

On deriving efficient information-seeking behaviour for intelligent autonomous systems

Applied Mathematics Seminar, McGill University

October 20, 2025

Wouter Kouw

Assistant professor — Bayesian Intelligent Autonomous Systems lab

Department of Electrical Engineering

Outline

- Problem: how to make robots more autonomous?

Outline

- Problem: how to make robots more autonomous?
- Technique: free energy minimization by message passing

Outline

- Problem: how to make robots more autonomous?
- Technique: free energy minimization by message passing
- Example: ARxl - Autoregressive Reactive Inference agent

Outline

- Problem: how to make robots more autonomous?
- Technique: free energy minimization by message passing
- Example: ARxl - Autoregressive Reactive Inference agent
- Discussion: effects of controlling information gain

Problem

How can we make robots more autonomous?



(a) MTC, <https://tinyurl.com/y45mksa6>

Problem

How can we make robots more autonomous?



(a) MTC, <https://tinyurl.com/y45mksa6>



(b) Postmates, <https://tinyurl.com/5442pph6>

Problem

How can we make robots more autonomous?



(a) MTC, <https://tinyurl.com/y45mksa6>



(b) Postmates, <https://tinyurl.com/5442pph6>



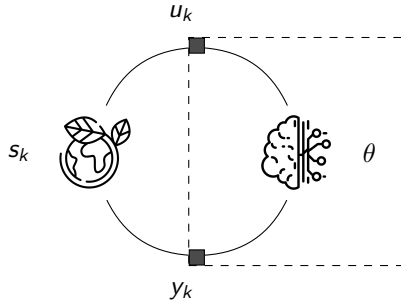
(c) MPI IS, <https://tinyurl.com/5n8sd6ps>

Problem

Let s_k be the state of a system at time t_k . An agent is an entity that models the system using parameters θ based on noisy partial observations y_k , and acts upon the world with u_k .

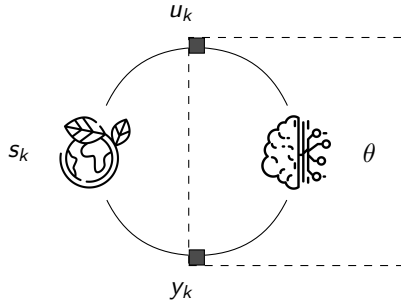
Problem

Let s_k be the state of a system at time t_k . An agent is an entity that models the system using parameters θ based on noisy partial observations y_k , and acts upon the world with u_k .



Problem

Let s_k be the state of a system at time t_k . An agent is an entity that models the system using parameters θ based on noisy partial observations y_k , and acts upon the world with u_k .



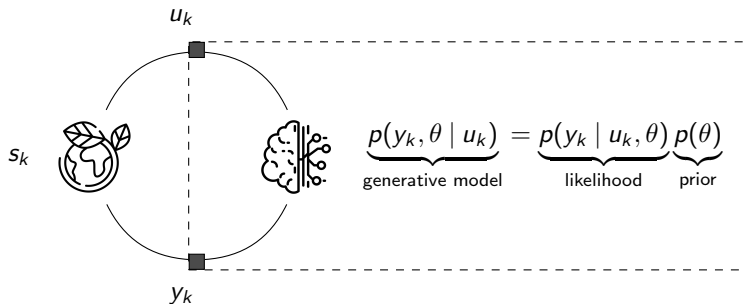
How can intelligent agents continually and efficiently adapt to uncertain environments?

Free Energy Principle

I approach this problem from a probabilistic perspective.

Free Energy Principle

I approach this problem from a probabilistic perspective.



Free Energy Principle

We find intelligent agents in nature. Modern theories on neural information processing argue that agents pursue a single objective function: a free energy functional.

Friston (2010), The free-energy principle: a unified brain theory?, Nature.

Free Energy Principle

We find intelligent agents in nature. Modern theories on neural information processing argue that agents pursue a single objective function: a free energy functional.

Free energy functional

Given the generative model p and a variational model q , the free energy functional is

$$\mathcal{F}[q] = \int q(\theta) \ln \frac{q(\theta)}{p(y_k, \theta | u_k)} d\theta .$$

Friston (2010), The free-energy principle: a unified brain theory?, Nature.

Free Energy Principle

We find intelligent agents in nature. Modern theories on neural information processing argue that agents pursue a single objective function: a free energy functional.

Free energy functional

Given the generative model p and a variational model q , the free energy functional is

$$\mathcal{F}[q] = \int q(\theta) \ln \frac{q(\theta)}{p(y_k, \theta | u_k)} d\theta .$$

The variational model q represents a trade-off between accuracy and computational effort.

Friston (2010), The free-energy principle: a unified brain theory?, Nature.

Variational free energy minimization

Inference is the process of finding the optimal variational distribution:

$$q^*(\theta) = \arg \min_{q \in Q} \mathcal{F}[q].$$

Variational free energy minimization

Inference is the process of finding the optimal variational distribution:

$$q^*(\theta) = \arg \min_{q \in Q} \mathcal{F}[q] .$$

This is done by minimizing the variational derivative over the free energy functional with a normalization constraint:

$$\mathcal{L}[q, \gamma] = \mathcal{F}[q] + \gamma \left(\int q(\theta) d\theta - 1 \right) .$$

Variational free energy minimization

Inference is the process of finding the optimal variational distribution:

$$q^*(\theta) = \arg \min_{q \in Q} \mathcal{F}[q] .$$

This is done by minimizing the variational derivative over the free energy functional with a normalization constraint:

$$\mathcal{L}[q, \gamma] = \mathcal{F}[q] + \gamma \left(\int q(\theta) d\theta - 1 \right) .$$

More complex models also have factorization, parametrization and/or chance constraints.

Senoz et al. (2021), Variational message passing and local constraint manipulation in factor graphs, Entropy.

Variational free energy minimization

Let $\delta q(\theta) = \epsilon \phi(\theta)$ be a variation with ϕ a continuous and differentiable test function. Then the variational derivative can be found with:

Variational free energy minimization

Let $\delta q(\theta) = \epsilon \phi(\theta)$ be a variation with ϕ a continuous and differentiable test function. Then the variational derivative can be found with:

$$\int \frac{\delta \mathcal{L}[q, \gamma]}{\delta q} \phi(\theta) d\theta = \left. \frac{d\mathcal{L}[q(\theta) + \epsilon \phi(\theta), \gamma]}{d\epsilon} \right|_{\epsilon=0}$$

Variational free energy minimization

Let $\delta q(\theta) = \epsilon \phi(\theta)$ be a variation with ϕ a continuous and differentiable test function. Then the variational derivative can be found with:

$$\begin{aligned} \int \frac{\delta \mathcal{L}[q, \gamma]}{\delta q} \phi(\theta) d\theta &= \left. \frac{d\mathcal{L}[q(\theta) + \epsilon \phi(\theta), \gamma]}{d\epsilon} \right|_{\epsilon=0} \\ &= \frac{d}{d\epsilon} \int (q(\theta) + \epsilon \phi(\theta)) \ln \frac{q(\theta) + \epsilon \phi(\theta)}{p(y_k, \theta | u_k)} d\theta + \frac{d}{d\epsilon} \gamma \int (q(\theta) + \epsilon \phi(\theta)) d\theta \Big|_{\epsilon=0} \end{aligned}$$

Variational free energy minimization

Let $\delta q(\theta) = \epsilon \phi(\theta)$ be a variation with ϕ a continuous and differentiable test function. Then the variational derivative can be found with:

$$\begin{aligned}\int \frac{\delta \mathcal{L}[q, \gamma]}{\delta q} \phi(\theta) d\theta &= \left. \frac{d\mathcal{L}[q(\theta) + \epsilon \phi(\theta), \gamma]}{d\epsilon} \right|_{\epsilon=0} \\ &= \frac{d}{d\epsilon} \int (q(\theta) + \epsilon \phi(\theta)) \ln \frac{q(\theta) + \epsilon \phi(\theta)}{p(y_k, \theta | u_k)} d\theta + \frac{d}{d\epsilon} \gamma \int (q(\theta) + \epsilon \phi(\theta)) d\theta \Big|_{\epsilon=0} \\ &= \int (\phi(\theta) \ln \frac{q(\theta)}{p(y_k, \theta | u_k)} + \phi(\theta)) d\theta + \gamma \int \phi(\theta) d\theta\end{aligned}$$

Variational free energy minimization

Let $\delta q(\theta) = \epsilon \phi(\theta)$ be a variation with ϕ a continuous and differentiable test function. Then the variational derivative can be found with:

$$\begin{aligned}\int \frac{\delta \mathcal{L}[q, \gamma]}{\delta q} \phi(\theta) d\theta &= \left. \frac{d\mathcal{L}[q(\theta) + \epsilon \phi(\theta), \gamma]}{d\epsilon} \right|_{\epsilon=0} \\&= \frac{d}{d\epsilon} \int (q(\theta) + \epsilon \phi(\theta)) \ln \frac{q(\theta) + \epsilon \phi(\theta)}{p(y_k, \theta | u_k)} d\theta + \frac{d}{d\epsilon} \gamma \int (q(\theta) + \epsilon \phi(\theta)) d\theta \Big|_{\epsilon=0} \\&= \int \left(\phi(\theta) \ln \frac{q(\theta)}{p(y_k, \theta | u_k)} + \phi(\theta) \right) d\theta + \gamma \int \phi(\theta) d\theta \\&= \int \left(\ln \frac{q(\theta)}{p(y_k, \theta | u_k)} + 1 + \gamma \right) \phi(\theta) d\theta.\end{aligned}$$

Variational free energy minimization

Setting the variational derivative to 0, yields

Variational free energy minimization

Setting the variational derivative to 0, yields

$$0 = \ln \frac{q(\theta)}{p(y_k, \theta | u_k)} + 1 + \gamma$$

Variational free energy minimization

Setting the variational derivative to 0, yields

$$0 = \ln \frac{q(\theta)}{p(y_k, \theta | u_k)} + 1 + \gamma$$
$$q(\theta) = \exp(-\gamma - 1)p(y_k, \theta | u_k).$$

Variational free energy minimization

Setting the variational derivative to 0, yields

$$0 = \ln \frac{q(\theta)}{p(y_k, \theta | u_k)} + 1 + \gamma$$
$$q(\theta) = \exp(-\gamma - 1) p(y_k, \theta | u_k).$$

Plugging this into the constraint gives $\exp(-\gamma - 1) = 1 / \int p(y_k, \theta | u_k) d\theta$, which means

$$q(\theta) = \frac{p(y_k | u_k, \theta) p(\theta)}{\int p(y_k | u_k, \theta) p(\theta) d\theta}.$$

Factor graphs

Factor graphs

Using the factorization structure, the free energy functional can be decomposed to:

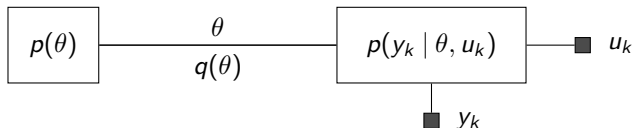
$$\mathcal{F}[q] = \int q(\theta) [-\ln p(y_k | u_k, \theta)] d\theta + \int q(\theta) \ln p(\theta) d\theta - \int q(\theta) \ln q(\theta) d\theta.$$

Factor graphs

Using the factorization structure, the free energy functional can be decomposed to:

$$\mathcal{F}[q] = \int q(\theta) [-\ln p(y_k | u_k, \theta)] d\theta + \int q(\theta) \ln p(\theta) d\theta - \int q(\theta) \ln q(\theta) d\theta.$$

This can be matched to a factor graph structure:

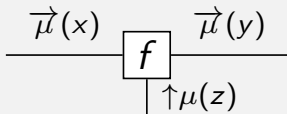


Loeliger (2004), An introduction to factor graphs, IEEE Signal Processing Magazine.

Message passing

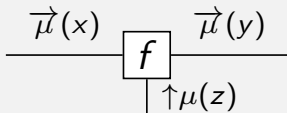
Message passing

Aggregating incoming messages:

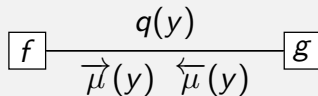


Message passing

Aggregating incoming messages:

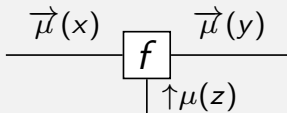


Updating marginals by product of messages:

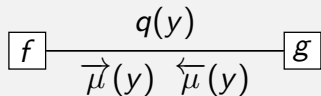


Message passing

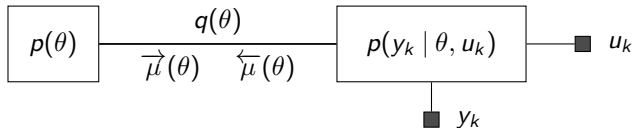
Aggregating incoming messages:



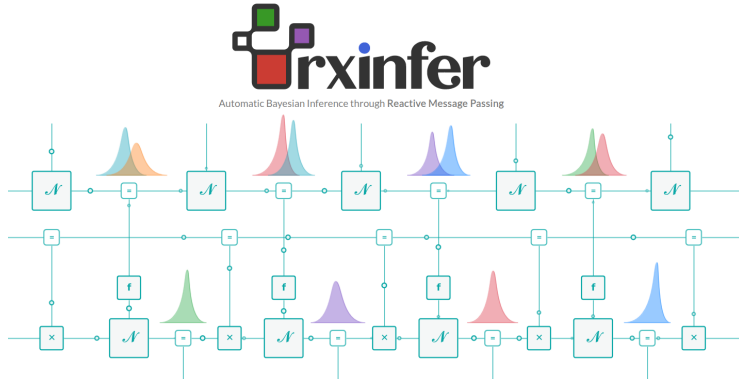
Updating marginals by product of messages:



For the example graph, this is:



At BIASlab in Eindhoven, we built an open-source state-of-the-art toolbox for this:

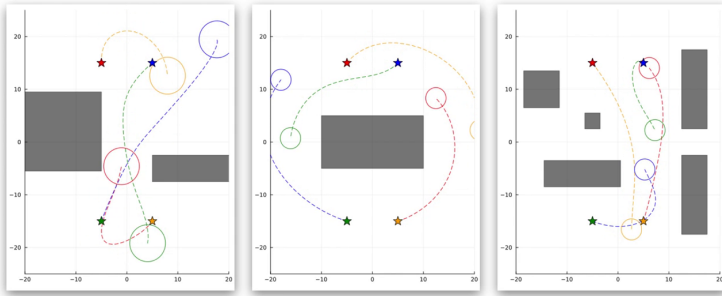


Free energy minimizing systems

Using these techniques, we construct all sorts of signal processing and control systems.

Free energy minimizing systems

Using these techniques, we construct all sorts of signal processing and control systems.



ARxl: Autoregressive Reactive Inference agent

ARxl is an example of an agent that minimizes free energy by message passing.

ARxI: Autoregressive Reactive Inference agent

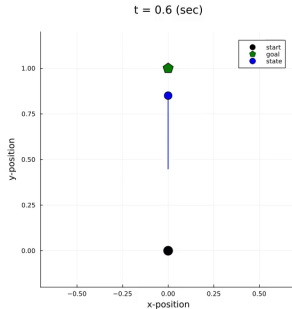
ARxI is an example of an agent that minimizes free energy by message passing.

I will show how ARxI plans and learns to navigate;

ARxl: Autoregressive Reactive Inference agent

ARxl is an example of an agent that minimizes free energy by message passing.

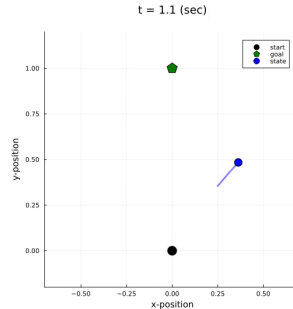
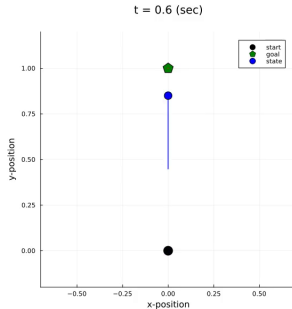
I will show how ARxl plans and learns to navigate;



ARxl: Autoregressive Reactive Inference agent

ARxl is an example of an agent that minimizes free energy by message passing.

I will show how ARxl plans and learns to navigate;



ARxl: Generative model specification

Let the likelihood of $y_k \in \mathbb{R}^{D_y}$, given $u_k \in \mathbb{R}^{D_u}$, be Gaussian distributed,

Nisslbeck & Kouw (2025), Online Bayesian system identification in multivariate autoregressive models via message passing, ECC.

ARxl: Generative model specification

Let the likelihood of $y_k \in \mathbb{R}^{D_y}$, given $u_k \in \mathbb{R}^{D_u}$, be Gaussian distributed,

$$p(y_k \mid \Theta, u_k, \bar{u}_k, \bar{y}_k) = \mathcal{N}(y_k \mid A^\top \begin{bmatrix} \bar{y}_k \\ u_k \\ \bar{u}_k \end{bmatrix}, W^{-1}),$$

Nisslbeck & Kouw (2025), Online Bayesian system identification in multivariate autoregressive models via message passing, ECC.

ARxl: Generative model specification

Let the likelihood of $y_k \in \mathbb{R}^{D_y}$, given $u_k \in \mathbb{R}^{D_u}$, be Gaussian distributed,

$$p(y_k \mid \Theta, u_k, \bar{u}_k, \bar{y}_k) = \mathcal{N}(y_k \mid A^\top \begin{bmatrix} \bar{y}_k \\ u_k \\ \bar{u}_k \end{bmatrix}, W^{-1}),$$

where $\Theta = \{A, W : A \in \mathbb{R}^{D_u \times D_y}, W \in \mathbb{R}^{D_y \times D_y}\}$, $\bar{u}_k = [u_{k-1} \dots u_{k-n_u}]^\top$, $\bar{y}_k = [y_{k-1} \dots y_{k-n_y}]^\top$.

Nisslbeck & Kouw (2025), Online Bayesian system identification in multivariate autoregressive models via message passing, ECC.

ARxl: Generative model specification

Let the likelihood of $y_k \in \mathbb{R}^{D_y}$, given $u_k \in \mathbb{R}^{D_u}$, be Gaussian distributed,

$$p(y_k \mid \Theta, u_k, \bar{u}_k, \bar{y}_k) = \mathcal{N}(y_k \mid A^\top \begin{bmatrix} \bar{y}_k \\ u_k \\ \bar{u}_k \end{bmatrix}, W^{-1}),$$

where $\Theta = \{A, W : A \in \mathbb{R}^{D_u \times D_y}, W \in \mathbb{R}^{D_y \times D_y}\}$, $\bar{u}_k = [u_{k-1} \dots u_{k-n_u}]^\top$, $\bar{y}_k = [y_{k-1} \dots y_{k-n_y}]^\top$.

For the prior distribution, we select a matrix Normal Wishart distribution:

Nisslbeck & Kouw (2025), Online Bayesian system identification in multivariate autoregressive models via message passing, ECC.

ARxl: Generative model specification

Let the likelihood of $y_k \in \mathbb{R}^{D_y}$, given $u_k \in \mathbb{R}^{D_u}$, be Gaussian distributed,

$$p(y_k \mid \Theta, u_k, \bar{u}_k, \bar{y}_k) = \mathcal{N}(y_k \mid A^\top \begin{bmatrix} \bar{y}_k \\ u_k \\ \bar{u}_k \end{bmatrix}, W^{-1}),$$

where $\Theta = \{A, W : A \in \mathbb{R}^{D_u \times D_y}, W \in \mathbb{R}^{D_y \times D_y}\}$, $\bar{u}_k = [u_{k-1} \dots u_{k-n_u}]^\top$, $\bar{y}_k = [y_{k-1} \dots y_{k-n_y}]^\top$.

For the prior distribution, we select a matrix Normal Wishart distribution:

$$p(\Theta) = \mathcal{MNW}(A, W \mid M_0, \Lambda_0^{-1}, \Omega_0^{-1}, \nu_0)$$

Nisslbeck & Kouw (2025), Online Bayesian system identification in multivariate autoregressive models via message passing, ECC.

ARxl: Generative model specification

Let the likelihood of $y_k \in \mathbb{R}^{D_y}$, given $u_k \in \mathbb{R}^{D_u}$, be Gaussian distributed,

$$p(y_k \mid \Theta, u_k, \bar{u}_k, \bar{y}_k) = \mathcal{N}(y_k \mid A^\top \begin{bmatrix} \bar{y}_k \\ u_k \\ \bar{u}_k \end{bmatrix}, W^{-1}),$$

where $\Theta = \{A, W : A \in \mathbb{R}^{D_u \times D_y}, W \in \mathbb{R}^{D_y \times D_y}\}$, $\bar{u}_k = [u_{k-1} \dots u_{k-n_u}]^\top$, $\bar{y}_k = [y_{k-1} \dots y_{k-n_y}]^\top$.

For the prior distribution, we select a matrix Normal Wishart distribution:

$$\begin{aligned} p(\Theta) &= \mathcal{MNW}(A, W \mid M_0, \Lambda_0^{-1}, \Omega_0^{-1}, \nu_0) \\ &= \mathcal{MN}(A \mid M_0, \Lambda_0^{-1}, \Omega_0^{-1}) \mathcal{W}(W \mid \Omega_0, \nu_0) \end{aligned}$$

Nisslbeck & Kouw (2025), Online Bayesian system identification in multivariate autoregressive models via message passing, ECC.

ARxl: Inferring parameters

The free energy functional for this model is:

$$\mathcal{F}[q] = \int q(\Theta) [-\ln p(y_k | \Theta, u_k, \bar{u}_k, \bar{y}_k)] d\Theta + \int q(\Theta) \ln p(\Theta) d\Theta - \int q(\Theta) \ln q(\Theta) d\Theta.$$

ARxl: Inferring parameters

The free energy functional for this model is:

$$\mathcal{F}[q] = \int q(\Theta) [-\ln p(y_k | \Theta, u_k, \bar{u}_k, \bar{y}_k)] d\Theta + \int q(\Theta) \ln p(\Theta) d\Theta - \int q(\Theta) \ln q(\Theta) d\Theta.$$

The optimal variational distribution is:

$$q(\Theta) = \frac{p(y_k | \Theta, u_k, \bar{u}_k, \bar{y}_k)}{p(y_k | u_k, \mathcal{D}_{k-1})} p(\Theta)$$

ARxl: Inferring parameters

The free energy functional for this model is:

$$\mathcal{F}[q] = \int q(\Theta) [-\ln p(y_k | \Theta, u_k, \bar{u}_k, \bar{y}_k)] d\Theta + \int q(\Theta) \ln p(\Theta) d\Theta - \int q(\Theta) \ln q(\Theta) d\Theta.$$

The optimal variational distribution is:

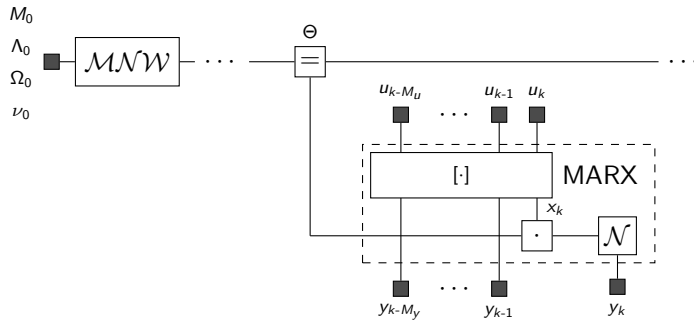
$$q(\Theta) = \frac{p(y_k | \Theta, u_k, \bar{u}_k, \bar{y}_k)}{p(y_k | u_k, \mathcal{D}_{k-1})} p(\Theta)$$

Here, $q(\Theta)$ recovers the true posterior distribution $p(\Theta | \mathcal{D}_k)$ exactly.

Särkkä, Svensson (2023). Bayesian filtering & Smoothing.

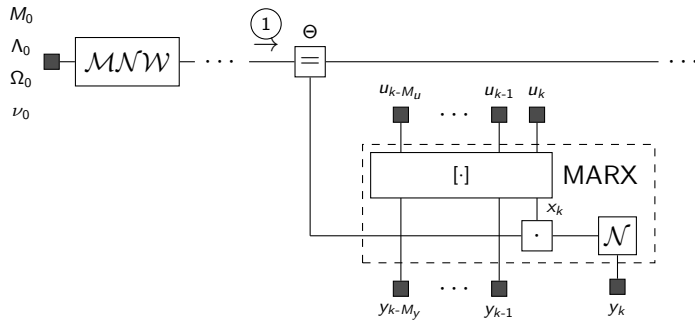
ARxl: Inferring parameters

In a factor graph, with messages, this is:



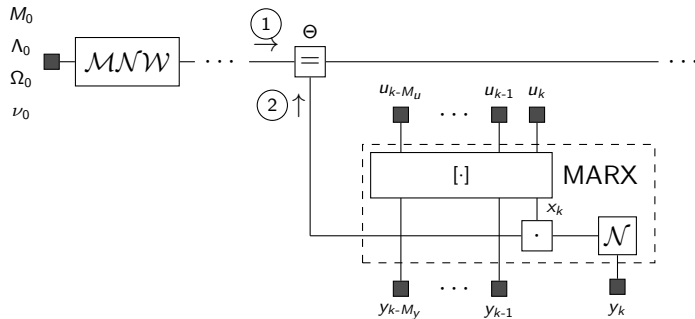
ARxl: Inferring parameters

In a factor graph, with messages, this is:



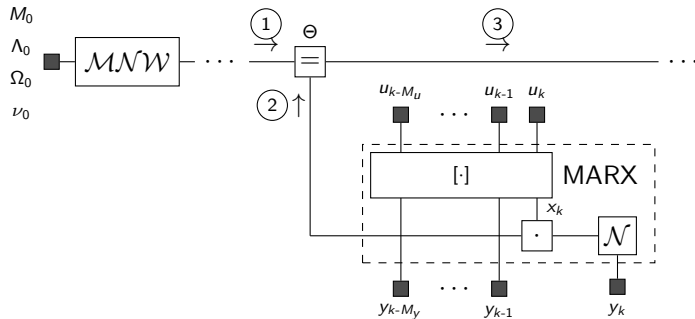
ARxl: Inferring parameters

In a factor graph, with messages, this is:



ARxl: Inferring parameters

In a factor graph, with messages, this is:



ARxl: Planning

For planning, we unroll the generative model into the future. For $t = k + 1$;

$$p(y_t, \Theta, u_t \mid \mathcal{D}_k) = p(y_t \mid \Theta, u_t, \bar{u}_t, \bar{y}_t) p(\Theta \mid \mathcal{D}_k) p(u_t)$$

ARxl: Planning

For planning, we unroll the generative model into the future. For $t = k + 1$;

$$p(y_t, \Theta, u_t \mid \mathcal{D}_k) = p(y_t \mid \Theta, u_t, \bar{u}_t, \bar{y}_t) p(\Theta \mid \mathcal{D}_k) p(u_t)$$

To incorporate a target, we first isolate the marginal distribution for the future output:

$$p(y_t \mid \Theta, u_t, \bar{u}_t, \bar{y}_t) p(\Theta \mid \mathcal{D}_k) = p(y_t, \Theta \mid u_t, \mathcal{D}_k) = p(\Theta \mid y_t, u_t, \mathcal{D}_k) p(y_t).$$

ARxl: Planning

For planning, we unroll the generative model into the future. For $t = k + 1$;

$$p(y_t, \Theta, u_t \mid \mathcal{D}_k) = p(y_t \mid \Theta, u_t, \bar{u}_t, \bar{y}_t) p(\Theta \mid \mathcal{D}_k) p(u_t)$$

To incorporate a target, we first isolate the marginal distribution for the future output:

$$p(y_t \mid \Theta, u_t, \bar{u}_t, \bar{y}_t) p(\Theta \mid \mathcal{D}_k) = p(y_t, \Theta \mid u_t, \mathcal{D}_k) = p(\Theta \mid y_t, u_t, \mathcal{D}_k) p(y_t).$$

We intervene on the marginal, $p(y_t) \rightarrow p(y_t \mid y_*) = \mathcal{N}(y_t \mid m_*, S_*)$, and rewrite

$$p(\Theta \mid y_t, u_t, \mathcal{D}_k) = \frac{p(y_t \mid \Theta, u_t, \bar{u}_t, \bar{y}_t) p(\Theta \mid \mathcal{D}_k)}{p(y_t \mid u_t, \mathcal{D}_k)}.$$

ARxl: Planning

For planning, we unroll the generative model into the future. For $t = k + 1$;

$$p(y_t, \Theta, u_t \mid \mathcal{D}_k) = p(y_t \mid \Theta, u_t, \bar{u}_t, \bar{y}_t) p(\Theta \mid \mathcal{D}_k) p(u_t)$$

To incorporate a target, we first isolate the marginal distribution for the future output:

$$p(y_t \mid \Theta, u_t, \bar{u}_t, \bar{y}_t) p(\Theta \mid \mathcal{D}_k) = p(y_t, \Theta \mid u_t, \mathcal{D}_k) = p(\Theta \mid y_t, u_t, \mathcal{D}_k) p(y_t).$$

We intervene on the marginal, $p(y_t) \rightarrow p(y_t \mid y_*) = \mathcal{N}(y_t \mid m_*, S_*)$, and rewrite

$$p(\Theta \mid y_t, u_t, \mathcal{D}_k) = \frac{p(y_t \mid \Theta, u_t, \bar{u}_t, \bar{y}_t) p(\Theta \mid \mathcal{D}_k)}{p(y_t \mid u_t, \mathcal{D}_k)}.$$

Overall, the generative model for the future becomes:

$$p(y_t, \Theta, u_t \mid y_*, \mathcal{D}_k) = \frac{p(y_t \mid \Theta, u_t, \bar{u}_t, \bar{y}_t) p(\Theta \mid \mathcal{D}_k)}{p(y_t \mid u_t, \mathcal{D}_k)} p(y_t \mid y_*) p(u_t).$$

ARxl: Predictions

The planning model requires $p(y_t \mid u_t, \mathcal{D}_k)$, i.e., the posterior predictive distribution.

Murphy (2022). Probabilistic machine learning: an introduction.

ARxl: Predictions

The planning model requires $p(y_t \mid u_t, \mathcal{D}_k)$, i.e., the posterior predictive distribution.

Under the autoregressive model specification, we can solve the integral analytically:

Murphy (2022). Probabilistic machine learning: an introduction.

ARxl: Predictions

The planning model requires $p(y_t | u_t, \mathcal{D}_k)$, i.e., the posterior predictive distribution.

Under the autoregressive model specification, we can solve the integral analytically:

$$p(y_t | u_t, \mathcal{D}_k) = \int p(y_t | \Theta, u_t, \bar{u}_t, \bar{y}_t) p(\Theta | \mathcal{D}_k) d\Theta$$

Murphy (2022). Probabilistic machine learning: an introduction.

ARxl: Predictions

The planning model requires $p(y_t | u_t, \mathcal{D}_k)$, i.e., the posterior predictive distribution.

Under the autoregressive model specification, we can solve the integral analytically:

$$\begin{aligned} p(y_t | u_t, \mathcal{D}_k) &= \int p(y_t | \Theta, u_t, \bar{u}_t, \bar{y}_t) p(\Theta | \mathcal{D}_k) d\Theta \\ &= \mathcal{T}_{\eta_k}(y_t | \mu_k(u_t), \Sigma_k(u_t)) , \end{aligned}$$

Murphy (2022). Probabilistic machine learning: an introduction.

ARxl: Predictions

The planning model requires $p(y_t | u_t, \mathcal{D}_k)$, i.e., the posterior predictive distribution.

Under the autoregressive model specification, we can solve the integral analytically:

$$\begin{aligned} p(y_t | u_t, \mathcal{D}_k) &= \int p(y_t | \Theta, u_t, \bar{u}_t, \bar{y}_t) p(\Theta | \mathcal{D}_k) d\Theta \\ &= \mathcal{T}_{\eta_k}(y_t | \mu_k(u_t), \Sigma_k(u_t)), \end{aligned}$$

$$\text{where } \eta_k = \nu_k - D_y + 1, \mu_k(u_t) = M_k^\top \begin{bmatrix} u_k \\ \bar{u}_t \\ \bar{y}_t \end{bmatrix} \text{ and } \Sigma_k(u_t) = \eta_k \Omega_k \left(1 + \begin{bmatrix} u_t \\ \bar{u}_t \\ \bar{y}_t \end{bmatrix}^\top \Lambda_k^{-1} \begin{bmatrix} u_t \\ \bar{u}_t \\ \bar{y}_t \end{bmatrix} \right).$$

Murphy (2022). Probabilistic machine learning: an introduction.

ARxl: Inferring controls

To infer a variational posterior over controls, we form a free energy functional;

Kouw et al. (2025), Message passing-based inference in an autoregressive active inference agent, IWAI.

ARxl: Inferring controls

To infer a variational posterior over controls, we form a free energy functional;

$$\mathcal{F}_k[q] = \mathbb{E}_{q(y_t, \Theta, u_t)} \left[\ln \frac{q(y_t, \Theta, u_t)}{p(y_t, \Theta, u_t \mid y_*, \mathcal{D}_k)} \right],$$

Kouw et al. (2025), Message passing-based inference in an autoregressive active inference agent, IWA1.

ARxl: Inferring controls

To infer a variational posterior over controls, we form a free energy functional;

$$\mathcal{F}_k[q] = \mathbb{E}_{q(y_t, \Theta, u_t)} \left[\ln \frac{q(y_t, \Theta, u_t)}{p(y_t, \Theta, u_t \mid y_*, \mathcal{D}_k)} \right],$$

where $q(y_t, \Theta, u_t) = p(y_t \mid \Theta, u_t, \bar{u}_t, \bar{y}_t) p(\Theta \mid \mathcal{D}_k) q(u_t)$.

Kouw et al. (2025), Message passing-based inference in an autoregressive active inference agent, IWA1.

ARxl: Inferring controls

To infer a variational posterior over controls, we form a free energy functional;

$$\mathcal{F}_k[q] = \mathbb{E}_{q(y_t, \Theta, u_t)} \left[\ln \frac{q(y_t, \Theta, u_t)}{p(y_t, \Theta, u_t \mid y_*, \mathcal{D}_k)} \right],$$

where $q(y_t, \Theta, u_t) = p(y_t \mid \Theta, u_t, \bar{u}_t, \bar{y}_t) p(\Theta \mid \mathcal{D}_k) q(u_t)$.

The solution is proportional to the prior times an energy-based pseudo-likelihood:

Kouw et al. (2025), Message passing-based inference in an autoregressive active inference agent, IWA1.

ARxl: Inferring controls

To infer a variational posterior over controls, we form a free energy functional;

$$\mathcal{F}_k[q] = \mathbb{E}_{q(y_t, \Theta, u_t)} \left[\ln \frac{q(y_t, \Theta, u_t)}{p(y_t, \Theta, u_t \mid y_*, \mathcal{D}_k)} \right],$$

where $q(y_t, \Theta, u_t) = p(y_t \mid \Theta, u_t, \bar{u}_t, \bar{y}_t) p(\Theta \mid \mathcal{D}_k) q(u_t)$.

The solution is proportional to the prior times an energy-based pseudo-likelihood:

$$q^*(u_t) = \arg \min_{q \in \mathcal{Q}} \mathcal{F}_k[q]$$

ARxl: Inferring controls

To infer a variational posterior over controls, we form a free energy functional;

$$\mathcal{F}_k[q] = \mathbb{E}_{q(y_t, \Theta, u_t)} \left[\ln \frac{q(y_t, \Theta, u_t)}{p(y_t, \Theta, u_t \mid y_*, \mathcal{D}_k)} \right],$$

where $q(y_t, \Theta, u_t) = p(y_t \mid \Theta, u_t, \bar{u}_t, \bar{y}_t) p(\Theta \mid \mathcal{D}_k) q(u_t)$.

The solution is proportional to the prior times an energy-based pseudo-likelihood:

$$\begin{aligned} q^*(u_t) &= \arg \min_{q \in \mathcal{Q}} \mathcal{F}_k[q] \\ &\propto p(u_t) \exp \left(\underbrace{\mathbb{E}_{p(y_t \mid u_t, \mathcal{D}_k)} \left[-\ln p(y_t \mid u_t, \mathcal{D}_k) \right]}_{\text{predictive entropy}} - \underbrace{\mathbb{E}_{p(y_t \mid u_t, \mathcal{D}_k)} \left[-\ln p(y_t \mid y_*) \right]}_{\text{cross-entropy to goal}} \right). \end{aligned}$$

Kouw et al. (2025), Message passing-based inference in an autoregressive active inference agent, IWA1.

ARxl: information gain

Under the given model specification, we can derive an insightful interpretation:

ARxl: information gain

Under the given model specification, we can derive an insightful interpretation:

Theorem

For particular p, q , the entropy of the posterior predictive is equivalent, up to additive constants, to the information gain from future observations to parameters given actions:

$$\mathbb{E}_{p(y_t | u_t, \mathcal{D}_k)} [-\ln p(y_t | u_t, \mathcal{D}_k)] + C = \mathbb{E}_{p(y_t, \Theta | u_t, \mathcal{D}_k)} \left[\ln \frac{p(y_t, \Theta | u_t, \mathcal{D}_k)}{p(y_t | u_t, \mathcal{D}_k) p(\Theta | \mathcal{D}_k)} \right]$$

where the constants C are terms not involving u_t .

ARxl: information gain

Under the given model specification, we can derive an insightful interpretation:

Theorem

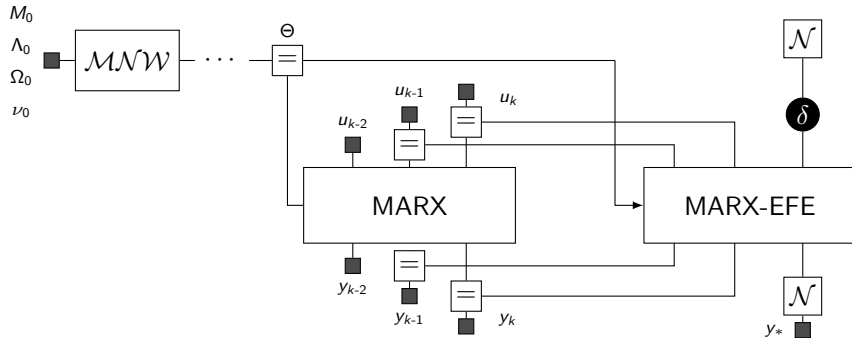
For particular p, q , the entropy of the posterior predictive is equivalent, up to additive constants, to the information gain from future observations to parameters given actions:

$$\mathbb{E}_{p(y_t | u_t, \mathcal{D}_k)} [-\ln p(y_t | u_t, \mathcal{D}_k)] + C = \mathbb{E}_{p(y_t, \Theta | u_t, \mathcal{D}_k)} \left[\ln \frac{p(y_t, \Theta | u_t, \mathcal{D}_k)}{p(y_t | u_t, \mathcal{D}_k)p(\Theta | \mathcal{D}_k)} \right]$$

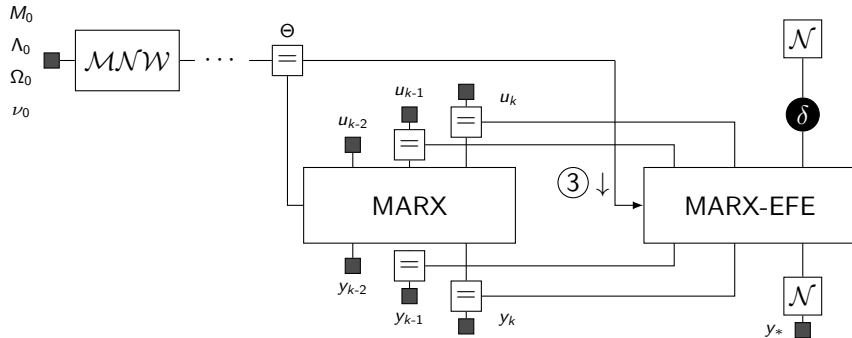
where the constants C are terms not involving u_t .

Here, maximizing predictive entropy, with respect to u_t , is maximizing information gain.

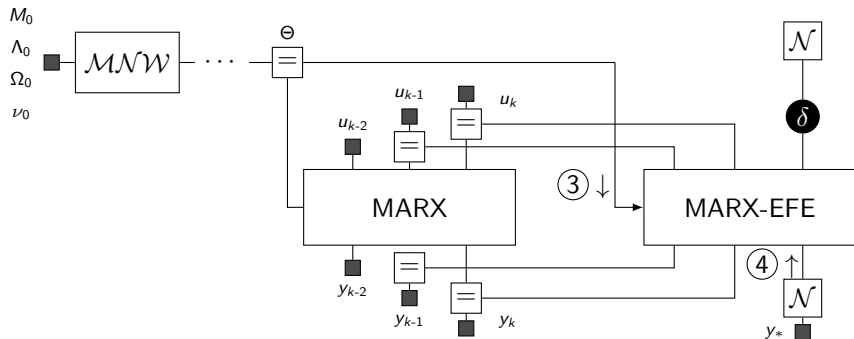
ARxl: 1-step ahead planning graph



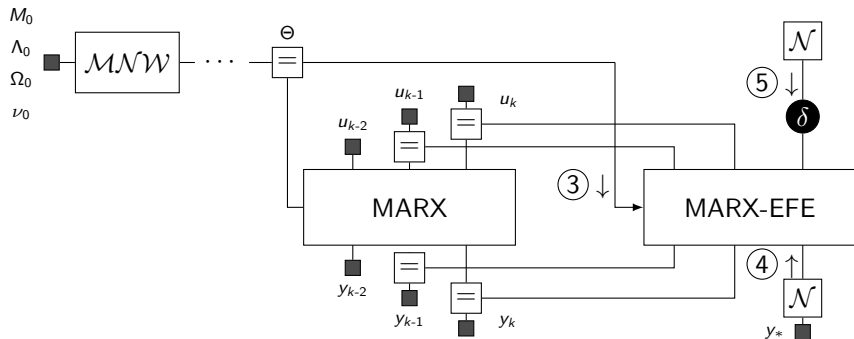
ARxl: 1-step ahead planning graph



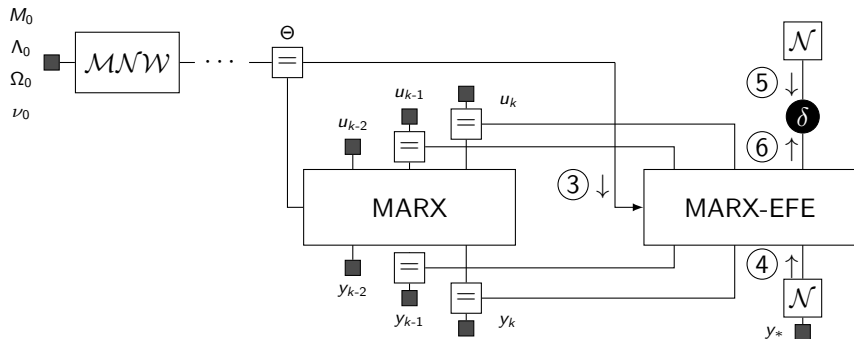
ARxl: 1-step ahead planning graph



ARxl: 1-step ahead planning graph



ARxl: 1-step ahead planning graph



ARxl: Extending the horizon

At $t \geq k + 2$, the buffers will contain latent instead of observed variables.

ARxl: Extending the horizon

At $t \geq k + 2$, the buffers will contain latent instead of observed variables.

Now, each node in the planning graph for $t \geq 2$ has to send backwards messages.

ARxl: Extending the horizon

At $t \geq k + 2$, the buffers will contain latent instead of observed variables.

Now, each node in the planning graph for $t \geq 2$ has to send backwards messages.

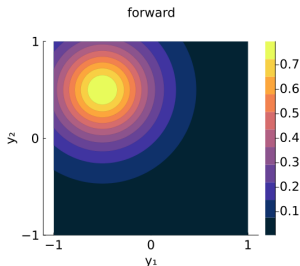
Forwards and backwards messages over y_t are combined to form marginals;

ARxl: Extending the horizon

At $t \geq k + 2$, the buffers will contain latent instead of observed variables.

Now, each node in the planning graph for $t \geq 2$ has to send backwards messages.

Forwards and backwards messages over y_t are combined to form marginals;

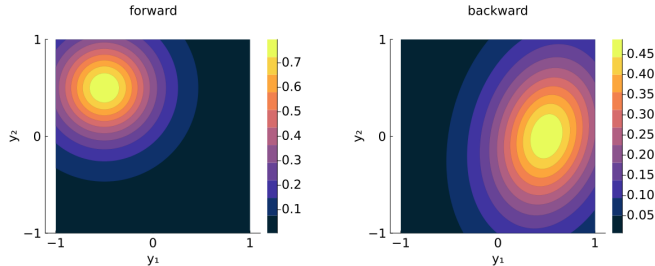


ARxl: Extending the horizon

At $t \geq k + 2$, the buffers will contain latent instead of observed variables.

Now, each node in the planning graph for $t \geq 2$ has to send backwards messages.

Forwards and backwards messages over y_t are combined to form marginals;

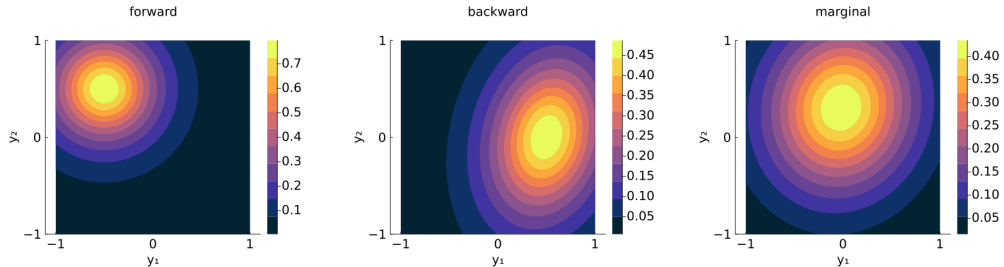


ARxl: Extending the horizon

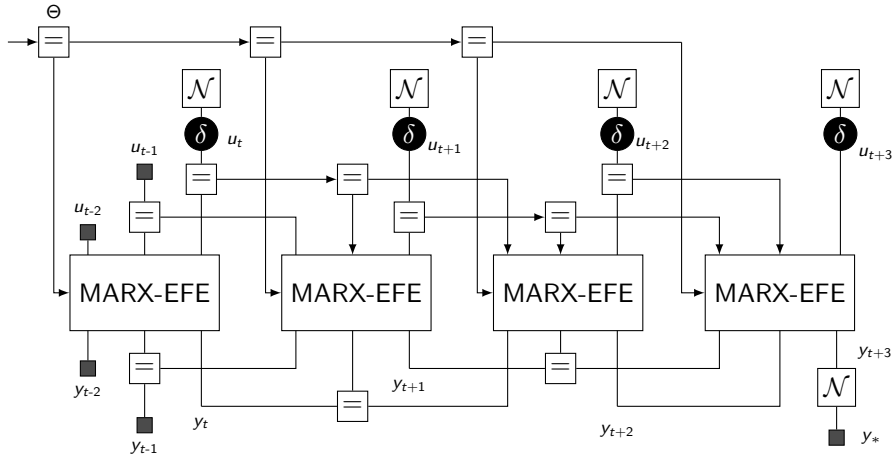
At $t \geq k + 2$, the buffers will contain latent instead of observed variables.

Now, each node in the planning graph for $t \geq 2$ has to send backwards messages.

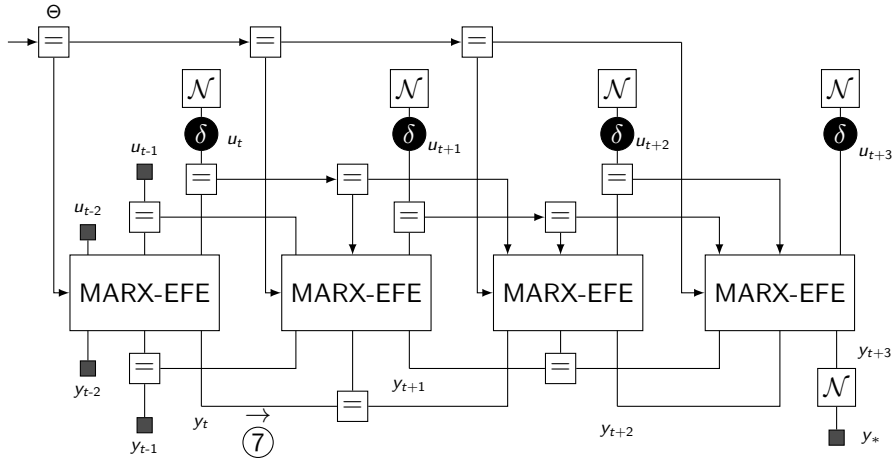
Forwards and backwards messages over y_t are combined to form marginals;



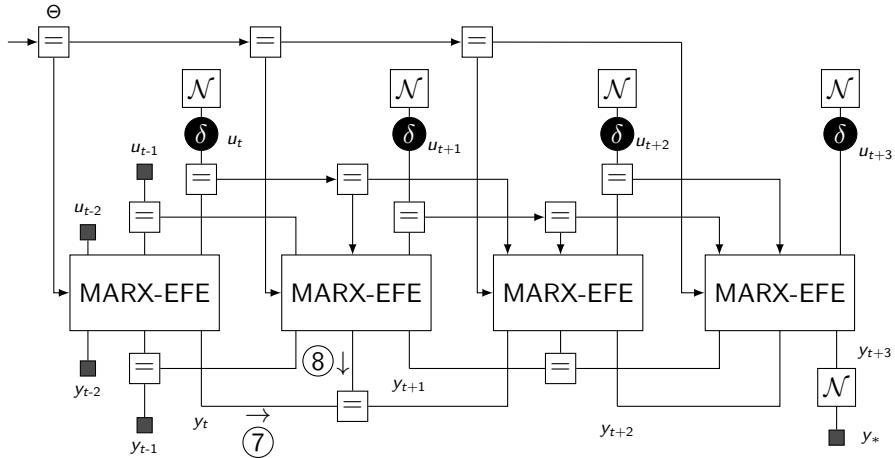
ARxl: T -step ahead planning



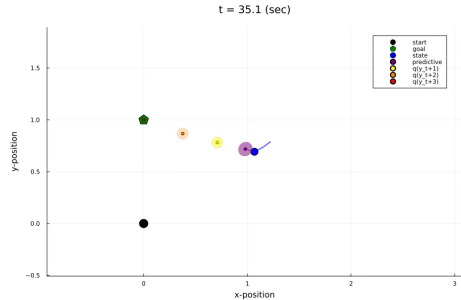
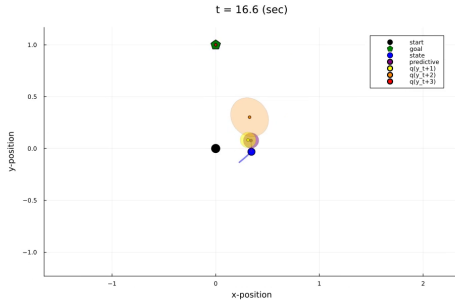
ARxl: T -step ahead planning



ARxl: T -step ahead planning



ARxl: Navigation trial



ARxl: T -step ahead planning

On the planning trajectory:

- The forward pass of predictions generates a trajectory, given prior actions.

ARxl: T -step ahead planning

On the planning trajectory:

- The forward pass of predictions generates a trajectory, given prior actions.
- The backwards pass from goal priors to previous time variables, gives direction.

ARxl: T -step ahead planning

On the planning trajectory:

- The forward pass of predictions generates a trajectory, given prior actions.
- The backwards pass from goal priors to previous time variables, gives direction.
- Together, the two passes generate a sequence of intermediate goal priors.

ARxl: T -step ahead planning

On the planning trajectory:

- The forward pass of predictions generates a trajectory, given prior actions.
- The backwards pass from goal priors to previous time variables, gives direction.
- Together, the two passes generate a sequence of intermediate goal priors.

ARxl: T -step ahead planning

On the planning trajectory:

- The forward pass of predictions generates a trajectory, given prior actions.
- The backwards pass from goal priors to previous time variables, gives direction.
- Together, the two passes generate a sequence of intermediate goal priors.

On the sequence of actions:

- Each planning node selects its own action based on balancing local info gain versus cross-entropy to target.

ARxl: T -step ahead planning

On the planning trajectory:

- The forward pass of predictions generates a trajectory, given prior actions.
- The backwards pass from goal priors to previous time variables, gives direction.
- Together, the two passes generate a sequence of intermediate goal priors.

On the sequence of actions:

- Each planning node selects its own action based on balancing local info gain versus cross-entropy to target.
- The ability to block backwards messages towards previous actions reduces computational cost, and improves scaling.

ARxl: T -step ahead planning

On the planning trajectory:

- The forward pass of predictions generates a trajectory, given prior actions.
- The backwards pass from goal priors to previous time variables, gives direction.
- Together, the two passes generate a sequence of intermediate goal priors.

On the sequence of actions:

- Each planning node selects its own action based on balancing local info gain versus cross-entropy to target.
- The ability to block backwards messages towards previous actions reduces computational cost, and improves scaling.
- Limitation: controls are not enforced to be smooth over time.

Comparing ARxI to MPC

Let's compare ARxI's controller with a standard model-predictive controller:

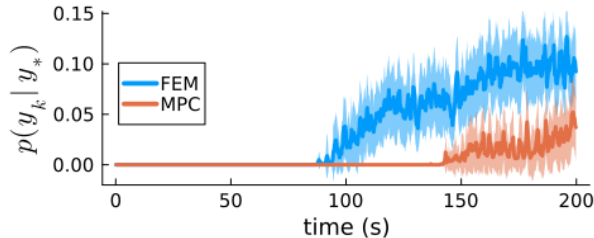
$$u_t^{\text{MPC}} = \arg \min u_t^\top \Upsilon u_t + (\mu_t(u_t) - m_*)^\top S_*^{-1} (\mu_t(u_t) - m_*).$$

Comparing ARxI to MPC

Let's compare ARxI's controller with a standard model-predictive controller:

$$u_t^{\text{MPC}} = \arg \min u_t^\top \Upsilon u_t + (\mu_t(u_t) - m_*)^\top S_*^{-1} (\mu_t(u_t) - m_*).$$

ARxI learns its dynamics more quickly and is able to reach the goal sooner:



Mutual information

Working out the two entropic terms in the variational posterior gives:

Mutual information

Working out the two entropic terms in the variational posterior gives:

$$E(u_t) = \mathbb{E}_{p(y_t | u_t, \mathcal{D}_k)} \left[-\ln p(y_t | u_t, \mathcal{D}_k) \right] - \mathbb{E}_{p(y_t | u_t, \mathcal{D}_k)} \left[-\ln p(y_t | y_*) \right]$$

Mutual information

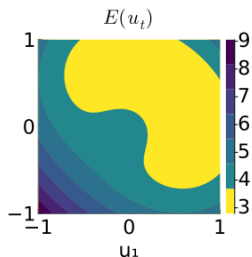
Working out the two entropic terms in the variational posterior gives:

$$\begin{aligned} E(u_t) &= \mathbb{E}_{p(y_t | u_t, \mathcal{D}_k)} \left[-\ln p(y_t | u_t, \mathcal{D}_k) \right] - \mathbb{E}_{p(y_t | u_t, \mathcal{D}_k)} \left[-\ln p(y_t | y_*) \right] \\ &= C + \underbrace{-\frac{1}{2} \ln |\Sigma_t(u_t)| + \frac{1}{2} \text{Tr} \left[S_*^{-1} \frac{\eta_t}{\eta_t - 2} \Sigma_t(u_t) \right]}_{I(u_t)} + \underbrace{\frac{1}{2} (\mu_t(u_t) - m_*)^\top S_*^{-1} (\mu_t(u_t) - m_*)}_{G(u_t)} . \end{aligned}$$

Mutual information

Working out the two entropic terms in the variational posterior gives:

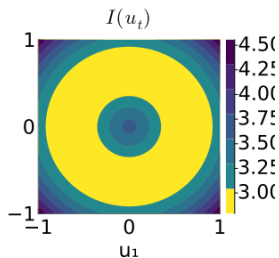
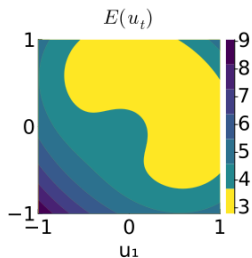
$$\begin{aligned} E(u_t) &= \mathbb{E}_{p(y_t | u_t, \mathcal{D}_k)} \left[-\ln p(y_t | u_t, \mathcal{D}_k) \right] - \mathbb{E}_{p(y_t | u_t, \mathcal{D}_k)} \left[-\ln p(y_t | y_*) \right] \\ &= C + \underbrace{-\frac{1}{2} \ln |\Sigma_t(u_t)| + \frac{1}{2} \text{Tr} \left[S_*^{-1} \frac{\eta_t}{\eta_t - 2} \Sigma_t(u_t) \right]}_{I(u_t)} + \underbrace{\frac{1}{2} (\mu_t(u_t) - m_*)^\top S_*^{-1} (\mu_t(u_t) - m_*)}_{G(u_t)} . \end{aligned}$$



Mutual information

Working out the two entropic terms in the variational posterior gives:

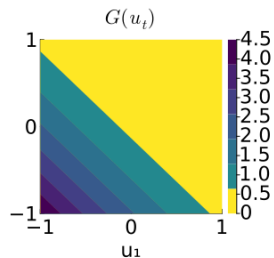
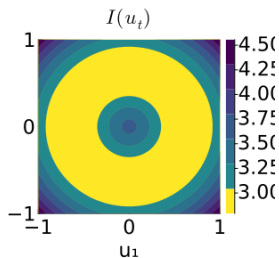
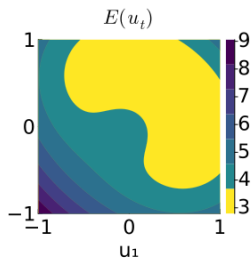
$$\begin{aligned}
 E(u_t) &= \mathbb{E}_{p(y_t | u_t, \mathcal{D}_k)} \left[-\ln p(y_t | u_t, \mathcal{D}_k) \right] - \mathbb{E}_{p(y_t | u_t, \mathcal{D}_k)} \left[-\ln p(y_t | y_*) \right] \\
 &= C + \underbrace{-\frac{1}{2} \ln |\Sigma_t(u_t)| + \frac{1}{2} \text{Tr} \left[S_*^{-1} \frac{\eta_t}{\eta_t - 2} \Sigma_t(u_t) \right]}_{I(u_t)} + \underbrace{\frac{1}{2} (\mu_t(u_t) - m_*)^\top S_*^{-1} (\mu_t(u_t) - m_*)}_{G(u_t)} .
 \end{aligned}$$



Mutual information

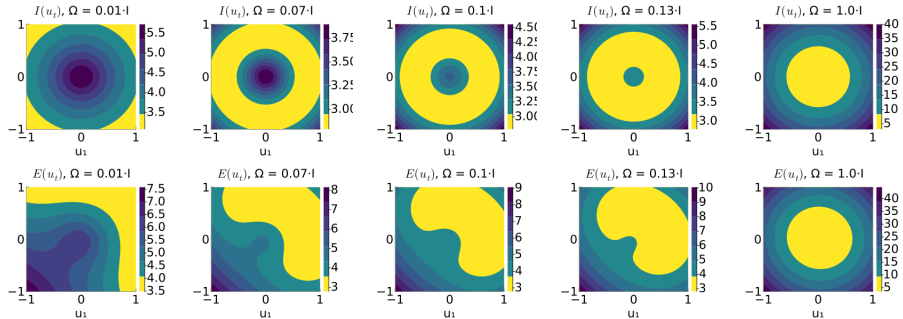
Working out the two entropic terms in the variational posterior gives:

$$\begin{aligned}
 E(u_t) &= \mathbb{E}_{p(y_t | u_t, \mathcal{D}_k)} \left[-\ln p(y_t | u_t, \mathcal{D}_k) \right] - \mathbb{E}_{p(y_t | u_t, \mathcal{D}_k)} \left[-\ln p(y_t | y_*) \right] \\
 &= C + \underbrace{-\frac{1}{2} \ln |\Sigma_t(u_t)| + \frac{1}{2} \text{Tr} \left[S_*^{-1} \frac{\eta_t}{\eta_t - 2} \Sigma_t(u_t) \right]}_{I(u_t)} + \underbrace{\frac{1}{2} (\mu_t(u_t) - m_*)^\top S_*^{-1} (\mu_t(u_t) - m_*)}_{G(u_t)} .
 \end{aligned}$$



Shape of information gain

Uncertainty affects the shape of the information gain, and the overall control posterior.



Controlling information gain

To be able to guide information gain, an agent must be able to *control* their predictive uncertainty, here $\Sigma_k(u_t)$.

Controlling information gain

To be able to guide information gain, an agent must be able to *control* their predictive uncertainty, here $\Sigma_k(u_t)$.

Perhaps obvious but note that some approximation techniques will drop controllable information gain.

Controlling information gain

To be able to guide information gain, an agent must be able to *control* their predictive uncertainty, here $\Sigma_k(u_t)$.

Perhaps obvious but note that some approximation techniques will drop controllable information gain.

For example, in deep reinforcement learning, one typically approximates expectations with sample averages, which leads to uncontrollable predictive entropy;

$$\int p(y_t \mid \Theta, u_t, \bar{u}_t, \bar{y}_t) p(\Theta \mid \mathcal{D}_k) d\Theta \approx \frac{1}{n} \sum_{i=1}^n \mathcal{N}\left(y_t \mid A_{(i)}^\top \begin{bmatrix} \bar{y}_t \\ u_t \\ \bar{u}_t \end{bmatrix}, W_{(i)}^{-1}\right) \text{ for } A_{(i)}, W_{(i)} \sim p(\Theta \mid \mathcal{D}_k).$$

Thanks for your attention



References

- Parr, Pezullo, Friston (2022), Active inference: the free energy principle in mind, brain, and behavior, MIT Press.
- Senoz et al. (2021), Variational message passing and local constraint manipulation in factor graphs, MDPI Entropy.
- Loeliger (2004), An introduction to factor graphs, IEEE Signal Processing Magazine.
- Nisslbeck & Kouw (2025), Online Bayesian system identification in multivariate autoregressive models via message passing, IEEE European Control Conference.
- Särkkä, Svensson (2023). Bayesian filtering & Smoothing. Cambridge University Press.
- Kouw et al. (2025), Message passing-based inference in an autoregressive active inference agent, International Workshop on Active Inference.
- Murphy (2022). Probabilistic machine learning: an introduction. MIT Press.

Tiny AI on robots

