K}$, positional encoding $\mathbf{r}_i$ and **extremal encoding**³ are added

$$\mathbf{s}_i = \mathbf{E}\tau(X_i) + \mathbf{r}_i + \mathbf{e}_i.$$

- ▶ Instead of modeling the joint density of $(X_1, \dots, X_d)$, we model the probability that each X_i falls in a set $\mathcal{C}_k$, using:

$$P(X_1 \in \mathcal{C}_{k_1}, \dots, X_d \in \mathcal{C}_{k_d}) = \prod_{i=1}^d P(X_i \in \mathcal{C}_{k_i} \mid X_{1:i-1}).$$

²Ongoing work with **Miguel de Carvalho** and **Ronny Vallejos**.

³e.g. based on extreme quantile regression models.

HOW DO WE APPLY GPT TO GENERATE SEQUENCES OF EXTREME EVENTS?²

- ▶ Let $(X_1, \dots, X_d) \in \mathcal{X}^d$ be a **chain of extreme events**, e.g. $\mathcal{X} = \mathbb{R}_+$.
- ▶ Define a **finite partition** $\{\mathcal{C}_k\}_{k=1}^K$ of $\mathcal{X}$, and a **one-hot tokenization function** $\tau : \mathcal{X} \rightarrow \{0, 1\}^K$ defined as

$$\tau(X_i) = (I(X_i \in \mathcal{C}_1), \dots, I(X_i \in \mathcal{C}_K)), \quad i = 1, \dots, d.$$

- ▶ Each token is then projected into an **embedding space** using a **learnable matrix** $\mathbf{E} \in \mathbb{R}^{D \times K}$, positional encoding $\mathbf{r}_i$ and **extremal encoding**³ are added

$$\mathbf{s}_i = \mathbf{E}\tau(X_i) + \mathbf{r}_i + \mathbf{e}_i.$$

- ▶ Instead of modeling the joint density of $(X_1, \dots, X_d)$, we model the probability that each X_i falls in a set $\mathcal{C}_k$, using:

$$P(X_1 \in \mathcal{C}_{k_1}, \dots, X_d \in \mathcal{C}_{k_d}) = \prod_{i=1}^d P(X_i \in \mathcal{C}_{k_i} \mid X_{1:i-1}).$$

²Ongoing work with **Miguel de Carvalho** and **Ronny Vallejos**.

³e.g. based on extreme quantile regression models.

HOW DO WE APPLY GPT TO GENERATE SEQUENCES OF EXTREME EVENTS?²

- ▶ Let $(X_1, \dots, X_d) \in \mathcal{X}^d$ be a **chain of extreme events**, e.g. $\mathcal{X} = \mathbb{R}_+$.
- ▶ Define a **finite partition** $\{\mathcal{C}_k\}_{k=1}^K$ of $\mathcal{X}$, and a **one-hot tokenization function** $\tau : \mathcal{X} \rightarrow \{0, 1\}^K$ defined as

$$\tau(X_i) = (I(X_i \in \mathcal{C}_1), \dots, I(X_i \in \mathcal{C}_K)), \quad i = 1, \dots, d.$$

- ▶ Each token is then projected into an **embedding space** using a **learnable matrix** $\mathbf{E} \in \mathbb{R}^{D \times K}$, positional encoding $\mathbf{r}_i$ and **extremal encoding**³ are added

$$\mathbf{s}_i = \mathbf{E}\tau(X_i) + \mathbf{r}_i + \mathbf{e}_i.$$

- ▶ Instead of modeling the joint density of $(X_1, \dots, X_d)$, we model the probability that each X_i falls in a set $\mathcal{C}_k$, using:

$$P(X_1 \in \mathcal{C}_{k_1}, \dots, X_d \in \mathcal{C}_{k_d}) = \prod_{i=1}^d P(X_i \in \mathcal{C}_{k_i} \mid X_{1:i-1}).$$

²Ongoing work with **Miguel de Carvalho** and **Ronny Vallejos**.

³e.g. based on extreme quantile regression models.

HOW DO WE APPLY GPT TO GENERATE SEQUENCES OF EXTREME EVENTS?²

- ▶ Let $(X_1, \dots, X_d) \in \mathcal{X}^d$ be a **chain of extreme events**, e.g. $\mathcal{X} = \mathbb{R}_+$.
- ▶ Define a **finite partition** $\{\mathcal{C}_k\}_{k=1}^K$ of $\mathcal{X}$, and a **one-hot tokenization function** $\tau : \mathcal{X} \rightarrow \{0, 1\}^K$ defined as

$$\tau(X_i) = (I(X_i \in \mathcal{C}_1), \dots, I(X_i \in \mathcal{C}_K)), \quad i = 1, \dots, d.$$

- ▶ Each token is then projected into an **embedding space** using a **learnable matrix** $\mathbf{E} \in \mathbb{R}^{D \times K}$, positional encoding $\mathbf{r}_i$ and **extremal encoding**³ are added

$$\mathbf{s}_i = \mathbf{E}\tau(X_i) + \mathbf{r}_i + \mathbf{e}_i.$$

- ▶ Instead of modeling the joint density of $(X_1, \dots, X_d)$, we model the probability that each X_i falls in a set $\mathcal{C}_k$, using:

$$P(X_1 \in \mathcal{C}_{k_1}, \dots, X_d \in \mathcal{C}_{k_d}) = \prod_{i=1}^d P(X_i \in \mathcal{C}_{k_i} \mid X_{1:i-1}).$$

²Ongoing work with **Miguel de Carvalho** and **Ronny Vallejos**.

³e.g. based on extreme quantile regression models.

HOW DO WE APPLY GPT TO GENERATE SEQUENCES OF EXTREME EVENTS?²

- ▶ Let $(X_1, \dots, X_d) \in \mathcal{X}^d$ be a **chain of extreme events**, e.g. $\mathcal{X} = \mathbb{R}_+$.
- ▶ Define a **finite partition** $\{\mathcal{C}_k\}_{k=1}^K$ of $\mathcal{X}$, and a **one-hot tokenization function** $\tau : \mathcal{X} \rightarrow \{0, 1\}^K$ defined as

$$\tau(X_i) = (I(X_i \in \mathcal{C}_1), \dots, I(X_i \in \mathcal{C}_K)), \quad i = 1, \dots, d.$$

- ▶ Each token is then projected into an **embedding space** using a **learnable matrix** $\mathbf{E} \in \mathbb{R}^{D \times K}$, positional encoding $\mathbf{r}_i$ and **extremal encoding**³ are added

$$\mathbf{s}_i = \mathbf{E}\tau(X_i) + \mathbf{r}_i + \mathbf{e}_i.$$

- ▶ Instead of modeling the joint density of $(X_1, \dots, X_d)$, we model the probability that each X_i falls in a set $\mathcal{C}_k$, using:

$$P(X_1 \in \mathcal{C}_{k_1}, \dots, X_d \in \mathcal{C}_{k_d}) = \prod_{i=1}^d P(X_i \in \mathcal{C}_{k_i} \mid X_{1:i-1}).$$

²Ongoing work with **Miguel de Carvalho** and **Ronny Vallejos**.

³e.g. based on extreme quantile regression models.

HOW DO WE APPLY GPT TO GENERATE SEQUENCES OF EXTREME EVENTS?

- ▶ Each conditional $P(X_i \in \mathcal{C}_k \mid X_{i-L:i-1})$ is approximated by a **decoder Transformer**⁴ TF_θ , which maps the input embeddings $(\mathbf{s}_{i-L}, \dots, \mathbf{s}_{i-1})$:

$$\tilde{\mathbf{S}}_i = \text{TF}_\theta(\mathbf{S}_i), \quad \text{with } \mathbf{S}_i = (\mathbf{s}_{i-L}, \dots, \mathbf{s}_{i-1})^\top,$$

- ▶ This representation is then passed through a **learnable linear projection** $\mathbf{W} \in \mathbb{R}^{K \times D}$ to obtain the logits $\mathbf{z}_i = (z_{i1}, \dots, z_{iK}) \in \mathbb{R}^K$:

$$\mathbf{z}_i = \mathbf{W}\tilde{\mathbf{S}}_i.$$

- ▶ Finally, the **softmax layer** converts the logits into probabilities:

$$P_\theta(X_i \in \mathcal{C}_k \mid X_{i-L:i-1}) = \frac{\exp(z_{ik})}{\sum_{j=1}^K \exp(z_{ij})}, \quad k = 1, \dots, K.$$

- ▶ The model is **trained by minimizing the empirical negative log-likelihood** over a large dataset $\mathcal{D} = \{(X_1^{(t)}, \dots, X_d^{(t)})\}_{t=1}^N$:

$$\hat{\ell}(\theta) = -\frac{1}{Nd} \sum_{t=1}^N \sum_{i=1}^d \log P_\theta(X_i^{(t)} \in \mathcal{C}_{k_i^{(t)}} \mid X_{i-L:i-1}^{(t)}).$$

⁴Just like ChatGPT!

HOW DO WE APPLY GPT TO GENERATE SEQUENCES OF EXTREME EVENTS?

- ▶ Each conditional $P(X_i \in \mathcal{C}_k \mid X_{i-L:i-1})$ is approximated by a **decoder Transformer**⁴ TF_θ , which maps the input embeddings $(\mathbf{s}_{i-L}, \dots, \mathbf{s}_{i-1})$:

$$\tilde{\mathbf{S}}_i = \text{TF}_\theta(\mathbf{S}_i), \quad \text{with } \mathbf{S}_i = (\mathbf{s}_{i-L}, \dots, \mathbf{s}_{i-1})^\top,$$

- ▶ This representation is then passed through a **learnable linear projection** $\mathbf{W} \in \mathbb{R}^{K \times D}$ to obtain the logits $\mathbf{z}_i = (z_{i1}, \dots, z_{iK}) \in \mathbb{R}^K$:

$$\mathbf{z}_i = \mathbf{W}\tilde{\mathbf{S}}_i.$$

- ▶ Finally, the **softmax layer** converts the logits into probabilities:

$$P_\theta(X_i \in \mathcal{C}_k \mid X_{i-L:i-1}) = \frac{\exp(z_{ik})}{\sum_{j=1}^K \exp(z_{ij})}, \quad k = 1, \dots, K.$$

- ▶ The model is **trained by minimizing the empirical negative log-likelihood** over a large dataset $\mathcal{D} = \{(X_1^{(t)}, \dots, X_d^{(t)})\}_{t=1}^N$:

$$\hat{\ell}(\theta) = -\frac{1}{Nd} \sum_{t=1}^N \sum_{i=1}^d \log P_\theta(X_i^{(t)} \in \mathcal{C}_{k_i^{(t)}} \mid X_{i-L:i-1}^{(t)}).$$

⁴Just like ChatGPT!

HOW DO WE APPLY GPT TO GENERATE SEQUENCES OF EXTREME EVENTS?

- ▶ Each conditional $P(X_i \in \mathcal{C}_k \mid X_{i-L:i-1})$ is approximated by a **decoder Transformer**⁴ TF_θ , which maps the input embeddings $(\mathbf{s}_{i-L}, \dots, \mathbf{s}_{i-1})$:

$$\tilde{\mathbf{S}}_i = \text{TF}_\theta(\mathbf{S}_i), \quad \text{with } \mathbf{S}_i = (\mathbf{s}_{i-L}, \dots, \mathbf{s}_{i-1})^\top,$$

- ▶ This representation is then passed through a **learnable linear projection** $\mathbf{W} \in \mathbb{R}^{K \times D}$ to obtain the logits $\mathbf{z}_i = (z_{i1}, \dots, z_{iK}) \in \mathbb{R}^K$:

$$\mathbf{z}_i = \mathbf{W}\tilde{\mathbf{S}}_i.$$

- ▶ Finally, the **softmax layer** converts the logits into probabilities:

$$P_\theta(X_i \in \mathcal{C}_k \mid X_{i-L:i-1}) = \frac{\exp(z_{ik})}{\sum_{j=1}^K \exp(z_{ij})}, \quad k = 1, \dots, K.$$

- ▶ The model is **trained by minimizing the empirical negative log-likelihood** over a large dataset $\mathcal{D} = \{(X_1^{(t)}, \dots, X_d^{(t)})\}_{t=1}^N$:

$$\hat{\ell}(\theta) = -\frac{1}{Nd} \sum_{t=1}^N \sum_{i=1}^d \log P_\theta(X_i^{(t)} \in \mathcal{C}_{k_i^{(t)}} \mid X_{i-L:i-1}^{(t)}).$$

⁴Just like ChatGPT!

HOW DO WE APPLY GPT TO GENERATE SEQUENCES OF EXTREME EVENTS?

- ▶ Each conditional $P(X_i \in \mathcal{C}_k \mid X_{i-L:i-1})$ is approximated by a **decoder Transformer**⁴ TF_θ , which maps the input embeddings $(\mathbf{s}_{i-L}, \dots, \mathbf{s}_{i-1})$:

$$\tilde{\mathbf{S}}_i = \text{TF}_\theta(\mathbf{S}_i), \quad \text{with } \mathbf{S}_i = (\mathbf{s}_{i-L}, \dots, \mathbf{s}_{i-1})^\top,$$

- ▶ This representation is then passed through a **learnable linear projection** $\mathbf{W} \in \mathbb{R}^{K \times D}$ to obtain the logits $\mathbf{z}_i = (z_{i1}, \dots, z_{iK}) \in \mathbb{R}^K$:

$$\mathbf{z}_i = \mathbf{W}\tilde{\mathbf{S}}_i.$$

- ▶ Finally, the **softmax layer** converts the logits into probabilities:

$$P_\theta(X_i \in \mathcal{C}_k \mid X_{i-L:i-1}) = \frac{\exp(z_{ik})}{\sum_{j=1}^K \exp(z_{ij})}, \quad k = 1, \dots, K.$$

- ▶ The model is **trained by minimizing the empirical negative log-likelihood** over a **large dataset** $\mathcal{D} = \{(X_1^{(t)}, \dots, X_d^{(t)})\}_{t=1}^N$:

$$\hat{\ell}(\theta) = -\frac{1}{Nd} \sum_{t=1}^N \sum_{i=1}^d \log P_\theta(X_i^{(t)} \in \mathcal{C}_{k_i^{(t)}} \mid X_{i-L:i-1}^{(t)}).$$

⁴Just like ChatGPT!

REFERENCES

- 1) Vaswani et al. (2017). *Attention is All You Need*. arXiv:1706.03762.
- 2) Ba et al. (2016). *Layer Normalization*. arXiv:1607.06450.
- 3) Garg et al. (2022). *What Can Transformers Learn In-Context? A Case Study of Simple Function Classes*. arXiv:2208.01066v3.
- 4) Akyürek et al. (2023). *What Learning Algorithm Is In-Context Learning? Investigations with Linear Models*. arXiv:2301.05207.
- 5) Bai et al. (2023). *Transformers as Statisticians: Provable In-Context Learning with In-Context Algorithm Selection*. arXiv:2306.11644
- 6) Bishop & Bishop (2024). *Deep Learning: Foundations and Concepts*. Springer.

Exploring Transformer-Based Models for Cascading Extremes

Clemente Ferrer

`clemente.ferrer@usm.cl`

Department of Mathematics
Universidad Técnica Federico Santa María

Joint work with **Miguel de Carvalho & Ronny Vallejos**

WORKSHOP ON GENERATIVE AI FOR EXTREME EVENTS

Edinburgh, UK
June 13, 2025