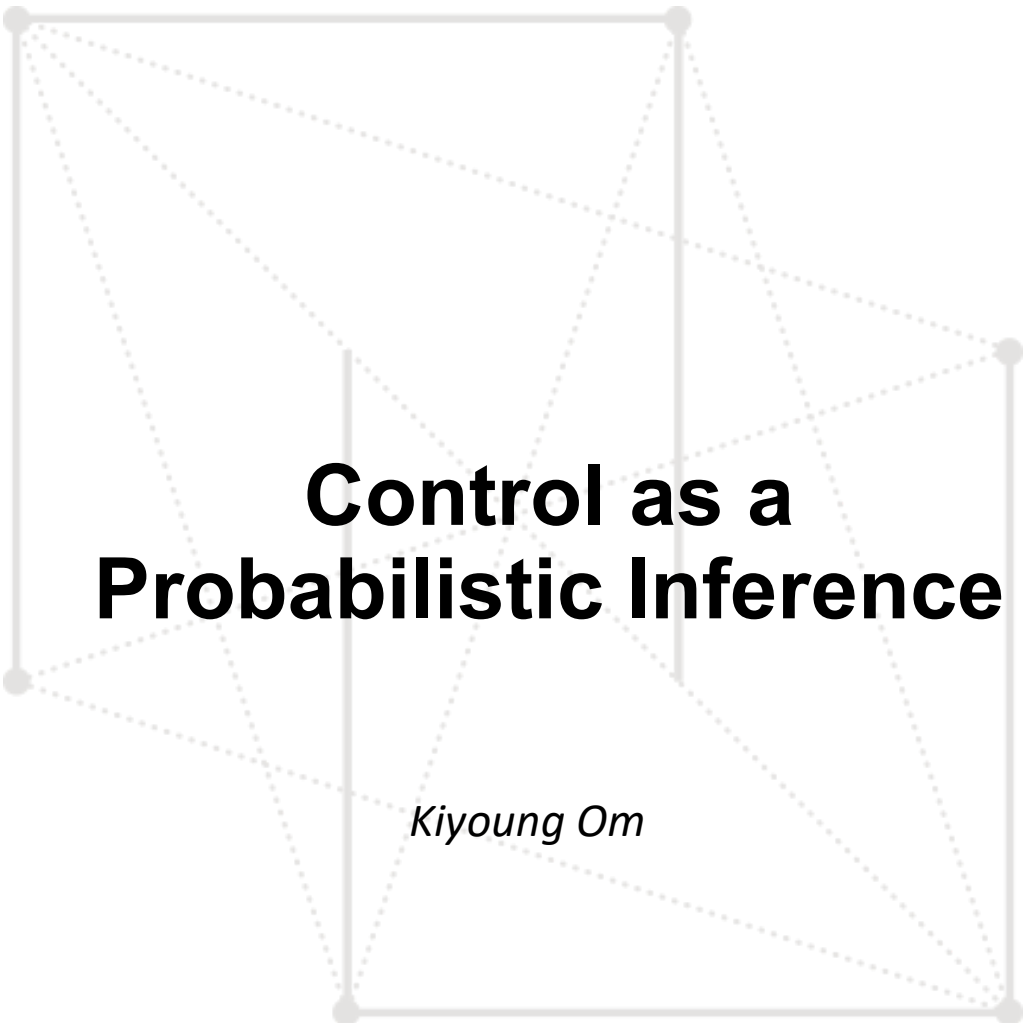




Control as a Probabilistic Inference

Kiyoung Om

Table of Contents

- 1. Motivation of understanding control as probabilistic inference framework**
- 2. Theory of control as probabilistic inference**
- 3. Applications, connection to various algorithm**
- 4. Future plan for research**

Motivation of understanding control as probabilistic inference framework

- **Maximum Entropy RL**

$$J_{\text{MaxEnt}}(\pi; p, r) \triangleq \mathbb{E}_{\mathbf{a}_t \sim \pi(\mathbf{a}_t | \mathbf{s}_t), \mathbf{s}_{t+1} \sim p(\mathbf{s}_{t+1} | \mathbf{s}_t, \mathbf{a}_t)} \left[\sum_{t=1}^T r(\mathbf{s}_t, \mathbf{a}_t) + \alpha \mathcal{H}_{\pi} [\mathbf{a}_t | \mathbf{s}_t] \right],$$

- Representative algorithms: Soft-Q Learning(**SQL**), Soft Actor Critic(**SAC**)

- **KL-Divergence Constraints for Policy Search**

$$\max_{\theta} \mathbb{E}_{s \sim \rho_{\theta_{\text{old}}}, a \sim \pi_{\theta_{\text{old}}}} \left[\frac{\pi_{\theta}(a|s)}{\pi_{\theta_{\text{old}}}(a|s)} A^{\pi_{\theta_{\text{old}}}}(s, a) \right] \quad \mathbb{E}_{s \sim \rho_{\theta_{\text{old}}}} [D_{\text{KL}}(\pi_{\theta_{\text{old}}}(\cdot|s) \parallel \pi_{\theta}(\cdot|s))] \leq \delta$$

- Representative algorithms: Trust Region Policy Optimization(**TRPO**), Maximum a posteriori policy optimization (**MPO**), Advantage weighted regression (**AWR**),

- **Return conditioned trajectory planning**

$$p(\tau | \mathcal{O}_{1:T}) \propto p(\tau) \exp \left(\sum_{t=1}^T r(s_t, a_t) \right)$$

- Representative algorithms:
Planning with diffusion for flexible behavior synthesis (**Diffuser**)

Theory of control as probabilistic inference

Standard RL Policy Search Problem

- Find θ that maximizes cumulative rewards, with parameterization: $p(a_t|s_t, \theta) = \pi_\theta(a_t|s_t)$

$$\theta^* = \arg \max_{\theta} \sum_{t=1}^T \mathbb{E}_{(s_t, a_t) \sim p(s_t, a_t | \theta)} [r(s_t, a_t)]$$

- In other words,

$$p(\tau | \theta) = p(s_1, a_1, \dots, s_T, a_T | \theta) = p(s_1) \prod_{t=1}^T p(a_t | s_t, \theta) p(s_{t+1} | s_t, a_t).$$

$$\theta^* = \arg \max_{\theta} \mathbb{E}_{\tau \sim p(\tau | \theta)} \left[\sum_{t=1}^T r(s_t, a_t) \right]$$

To make RL policy search problem into probabilistic inference framework:

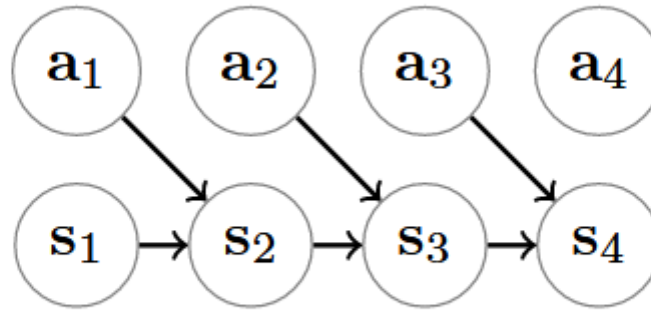
- Draw a **Probabilistic Graphical Model (PGM)**, where:

Most probable trajectory: Trajectory from the **optimal policy**.

Inference of posterior action conditional $p(a_t | s_t, \theta)$: Gives the **optimal policy**.

Theory of control as probabilistic inference

How to draw the PGM?



(a) graphical model with states and actions

- Simplest graphical model to express control problem.
- When we see the graphical model: **Decompose the Joint Distribution!**

$$p(s_1, a_1, s_2, a_2, \dots, a_T, s_{T+1}) = p(s_1) \prod_{t=1}^T p(a_t) \cdot p(s_{t+1} \mid s_t, a_t)$$

- Is it enough to satisfy below questions?

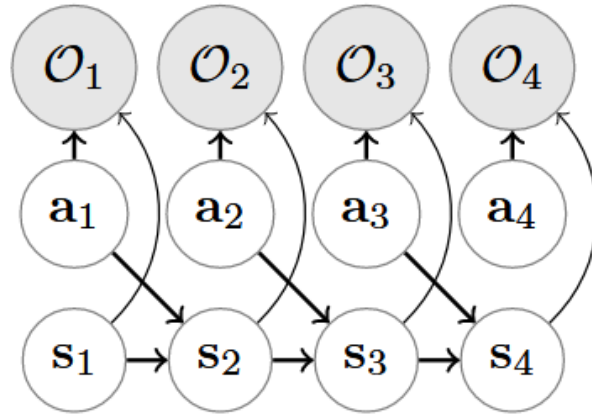
Most probable trajectory: Trajectory from the **optimal policy**.

Cannot induce reward/cost information here.

Inference of posterior action conditional $p(a_t \mid s_t, \theta)$: Gives the **optimal policy**.

Theory of control as probabilistic inference

How to draw the PGM?



(b) graphical model with optimality variables

- Introduced additional random variable $\mathcal{O}_t = \begin{cases} 1 & \text{if timestep } t \text{ is optimal} \\ 0 & \text{otherwise} \end{cases}$
- When we see the graphical model: **Decompose the Joint Distribution!**

$$p(\tau, \mathcal{O}_{1:T}) = p(s_1) \prod_{t=1}^T p(a_t) p(\mathcal{O}_t = 1 \mid s_t, a_t) p(s_{t+1} \mid s_t, a_t)$$

Decision Choice 1: As we were not interested in action prior (online), we can take it as simplest:

$$p(a_t) = p(a_t \mid s_t) = \frac{1}{|\mathcal{A}_t|}$$

Decision Choice 2: We define conditional distribution: $p(\mathcal{O}_t = 1 \mid s_t, a_t) = \exp(r(s_t, a_t))$

Theory of control as probabilistic inference

Most probable trajectory as trajectory from the optimal policy

- From the two Decision Choices, we can factorize posterior:

Decision 1 $\Rightarrow p(\tau \mid \mathcal{O}_{1:T}) \propto p(\tau, \mathcal{O}_{1:T}) = p(s_1) \prod_{t=1}^T p(\mathcal{O}_t = 1 \mid s_t, a_t) p(s_{t+1} \mid s_t, a_t)$

Decision 2 $\Rightarrow$

$$= p(s_1) \prod_{t=1}^T \exp(r(s_t, a_t)) p(s_{t+1} \mid s_t, a_t)$$
$$= \left[p(s_1) \prod_{t=1}^T p(s_{t+1} \mid s_t, a_t) \right] \exp \left(\sum_{t=1}^T r(s_t, a_t) \right).$$

RHS term should be $\propto$;
We can incorporate this
defining reward $r(s_t, a_t)$

So we can set the target distribution as:

- Product distribution of **Dynamics** and **exponential of Reward Sum (stochastic dynamics)**

$$p(\tau \mid \mathcal{O}_{1:T}) \propto \left[p(s_1) \prod_{t=1}^T p(s_{t+1} \mid s_t, a_t) \right] \exp \left(\sum_{t=1}^T r(s_t, a_t) \right).$$

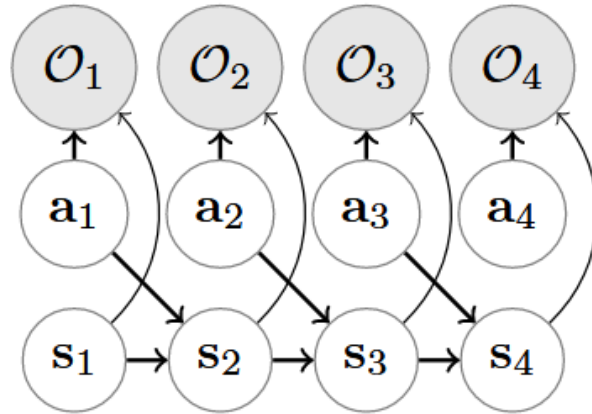
- Product distribution of **Feasible dynamics** and **exponential of Reward Sum (deterministic dynamics)**

$$p(\tau \mid \mathcal{O}_{1:T}) \propto \mathbb{1}[p(\tau) \neq 0] \exp \left(\sum_{t=1}^T r(s_t, a_t) \right)$$

Theory of control as probabilistic inference

Inference of posterior action conditional $p(a_t|s_t, \theta)$: Gives the optimal policy.

- Now let's find out how to recover **optimal policy** $p(a_t|s_t, \mathcal{O}_{t:T} = 1)$ from designed PGM.



(b) graphical model with optimality variables

- Definition 1. $\beta_t(s_t, a_t) = p(\mathcal{O}_{t:T}|s_t, a_t)$
- Definition 2. $\beta_t(s_t) = p(\mathcal{O}_{t:T}|s_t)$

Theory of control as probabilistic inference

Inference of posterior action conditional $p(a_t|s_t, \theta)$: Gives the optimal policy.

Definition 1. $\beta_t(s_t, a_t) = p(\mathcal{O}_{t:T}|s_t, a_t)$

Definition 2. $\beta_t(s_t) = p(\mathcal{O}_{t:T}|s_t)$

- **Formulation 1**; Marginalizing $\beta_t(s_t)$ with a_t

$$\beta_t(s_t) = p(\mathcal{O}_{t:T} | s_t) = \int_{\mathcal{A}} p(\mathcal{O}_{t:T}, a_t | s_t) da_t \quad (1)$$

$$= \int_{\mathcal{A}} \frac{p(\mathcal{O}_{t:T}, a_t | s_t)}{p(a_t | s_t)} p(a_t | s_t) da_t \quad (2)$$

$$= \int_{\mathcal{A}} p(\mathcal{O}_{t:T} | s_t, a_t) p(a_t | s_t) da_t \quad (3)$$

$$= \int_{\mathcal{A}} \beta_t(s_t, a_t) p(a_t | s_t) da_t. \quad (4)$$

Decision 1  $\beta_t(s_t) = \int_{\mathcal{A}} \beta_t(s_t, a_t) da_t$

$$p(a_t) = p(a_t | s_t) = \frac{1}{|\mathcal{A}_t|}$$

Theory of control as probabilistic inference

Inference of posterior action conditional $p(a_t|s_t, \theta)$: Gives the optimal policy.

Definition 1. $\beta_t(s_t, a_t) = p(\mathcal{O}_{t:T}|s_t, a_t)$

Definition 2. $\beta_t(s_t) = p(\mathcal{O}_{t:T}|s_t)$

- **Formulation 2;** Marginalizing $\beta_t(s_t, a_t)$ with s_{t+1}

$$\beta_t(s_t, a_t) = p(\mathcal{O}_{t:T} \mid s_t, a_t) \quad (1)$$

$$= \int_{\mathcal{S}} \beta_{t+1}(s_{t+1}) p(s_{t+1} \mid s_t, a_t) p(\mathcal{O}_t \mid s_t, a_t) ds_{t+1}. \quad (2)$$

- **Formulation 3;** Optimal target policy from $\beta_t(s_t, a_t)$, $\beta_t(s_t)$

$$\begin{aligned} p(a_t \mid s_t, \mathcal{O}_{t:T}) &= \frac{p(s_t, a_t \mid \mathcal{O}_{t:T})}{p(s_t \mid \mathcal{O}_{t:T})} = \frac{p(\mathcal{O}_{t:T} \mid s_t, a_t) p(a_t \mid s_t) p(s_t)}{p(\mathcal{O}_{t:T} \mid s_t) p(s_t)} \\ &\propto \frac{p(\mathcal{O}_{t:T} \mid s_t, a_t)}{p(\mathcal{O}_{t:T} \mid s_t)} = \frac{\beta_t(s_t, a_t)}{\beta_t(s_t)}, \end{aligned}$$

 **Decision 1**

Theory of control as probabilistic inference

Inference of posterior action conditional $p(a_t|s_t, \theta)$: Gives the optimal policy.

- Definition 3 (Soft Q): $Q(s_t, \mathbf{a}_t) = \log \beta_t(s_t, \mathbf{a}_t)$

- Definition 4 (Soft V): $V(s_t) = \log \beta_t(s_t)$

- From formulation 1:

$$\beta_t(s_t) = \int_{\mathcal{A}} \beta_t(s_t, a_t) da_t$$

If the scale of Q is large,
High Q-value dominates the other actions.



Soft Maximization

$$V(s_t) = \log \int_{\mathcal{A}} \exp(Q(s_t, \mathbf{a}_t)) d\mathbf{a}_t \approx \max_{\mathbf{a}_t} Q(s_t, \mathbf{a}_t).$$

- From formulation 2:

$$\beta_t(s_t, a_t) = \int_{\mathcal{S}} \beta_{t+1}(s_{t+1}) p(s_{t+1} | s_t, a_t) p(\mathcal{O}_t | s_t, a_t) ds_{t+1}.$$

$$Q(s_t, \mathbf{a}_t) = r(s_t, \mathbf{a}_t) + \log E_{s_{t+1} \sim p(s_{t+1} | s_t, \mathbf{a}_t)} [\exp(V(s_{t+1}))]$$

- From formulation 3:

$$p(a_t | s_t, \mathcal{O}_{t:T}) = \frac{\beta_t(s_t, a_t)}{\beta_t(s_t)} = \exp(Q(s_t, a_t) - V(s_t))$$

Theory of control as probabilistic inference

Problem of the drawn PGM; **Overly optimistic Q function**

- Stochastic dynamics Leads to **overly optimistic Q function**, risk seeking property.

$$V(s_t) = \log \int_{\mathcal{A}} \exp(Q(s_t, \mathbf{a}_t)) d\mathbf{a}_t \approx \max_{\mathbf{a}_t} Q(s_t, \mathbf{a}_t).$$

Another soft maximization term

$$Q(s_t, \mathbf{a}_t) = r(s_t, \mathbf{a}_t) + \log E_{s_{t+1} \sim p(s_{t+1} | s_t, \mathbf{a}_t)} [\exp(V(s_{t+1}))]$$

Where original bellman expectation equation:

$$V(s_t) = \max_{a_t} Q(s_t, a_t)$$

$$Q(s_t, \mathbf{a}_t) = r(s_t, \mathbf{a}_t) + E_{s_{t+1} \sim p(s_{t+1} | s_t, \mathbf{a}_t)} [V(s_{t+1})]$$

- EX): Two states for s_{t+1} ,
 - Win the lottery with $p(s_{t+1} = \text{win} | s_t, a_t = \text{buy}) = 0.000001$
 - Lose the lottery with $p(s_{t+1} = \text{lose} | s_t, a_t = \text{buy}) = 0.999999$

As $V(s_{t+1} = \text{win})$ dominates, always choose action to buy lottery.

Theory of control as probabilistic inference

Problem of the drawn PGM; Overly optimistic Q function

- Deterministic dynamics much more desirable.

$$V(s_t) = \log \int_{\mathcal{A}} \exp(Q(s_t, \mathbf{a}_t)) d\mathbf{a}_t \approx \max_{\mathbf{a}_t} Q(s_t, \mathbf{a}_t).$$

$$Q(s_t, \mathbf{a}_t) = r(s_t, \mathbf{a}_t) + V(s_{t+1})$$

Where original bellman expectation equation:

$$V(s_t) = \max_{a_t} Q(s_t, a_t)$$

$$Q(s_t, \mathbf{a}_t) = r(s_t, \mathbf{a}_t) + V(s_{t+1})$$

- We can **safely** use deterministic update of bellman equation, extract the policy.

$$p(a_t \mid s_t, \mathcal{O}_{t:T}) = \frac{\beta_t(s_t, a_t)}{\beta_t(s_t)} = \exp(Q(s_t, a_t) - V(s_t))$$

Then, is this framework only available in deterministic settings?

Theory of control as probabilistic inference

So what this inference of optimal conditioned policy exactly optimizes?

- Lets get back to the optimality conditioned trajectory distribution (**Target**).

$$p(\tau | \mathcal{O}_{1:T}) = \left[p(s_1) \prod_{t=1}^T p(s_{t+1} | s_t, \mathbf{a}_t) \right] \exp \left(\sum_{t=1}^T r(s_t, \mathbf{a}_t) \right) \quad (1)$$

- What we want **to model is policy**; So we formulate differently with above equation.

$$p(\tau | \mathcal{O}_{1:T}) = p(s_1 | \mathcal{O}_{1:T}) \cdot \prod_{t=1}^T p(a_t | s_t, \mathcal{O}_{1:T}) \cdot p(s_{t+1} | s_t, a_t, \mathcal{O}_{1:T})$$

- We can define approximate trajectory distribution **closest to the target distribution**.

$$\hat{p}(\tau) = p(s_1 | \mathcal{O}_{1:T}) \prod_{t=1}^T p(s_{t+1} | s_t, \mathbf{a}_t, \mathcal{O}_{1:T}) \pi(\mathbf{a}_t | s_t) \quad (2)$$

- Then make this optimization problem: (1), (2)

$$D_{\text{KL}}(\hat{p}(\tau) \| p(\tau)) = -E_{\tau \sim \hat{p}(\tau)} [\log p(\tau) - \log \hat{p}(\tau)].$$

Theory of control as probabilistic inference

So what this inference of optimal conditioned policy exactly optimizes?

- In deterministic case,

$$p(\tau | \mathcal{O}_{1:T}) = \left[p(\mathbf{s}_1) \prod_{t=1}^T p(\mathbf{s}_{t+1} | \mathbf{s}_t, \mathbf{a}_t) \right] \exp \left(\sum_{t=1}^T r(\mathbf{s}_t, \mathbf{a}_t) \right) \quad (1)$$

- Also most probable approximate distribution will be:

$$\begin{aligned} \hat{p}(\tau) &= p(\mathbf{s}_1 | \mathcal{O}_{1:T}) \prod_{t=1}^T p(\mathbf{s}_{t+1} | \mathbf{s}_t, \mathbf{a}_t, \mathcal{O}_{1:T}) \pi(\mathbf{a}_t | \mathbf{s}_t) \\ &= \hat{p}(\tau) = p(\mathbf{s}_1) \prod_{t=1}^T p(\mathbf{s}_{t+1} | \mathbf{s}_t, \mathbf{a}_t) \pi(\mathbf{a}_t | \mathbf{s}_t) \quad (3) \end{aligned}$$

- Then optimization problem solves: (1), (3)

$$D_{\text{KL}}(\hat{p}(\tau) \| p(\tau)) = -E_{\tau \sim \hat{p}(\tau)} [\log p(\tau) - \log \hat{p}(\tau)].$$

Theory of control as probabilistic inference

So what this inference of optimal conditioned policy exactly optimizes?

$$\begin{aligned} -D_{\text{KL}}(\hat{p}(\tau) \| p(\tau)) &= E_{\tau \sim \hat{p}(\tau)} \left[\log p(\mathbf{s}_1) + \sum_{t=1}^T (\log p(\mathbf{s}_{t+1} \mid \mathbf{s}_t, \mathbf{a}_t) + r(\mathbf{s}_t, \mathbf{a}_t)) \right. \\ &\quad \left. - \log p(\mathbf{s}_1) - \sum_{t=1}^T (\log p(\mathbf{s}_{t+1} \mid \mathbf{s}_t, \mathbf{a}_t) + \log \pi(\mathbf{a}_t \mid \mathbf{s}_t)) \right] \\ &= E_{\tau \sim \hat{p}(\tau)} \left[\sum_{t=1}^T r(\mathbf{s}_t, \mathbf{a}_t) - \log \pi(\mathbf{a}_t \mid \mathbf{s}_t) \right] \\ &= \sum_{t=1}^T E_{(\mathbf{s}_t, \mathbf{a}_t) \sim \hat{p}(\mathbf{s}_t, \mathbf{a}_t)} [r(\mathbf{s}_t, \mathbf{a}_t) - \log \pi(\mathbf{a}_t \mid \mathbf{s}_t)] \\ &= \sum_{t=1}^T E_{(\mathbf{s}_t, \mathbf{a}_t) \sim \hat{p}(\mathbf{s}_t, \mathbf{a}_t)} [r(\mathbf{s}_t, \mathbf{a}_t)] + E_{\mathbf{s}_t \sim \hat{p}(\mathbf{s}_t)} [\mathcal{H}(\pi(\mathbf{a}_t \mid \mathbf{s}_t))] \end{aligned}$$

In deterministic setting; Approximation of inference, equivalent to **max entropy RL**

Theory of control as probabilistic inference

So what this inference of optimal conditioned policy exactly optimizes?

- In stochastic case we cannot derive the same as deterministic.

$$p(\tau | \mathcal{O}_{1:T}) = \left[p(\mathbf{s}_1) \prod_{t=1}^T p(\mathbf{s}_{t+1} | \mathbf{s}_t, \mathbf{a}_t) \right] \exp \left(\sum_{t=1}^T r(\mathbf{s}_t, \mathbf{a}_t) \right) \quad (1)$$

$$\hat{p}(\tau) = p(\mathbf{s}_1 | \mathcal{O}_{1:T}) \prod_{t=1}^T p(\mathbf{s}_{t+1} | \mathbf{s}_t, \mathbf{a}_t, \mathcal{O}_{1:T}) \pi(\mathbf{a}_t | \mathbf{s}_t) \quad (2)$$

- We **cannot erase optimality conditioned dynamics** terms.
- Directly optimize : $D_{\text{KL}}(\hat{p}(\tau) \| p(\tau)) = -E_{\tau \sim \hat{p}(\tau)} [\log p(\tau) - \log \hat{p}(\tau)]$.
- Leads: $E_{\tau \sim \hat{p}(\tau)} \left[\log p(\mathbf{s}_1) + \sum_{t=1}^T r(\mathbf{s}_t, \mathbf{a}_t) + \log p(\mathbf{s}_{t+1} | \mathbf{s}_t, \mathbf{a}_t) \right] + \mathcal{H}(\hat{p}(\tau))$

So, to modeling the approximation as (2), we cannot use model-free algorithms.

Theory of control as probabilistic inference

Direct Solution! Abuse PGM.

- Why not we just model our approximate distribution **with (3)** rather than (2) ?

$$p(\tau | \mathcal{O}_{1:T}) = \left[p(\mathbf{s}_1) \prod_{t=1}^T p(\mathbf{s}_{t+1} | \mathbf{s}_t, \mathbf{a}_t) \right] \exp \left(\sum_{t=1}^T r(\mathbf{s}_t, \mathbf{a}_t) \right) \quad (1)$$

$$\hat{p}(\tau) = p(\mathbf{s}_1 | \mathcal{O}_{1:T}) \prod_{t=1}^T p(\mathbf{s}_{t+1} | \mathbf{s}_t, \mathbf{a}_t, \mathcal{O}_{1:T}) \pi(\mathbf{a}_t | \mathbf{s}_t) \quad (2)$$

$$= \hat{p}(\tau) = p(\mathbf{s}_1) \prod_{t=1}^T p(\mathbf{s}_{t+1} | \mathbf{s}_t, \mathbf{a}_t) \pi(\mathbf{a}_t | \mathbf{s}_t) \quad (3)$$

- Pros: We can **recover the maximum entropy RL objective**, like deterministic case.
Agents cannot control dynamical systems like (2), not intuitive.
(Is there anyone who can make environment to always win the lottery?)
- Cons: We abuse the PGM w.r.t. optimality -> No theoretical grounds (yet)

Theory of control as probabilistic inference

Connection to Structural Variational Inference

- The theoretical grounds of abused solution!
- Target distribution: $p(\tau|\mathcal{O}_{1:T}) = \left[p(\mathbf{s}_1) \prod_{t=1}^T p(\mathbf{s}_{t+1} | \mathbf{s}_t, \mathbf{a}_t) \right] \exp \left(\sum_{t=1}^T r(\mathbf{s}_t, \mathbf{a}_t) \right)$ (1)
- Variational approximate distribution: $\hat{q}(\tau) = q(\mathbf{s}_1) \prod_{t=1}^T q(\mathbf{s}_{t+1} | \mathbf{s}_t, \mathbf{a}_t) q(\mathbf{a}_t | \mathbf{s}_t)$. (2) *q, not p*
- With given observation, **ELBO with variational trajectory distribution:**

$$\begin{aligned} \log p(\mathcal{O}_{1:T}) &= \log \iint p(\mathcal{O}_{1:T}, \mathbf{s}_{1:T}, \mathbf{a}_{1:T}) d\mathbf{s}_{1:T} d\mathbf{a}_{1:T} \\ &= \log \iint p(\mathcal{O}_{1:T}, \mathbf{s}_{1:T}, \mathbf{a}_{1:T}) \frac{q(\mathbf{s}_{1:T}, \mathbf{a}_{1:T})}{q(\mathbf{s}_{1:T}, \mathbf{a}_{1:T})} d\mathbf{s}_{1:T} d\mathbf{a}_{1:T} \\ &= \log E_{(\mathbf{s}_{1:T}, \mathbf{a}_{1:T}) \sim q(\mathbf{s}_{1:T}, \mathbf{a}_{1:T})} \left[\frac{p(\mathcal{O}_{1:T}, \mathbf{s}_{1:T}, \mathbf{a}_{1:T})}{q(\mathbf{s}_{1:T}, \mathbf{a}_{1:T})} \right] \\ &\geq E_{(\mathbf{s}_{1:T}, \mathbf{a}_{1:T}) \sim q(\mathbf{s}_{1:T}, \mathbf{a}_{1:T})} [\log p(\mathcal{O}_{1:T}, \mathbf{s}_{1:T}, \mathbf{a}_{1:T}) - \log q(\mathbf{s}_{1:T}, \mathbf{a}_{1:T})] \end{aligned}$$

Theory of control as probabilistic inference

Connection to Structural Variational Inference

- The theoretical grounds of abused solution!

- Target distribution: $p(\tau|O_{1:T}) = \left[p(s_1) \prod_{t=1}^T p(s_{t+1} | s_t, a_t) \right] \exp \left(\sum_{t=1}^T r(s_t, a_t) \right)$ (1)

- Variational approximate distribution: $\hat{q}(\tau) = q(s_1) \prod_{t=1}^T q(s_{t+1} | s_t, a_t) q(a_t | s_t)$ (2) *q, not p*

$$\log p(\mathcal{O}_{1:T}) \geq E_{(s_{1:T}, a_{1:T}) \sim q(s_{1:T}, a_{1:T})} [\log p(\mathcal{O}_{1:T}, s_{1:T}, a_{1:T}) - \log q(s_{1:T}, a_{1:T})]$$



put (1), (2)

$$\log p(\mathcal{O}_{1:T}) \geq E_{(s_{1:T}, a_{1:T}) \sim q(s_{1:T}, a_{1:T})} \left[\sum_{t=1}^T r(s_t, a_t) - \log q(a_t | s_t) \right]$$

- Then we set dynamics distributions q with $p(s_1), p(s_{t+1}|s_t, a_t) \rightarrow$ **Recovers Max-Entropy RL**

$$= \sum_{t=1}^T E_{(s_t, a_t) \sim \hat{p}(s_t, a_t)} [r(s_t, a_t) + \mathcal{H}(\pi(a_t | s_t))]$$

- Abused distribution is one of the ELBO, derivation of our designed PGM (**Got Theoretical Grounds.**)

Theory of control as probabilistic inference

How can we extract policy from the abused settings (ELBO)

- We already discussed before **soft- Q, V** and **optimal policy** for **stochastic/deterministic** settings.
- What about **abused settings**, any difference with $\exp(Q(s_t, a_t) - V(s_t))$?

From the max-entropy objective function:

$$\sum_{t=1}^T \mathbb{E}_{(\mathbf{s}_t, \mathbf{a}_t) \sim \hat{p}(\mathbf{s}_t, \mathbf{a}_t)} [r(\mathbf{s}_t, \mathbf{a}_t) + \mathcal{H}(\pi(\mathbf{a}_t | \mathbf{s}_t))]$$

If $t = T$:

$$\begin{aligned} & \mathbb{E}_{(s_T, a_T) \sim \hat{p}(s_T, a_T)} [r(s_T, a_T) - \log \pi(a_T | s_T)] \\ &= \mathbb{E}_{s_T \sim \hat{p}(s_T)} [\mathbb{E}_{a_T \sim \pi(\cdot | s_T)} [r(s_T, a_T) - \log \pi(a_T | s_T)]] \\ &= \mathbb{E}_{s_T \sim \hat{p}(s_T)} [\mathbb{E}_{a_T \sim \pi(\cdot | s_T)} [-\log \pi(a_T | s_T) + r(s_T, a_T)]] \\ &= \mathbb{E}_{s_T \sim \hat{p}(s_T)} \left[-D_{\text{KL}} \left(\pi(a_T | s_T) \parallel \frac{1}{Z(s_T)} \exp(r(s_T, a_T)) \right) + \log Z(s_T) \right] \\ &= \mathbb{E}_{s_T \sim \hat{p}(s_T)} \left[-D_{\text{KL}} \left(\pi(a_T | s_T) \parallel \frac{1}{\exp(V(s_T))} \exp(r(s_T, a_T)) \right) + V(s_T) \right] \end{aligned}$$

Define:

$$V(s_T) = \log \int_{\mathcal{A}} \exp(r(s_T, a_t)) da_T$$

$$\pi^*(a_T | s_t) = \exp(r(s_T, a_T))$$

Theory of control as probabilistic inference

How can we extract policy from the abused settings (ELBO)

If $t = t$ (intermediate steps):

- Policy should also consider the **future expected value**. (Summation of $t = 1 \sim T$ recovers original objective.)

$$\begin{aligned} & \mathbb{E}_{(s_t, a_t) \sim \hat{p}(s_t, a_t)} \left[r(s_t, a_t) - \log \pi(a_t | s_t) + \underbrace{\mathbb{E}_{s_{t+1} \sim p(\cdot | s_t, a_t)} [V(s_{t+1})]} \right] \\ \stackrel{(1)}{=} & \mathbb{E}_{(s_t, a_t) \sim \hat{p}} \left[\underbrace{r(s_t, a_t) + \mathbb{E}_{s_{t+1}} [V(s_{t+1})]}_{:= Q(s_t, a_t)} - \log \pi(a_t | s_t) \right] \\ & \hspace{10em} \text{Define, same as Bellman Expectation Eq.} \\ & \hspace{10em} \text{We can mitigate risk-seeking } Q \text{ Problem!} \\ = & \mathbb{E}_{s_t \sim \hat{p}(s_t)} \left[\mathbb{E}_{a_t \sim \pi(\cdot | s_t)} [Q(s_t, a_t) - \log \pi(a_t | s_t)] \right] \\ \stackrel{(2)}{=} & \mathbb{E}_{s_t \sim \hat{p}(s_t)} \left[-D_{\text{KL}} \left(\pi(\cdot | s_t) \parallel \frac{1}{Z(s_t)} \exp(Q(s_t, \cdot)) \right) + \log Z(s_t) \right] \\ & \hspace{10em} \text{Now defined on } Q \\ & \hspace{10em} V(s_t) = \log \int_{\mathcal{A}} \exp(Q(s_t, a_t)) da_t \\ = & \boxed{\mathbb{E}_{s_t \sim \hat{p}(s_t)} \left[-D_{\text{KL}} \left(\pi(a_t | s_t) \parallel \frac{1}{\exp(V(s_t))} \exp(Q(s_t, a_t)) \right) + V(s_t) \right]} \\ & \hspace{15em} \pi^*(\mathbf{a}_t | \mathbf{s}_t) = \exp(Q(\mathbf{s}_t, \mathbf{a}_t) - V(\mathbf{s}_t)), \end{aligned}$$

In Summary:

1. We want to transform the control problem -> Probabilistic Inference.
2. Draw PGM which can infer $p(a_t|s_t, O_{1:T})$, $p(\tau|O_{1:T})$ posteriors.
3. Decompose and **got optimal target distribution**: $p(\tau|O_{1:T})$
4. To find out optimal policy, defined **soft Q, V** functions and found **analytical form of $\pi^*(a_t|s_t)$**
Stochastic dynamics has risk-seeking problem with Q function.
5. With defining optimal approximate distribution (closest distribution),
Transformed inference problem to optimization problem.
Again, problem with stochastic dynamics setting has problem. (Only model-based works)
6. We abused the PGM, and modeled approximate distribution differently.
This reveals that one of the **ELBO** of PGM := Abused distribution.
7. Abused distribution recovers **Max Entropy RL**, no problem of **risk-seeking Q function**.
Most of the Max-Entropy RL algorithms built on this setting.

Applications of Control as Probabilistic Inference

Maximum Entropy Policy Gradients

- Policy Gradients with the Max-Entropy Objective

$$\begin{aligned}\nabla_{\theta} J(\theta) &= \sum_{t=1}^T \nabla_{\theta} E_{(\mathbf{s}_t, \mathbf{a}_t) \sim q(\mathbf{s}_t, \mathbf{a}_t)} [r(\mathbf{s}_t, \mathbf{a}_t) + \mathcal{H}(q_{\theta}(\mathbf{a}_t | \mathbf{s}_t))] \\ &= \sum_{t=1}^T E_{(\mathbf{s}_t, \mathbf{a}_t) \sim q(\mathbf{s}_t, \mathbf{a}_t)} \left[\nabla_{\theta} \log q_{\theta}(\mathbf{a}_t | \mathbf{s}_t) \left(\sum_{t'=t}^T r(\mathbf{s}_{t'}, \mathbf{a}_{t'}) - \log q_{\theta}(\mathbf{a}_{t'} | \mathbf{s}_{t'}) - 1 \right) \right] \\ &= \sum_{t=1}^T E_{(\mathbf{s}_t, \mathbf{a}_t) \sim q(\mathbf{s}_t, \mathbf{a}_t)} \left[\nabla_{\theta} \log q_{\theta}(\mathbf{a}_t | \mathbf{s}_t) \left(\sum_{t'=t}^T r(\mathbf{s}_{t'}, \mathbf{a}_{t'}) - \log q_{\theta}(\mathbf{a}_{t'} | \mathbf{s}_{t'}) - b(\mathbf{s}_{t'}) \right) \right],\end{aligned}$$

$$\nabla_{\theta} J(\theta) = \sum_{t=1}^T E_{(\mathbf{s}_t, \mathbf{a}_t) \sim q(\mathbf{s}_t, \mathbf{a}_t)} \left[\nabla_{\theta} \log q_{\theta}(\mathbf{a}_t | \mathbf{s}_t) \hat{A}(\mathbf{s}_t, \mathbf{a}_t) \right]$$

- We can just put $-\log q_{\theta}(\mathbf{a}_t | \mathbf{s}_t)$ **penalty** to the original reward, and use the policy gradient algorithms.

Applications of Control as Probabilistic Inference

Soft Q Learning

- With below two equations, definition of soft V , and its optimal policy $q(a_t|s_t)$.

$$V(s_t) = \log \int_{\mathcal{A}} \exp(Q(s_t, \mathbf{a}_t)) d\mathbf{a}_t$$
$$q(\mathbf{a}_t | s_t) = \exp(Q(s_t, \mathbf{a}_t) - V(s_t))$$

- We can parameterize Q_ϕ only to express both equations.

- Training Q_ϕ :
$$\mathcal{E}(\phi) = E_{(\mathbf{s}_t, \mathbf{a}_t) \sim q(\mathbf{s}_t, \mathbf{a}_t)} \left[\left(r(\mathbf{s}_t, \mathbf{a}_t) + E_{q(\mathbf{s}_{t+1}|\mathbf{s}_t, \mathbf{a}_t)} [V_\psi(\mathbf{s}_{t+1})] - Q_\phi(\mathbf{s}_t, \mathbf{a}_t) \right)^2 \right]$$
$$\phi \leftarrow \phi - \alpha E \left[\frac{dQ_\phi}{d\phi}(\mathbf{s}_t, \mathbf{a}_t) \left(Q_\phi(\mathbf{s}_t, \mathbf{a}_t) - \underbrace{\left(r(\mathbf{s}_t, \mathbf{a}_t) + \log \int_{\mathcal{A}} \exp(Q(\mathbf{s}_{t+1}, \mathbf{a}_{t+1})) d\mathbf{a}_{t+1} \right)}_{\log \mathbb{E}_{q(a_t)} \frac{\exp(Q(s_{t+1}, a_{t+1}))}{q(a_t)}} \right) \right].$$

Importance Sampling with current policy

- Training π_θ :

Stein Variational Gradient Descent (**SVGD**) to match: $\pi_\theta(a_t|s_t) \propto \exp(Q_{soft}^\phi(s_t, \cdot))$

Applications of Control as Probabilistic Inference

Soft Actor Critic

- Explicit modeling of the policy network $\pi(a_t|s_t; \theta)$

$$\mathcal{T}^\pi Q(s_t, \mathbf{a}_t) \triangleq r(s_t, \mathbf{a}_t) + \gamma \mathbb{E}_{\mathbf{s}_{t+1} \sim p} [V(\mathbf{s}_{t+1})]$$

$$V(s_t) = \mathbb{E}_{\mathbf{a}_t \sim \pi} [Q(s_t, \mathbf{a}_t) - \log \pi(\mathbf{a}_t | s_t)] \quad \leftarrow \text{Check lemma 1 from Haarnoja (2018)}$$

- Train both soft V, Q with MSE loss:

$$\mathcal{E}(\phi) = E_{(s_t, \mathbf{a}_t) \sim q(s_t, \mathbf{a}_t)} \left[\left(r(s_t, \mathbf{a}_t) + E_{q(\mathbf{s}_{t+1} | s_t, \mathbf{a}_t)} [V_\psi(\mathbf{s}_{t+1})] - Q_\phi(s_t, \mathbf{a}_t) \right)^2 \right]$$

$$\mathcal{E}(\psi) = E_{s_t \sim q(s_t)} \left[\left(E_{\mathbf{a}_t \sim q(\mathbf{a}_t | s_t)} [Q_\phi(s_t, \mathbf{a}_t) - \log q(\mathbf{a}_t | s_t)] - V_\psi(s_t, \mathbf{a}_t) \right)^2 \right].$$

- Training π_θ :

$$J_\pi(\theta) = \mathbb{E}_{s_t \sim \mathcal{D}} \left[D_{\text{KL}} \left(\pi_\theta(\cdot | s_t) \parallel \frac{\exp(Q_\phi(s_t, \cdot))}{Z_\theta(s_t)} \right) \right].$$

$$J_\pi(\theta) = \mathbb{E}_{s_t \sim \mathcal{D}, \epsilon_t \sim \mathcal{N}} [\log \pi_\theta(f_\theta(\epsilon_t; s_t) | s_t) - Q_\phi(s_t, f_\theta(\epsilon_t; s_t))],$$

$$\hat{\nabla}_\theta J_\pi(\theta) = \nabla_\theta \log \pi_\theta(\mathbf{a}_t | s_t)$$

$$+ (\nabla_{\mathbf{a}_t} \log \pi_\theta(\mathbf{a}_t | s_t) - \nabla_{\mathbf{a}_t} Q(s_t, \mathbf{a}_t)) \nabla_\theta f_\theta(\epsilon_t; s_t)$$

Applications of Control as Probabilistic Inference

Replacing Prior $p(a_t|s_t)$

- **Decision Choice 1:** As we were not interested in action prior (online), we can take it as simplest:

$$p(a_t) = p(a_t|s_t) = \frac{1}{|\mathcal{A}_t|}$$

- What if we assume the prior as $p(a_t|s_t) = \pi_{old}(a_t|s_t)$ or $\pi_{behavior}(a_t|s_t)$?

$$J(\pi) = \mathbb{E}_{\tau \sim \pi} \left[\sum_{t=0}^T r(s_t, a_t) + \alpha \mathcal{H}(\pi(\cdot | s_t)) \right]$$

- Max Entropy objective leads to KL regularized objective. TRPO/PPO/AWR....

$$J(\pi) = \mathbb{E}_{\tau \sim \pi} \left[\sum_{t=0}^T r(s_t, a_t) \right] - \alpha \mathbb{E}_{s \sim \rho_\pi} [D_{KL}(\pi(\cdot | s) \parallel \pi_{old}(\cdot | s))]$$

- Trajectory planning. **Diffuser**

Modeling Entire trajectory distribution with offline dataset.

$$p(\tau | \mathcal{O}_{1:T}) = p(\tau) \exp \left(\sum_{t=1}^T r(s_t, a_t) \right)$$

Future Plan of Research

1. Constrained black-box optimization with posterior inference. (Taeyoung Yun, Kyuil Sim) ~NIPS(2025)
2. Probabilistic Inference for sequential decision making. (Diffusion meets control/Control meets diffusion)

EOD
