

Informed Asymmetric Actor-Critic: Theoretical Insights and Open Questions

EWRL 2025 | Tübingen, Germany

Daniel Ebi (daniel.ebi@kit.edu) | Gaspard Lambrechts | Damien Ernst | Klemens Böhm

Asymmetric Observability

In many RL applications, agents must act under **partial observability**. However, assuming **identical observability** during training and execution is often **too restrictive**.

Asymmetric RL relaxes this assumption: It leverages additional state information at training time to accelerate learning.

Informed Partially Observable MDP

We consider an **informed POMDP** [1] $(\mathcal{S}, \mathcal{A}, \mathcal{I}, \mathcal{O}, T, I, O, R, P, \gamma)$:

- States $\mathbf{s}_t \in \mathcal{S}$,
- Actions $\mathbf{a}_t \in \mathcal{A}$,
- Information $\mathbf{i}_t \in \mathcal{I}$,
- Observations $\mathbf{o}_t \in \mathcal{O}$,
- Transition $\mathbf{s}_{t+1} \sim T(\cdot | \mathbf{s}_t, \mathbf{a}_t)$,
- Cognition $\mathbf{i}_t \sim I(\cdot | \mathbf{s}_t)$,
- Perception $\mathbf{o}_t \sim O(\cdot | \mathbf{i}_t)$,
- Rewards $r_t = R(\mathbf{s}_t, \mathbf{a}_t)$,
- Initial state $\mathbf{s}_0 \sim P(\cdot)$,
- Discount $\gamma \in [0, 1)$.

At execution, the agent observes $\mathbf{o}_t \sim \tilde{O}(\cdot | \mathbf{s}_t) = \sum_{i \in \mathcal{I}} O(\cdot | \mathbf{i}) I(\mathbf{i} | \mathbf{s}_t)$.

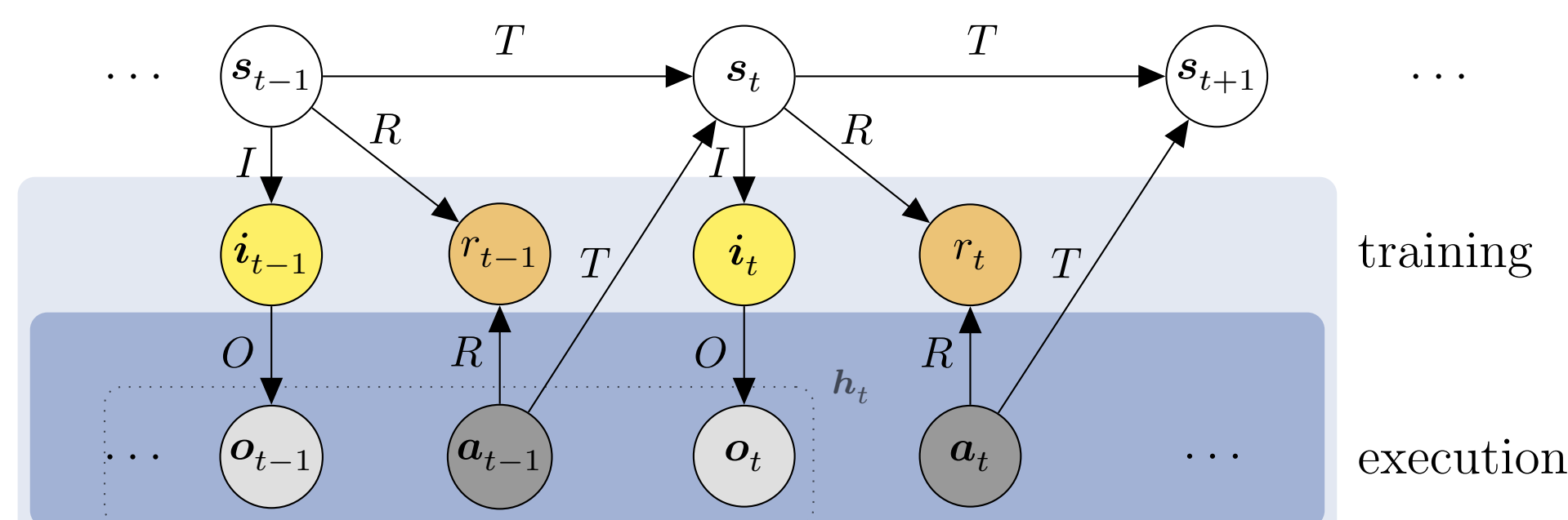


Figure 1 Bayesian network of an informed POMDP.

Objective: Given histories $\mathbf{h}_t = [\mathbf{o}_0 \mathbf{a}_0 \dots \mathbf{o}_t]^\top \in \mathcal{H}$, we want the agent to learn a policy $\pi : \mathcal{H} \rightarrow \Delta(\mathcal{A})$ that maximizes $J(\pi) = \mathbb{E}^\pi [\sum_{t=0}^{\infty} \gamma^t r_t]$.

Unbiased Learning with Informed Critic

In **actor-critic** methods, the critic is not used during execution. → It can leverage additional information, typically the true state $\mathbf{s}_t$ [2].

Informed history-dependent value functions

We introduce $Q^\pi(\mathbf{h}, \mathbf{i}, \mathbf{a})$ and $V^\pi(\mathbf{h}, \mathbf{i})$ and prove that they are unbiased, i.e., $\mathbb{E}_{\mathbf{i}|\mathbf{h}}[Q^\pi(\mathbf{h}, \mathbf{i}, \mathbf{a})] = Q^\pi(\mathbf{h}, \mathbf{a})$ and $\mathbb{E}_{\mathbf{i}|\mathbf{h}}[V^\pi(\mathbf{h}, \mathbf{i})] = V^\pi(\mathbf{h})$.

Informed asymmetric policy gradient

We derive the informed asymmetric policy gradient $\nabla_{\theta}^{\text{IAAC}} J(\pi_{\theta})$ and establish its equivalence to the symmetric policy gradient.

Theorem 1. Informed asymmetric policy gradient equivalence.

$$\begin{aligned} \nabla_{\theta}^{\text{IAAC}} J(\pi_{\theta}) &:= \mathbb{E} \left[\sum_{t=0}^{\infty} \gamma^t Q^{\pi_{\theta}}(\mathbf{h}_t, \mathbf{i}_t, \mathbf{a}_t) \nabla_{\theta} \log \pi_{\theta}(\mathbf{a}_t | \mathbf{h}_t) \right] \\ &\stackrel{(*)}{=} \mathbb{E} \left[\sum_{t=0}^{\infty} \gamma^t Q^{\pi_{\theta}}(\mathbf{h}_t, \mathbf{a}_t) \nabla_{\theta} \log \pi_{\theta}(\mathbf{a}_t | \mathbf{h}_t) \right] =: \nabla_{\theta} J(\pi_{\theta}) \end{aligned}$$

(*) Follows from expectation properties and our unbiasedness results.

Hence, informing the critic with state-conditioned information yields unbiased policy gradients beyond the special case $\mathbf{i}_t = \mathbf{s}_t$.

Informed Asymmetric Rec-NAC

We propose an **informed asymmetric variant** of the recurrent natural actor-critic (Rec-NAC) algorithm [3], where the **informed critic** receives additional input $\mathbf{i}_t$, breaking symmetry with the actor.

Algorithm 1. Recurrent Natural Actor-Critic (Rec-NAC)

1. Initialize actor parameters θ_0 and critic parameters ϑ_0 .
2. For $t = 0$ to $T-1$:
 - a. Estimate $\hat{Q}_{\vartheta}^{\pi_{\theta_t}}$ with $\hat{\vartheta} = \frac{1}{K} \sum_{k=0}^{K-1} \vartheta_k$. (via K steps of **Rec-TD**)
 - b. Estimate gradient $\hat{\mathbf{g}}_{t,n}$. (via N steps of **Rec-NPG**)
 - c. Update policy parameters: $\theta_{t+1} = \theta_t + \eta_{\text{npg}} \cdot \frac{1}{N} \sum_{n=0}^{N-1} \hat{\mathbf{g}}_{t,n}$.
3. Return final policy π_{θ_T} .

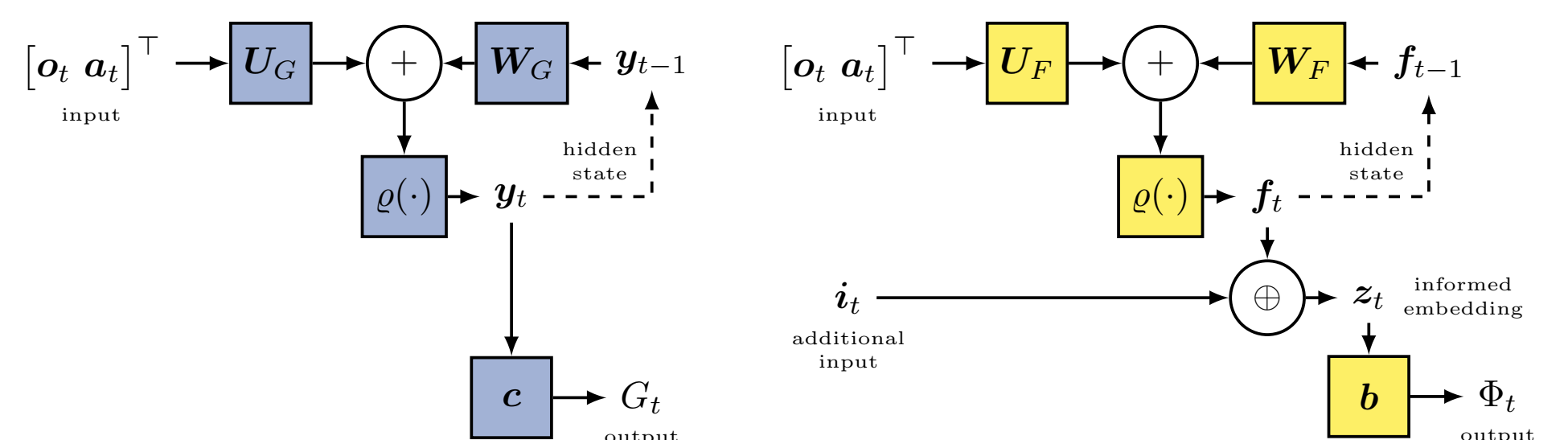


Figure 2 Architecture of the actor (left) and the informed critic (right).

Finite-Time and Finite-Width Analysis

We study the effect of an informed critic on Rec-TD convergence.

Theorem 2. Error bound of informed asymmetric Rec-TD.

Under standard smoothness and regularization assumptions, for a critic of width $m_F \in \mathbb{N}$ and $K \in \mathbb{N}$ gradient steps:

$$\begin{aligned} \mathbb{E} \left[\frac{1}{K} \sum_{k=0}^{K-1} \mathcal{R}_T^{\pi}(\vartheta_k) \right] &\leq \frac{\|\nu\|_2^2}{\eta_{td} (1-\gamma)\sqrt{K}} + \frac{\eta_{td} C_T^{(1)}}{(1-\gamma)^3\sqrt{K}} \\ &\quad + \frac{C_T^{(2)}}{(1-\gamma)^2\sqrt{m_F}} + \frac{\gamma^T}{(1-\gamma)K} \sum_{k=0}^{K-1} \omega_{T,k}^2. \end{aligned}$$

Trade-off: increased complexity vs. approximation gain

Rec-TD with informed critic achieves the **same asymptotic rate** as in the symmetric case [3], but the error bound reveals a key **trade-off**:

- | | |
|--|---|
| <p>↑ Complexity</p> <ul style="list-style-type: none"> ■ Generally larger Lipschitz and smoothness constants. ■ Increased dimensionality of parameter-related quantities. | <p>↓ Approximation error</p> <ul style="list-style-type: none"> ■ Richer function class from additional information. ■ Informative $\mathbf{i}_t$ may mitigate partial observability. |
|--|---|

Conclusion

Additional inputs improve learning only if their informativeness outweighs the added complexity.

Central question for future work:

How can we formalize this trade-off to guide auxiliary input design?

References

1. G. Lambrechts, A. Bolland, and D. Ernst. Informed POMDP: Leveraging additional information in model-based RL. Reinforcement Learning Journal, 2024.
2. A. Baisero and C. Amato. Unbiased asymmetric reinforcement learning under partial observability. 21st Intl. Conference on Autonomous Agents and Multiagent Systems, 2022.
3. S. Cayci and A. Eryilmaz. Recurrent natural policy gradient for POMDPs. 17th European Workshop of Reinforcement Learning (EWRL), 2024.

