

Active inference and predictive coding in human cognition

Chris Mathys



Machine Learning and Human Cognition summer school

Technical University of Denmark

17 August 2020

What do you see here?

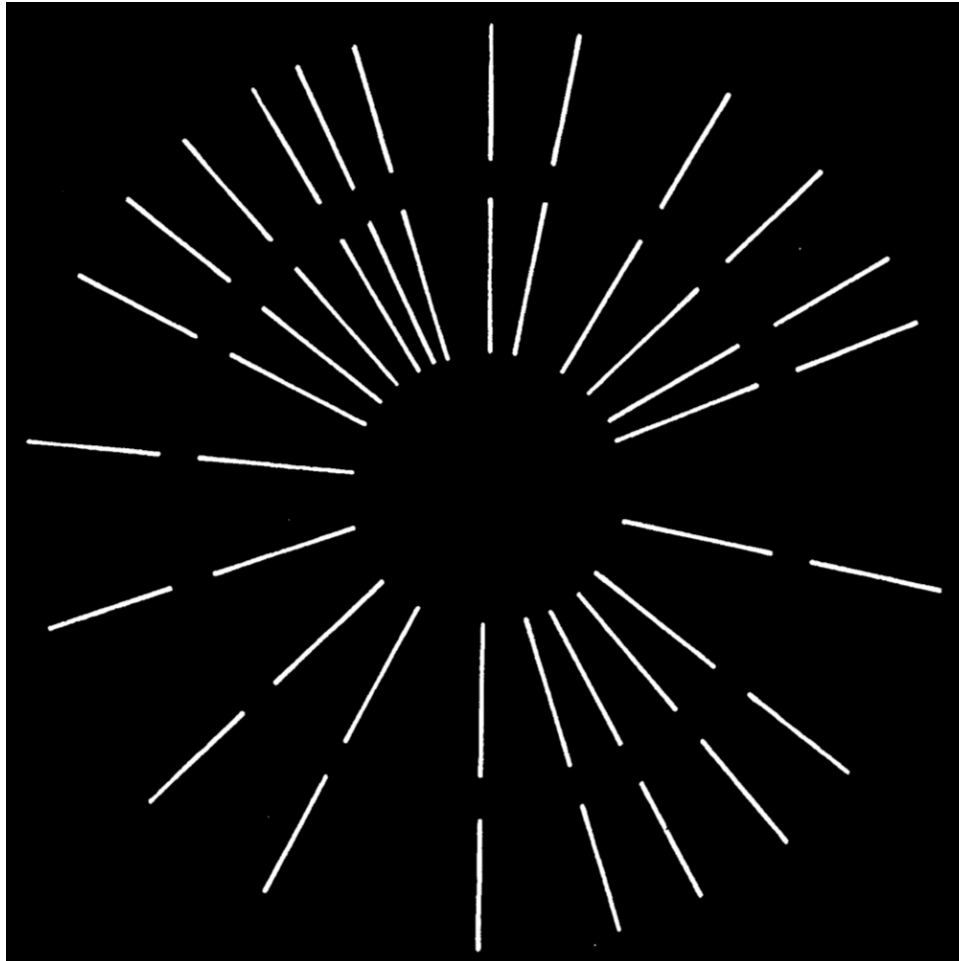


FIGURE 2. Illusory contours and regions of modified brightness seem to be postulates of nearer eclipsing objects, to 'explain' unlikely gaps. It seems that an essentially Bayesian strategy is responsible.

What do you see here?

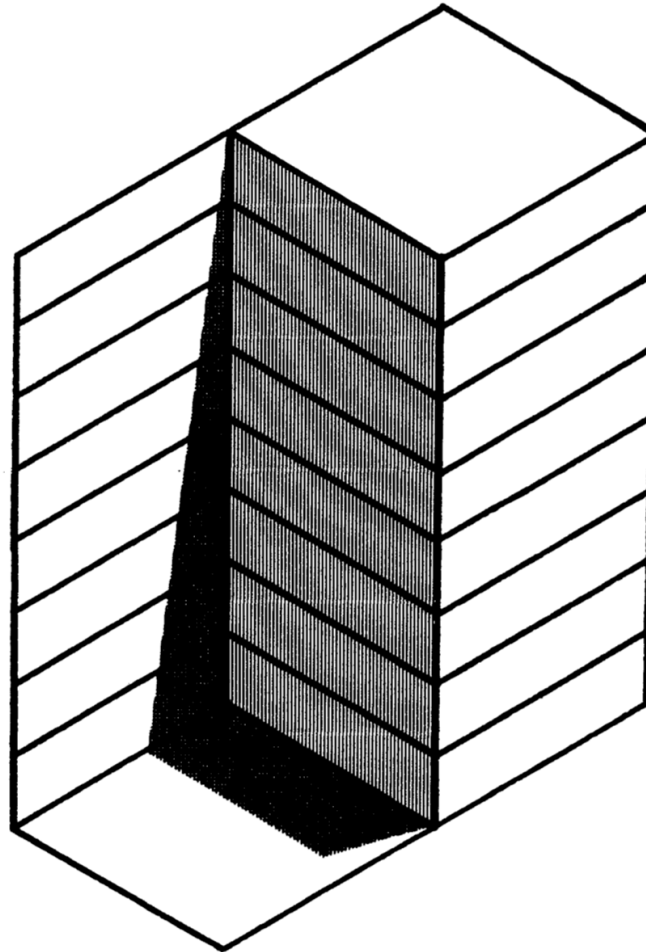


FIGURE 3. In this depth-ambiguous figure the grey rectangular region changes in brightness for most observers according to whether this area is a probable shadow. The systematic brightness change with depth-reversal is, clearly, centrally initiated and is a 'downwards' effect.

D. Mumford

Mathematics Department, Harvard University, 1 Oxford Street, Cambridge, MA 02138, USA

Biol. Cybern. 66, 241–251 (1992)

What do you see here?



Fig. 2. Top down processing reveals the dalmation dog (Rock 1984)

What do you see here?

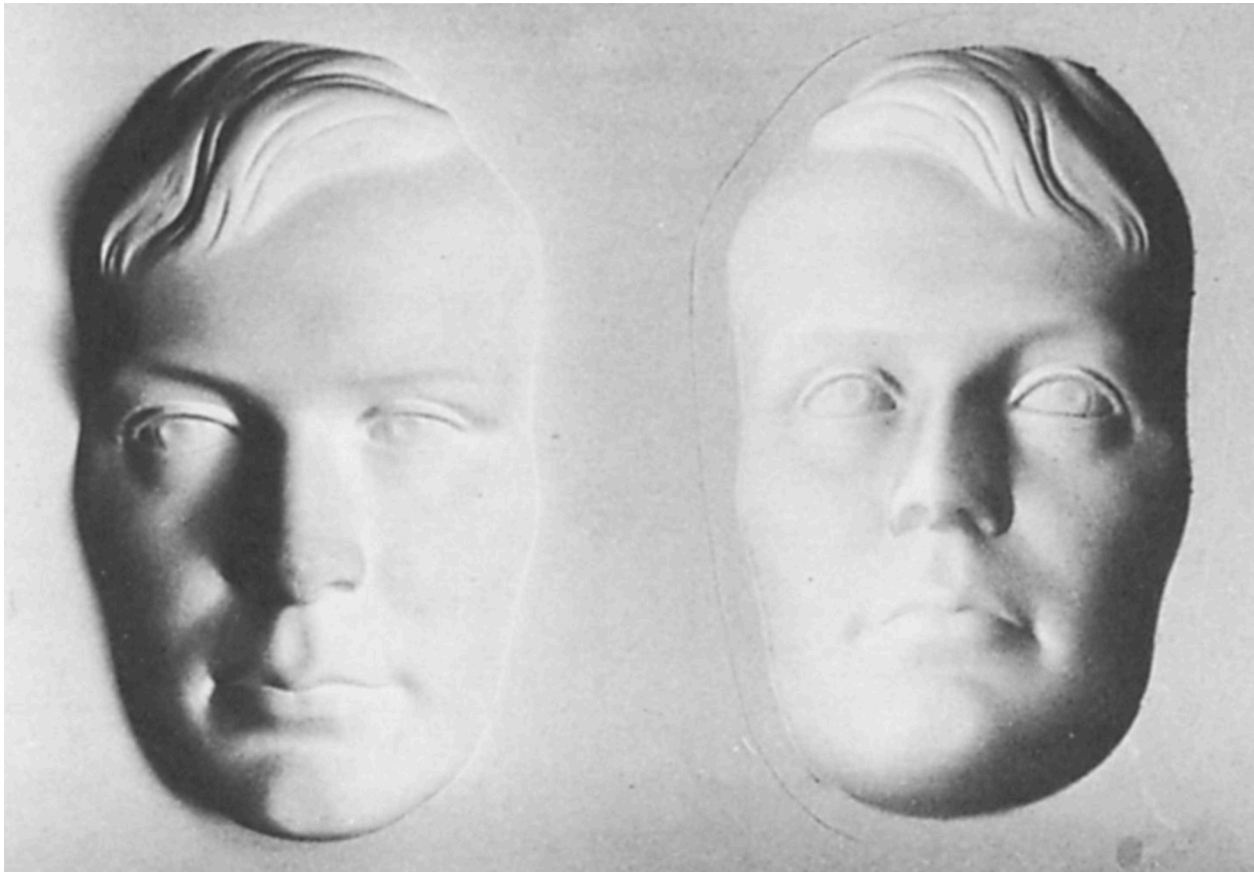


FIGURE 6. The 'face' on the right is in fact the hollow mould of the face on the left. Though hollow, the mould appears as a normal face. Texture and even stereoscopic counter-information are rejected to maintain this highly probable (but incorrect) hypothesis.

More on this:

https://www.giannisarcone.com/Muller_lyer_illusion.html

<https://www.youtube.com/watch?v=ORoTCBrCKIQ>

<https://www.youtube.com/watch?v=6YIPtJlCbIA>

In perceiving and acting, what problem are we solving?

According to Helmholtz, we are **minimizing our surprise** at the sensory signals reaching us:

‘Objects are always imagined as being present in the field of vision as would have to be there in order to produce the same impression on the nervous mechanism.’

— Helmholtz (1867), Handbuch der physiologischen Optik [Treatise on Physiological Optics], p. 428

The perceived (‘imagined’) object is the one that is least unlikely given the sensory data (‘impression on the nervous mechanism’).

Visually exploring a scene (i.e., acting to move our eyeballs) allows for testing hypotheses about the objects present in the field of vision.

In other words: the mind **actively infers** the objects of its perception in what Helmholtz calls ‘unconscious inference’ (p. 430).

If this is true, it means that **to understand the mind and the brain, we need to understand the mechanics of inference**, i.e. the mechanics of updating predictions.

Why is it important to minimize surprise?

- In order for us (or any biological organism) to survive and prosper, it is important that we understand how our environment works.
- The key measure of this understanding (or rather the *lack* of it) is the surprise we experience at the sensory input we receive over time.
- Given a probabilistic model of how the environment generates our sensory input, surprise can be defined mathematically and minimized algorithmically.

“Every good regulator of a system must be a model of that system”¹

Abstract:

«The design of a complex regulator often includes the making of a model of the system to be regulated. The making of such a model has hitherto been regarded as optional, as merely one of many possible ways.

In this paper a theorem is presented which shows, under very broad conditions, that any regulator that is maximally both successful and simple must be isomorphic with the system being regulated. (The exact assumptions are given.) Making a model is thus necessary.

*The theorem has the interesting corollary that **the living brain**, so far as it is to be successful and efficient as a regulator for survival, **must proceed**, in learning, **by the formation of a model (or models) of its environment.**»*

¹Conant, R. C., & Ashby, R. W. (1970). Every good regulator of a system must be a model of that system. *International Journal of Systems Science*, 1(2), 89–97. <https://doi.org/10.1080/00207727008920220>

What does ‘good regulation’ mean?

- It means keeping the entropy of inputs low
- Entropy is *expected surprise* (cf. later)
- Given a model m , under ergodic assumptions (i.e., time t spent in a state x corresponds to probability of that state), the **entropy** $H(x)$ is:

$$\begin{aligned} H(x) &:= - \int p(x|m) \log p(x|m) dx \\ &= - \lim_{T \rightarrow \infty} \frac{1}{T} \int_0^T \log p(x|m) dt \end{aligned}$$

- This means that in order to keep entropy of inputs low, our agent needs to **keep surprise at any one point in time low.**

How can we minimize surprise?

At the most general level, there are two ways:

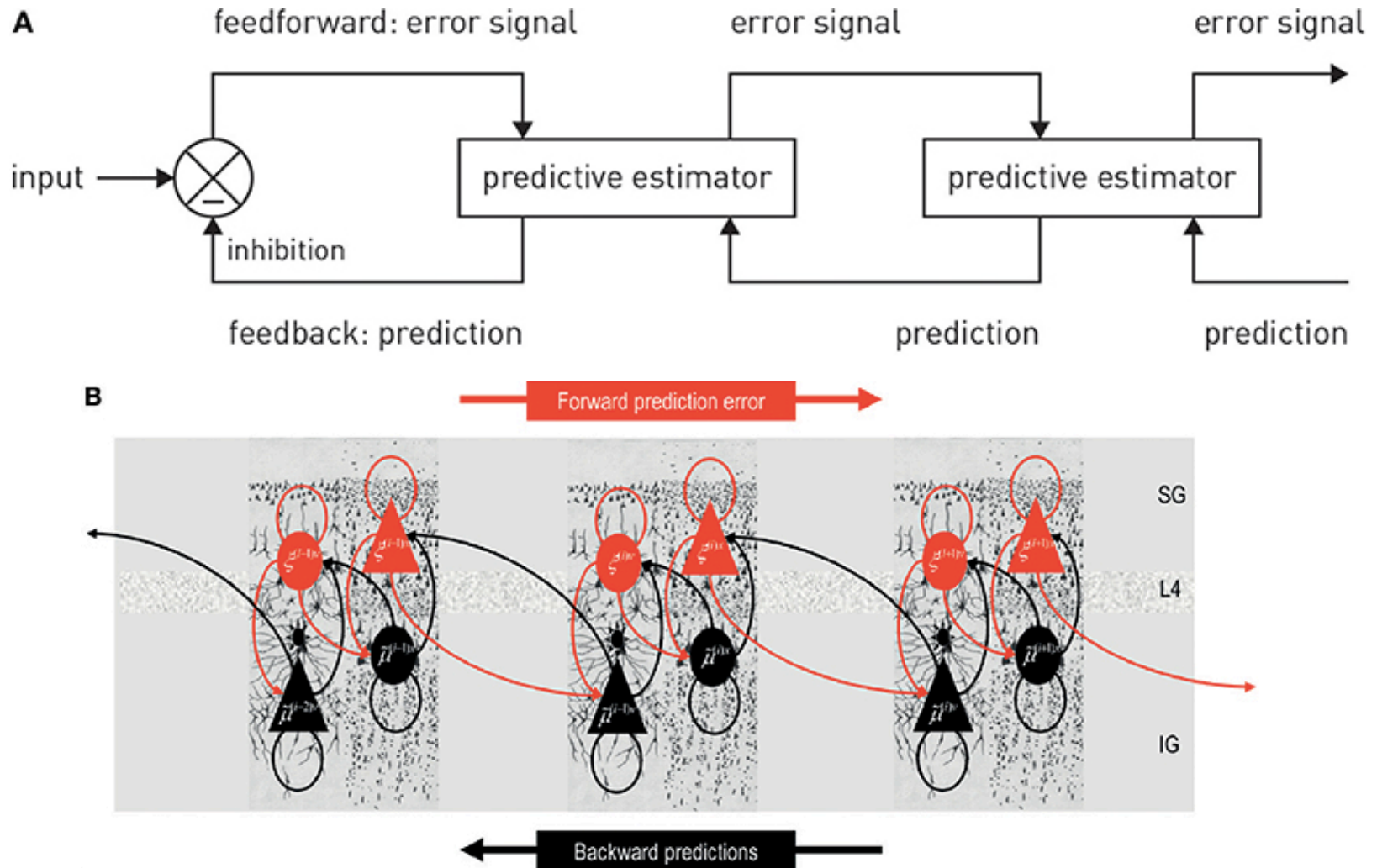
- **Adjust predictions to match inputs (inference and learning)**
- *Learning* means building a model of the time-invariant *structure* of the environment
- *Inference* means determining the time-dependent *state* of the environment.
- **Adjust inputs to match predictions (action)**
- Action that reduces surprise can be implemented by introducing the prediction that we will act such as to minimize surprise and then choosing actions such as to minimize surprise about these actions. This is *active inference*.

Predictive coding for learning and inference¹

- Biological organisms use sensory inputs to uncover (i.e., *infer*) the structure of their surroundings
- Sensory inputs are ambiguous and noisy
- Ambiguity and noise can be resolved using probabilistic predictions based on *prior information*
- These predictions help identify the *most likely* cause of sensory inputs (cf. Helmholtz)
- According to the *predictive coding framework* (sometimes also called the *predictive processing framework*), prediction based on prior information about the world is a *key feature of brain function*.

¹ Review: Teufel, C., & Fletcher, P. C. (2020). Forms of prediction in the nervous system. *Nature Reviews Neuroscience*, 1–12. <https://doi.org/10.1038/s41583-020-0275-5>

Hierarchical predictive coding in the brain¹



¹ Friston, K. (2005). A theory of cortical responses. *Philosophical Transactions of the Royal Society B: Biological Sciences*, 360(1456), 815–836. <https://doi.org/10.1098/rstb.2005.1622>

Hierarchical predictive coding in the brain

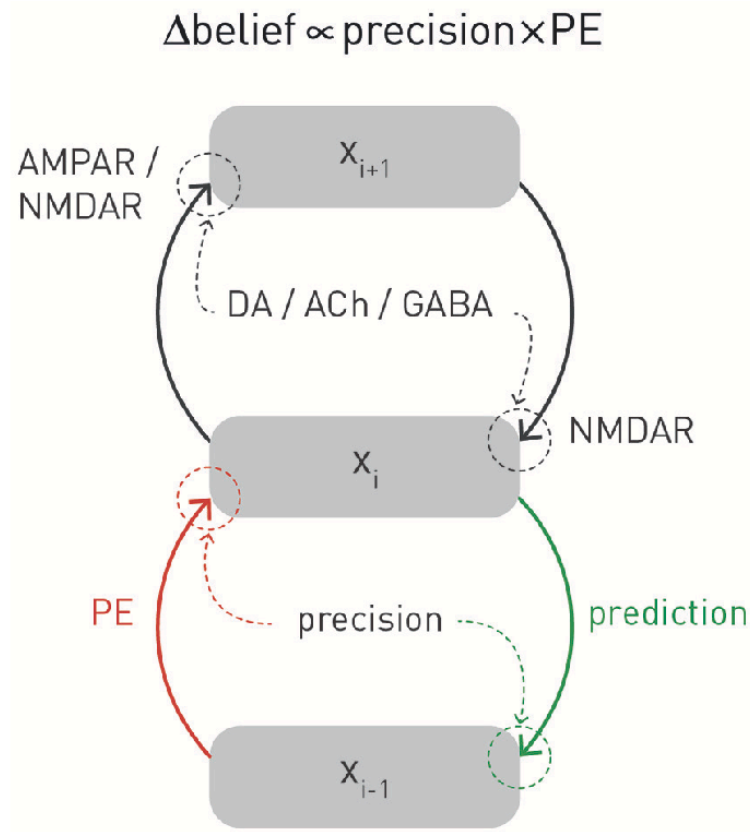
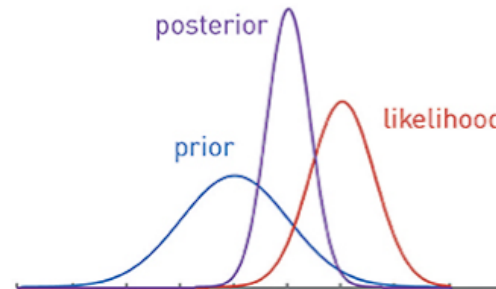


FIGURE 3 | A graphical summary of computational and physiological key components of hierarchical Bayesian inference. Computationally, prediction errors are conveyed by ascending or forward connections, while predictions are signaled via backward or descending connections. Critically, both experience a weighting by precision. Physiologically, the currently available evidence suggests that, in cortex, prediction errors are signaled via ionotropic glutamatergic receptors (AMPA and NMDA receptors), predictions via NMDA receptors, while precision-weighting is either implemented through neuromodulatory inputs (e.g., dopamine or acetylcholine) or by local GABAergic mechanisms. The figure is adapted, with permission, from Stephan et al. (2016b).

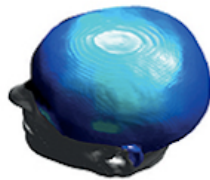
The rules of inference, plus action → surprise minimization

Bayes' Theorem

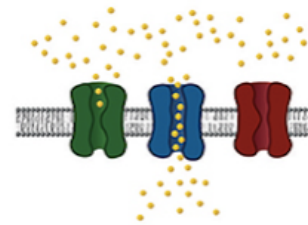
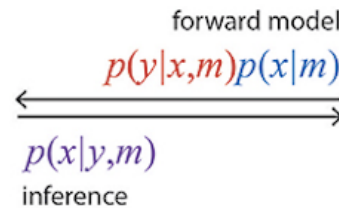
$$p(x|y,m) = \frac{p(y|x,m)p(x|m)}{p(y|m)}$$



Generative models as computational assays

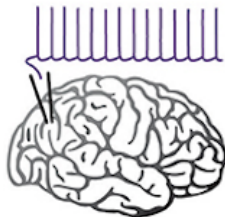


measurement y

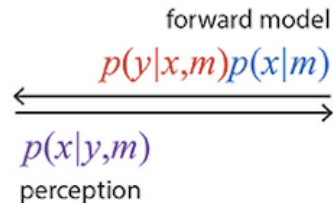


hidden system state x

Perception as the inversion of a generative model



neuronal activity y

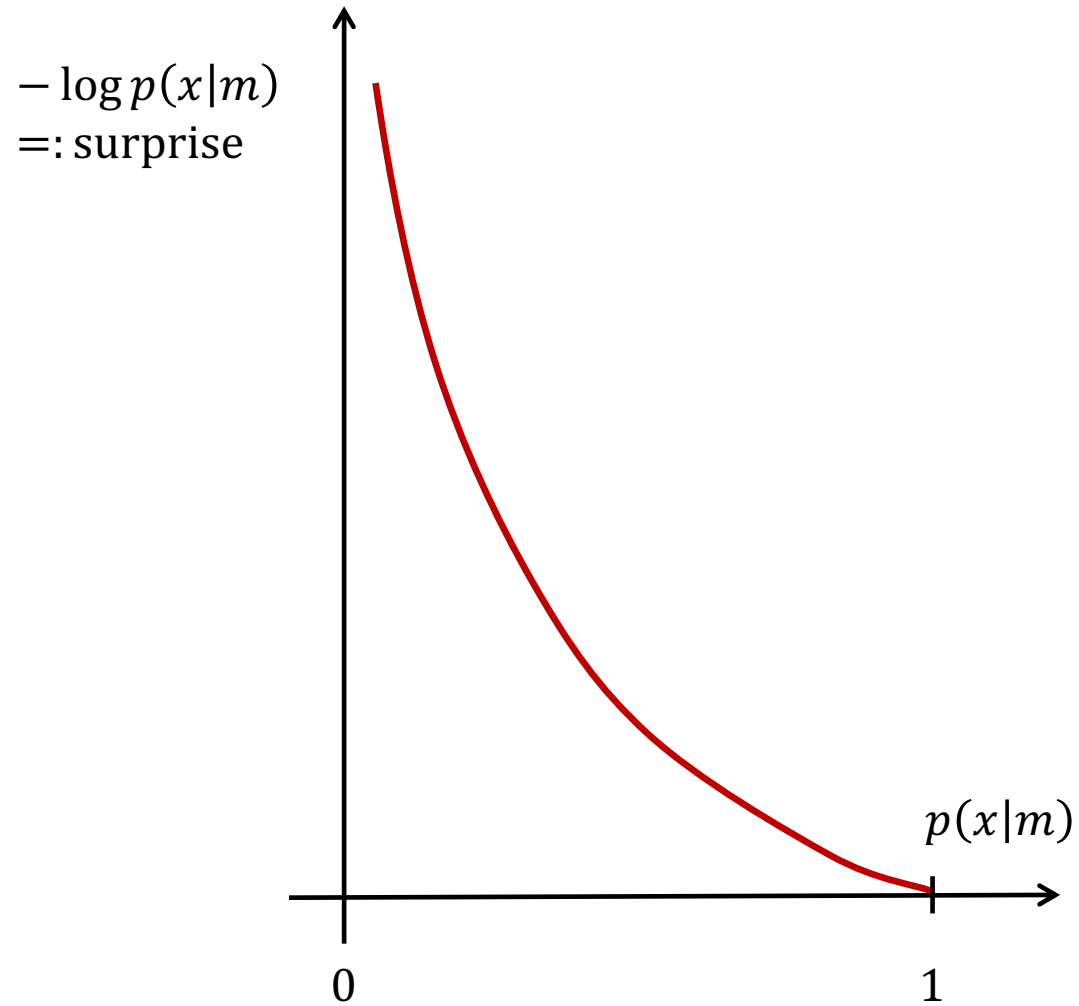


environmental state x

More on surprise

- How surprising an event is felt to be depends on its probability.
- It makes intuitive sense to take the negative logarithm of $p(x|m)$ as a measure of surprise.
- If $p(x|m) = 1$, the outcome was certain and there is no surprise at all ($-\log(p(x|m)) = 0$).
- If $p(x|m) = 0$, the outcome was impossible and surprise is infinite ($-\log(p(x|m)) = \infty$).
- In between, surprise is greater than zero and increases for less probable observations.

Surprise in a graph



Entropy

- A concept closely related to surprise is entropy
- The more ignorant we are about a quantity, the greater is the surprise we may expect when observing it.
- Expected surprise is called the **entropy** S of a probability distribution p :

$$S[p] := - \int p(x) \log p(x) \, dx$$

- Entropy is a **measure of ignorance**.
- Its name is due to an analogous quantity in thermodynamics.

Entropy example

- As a simple example, let's look at a coin toss.
- There are two possible outcomes: $x \in \{\text{heads}, \text{tails}\}$
- Since outcomes are discrete and binary, we use a sum instead of an integral and the binary logarithm to define the entropy:

$$S[p] := - \sum_y p(x) \log_2 p(x)$$

- For a fair coin (i.e., $p(\text{heads}) = p(\text{tails}) = \frac{1}{2}$), $S[p] = 1$
- However, for $p(\text{heads}) = \frac{9}{10}$, $p(\text{tails}) = \frac{1}{10}$, we get $S[p] \approx 0.47$ because expected surprise is much lower.

Free energy

In information theory, **free energy A is the surprise of a given model m at a particular (set of) observation(s) x :**

$$A := -\log p(x|m)$$

However, at least four kinds of free energy have to be kept apart:

- The free energy of thermodynamics
- The free energy of statistical physics
- Informational free energy (as above)
- Variational free energy (an upper bound on informational free energy)

Free energy in thermodynamics and statistical mechanics

Two kinds of **thermodynamic free energy**: *Gibbs* and *Helmholtz* (him again!)

Helmholtz free energy:

$$A := U - TS$$

U : internal energy; T : temperature; S : entropy

$$[\text{Gibbs free energy: } G := U + pV - TS]$$

In **statistical mechanics**, the Boltzmann distribution describes the **relation between energy and probability**. If particles in a system can be in states $s_1, s_2, s_3, \dots$ corresponding to energy levels $E_1, E_2, E_3, \dots$, then the probability p_i of finding a particle in state s_i is

$$p_i = \frac{\exp(-E_i/kT)}{\sum_j \exp(-E_j/kT)} = \frac{\exp(-E_i/kT)}{Z},$$

where T is *temperature* and k is *Boltzmann's constant*, and $Z := \sum_j \exp(-E_j/kT)$ is the *partition function*.

Free energy in thermodynamics and statistical mechanics

Taking the logarithm on both sides and rearranging gives us

$$-kT \log Z = E_i + kT \log p_i$$

Taking the expectation value on both sides, we get

$$\begin{aligned} \sum_i p_i (-kT \log Z) &= \sum_i p_i E_i + \sum_i p_i kT \log p_i \\ -kT \log Z &= \langle E \rangle - T \left(-k \sum_i p_i \log p_i \right) = \langle E \rangle - TS \end{aligned}$$

with entropy $S := -k \sum_i p_i \log p_i$.

In analogy to the definition $A := U - TS$ of Helmholtz free energy from thermodynamics, this motivates the definition

$$A := -kT \log Z$$

of free energy in statistical mechanics.

Informational free energy

Returning to information theory, we take the definition of informational free energy and perform a series of algebraic operations on it:

$$\begin{aligned} A &:= -\log p(x|m) = -\int p(\vartheta|x, m) \log p(x|m) d\vartheta \\ &= -\int p(\vartheta|x, m) \log \frac{p(x, \vartheta|m)}{p(\vartheta|x, m)} d\vartheta \\ &= \underbrace{-\int p(\vartheta|x, m) \log p(x, \vartheta|m) d\vartheta}_{=\langle E \rangle} - \underbrace{\left(-\int p(\vartheta|x, m) \log p(\vartheta|x, m) d\vartheta \right)}_{=S} \end{aligned}$$

This gives us an **information theoretic analogon** to the definition of Helmholtz free energy in statistical mechanics. Compare:

$$\begin{aligned} A &:= -kT \log Z = -kT \log \left(\sum_j \exp(-E_j/kT) \right) \\ A &:= -\log p(x|m) = -\log \int p(x, \vartheta|m) d\vartheta \end{aligned}$$

This is the same if we set $kT = 1$ and **interpret the negative logarithm of the joint probability distribution as an energy**:

$$E_{\vartheta} \equiv -\log p(x, \vartheta|m)$$

Variational free energy

The problem with informational free energy is that we cannot calculate it except in trivial cases. Whenever models are complicated enough to be interesting, the integrals involved are intractable.

$$A := \underbrace{- \int p(\vartheta|x, m) \log p(x, \vartheta|m) d\vartheta}_{=\langle E \rangle} - \underbrace{\left(- \int p(\vartheta|x, m) \log p(\vartheta|x, m) d\vartheta \right)}_{=S}$$

The solution to this is variational free energy, where we replace the true posterior $p(\vartheta|x, m)$ by an approximation $q(\vartheta)$:

$$A_v := \underbrace{- \int q(\vartheta) \log p(x, \vartheta|m) d\vartheta}_{:=E_v} - \underbrace{\left(- \int q(\vartheta) \log q(\vartheta) d\vartheta \right)}_{:=S_v}$$

Variational free energy

What makes variational free energy A_v such an extremely useful concept is the following theorem:

$$A_v \geq A \text{ for all } q(\vartheta)$$

This means that **whatever $q(\vartheta)$ we plug into A_v , we get an A_v that is greater than A** . So without having to know anything about A , we can vary $q(\vartheta)$ such that it minimizes A_v .

$$A_v := - \int q(\vartheta) \log p(x, \vartheta | m) d\vartheta + \int q(\vartheta) \log q(\vartheta) d\vartheta$$

The branch of mathematics that describes how to carry out the minimization of A_v with respect to $q(\vartheta)$ is called **variational calculus**, hence “variational” free energy.

Minimizing A_v with respect to $q(\vartheta)$ leads to an approximation of $p(\vartheta | x, m)$ by $q(\vartheta)$ because of the theorem above and because $A_v = A$ for $q(\vartheta) = p(\vartheta | x, m)$.

The remarkable thing here is that we can use variational calculus to find a $q(\vartheta)$ that approximates $p(\vartheta | x, m)$ **without ever having to know $p(\vartheta | x, m)$ itself**.

This is how the brain can build, update, and compare models of the world without ever “seeing behind the scenes” of its sensory input.

Variational free energy

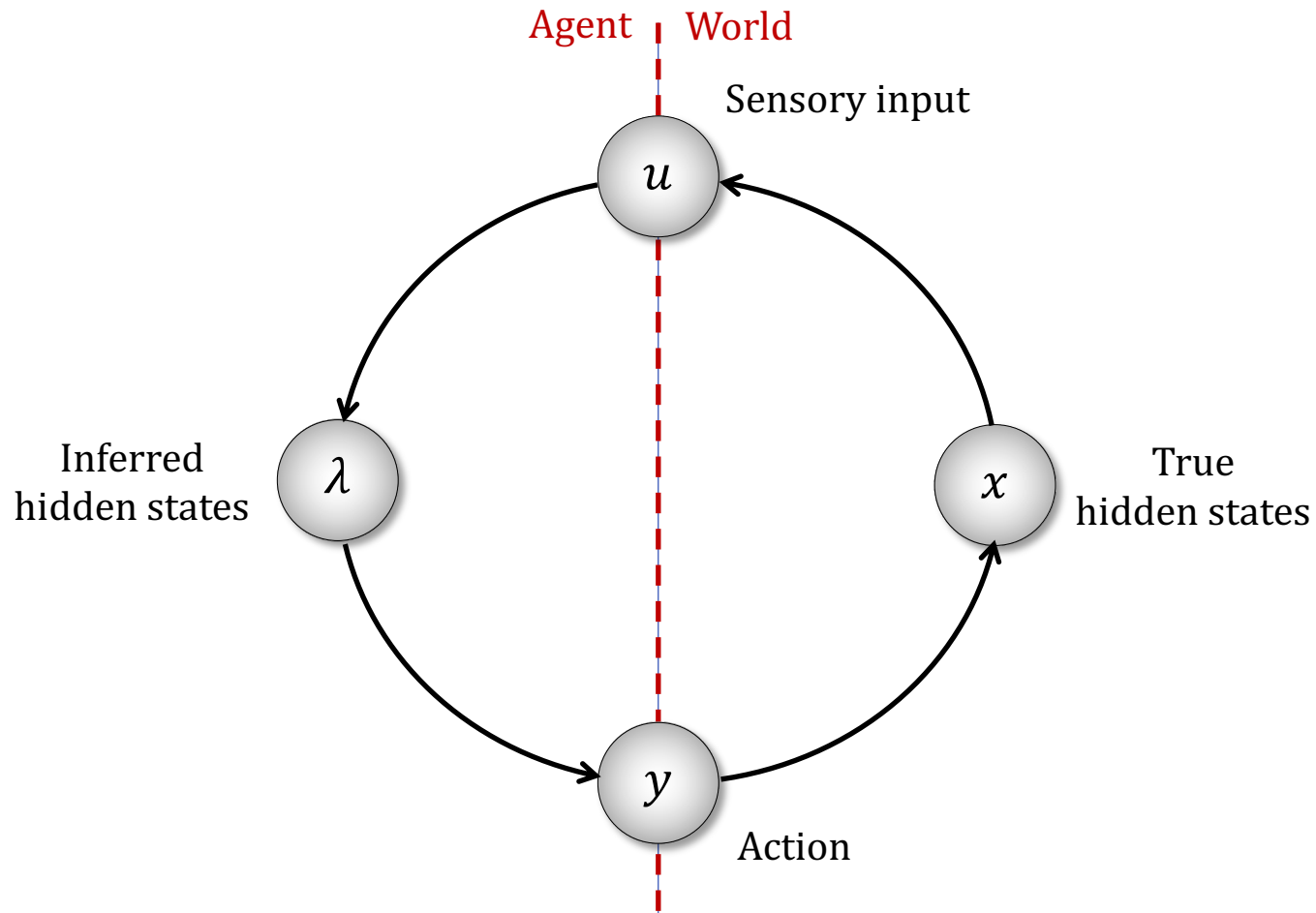
Proof that $A_v \geq A$ for all $q(\vartheta)$:

$$\begin{aligned} A &:= -\log p(x|m) \\ &= -\log \int p(x, \vartheta|m) d\vartheta \\ &= -\log \int q(\vartheta) \frac{p(x, \vartheta|m)}{q(\vartheta)} d\vartheta \\ &\stackrel{\text{Jensen's inequality}}{\leq} -\int q(\vartheta) \log \frac{p(x, \vartheta|m)}{q(\vartheta)} d\vartheta \\ &= -\int q(\vartheta) \log p(x, \vartheta|m) d\vartheta + \int q(\vartheta) \log q(\vartheta) d\vartheta \\ &=: A_v \end{aligned}$$

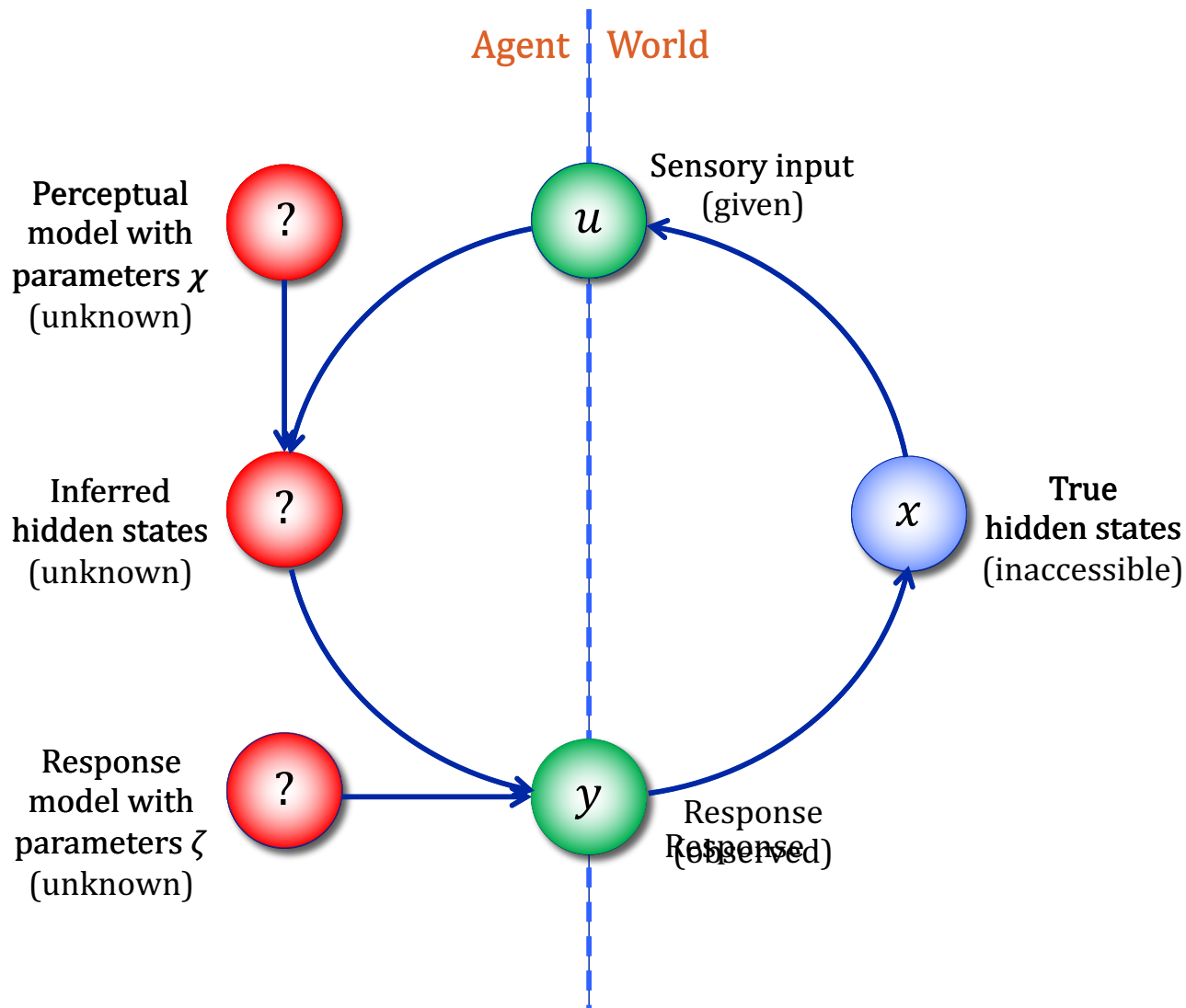
□

Jensen's inequality

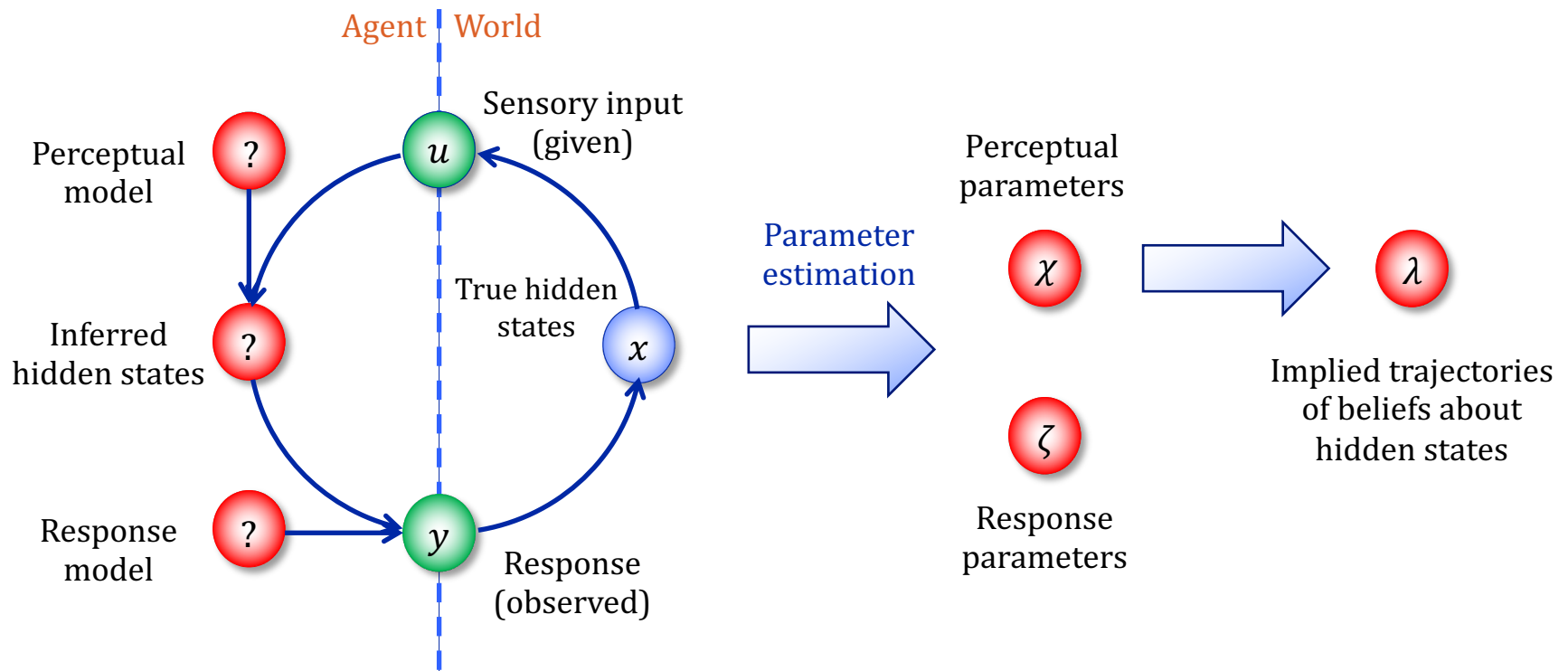
Agents and the world: active inference



Active inference: Minimizing surprise about the world *and* about my own actions



Application to experimental data: parameter estimation



Active inference: minimizing surprise/free energy *including about my own actions*

- We have seen that minimizing surprise is equivalent to doing Bayesian inference
- Often it is simpler to do inference (and act!) implicitly by minimizing surprise
- Active inference amounts to taking the actions that will make it most likely that I will receive the input I expect to receive. Other inputs would be surprising, and seeing myself act in a way that fails to minimize surprising inputs would also be surprising.
- This provides a single framework for inference and action, sidestepping the dichotomy between perceptual and decision models.
- The dark room problem doesn't arise because an absence of low-level surprises is surprising at the meta-level.

Remarks

- ‘But I like surprises!’ - The *dark room problem* doesn’t arise because receiving no inputs would be surprising. In a hierarchical formulation, lack of low-level surprise leads to high-level surprise, and to action leading to more low-level surprise. In a time-sensitive formulation, acting to decrease long-term surprise may lead to an increase in short-term surprise.
- This is a *framework*, not a theory. Trying to falsify it or criticizing it for being unfalsifiable is beside the point.
- Example of a framework: Newton’s 2nd law states that force is mass times acceleration. ‘Falsifying’ this is beside the point because its value does not lie in being true but in enabling us to do mechanics, i.e., to formulate mechanical theories which can be falsified. Progress in frameworks happens when more powerful frameworks are developed (Lagrangian mechanics, Hamiltonian mechanics, etc.).
- Active inference and the free energy principle are tools to describe and build agents negotiating an uncertain environment.

An example of active inference:

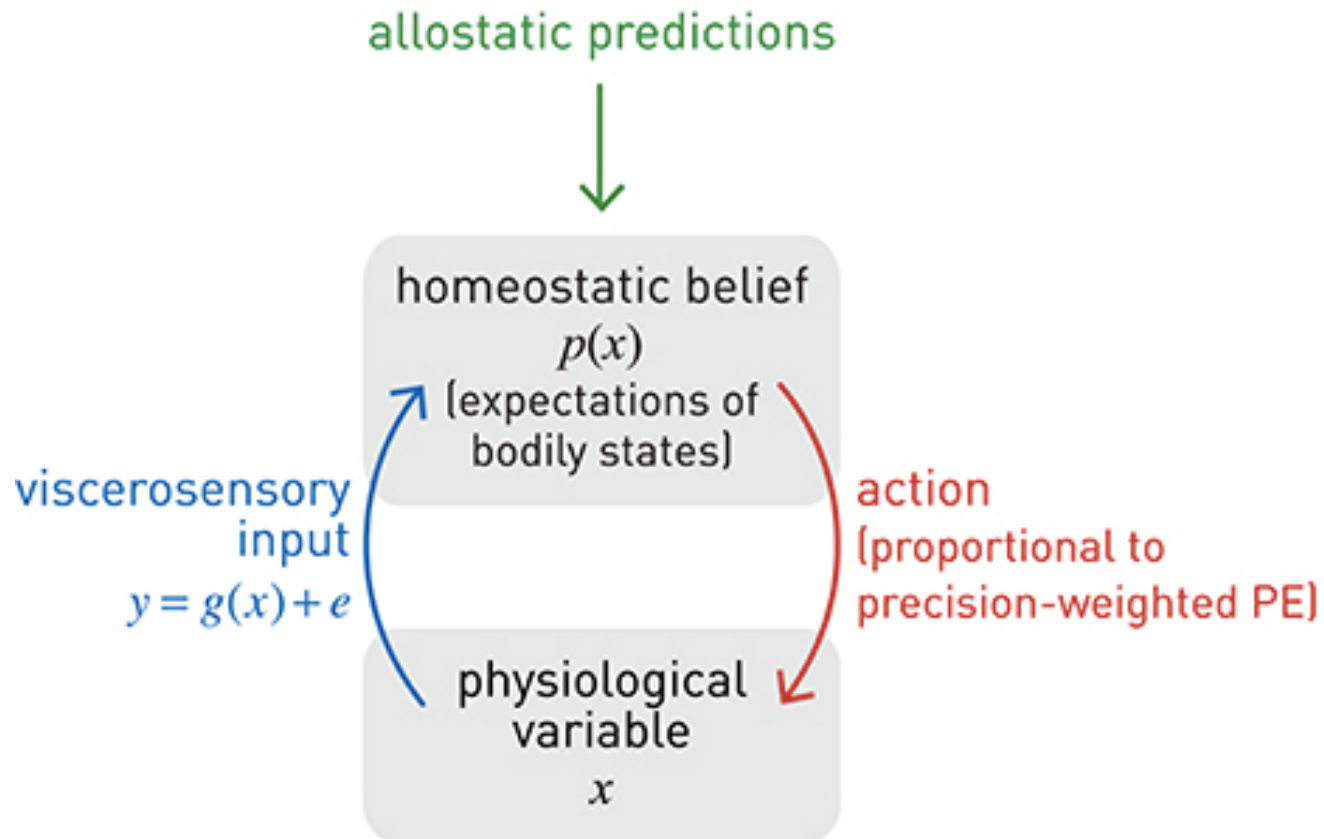
The allostatic reflex arc

Allostatic Self-efficacy: A Metacognitive Theory of Dyshomeostasis-Induced Fatigue and Depression

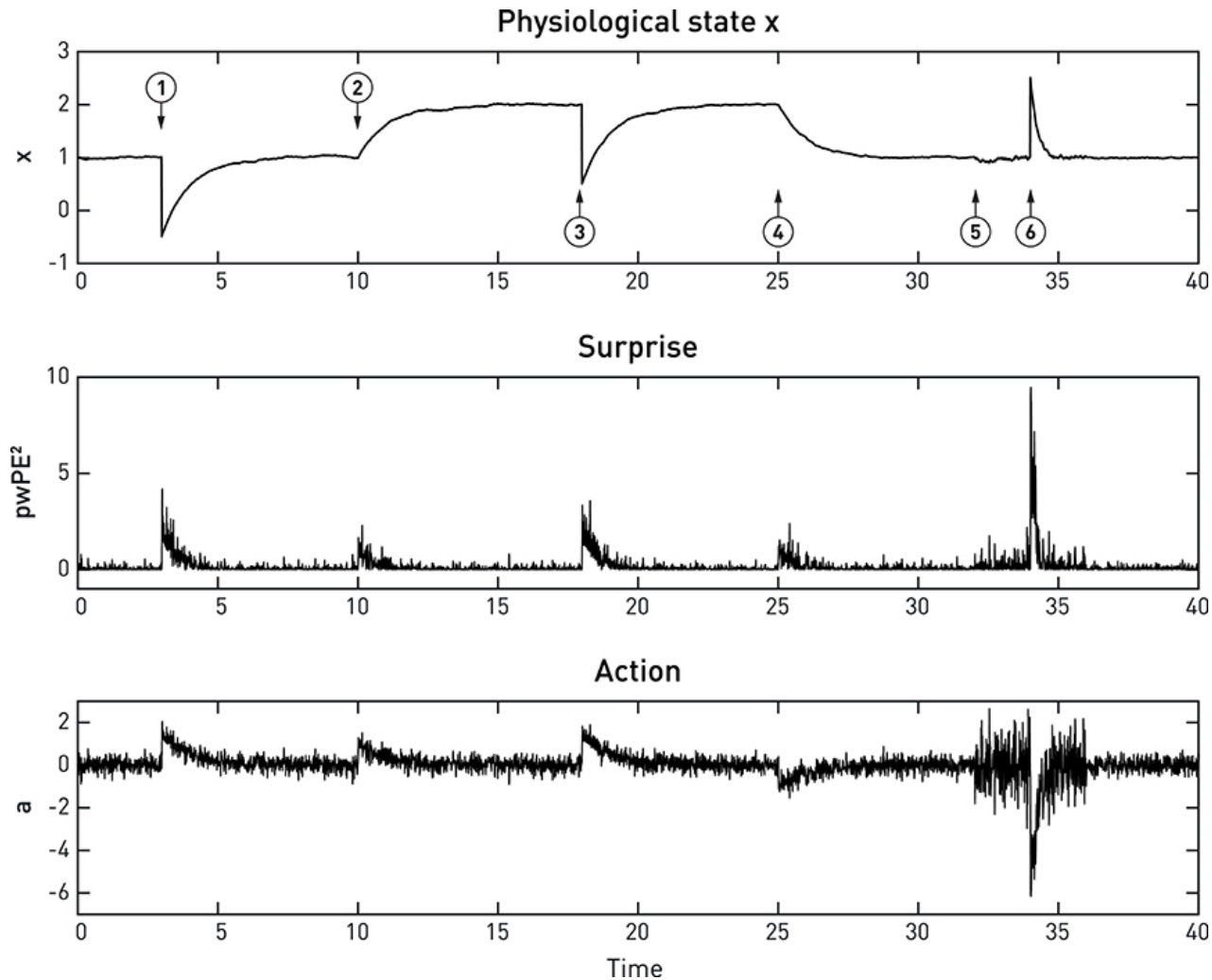
Klaas E. Stephan^{1,2,3*}, Zina M. Manjaly^{1,4}, Christoph D. Mathys², Lilian A. E. Weber¹, Saeed Paliwal¹, Tim Gard^{1,5}, Marc Tittgemeyer³, Stephen M. Fleming², Helene Haker¹, Anil K. Seth⁶ and Frederike H. Petzschner¹

¹ Translational Neuromodeling Unit, Institute for Biomedical Engineering, University of Zurich and ETH Zurich, Zurich, Switzerland, ² Wellcome Trust Centre for Neuroimaging, University College London, London, UK, ³ Max Planck Institute for Metabolism Research, Cologne, Germany, ⁴ Department of Neurology, Schulthess Clinic, Zurich, Switzerland, ⁵ Center for Complementary and Integrative Medicine, University Hospital Zurich, Zurich, Switzerland, ⁶ Sackler Centre for Consciousness Science, School of Engineering and Informatics, University of Sussex, Brighton, UK

Active inference: the allostatic reflex arc



Active inference: the allostatic reflex arc



Stephan et al. (2016). <https://doi.org/10.3389/fnhum.2016.00550>

Active inference: the allostatic reflex arc

Model:

$$p(x) = N(x; \mu_t, \pi_t^{-1})$$

$$p(y|x) = N(y; g(x), \pi_{data}^{-1})$$

Inference (i.e., belief update):

$$\mu_{t+1} = \mu_t + \frac{\pi_{data}}{\pi_t + \pi_{data}} (y_t - g(\mu_t))$$

$$\pi_{t+1} = \pi_t + \pi_{data}$$

Active inference: the allostatic reflex arc

Delta priors on mean and precision of state:

$$\begin{aligned} p(y|m_H) &= \int p(y|\mu_t, \pi_t) p(\mu_t) p(\pi_t) d\mu_t d\pi_t \\ &= \int N(y; g(\mu_t), \pi_t^{-1}) \delta(\mu_t - \mu_{prior}) \\ &\quad \delta(\pi_t - \pi_{prior}) d\mu_t d\pi_t \\ &= N(y; g(\mu_{prior}), \pi_{prior}^{-1}) \end{aligned}$$

Negative of log-evidence L is Shannon surprise S :

$$\begin{aligned} L &= \ln p(y|m_H) & S(y|m_H) &= -L \\ &= \frac{1}{2} \left(\ln \pi_{prior} - \pi_{prior} (y - g(\mu_{prior}))^2 \right) + c \\ &= \frac{1}{2} \left(\ln \pi_{prior} - \pi_{prior} (PE(y))^2 \right) + c \end{aligned}$$

Active inference: the allostatic reflex arc

Definition of action:

$$\begin{aligned} a(t) &= \frac{\partial L}{\partial x} \\ &= \frac{\partial \left[-\frac{1}{2} \pi_{prior} (y(x, t) - g(\mu_{prior}))^2 \right]}{\partial x} \\ &= \frac{\partial \left[-\frac{1}{2} \pi_{prior} PE(y)^2 \right]}{\partial x} \\ &= -\frac{\pi_{prior}}{2} \frac{\partial [PE(y)^2]}{\partial y} \frac{\partial y}{\partial x} \\ &= -\pi_{prior} PE(y) \frac{\partial g}{\partial x} \end{aligned}$$

precision-weighted prediction error

Action induces gradient descent on surprise S:

$$\frac{dx}{dt} = \lambda^{-1} f(a(t)) \quad \frac{dx}{dt} = -\lambda^{-1} f\left(\frac{\partial S}{\partial x}\right)$$

Exercise

- Reproduce the simulation in the simulation in the previous slide (Figure 6 from Stephan et al., 2016) in a programming language of your choice.
- You have the solution coded in Matlab. You can use this as the basis for your own code and to help your understanding of the simulation presented in the paper.
- What happens when you vary the input precision π_{data} ?
- What happens when you vary the mean and precision of the homeostatic setpoint?

The free-energy principle: a rough guide to the brain?

Karl Friston

The Wellcome Trust Centre for Neuroimaging, University College London, Queen Square, London WC1N 3BG, UK

This article reviews a free-energy formulation that advances Helmholtz's agenda to find principles of brain function based on conservation laws and neuronal energy. It rests on advances in statistical physics, theoretical biology and machine learning to explain a remarkable range of facts about brain structure and function. We could have just scratched the surface of what this formulation offers; for example, it is becoming clear that the Bayesian brain is just one facet of the free-energy principle and that perception is an inevitable consequence of active exchange with the environment. Furthermore, one can see easily how constructs like memory, attention, value, reinforcement and salience might disclose their simple relationships within this framework.

The motivation for the free-energy principle is again simple but fundamental. It rests upon the fact that self-organising biological agents resist a tendency to disorder and therefore minimize the entropy of their sensory states [4]. Under ergodic assumptions, this entropy is:

$$\begin{aligned} H(y) &= - \int p(y|m) \ln p(y|m) dy \\ &= \lim_{T \rightarrow \infty} \frac{1}{T} \int_0^T - \ln p(y|m) dt \end{aligned} \quad (\text{Equation 1})$$

See [Box 1](#) for an explanation of the variables and Ref. [5] for details. This equation (Equation 1) means that mini-

Reinforcement Learning or Active Inference?

Karl J. Friston*, Jean Daunizeau, Stefan J. Kiebel

The Wellcome Trust Centre for Neuroimaging, University College London, London, United Kingdom

Abstract

This paper questions the need for reinforcement learning or control theory when optimising behaviour. We show that it is fairly simple to teach an agent complicated and adaptive behaviours using a free-energy formulation of perception. In this formulation, agents adjust their internal states and sampling of the environment to minimize their free-energy. Such agents learn causal structure in the environment and sample it in an adaptive and self-supervised fashion. This results in behavioural policies that reproduce those optimised by reinforcement learning and dynamic programming. Critically, we do not need to invoke the notion of reward, value or utility. We illustrate these points by solving a benchmark problem in dynamic programming; namely the mountain-car problem, using active perception or inference under the free-energy principle. The ensuing proof-of-concept may be important because the free-energy formulation furnishes a unified account of both action and perception and may speak to a reappraisal of the role of dopamine in the brain.

Citation: Friston KJ, Daunizeau J, Kiebel SJ (2009) Reinforcement Learning or Active Inference?. PLoS ONE 4(7): e6421. doi:10.1371/journal.pone.0006421

Discussion Paper

Active inference and epistemic value

**Karl Friston¹, Francesco Rigoli¹, Dimitri Ognibene², Christoph Mathys^{1,3,4},
Thomas Fitzgerald¹, and Giovanni Pezzulo⁵**

¹The Wellcome Trust Centre for Neuroimaging, Institute of Neurology, London, UK

²Centre for Robotics Research, Department of Informatics, King's College London, London, UK

³Translational Neuromodeling Unit (TNU), Institute for Biomedical Engineering, University of Zürich and ETH Zürich, Zürich, Switzerland

⁴Laboratory for Social and Neural Systems Research (SNS Lab), Department of Economics, University of Zürich, Zürich, Switzerland

⁵Institute of Cognitive Sciences and Technologies, National Research Council, Rome, Italy

Thanks!