

# Transformers for Information Measures Estimation in Time Series Domain

RESEARCH PROPOSAL

OMER LUXEMBOURG · ADVISOR: PROF. HAIM PERMUTER ·  
AUGUST 2025

# Why Transfer Entropy (TE)?

---

TE measures directed information flow between two jointly stationary processes, for a fixed memory length  $k, l$ :

$$\text{TE}_{X \rightarrow Y}(k, l) := I(X^l; Y_l | Y_{l-k}^{l-1}).$$

Captures causal influence beyond correlation or Mutual Information (MI)

Useful in neuroscience, communications, IoT, finance

# Challenges in TE Estimation

---



Long contexts, continuous signals, non-Gaussian noise



Classical methods (KDE/kNN/copula) suffer from bias/variance trade-off



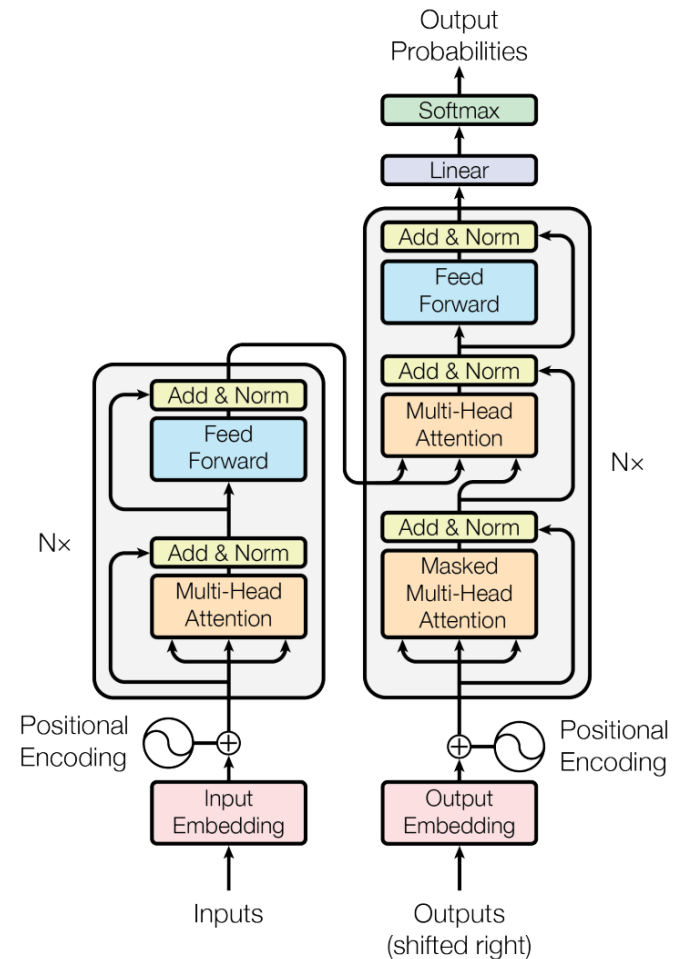
Generic neural MI estimators not tailored for finite-order TE

# Why Transformers?

Attention captures long-range dependencies efficiently

Causal masking aligns with TE's finite history requirement

Scales better than RNNs for long contexts



# TE vs. DI Rate

---

➤ TE: finite-order conditional MI (CMI)

➤ Directed Information (DI) rate: infinite memory CMI

$$\begin{aligned} I((\mathbb{X} \rightarrow \mathbb{Y}) &:= \lim_{n \rightarrow \infty} \frac{1}{n} I(X^n \rightarrow Y^n) \\ &= \lim_{n \rightarrow \infty} \frac{1}{n} \sum_{i=1}^n I(X^i; Y_i | Y^{i-1}) \\ &= \lim_{n \rightarrow \infty} I(X^n; Y_n | Y^{n-1}) \end{aligned}$$

➤ TE converges to DI rate under stationarity

# What is TREET?

---

- Transformer-based estimator of TE using Donsker-Vardhan (DV) representation

**Theorem 2 (DV representation)** *For any,  $P, Q \in \mathcal{P}(\mathcal{X})$ , we have*

$$D_{\text{KL}}(P\|Q) = \sup_{f:\mathcal{X}\rightarrow\mathbb{R}} \mathbb{E}_P[f] - \log(\mathbb{E}_Q[e^f]),$$

*where the supremum is taken over all measurable functions  $f$  with finite expectations.*

- Shared-weights dual potentials with tailored attention.
- Consistent for order- $l$  TE.

# DV Objective for TE

---

- TE can be expressed as difference of two KL divergences

**Lemma 2 (TE as KL Divergences)** *TE decomposes as*

$$\text{TE}_{X \rightarrow Y}(l) = D_{Y_l | Y^{l-1} X^l \| \tilde{Y}_l} - D_{Y_l | Y^{l-1} \| \tilde{Y}_l},$$

where

$$D_{Y_l | Y^{l-1} \| \tilde{Y}_l} := D_{\text{KL}} \left( P_{Y_l | Y^{l-1}} \| \tilde{P}_Y \middle| P_{Y^{l-1}} \right),$$

$$D_{Y_l | Y^{l-1} X^l \| \tilde{Y}_l} := D_{\text{KL}} \left( P_{Y_l | Y^{l-1} X^l} \| \tilde{P}_Y \middle| P_{Y^{l-1} X^l} \right)$$

- Each term is **optimized via DV representation**, where the potential function is a transformer.
- Reference sampling (Gaussian / Uniform) enables stable estimation of KL divergence terms

# Consistency Guarantee

---

➤ Under stationarity and ergodicity, TREET converges to true TE

➤ Proven for fixed memory parameter  $l$ .

➤  $\tilde{Y}$  is i.i.d. under absolutely continuous reference measure over the alphabet  $\mathcal{Y}$ .

➤ Each term is a lower bound estimator of the DKL in Lemma 2, respectively.

$$\begin{aligned} \widehat{\text{TE}}_{X \rightarrow Y}(D_n; l) &= \sup_{g_{xy} \in \mathcal{G}_{\text{ctf}}^{XY}} \widehat{\text{D}}_{Y_l | Y^{l-1} X^l \| \tilde{Y}_l}(D_n, g_{xy}) \\ &\quad - \sup_{g_y \in \mathcal{G}_{\text{ctf}}^Y} \widehat{\text{D}}_{Y_l | Y^{l-1} \| \tilde{Y}_l}(D_n, g_y), \end{aligned}$$

where

$$\begin{aligned} \widehat{\text{D}}_{Y_l | Y^{l-1} \| \tilde{Y}_l}(D_n, g_y) &:= \frac{1}{n} \sum_{i=1}^n g_y(Y_i^{i+l}) \\ &\quad - \log \left( \frac{1}{n} \sum_{i=1}^n e^{g_y(\tilde{Y}_{i+l}, Y_i^{i+l-1})} \right), \\ \widehat{\text{D}}_{Y_l | Y^{l-1} X^l \| \tilde{Y}_l}(D_n, g_{xy}) &:= \frac{1}{n} \sum_{i=1}^n g_{xy}(Y_i^{i+l}, X_i^{i+l}) \\ &\quad - \log \left( \frac{1}{n} \sum_{i=1}^n e^{g_{xy}(\tilde{Y}_{i+l}, Y_i^{i+l-1}, X_i^{i+l})} \right). \end{aligned}$$

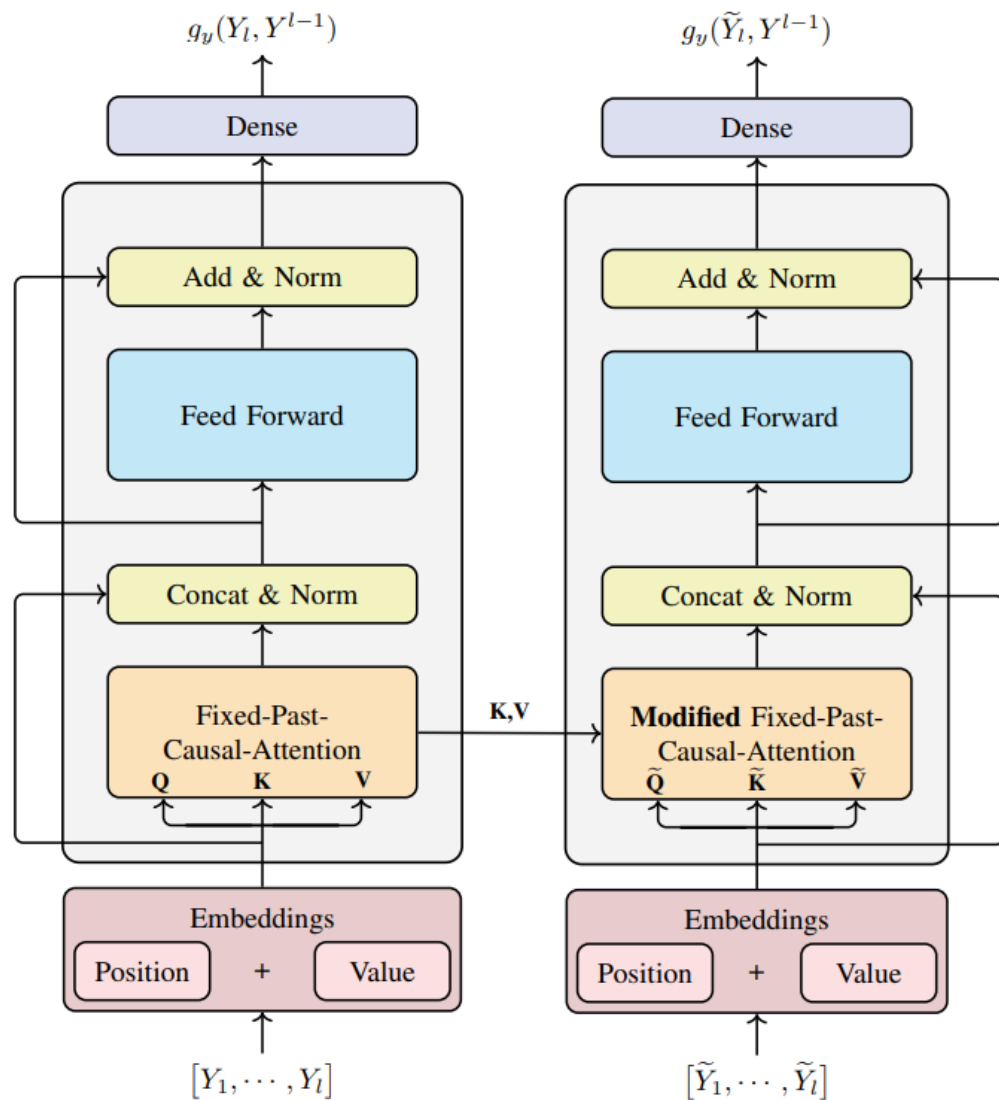


# TREET Architecture

Two passes through the same transformer (shared weights)

Fix-Past-Causal-Attention (FPCA) for main term; modified FPCA for reference term – reuses past keys and values to simulate the same conditionals

Stable training and interpretable lag-wise attention



# Training & Implementation

- Sampled order  $l$  sequences of  $\mathbb{Y}$  and joint  $\mathbb{X}, \mathbb{Y}$  process for each DKL term – unsupervised training for maximization of DV representation.
- Mini-batch optimization via gradient ascent. with reference sampling – tested on Uniform and Gaussian sampling
- FPCA and modified FPCA enforce causal structure to preserve no past beyond order  $l$  and future leakage. Parallel  $(L - l)$  outputs for sequence length  $L$ .

---

## Algorithm 1 TREET

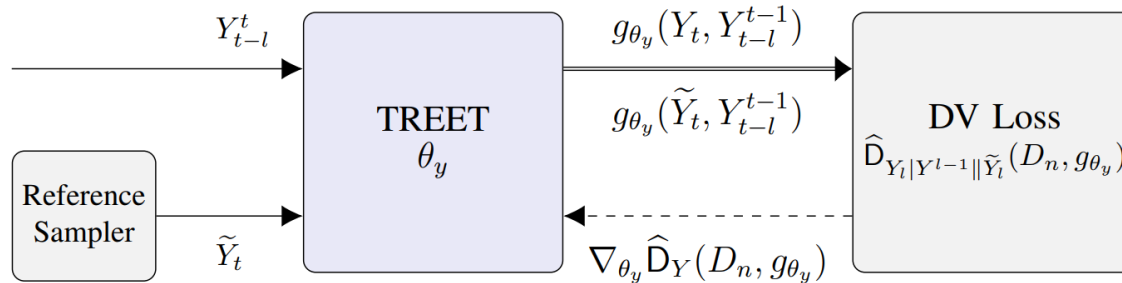
---

**Input:** Joint process samples  $D_n$ ; Observation length  $l \in \mathbb{N}$ .

**Output:**  $\widehat{\text{TE}}_{X \rightarrow Y}(D_n; l)$  - TE estimation.

---

- 1: NNs initialization  $g_{\theta_y}, g_{\theta_{xy}}$  with corresponding parameters  $\theta_y, \theta_{xy}$ .
  - 2: **Step 1 – Optimization:**
  - 3: **repeat**
  - 4: Draw a batch  $B_m$ :  $m < n$  sub-sequences, length  $L > l$  from  $D_n$ , with reference samples  $P_{\tilde{Y}}$  for each.
  - 5: Compute both potentials  $\widehat{D}_{Y_l|Y^{l-1}X^l\|\tilde{Y}_l}(B_m, g_{\theta_{xy}})$ ,  $\widehat{D}_{Y_l|Y^{l-1}\|\tilde{Y}_l}(B_m, g_{\theta_y})$  via (20).
  - 6: Update parameters:
  - 7:  $\theta_{xy} \leftarrow \theta_{xy} + \nabla_{\theta_{xy}} \widehat{D}_{Y_l|Y^{l-1}X^l\|\tilde{Y}_l}(B_m, g_{\theta_{xy}})$
  - 8:  $\theta_y \leftarrow \theta_y + \nabla_{\theta_y} \widehat{D}_{Y_l|Y^{l-1}\|\tilde{Y}_l}(B_m, g_{\theta_y})$
  - 9: **until** convergence criteria.
  - 10: **Step 2 – Evaluation:** Evaluate for a sub-sequence (20) and (19a) to obtain  $\widehat{\text{TE}}_{X \rightarrow Y}(D_n; l)$ .
- 



# TE Benchmark: Long-Memory Synthetic Data

---

➤ Compared TREET vs. TENE and Copnet – both ignoring time settings and using the sequence as a whole vector.

- **TENE** – DV based TE estimator,

$$\text{TE}_{X \rightarrow Y}(l) = I(X^l; Y_l | Y^{l-1}) = I(X^l; Y_l, Y^{l-1}) - I(X^l; Y^{l-1})$$

- **Copnet** – KNN based estimator – 4 coupla entropy estimators,

$$I(X) = -h_C(X);$$

$$\text{TE}_{X \rightarrow Y}(l) = h_C(Y_l, Y^{l-1}, X^l) + h_C(Y_l, Y^{l-1}) + h_C(Y^{l-1}, X^l) - h_C(Y^{l-1})$$

➤ TREET robust to long context lengths (up to  $l = 99$ )

For any order TE,  $l \geq 1$ , the value is constant since  $X$  is i.i.d. and  $Y$  is 1-order Markov.

[illegible]

# Channel Capacity Estimation

---

As stated before, TE converges to the DI rate, thus TREET can estimate the DI rate – which achieves the channel capacity (for optimized input distribution)

**Remark 1 (Channel capacity)** *Consider channels with and without feedback links from the channel output back to the encoder. The feedforward capacity of a channel sequence  $\{P_{Y^n||X^n}\}$ , for  $n \in \mathbb{N}$  is*

$$C_{\text{FF}} = \lim_{n \rightarrow \infty} \sup_{P_{X^n}} \frac{1}{n} I(X^n; Y^n),$$

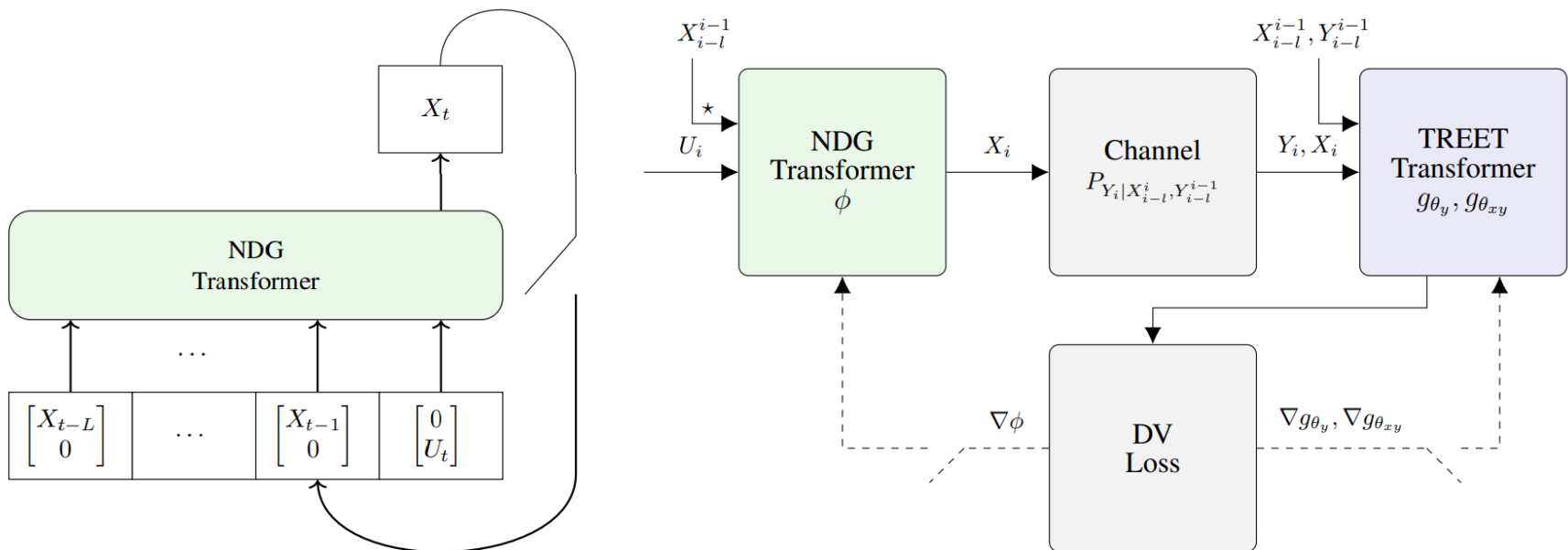
*and the feedback capacity is*

$$C_{\text{FB}} = \lim_{n \rightarrow \infty} \sup_{P_{X^n||Y^{n-1}}} \frac{1}{n} I(X^n \rightarrow Y^n). \quad (29)$$

*The achievability of the capacities is further discussed in [38], [39]. [34] showed that for non-feedback scenario, the optimization problem over  $P_{X^n||Y^n}$  can be translated to  $P_{X^n}$*

# Channel Capacity Estimation

Neural Density Generator (NDG) learns input distribution maximizing TE – thus, TE is optimized and estimated in an alternate procedure.



# Channel Capacity Estimation

---

➤ Applied to AWGN, AR(1), MA(1) channels with/without feedback.

➤ TREET matches theoretical capacities and the DI rate neural estimator - DINE

---

**Algorithm 2** Continuous TREET optimization

---

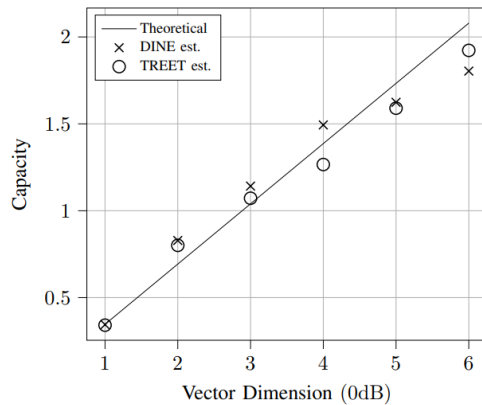
**Input:** Continuous sequence-to-sequence system  $\mathcal{S}$ ; Observation length  $l \in \mathbb{N}$ .

**Output:**  $\widehat{\text{TE}}_{X \rightarrow Y}^*(U^n; l)$ , optimized NDG.

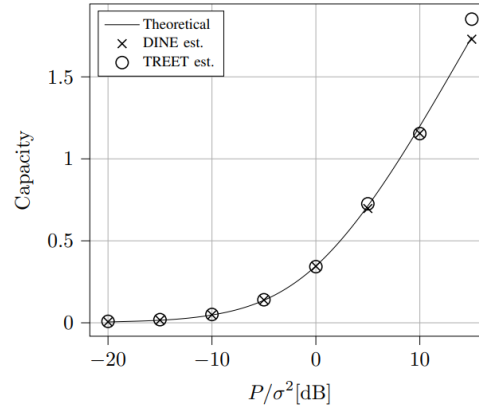
---

- 1: NNs initialization  $g_{\theta_y}$ ,  $g_{\theta_{xy}}$  and  $h_\phi$  with corresponding parameters  $\theta_y, \theta_{xy}, \phi$ .
  - 2: **repeat**
  - 3:   Draw noise  $U^m$ ,  $m < n$ .
  - 4:   Compute batch  $B_m^\phi$  sized  $m$  using NDG,  $\mathcal{S}$
  - 5:   **if** training TREET **then**
  - 6:     Perform TREET optimization - Step 1 in Algorithm 1.
  - 7:   **else** Train NDG
  - 8:     Compute  $\widehat{\text{TE}}_{X \rightarrow Y}(B_m^\phi, g_{\theta_y}, g_{\theta_{xy}}, h_\phi; l)$  using (19a).
  - 9:     Update NDG parameters:
  - 10:      $\phi \leftarrow \phi + \nabla_\phi \widehat{\text{TE}}_{X \rightarrow Y}(B_m^\phi, g_{\theta_y}, g_{\theta_{xy}}, h_\phi; l)$
  - 11: **until** convergence criteria.
  - 12: Draw  $U^m$  to produce  $l$  length sequence and evaluate  $\widehat{\text{TE}}_{X \rightarrow Y}(D_n^\phi; l)$ .
  - 13: **return**  $\widehat{\text{TE}}_{X \rightarrow Y}^*(U^n; l)$ , optimized NDG.
-

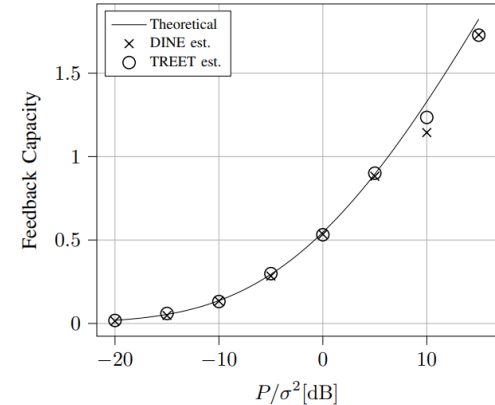
# Channel Capacity Estimation



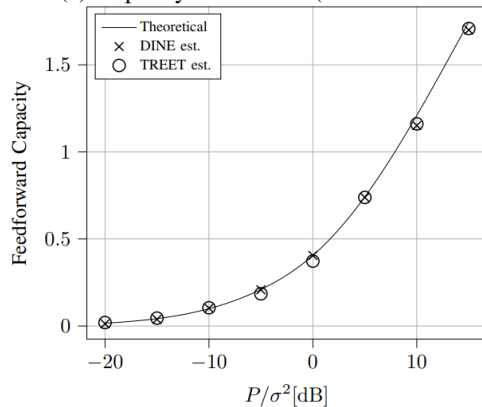
(a) Capacity of AWGN (variate dimension).



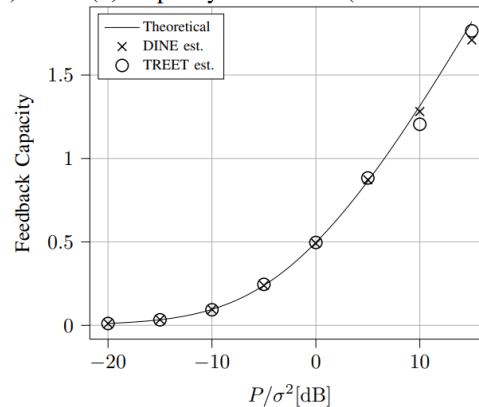
(b) Capacity of AWGN (variate SNR).



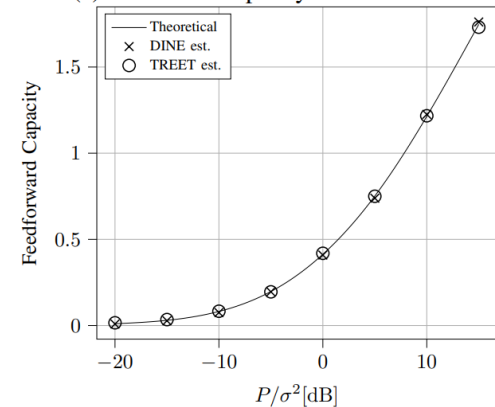
(c) Feedback Capacity of Gaussian MA(1).



(d) Feedforward Capacity of Gaussian MA(1).



(e) Feedback Capacity of Gaussian AR(1).



(f) Feedforward Capacity of Gaussian AR(1).



# Channel Capacity Estimation For Long Memory Analysis

Process capacity estimation – GMA(100),

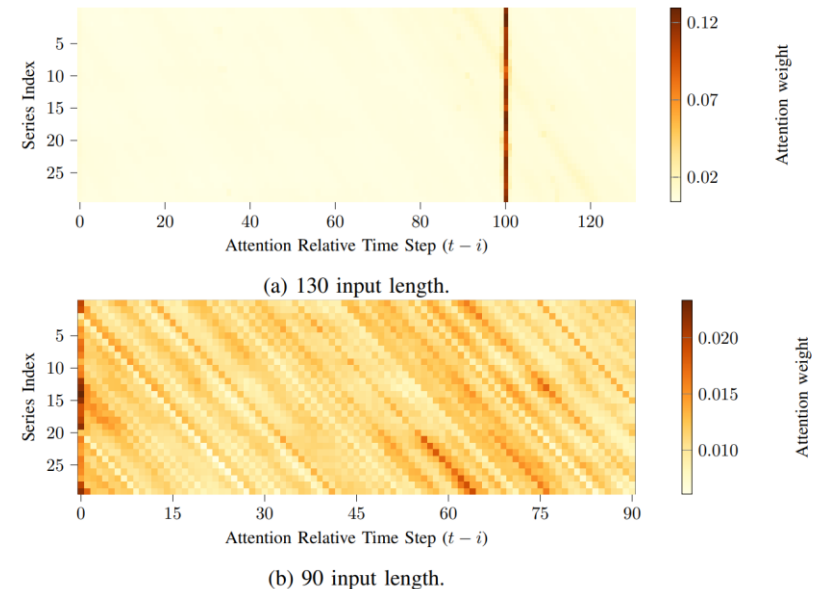
$$Z_t = N_t + \alpha N_{t-100},$$

$$Y_t = X_t + Z_t, \quad t \in \mathbb{Z}.$$

FPCA visualization of influential lags for shorter or larger memory than the process' order

| $l$ | TREET                    |                    | DINE                     |                    |
|-----|--------------------------|--------------------|--------------------------|--------------------|
|     | Estimated capacity [nat] | Absolute error (%) | Estimated capacity [nat] | Absolute error (%) |
| 60  | 0.19                     | 53                 | <b>0.30</b>              | <b>25</b>          |
| 70  | 0.29                     | 29                 | <b>0.35</b>              | <b>15</b>          |
| 80  | 0.28                     | 31                 | <b>0.34</b>              | <b>17</b>          |
| 90  | 0.29                     | 28                 | <b>0.30</b>              | <b>26</b>          |
| 100 | <b>0.37</b>              | <b>9</b>           | 0.36                     | 12                 |
| 110 | <b>0.38</b>              | <b>7</b>           | 0.33                     | 20                 |
| 120 | <b>0.36</b>              | <b>11</b>          | 0.33                     | 18                 |
| 130 | <b>0.36</b>              | <b>11</b>          | 0.34                     | 17                 |
| 140 | <b>0.35</b>              | <b>14</b>          | 0.26                     | 35                 |

Stable Estimation



True capacity = 0.405 [nat]

# Density Estimation via TREET

---

- Optimized potentials yield the following log-likelihood ratio (LLR) –

$$\begin{aligned} f_{y,l}^* &:= \log \left( \frac{dP_{Y^l}}{d(P_{X^l Y^{l-1}} \otimes \tilde{P}_{Y_l})} \right) & f_{y,l}^* &:= \log \left( \frac{dP_{Y^l}}{d(P_{Y^{l-1}} \otimes \tilde{P}_{Y_l})} \right) \\ &= \log p_{Y_l | X^l Y^{l-1}} - \log \tilde{p}_{Y_l} & &= \log p_{Y_l | Y^{l-1}} - \log \tilde{p}_{Y_l} \end{aligned}$$

- The conditional density estimator can be derived out-of-the-box –

$$P_{Y_t | Y^{t-1}}(y_t | y^{t-1}) \approx \exp \left( \hat{\mathbf{D}}_{Y_t | Y^{t-1} \| \tilde{Y}_t}(D_n, g_y) \right) \cdot \tilde{P}_{Y_t}(y_t)$$

- Entire distribution derived from normalized values from grid input of  $\{y_t\} \subset \mathcal{Y}$ .
- Competitive with MDN, KDE, Kalman on HMM-like processes and handles long delays and noise models robustly

Same can be derived for causally conditioned on X process

# Density Estimation via TREET

$$X_t = \alpha X_{t-1} + \beta X_{t-k} + W_t,$$

$$Y_t = \gamma X_t + V_t,$$

Where  $W_t, V_t$  are i.i.d noise,  $\mathcal{N}(0,1)$  or  $U[-1,1]$ .

| Model  | Gaussian     |              | Uniform      |              |
|--------|--------------|--------------|--------------|--------------|
|        | KLD          | TV           | KLD          | TV           |
| Kalman | 1.076        | 0.467        | 1.304        | 0.408        |
| KDE    | 1.135        | 0.608        | 0.847        | 0.417        |
| MDN    | 1.098        | 0.632        | 0.667        | 0.460        |
| DINE   | <b>0.795</b> | <u>0.525</u> | <b>0.475</b> | <b>0.291</b> |
| TREET  | <u>0.797</u> | <b>0.524</b> | <u>0.482</u> | <u>0.296</u> |

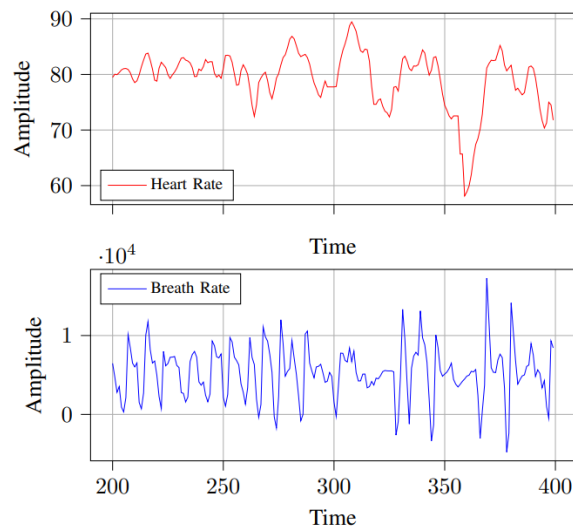
| $k$ | TREET        |              | DINE         |              | MDN   |       | Kalman       |              |
|-----|--------------|--------------|--------------|--------------|-------|-------|--------------|--------------|
|     | KL           | TV           | KL           | TV           | KL    | TV    | KL           | TV           |
| 2   | <u>0.933</u> | <u>0.569</u> | 0.934        | 0.570        | 1.460 | 0.708 | <b>0.636</b> | <b>0.405</b> |
| 5   | <u>1.036</u> | 0.601        | <b>0.875</b> | <b>0.556</b> | 1.454 | 0.707 | 1.284        | <u>0.599</u> |
| 10  | <b>0.984</b> | <b>0.585</b> | 1.299        | 0.662        | 1.450 | 0.706 | <u>1.239</u> | <u>0.603</u> |
| 15  | <b>0.992</b> | <b>0.587</b> | 1.441        | 0.704        | 1.448 | 0.707 | <u>1.232</u> | <u>0.602</u> |
| 25  | <b>0.993</b> | <b>0.587</b> | 1.451        | 0.706        | 1.443 | 0.704 | <u>1.196</u> | <u>0.598</u> |
| 50  | <b>0.993</b> | <b>0.589</b> | 1.452        | 0.707        | 1.453 | 0.707 | <u>1.193</u> | <u>0.597</u> |

1-order process -  $\beta = 0$

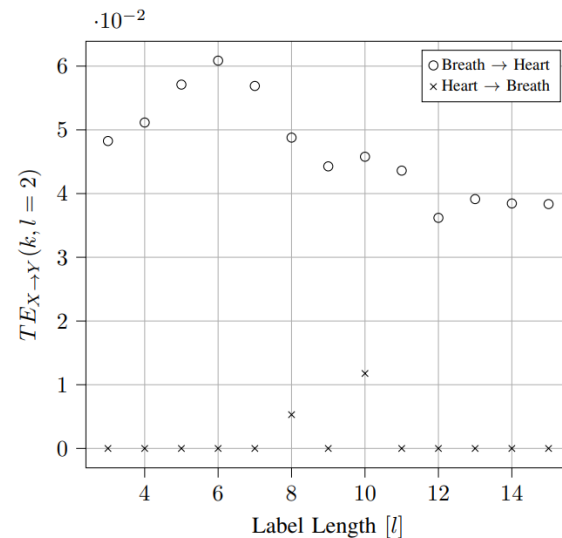
$k$ -order process -  $\beta \neq 0$  and  $k > 1$

# Apnea Case Study: Feature-Level Causal Analysis

- Analyzed **respiration**  $\rightleftarrows$  **heart-rate** influence in sleep apnea patients
- TE directionality aligns with clinical understanding – breathing affects heart rate
- Different orders of TE estimation reveals causal relationships



(a) Sampled sequence of Apnea dataset.



(b) TE estimation for variable length history of  $Y$ .

# Future Scope: TREET Time-Series Applications

---

## Cross-Domain Potential

Apply TREET across diverse fields such as:

**Communications** – analyzing feedback and memory in channels

**Physiology** – uncovering causal relationships in biomedical signals

**Finance** – detecting directional influence in market dynamics

## Advanced Applications of TREET

TREET enables powerful time-series analysis across several domains, including:

**Feature Selection:** Identifying causally relevant variables in complex datasets.

**Anomaly Detection:** Detecting irregularities in single or joint processes.

**Control & Decision Systems:** Enhancing decision-making through causal insights.

# Future Scope: Information Theory & Deep Learning Application

---

## 1. Diffusion Models for Improved KLD and MI Estimation:

- Current estimators struggle with high dimensional KLD and MI
- Diffusion models can break down each variational representation task (DV, MINE, NWJ, InfoNCE) to smaller ones

$$\mathsf{T}_{\text{sum}}^*(\mathbb{P}_0, \mathbb{P}_K) := \sum_{k=0}^{K-1} \mathsf{T}_k(\mathbb{P}_k \parallel \mathbb{P}_{k+1})$$

$$\begin{aligned} \mathsf{l}_{\text{VR,diffused}}(X; Y) &:= \frac{1}{N} \sum_{(X_i, Y_i)_{i=1}^N \sim \mathbb{P}_{XY}} \mathsf{T}_{\text{sum}}^*(X_i, Y_i) \\ &\stackrel{(a)}{=} \mathbb{E}_{\mathbb{P}_{XY}} \left[ \log \frac{\mathbb{P}_{XY}(X, Y)}{\mathbb{P}_X(X) \mathbb{P}_Y(Y)} \right] \\ &= \mathsf{l}(X; Y) \end{aligned}$$

# Future Scope: Information Theory & Deep Learning Application

---

## 1. Diffusion Models for Improved Mutual Information Estimation:

- Proved lower bound for estimating the KLD, from optimal set of functions  $\{T_k^\star\}_{k=0}^{K-1}$

**Bernoulli Diffusion Process:** Define  $\mathbb{P}_k$  as a probabilistic mixture between  $\mathbb{P}_0$  and  $\mathbb{P}_K$ , controlled by the parameter  $\alpha_k$ :

$$\mathbb{P}_k = (1 - \alpha_k)\mathbb{P}_0 + \alpha_k\mathbb{P}_K.$$

This process results in a stochastic switching between samples from  $\mathbb{P}_0$  and  $\mathbb{P}_K$ :

$$X_k = \begin{cases} X_0, & \text{with probability } 1 - \alpha_k, \\ X_K, & \text{with probability } \alpha_k. \end{cases}$$

**Lemma** ( $D_{\text{KL}}$  upper bound - Bernoulli)

$$D_{\text{KL}}(P_0 || P_K) \geq \sum_{k=0}^{K-1} D_{\text{KL}}(P_k || P_{k+1}).$$

# Future Scope: Information Theory & Deep Learning Application

---

## 2. Optimizing Inference of Discrete Diffusion Language Models Through MI:

- Diffusion models, discrete ones in particular, suffer from many function evaluations until resulting with a good quality generated content
- Dilated Unmasking Scheduler (DUS) - Assuming a Markov model, uncovering tokens at the same time can be done with minimal shared information between them, thus generation preserves its' quality

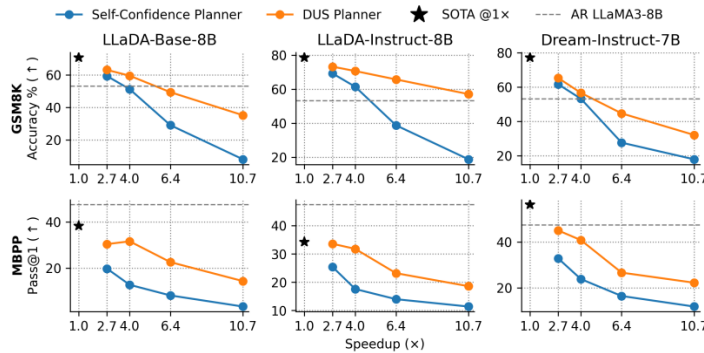
**Lemma 1** *Under a fast-mixing first-order Markov chain, let  $i_1 < \dots < i_k$  be the indices selected by DUS. Then*

$$H(X_{i_1}, \dots, X_{i_k} \mid \mathcal{S}_t) \approx \sum_{j=1}^k H(X_{i_j} \mid \mathcal{S}_t). \quad (8)$$



# Future Scope: Information Theory & Deep Learning Application

## 2. Optimizing Inference of Discrete Diffusion Language Models Through MI:



(a) Score vs. speedup - DUS and self-confidence MDLM planners.

Question: A car is on a road trip and drives 60 mph for 2 hours, and then 30 mph for 1 hour. What is the car's average speed in mph during this trip?  
 Answer: The car drives 2 \* 60 miles/hour = <<2\*60=60>>120 miles in the first part of the trip. The car drives 1 \* 30 miles/hour = <<1\*30=30>>30 miles in the second part of the trip. The total distance traveled is 120 + 30 miles = <<120+30=150>>150 miles. The total time taken is 2 + 1 hours = <<2+1=3>>3 hours. The average speed is 150 miles / 3 hours = <<150/3=50>>50 miles/hour. #### 50

(b) DUS planner

Question: A car is on a road trip and drives 60 mph for 2 hours, and then 30 mph for 1 hour. What is the car's average speed in mph during this trip?  
 Answer: The car went 60 30 = <<60+30=90>>90 miles. It took 2 1 = <<2+1=3>>3 hours. The average speed is 90 / 3 = <<90/3=30>>30 mph. #### 30

(c) Self-confidence planner

Generation Start  Generation End

| Model     |                     | B=8<br>×2.7 |              | B=16<br>×4 |              | B=32<br>×6.4 |              | B=64<br>×10.7 |              |
|-----------|---------------------|-------------|--------------|------------|--------------|--------------|--------------|---------------|--------------|
|           |                     | Conf.       | DUS          | Conf.      | DUS          | Conf.        | DUS          | Conf.         | DUS          |
| GSM8K     | LLaDA-Base          | 59.29       | <b>63.08</b> | 51.23      | <b>59.51</b> | 29.04        | <b>49.36</b> | 8.04          | <b>35.18</b> |
|           | LLaDA-Instruct      | 69.22       | <b>73.24</b> | 61.41      | <b>70.66</b> | 38.74        | <b>65.73</b> | 18.73         | <b>57.09</b> |
|           | Dream-Instruct      | 61.64       | <b>65.28</b> | 53.22      | <b>56.63</b> | 27.60        | <b>44.66</b> | 17.89         | <b>32.07</b> |
| MATH500   | LLaDA-Base          | 16.6        | <b>21.4</b>  | 11.2       | <b>19.2</b>  | 6.0          | <b>13.6</b>  | 2.6           | <b>10.2</b>  |
|           | LLaDA-Instruct      | 21.4        | <b>23.8</b>  | 15.4       | <b>22.8</b>  | 10.8         | <b>19.2</b>  | 8.0           | <b>14.8</b>  |
|           | Dream-Instruct      | 22.4        | <b>27.0</b>  | 15.4       | <b>19.8</b>  | 7.2          | <b>13.2</b>  | 4.0           | <b>11.6</b>  |
| Humaneval | LLaDA-Base          | 15.85       | <b>25.61</b> | 12.8       | <b>19.51</b> | 4.88         | <b>14.02</b> | 4.88          | <b>6.71</b>  |
|           | LLaDA-Instruct      | 21.95       | <b>28.05</b> | 14.02      | <b>23.17</b> | 9.76         | <b>10.37</b> | 10.98         | <b>11.59</b> |
|           | Dream-Instruct      | 8.54        | <b>14.63</b> | 5.49       | <b>11.59</b> | <b>6.71</b>  | <b>6.71</b>  | 6.10          | <b>9.15</b>  |
|           | DiffuCoder-Base     | 17.07       | <b>28.66</b> | 6.71       | <b>38.41</b> | 2.44         | <b>21.95</b> | 0.61          | <b>6.10</b>  |
|           | DiffuCoder-Instruct | 7.93        | <b>22.56</b> | 14.02      | <b>20.12</b> | <b>13.41</b> | 12.80        | 11.59         | <b>8.54</b>  |
| MBPP      | LLaDA-Base          | 19.8        | <b>30.4</b>  | 12.8       | <b>31.6</b>  | 8.2          | <b>22.6</b>  | 3.4           | <b>14.4</b>  |
|           | LLaDA-Instruct      | 25.4        | <b>33.6</b>  | 17.6       | <b>31.8</b>  | 14.0         | <b>23.2</b>  | 11.4          | <b>18.6</b>  |
|           | Dream-Instruct      | 32.8        | <b>45.0</b>  | 23.8       | <b>40.8</b>  | 16.4         | <b>26.6</b>  | 11.8          | <b>22.2</b>  |
|           | DiffuCoder-Base     | 29.2        | <b>48.6</b>  | 17.4       | <b>43.0</b>  | 10.2         | <b>27.4</b>  | 3.4           | <b>17.2</b>  |
|           | DiffuCoder-Instruct | 31.8        | <b>46.4</b>  | 25.6       | <b>43.6</b>  | 21.0         | <b>26.6</b>  | 13.0          | <b>18.2</b>  |

# Thank you!

- [TREET: TRansfer Entropy Estimation via Transformers, IEEE Access, 2025 vol. 13, pp. 126477-126495, 2025  
Omer Luxembourg, Dor Tsur, Haim Permuter](#)
- [Plan for Speed: Dilated Scheduling for Masked Diffusion Language Models, arXiv preprint arXiv:2506.19037, 2025  
Omer Luxembourg, Haim Permuter, Eliya Nachmani](#)