

# An Introduction to Rate-Distortion-Perception Theory

Jun Chen  
McMaster University

North American School of Information Theory  
June 22, 2026

# Fidelity vs. Realism



# Distortion Measure vs. Perception Measure

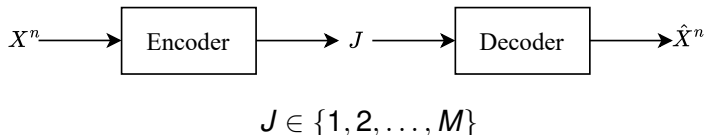
- Distortion measure:  $d(x, \hat{x})$

$$d(x, \hat{x}) = \|x - \hat{x}\|^2$$

- Perception measure:  $\phi(p_X, p_{\hat{X}})$

$$\phi(p_X, p_{\hat{X}}) = W_2^2(p_X, p_{\hat{X}}) = \inf_{p_{X\hat{X}} \in \Pi(p_X, p_{\hat{X}})} \mathbb{E}[\|X - \hat{X}\|^2]$$

# Rate-Distortion Theory: Review

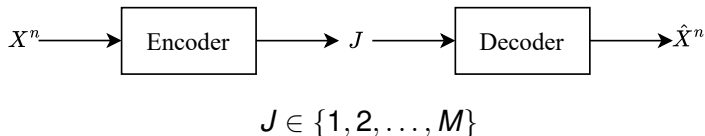


Assume  $X^n \sim p_X^n$ :

$$D^{(\infty)}(R) = \lim_{n \rightarrow \infty} \inf_{\mathbf{E}, \mathbf{D}} \frac{1}{n} \sum_{t=1}^n \mathbb{E}[\|X_t - \hat{X}_t\|^2]$$

s.t.  $\frac{1}{n} \log M \leq R$

# Rate-Distortion Theory: Review



$$D^{(\infty)}(R) = \inf_{p_{\hat{X}|X}} \mathbb{E}[\|X - \hat{X}\|^2]$$
$$\text{s.t. } I(X; \hat{X}) \leq R$$

Shannon's rate-distortion theorem

# Rate-Distortion Theory: Review

- Test channel  $p_{\hat{X}|X}$  satisfying  $I(X; \hat{X}) \leq R$ 
  - Joint distribution  $p_{X\hat{X}} = p_X p_{\hat{X}|X}$

$$\mathcal{T}(p_{X\hat{X}}) = \{(x^n, \hat{x}^n) : \text{joint empirical distribution } \gamma_{x^n \hat{x}^n} \approx p_{X\hat{X}}\}$$

joint-typicality encoding

- Output distribution  $p_{\hat{X}}$

codebook generation

# Rate-Distortion Theory: Review

- Codebook generation:  $\hat{X}^n(j) \sim p_{\hat{X}}^n, j = 1, 2, \dots, M \approx 2^{nI(X; \hat{X})}$

|                |                |     |                |
|----------------|----------------|-----|----------------|
| $\hat{X}_1(1)$ | $\hat{X}_2(1)$ | ... | $\hat{X}_n(1)$ |
| $\vdots$       | $\vdots$       |     | $\vdots$       |
| $\hat{X}_1(j)$ | $\hat{X}_2(j)$ | ... | $\hat{X}_n(j)$ |
| $\vdots$       | $\vdots$       |     | $\vdots$       |
| $\hat{X}_1(M)$ | $\hat{X}_2(M)$ | ... | $\hat{X}_n(M)$ |

# Rate-Distortion Theory: Review

- Encoding:  $(X^n, \hat{X}^n(j)) \in \mathcal{T}(p_{X\hat{X}})$

$X_1 \ X_2 \ \dots \ X_n$

|                |                |     |                |
|----------------|----------------|-----|----------------|
| $\hat{X}_1(1)$ | $\hat{X}_2(1)$ | ... | $\hat{X}_n(1)$ |
| $\vdots$       | $\vdots$       |     | $\vdots$       |
| $\hat{X}_1(j)$ | $\hat{X}_2(j)$ | ... | $\hat{X}_n(j)$ |
| $\vdots$       | $\vdots$       |     | $\vdots$       |
| $\hat{X}_1(M)$ | $\hat{X}_2(M)$ | ... | $\hat{X}_n(M)$ |

# Rate-Distortion Theory: Review

- Decoding:

$X_1 \ X_2 \ \dots \ X_n$

|                |                |         |                |
|----------------|----------------|---------|----------------|
| $\hat{X}_1(1)$ | $\hat{X}_2(1)$ | $\dots$ | $\hat{X}_n(1)$ |
| $\vdots$       | $\vdots$       |         | $\vdots$       |
| $\hat{X}_1(j)$ | $\hat{X}_2(j)$ | $\dots$ | $\hat{X}_n(j)$ |
| $\vdots$       | $\vdots$       |         | $\vdots$       |
| $\hat{X}_1(M)$ | $\hat{X}_2(M)$ | $\dots$ | $\hat{X}_n(M)$ |

$$\left. \vphantom{\begin{matrix} \hat{X}_1(1) & \hat{X}_2(1) & \dots & \hat{X}_n(1) \\ \vdots & \vdots & & \vdots \\ \hat{X}_1(j) & \hat{X}_2(j) & \dots & \hat{X}_n(j) \\ \vdots & \vdots & & \vdots \\ \hat{X}_1(M) & \hat{X}_2(M) & \dots & \hat{X}_n(M) \end{matrix}} \right\} \begin{aligned} & \frac{1}{n} \sum_{t=1}^n \mathbb{E} \left[ \|X_t - \hat{X}_t(J)\|^2 \right] \\ & \approx \mathbb{E} \left[ \|X - \hat{X}\|^2 \right] \end{aligned}$$

# Distortion-Rate-Perception Function

- Blau and Michaeli 2019

$$\begin{aligned} D^{(\infty)}(R, P) &= \inf_{p_{\hat{X}|X}} \mathbb{E}[\|X - \hat{X}\|^2] \\ \text{s.t. } & I(X; \hat{X}) \leq R \\ & W_2^2(p_X, p_{\hat{X}}) \leq P \end{aligned}$$

# Problem Formulation

- Test channel  $p_{\hat{X}|X}$  satisfying  $I(X; \hat{X}) \leq R$  and  $W_2^2(p_X, p_{\hat{X}}) \leq P$ 
  - Joint distribution  $p_{X\hat{X}} = p_X p_{\hat{X}|X}$

$$\mathcal{T}(p_{X\hat{X}}) = \{(x^n, \hat{x}^n) : \text{joint empirical distribution } \gamma_{x^n \hat{x}^n} \approx p_{X\hat{X}}\}$$

joint-typicality encoding

- Output distribution  $p_{\hat{X}}$

codebook generation

# Problem Formulation

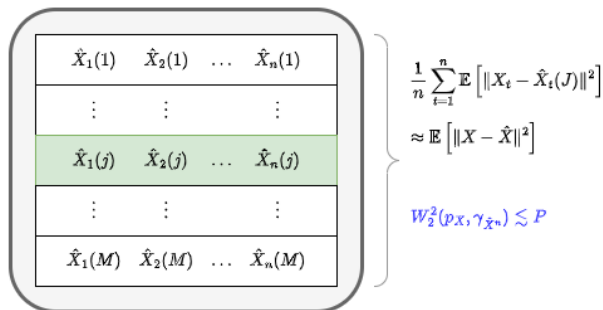
- Test channel  $p_{\hat{X}|X}$  satisfying  $I(X; \hat{X}) \leq R$  and  $W_2^2(p_X, p_{\hat{X}}) \leq P$

$$\begin{array}{cccc}
 X_1 & X_2 & \dots & X_n
 \end{array}
 \quad
 \left.
 \begin{array}{|c|}
 \hline
 \hat{X}_1(1) \quad \hat{X}_2(1) \quad \dots \quad \hat{X}_n(1) \\
 \hline
 \vdots \quad \quad \quad \vdots \quad \quad \quad \vdots \\
 \hline
 \hat{X}_1(j) \quad \hat{X}_2(j) \quad \dots \quad \hat{X}_n(j) \\
 \hline
 \vdots \quad \quad \quad \vdots \quad \quad \quad \vdots \\
 \hline
 \hat{X}_1(M) \quad \hat{X}_2(M) \quad \dots \quad \hat{X}_n(M) \\
 \hline
 \end{array}
 \right\}
 \begin{aligned}
 & \frac{1}{n} \sum_{t=1}^n \mathbb{E} \left[ \|X_t - \hat{X}_t(J)\|^2 \right] \\
 & \approx \mathbb{E} \left[ \|X - \hat{X}\|^2 \right]
 \end{aligned}$$

# Problem Formulation

- Test channel  $p_{\hat{X}|X}$  satisfying  $I(X; \hat{X}) \leq R$  and  $W_2^2(p_X, p_{\hat{X}}) \leq P$

$X_1 \ X_2 \ \dots \ X_n$



# Problem Formulation

- Test channel  $p_{\hat{X}|X}$  satisfying  $I(X; \hat{X}) \leq R$  and  $W_2^2(p_X, p_{\hat{X}}) \leq P$

$X_1 \ X_2 \ \dots \ X_n$

|                |                |     |                |
|----------------|----------------|-----|----------------|
| $\hat{X}_1(1)$ | $\hat{X}_2(1)$ | ... | $\hat{X}_n(1)$ |
| $\vdots$       | $\vdots$       |     | $\vdots$       |
| $\hat{X}_1(j)$ | $\hat{X}_2(j)$ | ... | $\hat{X}_n(j)$ |
| $\vdots$       | $\vdots$       |     | $\vdots$       |
| $\hat{X}_1(M)$ | $\hat{X}_2(M)$ | ... | $\hat{X}_n(M)$ |

$$\frac{1}{n} \sum_{t=1}^n \mathbb{E} [\|X_t - \hat{X}_t(j)\|^2]$$

$$\approx \mathbb{E} [\|X - \hat{X}\|^2]$$

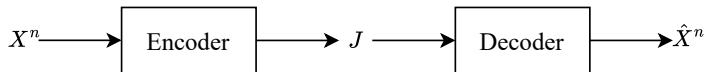
$$W_2^2 \left( p_X, \frac{1}{n} \sum_{t=1}^n p_{\hat{X}_t} \right) \lesssim P$$

# Problem Formulation

- Test channel  $p_{\hat{X}|X}$  satisfying  $I(X; \hat{X}) \leq R$  and  $W_2^2(p_X, p_{\hat{X}}) \leq P$

$$\begin{array}{cccc}
 X_1 & X_2 & \dots & X_n
 \end{array}
 \left.
 \begin{array}{c}
 \begin{array}{cccc}
 \hat{X}_1(1) & \hat{X}_2(1) & \dots & \hat{X}_n(1) \\
 \vdots & \vdots & & \vdots \\
 \hat{X}_1(j) & \hat{X}_2(j) & \dots & \hat{X}_n(j) \\
 \vdots & \vdots & & \vdots \\
 \hat{X}_1(M) & \hat{X}_2(M) & \dots & \hat{X}_n(M)
 \end{array}
 \end{array}
 \right\}
 \begin{array}{l}
 \frac{1}{n} \sum_{t=1}^n \mathbb{E} \left[ \|X_t - \hat{X}_t(J)\|^2 \right] \\
 \approx \mathbb{E} \left[ \|X - \hat{X}\|^2 \right] \\
 \frac{1}{n} \sum_{t=1}^n W_2^2(p_X, p_{\hat{X}_t}) \lesssim P
 \end{array}$$

# Problem Formulation



$$J \in \{1, 2, \dots, M\}$$

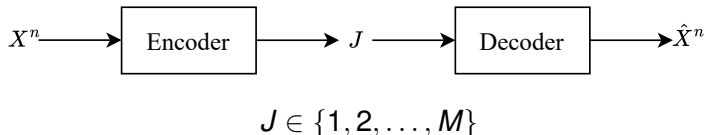
Assume  $X^n \sim p_X^n$ :

$$D^{(\infty)}(R, P) = \lim_{n \rightarrow \infty} \inf_{E, D} \frac{1}{n} \sum_{t=1}^n \mathbb{E}[\|X_t - \hat{X}_t\|^2]$$

s.t.

$$\frac{1}{n} \log M \leq R$$
$$\frac{1}{n} \sum_{t=1}^n W_2^2(p_X, p_{\hat{X}_t}) \leq P$$

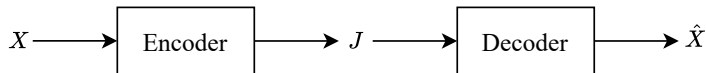
# Coding Theorem



- Blau and Michaeli 2019, Theis and Wagner 2021, Chen et al. 2022

$$\begin{aligned} D^{(\infty)}(R, P) &= \inf_{p_{\hat{X}|X}} \mathbb{E}[\|X - \hat{X}\|^2] \\ \text{s.t. } I(X; \hat{X}) &\leq R \\ W_2^2(p_X, p_{\hat{X}}) &\leq P \end{aligned}$$

# One-Shot Setting



$$J \in \{1, 2, \dots, M\}$$

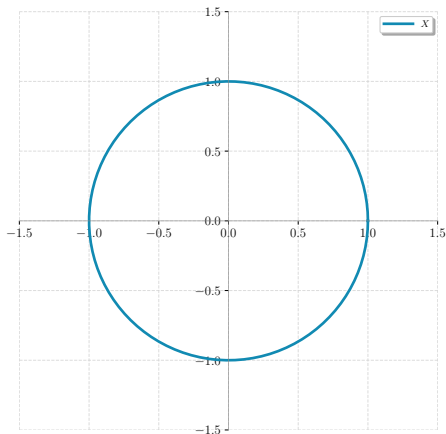
Assume  $X \sim p_X$ :

$$D^{(1)}(R, P) = \inf_{E, D} \mathbb{E}[\|X - \hat{X}\|^2]$$

$$\text{s.t.} \quad \log M \leq R$$

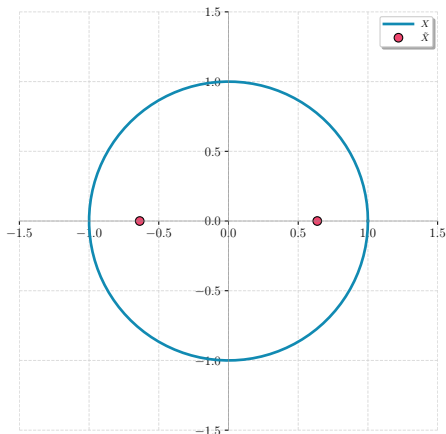
$$W_2^2(p_X, p_{\hat{X}}) \leq P$$

# Quantizing the Unit Circle



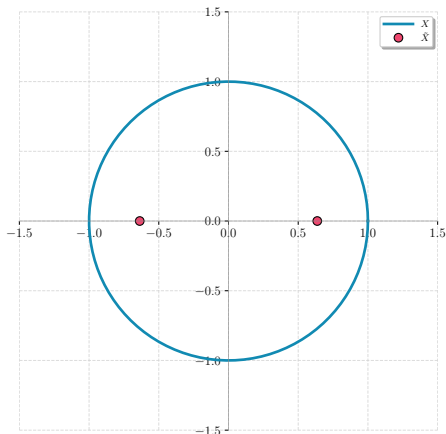
1-bit Quantization

# Quantizing the Unit Circle: No Perception Constraint



$$\mathbb{E}[\|X - \tilde{X}\|^2] = 1 - \frac{4}{\pi^2}$$

# Quantizing the Unit Circle: No Perception Constraint

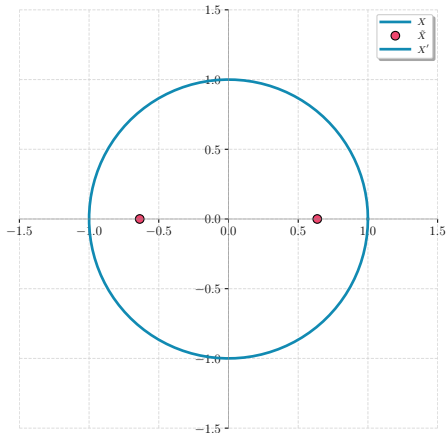


$$W_2^2(p_X, p_{\tilde{X}}) = 1 - \frac{4}{\pi^2}$$

- Yan et al. 2022

$$D^{(1)}(R, P) = D^{(1)}(R, \infty) \text{ if and only if } P \geq D^{(1)}(R, \infty)$$

# Quantizing the Unit Circle: Perfect Perception



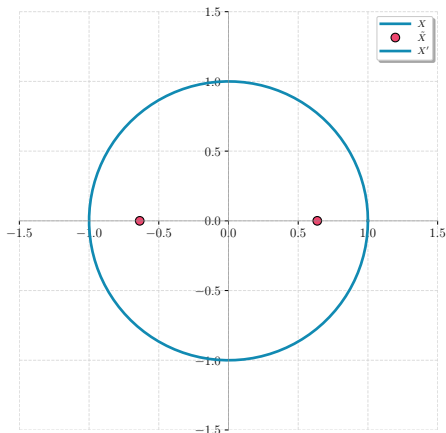
$$\mathbb{E}[\|X' - \tilde{X}\|^2] = W_2^2(p_X, p_{\tilde{X}}) = \mathbb{E}[\|X - \tilde{X}\|^2]$$

$$\mathbb{E}[\|X - X'\|^2] = \mathbb{E}[\|X - \tilde{X}\|^2] + \mathbb{E}[\|X' - \tilde{X}\|^2] = 2\mathbb{E}[\|X - \tilde{X}\|^2]$$

- Yan et al. 2021, Theis and Agustsson 2021

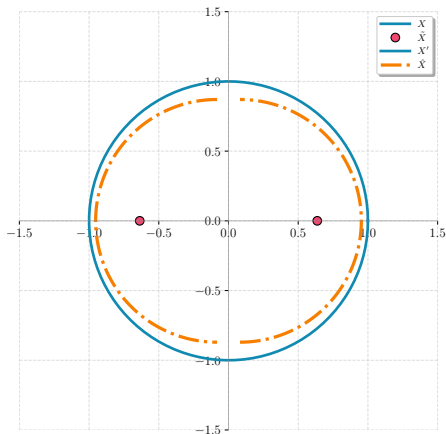
$$D^{(1)}(R, 0) = 2D^{(1)}(R, \infty)$$

# Quantizing the Unit Circle: Linear Interpolation



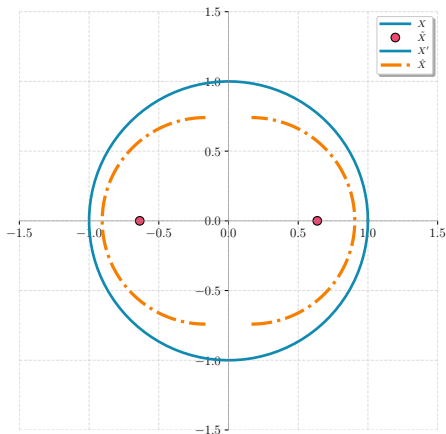
$$\hat{X} = (1 - \alpha)X' + \alpha\tilde{X}$$

# Quantizing the Unit Circle: Linear Interpolation



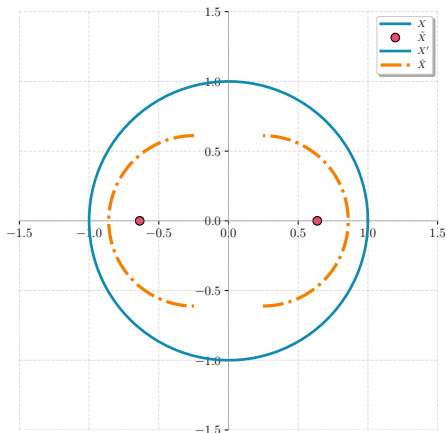
$$P = 0.01$$

# Quantizing the Unit Circle: Linear Interpolation



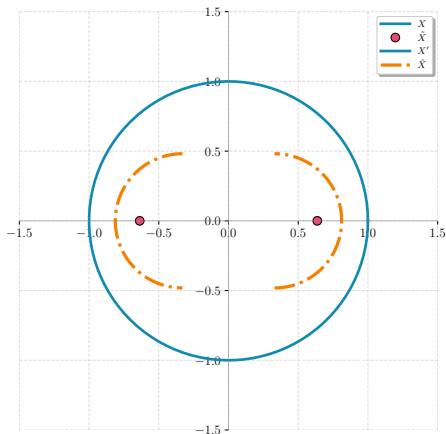
$$P = 0.04$$

# Quantizing the Unit Circle: Linear Interpolation



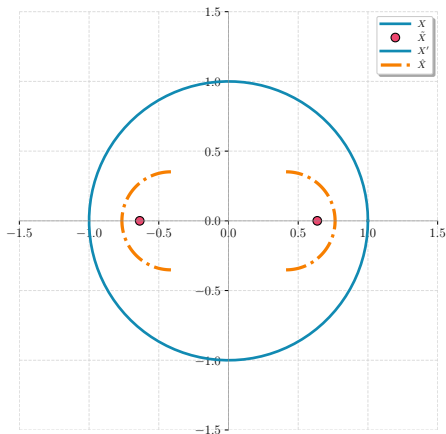
$$P = 0.09$$

# Quantizing the Unit Circle: Linear Interpolation



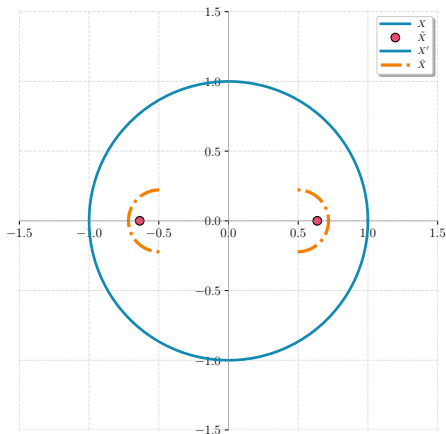
$$P = 0.16$$

# Quantizing the Unit Circle: Linear Interpolation



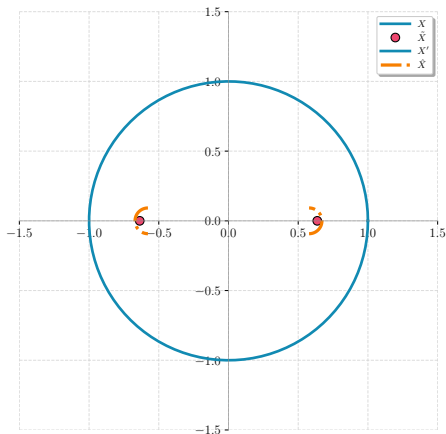
$$P = 0.25$$

# Quantizing the Unit Circle: Linear Interpolation



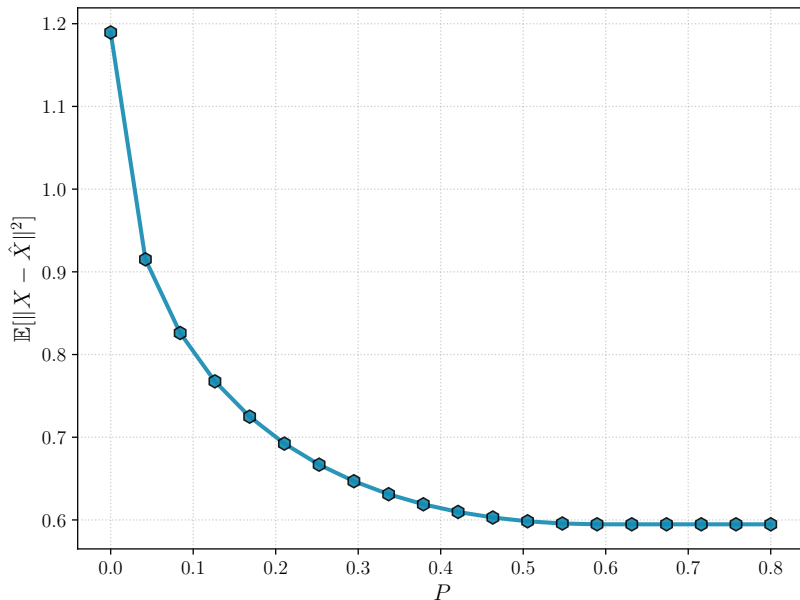
$$P = 0.36$$

# Quantizing the Unit Circle: Linear Interpolation

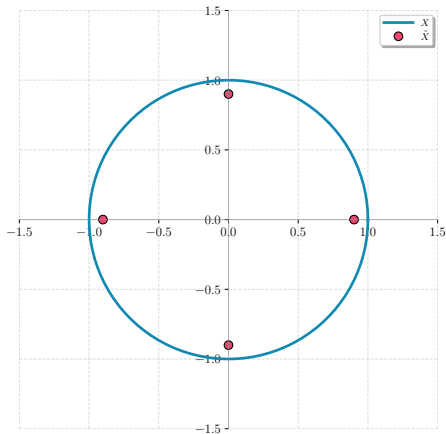


$$P = 0.49$$

# Distortion-Perception Tradeoff

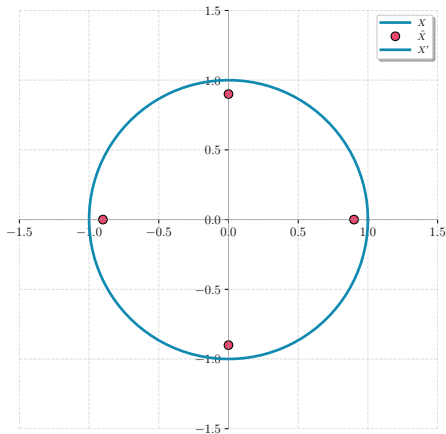


# Quantizing the Unit Circle: No Perception Constraint



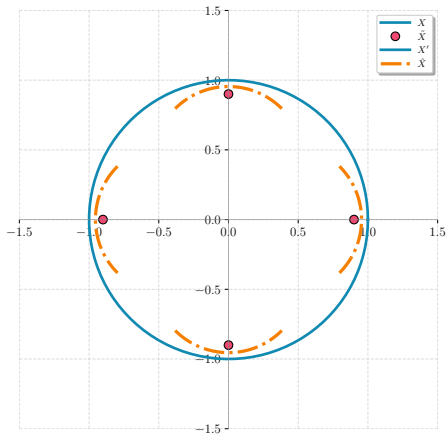
2-bit Quantization

# Quantizing the Unit Circle: Perfect Perception



2-bit Quantization

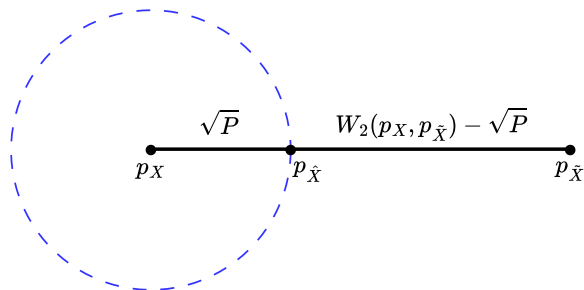
# Quantizing the Unit Circle: Linear Interpolation



2-bit Quantization

# Optimality of Linear Interpolation

- Zhang et al. 2021, Freirich et al. 2021



$$W_2(p_X, p_{\hat{X}}) \leq \sqrt{P}$$

$$\begin{aligned}\mathbb{E}[\|X - \hat{X}\|^2] &= \mathbb{E}[\|X - \tilde{X}\|^2] + \mathbb{E}[\|\tilde{X} - \hat{X}\|^2] \\ &\geq \mathbb{E}[\|X - \tilde{X}\|^2] + W_2^2(p_{\tilde{X}}, p_{\hat{X}}) \\ &\geq \mathbb{E}[\|X - \tilde{X}\|^2] + [(W_2(p_X, p_{\tilde{X}}) - \sqrt{P})_+]^2\end{aligned}$$

# Optimality of Linear Interpolation

Consider  $\hat{X} = (1 - \alpha)X' + \alpha\tilde{X}$  with  $\alpha = \frac{\sqrt{P}}{W_2(p_X, p_{\tilde{X}})} \wedge 1$ .

- Distortion loss

$$\begin{aligned}\mathbb{E}[\|X - \hat{X}\|^2] &= \mathbb{E}[\|X - \tilde{X} + (1 - \alpha)(\tilde{X} - X')\|^2] \\ &= \mathbb{E}[\|X - \tilde{X}\|^2] + (1 - \alpha)^2 \mathbb{E}[\|X' - \tilde{X}\|^2] \\ &= \mathbb{E}[\|X - \tilde{X}\|^2] + (1 - \alpha)^2 W_2^2(p_X, p_{\tilde{X}}) \\ &= \mathbb{E}[\|X - \tilde{X}\|^2] + [(\color{red}{W_2(P_X, p_{\tilde{X}})} - \sqrt{P})_+]^2\end{aligned}$$

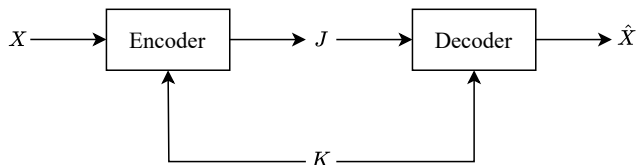
- Perception loss

$$\begin{aligned}W_2^2(p_X, p_{\hat{X}}) &\leq \mathbb{E}[\|X' - \hat{X}\|^2] \\ &= \mathbb{E}[\|\alpha(X' - \tilde{X})\|^2] \\ &= \alpha^2 \mathbb{E}[\|X' - \tilde{X}\|^2] \\ &= \alpha^2 W_2^2(p_X, p_{\tilde{X}}) \\ &= P \wedge W_2^2(p_X, p_{\tilde{X}})\end{aligned}$$

- Yan et al. 2022

$$D^{(1)}(R, P) = D^{(1)}(R, \infty) + [(\sqrt{D^{(1)}(R, \infty)} - \sqrt{P})_+]^2$$

# Common Randomness



$$J \in \{1, 2, \dots, M\} \quad K \in \{1, 2, \dots, N\}$$

Assume  $X \sim p_X$ :

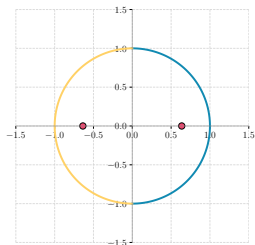
$$D^{(1)}(R, C, P) = \inf_{E, D} \mathbb{E}[\|X - \hat{X}\|^2]$$

$$\text{s.t.} \quad \log M \leq R$$

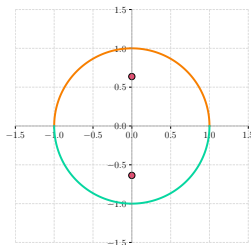
$$\log N \leq C$$

$$W_2^2(p_X, p_{\hat{X}}) \leq P$$

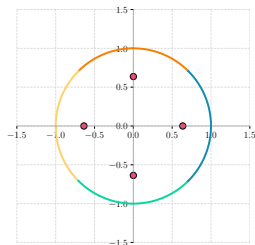
# Quantizing the Unit Circle: No Perception Constraint



$$\mathbb{E}[\|X - \tilde{X}\|^2 | K=1]$$



$$\mathbb{E}[\|X - \tilde{X}\|^2 | K=2]$$

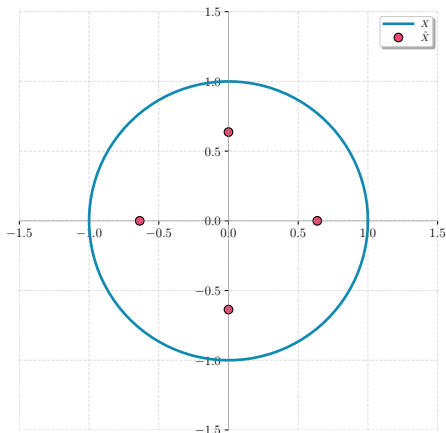


$$W_2^2(p_X, p_{\tilde{X}})$$

1-bit Quantization with 1-bit of Common Randomness

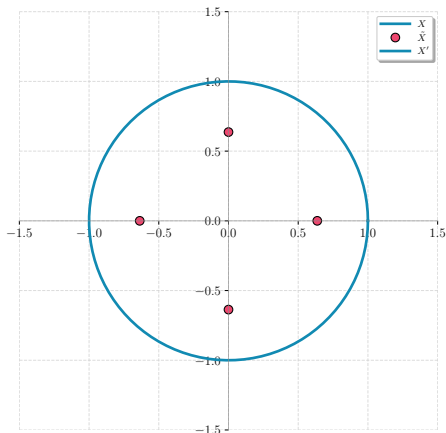
$$\mathbb{E}[\|X - \tilde{X}\|^2] = 1 - \frac{4}{\pi^2} \quad W_2^2(p_X, p_{\tilde{X}}) = 1 - \frac{4(2\sqrt{2} - 1)}{\pi^2}$$

# Quantizing the Unit Circle: No Perception Constraint



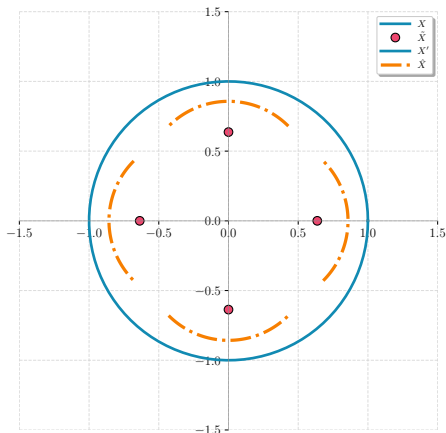
$$\mathbb{E}[\|X - \tilde{X}\|^2] > W_2^2(p_X, p_{\tilde{X}})$$

# Quantizing the Unit Circle: Perfect Perception



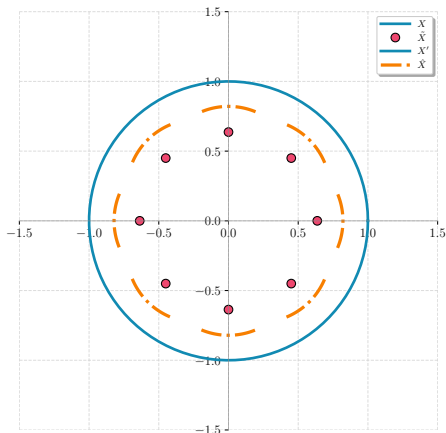
$$\mathbb{E}[\|X - X'\|^2] = \mathbb{E}[\|X - \tilde{X}\|^2] + W_2^2(p_X, p_{\tilde{X}}) < 2\mathbb{E}[\|X - \tilde{X}\|^2]$$

# Quantizing the Unit Circle: Linear Interpolation



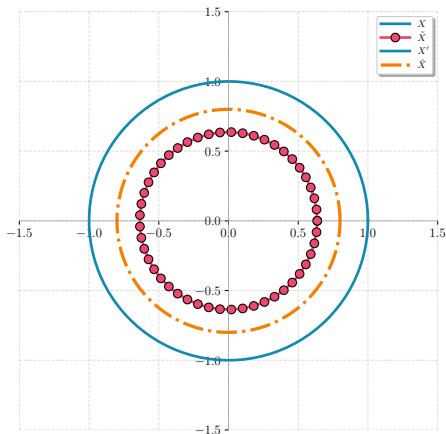
$$\mathbb{E}[\|X - \hat{X}\|^2] = \mathbb{E}[\|X - \tilde{X}\|^2] + [(W_2(p_X, p_{\tilde{X}}) - \sqrt{P})_+]^2$$

# 2-Bits of Common Randomness



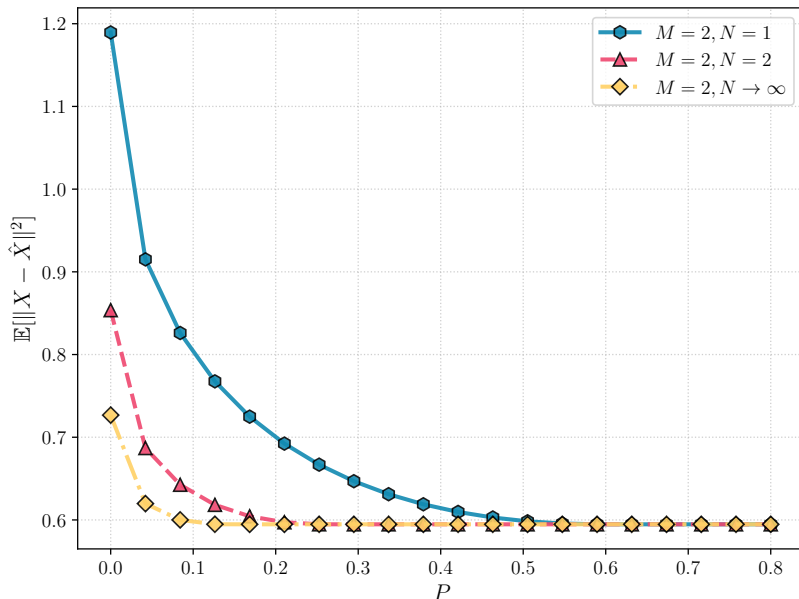
1-bit Quantization with 2-bits of Common Randomness

# Unlimited Common Randomness

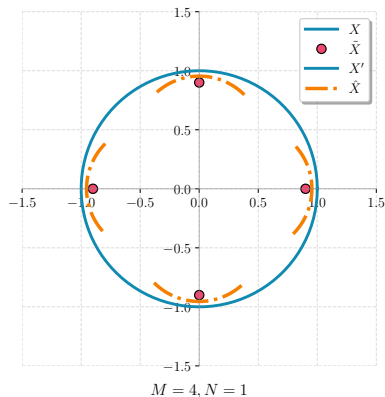
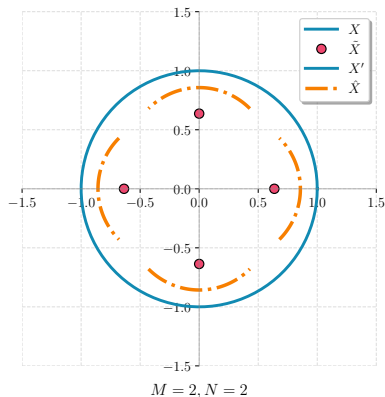


dithered quantization

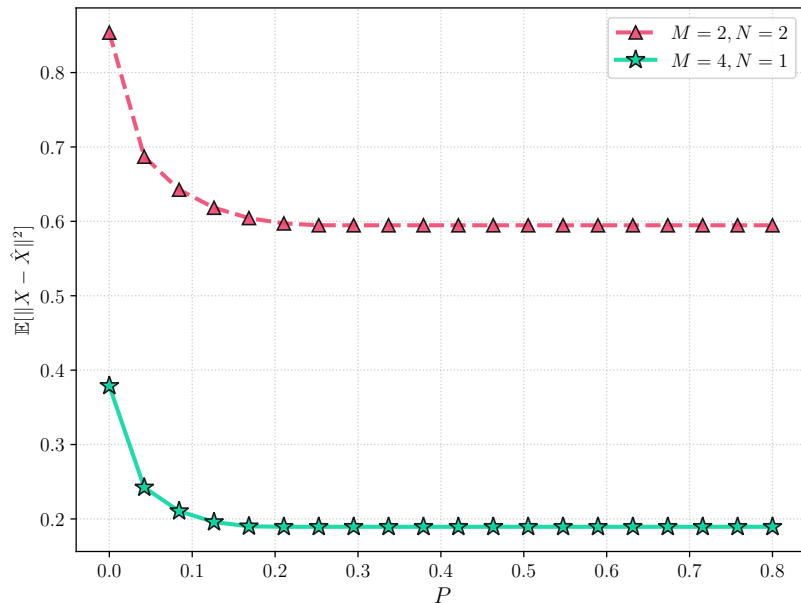
# Distortion-Perception Tradeoff



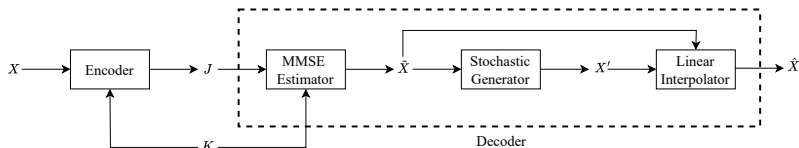
# Information Bits vs. Shared Bits



# Information Bits vs. Shared Bits



# Architectural Principles



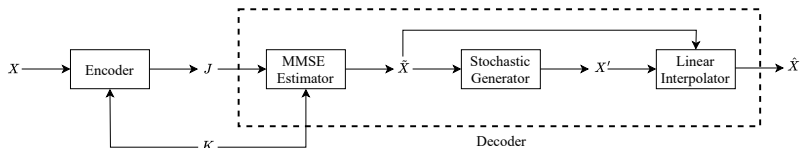
$$J \in \{1, 2, \dots, M\} \quad K \in \{1, 2, \dots, N\}$$

- Freirich et al. 2021, Chen et al. 2026

$$D^{(1)}(R, C, P) = \inf_E \mathbb{E}[\|X - \tilde{X}\|^2] + [(W_2(p_X, p_{\tilde{X}}) - \sqrt{P})_+]^2$$

s.t.  $\log M \leq R$   
 $\log N \leq C$   
 $\tilde{X} = \mathbb{E}[X|J, K]$

# MMSE Estimator

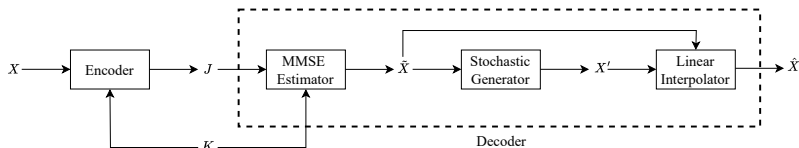


$$J \in \{1, 2, \dots, M\} \quad K \in \{1, 2, \dots, N\}$$

- MMSE estimator:  $\tilde{X} = \mathbb{E}[X|J, K]$

supervised learning

# Stochastic Generator



$$J \in \{1, 2, \dots, M\} \quad K \in \{1, 2, \dots, N\}$$

- Stochastic generator: Simulate the optimal transport plan  $p_{X'|\tilde{X}}$  that achieves  $W_2^2(p_X, p_{\tilde{X}})$

- Without common randomness

$$\begin{aligned}\mathbb{E}[\|X - \tilde{X}\|^2] &= W_2^2(p_X, p_{\tilde{X}}) \\ \implies p_{X'|X} &= p_{X|\tilde{X}}\end{aligned}$$

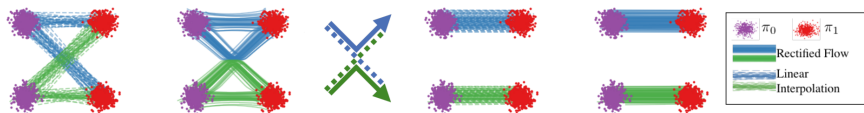
- Posterior sampling

# Stochastic Generator

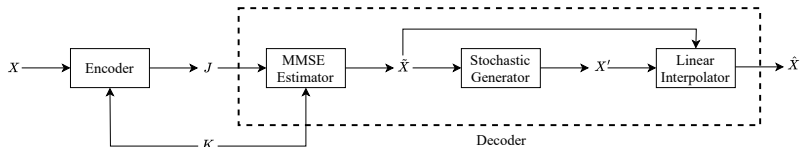
- With common randomness

$$\mathbb{E}[\|X - \tilde{X}\|^2] > W_2^2(p_X, p_{\tilde{X}})$$
$$\implies p_{X'|X} \neq p_{X|\tilde{X}}$$

- Rectified flow



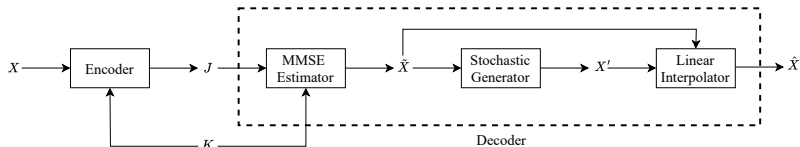
# Linear Interpolator



$$J \in \{1, 2, \dots, M\} \quad K \in \{1, 2, \dots, N\}$$

- Linear interpolator:  $\hat{X} = (1 - \alpha)X' + \alpha\tilde{X}$  with  $\alpha = \frac{\sqrt{P}}{W_2(p_X, p_{\tilde{X}})} \wedge 1$

# Encoder

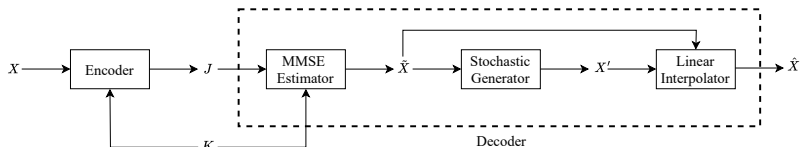


$$J \in \{1, 2, \dots, M\} \quad K \in \{1, 2, \dots, N\}$$

- Encoder:  $(X, K) \rightarrow J$

nonlinear transformation + quantization

# Universal Representation



$$J \in \{1, 2, \dots, M\} \quad K \in \{1, 2, \dots, N\}$$

A representation is said to be universal if the choice of the optimal encoder does not depend on the specific perception constraint  $P$ .

# Universal Representation

$$\begin{aligned} D^{(1)}(R, C, P) = \inf_E \mathbb{E}[\|X - \tilde{X}\|^2] + [(\mathcal{W}_2(p_X, p_{\tilde{X}}) - \sqrt{P})_+]^2 \\ \text{s.t.} \quad \log M \leq R \\ \log N \leq C \\ \tilde{X} = \mathbb{E}[X|J, K] \end{aligned}$$

Universal representation exists in the absence of common randomness ( $C = 0$ ):

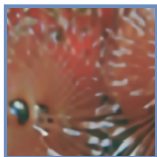
$$\begin{aligned} D^{(1)}(R, 0, P) = \inf_E \mathbb{E}[\|X - \tilde{X}\|^2] + [(\sqrt{\mathbb{E}[\|X - \tilde{X}\|^2]} - \sqrt{P})_+]^2 \\ \text{s.t.} \quad \log M \leq R \\ \tilde{X} = \mathbb{E}[X|J] \end{aligned}$$

# Experimental Results

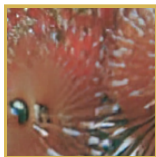
$X$



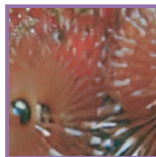
$\tilde{X}$



$X'$

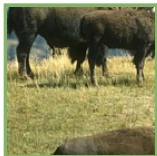


$\hat{X}_{0.12}$

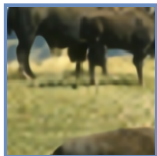


# Experimental Results

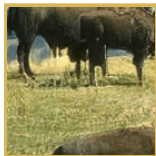
$X$



$\tilde{X}$



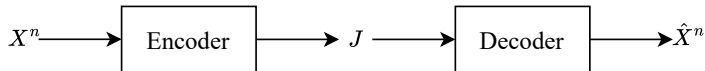
$X'$



$\hat{X}_{0.3}$



# Asymptotic Setting

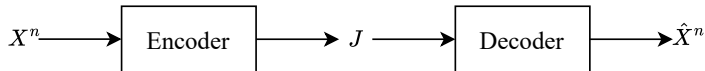


$$J \in \{1, 2, \dots, M\}$$

Assume  $X^n \sim p_X^n$ :

$$\begin{aligned} D^{(\infty)}(R, P) &= \lim_{n \rightarrow \infty} \inf_{\mathbf{E}, \mathbf{D}} \frac{1}{n} \sum_{t=1}^n \mathbb{E}[\|X_t - \hat{X}_t\|^2] \\ \text{s.t. } &\frac{1}{n} \log M \leq R \\ &\frac{1}{n} \sum_{t=1}^n W_2^2(p_X, p_{\hat{X}_t}) \leq P \end{aligned}$$

# Asymptotic Setting



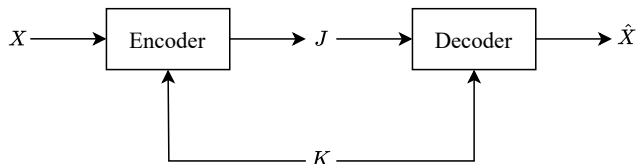
$$J \in \{1, 2, \dots, M\}$$

- Blau and Michaeli 2019, Theis and Wagner 2021, Chen et al. 2022

$$\begin{aligned} D^{(\infty)}(R, P) &= \inf_{p_{\hat{X}|X}} \mathbb{E}[\|X - \hat{X}\|^2] \\ \text{s.t. } & I(X; \hat{X}) \leq R \\ & W_2^2(p_X, p_{\hat{X}}) \leq P \end{aligned}$$

deterministic encoding and decoding suffices!

# One-Shot Setting



$$J \in \{1, 2, \dots, M\} \quad K \in \{1, 2, \dots, N\}$$

Assume  $X \sim p_X$ :

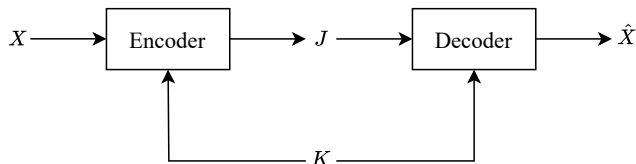
$$D^{(1)}(R, C, P) = \inf_{E, D} \mathbb{E}[\|X - \hat{X}\|^2]$$

$$\text{s.t.} \quad \log M \leq R$$

$$\log N \leq C$$

$$W_2^2(p_X, p_{\hat{X}}) \leq P$$

# One-Shot Setting



$$J \in \{1, 2, \dots, M\} \quad K \in \{1, 2, \dots, N\}$$

- Freirich et al. 2021, Chen et al. 2026

$$D^{(1)}(R, C, P) = \inf_E \mathbb{E}[\|X - \tilde{X}\|^2] + [(W_2(p_X, p_{\tilde{X}}) - \sqrt{P})_+]^2$$

$$\text{s.t.} \quad \log M \leq R$$

$$\log N \leq C$$

$$\tilde{X} = \mathbb{E}[X|J, K]$$

(common) randomness is essential!

# Random vs. Deterministic

- Asymptotic setting ( $R = 1, P = 0$ )

$$\begin{aligned} D^{(\infty)}(1, 0) &= \lim_{n \rightarrow \infty} \inf_{E, D} \frac{1}{n} \sum_{t=1}^n \mathbb{E}[\|X_t - \hat{X}_t\|^2] \\ \text{s.t. } &\frac{1}{n} \log M \leq 1 \\ &\frac{1}{n} \sum_{t=1}^n W_2^2(p_X, p_{\hat{X}_t}) \leq 0 \end{aligned}$$

- One-shot setting ( $R = 1, P = 0$ )

$$\begin{aligned} D^{(1)}(1, C, 0) &= \inf_{E, D} \mathbb{E}[\|X - \hat{X}\|^2] \\ \text{s.t. } &\log M \leq 1 \\ &\log N \leq C \\ &W_2^2(p_X, p_{\hat{X}}) \leq 0 \end{aligned}$$

# Symbol-Level vs. Sequence-Level

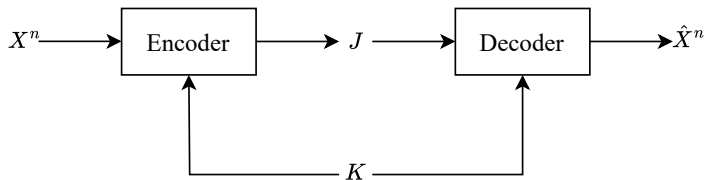
- Symbol-level perception constraint

$$\frac{1}{n} \sum_{t=1}^n W_2^2(p_{X_t}, p_{\hat{X}_t}) \leq P$$

- Sequence-level perception constraint

$$\frac{1}{n} W_2^2(p_{X^n}, p_{\hat{X}^n}) \leq P$$

# An Alternative Formulation



$$J \in \{1, 2, \dots, M\} \quad K \in \{1, 2, \dots, N\}$$

Assume  $X^n \sim p_X^n$ :

$$D^{(\infty)}(R, C, P) = \lim_{n \rightarrow \infty} \inf_{E, D} \frac{1}{n} \sum_{t=1}^n \mathbb{E}[\|X_t - \hat{X}_t\|^2]$$

$$\text{s.t.} \quad \frac{1}{n} \log M \leq R$$

$$\frac{1}{n} \log N \leq C$$

$$\frac{1}{n} W_2^2(p_{X^n}, p_{\hat{X}^n}) \leq P$$

# Fundamental Limits

- Asymptotic setting

$$\begin{aligned} D^{(\infty)}(R, C, P) = & \inf_{\substack{p_{\tilde{X}} \\ \mu, \nu \in \Pi(p_X, p_{\tilde{X}})}} \mathbb{E}_{\mu}[\|X - \tilde{X}\|^2] + [\sqrt{\mathbb{E}_{\nu}[\|X - \tilde{X}\|^2]} - \sqrt{P}]_+^2 \\ \text{s.t. } & I_{\mu}(X; \tilde{X}) \leq R \\ & I_{\nu}(X; \tilde{X}) \leq R + C \\ & \tilde{X} = \mathbb{E}_{\mu}[X|\tilde{X}] \end{aligned}$$

- One-shot setting

$$\begin{aligned} D^{(1)}(R, C, P) = & \inf_E \mathbb{E}[\|X - \tilde{X}\|^2] + [(\mathcal{W}_2(p_X, p_{\tilde{X}}) - \sqrt{P})_+]^2 \\ \text{s.t. } & \log M \leq R \\ & \log N \leq C \\ & \tilde{X} = \mathbb{E}[X|J, K] \end{aligned}$$

# Fundamental Limits

- Asymptotic setting

$$\begin{aligned} D^{(\infty)}(R, C, P) = & \inf_{\substack{p_{\tilde{X}} \\ \mu, \nu \in \Pi(p_X, p_{\tilde{X}})}} \mathbb{E}_{\mu}[\|X - \tilde{X}\|^2] + [\sqrt{\mathbb{E}_{\nu}[\|X - \tilde{X}\|^2]} - \sqrt{P}]_+^2 \\ \text{s.t. } & I_{\mu}(X; \tilde{X}) \leq R \\ & I_{\nu}(X; \tilde{X}) \leq R + C \\ & \tilde{X} = \mathbb{E}_{\mu}[X|\tilde{X}] \end{aligned}$$

- One-shot setting

$$\begin{aligned} D^{(1)}(R, C, P) = & \inf_E \mathbb{E}[\|X - \tilde{X}\|^2] + [(W_2(p_X, p_{\tilde{X}}) - \sqrt{P})_+]^2 \\ \text{s.t. } & \log M \leq R \\ & \log N \leq C \\ & \tilde{X} = \mathbb{E}[X|J, K] \end{aligned}$$

# Operational Encoding vs. Virtual Encoding

|                      |                      |     |                      |
|----------------------|----------------------|-----|----------------------|
| $\hat{x}_1^{(1)}(1)$ | $\hat{x}_2^{(1)}(1)$ | ... | $\hat{x}_n^{(1)}(1)$ |
| $\vdots$             | $\vdots$             |     | $\vdots$             |
| $\hat{x}_1^{(1)}(j)$ | $\hat{x}_2^{(1)}(j)$ | ... | $\hat{x}_n^{(1)}(j)$ |
| $\vdots$             | $\vdots$             |     | $\vdots$             |
| $\hat{x}_1^{(1)}(M)$ | $\hat{x}_2^{(1)}(M)$ | ... | $\hat{x}_n^{(1)}(M)$ |

|                      |                      |     |                      |
|----------------------|----------------------|-----|----------------------|
| $\hat{x}_1^{(2)}(1)$ | $\hat{x}_2^{(2)}(1)$ | ... | $\hat{x}_n^{(2)}(1)$ |
| $\vdots$             | $\vdots$             |     | $\vdots$             |
| $\hat{x}_1^{(2)}(j)$ | $\hat{x}_2^{(2)}(j)$ | ... | $\hat{x}_n^{(2)}(j)$ |
| $\vdots$             | $\vdots$             |     | $\vdots$             |
| $\hat{x}_1^{(2)}(M)$ | $\hat{x}_2^{(2)}(M)$ | ... | $\hat{x}_n^{(2)}(M)$ |

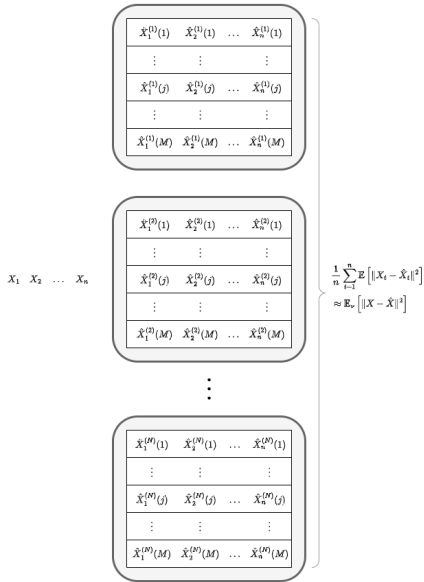
$\vdots$

|                      |                      |     |                      |
|----------------------|----------------------|-----|----------------------|
| $\hat{x}_1^{(N)}(1)$ | $\hat{x}_2^{(N)}(1)$ | ... | $\hat{x}_n^{(N)}(1)$ |
| $\vdots$             | $\vdots$             |     | $\vdots$             |
| $\hat{x}_1^{(N)}(j)$ | $\hat{x}_2^{(N)}(j)$ | ... | $\hat{x}_n^{(N)}(j)$ |
| $\vdots$             | $\vdots$             |     | $\vdots$             |
| $\hat{x}_1^{(N)}(M)$ | $\hat{x}_2^{(N)}(M)$ | ... | $\hat{x}_n^{(N)}(M)$ |

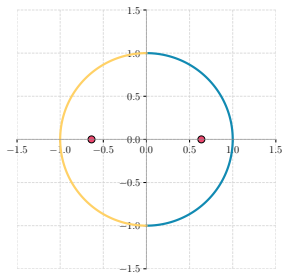
# Operational Encoding vs. Virtual Encoding

$$\begin{array}{c}
 \begin{array}{c} X_1 \quad X_2 \quad \dots \quad X_n \end{array} \\
 \left\{ \begin{array}{|c|} \hline \begin{array}{ccc} \hat{X}_1^{(1)}(1) & \hat{X}_2^{(1)}(1) & \dots & \hat{X}_n^{(1)}(1) \\ \vdots & \vdots & & \vdots \\ \hat{X}_1^{(1)}(j) & \hat{X}_2^{(1)}(j) & \dots & \hat{X}_n^{(1)}(j) \\ \vdots & \vdots & & \vdots \\ \hat{X}_1^{(1)}(M) & \hat{X}_2^{(1)}(M) & \dots & \hat{X}_n^{(1)}(M) \end{array} \\ \hline \end{array} \right\} \begin{array}{l} \frac{1}{n} \sum_{i=1}^n \mathbb{E} \left[ \|X_i - \hat{X}_i^{(1)}(J)\|^2 \right] \\ \approx \mathbb{E}_\mu \left[ \|X - \hat{X}\|^2 \right] \end{array} \\
 \\
 \left\{ \begin{array}{|c|} \hline \begin{array}{ccc} \hat{X}_1^{(2)}(1) & \hat{X}_2^{(2)}(1) & \dots & \hat{X}_n^{(2)}(1) \\ \vdots & \vdots & & \vdots \\ \hat{X}_1^{(2)}(j) & \hat{X}_2^{(2)}(j) & \dots & \hat{X}_n^{(2)}(j) \\ \vdots & \vdots & & \vdots \\ \hat{X}_1^{(2)}(M) & \hat{X}_2^{(2)}(M) & \dots & \hat{X}_n^{(2)}(M) \end{array} \\ \hline \end{array} \right\} \begin{array}{l} \frac{1}{n} \sum_{i=1}^n \mathbb{E} \left[ \|X_i - \hat{X}_i^{(2)}(J)\|^2 \right] \\ \approx \mathbb{E}_\mu \left[ \|X - \hat{X}\|^2 \right] \end{array} \\
 \\
 \vdots \\
 \\
 \left\{ \begin{array}{|c|} \hline \begin{array}{ccc} \hat{X}_1^{(N)}(1) & \hat{X}_2^{(N)}(1) & \dots & \hat{X}_n^{(N)}(1) \\ \vdots & \vdots & & \vdots \\ \hat{X}_1^{(N)}(j) & \hat{X}_2^{(N)}(j) & \dots & \hat{X}_n^{(N)}(j) \\ \vdots & \vdots & & \vdots \\ \hat{X}_1^{(N)}(M) & \hat{X}_2^{(N)}(M) & \dots & \hat{X}_n^{(N)}(M) \end{array} \\ \hline \end{array} \right\} \begin{array}{l} \frac{1}{n} \sum_{i=1}^n \mathbb{E} \left[ \|X_i - \hat{X}_i^{(N)}(J)\|^2 \right] \\ \approx \mathbb{E}_\mu \left[ \|X - \hat{X}\|^2 \right] \end{array}
 \end{array}$$

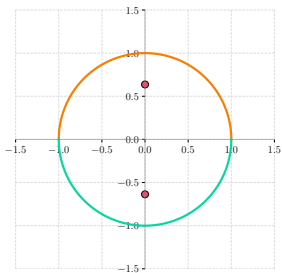
# Operational Encoding vs. Virtual Encoding



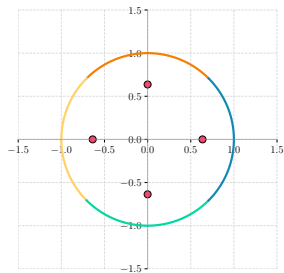
# Operational Encoding vs. Virtual Encoding



$$\mathbb{E}[\|X - \tilde{X}\|^2 | K = 1]$$

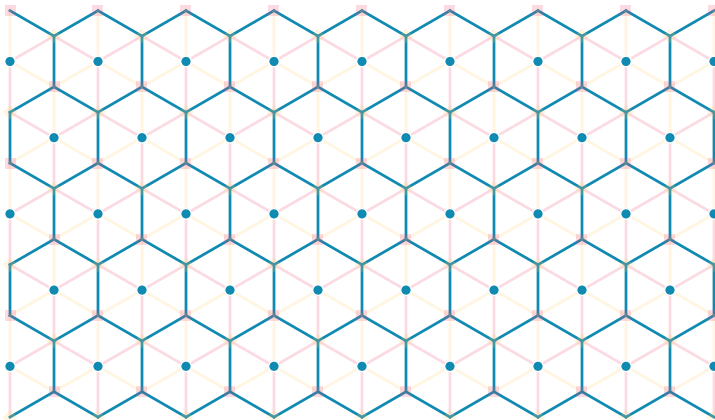


$$\mathbb{E}[\|X - \tilde{X}\|^2 | K = 2]$$

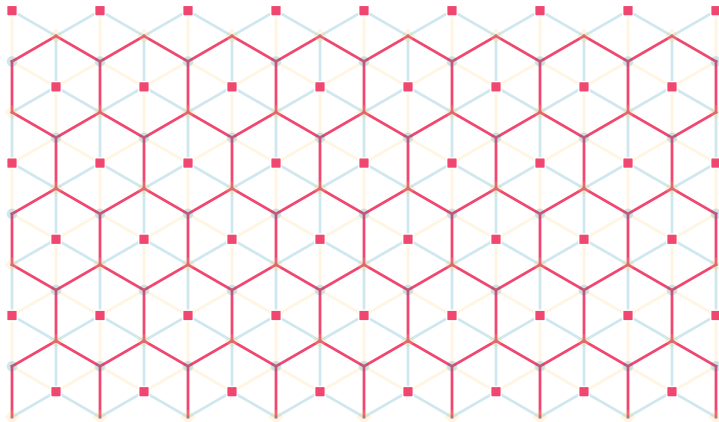


$$W_2^2(p_X, p_{\tilde{X}})$$

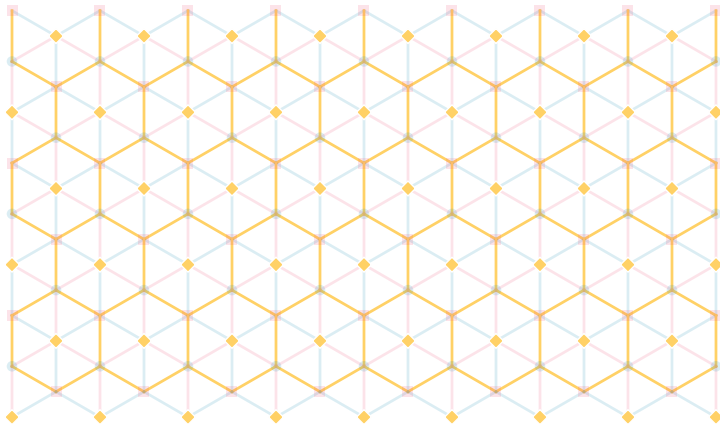
# Operational Encoding vs. Virtual Encoding



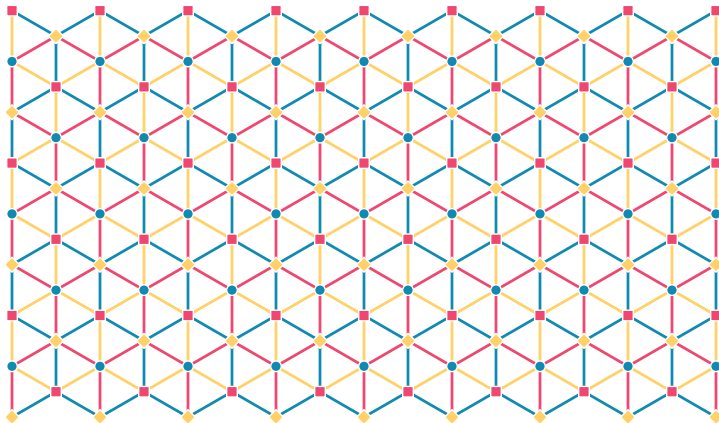
# Operational Encoding vs. Virtual Encoding



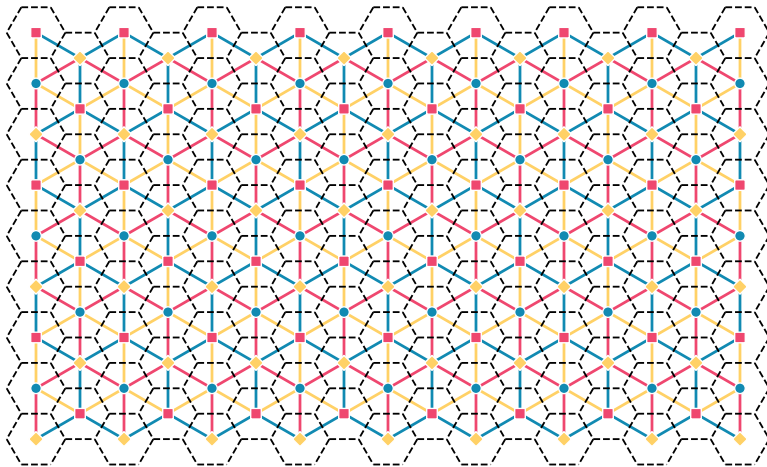
# Operational Encoding vs. Virtual Encoding



# Operational Encoding vs. Virtual Encoding



# Operational Encoding vs. Virtual Encoding



# Symbol-Level vs. Sequence-Level

- Symbol-level perception constraint:  $\frac{1}{n} \sum_{t=1}^n W_2^2(p_X, p_{\hat{X}_t}) \leq P$

$$D^{(\infty)}(R, P) = \inf_{p_{\tilde{X}|X}} \mathbb{E}_{\mu}[\|X - \tilde{X}\|^2]$$

$$\text{s.t. } I(X; \tilde{X}) \leq R \quad \text{Blau-Michaeli}$$

$$W_2^2(p_X, p_{\tilde{X}}) \leq P$$

- Sequence-level perception constraint:  $\frac{1}{n} W_2^2(p_{X^n}, p_{\hat{X}^n}) \leq P$

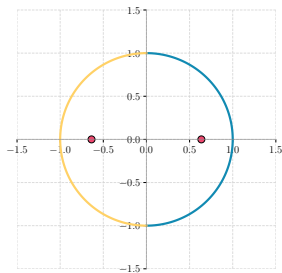
$$D^{(\infty)}(R, C, P) = \inf_{\substack{p_{\tilde{X}} \\ \mu, \nu \in \Pi(p_X, p_{\tilde{X}})}} \mathbb{E}_{\mu}[\|X - \tilde{X}\|^2] + [\sqrt{\mathbb{E}_{\nu}[\|X - \tilde{X}\|^2]} - \sqrt{P}]_+^2$$

$$\text{s.t. } I_{\mu}(X; \tilde{X}) \leq R \quad \text{reduces to Blau-Michaeli}$$

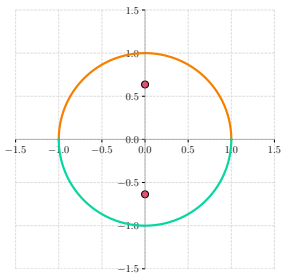
$$(X; \tilde{X}) \leq R + C \quad \text{when } C = \infty!$$

$$\tilde{X} = \mathbb{E}_{\mu}[X|\tilde{X}]$$

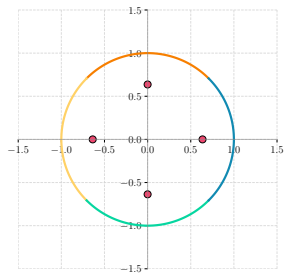
# A New Perspective



$$\mathbb{E}[\|X - \tilde{X}\|^2 | K = 1]$$

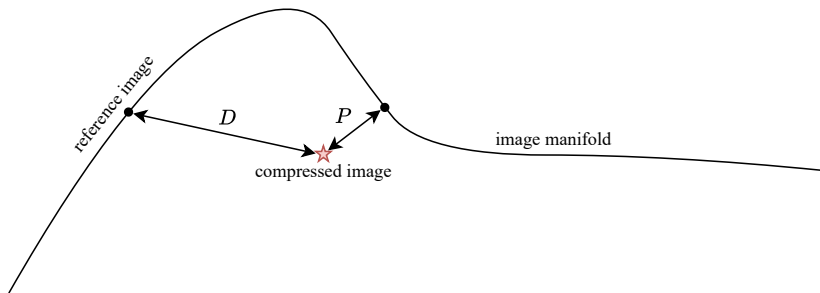


$$\mathbb{E}[\|X - \tilde{X}\|^2 | K = 2]$$

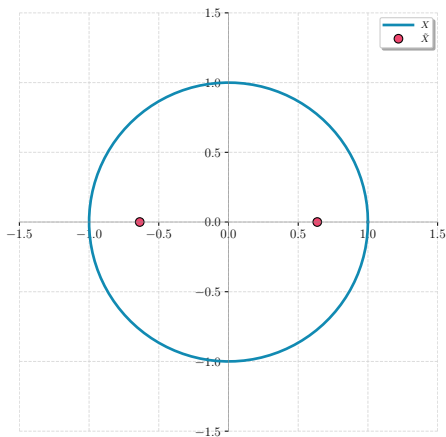


$$W_2^2(p_X, p_{\tilde{X}})$$

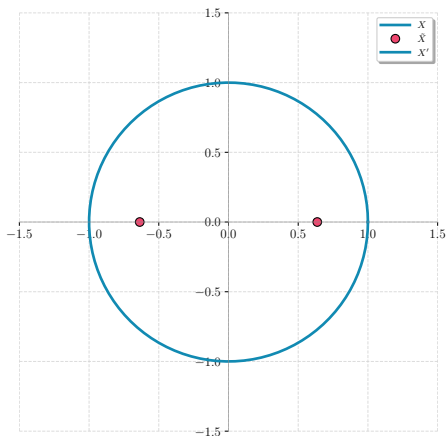
# A New Perspective



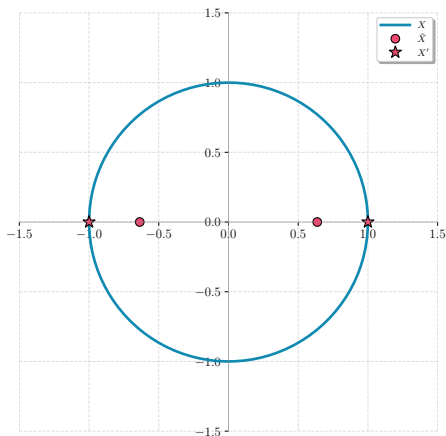
# Quantizing the Unit Circle: No Perception Constraint



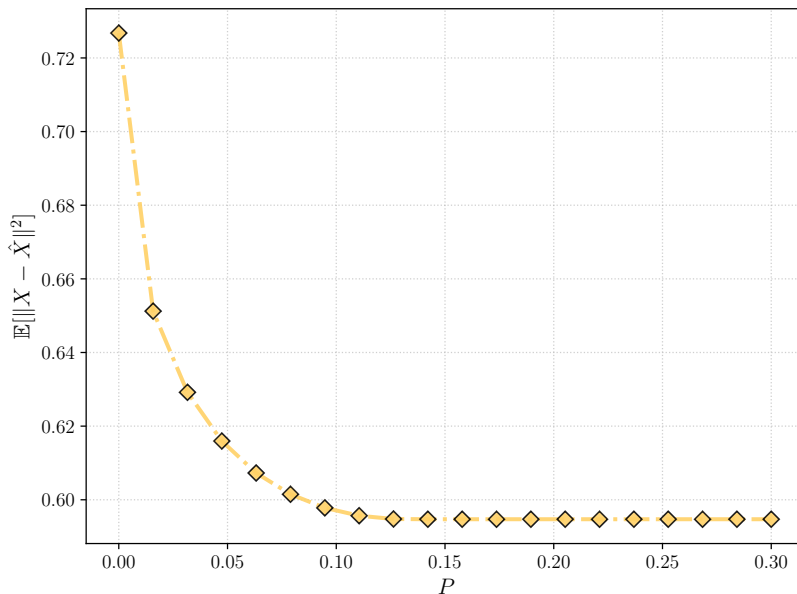
# Quantizing the Unit Circle: Perfect Perception (I)



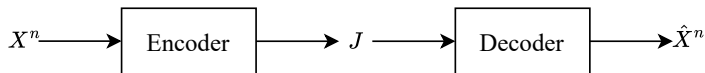
# Quantizing the Unit Circle: Perfect Perception (II)



# Distortion-Perception Tradeoff



# Yet Another Formulation



$$J \in \{1, 2, \dots, M\}$$

Assume  $X^n$  is uniformly distributed over image manifold  $\mathcal{M}_I$ :

$$D^{(\infty)}(R, P) = \lim_{n \rightarrow \infty} \inf_{E, D} \frac{1}{n} \sum_{t=1}^n \mathbb{E}[\|X_t - \hat{X}_t\|^2]$$

s.t.

$$\frac{1}{n} \log M \leq R$$
$$\frac{1}{n} \|\hat{X}^n - \mathcal{M}_I\|^2 \leq P$$

# Coding Theorem

Assume  $\mathcal{M}_I = \mathcal{T}(p_X)$ :

$$\begin{aligned} D^{(\infty)}(R, P) &= \inf_{p_{\hat{X}|X}} \mathbb{E}[\|X - \hat{X}\|^2] \\ \text{s.t. } I(X; \hat{X}) &\leq R \\ W_2^2(p_X, p_{\hat{X}}) &\leq P \end{aligned}$$

Blau-Michaeli again!

# Is the Image Manifold a Typical Set?

- No  
Permuting a typical sequence yields another typical sequence; however, permuting the pixels of a natural image does not guarantee a natural image.
- Yes  
Natural images reside on a low-dimensional manifold in a high-dimensional space. This represents a concentration phenomenon, analogous to the behavior of a typical set.

- Rate-distortion-perception theory: An evolving framework
  - Nature of perception measures
  - Role of common randomness
  - Computation of fundamental limits
  - Principles of system design

- A. Hussein, J. Chen, C. Tian, S. S. Pradhan, "Constructive approaches to perception-aware lossy source coding: Information-theoretic guidelines," arXiv:2604.19515.

## The End

## Questions? Comments?