

Decision theory

-1-

action $a \in A$ = set of possible actions

$$\text{utility: } u(s, a) = -l(s, a)$$

$\uparrow$ situation $\uparrow$ loss

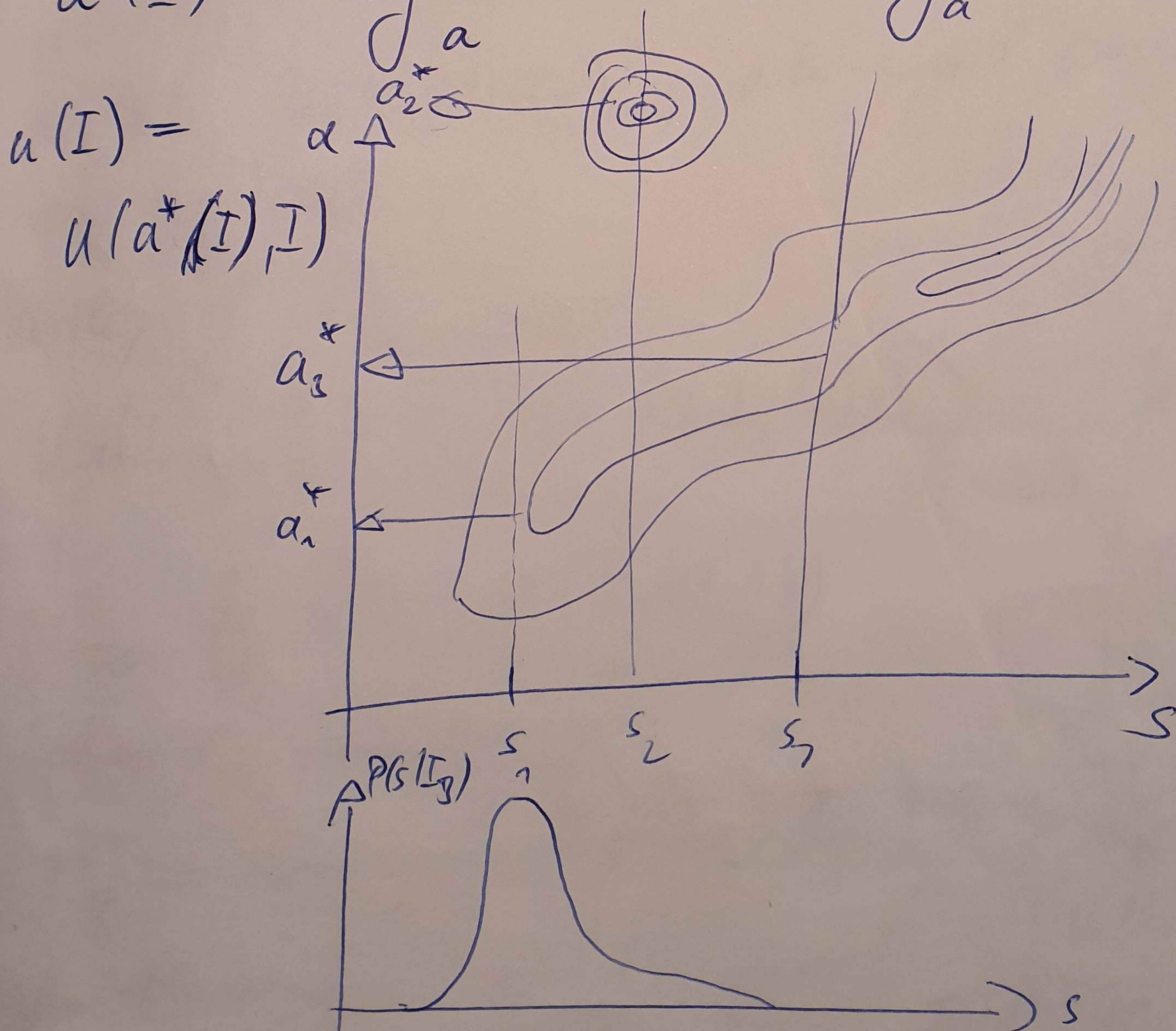
Expected utility given I on $s \in S'$

$$u(a, I) := \langle u(s, a) \rangle_{(s|I)} = \int ds P(s|I) u(s, a)$$

$S' = \{s | l(s, a, I) \}$

optimal action given I

$$a^*(I) = \underset{a}{\operatorname{argmax}} u(a, I) = \underset{a}{\operatorname{argmin}} l(a, I)$$



- 1,5 -

- quadratic loss

- linear loss

- delta losses

→ see script!

Optimal action

$$a_s^* = \underset{a \in A}{\operatorname{argmax}} u_s(a | I_s)$$

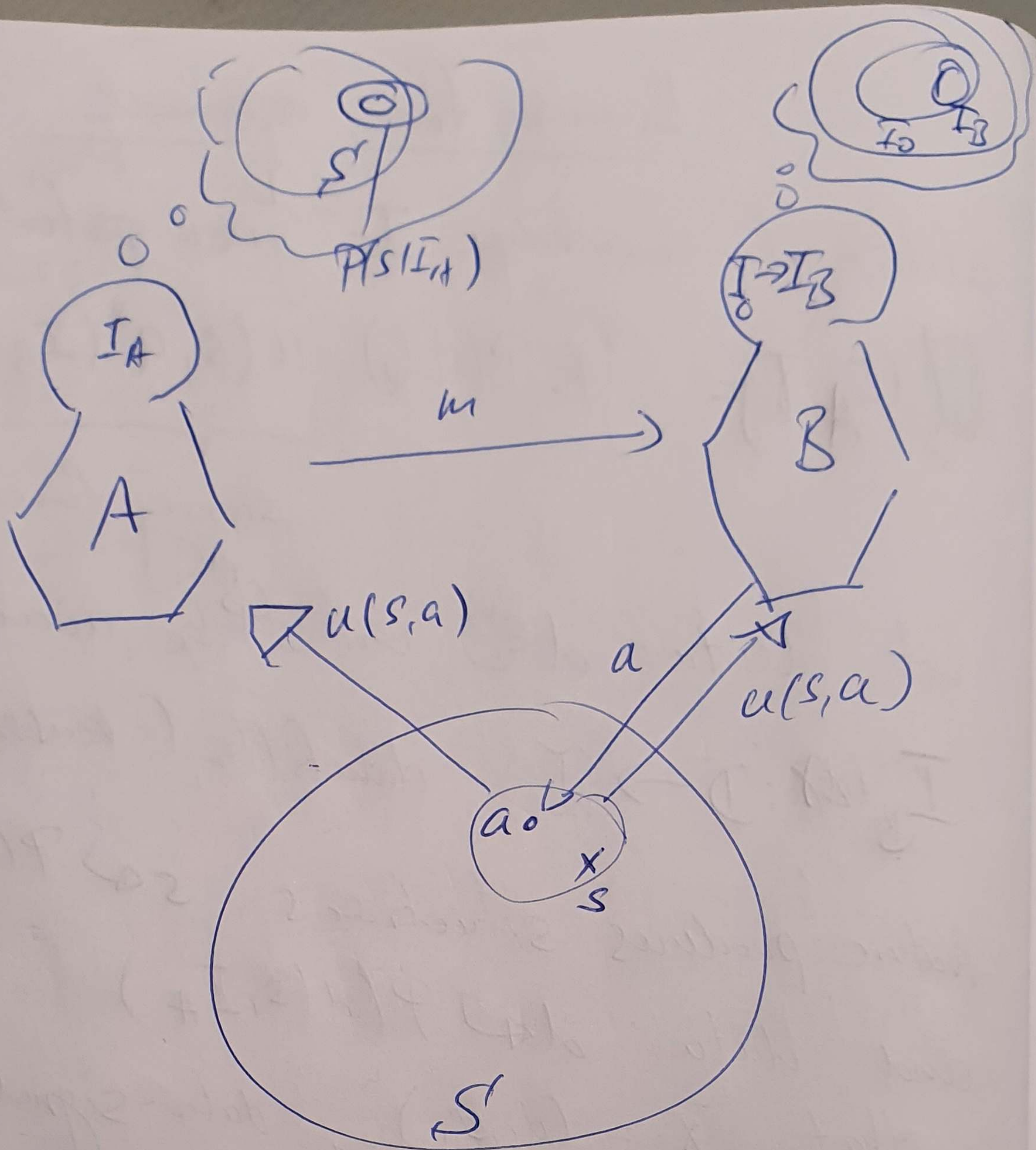
$$u_s(a | I) = \langle u(s | a) \rangle_{(s | I)}$$

assume: a^* is uniquely determined
 $u(s, a)$ differentiable w.r.t. a

$$\Rightarrow 0 = \frac{\partial u(a | I)}{\partial a} = \langle g(s, a) \rangle_{(s | I)}$$

$$g(s, a) = \frac{\partial u(s, a)}{\partial a}$$

-3-



Alice : message m about $s \in S'$

Bob : update $I_0 \rightarrow I_B$

action a

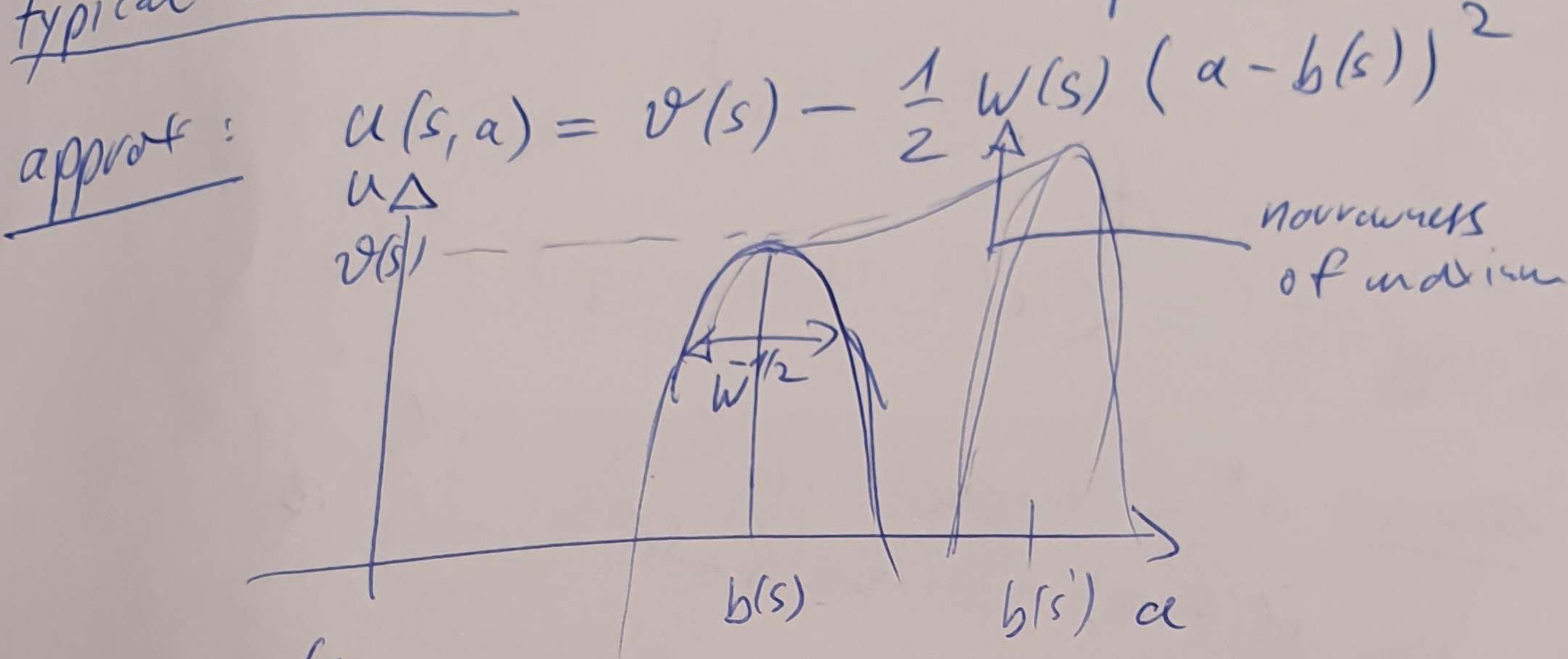
both : experience utility $u(s, a)$

See: Eyblin et al 2024

Attention to Entropic Communication

Ann Phys. 2024, 536, 2300334

typical situation s has one optimal action $a^*(s)$



Bob's action

$$a_B^*(I_B) = \arg \max_a \langle u(s, a) \rangle_{(s|I_B)}$$

$$= a : \left[\left\langle \frac{\partial u(s, a)}{\partial a} \right\rangle_{(s|I_B)} = 0 \right]$$

such that

$$= a : \left[- \langle (a - b(s)) w(s) \rangle_{(s|I_B)} = 0 \right]$$

$$= a : \left[\langle b(s) w(s) \rangle_{(s|I_B)} = a \langle w(s) \rangle_{(s|I_B)} \right]$$

$$= \frac{\langle b(s) w(s) \rangle_{(s|I_B)}}{\langle w(s) \rangle_{(s|I_B)}} =: \langle b(s) \rangle_{(s|I_B)}^{(w)}$$

$$\langle b(s) \rangle_{(s|I)}^{(w)} =: \int ds \mathcal{A}(s|I_B) b(s)$$

$$\mathcal{A}(s|I_B)^{(w)} =: \frac{w(s) P(s|I_B)}{\int ds w(s) P(s|I_B)}$$

Attention!

Learning from experience

utility of knowledge $I_B \in \mathcal{I}$ given perfect knowledge I_A

$$U_{\mathcal{S}}(I_A, I_B) = \int ds \, P(s|I_A) \underbrace{u(s, a^*(I_B))}_{\text{scoring function}}$$

Let data $d \in \mathcal{D}$ on $\mathcal{S}^{\mathcal{D}}$ be available to $\mathcal{I}$

$\mathcal{I}_B: \mathcal{D} \rightarrow \mathcal{I}$ data filter (= reasoning system)

Nature produces situations $s \in \mathcal{S} \sim P(s|I_A)$

and data $d \in \mathcal{D} \sim P(d|s, I_A)$

so that $\bar{\mathbf{x}}_N = (d_i, s_i)_{i=1}^N$ data-signal pairs are available

Utility of $\mathcal{I}_B$

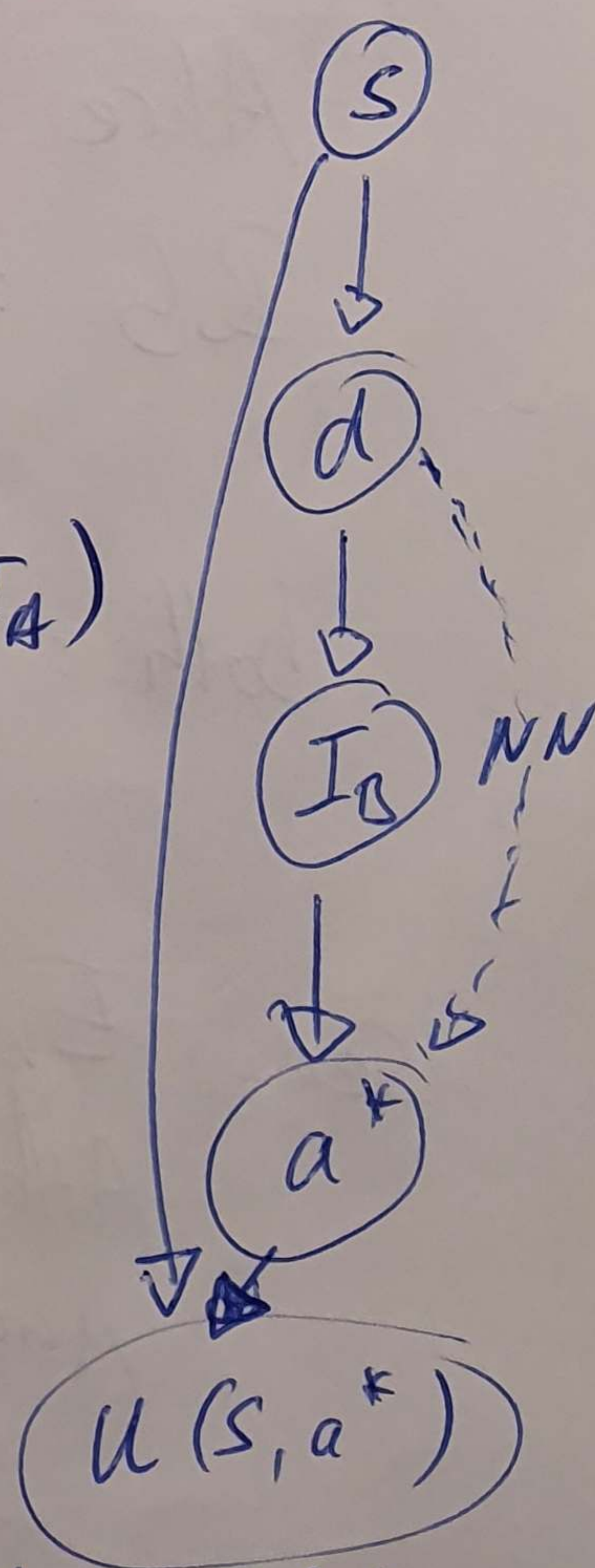
$$U_{\mathcal{S}}(I_A, \mathcal{I}_B) = \left\langle u(s, a^*(\mathcal{I}_B(d))) \right\rangle_{d \sim P(d|s, I_A)}$$

$$\approx \frac{1}{N} \sum u(s_i, a^*(\mathcal{I}_B(d_i)))$$

optional reasoning:

$$\mathcal{I}_B^* = \arg \max_{\mathcal{I}_B} U_{\mathcal{S}}(I_A, \mathcal{I}_B)$$

Pavlov-like learning: $a^*(d)$, abstract learning $\mathcal{I}_B(d)$, $a^*(s)$



Achieved information gain (AIG)

-6-

I_A
Alice

$\xrightarrow{m}$

$I_0 \rightarrow I_B$
Bob

Epfl et al
(2025)

Wanted: gain $G_S(I_A, I_B, I_0)$ that Alice uses to decide about which I_B she gives Bob

Axioms:

1) Additivity:

if Alice knowledge is hierarchical
$$P(s | I_A) = \sum_{i=1}^n P_{A,i} P(s | I_{A,i})$$

and $I_A = (I_{A,i}, P_{A,i})_{i=1}^n$

Knowledge sub states, their probabilities

then

$$G_S(I_A, I_B, I_0) = \sum_{i=1}^n P_{A,i} G_S(I_{A,i}, I_B, I_0)$$

this should also hold in the continuous

$$P(s | I_A) = \int dx P_A(x) P(s | I_A(x))$$

$$\Rightarrow G_S(I_A, I_B, I_0) = \int dx P_A(x) G_S(I_A(x), I_B, I_0)$$

2) Locality: if Alice knows which $s' \in S$ happens, ⁻⁷
 $I_A = I_{s=s'}$, $P(s | I_A) = \delta(s-s')$, gain
 should only depend on $P(s' | I_B)$, but
 not on $P(s'' | I_B)$ for $s'' \neq s'$

3) Analytical: $G_s(I_A, I_B, I_0) \in C^{\infty}(\mathcal{I})$ w.r.t I_B

4) Proper: $G_s(I_A, I_B, I_0)$ be maximal w.r.t I_B
 for $I_B = I_A$

5) Calibration: no update, no gain
 $G_s(I_A, I_0, I_0) = 0$

Derivation

decomposition of I_A into atomic beliefs by
 identifying $x=s$, $\bar{X}=S'$ such that

$$P_A(x) := P(s=x | I_A), \quad I_A(x) = I_B = x$$

$$P(s | I_A) = \int dx \underbrace{P_A(x)}_{P(s=x | I_A)} \underbrace{P(s | I_{s=x})}_{\delta(s-x)} = P(s | I_A) \quad \checkmark$$

additive $\Rightarrow$

$$G_s(I_A, I_B, I_0) = \int dx P_A(x) G_s(I_A(x), I_B, I_0)$$

$$= \int_S ds' P(s' | I_A) G_s(I_{s=s'}, I_B, I_0)$$

$$= \langle G_s(I_{s=s'}, I_B, I_0) \rangle_{(s' | I_A)}$$

Locality $\Rightarrow G_{s'}(I_{s=s'}, I_B, I_0) = g(s', P(s'|I_B), I_0) - \lambda \left(1 - \int ds'' P(s''|I_B) \right)$

to be found function

Lagrange multiplier to ensure normalization of $P(s''|I_B)$

maximizing w.r.t $q(s) := P(s|I_B)$

Properness $0 = \frac{\delta G_s(I_A, I_B, I_0)}{\delta P(s'|I_B)} \Big|_{I_B=I_A}$

$\frac{\partial G_s}{\partial \lambda} = 0 \Rightarrow \int ds'' P(s''|I_B) = 1$

properness: maximum at $I_B = I_A$!

$= \frac{\delta}{\delta q(s')} \left\langle g(s, q(s), I_0) + \int_s (1 - \int ds'' q(s'')) \lambda \right\rangle_{(s|I_A)} \Big|_{q=P(\cdot|I_A)}$

$= \frac{\delta}{\delta q(s')} \left[\left\langle g(s, q(s), I_0) \right\rangle_{(s|I_A)} + (1 - \int ds'' q(s'')) \lambda \right]_{q=P(\cdot|I_A)}$

$= \frac{\partial g(s', q(s'), I_0)}{\partial q(s')} P(s|I_A) - \lambda \Big|_{q=P(\cdot|I_A)}$

$= \frac{\partial g(s', \tilde{q}, I_0)}{\partial \tilde{q}} \tilde{q} - \lambda \Big|_{\tilde{q}=P(s'|I_A)}$

$\Rightarrow \frac{\partial g(s', \tilde{q}, I_0)}{\partial \tilde{q}} = \frac{\lambda}{\tilde{q}} \Rightarrow g(s, \tilde{q}, I_0) = \lambda \ln \tilde{q} + c(s, I_0)$

Calibration:

$$0 = G_S(I_A, I_0, I_0) = \langle g(s, P(s|I_0), I_0) \rangle_{(s|I_A)}$$

$$= \langle \lambda \ln P(s|I_0) + c(s, I_0) \rangle_{(s|I_A)}$$

$$\Rightarrow \langle c(s, I_0) \rangle_{(s|I_A)} = -\lambda \langle \ln P(s|I_0) \rangle$$

$$\Rightarrow G_S(I_A, I_B, I_0) = \langle g(s, P(s|I_B), I_0) \rangle_{(s|I_A)}$$

$$= \langle \lambda \ln P(s|I_B) + c(s, I_0) \rangle_{(s|I_A)}$$

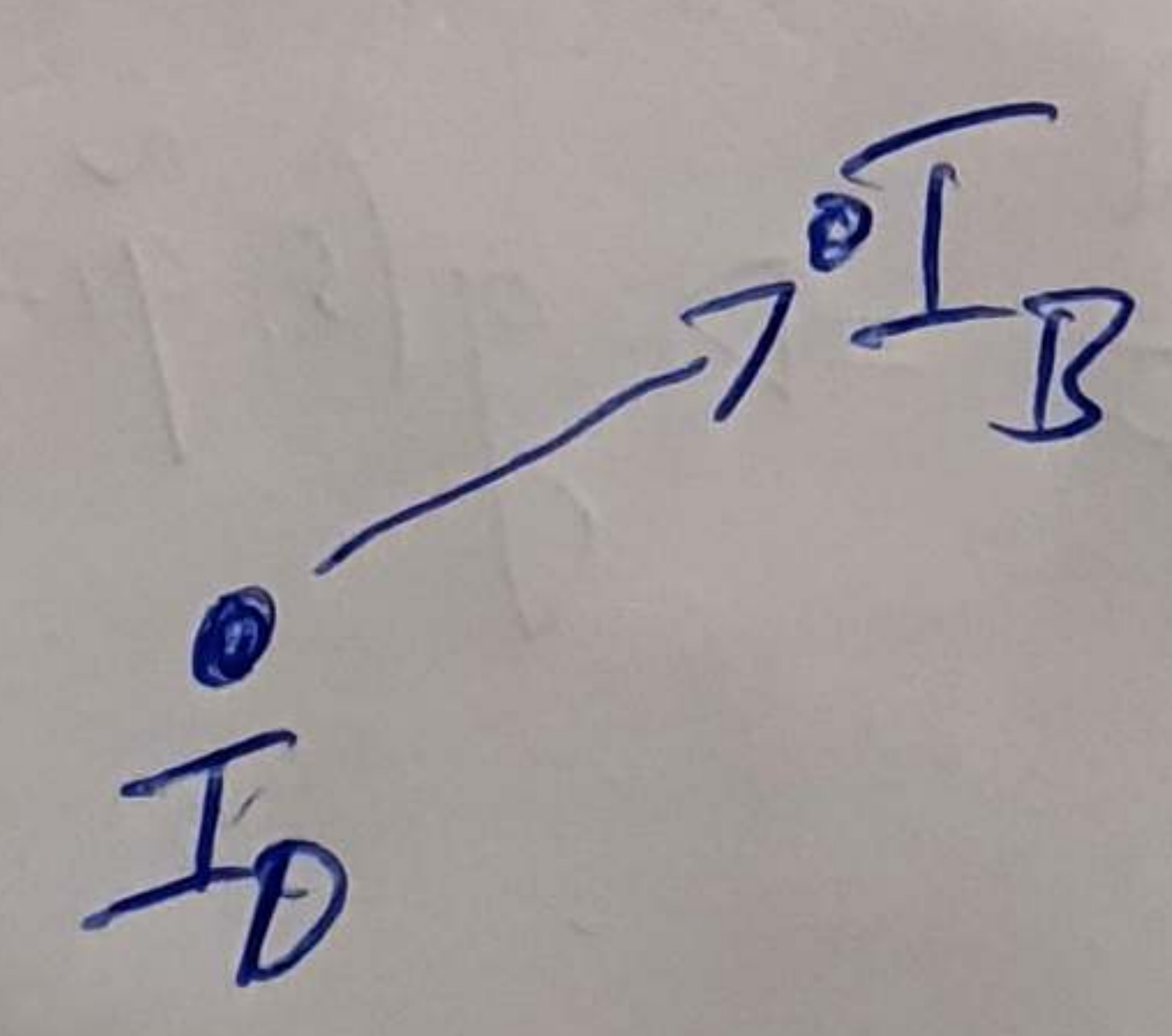
$$= \lambda \langle \ln P(s|I_B) - \ln P(s|I_0) \rangle_{(s|I_A)}$$

$$G_S(I_A, I_B, I_0) = \lambda \int_{\mathcal{S}} ds P(s|I_A) \ln \frac{P(s|I_B)}{P(s|I_0)}$$

$$=: D_{S^1}(I_A, I_B, I_0)$$

Achieved information gain for the update
 $I_0 \rightarrow I_B$ from the perspective of I_A

• I_A



units: nits if
 $\ln$
 bits if $\log_2$
 is used

Properties of AIG

ideal AIG: $I_B = I_A$

$$D_S(I_A, I_A, I_0) = \int_{S'} ds P(s|I_A) \ln \frac{P(s|I_A)}{P(s|I_0)}$$

$$=: D_{S'}(I_A, I_0)$$

$$\equiv D_{KL}(P(s|I_A) \parallel P(s|I_B))$$

Kullback-Leibler Divergence or
cross-entropy, $D_S(I_A, I_0) \geq 0$ if $I_A \neq I_0$

no update: $I_B = I_0$

$$D_S(I_A, I_0, I_0) = \int_{S'} ds P(s|I_A) \ln \frac{P(s|I_0)}{P(s|I_0)} = 0$$

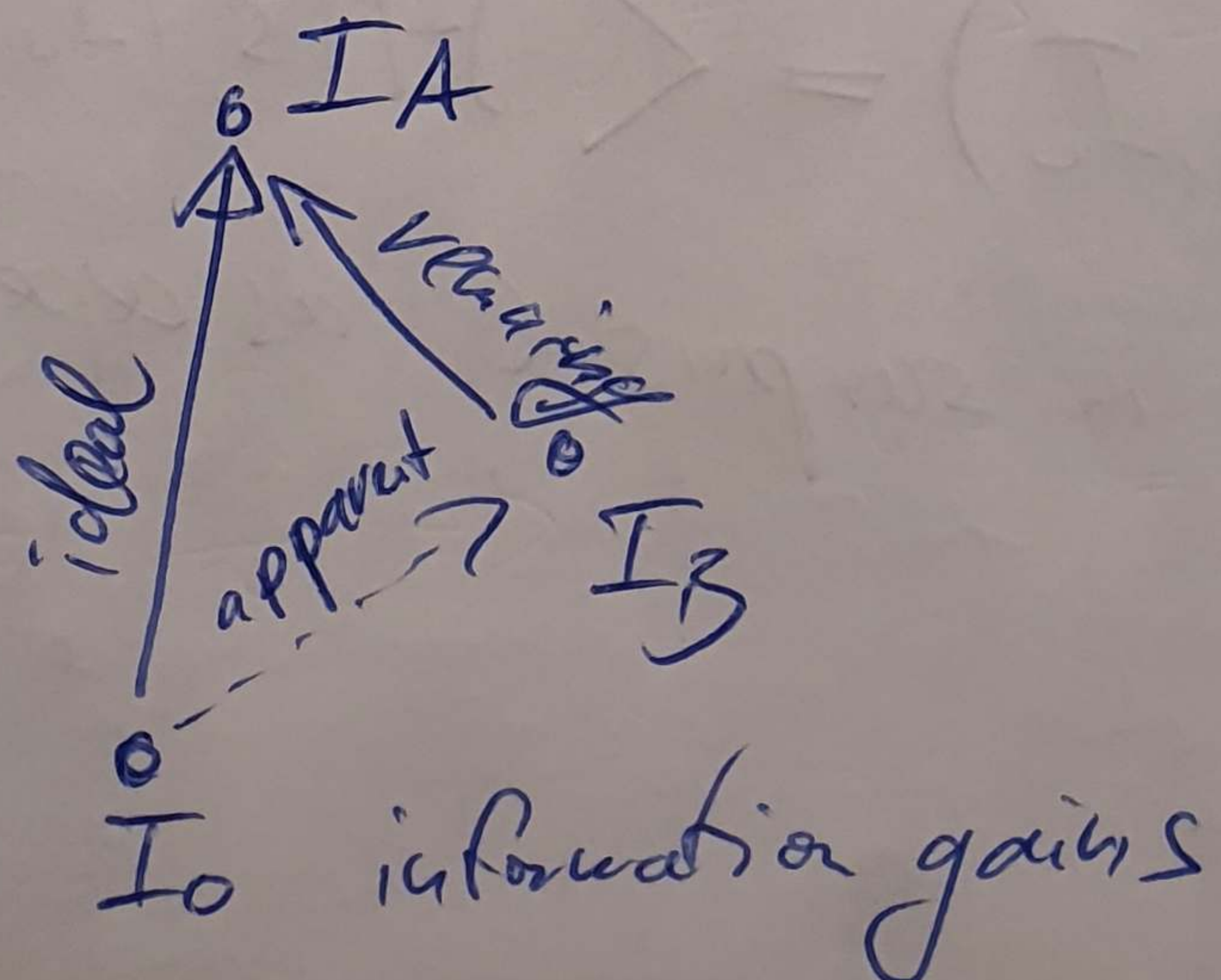
as requested by calibration.

relation to KL

$$D_S(I_A, I_B, I_0) = \int ds P(s|I_A) \left[\ln \frac{P(s|I_A)}{P(s|I_0)} - \ln \frac{P(s|I_A)}{P(s|I_B)} \right]$$

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

= (ideal -
remaining)
information
gain



AIG for perfect knowledge (Alice is infinitely informed) -11-

$$I_A = I_S = s'$$

$$D_S(I_A, I_B, I_0) = \int_{s'} ds \delta(s-s') \ln \frac{P(s|I_B)}{P(s|I_0)}$$

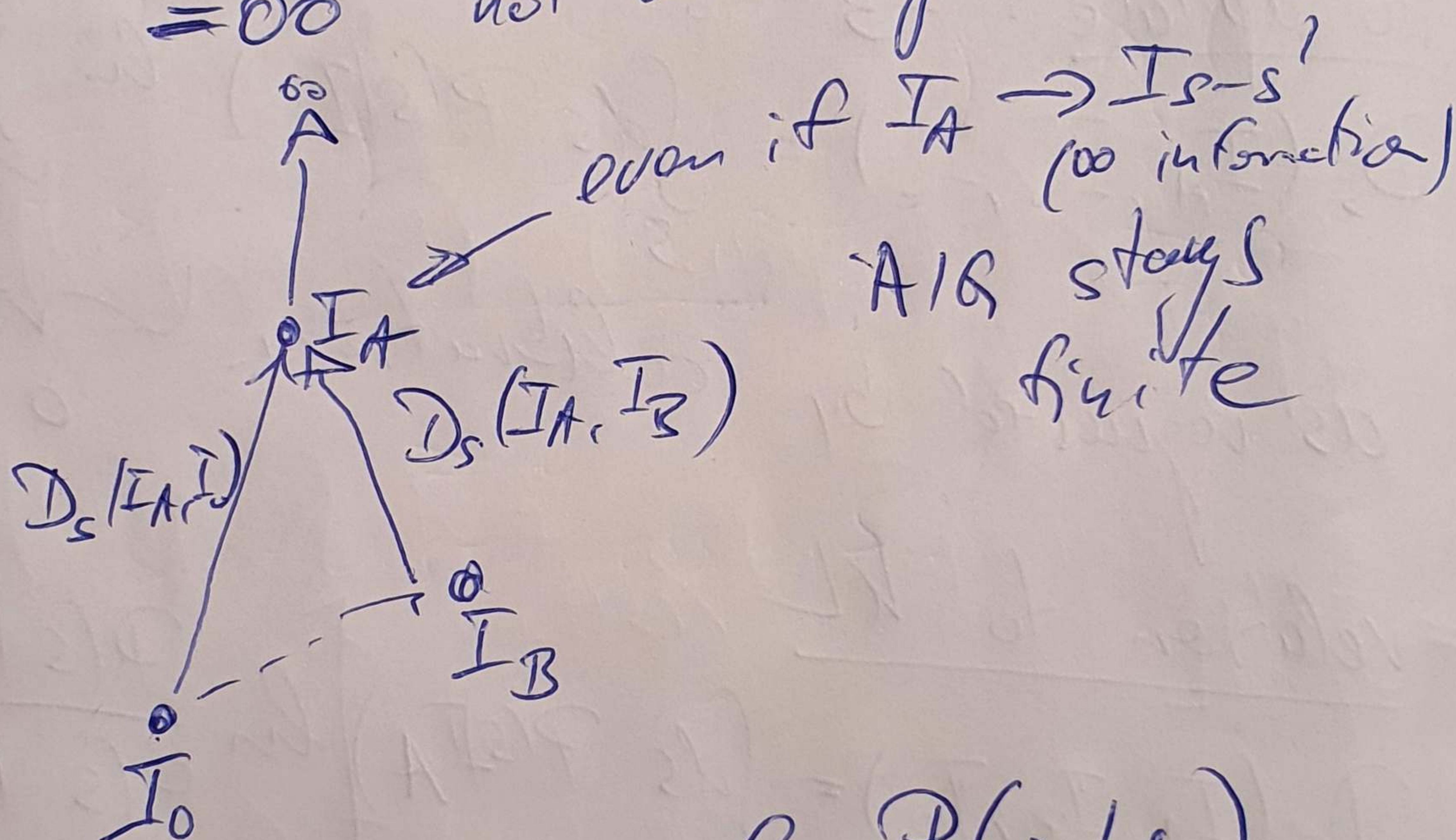
$$= \ln \frac{P(s'|I_B)}{P(s'|I_0)}$$

well defined as long
none of these is 0

However:

$$D_S(I_A, I_{B/0}) = \int ds \delta(s-s') \ln \frac{\delta(s-s')}{P(s|I_{B/0})}$$

$$= \infty \quad \text{not well defined}$$



surprise $H(\cdot | \cdot) := -\ln P(\cdot | \cdot)$

$$D_S(I_A, I_B, I_0) = \langle H(s|I_0) - H(s|I_B) \rangle_{(s|I_A)}$$

Bb's reduction in surprise as expected by Alice

Separability

be $S = S_1 \otimes S_2$, $s = \begin{pmatrix} s_1 \\ s_2 \end{pmatrix}$ and

I_0 's knowledge I_0, I_B separable,

$$P(s | I_{0/B}) = P(s_1 | I_{0/B}) P(s_2 | I_{0/B})$$

$$\Rightarrow D_S(I_A, I_B, I_0) = \left\langle \ln \frac{P(s_1 | I_B) P(s_2 | I_B)}{P(s_1 | I_0) P(s_2 | I_0)} \right\rangle_{(s | I_A)}$$

$$= \left\langle \ln \frac{P(s_1 | I_B)}{P(s_1 | I_0)} \right\rangle_{(s | I_A)} + \left\langle \ln \frac{P(s_2 | I_B)}{P(s_2 | I_0)} \right\rangle_{(s | I_A)}$$

$$= D_{S_1}(I_A, I_B, I_0) + D_{S_2}(I_A, I_B, I_0)$$

Antisymmetry

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

Learning a bit

$S = \{\text{head}, \text{tail}\}$, $I_A = I_{S=\text{head}} = \text{head}$

$I_0: P(s | I_0) = 1/2 \rightarrow I_B = I_A$

$$D_S(I_A, I_A, I_0) = \sum_{s \in S} P(s | I_A) \log_2 \frac{P(s | I_A)}{P(s | I_0)}$$

$$= P(\text{head} | \text{head}) \log_2 \frac{P(\text{head} | \text{head})}{1/2} + P(\text{tail} | \text{head}) \log_2 \frac{0}{1/2}$$

$$1 = \log_2 2 = 1 \quad \checkmark$$

Cognitive operations

for communicating $A \xrightarrow{m} B$
 $I_A \quad I_0 \rightarrow I_B$
 $D_S(I_A, I_B, I_0)$ is AIG of communication act

memorizing

$I_0 = I_A$ Alice's initial memory
 I_A " perfect update
 I_B " imperfect memory update (future self)

$$D_S(I_A, I_B, I_0) \text{ is AIG of memory update}$$

$$= D_S(I_A, I_B, I_A) = \underbrace{D_S(I_A, I_A)}_{=0} - \underbrace{D_S(I_A, I_B)}_{\geq 0}$$

memorizing is at best lossless!

inference

prior $P(s | I_0)$

message data $d \rightarrow$ ideal update to $\underbrace{P(s | d, I_0)}_{I_A}$

imperfect update $P(s | I_B)$

$D_S(d, I_0, I_B)$ is AIG of approximate inference.

Cognitive fidelity:

$$CF(I_A, I_B, I_0) := \frac{D_S(I_A, I_B, I_0)}{D_S(I_A, I_A, I_0)}$$

$$= \frac{D_S(I_A, I_0) - D_S(I_A, I_B)}{D_S(I_A, I_0)}$$

$$= 1 - \frac{D_S(I_A, I_B)}{D_S(I_A, I_0)} \in [-\infty, 1]$$

$$CF(I_A, I_A, I_0) = 1$$

for perfect update

$$CF(I_A, I_0, I_0) = 0$$

for no update

$$CF(I_A, I_B, I_0) < 0 \quad \text{if } D_S(I_A, I_0) < D_S(I_A, I_B)$$

$\bullet I_A$

$\bullet I_0$

$\searrow$
 $\bullet I_B$

Cognitive efficiency

$$CE(I_A, I_B, I_0) = \frac{D_S(I_A, I_B, I_0)}{C(I_A, I_B, I_0)}$$

costs $\rightarrow$

units: bits/Euro, Energy, time, ...

AIG for Bernoulli experiment

$$S = \{0, 1\} \quad P_{A,B,0} := P(S=0 | I_A, B, C)$$

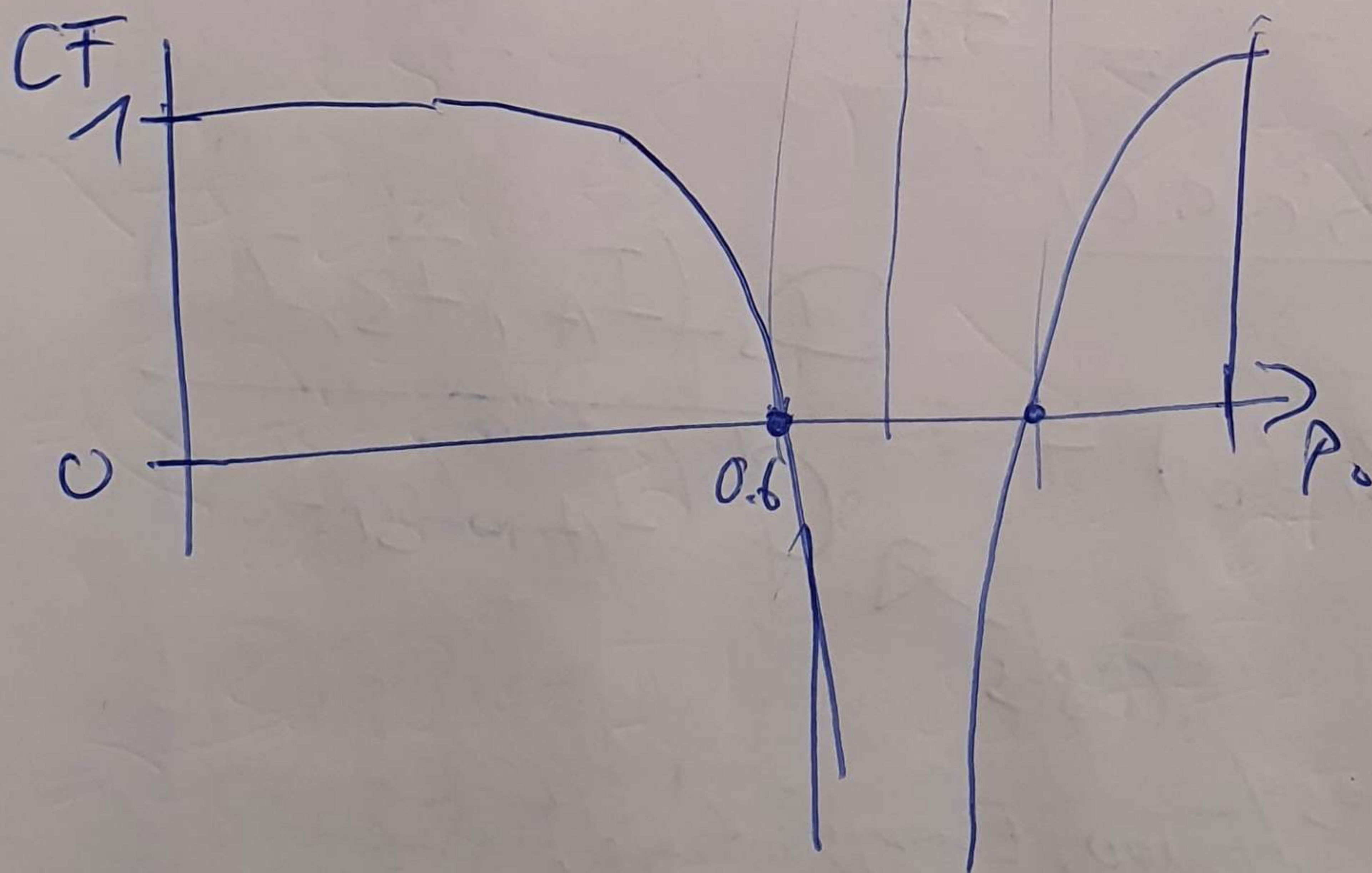
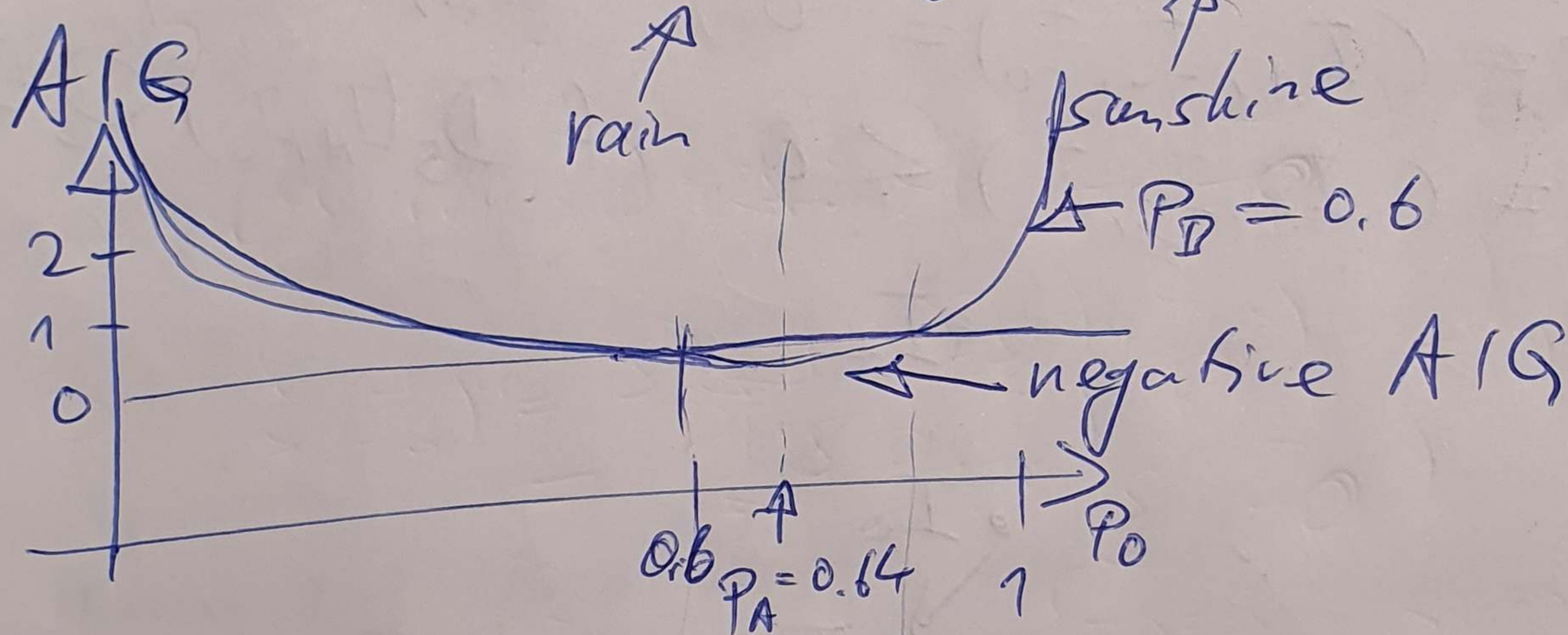
rain sunshine

Alice = weather reporter

Bob = audience

Should A update B when she can only communicate ~~0.6~~ P_A rounded to 10%?

$$D_S(I_A, I_B, I_0) = P_A \ln \frac{P_B}{P_0} + (1-P_A) \ln \frac{1-P_B}{1-P_0}$$



expected AIG from generative model - 16

Scenario: update after receiving data d

prior $P(s | I_0)$

posterior $P(s | d, I_0) =: P(s | I_A)$ ideal update

approx posterior $P(s | I_B(d), I_0)$

how to get ~~$P(s | I_B(d), I_0)$~~ calculate

AIG for $I_B(d)$ when $P(s | d, I_0)$ is intractable?

Assume $(d_i, s_i)_i \leftarrow P(d, s | I_0)$ are

available via $s_i \leftarrow P(s_i | I_0)$ and

$d_i \leftarrow P(d | s_i, I_0)$, then expected

AIG of perfect knowledge $I_{s=s_i}$ averaged over i

$$\langle D_S(I_{s=s_i}, I_B(d_i), I_0) \rangle_i \approx \langle D_S(I_{s=s'_i}, I_B(d'), I_0) \rangle_{(d', s' | I_0)}$$

$$= \langle D_S(I_{s=s'}, I_B(d'), I_0) \rangle_{(s' | d', I_0)} \rangle_{(d' | I_0)}$$

$$= \left\langle \int ds' P(s' | d, I_0) \int ds S(s-s') \frac{P(s' | I_A(d))}{P(s | I_B(d))} \right\rangle_{(d | I_0)}$$

$$= \left\langle \int ds' P(s' | d, I_0) \ln \frac{P(s' | I_B(d))}{P(s' | I_0)} \right\rangle_{(d | I_0)}$$

$$= \langle D_S(I_A, I_B(d), I_0) \rangle_{(d | I_0)}$$

- 17 -

Recipe to get expected AIB of a
cognitive operation $I_B(d)$:

draw samples $s_i \leftarrow P(s_i | I_0)$
 $d_i \leftarrow P(d_i | s_i, I_0)$

$$\left\langle D_S(I_A, I_B(d), I_0) \right\rangle_{(d, I_0)} \approx \sum_{i=1}^n \ln \frac{P(s_i | I_B(d_i))}{P(s_i | I_0)} \\ = \frac{1}{n} \sum_{i=1}^n \left[H(s_i | I_0) - H(s_i | I_B(d_i)) \right]$$

why does this work? product rule!

$$P(d, s | I_0) = P(d | s, I_0) P(s | I_0) \\ = P(s | d, I_0) P(d | I_0)$$

If AIB is of value to a cognitive system
and evolution shaped it, it will aim
to maximize it, which means there is
a selection force pushing $I_B(d)$ towards
the Bayesian optimum $I_A(d) = (I_0, d)$

~~no let's keep it as is~~