

09/27/2016
Tuesday

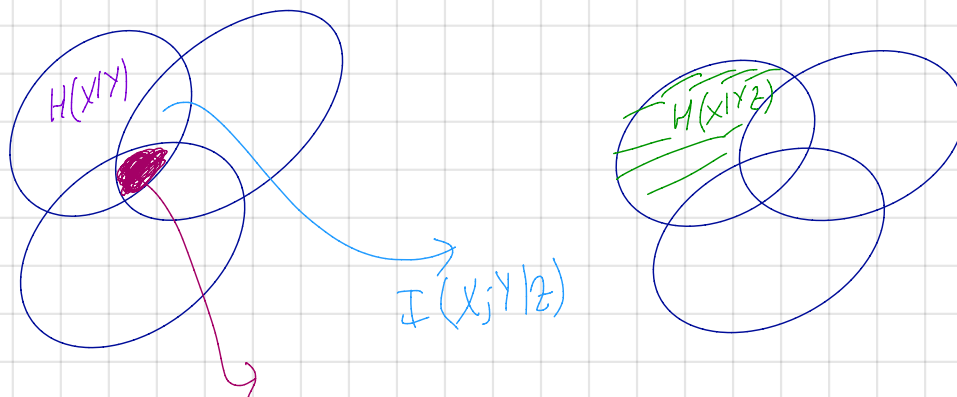
Last time:

$$I(X;Y) \leq I(X,Y|Z)$$

$$\text{if } X \perp\!\!\!\perp Z \leq$$

$$\text{if } X-Y-Z \geq$$

How to remember:



$$\begin{aligned} H(X,Y,Z) - H(X,Y) - H(X,Z) - H(Y,Z) + H(X) + H(Y) + H(Z) \\ &= I(X;Y) - I(X,Y|Z) \\ &= I(Y;Z) - I(Y,Z|X) \\ &= I(X;Z) - I(X,Z|Y) \end{aligned}$$

$I(X;Y)$ is concave in $P_X \rightarrow$ How to show this using $I(X;Y) \geq I(X,Y|Z)$ when $X-Y-Z$?

$I(X;Y)$ is convex in $P_{Y|X}$

$$\begin{aligned} W &\sim \text{Ber}(p) \\ P_{X|W=1} &= P_1 \\ P_{X|W=0} &= P_0 \end{aligned}$$

$$W-X-Y \text{ fix } P_{Y|X}$$

$$I(X;Y) \geq I(X,Y|W)$$

$\downarrow$ mixture P_X (input) $\downarrow$ convex combination of result.

Fano's Inequality

Suppose you can estimate X from Y with high probability ($H(X|Y)$ is small.)

Suppose $X-Y-\hat{X}$ and $P[X \neq \hat{X}] = P_e$

$$H(X|Y) \leq H(P_e) + P_e \log |X|$$

proof

$$\text{Define } E = \mathbb{1}_{X \neq \hat{X}} = \begin{cases} 1 & \text{when } X \neq \hat{X} \\ 0 & \text{when } X = \hat{X} \end{cases}$$

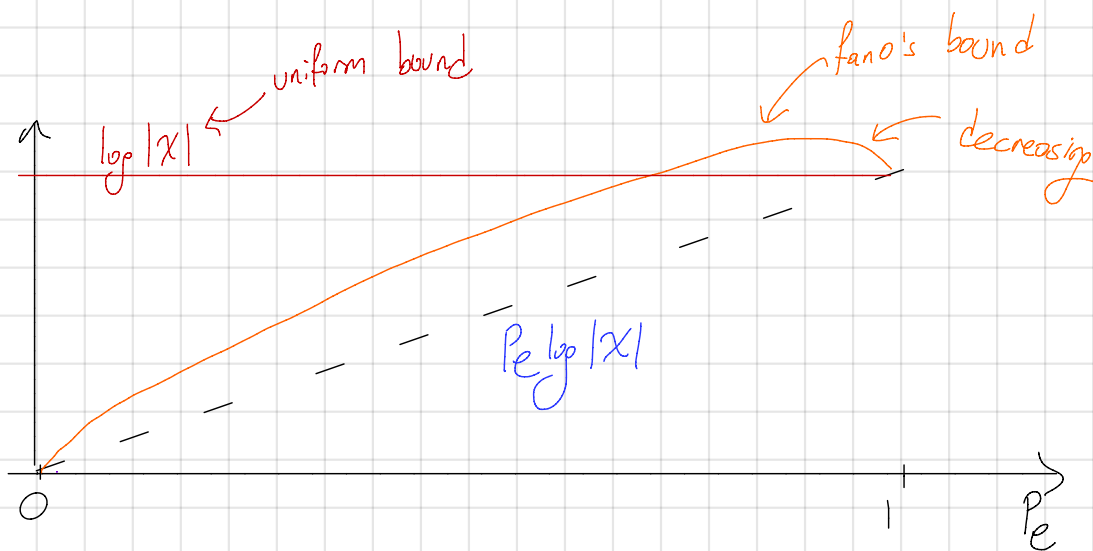
$$\begin{aligned} H(X, E | \hat{X}) &= H(E | \hat{X}) + H(X | E, \hat{X}) \\ &= H(X | \hat{X}) + H(E | X, \hat{X}) \end{aligned}$$

$$H(X | \hat{X}) = H(E | \hat{X}) + H(X | E, \hat{X})$$

$$\leq \underbrace{H(E)}_{H(P_e)} + \underbrace{P[E=0] \cdot 0}_{0} + \underbrace{P[E=1]}_{P_e} \log |X|$$

we are being slightly loose could write $\log(|X|-1)$





Asymptotic Equipartition Property (A.E.P)

L.N. If $X_1, \dots, X_n$ i.i.d $\frac{1}{n} \sum_{i=1}^n X_i \xrightarrow{P} E[X]$

A.E.P $\frac{1}{n} \log \frac{1}{P_{X_1 \dots X_n}(X_1, \dots, X_n)} \xrightarrow{P} H(X)$

in other words: $P_{X_1 \dots X_n}(X_1, \dots, X_n) = 2^{-nH(X)}$
 "almost all events are almost equally surprising."

Types of Convergence:

Convergence in probability: $Z_n \xrightarrow{P} Z : \forall \epsilon > 0 \quad P[|Z_n - Z| > \epsilon] \rightarrow 0 \text{ as } n \rightarrow \infty$

Almost sure convergence: $Z_n \xrightarrow{a.s.} Z : P[\lim_{n \rightarrow \infty} Z_n \neq Z] = 0$

L_2 convergence: $Z_n \xrightarrow{L_2} Z : E[(Z_n - Z)^2] \rightarrow 0$

proof of AEP:

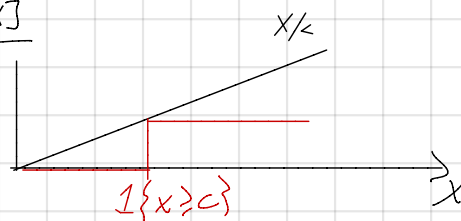
$$\begin{aligned} \frac{1}{n} \log \frac{1}{P_{X_1 \dots X_n}(X_1, \dots, X_n)} &= \frac{1}{n} \log \frac{1}{\prod_{i=1}^n P_X(X_i)} \\ &= \frac{1}{n} \sum_{i=1}^n \log \frac{1}{P_X(X_i)} \\ &\xrightarrow[\text{by L.N.}]{P} E \log \frac{1}{P_X(X)} = H(X) \end{aligned}$$

Markov inequality: If $X \geq 0$ w.p. 1

$$P[X \geq c] \leq \frac{E[X]}{c}$$

proof:

$$E[1_{\{X \geq c\}}] \leq E\left[\frac{X}{c}\right]$$



Chebyshev's Inequality: $P[|X - E[X]| \geq c] \leq \frac{\sigma^2}{c^2}$

proof: Let $Y = (X - E[X])^2$ $Y \geq 0$

$$\begin{aligned} P[Y \geq d] &\leq \frac{E[Y]}{d} = \frac{\sigma_X^2}{d} \\ &= P[(X - E[X])^2 \geq d] \\ &= P[|X - E[X]| \geq \sqrt{d}] \end{aligned}$$

proof of LLN: Assume: X has finite variance. Use Chebyshev's Inequality.

$$\mathbb{P}\left[\left|\frac{1}{n}\sum X_i - \mathbb{E}[X]\right| > \epsilon\right] \leq \frac{\frac{1}{n^2} \cdot n \cdot \text{Var}(X)}{\epsilon^2} \rightarrow 0$$

Typical Sets: Let $X^n = (X_1, \dots, X_n)$ (write $X_i^j = (X_i, \dots, X_j)$)

$$A_\epsilon^{(n)} = \left\{ X^n : \left| \frac{1}{n} \log \frac{1}{P_{X^n}(X^n)} - H(X) \right| < \epsilon \right\} \quad \left(\text{implicitly a function of } P_X \text{ and } P_{X^n} = \prod P_X \right)$$

Thm: 1) $\forall X^n \in A_\epsilon^{(n)} : P_{X^n}(X^n) \in [2^{-n(H(X)+\epsilon)}, 2^{-n(H(X)-\epsilon)}]$

2) $P_{X^n}(A_\epsilon^{(n)}) = \mathbb{P}[X^n \in A_\epsilon^{(n)}] \rightarrow 1 \quad \text{as } n \rightarrow \infty \quad \forall \epsilon > 0$

3) $|A_\epsilon^{(n)}| \leq 2^{n(H(X)+\epsilon)}$

4) $|A_\epsilon^{(n)}| \geq (1-\epsilon) 2^{n(H(X)-\epsilon)}$ for n large enough.

proof:

1.) By defn of $A_\epsilon^{(n)}$

2.) By AEP theorem

3) $1 = \sum_{X^n \in \mathcal{X}^n} P_{X^n}(X^n)$

$$\geq \sum_{X^n \in A_\epsilon^{(n)}} P_{X^n}(X^n)$$

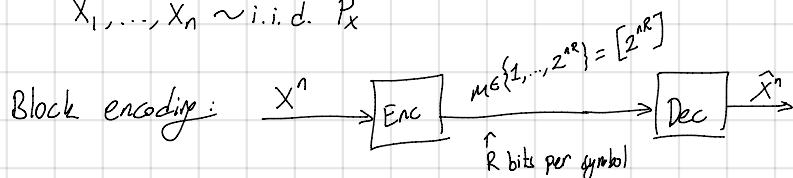
$$\geq \sum_{X^n \in A_\epsilon^{(n)}} 2^{-n(H(X)+\epsilon)} = |A_\epsilon^{(n)}| 2^{-n(H(X)+\epsilon)}$$

4) For n large enough: $\sum_{X^n \in A_\epsilon^{(n)}} P_{X^n}(X^n) = P_{X^n}(A_\epsilon^{(n)}) \geq 1-\epsilon$

$$\sum_{X^n \in A_\epsilon^{(n)}} 2^{-n(H(X)-\epsilon)} = |A_\epsilon^{(n)}| 2^{-n(H(X)-\epsilon)}$$

Lossless Data Compression:

$$X_1, \dots, X_n \sim \text{i.i.d. } P_X$$

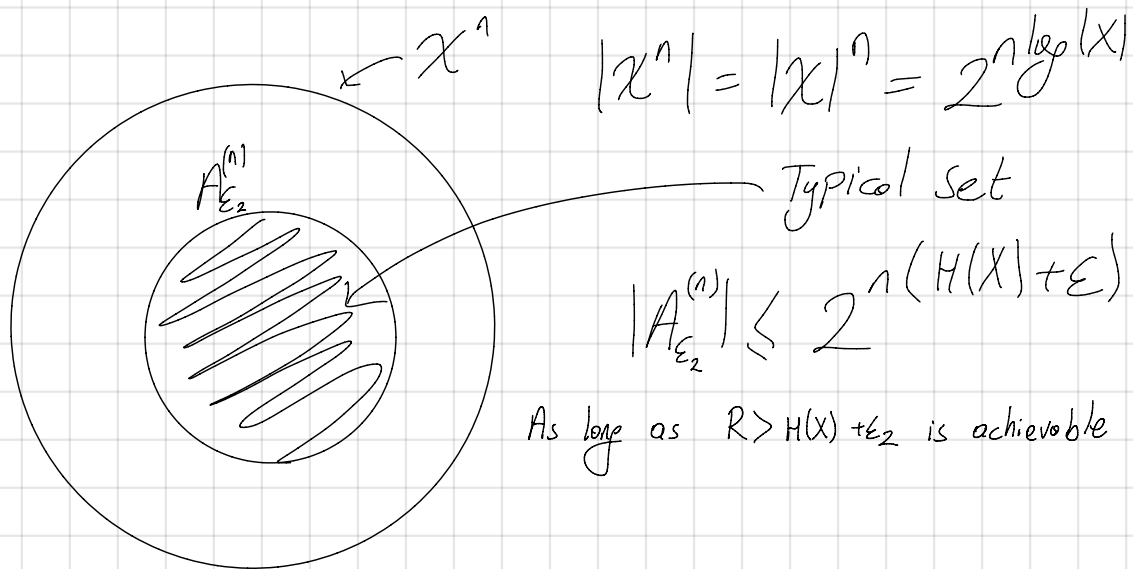


Goal: $\mathbb{P}[X \neq \hat{X}] < \epsilon$

Achievability: R is achievable if $\forall \epsilon > 0$, exists n , encoder, decoder such that $\mathbb{P}[X^n \neq \hat{X}^n] < \epsilon$

Theorem: Minimum achievable rate for lossless compression is $H(X)$





→ You can turn this into 0-error scheme by putting an indicator bit: If typical indicator bit is 1 and we do AEP encoding. If indicator is 0 then we encode the whole sequence x^n with $n \log |X|$ and Average length: $P(A_{\epsilon}^{(n)}) (n(H(X) + \epsilon)) + (1 - P(A_{\epsilon}^{(n)})) (n \log |X|)$

What about a converse result?

Given n , enc., dec. s.t. $P[X^n \neq \hat{X}^n] < \epsilon$ for ϵ arbitrarily small.

$$X^n - M - \hat{X}^n$$

$$nR \geq H(M) \geq I(X^n; M) \stackrel{\text{Data Proc. Ineq.}}{\geq} I(X^n; \hat{X}^n) = H(X^n) - H(X^n | \hat{X}^n) = nH(X) - H(X^n | \hat{X}^n) \geq nH(X) - H(\epsilon) - \epsilon n \log |X|$$

$$\Rightarrow R \geq H(X) - \epsilon \log |X| - \frac{H(\epsilon)}{n} \quad \text{as } \epsilon \rightarrow 0 \quad R \geq \underline{H(X)}$$