

PRINCIPLE OF MAXIMUM ENTROPY

i) Surprise (uncertainty): how surprised am I by what happened?

$$H = -\ln(p) \Rightarrow H(s|I) = -\ln(P(s|I))$$

Smaller $p \Rightarrow$ more surprise.

event at probability 0.5 happens $\Rightarrow$ some surprise
event at probability 5×10^{-10} happens $\Rightarrow$ ^{much} ~~lot~~ surprise!

Why $\ln$? It makes surprise additive but also other reasons tied to additivity:

$$\begin{aligned} H(s|d|I) &= -\ln(P(s|d|I)) = -\ln(P(s|d,I)) - \ln(P(d|I)) \\ &= H(s|d,I) + H(d|I). \end{aligned}$$

Entropy: expected surprise $= -\sum_s P(s|I) \ln(P(s|I))$.

More entropy $\Rightarrow$ more likely to be surprised.

ii) What is 1 bit?

Coin toss: $\begin{cases} \frac{1}{2} & H \\ \frac{1}{2} & + \end{cases}$ Say head occurs then surprise $= -\ln\left(\frac{1}{2}\right) = \ln(2)$ nits

$$\begin{aligned} \frac{1}{\ln(2)} \frac{\text{bits}}{\text{nits}} &\Rightarrow \ln(N) \text{ nits} = \frac{\ln(N)}{\ln(2)} \text{ bits} = \log_2(N) = 1 \text{ bit.} \\ &= \log_2(N) \text{ bits} \end{aligned}$$

Expected surprise before learning outcome.

$$S(\{p_i\}) = - \sum_i p_i \ln(p_i) = \underbrace{-\frac{1}{2} \ln\left(\frac{1}{2}\right)}_{H} - \underbrace{\frac{1}{2} \ln\left(\frac{1}{2}\right)}_{T}$$

$$= \ln(2) = 1 \text{ bit}$$

$H \rightarrow \frac{1}{2} \text{ prob, 1 bit}$
 $T \rightarrow \frac{1}{2} \text{ prob, 1 bit}$

Cross entropy: ~~both~~ How much does Alice expect Bob to be surprised.

$$S(A||B) = - \sum_s \underbrace{P(s|I_A)}_{\text{how likely Alice thinks s is}} \ln \left(\underbrace{P(s|I_B)}_{\text{how much she thinks Bob is surprised by this s.}} \right)$$

Two cases: Alice ~~knows~~ Bob gives Alice a fair win.



i) Bob gives Alice a fair win. Alice tosses coin and finds out "H". What is ^{Alice}Bob's expected surprise of Bob?

$$S(P_A||P_B) = \underbrace{-\frac{1}{2} \ln\left(\frac{1}{2}\right)}_{P_{A,H}} - \underbrace{0 \ln\left(\frac{1}{2}\right)}_{P_{A,T}} = 1 \text{ bit.}$$

$P_{B,H}$ $P_{B,T}$

ii) Bob gives Alice a biased win (both sides H).

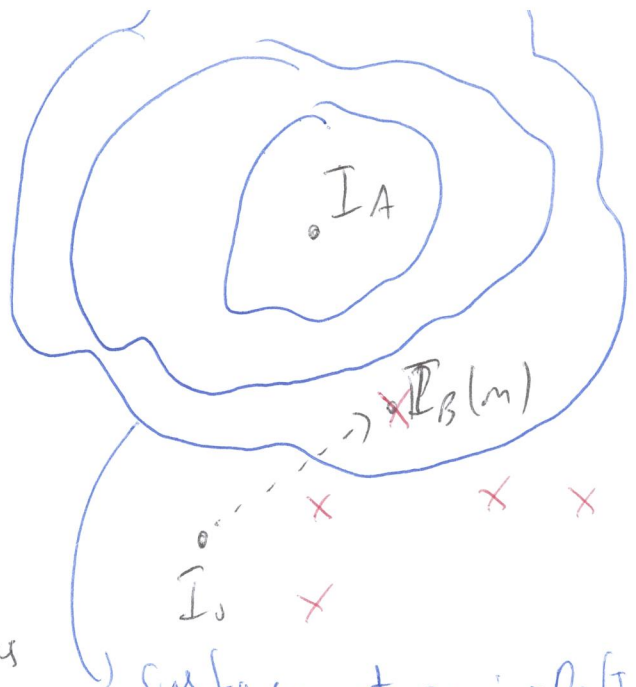
At Coin fairy replaces biased coin with fair one.

Alice tosses coin. What is Alice expected surprise at Bob?

$$S(P_A||P_B) = -\frac{1}{2} \ln(1) - \frac{1}{2} \ln(0) = \infty \text{ bits!}$$

iii) COMMUNICATION:

a) Encoding a message



Alice choose $m \in M$ that maximises the achieved information gain:

$$D_S(I_A, I_B(m), I_0) = D_S(I_A, I_0) - D_S(I_A, I_B(m)).$$

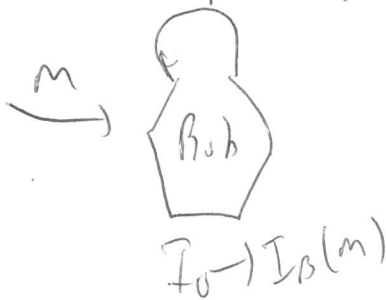
Not a metric space! Still surface can be drawn

$$M = \arg \max_{m \in M} [D_S(I_A, I_0) - D_S(I_A, I_B(m))]$$

$$= \arg \min_{m \in M} (D_S(I_A, I_B(m))).$$

As $D_S(I_A, I_0)$ does not depend on I_B , only second term matters here.

b) Decoding a message: No I_A ! There is no longer an ideal probability distribution



Bob assumes I_B is correct.

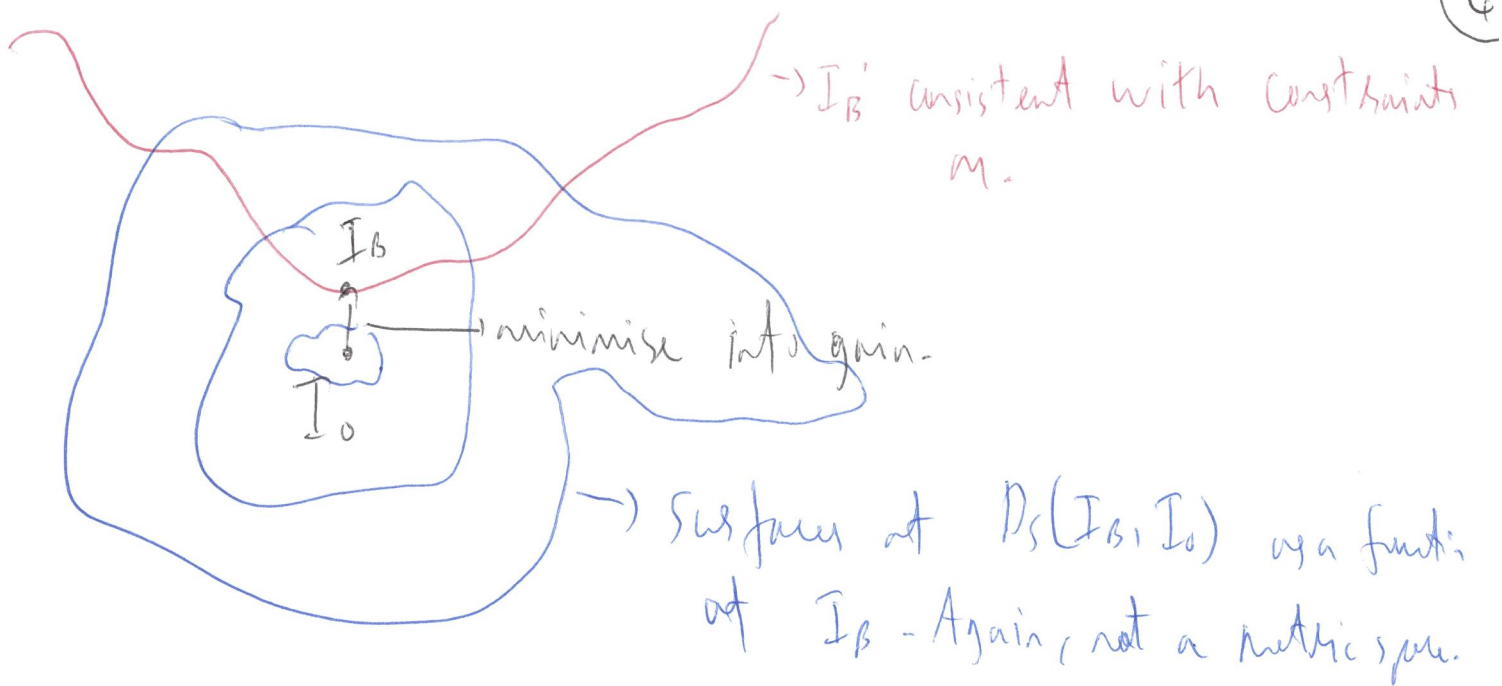
Minimise information gain from I_0 consistent with m — Max Entropy principle

$$I_B = \arg \min_{I_B'} (D_S(I_B', I_0) + \lambda(m))$$

Minimise spurious information

Constraints from m as Lagrange multiplier terms

Max Ent principle

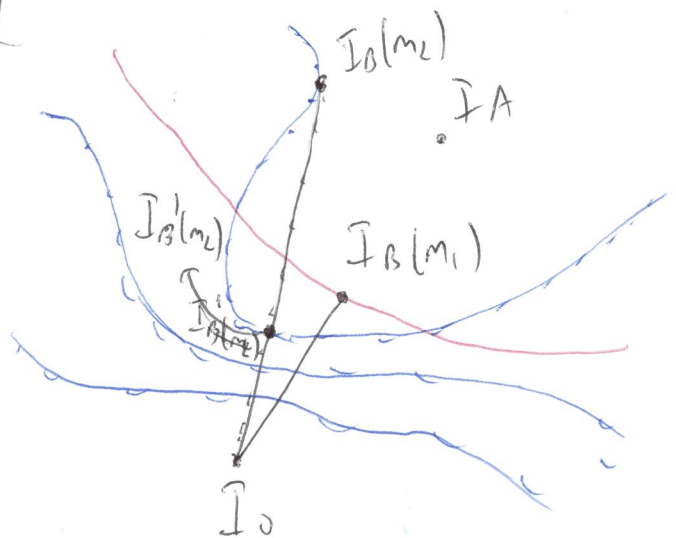
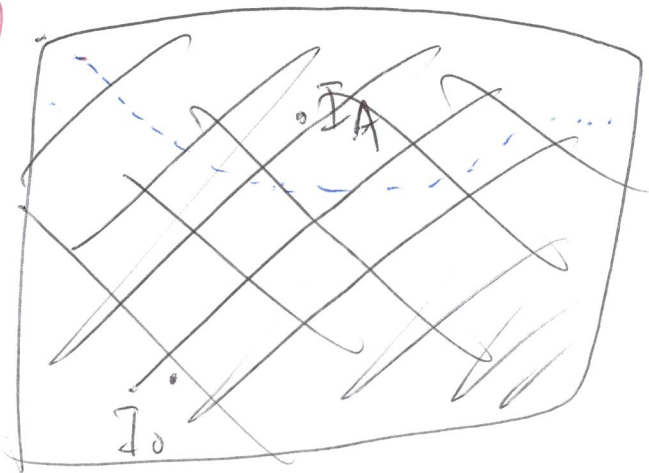


c) Sending a message to a Cognizant agent:

$$m = \underset{m' \in \mu}{\operatorname{argmin}} \left[D_S(I_A, \underset{I_B'}{\operatorname{argmin}} [D_S(I_B'(m'), I_0)]) \right]$$

(i) (iii)

minimization (i) must be done by Alice first before (ii)



While $I_0(m_2)$ is closer to I_A than $I_0(m_1)$, Bob would choose $I_B'(m_2)$ based on Max Ent Principle so Alice should choose message m_1 over m_2 as Bob is Cognizant!

Qiv) Max Ent principle : Action principle.

Newtonian mechanics $\rightarrow$ Lagrangian / Hamiltonian mechanics



force drives dynamics

minimise energy in gravitational potential



pressure drives gas

maximize free energy in a thermal state

v) Soft data Constraint
 $\langle f(s) \rangle = d$

Maximise $-D_S(I_S | I_0)$ with $\langle f(s) \rangle_{p(s|I_0)} \stackrel{!}{=} d$.



$$p(s|I_0) =: p(s) \quad p(s|I_A) =: q(s)$$

$$\Rightarrow \max \left[-\int ds p(s) \ln \left(\frac{p(s)}{q(s)} \right) + \lambda \left(\int ds p(s) - 1 \right) + \mu \left(\int ds f(s) p(s) - d \right) \right]$$

constraints

analyticity $\rightarrow$

$$0 \stackrel{!}{=} \frac{\delta}{\delta p(s)} [\dots]$$

$$= -\ln \left(\frac{p(s)}{q(s)} \right) - 1 + \lambda + \mu f(s)$$

$$\Rightarrow p(s) = q(s) e^{\lambda - 1 + \mu f(s)}$$