

Conditional Entropy:

$$\begin{aligned}
 H(X|Y) &= E_Y H(X) = E_Y E_{X|Y} \log \frac{1}{p_{X|Y}(X|Y)} \\
 &= E \log \frac{1}{p_{X|Y}(X|Y)} \leftarrow \text{Towering Property}
 \end{aligned}$$

Average conditional entropy, averaged over Y .

Joint Entropy: (same as entropy, but random variable is a vector)

$$H(X, Y) = E \log \frac{1}{p_{X,Y}(X, Y)}$$

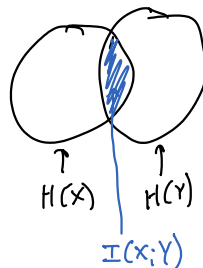
Chain Rule:

$$\begin{aligned}
 H(X, Y) &= E \log \frac{1}{p(X, Y)} \\
 &= E \log \frac{1}{p(X) p(Y|X)} \\
 &= E \log \frac{1}{p(X)} + E \log \frac{1}{p(Y|X)} \\
 &= H(X) + H(Y|X)
 \end{aligned}$$

$$\Rightarrow H(X, Y) \geq H(X)$$

Property: $H(X, Y) \leq H(X) + H(Y)$

Proof momentarily.



$$\begin{aligned}
 \Rightarrow H(Y) &\geq H(Y|X) \text{ by the chain.} \\
 H(X) &\geq H(X|Y)
 \end{aligned}$$

*only true because this
an average quantity*

Mutual Information:

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

Property: $0 \leq I(X; Y) \leq \min \{ H(X), H(Y) \}$

Property: $0 \leq I(X;Y) \leq \min \{H(X), H(Y)\}$

Equality only if X is a function of Y (i.e. $H(X|Y)=0$) or Y is a function of X (i.e. $H(Y|X)=0$)

Relative entropy:
Kullback-Leibler divergence
K-L divergence
Divergence

$$d_{KL}(p, q) = E_p \log \frac{p(x)}{q(x)} \\ = \sum_x p(x) \log \frac{p(x)}{q(x)}$$

This is another metric of distance, but not a true distance.
- Not symmetric
- No triangle inequality.

Important to note that $d_{KL}(p, q)$ is discontinuous w.r.t. $\|p - q\|_1$.
If $p \not\sim q$ then $d_{KL}(p, q) = \infty$, even though $\|p - q\|_1$ may be arbitrarily small.

K-L divergence plays many important roles.

- Fundamental limits of detection theory
- Penalty for encoding $\sim q$ when true distribution is $\sim p$.
- Pinsker inequality: $\|p - q\|_1 \leq \sqrt{2 d_{KL}(p, q)}$

Property: $d_{KL}(p, q) \geq 0$ with equality iff $p = q$

Proof: $d_{KL}(p, q) = E_p \log \frac{p(x)}{q(x)} \\ = -E_p \log \frac{q(x)}{p(x)} \\ \geq -\log E_p \frac{q(x)}{p(x)} \\ = -\log \sum_x p(x) \frac{q(x)}{p(x)} \\ = -\log \sum_{x \in \mathcal{X}_p} q(x) \\ \geq -\log 1 \\ = 0$

Jensen's Ineq.

$-\log$ is strictly convex

Equality iff $\frac{q(x)}{p(x)}$ is deterministic.

$$\Rightarrow p(x) = \alpha q(x) \quad \forall x \\ \alpha = 1$$

$\Rightarrow$ Equality iff $p = q$

Mutual Information as Relative Entropy;

$$\begin{aligned} I(X; Y) &= H(X) + H(Y) - H(X, Y) \\ &= E \left(\log \frac{1}{p(x)} + \log \frac{1}{p(y)} - \log \frac{1}{p(x, y)} \right) \\ &= E_{p(x, y)} \log \frac{p(x, y)}{p(x)p(y)} \\ &= d_{KL}(p(x, y), p(x)p(y)) \end{aligned}$$

Measures how far the distribution is from being independent

Consequences: $I(X; Y) \geq 0$ with equality iff $X \perp\!\!\!\perp Y$
↑
independent

Conditional Mutual Information:

$$I(X; Y | Z) = E_Z \frac{I(X; Y)}{p(x, y|z)}$$

Properties (proven by simple algebra)

1.) $I(X; Y | Z) = H(X|Z) + H(Y|Z) - H(X, Y|Z)$

2.) Chain rule: $I(X; Y, Z) = I(X; Z) + I(X; Y | Z)$

Same definition
for M.I.
Interpret as a vector

Data Processing Inequality:

If $X - Y - Z$ form a Markov chain.

(i.e. X and Z are conditionally independent given $Y \Leftrightarrow p(x, y, z) = p(x, y)p(z|y)$
 $= p(y)p(x|y)p(z|y)$)

then $I(X; Y) \geq I(X; Z)$

Proof: Notice that $I(X; Z|Y) = 0 \leftarrow$ Equivalent to Markov chain condition.

$$\begin{aligned} \Rightarrow I(X; Y) &= I(X; Y) + I(X; Z|Y) && \text{Markovity} \\ &= I(X; YZ) && \text{Chain rule} \\ &= I(X; Z) + I(X; Y|Z) && \text{Chain rule} \\ &\geq I(X; Z) && \text{non-negativity.} \end{aligned}$$

$I(X; Y) \geq I(X; f(Y)) \quad \text{because } X - Y - f(Y)$

Notice: $I(X; Y) \geq I(X; Y|Z)$ from same proof
Not true in general.

Also notice that if $X \perp\!\!\!\perp Z$ then $I(X; Y) \leq I(X; Y|Z)$

Fano's inequality:

Entropy continuously goes to zero as uncertainty becomes less and less.
Exact bound depends on cardinality, which must be accounted for in proof.

Let p_e be the probability of error in estimating X from Y .

$$\begin{aligned} H(X|Y) &\leq h(p_e) + p_e \log(|\mathcal{X}| - 1) \leftarrow \text{Maximized by spreading probability uniformly on remaining support} \\ &\leq 1 + p_e \log |\mathcal{X}| \end{aligned}$$

Application: Lower bound on rate for lossless compression:

Thm Minimum data needed for compressing an iid signal $\sim P_X$ is $H(X)$.

Lower bound: Denote $X_1, X_2, \dots, X_n$ as X^n
Assume encoding $X^n \rightarrow M \rightarrow \hat{X}^n$ where $M \in [2^{nR}] \triangleq \{1, 2, \dots, 2^{nR}\}$
and $P(X^n \neq \hat{X}^n) < \varepsilon$ for arbitrarily small $\varepsilon > 0$.

$$\begin{aligned} nR &\geq H(M) && \text{Entropy upper bound} \\ &\geq I(X^n; M) && \text{Def., non-negativity} \\ &\geq I(X^n; \hat{X}^n) && \text{Data Processing Ineq.} \\ &= H(X^n) - H(X^n|\hat{X}^n) && \text{Def.} \end{aligned}$$

$$\geq I(x^n; \hat{x}^n) \quad \text{Data Processing Ineq.}$$

$$= H(x^n) - H(x^n | \hat{x}^n) \quad \text{Def.}$$

$$= n H(X) - H(x^n | \hat{x}^n) \quad \text{Prob. set - IID}$$

$$\geq n H(X) - h(\epsilon) - \epsilon \log |X^n| \quad \text{Fano's Ineq.}$$

$$= n \left(H(X) - \epsilon \log |X| - \frac{h(\epsilon)}{n} \right)$$

$$\Rightarrow R \geq H(X) - \epsilon \log |X| - \frac{h(\epsilon)}{n} \rightarrow 0$$

Quantum Information and Entropy: Ch. 11

von Neumann Entropy: $H(A) = -\text{Tr}(\rho^A \log \rho^A)$

Notice that if ρ is diagonal this is Shannon entropy.

Canonical Decomposition: $\rho^A = \sum_x p(x) |\varphi_x\rangle \langle \varphi_x|^A$

orthonormal basis
Shannon entropy of $p(x) =$ von Neumann entropy.

For any ensemble: This implies an ensemble preparation $\{(p(x), |\varphi_x\rangle)\}$
If $|\varphi_x\rangle$ are orthogonal basis: $H_S = H_{VN}$

If $|\varphi_x\rangle$ are not orthogonal: $H_S > H_{VN}$

deterministic

Proof later: Think of the ensemble as a function

$$H(x) \geq H(f(x))$$

Example of non-orthogonal ensemble:

$$\left\{ \left(\frac{1}{4}, |0\rangle \right), \left(\frac{1}{4}, |1\rangle \right), \left(\frac{1}{4}, |+\rangle \right), \left(\frac{1}{4}, |-\rangle \right) \right\}$$

$$\rho = \pi \quad H(\rho) = 1 \text{ bit}$$

$$H(x) = 2 \text{ bits}$$

(von Neumann) Entropy gives the fundamental limit for (Schumacher) compression

... then: $1 \setminus H(\rho) > 0$... the ensemble: if ρ is a pure state

Properties :

- 1.) $H(\rho) \geq 0$ with equality iff ρ is a pure state
- 2.) $H(\rho) \leq \log D$
- 3.) Concavity: $H(\rho) \geq \sum_x p(x) H(\rho_x)$ *Proof later*
- 4.) Unitary invariant: $H(u \rho u^\dagger) = H(\rho)$
 $\uparrow$
isometric

- Think of this like a classical one-to-one mapping.

von Neumann entropy and measurement

$$H(\rho) = \min_{\substack{\Lambda_y \\ \text{rank } 1}} H_S(\text{Tr}(\Lambda_y \rho))$$

↑
Shannon entropy
↑
Distribution induced by measurement

How do we define conditional entropy?

No natural notion of conditional density.

↑
Unless conditioning on classical information

This is the first major anomaly of quantum information theory!

$$H(B|A) = ? \quad \text{given } \rho^{AB}$$