

What we have seen:

- RL MPC, why does it work & different flavors

What we have seen:

- RL MPC, why does it work & different flavors

What we will do now:

- The theory is not about MPC, it is about MDPs (MPC is a special case)
- It has broader implications for AI and model-based decision making
- Provide (tentative) practical ideas for in Sim2Real

AI for Operational Decisions

Some perspectives from recent research

Prof. Sebastien Gros

Department of Cybernetic
Faculty of Information Technology
NTNU

Freiburg PhD School

Outline

- 1 Decisions from data
- 2 AI models for Decision



Model-Free Pathway - “Pure” Reinforcement Learning



- Define policy π_{θ} as a parametrized function
- Optimize policy as:

$$\min_{\theta} J(\pi_{\theta})$$

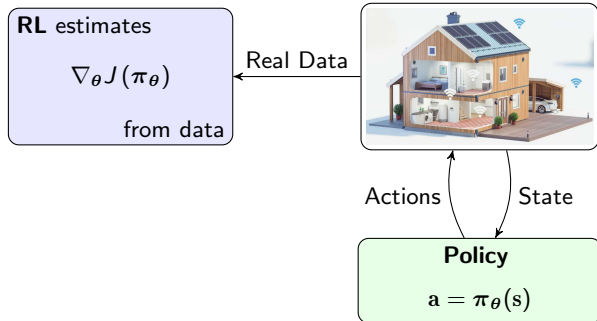
Actions

State

Policy

$$\mathbf{a} = \pi_{\theta}(\mathbf{s})$$

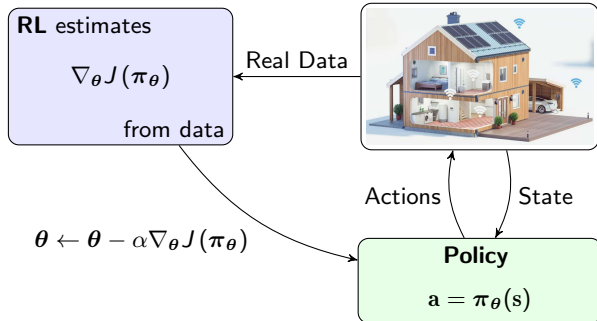
Model-Free Pathway - “Pure” Reinforcement Learning



- Define policy π_{θ} as a parametrized function
- Optimize policy as:

$$\min_{\theta} J(\pi_{\theta})$$

Model-Free Pathway - “Pure” Reinforcement Learning



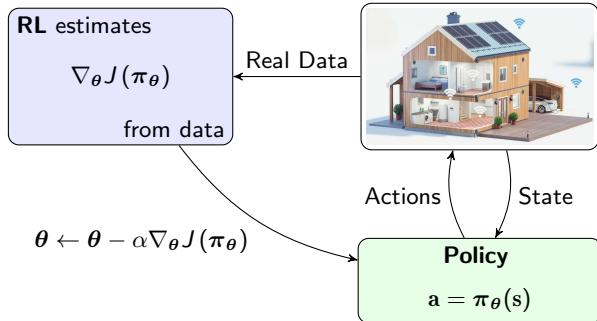
- Define policy π_{θ} as a parametrized function

- Optimize policy as:

$$\min_{\theta} J(\pi_{\theta})$$

- Gradient descent can be used to optimize policy

Model-Free Pathway - “Pure” Reinforcement Learning



- Define policy π_{θ} as a parametrized function
- Optimize policy as:

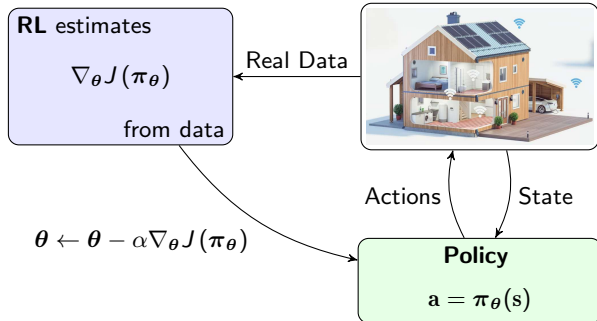
$$\min_{\theta} J(\pi_{\theta})$$

- Gradient descent can be used to optimize policy

Remarks

- $\pi_{\theta}(s)$ often from DNN
- End-to-end training requires a lot of data and effort
- Difficult to provide explainability and guarantees on the policy we obtain
- Poor track record of industrial adoption (exception chatbots)

Model-Free Pathway - “Pure” Reinforcement Learning



- Define policy π_{θ} as a parametrized function

- Optimize policy as:

$$\min_{\theta} J(\pi_{\theta})$$

- Gradient descent can be used to optimize policy

Remarks

- $\pi_{\theta}(s)$ often from DNN
- End-to-end training requires a lot of data and effort
- Difficult to provide explainability and guarantees on the policy we obtain
- Poor track record of industrial adoption (exception chatbots)

→ policy is often trained on **high-fidelity simulations**: embed knowledge, reduce costs, manage safety, promote explainability, control the training

Model-based Pathway - Predictive AI Models (data-driven, ML, etc)

Real world: $s_{k+1} \sim \rho(\cdot | s_k, a_k)$

- Data with enough “richness”

Model-based Pathway - Predictive AI Models (data-driven, ML, etc)

One-step model

AI model e.g.

$$\hat{s}_{k+1} \sim \rho_{\theta}(\cdot | s_k, a_k)$$

fitting data

Real world: $s_{k+1} \sim \rho(\cdot | s_k, a_k)$

- Data with enough “richness”

Here we no longer need the model ρ_{θ} to be “optimization-friendly”

Model-based Pathway - Predictive AI Models (data-driven, ML, etc)

One-step model

AI model e.g.

$$\hat{s}_{k+1} \sim \rho_{\theta}(\cdot | s_k, a_k)$$

fitting data

Real world: $s_{k+1} \sim \rho(\cdot | s_k, a_k)$

- Data with enough “richness”

Here we no longer need the model ρ_{θ} to be “optimization-friendly”

Model parameters θ

- Selected such that model ρ_{θ} “resembles” reality ρ , using data (+ physics)
- Deterministic models are a special case of ρ_{θ}
- Many methods, from Least Squares and MLE to Bayesian and Adversarial Learning
- Still, in most cases ρ_{θ} “simplifies” ρ because
 - ▶ need for huge amounts of data to decide θ if model is very rich
 - ▶ computationally demanding simulations if ρ_{θ} is expensive to sample from
- Then
 - ▶ biases in case of insufficiently rich model structure
 - ▶ validity is limited to some parts of the state-action space
 - ▶ only few first moments are correct (typ. mean + variance)

Model-based Pathway - Predictive AI Models (data-driven, ML, etc)

One-step model

AI model e.g.

$$\hat{s}_{k+1} \sim \rho_{\theta}(\cdot | s_k, a_k)$$

fitting data

Virtual data from “simulating”
 $\hat{s}_{k+1} \sim \rho_{\theta}(\cdot | s_k, a_k)$ forward (Monte Carlo)

Here we no longer need the model ρ_{θ} to be “optimization-friendly”

Model parameters θ

- Selected such that model ρ_{θ} “resembles” reality ρ , using data (+ physics)
- Deterministic models are a special case of ρ_{θ}
- Many methods, from Least Squares and MLE to Bayesian and Adversarial Learning
- Still, in most cases ρ_{θ} “simplifies” ρ because
 - ▶ need for huge amounts of data to decide θ if model is very rich
 - ▶ computationally demanding simulations if ρ_{θ} is expensive to sample from
- Then
 - ▶ biases in case of insufficiently rich model structure
 - ▶ validity is limited to some parts of the state-action space
 - ▶ only few first moments are correct (typ. mean + variance)

Model-based Pathway - Predictive AI Models (data-driven, ML, etc)

One-step model

AI model e.g.

$$\hat{s}_{k+1} \sim \rho_{\theta}(\cdot | s_k, a_k)$$

fitting data

Virtual data from “simulating”
 $\hat{s}_{k+1} \sim \rho_{\theta}(\cdot | s_k, a_k)$ forward (Monte Carlo)

In most cases, ρ_{θ} is a fairly simple representation of reality ρ , easy to sample from

Here we no longer need the model ρ_{θ} to be “optimization-friendly”

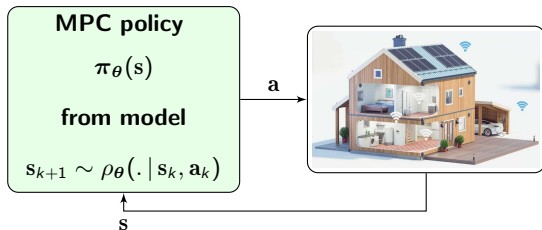
Model parameters θ

- Selected such that model ρ_{θ} “resembles” reality ρ , using data (+ physics)
- Deterministic models are a special case of ρ_{θ}
- Many methods, from Least Squares and MLE to Bayesian and Adversarial Learning
- Still, in most cases ρ_{θ} “simplifies” ρ because
 - ▶ need for huge amounts of data to decide θ if model is very rich
 - ▶ computationally demanding simulations if ρ_{θ} is expensive to sample from
- Then
 - ▶ biases in case of insufficiently rich model structure
 - ▶ validity is limited to some parts of the state-action space
 - ▶ only few first moments are correct (typ. mean + variance)

Back to MPC - A “Classical” View on Decision Making

Policy $\pi_{\theta}(s)$ given implicitly by

$$\begin{aligned} \min_{\mathbf{x}, \mathbf{u}} \quad & \mathbb{E} \left[T(\mathbf{x}_N) + \sum_{k=0}^{N-1} L(\mathbf{x}_k, \mathbf{u}_k) \right] \\ \text{s.t.} \quad & \mathbf{x}_{k+1} \sim \rho_{\theta}(\mathbf{x}_k, \mathbf{u}_k) \\ & \mathbf{x}_0 = \mathbf{s} \\ & \text{use } \pi_{\theta}(\mathbf{s}) = \mathbf{u}_0^* \text{ on the system} \end{aligned}$$



Back to MPC - A “Classical” View on Decision Making

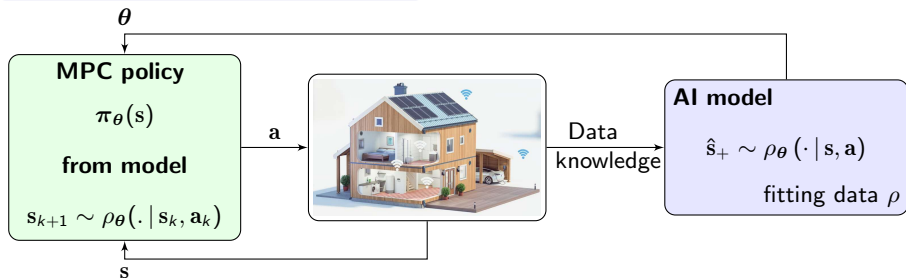
Policy $\pi_{\theta}(s)$ given implicitly by

$$\min_{\mathbf{x}, \mathbf{u}} \quad \mathbb{E} \left[T(\mathbf{x}_N) + \sum_{k=0}^{N-1} L(\mathbf{x}_k, \mathbf{u}_k) \right]$$

$$\text{s.t.} \quad \mathbf{x}_{k+1} \sim \rho_{\theta}(\mathbf{x}_k, \mathbf{u}_k)$$

$$\mathbf{x}_0 = \mathbf{s}$$

use $\pi_{\theta}(s) = \mathbf{u}_0^*$ on the system



Back to MPC - A “Classical” View on Decision Making

Policy $\pi_{\theta}(s)$ given implicitly by

$$\min_{\mathbf{x}, \mathbf{u}} \quad \mathbb{E} \left[T(\mathbf{x}_N) + \sum_{k=0}^{N-1} L(\mathbf{x}_k, \mathbf{u}_k) \right]$$

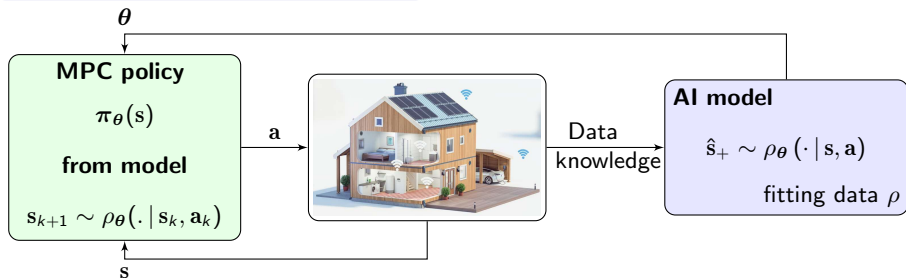
$$\text{s.t.} \quad \mathbf{x}_{k+1} \sim \rho_{\theta}(\mathbf{x}_k, \mathbf{u}_k)$$

$$\mathbf{x}_0 = \mathbf{s}$$

use $\pi_{\theta}(s) = \mathbf{u}_0^*$ on the system

Defines a paradigm...

- Performance from model accuracy
i.e. ρ_{θ} “close enough” to ρ
- “Ignore” that we replan all the time



Back to MPC - A “Classical” View on Decision Making

Policy $\pi_{\theta}(s)$ given implicitly by

$$\min_{\mathbf{x}, \mathbf{u}} \quad \mathbb{E} \left[T(\mathbf{x}_N) + \sum_{k=0}^{N-1} L(\mathbf{x}_k, \mathbf{u}_k) \right]$$

$$\text{s.t.} \quad \mathbf{x}_{k+1} \sim \rho_{\theta}(\mathbf{x}_k, \mathbf{u}_k)$$

$$\mathbf{x}_0 = \mathbf{s}$$

use $\pi_{\theta}(s) = \mathbf{u}_0^*$ on the system

Relationship π_{θ} to π^* ??

- They match if ρ_{θ} is exact and deterministic
- But it usually cannot be...



Back to MPC - A “Classical” View on Decision Making

Policy $\pi_\theta(s)$ given implicitly by

$$\min_{\mathbf{x}, \mathbf{u}} \quad \mathbb{E} \left[T(\mathbf{x}_N) + \sum_{k=0}^{N-1} L(\mathbf{x}_k, \mathbf{u}_k) \right]$$

$$\text{s.t.} \quad \mathbf{x}_{k+1} \sim \rho_\theta(\mathbf{x}_k, \mathbf{u}_k)$$

$$\mathbf{x}_0 = \mathbf{s}$$

use $\pi_\theta(s) = \mathbf{u}_0^*$ on the system

Relationship π_θ to π^* ??

- They match if ρ_θ is exact and deterministic
- But it usually cannot be...

“Standard methods” for choosing θ in general do not yield the best π_θ .
RL over MPC can fix that, L becomes part of the model



Back to MPC - A “Classical” View on Decision Making

Policy $\pi_\theta(s)$ given implicitly by

$$\min_{\mathbf{x}, \mathbf{u}} \quad \mathbb{E} \left[T(\mathbf{x}_N) + \sum_{k=0}^{N-1} L(\mathbf{x}_k, \mathbf{u}_k) \right]$$

$$\text{s.t.} \quad \mathbf{x}_{k+1} \sim \rho_\theta(\mathbf{x}_k, \mathbf{u}_k)$$

$$\mathbf{x}_0 = \mathbf{s}$$

use $\pi_\theta(s) = \mathbf{u}_0^*$ on the system

Relationship π_θ to π^* ??

- They match if ρ_θ is exact and deterministic
- But it usually cannot be...

“Standard methods” for choosing θ in general do not yield the best π_θ .

RL over MPC can fix that, L becomes part of the model



This is not the only way of taking decisions from models

Decision policy from AI models (Sim2Real)

AI model

$$\hat{s}_+ \sim \rho_{\theta}(\cdot | s, a)$$

fitting data ρ

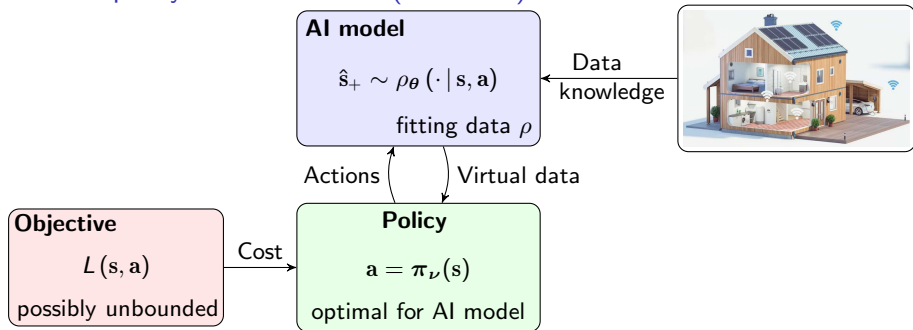
Data
knowledge



Classical Process:

- Fit AI model ρ_{θ} to real data

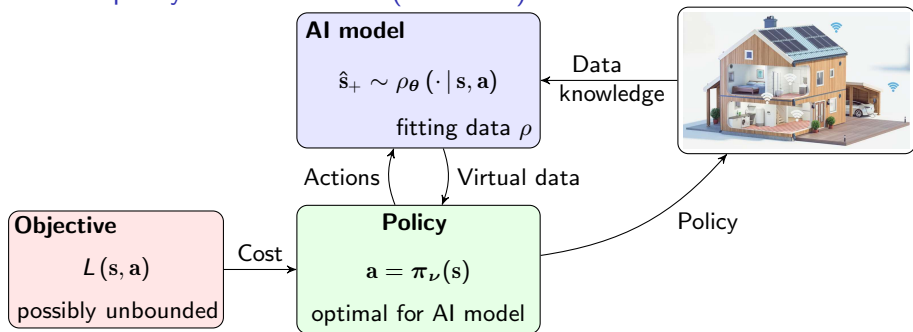
Decision policy from AI models (Sim2Real)



Classical Process:

- Fit AI model ρ_θ to real data
- Develop optimal policy for L and ρ_θ from AI model, e.g. using In-Sim RL
 - ▶ Define parametrized policy π_ν
 - ▶ Optimize policy parameters ν for performance w.r.t. AI model: $\rho_\theta, L \rightarrow \nu^*$

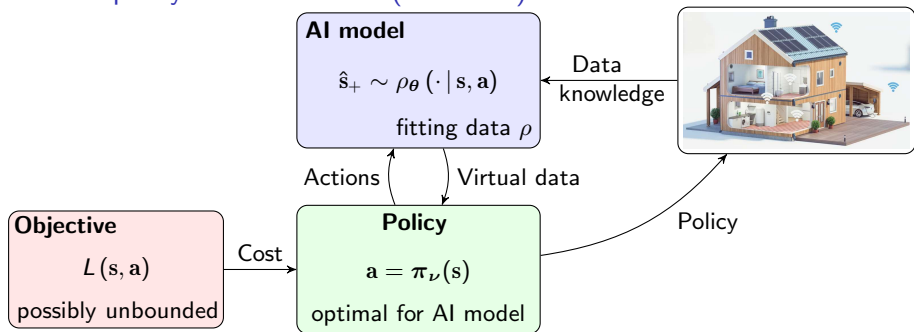
Decision policy from AI models (Sim2Real)



Classical Process:

- Fit AI model ρ_θ to real data
- Develop optimal policy for L and ρ_θ from AI model, e.g. using In-Sim RL
 - ▶ Define parametrized policy π_ν
 - ▶ Optimize policy parameters ν for performance w.r.t. AI model: $\rho_\theta, L \rightarrow \nu^*$
- Transfer policy π_{ν^*} into the real world

Decision policy from AI models (Sim2Real)

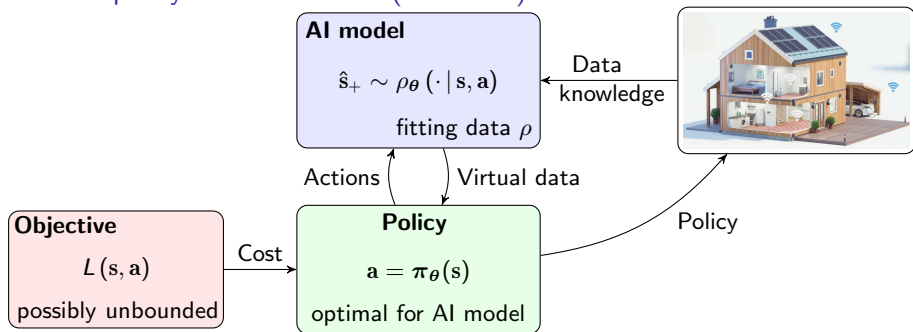


Classical Process:

- Fit AI model ρ_θ to real data
- Develop optimal policy for L and ρ_θ from AI model, e.g. using In-Sim RL
 - ▶ Define parametrized policy π_ν
 - ▶ Optimize policy parameters ν for performance w.r.t. AI model: $\rho_\theta, L \rightarrow \nu^*$
- Transfer policy π_{ν^*} into the real world

Note: policy parameters ν^* optimal for AI model become function of θ
Let's then label $\pi_\theta = \pi_{\nu^*}$ for simplicity

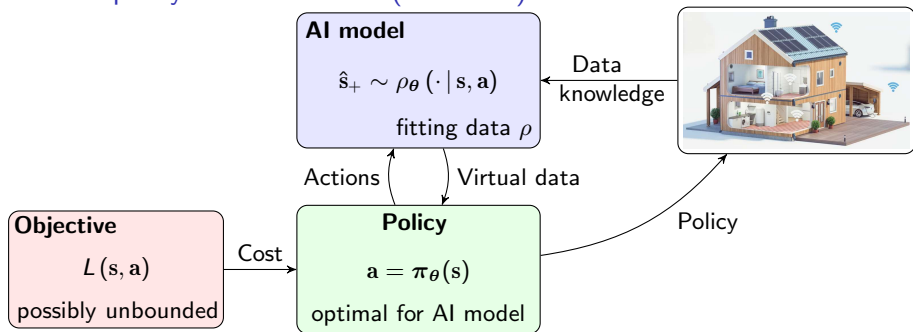
Decision policy from AI models (Sim2Real)



Remarks

- Relationship π_θ to π^* ??
 - ▶ They match if $\rho_\theta = \rho$, i.e. if model is exact
 - ▶ But ρ_θ is (almost always) an approximation

Decision policy from AI models (Sim2Real)



Remarks

- Relationship π_θ to π^* ??
 - ▶ They match if $\rho_\theta = \rho$, i.e. if model is exact
 - ▶ But ρ_θ is (almost always) an approximation

Are “standard methods” for choosing θ resulting in a good policy π_θ ??
Empirically the answer is “no”. Then how to choose θ to get a good policy?

MPC or Sim2Real - It's the same question...

MPC or Sim2Real - It's the same question...

The theory is about model-based decisions, and equivalences between MDPs

MPC or Sim2Real - It's the same question...

The theory is about model-based decisions, and equivalences between MDPs

World MDP

World MDP definition

- States s and actions a
- Cost $L(s, a)$
- Transition $s_+ \sim \rho[\cdot | s, a]$

MPC or Sim2Real - It's the same question...

The theory is about model-based decisions, and equivalences between MDPs

World MDP

World MDP definition

- States s and actions a
- Cost $L(s, a)$
- Transition $s_+ \sim \rho[\cdot | s, a]$

Optimal value functions

$$V^*(s) = \min_a Q^*(s, a)$$

$$Q^*(s, a) = L(s, a) + \gamma \mathbb{E}_{s_+ \sim \rho} [V^*(s_+) | s, a]$$

MPC or Sim2Real - It's the same question...

The theory is about model-based decisions, and equivalences between MDPs

World MDP

World MDP definition

- States s and actions a
- Cost $L(s, a)$
- Transition $s_+ \sim \rho[\cdot | s, a]$

Optimal value functions

$$V^*(s) = \min_a Q^*(s, a)$$

$$Q^*(s, a) = L(s, a) + \gamma \mathbb{E}_{s_+ \sim \rho} [V^*(s_+) | s, a]$$

Optimal policy

$$\pi^*(s) = \arg \min_a Q^*(s, a)$$

MPC or Sim2Real - It's the same question...

The theory is about model-based decisions, and equivalences between MDPs

World MDP

World MDP definition

- States s and actions a
- Cost $L(s, a)$
- Transition $s_+ \sim \rho[\cdot | s, a]$

Optimal value functions

$$V^*(s) = \min_a Q^*(s, a)$$

$$Q^*(s, a) = L(s, a) + \gamma \mathbb{E}_{s_+ \sim \rho} [V^*(s_+) | s, a]$$

Optimal policy

$$\pi^*(s) = \arg \min_a Q^*(s, a)$$

Model MDP (simulation environment)

Model MDP definition

- States s and actions a
- Cost $\hat{L}(s, a)$
- Transition $s_+ \sim \hat{\rho}[\cdot | s, a]$

Optimal value functions

$$\hat{V}^*(s) = \min_a \hat{Q}^*(s, a)$$

$$\hat{Q}^*(s, a) = \hat{L}(s, a) + \gamma \mathbb{E}_{s_+ \sim \hat{\rho}} [V^*(\hat{s}_+) | s, a]$$

Optimal policy

$$\hat{\pi}^*(s) = \arg \min_a \hat{Q}^*(s, a)$$

MPC or Sim2Real - It's the same question...

The theory is about model-based decisions, and equivalences between MDPs

World MDP

World MDP definition

- States s and actions a
- Cost $L(s, a)$
- Transition $s_+ \sim \rho[\cdot | s, a]$

Optimal value functions

$$V^*(s) = \min_a Q^*(s, a)$$

$$Q^*(s, a) = L(s, a) + \gamma \mathbb{E}_{s_+ \sim \rho} [V^*(s_+) | s, a]$$

Optimal policy

$$\pi^*(s) = \arg \min_a Q^*(s, a)$$

Model MDP (simulation environment)

Model MDP definition

- States s and actions a
- Cost $\hat{L}(s, a)$
- Transition $s_+ \sim \hat{\rho}[\cdot | s, a]$

Optimal value functions

$$\hat{V}^*(s) = \min_a \hat{Q}^*(s, a)$$

$$\hat{Q}^*(s, a) = \hat{L}(s, a) + \gamma \mathbb{E}_{s_+ \sim \hat{\rho}} [V^*(\hat{s}_+) | s, a]$$

Optimal policy

$$\hat{\pi}^*(s) = \arg \min_a \hat{Q}^*(s, a)$$

Theory says that - under some technical conditions - there is a $\hat{L}$ such that $\hat{Q}^* = Q^*$

MPC or Sim2Real - It's the same question...

The theory is about model-based decisions, and equivalences between MDPs

World MDP

World MDP definition

- States s and actions a
- Cost $L(s, a)$
- Transition $s_+ \sim \rho[\cdot | s, a]$

Optimal value functions

$$V^*(s) = \min_a Q^*(s, a)$$

$$Q^*(s, a) = L(s, a) + \gamma \mathbb{E}_{s_+ \sim \rho} [V^*(s_+) | s, a]$$

Optimal policy

$$\pi^*(s) = \arg \min_a Q^*(s, a)$$

Model MDP (simulation environment)

Model MDP definition

- States s and actions a
- Cost $\hat{L}(s, a)$
- Transition $s_+ \sim \hat{\rho}[\cdot | s, a]$

Optimal value functions

$$\hat{V}^*(s) = \min_a \hat{Q}^*(s, a)$$

$$\hat{Q}^*(s, a) = \hat{L}(s, a) + \gamma \mathbb{E}_{s_+ \sim \hat{\rho}} [V^*(\hat{s}_+) | s, a]$$

Optimal policy

$$\hat{\pi}^*(s) = \arg \min_a \hat{Q}^*(s, a)$$

Theory says that - under some technical conditions - there is a $\hat{L}$ such that $\hat{Q}^* = Q^*$

Proof: telescopic sum, some non-trivial assumptions to prevent $\infty - \infty$ cancellations

More on MDP Equivalence

World MDP

- States s and actions a
- Cost $L(s, a)$
- Transition $s_+ \sim \rho[\cdot | s, a]$

Model MDP

- States s and actions a
- Cost $\hat{L}(s, a)$
- Transition $s_+ \sim \hat{\rho}[\cdot | s, a]$

More on MDP Equivalence

World MDP

- States s and actions a
- Cost $L(s, a)$
- Transition $s_+ \sim \rho[\cdot | s, a]$

Model MDP

- States s and actions a
- Cost $\hat{L}(s, a)$
- Transition $s_+ \sim \hat{\rho}[\cdot | s, a]$

Theory says that

- Under some technical conditions there is a $\hat{L}$ such that $\hat{Q}^* = Q^*$

More on MDP Equivalence

World MDP

- States s and actions a
- Cost $L(s, a)$
- Transition $s_+ \sim \rho[\cdot | s, a]$

Model MDP

- States s and actions a
- Cost $\hat{L}(s, a)$
- Transition $s_+ \sim \hat{\rho}[\cdot | s, a]$

Theory says that

- Under some technical conditions there is a $\hat{L}$ such that $\hat{Q}^* = Q^*$
- For $\hat{L} = L$ there is a (non-unique) “optimal” model $\hat{\rho}$ such that $\hat{Q}^* = Q^*$

More on MDP Equivalence

World MDP

- States s and actions a
- Cost $L(s, a)$
- Transition $s_+ \sim \rho[\cdot | s, a]$

Model MDP

- States s and actions a
- Cost $\hat{L}(s, a)$
- Transition $s_+ \sim \hat{\rho}[\cdot | s, a]$

Theory says that

- Under some technical conditions there is a $\hat{L}$ such that $\hat{Q}^* = Q^*$
- For $\hat{L} = L$ there is a (non-unique) “optimal” model $\hat{\rho}$ such that $\hat{Q}^* = Q^*$
- Conditions for model $\hat{\rho}$ “optimality” $\neq$ min of classical loss functions (except. LQR)

More on MDP Equivalence

World MDP

- States s and actions a
- Cost $L(s, a)$
- Transition $s_+ \sim \rho[\cdot | s, a]$

Model MDP

- States s and actions a
- Cost $\hat{L}(s, a)$
- Transition $s_+ \sim \hat{\rho}[\cdot | s, a]$

Theory says that

- Under some technical conditions there is a $\hat{L}$ such that $\hat{Q}^* = Q^*$
- For $\hat{L} = L$ there is a (non-unique) “optimal” model $\hat{\rho}$ such that $\hat{Q}^* = Q^*$
- Conditions for model $\hat{\rho}$ “optimality” $\neq$ min of classical loss functions (except. LQR)
- World MDP and Model MDP do not need to use the same discount γ

More on MDP Equivalence

World MDP

- States s and actions a
- Cost $L(s, a)$
- Transition $s_+ \sim \rho[\cdot | s, a]$

Model MDP

- States s and actions a
- Cost $\hat{L}(s, a)$
- Transition $s_+ \sim \hat{\rho}[\cdot | s, a]$

Theory says that

- Under some technical conditions there is a $\hat{L}$ such that $\hat{Q}^* = Q^*$
- For $\hat{L} = L$ there is a (non-unique) “optimal” model $\hat{\rho}$ such that $\hat{Q}^* = Q^*$
- Conditions for model $\hat{\rho}$ “optimality” $\neq$ min of classical loss functions (except. LQR)
- World MDP and Model MDP do not need to use the same discount γ

In-Sim policy training uses a Model MDP

- Practitioners mostly work on model $\hat{\rho}$, but “unsure” on how to tune it
- In the ML communities, people talk about “value alignment”, we are trying to make this “MDP equivalence” understood

More on MDP Equivalence

World MDP

- States s and actions a
- Cost $L(s, a)$
- Transition $s_+ \sim \rho[\cdot | s, a]$

Model MDP

- States s and actions a
- Cost $\hat{L}(s, a)$
- Transition $s_+ \sim \hat{\rho}[\cdot | s, a]$

Theory says that

- Under some technical conditions there is a $\hat{L}$ such that $\hat{Q}^* = Q^*$
- For $\hat{L} = L$ there is a (non-unique) “optimal” model $\hat{\rho}$ such that $\hat{Q}^* = Q^*$
- Conditions for model $\hat{\rho}$ “optimality” $\neq$ min of classical loss functions (except. LQR)
- World MDP and Model MDP do not need to use the same discount γ

Remarks

- Non-unique optimal model $\hat{\rho}$ leaves room for aligning it to the real world (classical fitting)

More on MDP Equivalence

World MDP

- States s and actions a
- Cost $L(s, a)$
- Transition $s_+ \sim \rho[\cdot | s, a]$

Model MDP

- States s and actions a
- Cost $\hat{L}(s, a)$
- Transition $s_+ \sim \hat{\rho}[\cdot | s, a]$

Theory says that

- Under some technical conditions there is a $\hat{L}$ such that $\hat{Q}^* = Q^*$
- For $\hat{L} = L$ there is a (non-unique) “optimal” model $\hat{\rho}$ such that $\hat{Q}^* = Q^*$
- Conditions for model $\hat{\rho}$ “optimality” $\neq$ min of classical loss functions (except. LQR)
- World MDP and Model MDP do not need to use the same discount γ

Remarks

- Non-unique optimal model $\hat{\rho}$ leaves room for aligning it to the real world (classical fitting)
- If V^* is continuous and support of ρ is bounded(?) and path-connected, then we can have support $\hat{\rho} \subset$ support of ρ

More on MDP Equivalence

World MDP

- States s and actions a
- Cost $L(s, a)$
- Transition $s_+ \sim \rho[\cdot | s, a]$

Model MDP

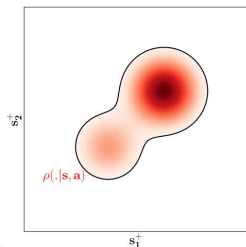
- States s and actions a
- Cost $\hat{L}(s, a)$
- Transition $s_+ \sim \hat{\rho}[\cdot | s, a]$

Theory says that

- Under some technical conditions there is a $\hat{L}$ such that $\hat{Q}^* = Q^*$
- For $\hat{L} = L$ there is a (non-unique) “optimal” model $\hat{\rho}$ such that $\hat{Q}^* = Q^*$
- Conditions for model $\hat{\rho}$ “optimality” $\neq$ min of classical loss functions (except. LQR)
- World MDP and Model MDP do not need to use the same discount γ

Remarks

- Non-unique optimal model $\hat{\rho}$ leaves room for aligning it to the real world (classical fitting)
- If V^* is continuous and support of ρ is bounded(?) and path-connected, then we can have support $\hat{\rho} \subset$ support of ρ



More on MDP Equivalence

World MDP

- States s and actions a
- Cost $L(s, a)$
- Transition $s_+ \sim \rho[\cdot | s, a]$

Model MDP

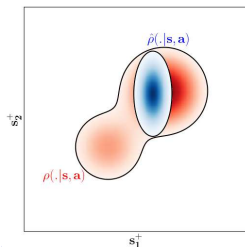
- States s and actions a
- Cost $\hat{L}(s, a)$
- Transition $s_+ \sim \hat{\rho}[\cdot | s, a]$

Theory says that

- Under some technical conditions there is a $\hat{L}$ such that $\hat{Q}^* = Q^*$
- For $\hat{L} = L$ there is a (non-unique) “optimal” model $\hat{\rho}$ such that $\hat{Q}^* = Q^*$
- Conditions for model $\hat{\rho}$ “optimality” $\neq$ min of classical loss functions (except. LQR)
- World MDP and Model MDP do not need to use the same discount γ

Remarks

- Non-unique optimal model $\hat{\rho}$ leaves room for aligning it to the real world (classical fitting)
- If V^* is continuous and support of ρ is bounded(?) and path-connected, then we can have support $\hat{\rho} \subset$ support of ρ



More on MDP Equivalence

World MDP

- States s and actions a
- Cost $L(s, a)$
- Transition $s_+ \sim \rho[\cdot | s, a]$

Model MDP

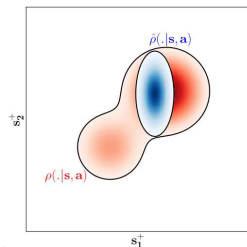
- States s and actions a
- Cost $\hat{L}(s, a)$
- Transition $s_+ \sim \hat{\rho}[\cdot | s, a]$

Theory says that

- Under some technical conditions there is a $\hat{L}$ such that $\hat{Q}^* = Q^*$
- For $\hat{L} = L$ there is a (non-unique) “optimal” model $\hat{\rho}$ such that $\hat{Q}^* = Q^*$
- Conditions for model $\hat{\rho}$ “optimality” $\neq$ min of classical loss functions (except. LQR)
- World MDP and Model MDP do not need to use the same discount γ

Remarks

- Non-unique optimal model $\hat{\rho}$ leaves room for aligning it to the real world (classical fitting)
- If V^* is continuous and support of ρ is bounded(?) and path-connected, then we can have support $\hat{\rho} \subset$ support of ρ
- *More simply said:* we can build optimal models $\hat{\rho}$ that make “plausible” predictions about the real world.

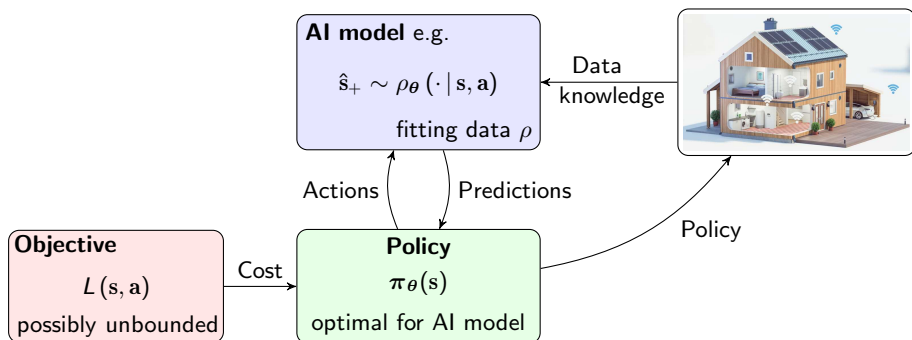


Outline

- 1 Decisions from data
- 2 AI models for Decision



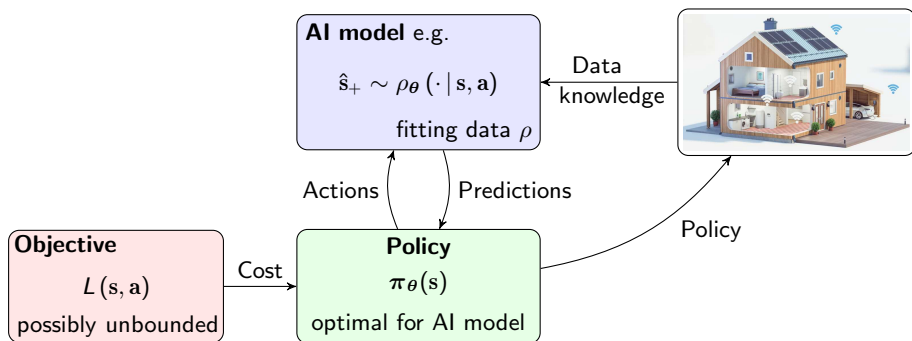
AI models for Decisions?



Learning model ρ_{θ} should aim at identifying

$$\theta^* = \arg \min_{\theta} J(\pi_{\theta}) \quad (1)$$

AI models for Decisions?

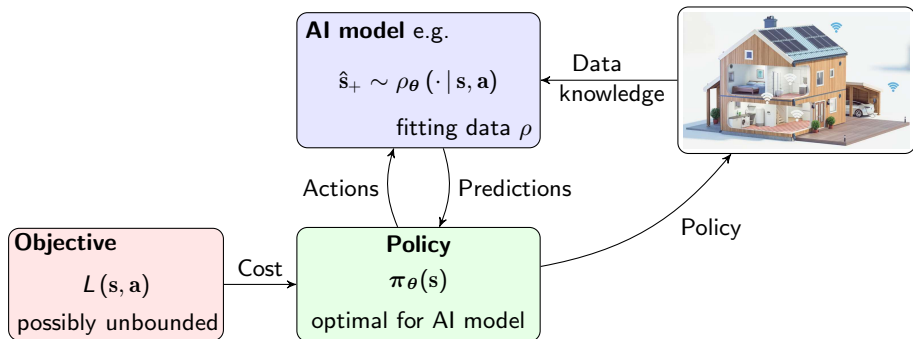


Learning model ρ_{θ} should aim at identifying

$$\theta^* = \arg \min_{\theta} J(\pi_{\theta}) \quad (1)$$

- Closed-loop performance J is intricate ($\theta \rightarrow \text{simulations} \rightarrow \text{policy} \rightarrow \text{real system}$)

AI models for Decisions?

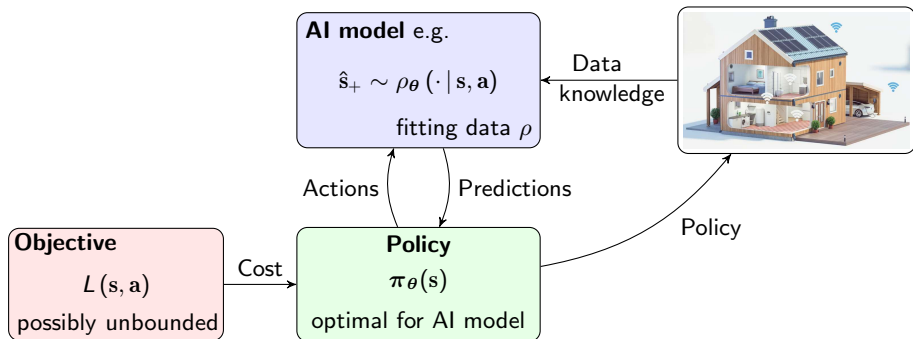


Learning model ρ_θ should aim at identifying

$$\theta^* = \arg \min_{\theta} J(\pi_\theta) \quad (1)$$

- Closed-loop performance J is intricate ($\theta \rightarrow \text{simulations} \rightarrow \text{policy} \rightarrow \text{real system}$)
- $\nabla_{\theta} J(\pi_\theta)$ is to be estimated from data, i.e. data hungry, noisy, possible biases

AI models for Decisions?

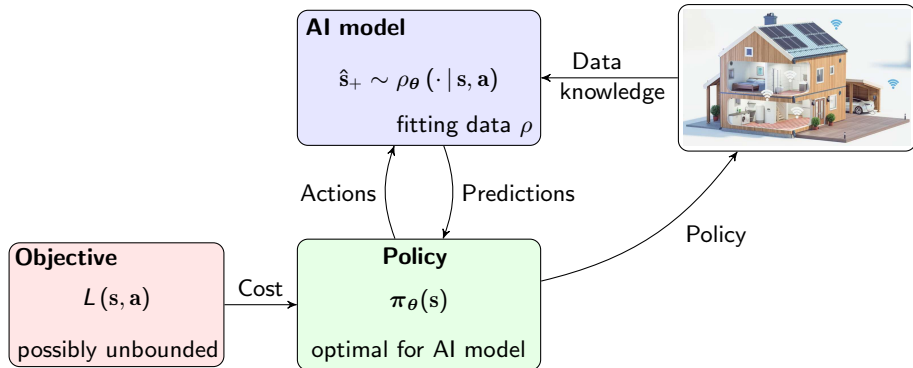


Learning model ρ_θ should aim at identifying

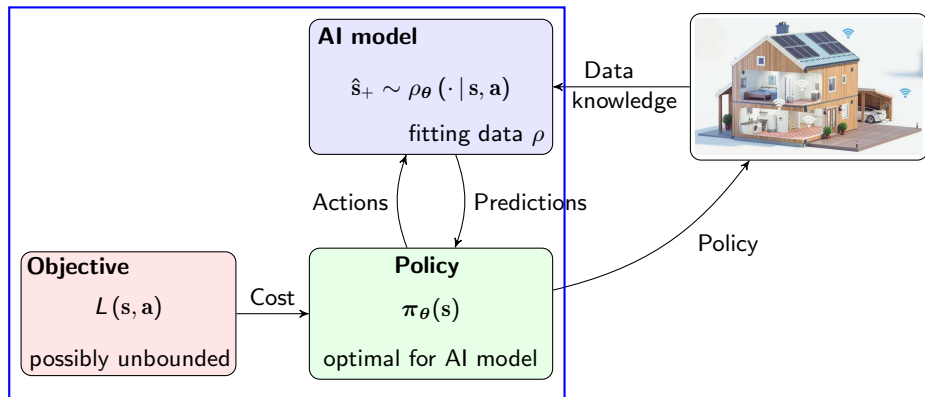
$$\theta^* = \arg \min_{\theta} J(\pi_\theta) \quad (1)$$

- Closed-loop performance J is intricate ($\theta \rightarrow \text{simulations} \rightarrow \text{policy} \rightarrow \text{real system}$)
- $\nabla_{\theta} J(\pi_\theta)$ is to be estimated from data, i.e. data hungry, noisy, possible biases
- (1) is in general different than model fitting, i.e. no loss function does (1)

Optimal AI models for Decision - A Paradigm Shift

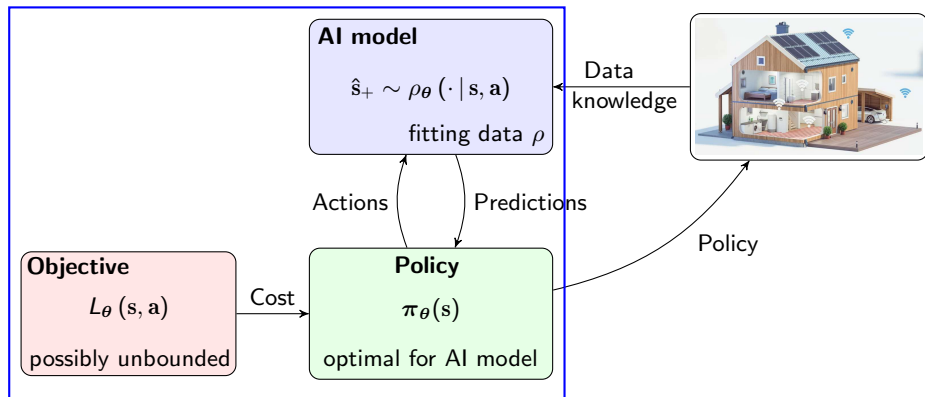


Optimal AI models for Decision - A Paradigm Shift



- The model is not ρ_{θ} , it is the entire "decision-making box"

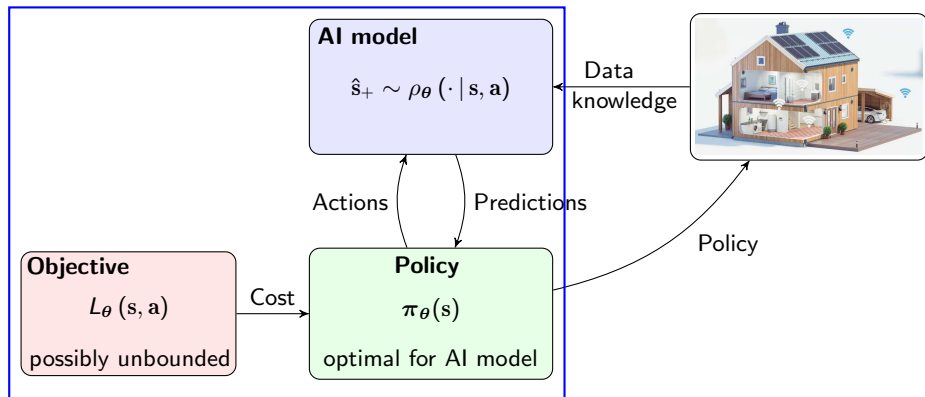
Optimal AI models for Decision - A Paradigm Shift



- The model is not ρ_{θ} , it is the entire “decision-making box”
- The model includes objective L_{θ} used to build the policy*

* policy performance on real system still assessed via L

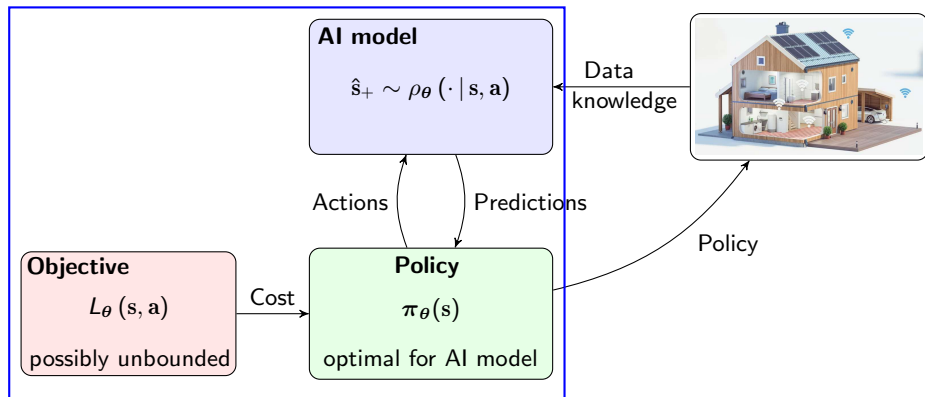
Optimal AI models for Decision - A Paradigm Shift



- The model is not ρ_{θ} , it is the entire “decision-making box”
- The model includes objective L_{θ} used to build the policy*
- Best model ρ_{θ} should *not necessarily* represent the data in a “classical sense”

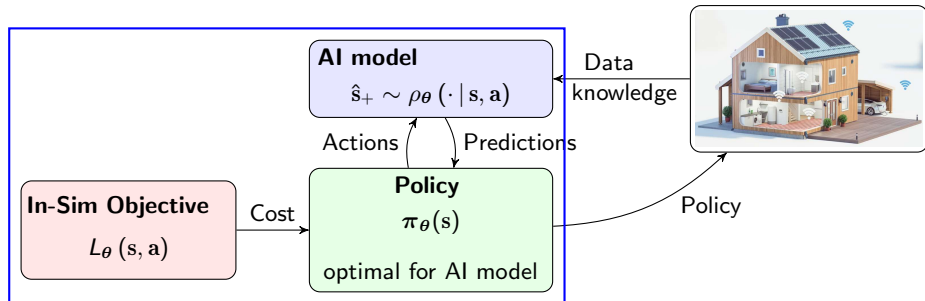
* policy performance on real system still assessed via L

Optimal AI models for Decision - A Paradigm Shift

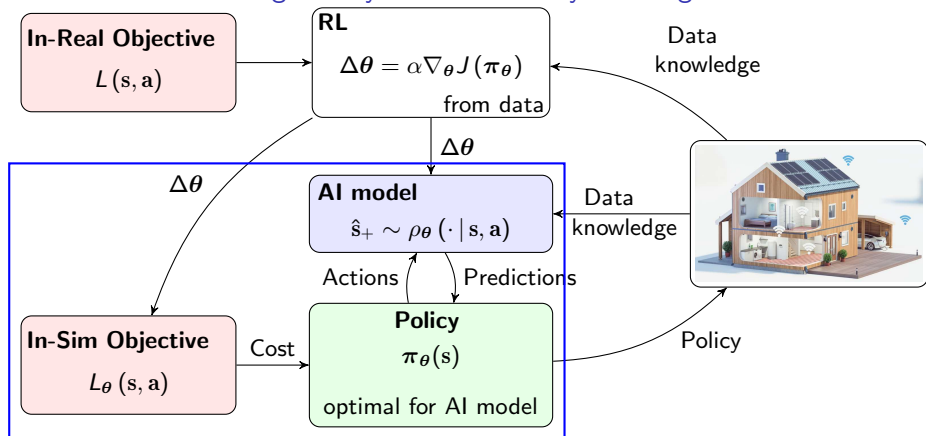


Theorem: under some (technical) assumptions, there is a L_θ such that $\pi_\theta = \pi_*$, even if ρ_θ does not represent real world ρ correctly

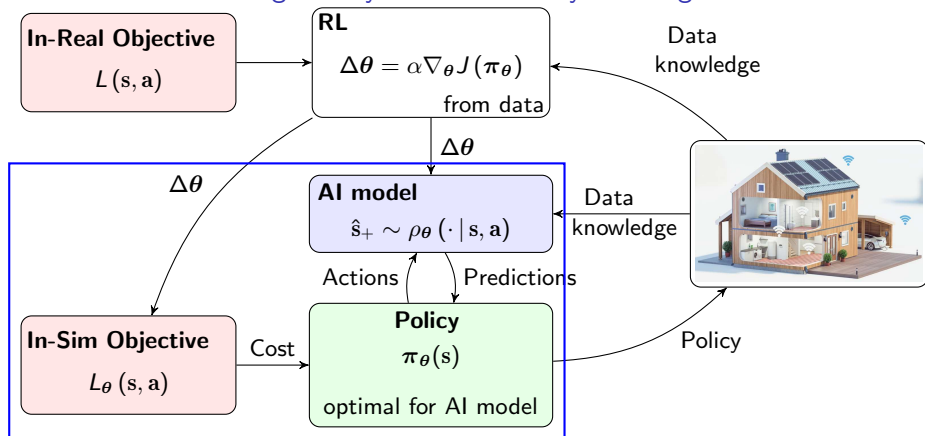
How to use this more generally? RL over Policy Training



How to use this more generally? RL over Policy Training



How to use this more generally? RL over Policy Training

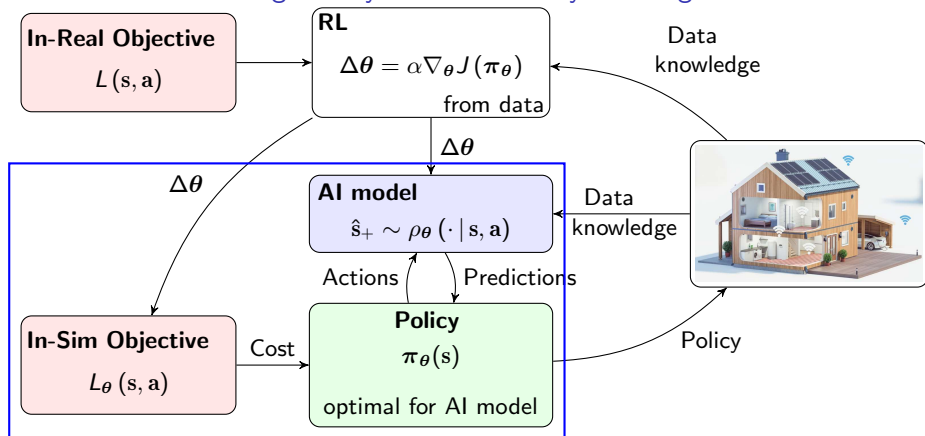


Policy gradient

$$\nabla_{\theta} J(\pi_{\theta}) = \mathbb{E} [\nabla_{\theta} \pi_{\theta} \nabla_a Q^{\pi_{\theta}}]$$

- $Q^{\pi_{\theta}}$ is the critic, well-established RL tool
- $\nabla_{\theta} \pi_{\theta}$ requires differentiating the closed-loop simulation...

How to use this more generally? RL over Policy Training



Difficulty: computing $\nabla_\theta \pi_\theta$ requires total differentiation through the simulation-based policy optimization process:

$\theta \rightarrow$ simulations $\rightarrow$ optimal policy for ρ_θ

Differentiating through policy optimization can be **computationally heavy**

Policy Differentiation

Performance for a stochastic policy π_{ν} , model MDP defined by L_{θ} , $\hat{\rho}_{\theta}$, γ

$$\hat{J}_{\theta}(\pi_{\nu}) = \mathbb{E}_{\hat{\rho}_{\theta}} \left[\sum_{k=0}^N \gamma^k L_{\theta}(\mathbf{s}_k, \mathbf{a}_k) \mid \mathbf{a}_k \sim \pi_{\nu}[\cdot | \mathbf{s}_k] \right]$$

Policy Differentiation

Performance for a stochastic policy $\pi_{\boldsymbol{\nu}}$, model MDP defined by $L_{\boldsymbol{\theta}}$, $\hat{\rho}_{\boldsymbol{\theta}}$, γ

$$\hat{J}_{\boldsymbol{\theta}}(\pi_{\boldsymbol{\nu}}) = \mathbb{E}_{\hat{\rho}_{\boldsymbol{\theta}}} \left[\sum_{k=0}^N \gamma^k L_{\boldsymbol{\theta}}(\mathbf{s}_k, \mathbf{a}_k) \mid \mathbf{a}_k \sim \pi_{\boldsymbol{\nu}}[\cdot | \mathbf{s}_k] \right]$$

Policy given by

$$\boldsymbol{\nu}^* = \arg \min_{\boldsymbol{\nu}} \hat{J}_{\boldsymbol{\theta}}(\pi_{\boldsymbol{\nu}}) \quad \text{or equivalently} \quad \frac{\partial}{\partial \boldsymbol{\nu}} \hat{J}_{\boldsymbol{\theta}}(\pi_{\boldsymbol{\nu}}) = 0$$

Policy Differentiation

Performance for a stochastic policy π_{ν} , model MDP defined by L_{θ} , $\hat{\rho}_{\theta}$, γ

$$\hat{J}_{\theta}(\pi_{\nu}) = \mathbb{E}_{\hat{\rho}_{\theta}} \left[\sum_{k=0}^N \gamma^k L_{\theta}(\mathbf{s}_k, \mathbf{a}_k) \mid \mathbf{a}_k \sim \pi_{\nu}[\cdot | \mathbf{s}_k] \right]$$

Policy given by

$$\nu^* = \arg \min_{\nu} \hat{J}_{\theta}(\pi_{\nu}) \quad \text{or equivalently} \quad \frac{\partial}{\partial \nu} \hat{J}_{\theta}(\pi_{\nu}) = 0$$

Policy sensitivity: if we change θ (cost L_{θ} and simulation $\hat{\rho}_{\theta}$) how does π_{ν} change?

$$\frac{\partial^2}{\partial \nu^2} \hat{J}_{\theta}(\pi_{\nu}) \underbrace{\frac{d\pi_{\nu}}{d\theta}}_{=\nabla_{\theta} \pi_{\theta}} + \frac{\partial^2}{\partial \nu \partial \theta} \hat{J}_{\theta}(\pi_{\nu}) = 0$$

Policy Differentiation

Performance for a stochastic policy π_{ν} , model MDP defined by L_{θ} , $\hat{\rho}_{\theta}$, γ

$$\hat{J}_{\theta}(\pi_{\nu}) = \mathbb{E}_{\hat{\rho}_{\theta}} \left[\sum_{k=0}^N \gamma^k L_{\theta}(\mathbf{s}_k, \mathbf{a}_k) \mid \mathbf{a}_k \sim \pi_{\nu}[\cdot | \mathbf{s}_k] \right]$$

Policy given by

$$\nu^* = \arg \min_{\nu} \hat{J}_{\theta}(\pi_{\nu}) \quad \text{or equivalently} \quad \frac{\partial}{\partial \nu} \hat{J}_{\theta}(\pi_{\nu}) = 0$$

Policy sensitivity: if we change θ (cost L_{θ} and simulation $\hat{\rho}_{\theta}$) how does π_{ν} change?

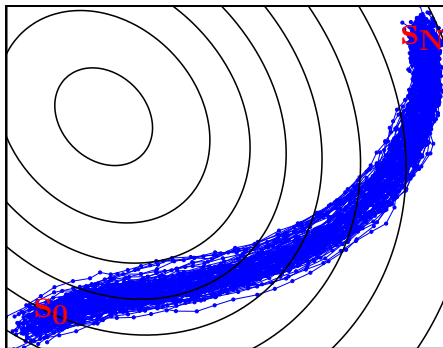
$$\frac{\partial^2}{\partial \nu^2} \hat{J}_{\theta}(\pi_{\nu}) \underbrace{\frac{d\pi_{\nu}}{d\theta}}_{=\nabla_{\theta} \pi_{\theta}} + \frac{\partial^2}{\partial \nu \partial \theta} \hat{J}_{\theta}(\pi_{\nu}) = 0$$

Simulation: (with $\mathbf{s} \equiv \mathbf{s}_0, \dots, \mathbf{s}_N$ and $\mathbf{a} \equiv \mathbf{a}_0, \dots, \mathbf{a}_N$)

$$\hat{J}_{\theta}(\pi_{\nu}) = \int \sum_{k=0}^N \gamma^k L_{\theta}(\mathbf{s}_k, \mathbf{a}_k) \varphi(\mathbf{s}, \mathbf{a}) d\mathbf{s} d\mathbf{a}$$

$$\varphi(\mathbf{s}, \mathbf{a}) = \rho_0[\mathbf{s}_0] \prod_{k=0}^{N-1} \hat{\rho}_{\theta}[\mathbf{s}_{k+1} | \mathbf{s}_k, \mathbf{a}_k] \prod_{k=0}^N \pi_{\nu}[\mathbf{a}_k | \mathbf{s}_k] \quad (\text{density of the simulation})$$

Policy Differentiation



- Blue trajectories are realization of φ :

$$\mathbf{s}^i, \mathbf{a}^i \sim \varphi(\cdot, \cdot)$$

- Sample-based evaluation (n samples)

$$\hat{J}_{\theta}(\pi_{\nu}) \approx \frac{1}{n} \sum_{k=0}^n \sum_{i=0}^N \gamma^k L_{\theta}(\mathbf{s}_k^i, \mathbf{a}_k^i)$$

Simulation: (with $\mathbf{s} \equiv \mathbf{s}_{0,\dots,N}$ and $\mathbf{a} \equiv \mathbf{a}_{0,\dots,N}$)

$$\hat{J}_{\theta}(\pi_{\nu}) = \int \sum_{k=0}^N \gamma^k L_{\theta}(\mathbf{s}_k, \mathbf{a}_k) \varphi(\mathbf{s}, \mathbf{a}) d\mathbf{s} d\mathbf{a}$$

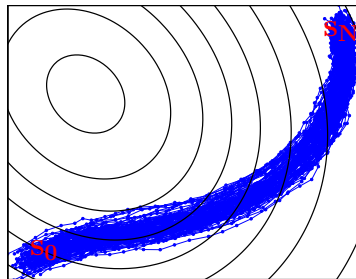
$$\varphi(\mathbf{s}, \mathbf{a}) = \rho_0[\mathbf{s}_0] \prod_{k=0}^{N-1} \hat{\rho}_{\theta}[\mathbf{s}_{k+1}|\mathbf{s}_k, \mathbf{a}_k] \prod_{k=0}^N \pi_{\nu}[\mathbf{a}_k|\mathbf{s}_k] \quad (\text{density of the simulation})$$

Direct Simulation Differentiation

Simulation

$$\hat{J}_{\theta}(\pi_{\nu}) = \int \sum_{k=0}^N \gamma^k L_{\theta}(\mathbf{s}_k, \mathbf{a}_k) \varphi(\mathbf{s}, \mathbf{a}) d\mathbf{s} d\mathbf{a}$$

$$\varphi(\mathbf{s}, \mathbf{a}) = \rho_0[\mathbf{s}_0] \prod_{k=0}^{N-1} \hat{p}_{\theta}[\mathbf{s}_{k+1} | \mathbf{s}_k, \mathbf{a}_k] \prod_{k=0}^N \pi_{\nu}[\mathbf{a}_k | \mathbf{s}_k]$$

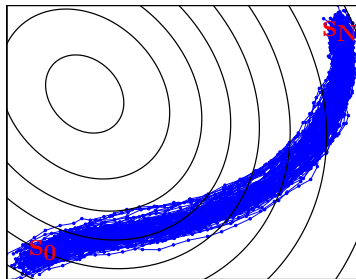


Direct Simulation Differentiation

Simulation

$$\hat{J}_{\theta}(\pi_{\nu}) = \int \sum_{k=0}^N \gamma^k L_{\theta}(\mathbf{s}_k, \mathbf{a}_k) \varphi(\mathbf{s}, \mathbf{a}) d\mathbf{s} d\mathbf{a}$$

$$\varphi(\mathbf{s}, \mathbf{a}) = \rho_0[\mathbf{s}_0] \prod_{k=0}^{N-1} \hat{p}_{\theta}[\mathbf{s}_{k+1} | \mathbf{s}_k, \mathbf{a}_k] \prod_{k=0}^N \pi_{\nu}[\mathbf{a}_k | \mathbf{s}_k]$$



Direct Differentiation:

$$\partial \hat{J}_{\theta}(\pi_{\nu}) = \mathbb{E}_{\hat{p}_{\theta}} \left[\sum_{k=0}^N \gamma^k \left(\partial L_{\theta}(\mathbf{s}_k, \mathbf{a}_k) + L_{\theta}(\mathbf{s}_k, \mathbf{a}_k) \partial \log \varphi(\mathbf{s}, \mathbf{a}) \right) \middle| \mathbf{a}_k \sim \pi_{\nu}[\cdot | \mathbf{s}_k] \right]$$

can be evaluated from data

Direct Simulation Differentiation

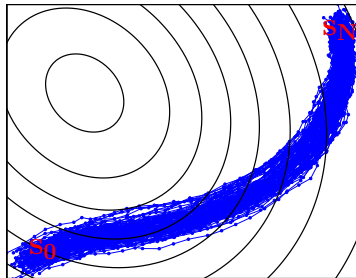
Simulation

$$\hat{J}_{\theta}(\pi_{\nu}) = \int \sum_{k=0}^N \gamma^k L_{\theta}(s_k, a_k) \varphi(s, a) ds da$$

$$\varphi(s, a) = \rho_0[s_0] \prod_{k=0}^{N-1} \hat{p}_{\theta}[s_{k+1}|s_k, a_k] \prod_{k=0}^N \pi_{\nu}[a_k|s_k]$$

Direct Differentiation:

$$\partial \hat{J}_{\theta}(\pi_{\nu}) \approx \frac{1}{n} \sum_{i=1}^n \sum_{k=0}^N \gamma^k \left(\partial L_{\theta}(s_k^i, a_k^i) + L_{\theta}(s_k^i, a_k^i) \partial \log \varphi(s_k^i, a_k^i) \right)$$

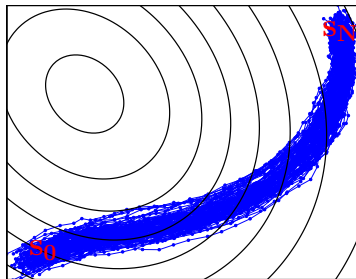


Direct Simulation Differentiation

Simulation

$$\hat{J}_{\theta}(\pi_{\nu}) = \int \sum_{k=0}^N \gamma^k L_{\theta}(s_k, a_k) \varphi(s, a) ds da$$

$$\varphi(s, a) = \rho_0[s_0] \prod_{k=0}^{N-1} \hat{p}_{\theta}[s_{k+1}|s_k, a_k] \prod_{k=0}^N \pi_{\nu}[a_k|s_k]$$



Direct Differentiation:

$$\partial \hat{J}_{\theta}(\pi_{\nu}) \approx \frac{1}{n} \sum_{i=1}^n \sum_{k=0}^N \gamma^k \left(\partial L_{\theta}(s_k^i, a_k^i) + L_{\theta}(s_k^i, a_k^i) \partial \log \varphi(s_k^i, a_k^i) \right)$$

Remarks

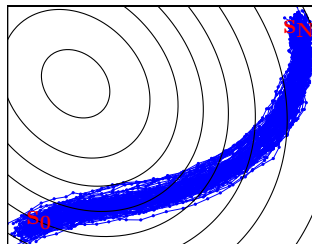
- ✓ Second-order sensitivities are similar
- ✓ Simple and computationally very efficient
- ✓ Can differentiate through discrete state and action sets
- ✗ Sample-based estimations are very noisy

Critic-Based Differentiation

In-Sim Actor-critic does

$$\nabla_{\boldsymbol{\nu}} \hat{J}_{\boldsymbol{\theta}}(\pi_{\boldsymbol{\nu}}) = \mathbb{E}_{\substack{\mathbf{s} \sim \rho_{\boldsymbol{\theta}} \\ \mathbf{a} \sim \pi_{\boldsymbol{\nu}}}} [\nabla_{\boldsymbol{\nu}} \log \pi_{\boldsymbol{\nu}}(\mathbf{a}|\mathbf{s}) A^{\pi_{\boldsymbol{\nu}}}(\mathbf{s}, \mathbf{a})] = 0$$

defines relationship $\pi_{\boldsymbol{\nu}^*}$ with $\boldsymbol{\nu}^*$ function of $\boldsymbol{\theta}$



Critic-Based Differentiation

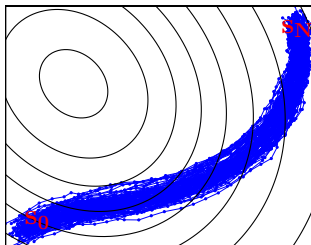
In-Sim Actor-critic does

$$\nabla_{\boldsymbol{\nu}} \hat{J}_{\boldsymbol{\theta}}(\pi_{\boldsymbol{\nu}}) = \mathbb{E}_{\substack{\mathbf{s} \sim \rho_{\boldsymbol{\theta}} \\ \mathbf{a} \sim \pi_{\boldsymbol{\nu}}}} [\nabla_{\boldsymbol{\nu}} \log \pi_{\boldsymbol{\nu}}(\mathbf{a}|\mathbf{s}) A^{\pi_{\boldsymbol{\nu}}}(\mathbf{s}, \mathbf{a})] = 0$$

defines relationship $\pi_{\boldsymbol{\nu}^*}$ with $\boldsymbol{\nu}^*$ function of $\boldsymbol{\theta}$

In-Real Actor-critic evaluates

$$\nabla_{\boldsymbol{\theta}} J(\pi_{\boldsymbol{\nu}^*}) = \mathbb{E}_{\text{Real World}} [\nabla_{\boldsymbol{\theta}} \log \pi_{\boldsymbol{\nu}^*}(\mathbf{a}|\mathbf{s}) A^{\boldsymbol{\nu}^*}(\mathbf{s}, \mathbf{a})]$$



Critic-Based Differentiation

In-Sim Actor-critic does

$$\nabla_{\boldsymbol{\nu}} \hat{J}_{\boldsymbol{\theta}}(\pi_{\boldsymbol{\nu}}) = \mathbb{E}_{\substack{\mathbf{s} \sim \rho_{\boldsymbol{\theta}} \\ \mathbf{a} \sim \pi_{\boldsymbol{\nu}}}} [\nabla_{\boldsymbol{\nu}} \log \pi_{\boldsymbol{\nu}}(\mathbf{a}|\mathbf{s}) A^{\pi_{\boldsymbol{\nu}}}(\mathbf{s}, \mathbf{a})] = 0$$

defines relationship $\pi_{\boldsymbol{\nu}^*}$ with $\boldsymbol{\nu}^*$ function of $\boldsymbol{\theta}$

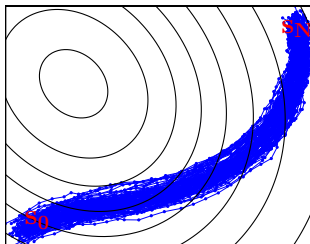
In-Real Actor-critic evaluates

$$\nabla_{\boldsymbol{\theta}} J(\pi_{\boldsymbol{\nu}^*}) = \mathbb{E}_{\text{Real World}} \left[\nabla_{\boldsymbol{\theta}} \log \pi_{\boldsymbol{\nu}^*}(\mathbf{a}|\mathbf{s}) A^{\boldsymbol{\nu}^*}(\mathbf{s}, \mathbf{a}) \right]$$

where $\nabla_{\boldsymbol{\theta}} \log \pi_{\boldsymbol{\nu}^*}(\mathbf{a}|\mathbf{s})$ requires $\frac{d\boldsymbol{\nu}^*}{d\boldsymbol{\theta}}$

Sensitivity

$$\frac{\partial^2}{\partial \boldsymbol{\nu}^2} \hat{J}_{\boldsymbol{\theta}}(\pi_{\boldsymbol{\nu}^*}) \frac{d\boldsymbol{\nu}^*}{d\boldsymbol{\theta}} + \frac{\partial^2}{\partial \boldsymbol{\nu} \partial \boldsymbol{\theta}} \hat{J}_{\boldsymbol{\theta}}(\pi_{\boldsymbol{\nu}^*}) = 0$$



Critic-Based Differentiation

In-Sim Actor-critic does

$$\nabla_{\nu} \hat{J}_{\theta}(\pi_{\nu}) = \mathbb{E}_{\substack{s \sim \rho_{\theta} \\ a \sim \pi_{\nu}}} [\nabla_{\nu} \log \pi_{\nu}(a|s) A^{\pi_{\nu}}(s, a)] = 0$$

defines relationship π_{ν^*} with ν^* function of θ

In-Real Actor-critic evaluates

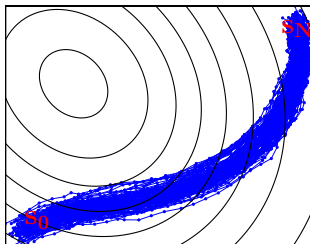
$$\nabla_{\theta} J(\pi_{\nu^*}) = \mathbb{E}_{\text{Real World}} [\nabla_{\theta} \log \pi_{\nu^*}(a|s) A^{\nu^*}(s, a)]$$

where $\nabla_{\theta} \log \pi_{\nu^*}(a|s)$ requires $\frac{d\nu^*}{d\theta}$

Sensitivity

$$\frac{\partial^2}{\partial \nu^2} \hat{J}_{\theta}(\pi_{\nu^*}) \frac{d\nu^*}{d\theta} + \frac{\partial^2}{\partial \nu \partial \theta} \hat{J}_{\theta}(\pi_{\nu^*}) = 0$$

Idea: compute $\frac{d\nu^*}{d\theta}$ from differentiating the in-sim actor-critic w.r.t. $\frac{\partial}{\partial \nu}$ and $\frac{\partial}{\partial \theta}$



Critic-Based Differentiation

In-Sim Actor-critic does

$$\nabla_{\nu} \hat{J}_{\theta}(\pi_{\nu}) = \mathbb{E}_{\substack{s \sim \rho_{\theta} \\ a \sim \pi_{\nu}}} [\nabla_{\nu} \log \pi_{\nu}(a|s) A^{\pi_{\nu}}(s, a)] = 0$$

defines relationship π_{ν^*} with ν^* function of θ

In-Real Actor-critic evaluates

$$\nabla_{\theta} J(\pi_{\nu^*}) = \mathbb{E}_{\text{Real World}} [\nabla_{\theta} \log \pi_{\nu^*}(a|s) A^{\nu^*}(s, a)]$$

where $\nabla_{\theta} \log \pi_{\nu^*}(a|s)$ requires $\frac{d\nu^*}{d\theta}$

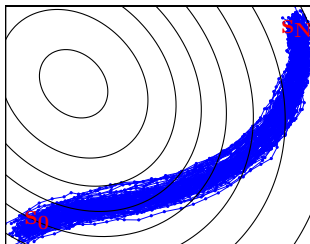
Sensitivity

$$\frac{\partial^2}{\partial \nu^2} \hat{J}_{\theta}(\pi_{\nu^*}) \frac{d\nu^*}{d\theta} + \frac{\partial^2}{\partial \nu \partial \theta} \hat{J}_{\theta}(\pi_{\nu^*}) = 0$$

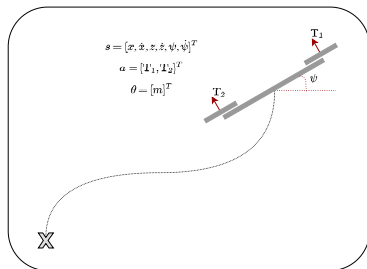
Idea: compute $\frac{d\nu^*}{d\theta}$ from differentiating the in-sim actor-critic w.r.t. $\frac{\partial}{\partial \nu}$ and $\frac{\partial}{\partial \theta}$

Proposed first algorithm to do that.

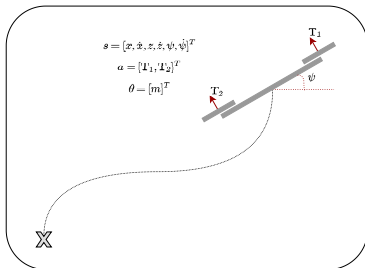
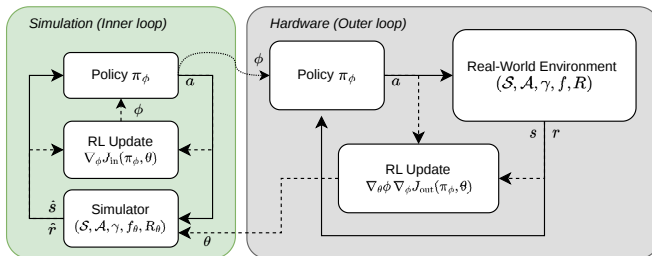
Computationally heavy, but we have “missed” some simplifications, testing further



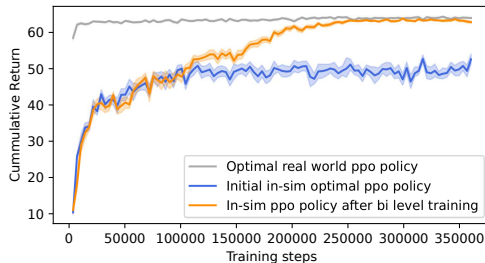
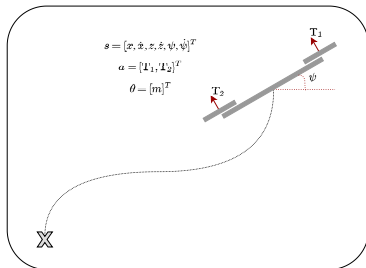
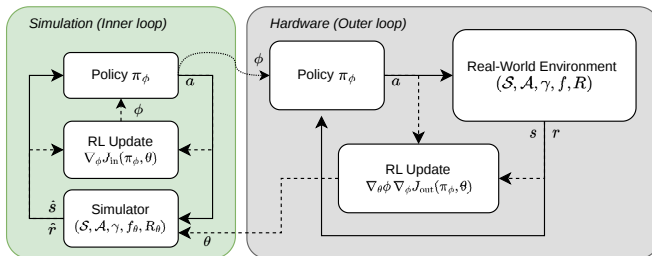
Example - In-Real RL over In-Sim RL



Example - In-Real RL over In-Sim RL



Example - In-Real RL over In-Sim RL





TORCG

-

◀ ◻ ▶ ◀ ◻ ▶ ◀ ≡ ▶ ◀ ≡ ▶ ≡