

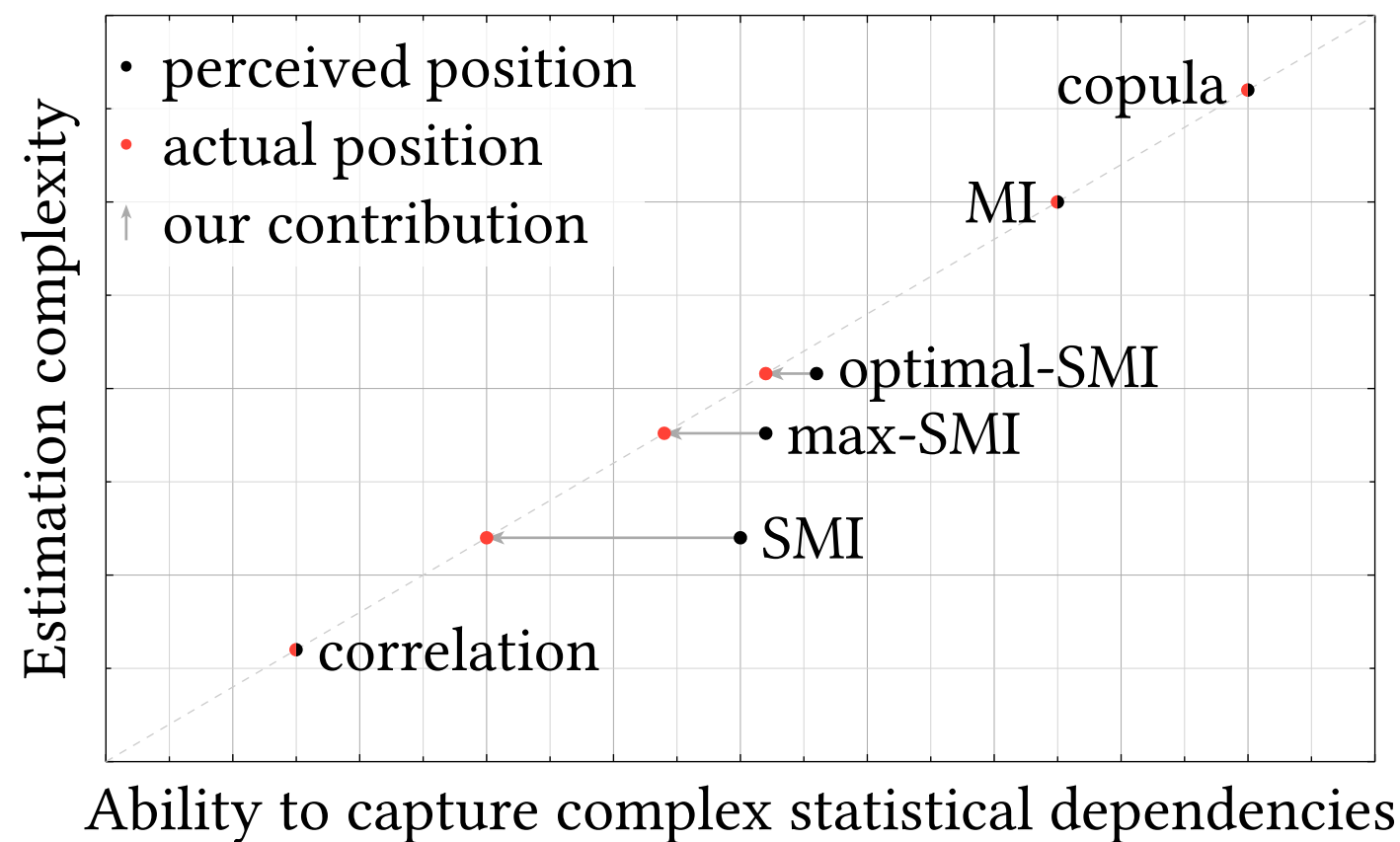
Curse of Slicing: Why Sliced Mutual Information is a Deceptive Measure of Statistical Dependence

Alexander Semenenko, Ivan Butakov, Alexey Frolov, Ivan Oseledets

Skoltech, MIPT, AIRI, Applied AI Institute, INM RAS

Skoltech
Wireless

Project Center for
Next Generation
Wireless and IoT



Introduction

Mutual Information (MI) is a universal measure of non-linear dependence between random X and Y .

Applications:

- Independence testing
- Feature selection
- Representation learning/disentanglement
- Analysis of DNNs

Problem: MI is hard to estimate.

Solution (?): use random projections — **Sliced Mutual Information (SMI)**.

Is this approach actually good?

Formal Definitions

Mutual Information:

$$I(X; Y) = \mathbb{E} \log \frac{d\mathbb{P}_{X,Y}}{d\mathbb{P}_X \otimes \mathbb{P}_Y} = D_{\text{KL}}(\mathbb{P}_{X,Y} \parallel \mathbb{P}_X \otimes \mathbb{P}_Y).$$

- Equals zero iff X and Y are independent
- Invariant under bijections

Conditional Mutual Information:

$$I(X; Y | Z) = \mathbb{E}_{z \sim Z} I(X; Y | Z = z)$$

Sliced Mutual Information:

$$SI_k(X; Y) = I(\Theta^T X; \Phi^T Y | \Theta, \Phi),$$

where Θ, Φ are random uniformly distributed projectors on k -dimensional subspaces.

Key Findings

We analyzed SMI across diverse settings and demonstrated

1. **SMI saturates prematurely and fails to detect significant increases in statistical dependence.**
2. **SMI prioritizes information redundancy over information content.**
3. **SMI is, in fact, cursed with the curse of dimensionality.**
4. **SMI follows the universal scalability-utility trade-off in statistical dependence measures.**

Theoretical analysis

Lemma. For $(X, Y) \sim \mathcal{N}\left(0, \begin{pmatrix} I & \rho I \\ \rho I & I \end{pmatrix}\right)$, $\rho \in (-1; 1)$, MI and SMI have analytical form

$$I(X; Y) = -\frac{d}{2} \log(1 - \rho^2),$$

$$SI(X; Y) = \frac{\rho^2}{2d} {}_3F_2\left(1, 1, \frac{3}{2}; \frac{d}{2} + 1, 2; \rho^2\right),$$

where ${}_3F_2$ is the *generalized hypergeometric function*. Additionally, the following limits hold:

$$\lim_{d \rightarrow \infty} I(X; Y) = +\infty \quad \lim_{d \rightarrow \infty} SI(X; Y) = 0$$

$$\lim_{\rho^2 \rightarrow 1} I(X; Y) = +\infty \quad \lim_{\rho^2 \rightarrow 1} SI(X; Y) \leq \frac{3}{d-1}.$$

This example shows that

- SMI suffers the curse of dimensionality
- SMI rapidly saturates

Remarks.

1. k -SMI inherits the same issues.
2. Applying $V = \mathbf{1} \cdot e_1^T$ to the random vectors from Lemma, one can observe

$$SI_k(VX; VY) > SI_k(X; Y).$$

In contrast,

$$I(VX; VY) < I(X; Y).$$

Synthetic Experiments

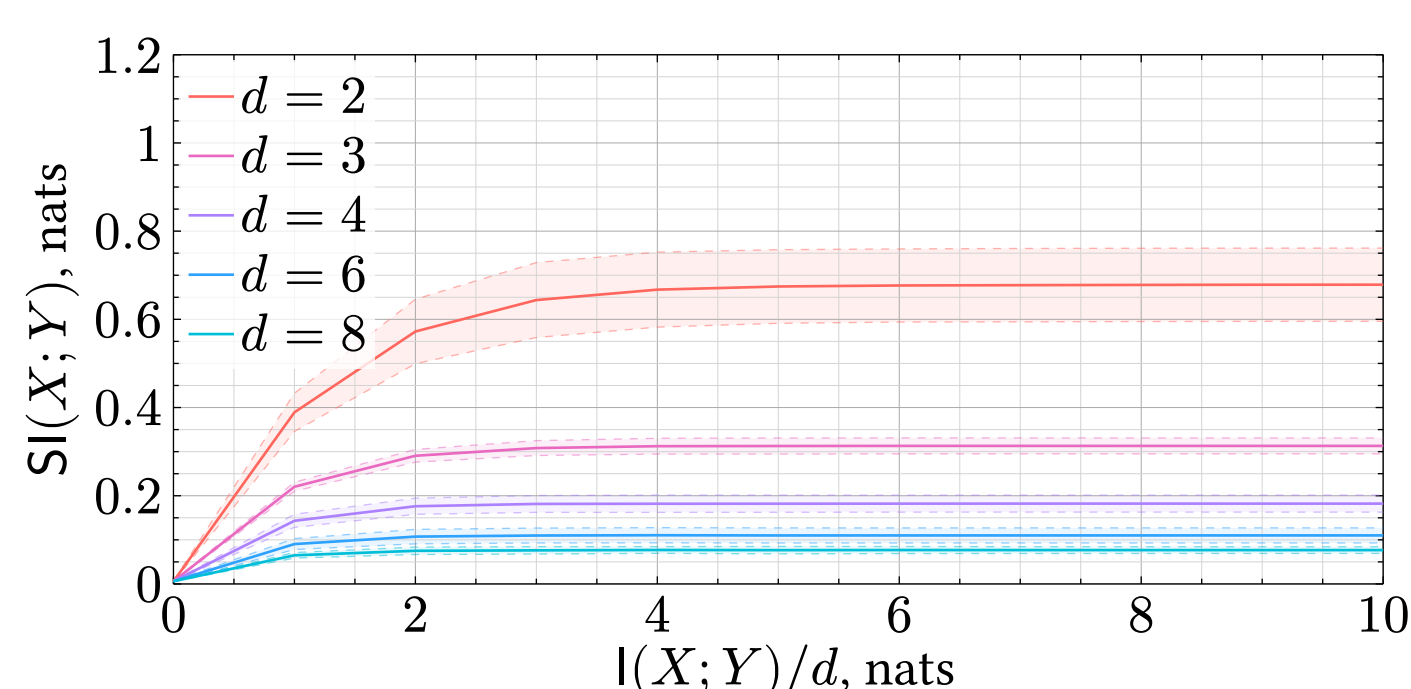


Fig. 2: Correlated Normal

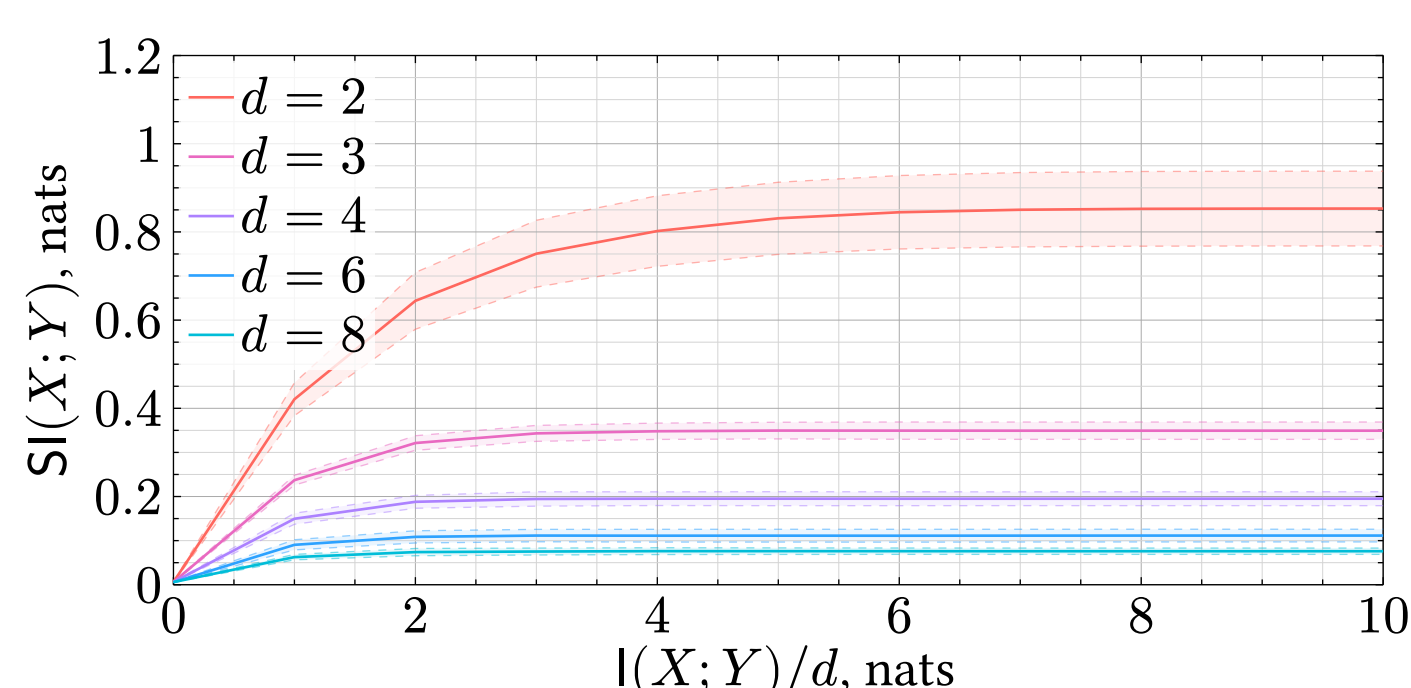


Fig. 3: Correlated Uniform

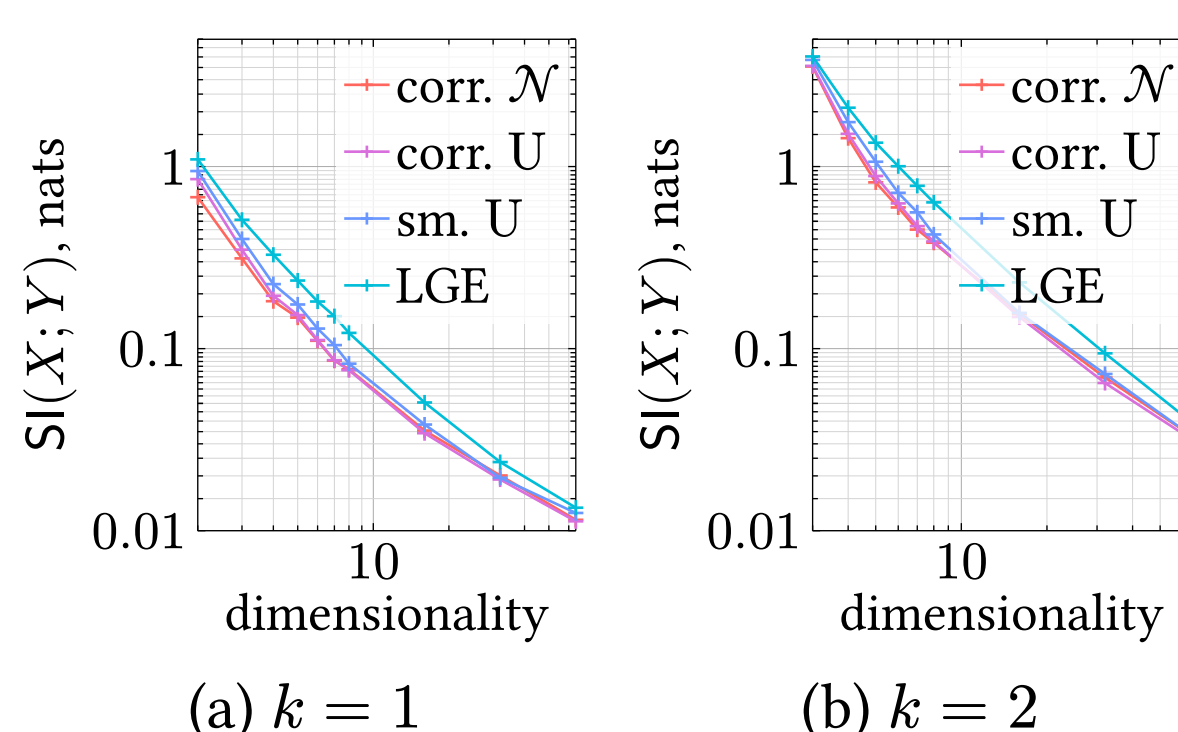


Fig. 4: Decaying trends of k -SMI. We plot saturated values of k -SMI against data dim d .

SMI: Gaussian channel

Let X be a zero-mean d -dimensional random vector, and let $Z \sim \mathcal{N}(0, \sigma^2 I)$ be an independent noise. Consider AWGN channel $X \rightarrow X + Z$ and the problem,

$$\sup_{\mathbb{E} X_i^2 = 1} I(X; X + Z) = \frac{d}{2} \log\left(1 + \frac{1}{\sigma^2}\right),$$

where $X_{\text{opt}} \sim \mathcal{N}(0, I)$.

1. *Whitening:* $\Sigma^{-1/2} X$;
2. *Standardization:* $D^{-1/2} X$, $D = \text{diag}(\Sigma)$.

	MI		2-SMI	
	$\Sigma^{-1/2}$	$D^{-1/2}$	$\Sigma^{-1/2}$	$D^{-1/2}$
X'	7.48 ± 0.01	3.04 ± 0.01	0.96 ± 0.04	2.46 ± 0.03
X''	7.49 ± 0.02	3.04 ± 0.01	0.82 ± 0.05	2.49 ± 0.05

SMI: Representation Learning

Let X be a random vector and f be an encoder (modeled by a NN).

■ In Deep InfoMax we learn most informative embeddings by

$$I(X; f(X)) \rightarrow \max$$

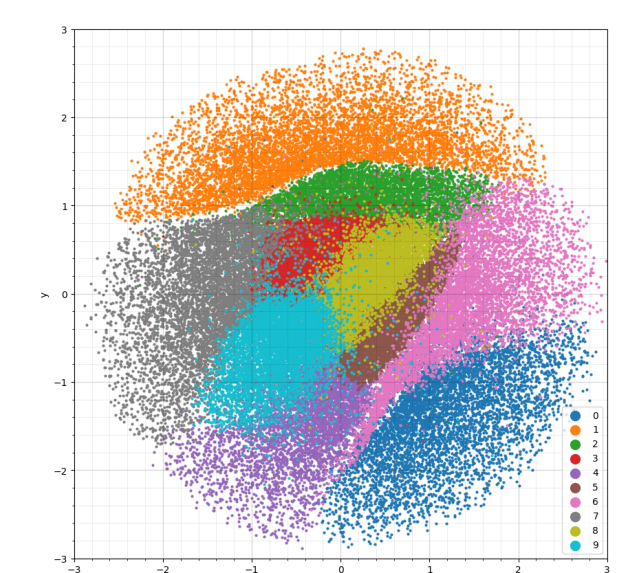
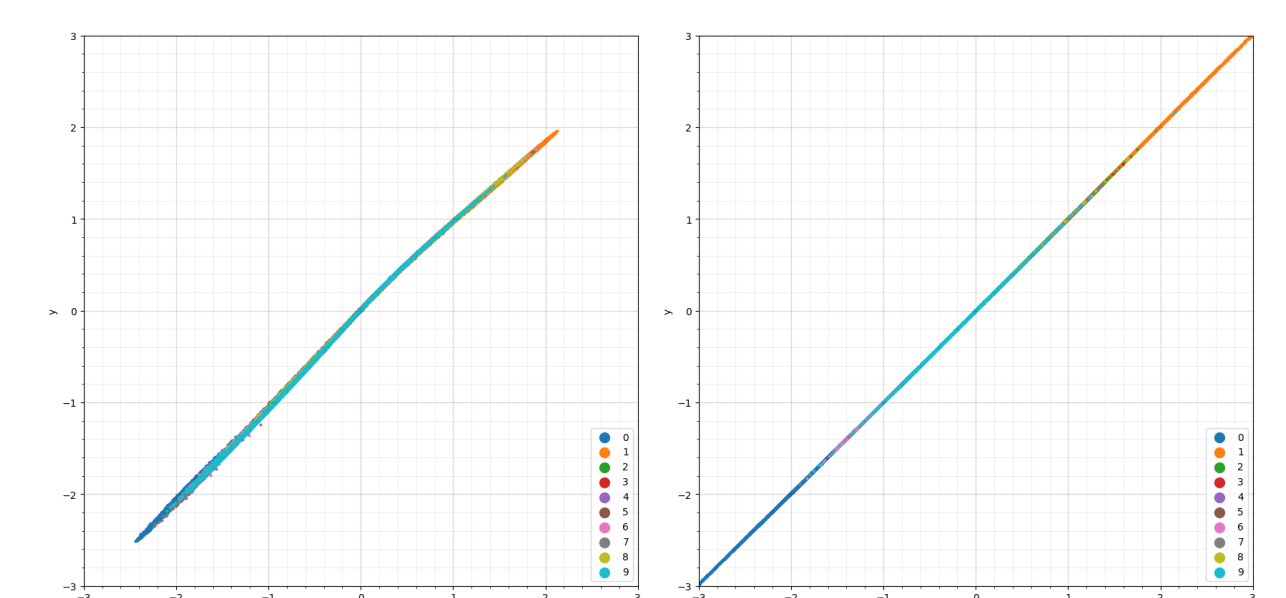


Fig. 5: MI $\rightarrow$ max, 2000 epochs.

■ Replace MI with SMI

$$SI(X; f(X)) \rightarrow \max$$



(a) SMI $\rightarrow$ max, 10 epochs.

(b) SMI $\rightarrow$ max, 2000 epochs.

- MI produces clustered low-redundancy representations
- SMI maximization results in immediate (after 10 epochs) collapse

Takeaways

1. **No Free Lunch.**
2. **Existing works that utilize SMI should be reevaluated to account for the biases uncovered.**
3. **Similar biases might arise in other slicing-based methods.**