

Overview

- I. Intro + Motivation
- II. Entropy + noiseless coding
- III. Conditional entropy + mutual info.
- IV. Noisy channel coding examples

I. Intro + motivation

- Classical Inf. theory is about fundamental limits of inf. processing tasks - e.g., sending, storing or ~~transmitting~~ learning info. - in the asymptotic setting.

↳ many uses/copies of same resource, some small error prob. going to zero.

- Story becomes more complicated in the quantum setting, but the most basic theorem statements carry over in some way from the classical cases. In particular, we can summarize some of these in the following table.

	Application	Classical	Quantum
this lecture	Data compr.	Shannon entropy	von Neumann entropy
	Channel coding	Mutual info.	Quantum mutual info. (for cl. info. over q. channel) or coherent info. (for q. info. over q. channel)
	Hypothesis testing	Relative entropy	Quantum relative entropy

II. Entropy + noiseless source coding

- entropy is a measure of uncertainty; or the expected amount of surprise upon learning the value of a random variable.

Defn: Let $p: \mathcal{X} \rightarrow [0, 1]$ be a distr. over some finite alphabet $\mathcal{X}$, let $X \sim p$ be a random variable. The Shannon entropy of X is

$$H(X) \equiv H(p) = - \sum_{x \in \mathcal{X}} p(x) \log_2 p(x). \quad 0 \log 0 = 0.$$

- Some properties:

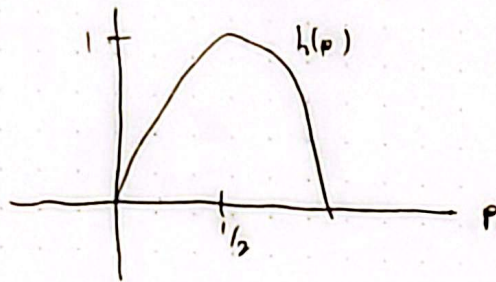
$$1. \text{ (Bounds)} \quad 0 \leq H(X) \leq \log_2 d \quad \text{for } d = |\mathcal{X}|$$

$\downarrow$ $\downarrow$
 deterministic unif.

$$2. \text{ (Concavity)} \quad \forall \quad 0 \leq \lambda \leq 1 \text{ we have}$$

$$H(\lambda p_1 + (1-\lambda)p_2) \geq \lambda H(p_1) + (1-\lambda)H(p_2).$$

- Can plot binary entropy $h(p) \equiv -p \log p - (1-p) \log(1-p)$.



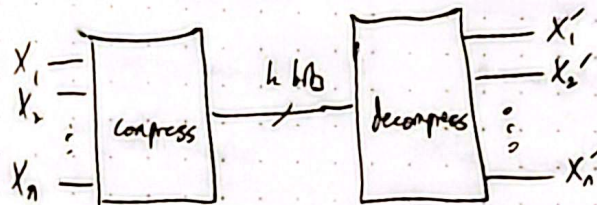
- For multiple random variables $p: \mathcal{X} \times \mathcal{Y} \rightarrow [0,1]$, $(X,Y) \sim p$, nothing changes:

$$H(X,Y) = - \sum_{x \in \mathcal{X}} \sum_{y \in \mathcal{Y}} p(x,y) \log p(x,y).$$

- Note: If p is a product distr. $H(X,Y) = H(X) + H(Y)$.

- Noiseless source coding setup:

$$\begin{aligned} X &\sim p \\ X_i &\sim p \text{ i.i.d.} \end{aligned}$$



- We say a rate R is achievable if the message can be recovered w/ vanishing prob. of error as $n \rightarrow \infty$ when $k = n(R \pm \delta)$ for any $\delta > 0$.

Theorem (noiseless source coding): The optimal rate of compression $R^* := \inf \{R: R \text{ is achievable}\}$ is given by $H(X)$, i.e., $R^* = H(X)$.

Pf sketch: Consider the distr. $p^n: \mathcal{X}^n \rightarrow [0,1]$ over n -letter words, & suppose w.l.o.g. $\mathcal{X} = \{1, 2, \dots, d\}$. Since letters are i.i.d.,

$$p^n(x_1, x_2, \dots, x_n) = \prod_{i=1}^n p(x_i).$$

By LLN, any letter $l \in [d]$ occurs as $n p(l)$ times, ^{typically} as $n \rightarrow \infty$.

$\therefore$ by Stirling, # of typical strings is $\frac{n!}{\prod_{l=1}^d (n p(l))!} \approx 2^{n H(X)}$.

Why?

$$\approx \underbrace{1 \dots 1}_{n p(1)} \underbrace{2 \dots 2}_{n p(2)} \dots \underbrace{d \dots d}_{n p(d)}$$

and all unique strings formed by permuting elements are the typical strings.

III. Conditional Entropy & Mutual Info.

Def. For random variables $(X, Y) \sim p$ the conditional entropy of Y given X is

$$\begin{aligned} H(Y|X) &= \sum_{x \in \mathcal{X}} p(x) H(Y|X=x) \\ &= - \sum_{x \in \mathcal{X}} p(x) \sum_{y \in \mathcal{Y}} p(y|x) \log p(y|x) \\ &= H(X, Y) - H(X). \end{aligned}$$

→ $H(Y|X)$ is the uncertainty of Y once we know X . Think of the case where X & Y are perfectly correlated, for example. Then it is straightforward to see that $H(Y|X) = 0$.

Mutual Info. The mutual information between X and Y is

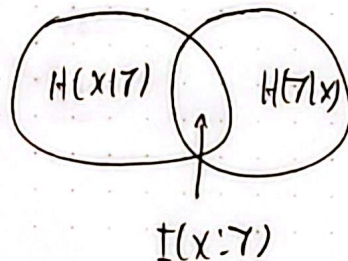
$$\begin{aligned} I(X:Y) &= H(X) + H(Y) - H(X, Y) \\ &= H(X) - H(X|Y) \\ &= H(Y) - H(Y|X) \quad \text{note the symmetry in } X \text{ & } Y \end{aligned}$$

• Remark: mutual info. is nonnegative, both classically and quantumly.

$$\begin{aligned} I(X:Y) \geq 0 &\Leftrightarrow H(X) + H(Y) \geq H(X, Y) \\ &\Leftrightarrow H(Y|X) \leq H(Y). \end{aligned}$$

→ we'll see why later.

• Picture:



• Information channel capacity.

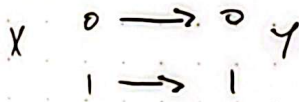
• Let N be a noisy channel, $X \rightarrow \boxed{N} \rightarrow Y$, which specifies $p(y|x)$.

• The "information channel capacity" of N is defined as

$$C = \max_{p(x)} I(X:Y).$$

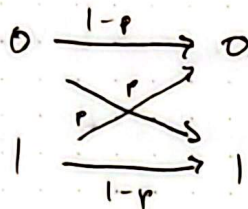
Examples

Noiseless



$$\begin{aligned} \max_X I(X; Y) &= \max_X [H(X) - H(X|Y)] \\ &= \boxed{1} \end{aligned}$$

Binary symmetric channel

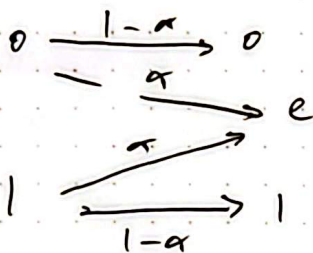


$$\begin{aligned} I(X; Y) &\leq H(Y) - \sum_{x \in \mathcal{X}} p(x) H(Y|X=x) \\ &\leq H(Y) - H(p) \\ &\leq 1 - h(p) \end{aligned}$$

If we take X uniform, equality attained

$$\Rightarrow I(X; Y) = \boxed{1 - h(p)}$$

Erasures

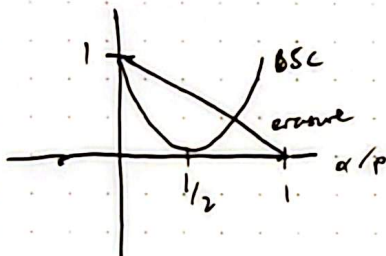


$$\begin{aligned} C &= \max_X H(Y) - H(Y|X) \\ &= \max_{p(x)} H(Y) - H(\alpha) \end{aligned}$$

Let E denote r.v. indicating if erasure occurred.

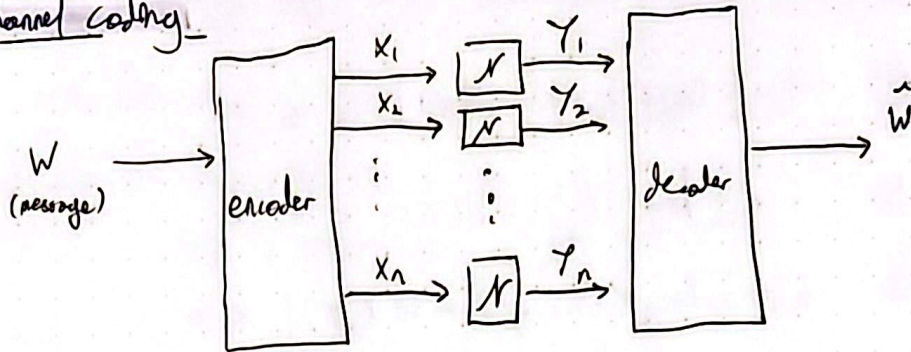
$$\begin{aligned} H(Y) &= H(E) + H(Y|E) \\ &= H(\alpha) + (1-\alpha) H(p), \quad p \equiv P_i[X=1] \end{aligned}$$

$$\begin{aligned} \therefore C &= \max_p (1-\alpha) h(p) + H(\alpha) - h(\alpha) \\ &= \boxed{1 - \alpha} \end{aligned}$$



IV. Noisy channel coding.

• Setup is



• We say this is an (M, n) code if the message W can take on M possible values. We say the code has rate R if $M = 2^{nR}$, and R is achievable if prob. error in decoding i.e., $P[\hat{W} \neq W]$, goes to zero as $n \rightarrow \infty$.

Theorem (noisy channel coding): The optimal rate $R^* := \sup \{ R : R \text{ is achievable} \}$, known as the channel capacity of N , is equal to the information channel capacity of N . I.e.,

$$R^* = C = \max_x I(X; Y).$$

($\Rightarrow$ achievability only)

Pf. Generate a $(2^{nR}, n)$ code at random according to specified $p(x)$. The code is

$$C = \begin{bmatrix} x_1(1) & \dots & x_n(1) \\ \vdots & & \vdots \\ x_1(2^{nR}) & \dots & x_n(2^{nR}) \end{bmatrix}.$$

Each entry is generated i.i.d according to $p(x)$. Thus, prob. that we generate a particular code is $P[C] = \prod_{w=1}^{2^{nR}} \prod_{i=1}^n p(x_i(w))$.

Receiver gets Y^n according to

$$p(Y^n | X^n(w)) = \prod_{i=1}^n p(y_i | x_i(w)).$$

Strategy: use jointly-typical decoding:

$\rightarrow$ declare " $\hat{W}$ sent" if $(X^n(\hat{W}), Y^n)$ is jointly typical, & there is no other index W' s.t. $(X^n(W'), Y^n)$ is jointly typical.

but first: what is jointly typical?

$\rightarrow$

→ "typical set", "jointly typical set"

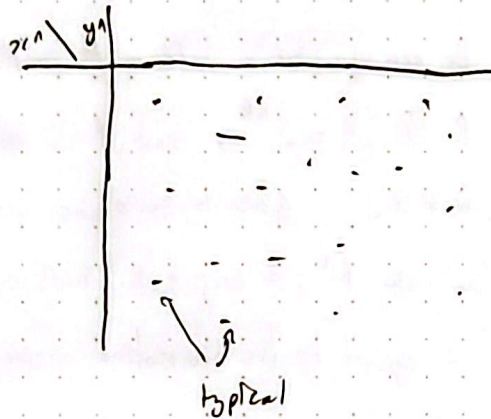
$$\text{Set } A_\epsilon^{(n)} := \{ (x^n, y^n) \in \mathcal{X}^n \times \mathcal{Y}^n \text{ s.t. } \textcircled{1}, \textcircled{2}, \textcircled{3} \text{ hold} \}$$

$$\text{where } \textcircled{1} \quad \left| -\frac{1}{n} \log p^n(x^n) - H(X) \right| < \epsilon$$

$$\textcircled{2} \quad \left| -\frac{1}{n} \log p^n(y^n) - H(Y) \right| < \epsilon$$

$$\textcircled{3} \quad \left| -\frac{1}{n} \log p^n(x^n, y^n) - H(X, Y) \right| < \epsilon$$

Intuition:



Most pairs (x^n, y^n) are not typical. But if $(X^n, Y^n) \sim p$, then typicality is overwhelmingly likely.

Next we will bound the average prob. of error, averaged over randomness due to

(i) choice of code C , (ii) choice of message w , & (iii) ~~choice of~~ randomness due to N .

→ if error is low, we can convert ~~(i) & (ii)~~ (iii) (average case statement)

to worst case statements via probabilistic existence argument discarding some messages.

→ if (i) can be converted to statement of existence of some fixed code C that succeeds, via probabilistic existence.

$$\begin{aligned} P_e[\text{error}] &= \sum_C P_e[C] \cdot (\text{avg. error w/ code } C) \\ &= \sum_C P_e[C] \frac{1}{2^{nR}} \sum_{w=1}^{2^{nR}} \lambda_w(C) \\ &= \frac{1}{2^{nR}} \sum_{w=1}^{2^{nR}} \sum_C P_e[C] \lambda_w(C) \quad \text{error on input } w \text{ w/ code } C \\ &= P_e[\text{error} | W=1] \quad \text{no } w\text{-dependence, by symmetry of } C \text{ randomness.} \end{aligned}$$

$$\therefore P_e[\text{error} | W=1] \leq P_e[(X^n(1), Y^n) \notin A_\epsilon^{(n)}] + \sum_{i=2}^{2^{nR}} P_e[(X^n(i), Y^n) \in A_\epsilon^{(n)}]$$

(A)

(B)

for ①, $\Pr[(X^n, Y^n) \in A_\epsilon^{(n)}] \rightarrow 1$ as $n \rightarrow \infty$ is just the LLN.

for ②, note that since $X^n(2)$ indep. of $X^n(1)$, $X^n(2), Y^n$ are indep. as well, & so on for $X^n(3), X^n(4)$, etc.

$$\therefore \Pr[(X^n(2), Y^n) \in A_\epsilon^{(n)}] = \sum_{(x^n, y^n) \in A_\epsilon^{(n)}} p(x^n) p(y^n)$$

joint typically $\rightarrow$

$$\leq \sum_{(x^n, y^n) \in A_\epsilon^{(n)}} 2^{-n[H(X) - \epsilon]} 2^{-n[H(Y) - \epsilon]}$$

$$\leq 2^{n[H(X, Y) + \epsilon]} 2^{-n[H(X) - \epsilon]} 2^{-n[H(Y) - \epsilon]}$$

$$= 2^{-n[I(X:Y) - 3\epsilon]}$$

so overall ② $\leq 2^{nR} \cdot 2^{-n[I(X:Y) - 3\epsilon]}$

$$\approx 2^{-n[I(X:Y) - R]}$$

which is small so long as $R < I(X:Y)$. Take sup over $p(x)$. This proves achievability.