

How Smooth Is Attention?

Valérie Castin¹, Pierre Ablin² and Gabriel Peyré¹

¹Ecole Normale Supérieure PSL, Paris

²Apple

ENS de Lyon, 21 January 2025

Transformers process tuples

Transformers are state-of-the-art in several fields: NLP, computer vision, multimodal learning, generative modeling...

Transformers process tuples

Transformers are state-of-the-art in several fields: NLP, computer vision, multimodal learning, generative modeling...

- Important feature: data represented as **tuples**
 $(x_1, \dots, x_n) \in (\mathbb{R}^d)^n$ of tokens

Transformers process tuples

Transformers are state-of-the-art in several fields: NLP, computer vision, multimodal learning, generative modeling...

- Important feature: data represented as **tuples**

$(x_1, \dots, x_n) \in (\mathbb{R}^d)^n$ of tokens

This is how GPT-3 tokenizes this sentence.

Tokenization of text



Tokenization of images

Transformers process tuples

Transformers are state-of-the-art in several fields: NLP, computer vision, multimodal learning, generative modeling...

- Important feature: data represented as **tuples**

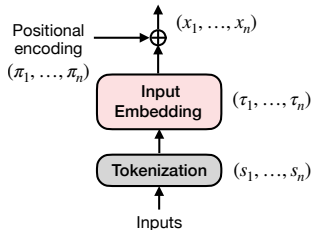
$(x_1, \dots, x_n) \in (\mathbb{R}^d)^n$ of tokens

This is how GPT-3 tokenizes this sentence.

Tokenization of text



Tokenization of images



- Positional encoding encodes the order of tokens

Transformers process tuples

Transformers are state-of-the-art in several fields: NLP, computer vision, multimodal learning, generative modeling...

- Important feature: data represented as **tuples**

$(x_1, \dots, x_n) \in (\mathbb{R}^d)^n$ of tokens

This is how GPT-3 tokenizes this sentence.

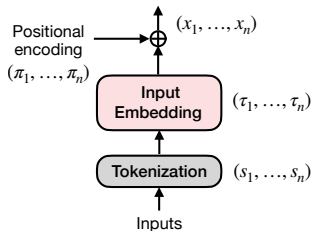
Tokenization of text



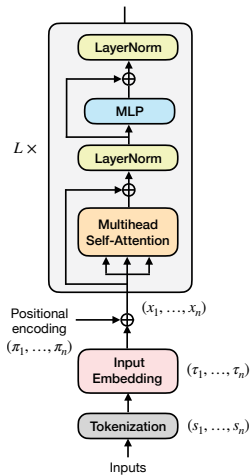
Tokenization of images

- Positional encoding encodes the order of tokens

Allows to learn local dependencies!



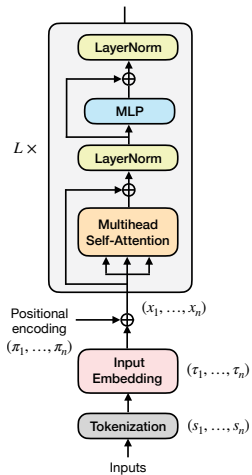
Attention is the core component of Transformers



Architecture of Transformers:

- Attention layers: $\cup_n(\mathbb{R}^d)^n \rightarrow \cup_n(\mathbb{R}^d)^n$ learn dependencies between tokens
- Multilayer perceptrons: $\mathbb{R}^d \rightarrow \mathbb{R}^d$ (applied token-wise)
- Layer normalization: project each token on an ellipsis

Attention is the core component of Transformers

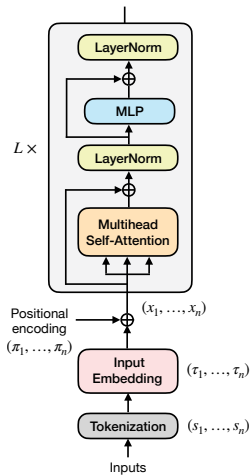


Architecture of Transformers:

- Attention layers: $\cup_n(\mathbb{R}^d)^n \rightarrow \cup_n(\mathbb{R}^d)^n$ learn dependencies between tokens
- Multilayer perceptrons: $\mathbb{R}^d \rightarrow \mathbb{R}^d$ (applied token-wise)
- Layer normalization: project each token on an ellipsis

How much can the output change when slightly perturbing the input?

Attention is the core component of Transformers



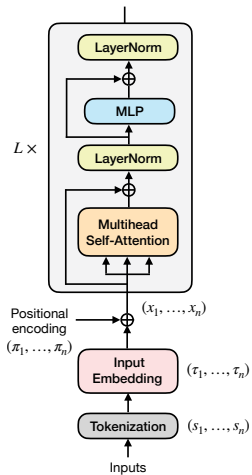
Architecture of Transformers:

- Attention layers: $\cup_n(\mathbb{R}^d)^n \rightarrow \cup_n(\mathbb{R}^d)^n$ learn dependencies between tokens
- Multilayer perceptrons: $\mathbb{R}^d \rightarrow \mathbb{R}^d$ (applied token-wise)
- Layer normalization: project each token on an ellipsis

How much can the output change when slightly perturbing the input?

- controls robustness and expressive power

Attention is the core component of Transformers



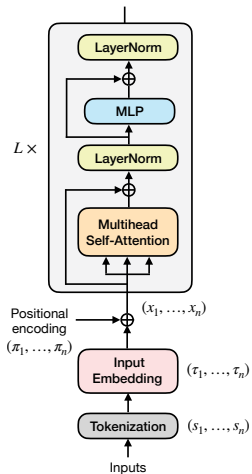
Architecture of Transformers:

- Attention layers: $\cup_n(\mathbb{R}^d)^n \rightarrow \cup_n(\mathbb{R}^d)^n$ learn dependencies between tokens
- Multilayer perceptrons: $\mathbb{R}^d \rightarrow \mathbb{R}^d$ (applied token-wise)
- Layer normalization: project each token on an ellipsis

How much can the output change when slightly perturbing the input?

- controls robustness and expressive power
- parameters are fixed

Attention is the core component of Transformers



Architecture of Transformers:

- Attention layers: $\cup_n(\mathbb{R}^d)^n \rightarrow \cup_n(\mathbb{R}^d)^n$ learn dependencies between tokens
- Multilayer perceptrons: $\mathbb{R}^d \rightarrow \mathbb{R}^d$ (applied token-wise)
- Layer normalization: project each token on an ellipsis

How much can the output change when slightly perturbing the input?

- controls robustness and expressive power
- parameters are fixed
- we analyze only one attention layer

The self-attention map

- Parameters: $Q, K \in \mathbb{R}^{k \times d}$ and $V \in \mathbb{R}^{d \times d}$

The self-attention map

- **Parameters:** $Q, K \in \mathbb{R}^{k \times d}$ and $V \in \mathbb{R}^{d \times d}$
- Attention of $X = (x_1, \dots, x_n) \in (\mathbb{R}^d)^n$ w.r.t. $z \in \mathbb{R}^d$:

$$\Gamma_X(z) := \sum_{j=1}^n p_j V x_j \quad \text{with} \quad p_j := \exp(\langle Qz, Kx_j \rangle) / \sum_{\ell=1}^n \exp(\langle Qz, Kx_\ell \rangle)$$

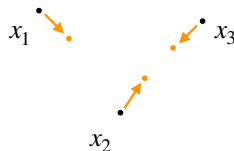
The self-attention map

- **Parameters:** $Q, K \in \mathbb{R}^{k \times d}$ and $V \in \mathbb{R}^{d \times d}$
- Attention of $X = (x_1, \dots, x_n) \in (\mathbb{R}^d)^n$ w.r.t. $z \in \mathbb{R}^d$:

$$\Gamma_X(z) := \sum_{j=1}^n p_j V x_j \quad \text{with} \quad p_j := \exp(\langle Qz, Kx_j \rangle) / \sum_{\ell=1}^n \exp(\langle Qz, Kx_\ell \rangle)$$

- **Self-attention** of $X = (x_1, \dots, x_n)$:

$$f: X \in (\mathbb{R}^d)^n \mapsto (\Gamma_X(x_1), \dots, \Gamma_X(x_n)) \in (\mathbb{R}^d)^n$$



The self-attention map

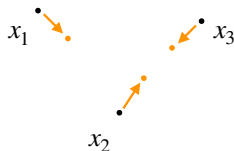
- Parameters: $Q, K \in \mathbb{R}^{k \times d}$ and $V \in \mathbb{R}^{d \times d}$
- Attention of $X = (x_1, \dots, x_n) \in (\mathbb{R}^d)^n$ w.r.t. $z \in \mathbb{R}^d$:

$$\Gamma_X(z) := \sum_{j=1}^n p_j V x_j \quad \text{with} \quad p_j := \exp(\langle Qz, Kx_j \rangle) / \sum_{\ell=1}^n \exp(\langle Qz, Kx_\ell \rangle)$$

- Self-attention of $X = (x_1, \dots, x_n)$:

$$f: X \in (\mathbb{R}^d)^n \mapsto (\Gamma_X(x_1), \dots, \Gamma_X(x_n)) \in (\mathbb{R}^d)^n$$

- Depends only on $A := K^\top Q$



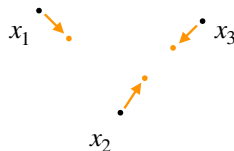
The self-attention map

- Parameters: $Q, K \in \mathbb{R}^{k \times d}$ and $V \in \mathbb{R}^{d \times d}$
- Attention of $X = (x_1, \dots, x_n) \in (\mathbb{R}^d)^n$ w.r.t. $z \in \mathbb{R}^d$:

$$\Gamma_X(z) := \sum_{j=1}^n p_j V x_j \quad \text{with} \quad p_j := \exp(\langle Qz, Kx_j \rangle) / \sum_{\ell=1}^n \exp(\langle Qz, Kx_\ell \rangle)$$

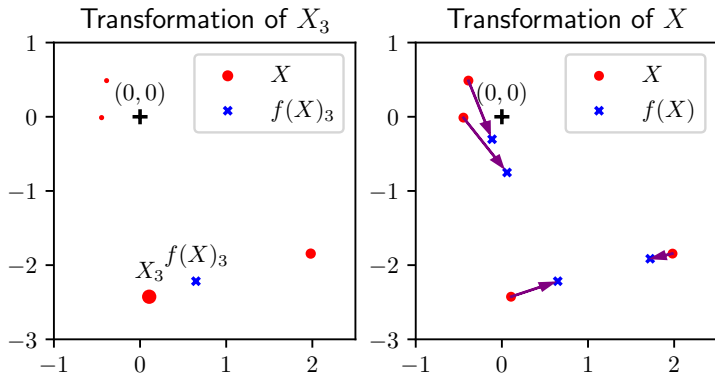
- Self-attention of $X = (x_1, \dots, x_n)$:

$$f: X \in (\mathbb{R}^d)^n \mapsto (\Gamma_X(x_1), \dots, \Gamma_X(x_n)) \in (\mathbb{R}^d)^n$$



- Depends only on $A := K^\top Q$
- Multi-head attention: linear combination $\sum_{h=1}^H W^{(h)} f^{(h)}$

The self-attention map



$$f(X)_i := \sum_{j=1}^n P_{ij} V_{x_j} \quad \text{with} \quad P_{ij} := \exp(\langle Q_{x_i}, K_{x_j} \rangle) / \sum_{\ell=1}^n \exp(\langle Q_{x_i}, K_{x_\ell} \rangle)$$

Measuring regularity with the local Lipschitz constant

$f: (\mathbb{R}^d)^n \rightarrow (\mathbb{R}^d)^n$ self-attention

Local Lipschitz constant

Norm on $(\mathbb{R}^d)^n$: $\|X\|^2 := \sum_{i=1}^n |x_i|^2$

Measuring regularity with the local Lipschitz constant

$f: (\mathbb{R}^d)^n \rightarrow (\mathbb{R}^d)^n$ self-attention

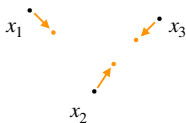
Local Lipschitz constant

Norm on $(\mathbb{R}^d)^n$: $\|X\|^2 := \sum_{i=1}^n |x_i|^2$

Local Lipschitz constant of f at X :

$$\text{Lip}_X(f) := \|D_X f\|_2 = \sup_{\|\varepsilon\|=1} \|D_X f(\varepsilon)\|$$

where $D_X f: (\mathbb{R}^d)^n \rightarrow (\mathbb{R}^d)^n$ Jacobian of f



Measuring regularity with the local Lipschitz constant

$f: (\mathbb{R}^d)^n \rightarrow (\mathbb{R}^d)^n$ self-attention

Local Lipschitz constant

Norm on $(\mathbb{R}^d)^n$: $\|X\|^2 := \sum_{i=1}^n |x_i|^2$

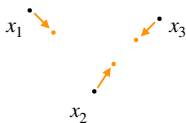
Local Lipschitz constant of f at X :

$$\text{Lip}_X(f) := \|D_X f\|_2 = \sup_{\|\varepsilon\|=1} \|D_X f(\varepsilon)\|$$

where $D_X f: (\mathbb{R}^d)^n \rightarrow (\mathbb{R}^d)^n$ Jacobian of f

- Gives global guarantees:

$$\sup_{X \neq Y \in B_R^n} \frac{\|f(X) - f(Y)\|}{\|X - Y\|} = \sup_{X \in B_R^n} \text{Lip}_X(f)$$



Measuring regularity with the local Lipschitz constant

$f: (\mathbb{R}^d)^n \rightarrow (\mathbb{R}^d)^n$ self-attention

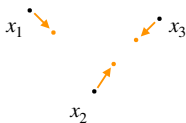
Local Lipschitz constant

Norm on $(\mathbb{R}^d)^n$: $\|X\|^2 := \sum_{i=1}^n |x_i|^2$

Local Lipschitz constant of f at X :

$$\text{Lip}_X(f) := \|D_X f\|_2 = \sup_{\|\varepsilon\|=1} \|D_X f(\varepsilon)\|$$

where $D_X f: (\mathbb{R}^d)^n \rightarrow (\mathbb{R}^d)^n$ Jacobian of f



- Gives global guarantees:

$$\sup_{X \neq Y \in B_R^n} \frac{\|f(X) - f(Y)\|}{\|X - Y\|} = \sup_{X \in B_R^n} \text{Lip}_X(f)$$

- Direction of maximal change:
 $U = (u_1, \dots, u_n)$ first singular vector of $D_X f$

$$\frac{\|f(X + \eta U) - f(X)\|}{\|U\|} \xrightarrow{\eta \rightarrow 0} \text{Lip}_X(f)$$

Presentation of the problem

- Self-attention is not globally Lipschitz continuous [Kim et al., 2021]

$$\text{Lip}(f|_{B_R^n}) \geq c(\textcolor{red}{A}, \textcolor{red}{V})R^2$$

→ set one particle at zero + spread the others

Presentation of the problem

- Self-attention is not globally Lipschitz continuous [Kim et al., 2021]

$$\text{Lip}(f|_{B_R^n}) \geq c(\textcolor{red}{A}, \textcolor{red}{V})R^2$$

→ set one particle at zero + spread the others

- Bound independent of n [Geshkovski et al., 2024]

$$\text{Lip}(f|_{B_R^n}) \leq \|\textcolor{red}{V}\|_2 (1 + 3 \|\textcolor{red}{A}\|_2 R^2) e^{2\|\textcolor{red}{A}\|_2 R^2}$$

Presentation of the problem

- Self-attention is not globally Lipschitz continuous [Kim et al., 2021]

$$\text{Lip}(f|_{B_R^n}) \geq c(\mathbf{A}, \mathbf{V}) R^2$$

→ set one particle at zero + spread the others

- Bound independent of n [Geshkovski et al., 2024]

$$\text{Lip}(f|_{B_R^n}) \leq \|\mathbf{V}\|_2 (1 + 3 \|\mathbf{A}\|_2 R^2) e^{2\|\mathbf{A}\|_2 R^2}$$

Questions

- Big discrepancy! Which bound is tighter? Dependency on n ?
- Characterize adversarial configurations?
- Local Lipschitz constant of real data?

I - Dependency on n of the Lipschitz constant of self-attention

Theoretical dependency on the sequence length n

Theorem 1 [Castin et al., 2024]

$$\text{Lip}(f|_{B_R^n}) \leq \sqrt{3} \|\mathbf{V}\|_2 \left(\|\mathbf{A}\|_2^2 R^4 (4n + 1) + n \right)^{1/2} \approx R^2 \sqrt{n}$$

Theoretical dependency on the sequence length n

Theorem 1 [Castin et al., 2024]

$$\text{Lip}(f|_{B_R^n}) \leq \sqrt{3} \|\mathbf{V}\|_2 \left(\|\mathbf{A}\|_2^2 R^4 (4n + 1) + n \right)^{1/2} \approx R^2 \sqrt{n}$$

and if $V = I_d$,

$$\text{Lip}(f|_{B_R^n}) \geq \frac{1}{1 + (n - 1)e^{-2R^2\gamma}} \sqrt{n - 1}$$

where $R^2\gamma \approx 10^{2-3}$ in practical Transformers.

Theoretical dependency on the sequence length n

Theorem 1 [Castin et al., 2024]

$$\text{Lip}(f|_{B_R^n}) \leq \sqrt{3} \|\mathbf{V}\|_2 \left(\|\mathbf{A}\|_2^2 R^4 (4n + 1) + n \right)^{1/2} \approx R^2 \sqrt{n}$$

and if $V = I_d$,

$$\text{Lip}(f|_{B_R^n}) \geq \frac{1}{1 + (n - 1)e^{-2R^2\gamma}} \sqrt{n - 1}$$

where $R^2\gamma \approx 10^{2-3}$ in practical Transformers.

- R fixed by layer normalization

Theoretical dependency on the sequence length n

Theorem 1 [Castin et al., 2024]

$$\text{Lip}(f|_{B_R^n}) \leq \sqrt{3} \|\mathbf{V}\|_2 \left(\|\mathbf{A}\|_2^2 R^4 (4n + 1) + n \right)^{1/2} \approx R^2 \sqrt{n}$$

and if $V = I_d$,

$$\text{Lip}(f|_{B_R^n}) \geq \frac{1}{1 + (n - 1)e^{-2R^2\gamma}} \sqrt{n - 1}$$

where $R^2\gamma \approx 10^{2-3}$ in practical Transformers.

- R fixed by layer normalization
- n not too large: $\text{Lip}(f|_{B_R^n})$ grows like $C\sqrt{n}$

Theoretical dependency on the sequence length n

Theorem 1 [Castin et al., 2024]

$$\text{Lip}(f|_{B_R^n}) \leq \sqrt{3} \|\mathbf{V}\|_2 \left(\|\mathbf{A}\|_2^2 R^4 (4n + 1) + n \right)^{1/2} \approx R^2 \sqrt{n}$$

and if $V = I_d$,

$$\text{Lip}(f|_{B_R^n}) \geq \frac{1}{1 + (n - 1)e^{-2R^2\gamma}} \sqrt{n - 1}$$

where $R^2\gamma \approx 10^{2-3}$ in practical Transformers.

- R fixed by layer normalization
- n not too large: $\text{Lip}(f|_{B_R^n})$ grows like $C\sqrt{n}$
- n exponentially large: analyzed separately (mean-field regime)

Theoretical dependency on the sequence length n

Theorem 1 [Castin et al., 2024]

$$\text{Lip}(f|_{B_R^n}) \leq \sqrt{3} \|\mathbf{V}\|_2 \left(\|\mathbf{A}\|_2^2 R^4 (4n + 1) + n \right)^{1/2} \approx R^2 \sqrt{n}$$

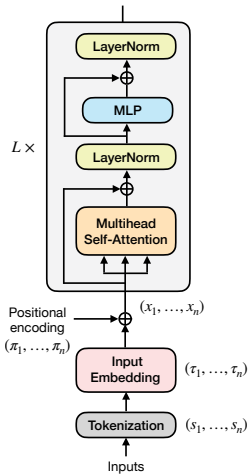
and if $V = I_d$,

$$\text{Lip}(f|_{B_R^n}) \geq \frac{1}{1 + (n - 1)e^{-2R^2\gamma}} \sqrt{n - 1}$$

where $R^2\gamma \approx 10^{2-3}$ in practical Transformers.

- Configuration for lower bound: $(Ru, -Ru, \dots, -Ru)$ or $(Ru, Ru/2, \dots, Ru/2)$ with u eigenvector of $\mathbf{A}$
- Similar bound for multi-head

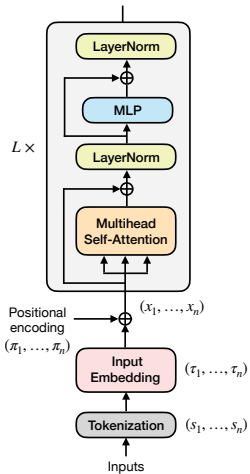
Experiments: typical case and worst case



Plot dependency of $\text{Lip}_X(f)$ in number of tokens n :

- for X obtained from real text (Alice in Wonderland, AG_NEWS)

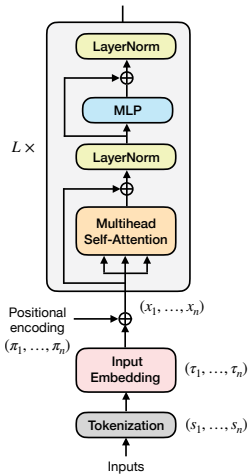
Experiments: typical case and worst case



Plot dependency of $\text{Lip}_X(f)$ in number of tokens n :

- for X obtained from real text (Alice in Wonderland, AG_NEWS)
- for adversarial data of the lower bound

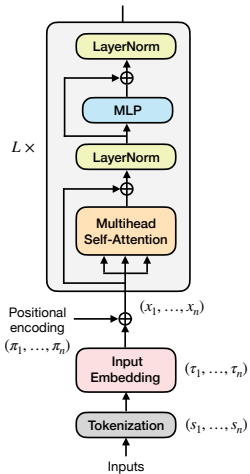
Experiments: typical case and worst case



Plot dependency of $\text{Lip}_X(f)$ in number of tokens n :

- for X obtained from real text (Alice in Wonderland, AG_NEWS)
- for adversarial data of the lower bound
- $\|D_X f\|_2$ computed with power method ($d = 768$)

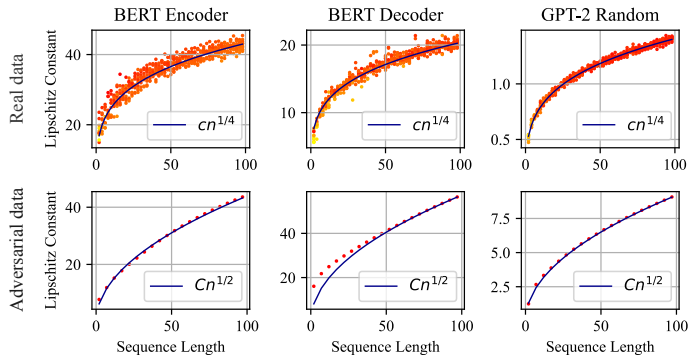
Experiments: typical case and worst case



Plot dependency of $\text{Lip}_X(f)$ in number of tokens n :

- for X obtained from real text (Alice in Wonderland, AG_NEWS)
- for adversarial data of the lower bound
- $\|D_X f\|_2$ computed with power method ($d = 768$)
- Same theory for masked self-attention

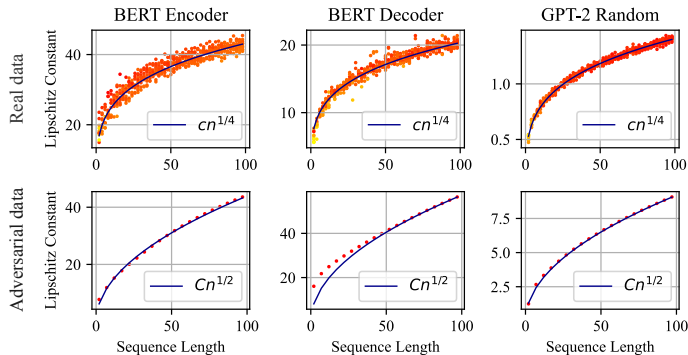
Experiments: typical case and worst case



- Growth in $Cn^{1/4}$ for real data

Local Lipschitz constant of real vs. adversarial data

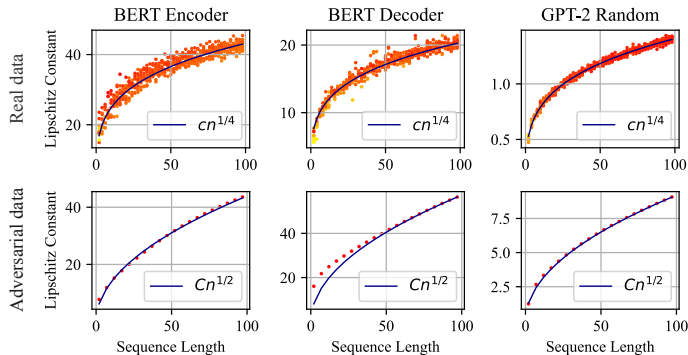
Experiments: typical case and worst case



- Growth in $Cn^{1/4}$ for real data
- Growth in $C\sqrt{n}$ for adv. data $\rightarrow$ matches lower bound

Local Lipschitz constant of real vs. adversarial data

Experiments: typical case and worst case

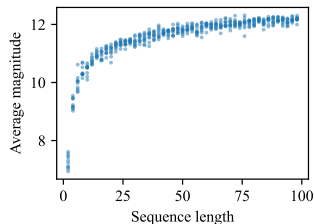


- Growth in $Cn^{1/4}$ for real data
- Growth in $C\sqrt{n}$ for adv. data $\rightarrow$ matches lower bound

Obstacle to Lipschitz attention (GPT-4o context window: 128k)

Local Lipschitz constant of real vs. adversarial data

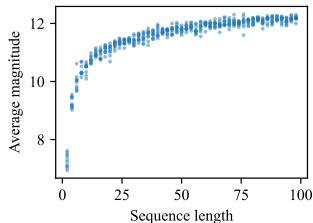
The special case of learnt positional encoding



GPT-2 pretrained: magnitude of tokens R grows with number of tokens n

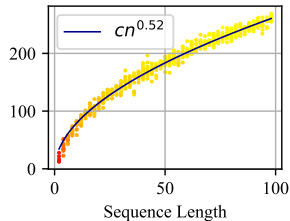
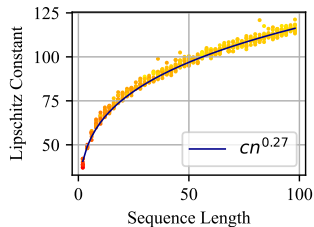
$$\text{Lip}(f|_{B_R^n}) \leq R^2 \sqrt{n}$$

The special case of learnt positional encoding



GPT-2 pretrained: magnitude of tokens R grows with number of tokens n

$$\text{Lip}(f|_{B_R^n}) \leq R^2 \sqrt{n}$$



Gaining intuition with the large radius regime

Large radius regime

$R \rightarrow +\infty$ while n fixed

Denote $m_i = \operatorname{argmax}_{1 \leq j \leq n} \langle \mathbf{A} \mathbf{x}_i, \mathbf{x}_j \rangle$ (a.s. unique)

$$(D_{RX}f)(\varepsilon) \rightarrow_{R \rightarrow +\infty} (\mathbf{V}_{\varepsilon_{m_1}}, \dots, \mathbf{V}_{\varepsilon_{m_n}}) \quad \varepsilon \in (\mathbb{R}^d)^n$$

Proposition [Castin et al., 2024]

In the large radius regime:

$$\operatorname{Lip}(f|_{B_R^n}) \lesssim_{R \rightarrow +\infty} \|\mathbf{V}\|_2 \sqrt{n}$$

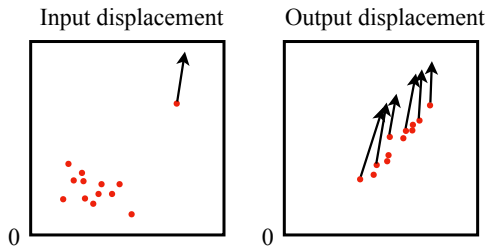
reached if $m_1 = \dots = m_n$

Adversarial configurations in the large radius regime

One token x_j far away from the others s.t.

$$\forall i = 1, \dots, n, \quad \langle Ax_i, x_j \rangle \approx \max_k \langle Ax_i, x_k \rangle$$

→ local Lipschitz constant **proportional to** $\sqrt{n}$



II - The mean-field regime (n exponentially large)

Generalizing self-attention to probability measures

Self-attention $f(X) := (\Gamma_X(x_1), \dots, \Gamma_X(x_n))$ with

$$\Gamma_X: x \in \mathbb{R}^d \mapsto \sum_{j=1}^n p_j \textcolor{red}{V} x_j \quad \text{with} \quad p_j := \exp(\langle \textcolor{red}{A} x, x_j \rangle) / \sum_{\ell=1}^n \exp(\langle \textcolor{red}{A} x, x_\ell \rangle)$$

Generalizing self-attention to probability measures

Self-attention $f(X) := (\Gamma_X(x_1), \dots, \Gamma_X(x_n))$ with

$$\Gamma_X: x \in \mathbb{R}^d \mapsto \sum_{j=1}^n p_j \textcolor{red}{V} x_j \quad \text{with} \quad p_j := \exp(\langle \textcolor{red}{A} x, x_j \rangle) / \sum_{\ell=1}^n \exp(\langle \textcolor{red}{A} x, x_\ell \rangle)$$

Permutation equivariant: $x_i \leftrightarrow x_j \Rightarrow f(X)_i \leftrightarrow f(X)_j$

Generalizing self-attention to probability measures

Self-attention $f(X) := (\Gamma_X(x_1), \dots, \Gamma_X(x_n))$ with

$$\Gamma_X: x \in \mathbb{R}^d \mapsto \sum_{j=1}^n p_j \textcolor{red}{V} x_j \quad \text{with} \quad p_j := \exp(\langle \textcolor{red}{A} x, x_j \rangle) / \sum_{\ell=1}^n \exp(\langle \textcolor{red}{A} x, x_\ell \rangle)$$

Permutation equivariant: $x_i \leftrightarrow x_j \Rightarrow f(X)_i \leftrightarrow f(X)_j$

Mean-field self-attention [Sander et al., 2022]

$F: \mu \in \mathcal{P}(\mathbb{R}^d) \mapsto (\Gamma_\mu)_\# \mu$ with

$$\Gamma_\mu: x \in \mathbb{R}^d \mapsto \int k(x, y) \textcolor{red}{V} y d\mu(y) \quad \text{with} \quad k(x, y) := e^{\langle \textcolor{red}{A} x, y \rangle} / \int e^{\langle \textcolor{red}{A} x, z \rangle} d\mu(z)$$

Generalizing self-attention to probability measures

Self-attention $f(X) := (\Gamma_X(x_1), \dots, \Gamma_X(x_n))$ with

$$\Gamma_X: x \in \mathbb{R}^d \mapsto \sum_{j=1}^n p_j \mathbf{V} x_j \quad \text{with} \quad p_j := \exp(\langle \mathbf{A}x, x_j \rangle) / \sum_{\ell=1}^n \exp(\langle \mathbf{A}x, x_\ell \rangle)$$

Permutation equivariant: $x_i \leftrightarrow x_j \Rightarrow f(X)_i \leftrightarrow f(X)_j$

Mean-field self-attention [Sander et al., 2022]

$F: \mu \in \mathcal{P}(\mathbb{R}^d) \mapsto (\Gamma_\mu)_\# \mu$ with

$$\Gamma_\mu: x \in \mathbb{R}^d \mapsto \int k(x, y) \mathbf{V} y d\mu(y) \quad \text{with} \quad k(x, y) := e^{\langle \mathbf{A}x, y \rangle} / \int e^{\langle \mathbf{A}x, z \rangle} d\mu(z)$$

Covers discrete case representing X as $\frac{1}{n} \sum_{i=1}^n \delta_{x_i}$

Generalizing self-attention to probability measures

Mean-field self-attention [Sander et al., 2022]

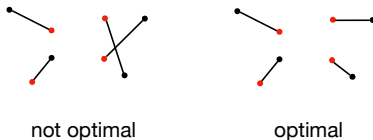
$F: \mu \in \mathcal{P}_c(\mathbb{R}^d) \mapsto (\Gamma_\mu)_\# \mu$ with

$$\Gamma_\mu: x \in \mathbb{R}^d \mapsto \int k(x, y) \mathbf{V} y d\mu(y) \quad \text{with} \quad k(x, y) := e^{\langle \mathbf{A}x, y \rangle} / \int e^{\langle \mathbf{A}x, z \rangle} d\mu(z)$$

Lipschitz constant measured with **Wasserstein distance**

$$W_2(\mu, \nu) := \left(\inf_{\pi} \int |x - y|^2 d\pi(x, y) \right)^{1/2}$$

where $\int d\pi(x, \cdot) = d\mu(x)$ and $\int d\pi(\cdot, y) = d\nu(y)$



Generalizing self-attention to probability measures

Mean-field self-attention [Sander et al., 2022]

$F: \mu \in \mathcal{P}_c(\mathbb{R}^d) \mapsto (\Gamma_\mu)_\# \mu$ with

$$\Gamma_\mu: x \in \mathbb{R}^d \mapsto \int k(x, y) \mathbf{V} y d\mu(y) \quad \text{with} \quad k(x, y) := e^{\langle \mathbf{A}x, y \rangle} / \int e^{\langle \mathbf{A}x, z \rangle} d\mu(z)$$

Lipschitz constant measured with **Wasserstein distance**

$$W_2(\mu, \nu) := \left(\inf_{\pi} \int |x - y|^2 d\pi(x, y) \right)^{1/2}$$

where $\int d\pi(x, \cdot) = d\mu(x)$ and $\int d\pi(\cdot, y) = d\nu(y)$



not optimal

optimal

$$\text{Lip}(F|_{\mathcal{P}(B_R)}) := \sup_{\mu \neq \nu \in \mathcal{P}(B_R)} \frac{W_2(F(\mu), F(\nu))}{W_2(\mu, \nu)}$$

Generalizing self-attention to probability measures

Mean-field self-attention [Sander et al., 2022]

$F: \mu \in \mathcal{P}_c(\mathbb{R}^d) \mapsto (\Gamma_\mu)_\# \mu$ with

$$\Gamma_\mu: x \in \mathbb{R}^d \mapsto \int k(x, y) \mathbf{V} y d\mu(y) \quad \text{with} \quad k(x, y) := e^{\langle \mathbf{A}x, y \rangle} / \int e^{\langle \mathbf{A}x, z \rangle} d\mu(z)$$

Lipschitz constant measured with **Wasserstein distance**

$$W_2(\mu, \nu) := \left(\inf_{\pi} \int |x - y|^2 d\pi(x, y) \right)^{1/2}$$

$\rightarrow$ upp. bound $R^2 e^{R^2}$
[Geshkovski et al., 2024]

where $\int d\pi(x, \cdot) = d\mu(x)$ and $\int d\pi(\cdot, y) = d\nu(y)$

$$\text{Lip}(F|_{\mathcal{P}(B_R)}) := \sup_{\mu \neq \nu \in \mathcal{P}(B_R)} \frac{W_2(F(\mu), F(\nu))}{W_2(\mu, \nu)}$$

Lower bound in the mean-field regime

Upper bound [Geshkovski et al., 2024]

$$\text{Lip}(f|_{B_R^n}) \leq \|V\|_2 (1 + 3 \|A\|_2 R^2) e^{2\|A\|_2 R^2}$$

Lower bound in the mean-field regime

Upper bound [Geshkovski et al., 2024]

$$\text{Lip}(f_{|B_R^n}) \leq \|V\|_2 (1 + 3 \|A\|_2 R^2) e^{2\|A\|_2 R^2}$$

Proposition 2 [Castin et al., 2024]

If $V = I_d$ and $n \sim e^{2\gamma R^2}$:

$$\text{Lip}(f_{|B_R^n}) \gtrsim \frac{\gamma}{2} R^2 e^{\gamma R^2}$$

Lower bound in the mean-field regime

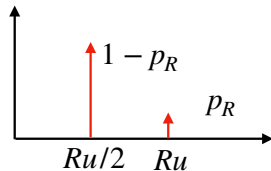
Upper bound [Geshkovski et al., 2024]

$$\text{Lip}(f|_{B_R^n}) \leq \|\mathbf{V}\|_2 (1 + 3 \|\mathbf{A}\|_2 R^2) e^{2\|\mathbf{A}\|_2 R^2}$$

Proposition 2 [Castin et al., 2024]

If $\mathbf{V} = I_d$ and $n \sim e^{2\gamma R^2}$:

$$\text{Lip}(f|_{B_R^n}) \gtrsim \frac{\gamma}{2} R^2 e^{\gamma R^2}$$



Lower bound in the mean-field regime

Upper bound [Geshkovski et al., 2024]

$$\text{Lip}(f_{|B_R^n}) \leq \|V\|_2 (1 + 3 \|A\|_2 R^2) e^{2\|A\|_2 R^2}$$

Proposition 2 [Castin et al., 2024]

If $V = I_d$ and $n \sim e^{2\gamma R^2}$:

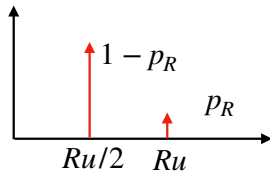
$$\text{Lip}(f_{|B_R^n}) \gtrsim \frac{\gamma}{2} R^2 e^{\gamma R^2}$$

Probability measures for lower bound:

$$p_R \delta_{Ru} + (1 - p_R) \delta_{Ru/2} \quad \text{or} \quad p_R \delta_{Ru} + (1 - p_R) \delta_{-Ru}$$

with

$$p_R = e^{-2\gamma R^2} \rightarrow_{R \rightarrow +\infty} 0$$



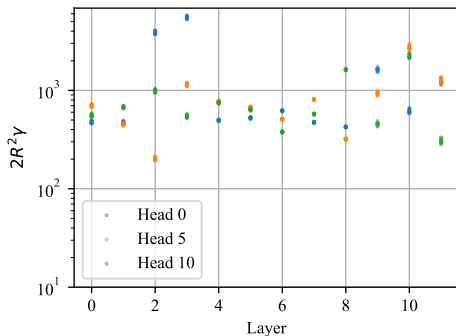
The mean-field regime is not practically relevant

Proposition 2 [Castin et al., 2024]

If $V = I_d$ and $n \sim e^{2\gamma R^2}$:

$$\text{Lip}(f|_{B_R^n}) \gtrsim \frac{\gamma}{2} R^2 e^{\gamma R^2}$$

In practice $2\gamma R^2 \approx 10^3 \rightarrow$ not realistic



III - Masked self-attention

Masked self-attention processes tokens sequentially

Masked self-attention $f^m: (\mathbb{R}^d)^n \rightarrow (\mathbb{R}^d)^n$ such that

$$f^m(X)_i := f(x_1, \dots, x_i)_i$$

with params $A, V \in \mathbb{R}^{d \times d}$

Theorem 2 [Castin et al., 2024]

The upper and lower bounds of Theorem 1 on $\text{Lip}(f|_{B_R^n})$ also hold for $\text{Lip}(f^m|_{B_R^n})$

Mean-field regime? $\rightarrow$ not permutation equivariant!

Generalizing masked self-attention to measures

Mean-field self-attention $F: \mu \in \mathcal{P}_c(\mathbb{R}^d) \mapsto (\Gamma_\mu)_\# \mu$ where

$$\Gamma_\mu: x \in \mathbb{R}^d \mapsto \int_{\mathbb{R}^d} k(x, y) \textcolor{red}{V} y d\mu(y) \quad \text{with} \quad k(x, y) := e^{\langle \textcolor{red}{A}x, y \rangle} / \int e^{\langle \textcolor{red}{A}x, z \rangle} d\mu(z)$$

Generalizing masked self-attention to measures

Mean-field self-attention $F: \mu \in \mathcal{P}_c(\mathbb{R}^d) \mapsto (\Gamma_\mu)_\# \mu$ where

$$\Gamma_\mu: x \in \mathbb{R}^d \mapsto \int_{\mathbb{R}^d} k(x, y) \textcolor{red}{V} y d\mu(y) \quad \text{with} \quad k(x, y) := e^{\langle \textcolor{red}{A}x, y \rangle} / \int e^{\langle \textcolor{red}{A}x, z \rangle} d\mu(z)$$

Mean-field masked self-attention [Castin et al., 2024]

Replace $\mu \in \mathcal{P}_c(\mathbb{R}^d)$ by $\bar{\mu} \in \mathcal{P}_c([0, 1] \times \mathbb{R}^d)$:

$$F^m: \bar{\mu} \mapsto (\Gamma_{\bar{\mu}})_\# \bar{\mu} \quad \text{where} \quad \Gamma_{\bar{\mu}}(s, x) := \left(s, \int_{[0, 1] \times \mathbb{R}^d} \textcolor{red}{V} y k_s(x, y) d\bar{\mu}(\tau, y) \right)$$

with

$$k_s(x, y) := e^{\langle \textcolor{red}{A}x, y \rangle} \textcolor{blue}{1}_{\tau \leq s} / \int_{[0, 1] \times \mathbb{R}^d} e^{\langle \textcolor{red}{A}x, y \rangle} \textcolor{blue}{1}_{\tau \leq s} d\bar{\mu}(\tau, y)$$

Generalizing masked self-attention to measures

Mean-field self-attention $F: \mu \in \mathcal{P}_c(\mathbb{R}^d) \mapsto (\Gamma_\mu)_\# \mu$ where

$$\Gamma_\mu: x \in \mathbb{R}^d \mapsto \int_{\mathbb{R}^d} k(x, y) \textcolor{red}{V} y d\mu(y) \quad \text{with} \quad k(x, y) := e^{\langle \textcolor{red}{A}x, y \rangle} / \int e^{\langle \textcolor{red}{A}x, z \rangle} d\mu(z)$$

Mean-field masked self-attention [Castin et al., 2024]

Replace $\mu \in \mathcal{P}_c(\mathbb{R}^d)$ by $\bar{\mu} \in \mathcal{P}_c([0, 1] \times \mathbb{R}^d)$:

$$F^m: \bar{\mu} \mapsto (\Gamma_{\bar{\mu}})_\# \bar{\mu} \quad \text{where} \quad \Gamma_{\bar{\mu}}(s, x) := \left(s, \int_{[0, 1] \times \mathbb{R}^d} \textcolor{red}{V} y k_s(x, y) d\bar{\mu}(\tau, y) \right)$$

with

$$k_s(x, y) := e^{\langle \textcolor{red}{A}x, y \rangle} \textcolor{blue}{1}_{\tau \leq s} / \int_{[0, 1] \times \mathbb{R}^d} e^{\langle \textcolor{red}{A}x, y \rangle} \textcolor{blue}{1}_{\tau \leq s} d\bar{\mu}(\tau, y)$$

Same upper bound as unmasked mean-field self-attention!

IV - Application of the mean-field framework: modeling Transformers as PDEs

Modeling an infinitely deep Transformer as a PDE

- Simplified Transformer with only attention layers:

$$f = f^L \circ \dots \circ f^1 \quad \text{with} \quad f^\ell(X) := X + \frac{1}{L}(\Gamma_X^\ell(x_1), \dots, \Gamma_X^\ell(x_n))$$

Modeling an infinitely deep Transformer as a PDE

- Simplified Transformer with only attention layers:

$$f = f^L \circ \dots \circ f^1 \quad \text{with} \quad f^\ell(X) := X + \frac{1}{L}(\Gamma_X^\ell(x_1), \dots, \Gamma_X^\ell(x_n))$$

- Discretizes

$$\dot{x}_i(t) = \Gamma_{\mu(t)}(x_i(t)), \quad 1 \leq i \leq n$$

with $\mu(t) := \frac{1}{n} \sum_{i=1}^n \delta_{x_i(t)}$ and $\Gamma_{\mu(t)}: \mathbb{R}^d \rightarrow \mathbb{R}^d$ velocity field

Modeling an infinitely deep Transformer as a PDE

- Simplified Transformer with only attention layers:

$$f = f^L \circ \dots \circ f^1 \quad \text{with} \quad f^\ell(X) := X + \frac{1}{L}(\Gamma_X^\ell(x_1), \dots, \Gamma_X^\ell(x_n))$$

- Discretizes

$$\dot{x}_i(t) = \Gamma_{\mu(t)}(x_i(t)), \quad 1 \leq i \leq n$$

with $\mu(t) := \frac{1}{n} \sum_{i=1}^n \delta_{x_i(t)}$ and $\Gamma_{\mu(t)}: \mathbb{R}^d \rightarrow \mathbb{R}^d$ velocity field

Corresponding PDE:

$$\partial_t \mu + \operatorname{div}(\mu \Gamma_\mu) = 0$$

also on continuous measures

Tokens cluster after renormalization

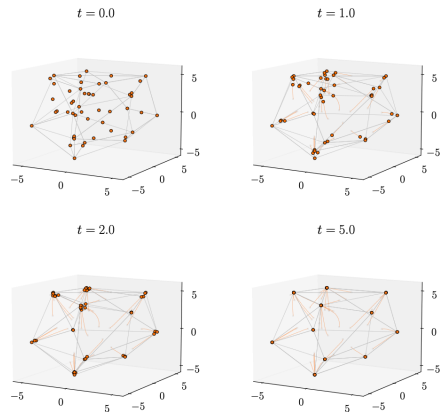
For discrete tokens $x_1, \dots, x_n$ following

$$\dot{x}_i(t) = \Gamma_{X(t)}(x_i(t))$$

and renormalizing

$$y_i(t) := e^{-tV} x_i(t)$$

→ **clusters emerge** [Geshkovski et al., 2024]



Tokens cluster after renormalization

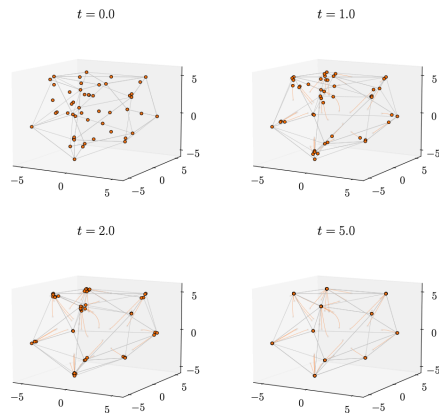
For discrete tokens $x_1, \dots, x_n$ following

$$\dot{x}_i(t) = \Gamma_{X(t)}(x_i(t))$$

and renormalizing

$$y_i(t) := e^{-t\mathbf{V}} x_i(t)$$

→ **clusters emerge** [Geshkovski et al., 2024]



Similar phenomenon for Gaussian inputs! (Castin, Carrillo, Peyré, Ablin, in preparation)

References



Castin, V., Ablin, P., and Peyré, G. (2024).

How smooth is attention?

In *ICML 2024*.



Geshkovski, B., Letrouit, C., Polyanskiy, Y., and Rigollet, P. (2024).

The emergence of clusters in self-attention dynamics.

Advances in Neural Information Processing Systems, 36.



Kim, H., Papamakarios, G., and Mnih, A. (2021).

The lipschitz constant of self-attention.

In *International Conference on Machine Learning*, pages 5562–5571. PMLR.



Sander, M. E., Ablin, P., Blondel, M., and Peyré, G. (2022).

Sinkformers: Transformers with doubly stochastic attention.

In *International Conference on Artificial Intelligence and Statistics*, pages 3515–3530. PMLR.