

So far we have seen...

- The building blocks
 - ▶ MDPs, DP, Bellman, Policy and Value Iteration, Reinforcement Learning, Actor-Critic
 - ▶ LQR, Dynamic Optimization, Nonlinear Programming, MPC, Sensitivities
- MPC & RL
 - ▶ Synthesis - Different flavors of ML / RL and MPC
 - ▶ Software tool leap-c

So far we have seen...

- The building blocks
 - ▶ MDPs, DP, Bellman, Policy and Value Iteration, Reinforcement Learning, Actor-Critic
 - ▶ LQR, Dynamic Optimization, Nonlinear Programming, MPC, Sensitivities
- MPC & RL
 - ▶ Synthesis - Different flavors of ML / RL and MPC
 - ▶ Software tool leap-c

What we will discuss: parametrized / hierarchical RL over MPC

- The theory that supports it and what does it tell us? (*Now*)

So far we have seen...

- The building blocks
 - ▶ MDPs, DP, Bellman, Policy and Value Iteration, Reinforcement Learning, Actor-Critic
 - ▶ LQR, Dynamic Optimization, Nonlinear Programming, MPC, Sensitivities
- MPC & RL
 - ▶ Synthesis - Different flavors of ML / RL and MPC
 - ▶ Software tool leap-c

What we will discuss: parametrized / hierarchical RL over MPC

- The theory that supports it and what does it tell us? (*Now*)
- Safe & stable RL over MPC (*In the afternoon*)

So far we have seen...

- The building blocks
 - ▶ MDPs, DP, Bellman, Policy and Value Iteration, Reinforcement Learning, Actor-Critic
 - ▶ LQR, Dynamic Optimization, Nonlinear Programming, MPC, Sensitivities
- MPC & RL
 - ▶ Synthesis - Different flavors of ML / RL and MPC
 - ▶ Software tool leap-c

What we will discuss: parametrized / hierarchical RL over MPC

- The theory that supports it and what does it tell us? (*Now*)
- Safe & stable RL over MPC (*In the afternoon*)
- RL over MPC with belief states – a future prospect (*In the afternoon*)

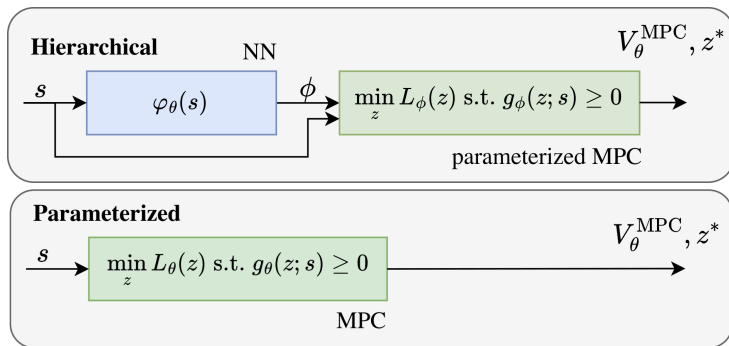
So far we have seen...

- The building blocks
 - ▶ MDPs, DP, Bellman, Policy and Value Iteration, Reinforcement Learning, Actor-Critic
 - ▶ LQR, Dynamic Optimization, Nonlinear Programming, MPC, Sensitivities
- MPC & RL
 - ▶ Synthesis - Different flavors of ML / RL and MPC
 - ▶ Software tool leap-c

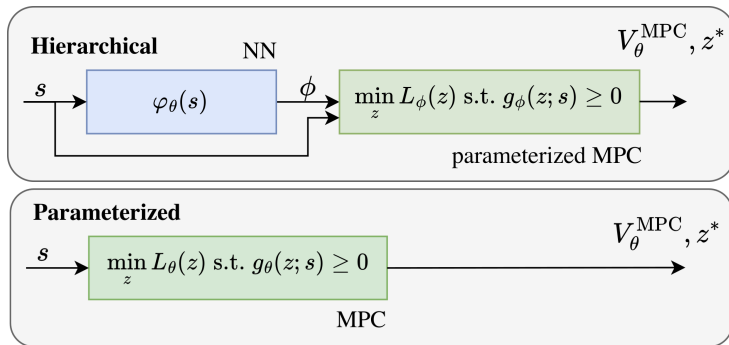
What we will discuss: parametrized / hierarchical RL over MPC

- The theory that supports it and what does it tell us? (*Now*)
- Safe & stable RL over MPC (*In the afternoon*)
- RL over MPC with belief states – a future prospect (*In the afternoon*)
- Beyond MPC – Model-based Decisions and AI for decisions (*Tomorrow*)

Focus of today's lectures (Leap-c structure)



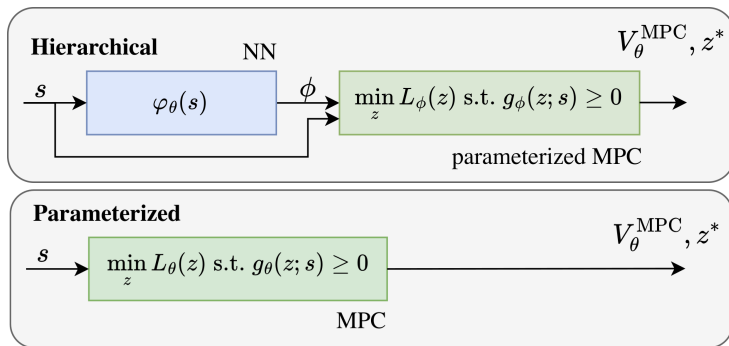
Focus of today's lectures (Leap-c structure)



Normal thinking when using MPC

- Fit MPC model to reality as good as possible (SYSID)
- MPC cost is what we want to minimize (energy, time, money, reference)
- MPC state constraints are what we need to respect (safety)

Focus of today's lectures (Leap-c structure)



Normal thinking when using MPC

- Fit MPC model to reality as good as possible (SYSID)
- MPC cost is what we want to minimize (energy, time, money, reference)
- MPC state constraints are what we need to respect (safety)

We are talking about changing everything!!

Reinforcement Learning Over MPC

Why Does It Work?

Sebastien Gros

Dept. of Cybernetic, NTNU
Faculty of Information Tech.

Freiburg PhD School

Outline

- 1 Let's rebuild some background – MPC and MDPs
- 2 MPC as a solution to MDPs
- 3 When is RL (most) beneficial for MPC?
- 4 A Deeper Look at the Theory

Model Predictive Control (MPC)

Optimize a plan over finite horizon, apply first move, repeat

MPC: at **current state** $\mathbf{s}$

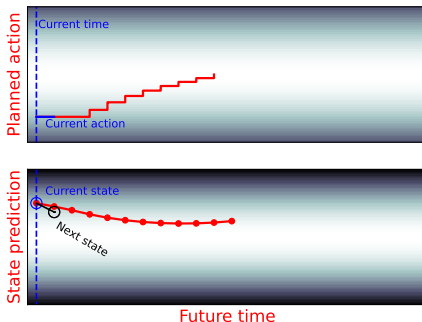
$$\min_{\mathbf{x}, \mathbf{u}} \quad T(\mathbf{x}_N) + \sum_{k=0}^{N-1} L(\mathbf{x}_k, \mathbf{u}_k)$$

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

$$\mathbf{h}(\mathbf{x}_k, \mathbf{u}_k) \leq 0$$

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

apply **action** $\mathbf{a} = \mathbf{u}_0^*$ to the system



Model Predictive Control (MPC)

Optimize a plan over finite horizon, apply first move, repeat

MPC: at **current state** $\mathbf{s}$

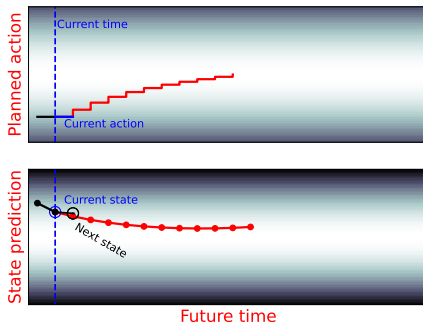
$$\min_{\mathbf{x}, \mathbf{u}} \quad T(\mathbf{x}_N) + \sum_{k=0}^{N-1} L(\mathbf{x}_k, \mathbf{u}_k)$$

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

$$\mathbf{h}(\mathbf{x}_k, \mathbf{u}_k) \leq 0$$

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

apply **action** $\mathbf{a} = \mathbf{u}_0^*$ to the system



Model Predictive Control (MPC)

Optimize a plan over finite horizon, apply first move, repeat

MPC: at **current state** $\mathbf{s}$

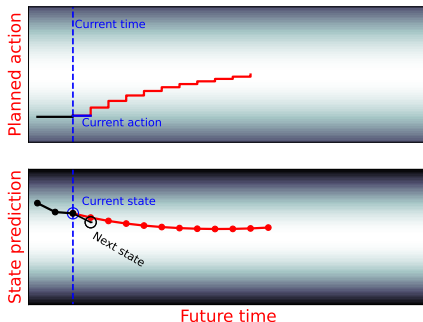
$$\min_{\mathbf{x}, \mathbf{u}} \quad T(\mathbf{x}_N) + \sum_{k=0}^{N-1} L(\mathbf{x}_k, \mathbf{u}_k)$$

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

$$\mathbf{h}(\mathbf{x}_k, \mathbf{u}_k) \leq 0$$

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

apply **action** $\mathbf{a} = \mathbf{u}_0^*$ to the system



Model Predictive Control (MPC)

Optimize a plan over finite horizon, apply first move, repeat

MPC: at **current state** $\mathbf{s}$

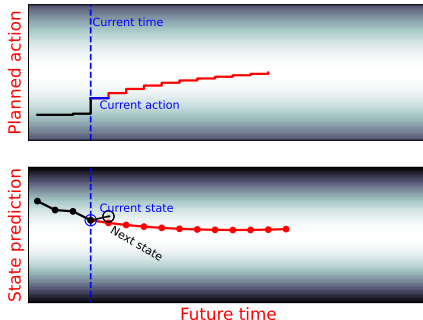
$$\min_{\mathbf{x}, \mathbf{u}} \quad T(\mathbf{x}_N) + \sum_{k=0}^{N-1} L(\mathbf{x}_k, \mathbf{u}_k)$$

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

$$\mathbf{h}(\mathbf{x}_k, \mathbf{u}_k) \leq 0$$

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

apply **action** $\mathbf{a} = \mathbf{u}_0^*$ to the system



Model Predictive Control (MPC)

Optimize a plan over finite horizon, apply first move, repeat

MPC: at **current state** $\mathbf{s}$

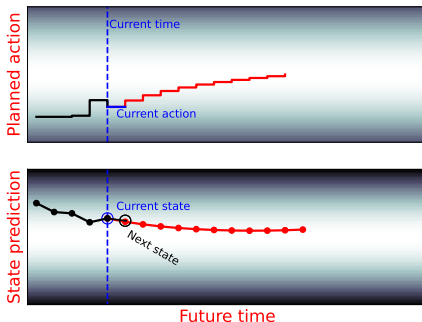
$$\min_{\mathbf{x}, \mathbf{u}} \quad T(\mathbf{x}_N) + \sum_{k=0}^{N-1} L(\mathbf{x}_k, \mathbf{u}_k)$$

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

$$\mathbf{h}(\mathbf{x}_k, \mathbf{u}_k) \leq 0$$

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

apply **action** $\mathbf{a} = \mathbf{u}_0^*$ to the system



Model Predictive Control (MPC)

Optimize a plan over finite horizon, apply first move, repeat

MPC: at current state $\mathbf{s}$

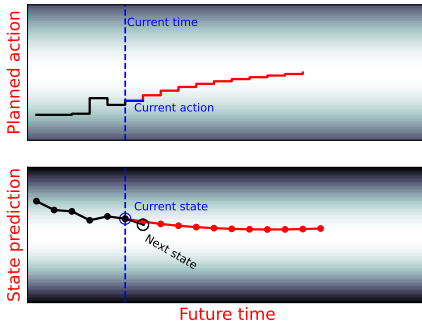
$$\min_{\mathbf{x}, \mathbf{u}} \quad T(\mathbf{x}_N) + \sum_{k=0}^{N-1} L(\mathbf{x}_k, \mathbf{u}_k)$$

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

$$\mathbf{h}(\mathbf{x}_k, \mathbf{u}_k) \leq 0$$

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

apply action $\mathbf{a} = \mathbf{u}_0^*$ to the system



Model Predictive Control (MPC)

Optimize a plan over finite horizon, apply first move, repeat

MPC: at **current state** $\mathbf{s}$

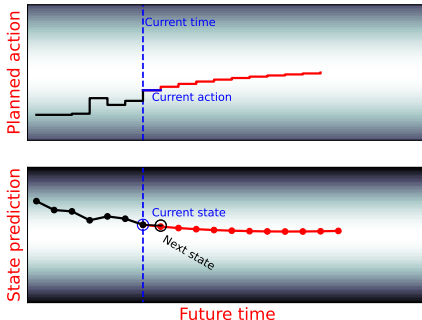
$$\min_{\mathbf{x}, \mathbf{u}} \quad T(\mathbf{x}_N) + \sum_{k=0}^{N-1} L(\mathbf{x}_k, \mathbf{u}_k)$$

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

$$\mathbf{h}(\mathbf{x}_k, \mathbf{u}_k) \leq 0$$

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

apply **action** $\mathbf{a} = \mathbf{u}_0^*$ to the system



Model Predictive Control (MPC)

Optimize a plan over finite horizon, apply first move, repeat

MPC: at current state $\mathbf{s}$

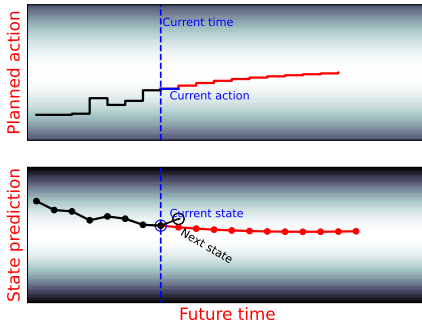
$$\min_{\mathbf{x}, \mathbf{u}} \quad T(\mathbf{x}_N) + \sum_{k=0}^{N-1} L(\mathbf{x}_k, \mathbf{u}_k)$$

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

$$\mathbf{h}(\mathbf{x}_k, \mathbf{u}_k) \leq 0$$

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

apply action $\mathbf{a} = \mathbf{u}_0^*$ to the system



Model Predictive Control (MPC)

Optimize a plan over finite horizon, apply first move, repeat

MPC: at **current state** $\mathbf{s}$

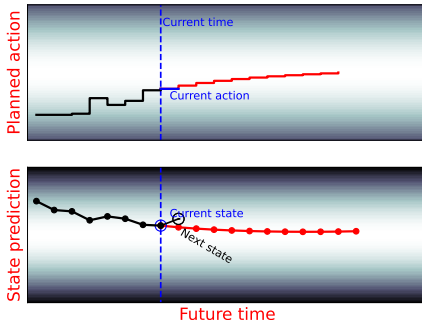
$$\min_{\mathbf{x}, \mathbf{u}} \quad T(\mathbf{x}_N) + \sum_{k=0}^{N-1} L(\mathbf{x}_k, \mathbf{u}_k)$$

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

$$\mathbf{h}(\mathbf{x}_k, \mathbf{u}_k) \leq 0$$

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

apply **action** $\mathbf{a} = \mathbf{u}_0^*$ to the system



MPC

- is **based on planning** the future
- **Policy** from repeated planning

$$\pi^{\text{MPC}}(\mathbf{s}) = \mathbf{u}_0^*$$

Model Predictive Control (MPC)

Optimize a plan over finite horizon, apply first move, repeat

MPC: at **current state** $\mathbf{s}$

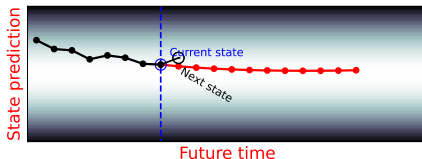
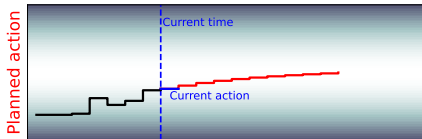
$$\min_{\mathbf{x}, \mathbf{u}} \quad T(\mathbf{x}_N) + \sum_{k=0}^{N-1} L(\mathbf{x}_k, \mathbf{u}_k)$$

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

$$\mathbf{h}(\mathbf{x}_k, \mathbf{u}_k) \leq 0$$

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

apply **action** $\mathbf{a} = \mathbf{u}_0^*$ to the system



MPC

- is **based on planning** the future
- **Policy** from repeated planning

$$\pi^{\text{MPC}}(\mathbf{s}) = \mathbf{u}_0^*$$

What an odd thing to do: we build and throw away plans all the time knowing that they are wrong all along, but kinda use them anyway...

Model Predictive Control (MPC)

Optimize a plan over finite horizon, apply first move, repeat

MPC: at **current state** s

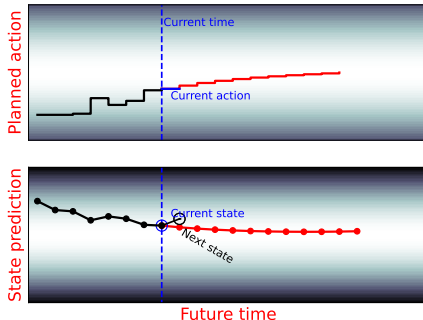
$$\min_{\mathbf{x}, \mathbf{u}} \quad T(\mathbf{x}_N) + \sum_{k=0}^{N-1} L(\mathbf{x}_k, \mathbf{u}_k)$$

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

$$\mathbf{h}(\mathbf{x}_k, \mathbf{u}_k) \leq 0$$

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

apply **action** $\mathbf{a} = \mathbf{u}_0^*$ to the system



MPC

- is **based on planning** the future
- **Policy** from repeated planning

$$\pi^{\text{MPC}}(s) = \mathbf{u}_0^*$$

MPC is a powerful tool to control constrained systems, increasingly used as a practical way of building optimal policies

Theoretical Framework to connect RL and MPC

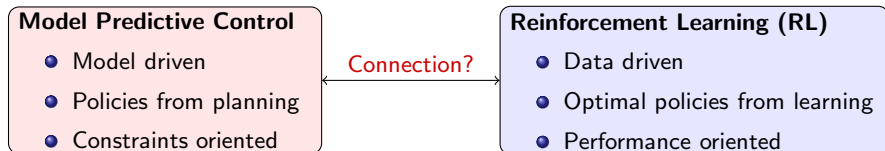
Model Predictive Control

- Model driven
- Policies from planning
- Constraints oriented

Reinforcement Learning (RL)

- Data driven
- Optimal policies from learning
- Performance oriented

Theoretical Framework to connect RL and MPC



Theoretical Framework to connect RL and MPC

Model Predictive Control

- Model driven
- Policies from planning
- Constraints oriented

Connection?

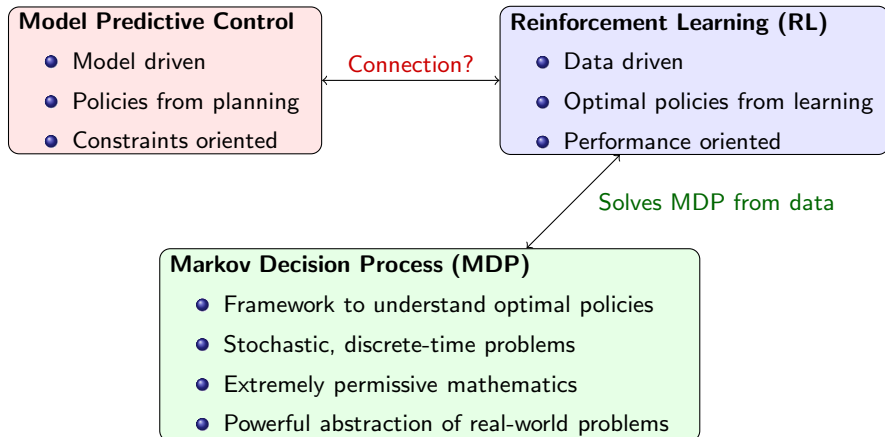
Reinforcement Learning (RL)

- Data driven
- Optimal policies from learning
- Performance oriented

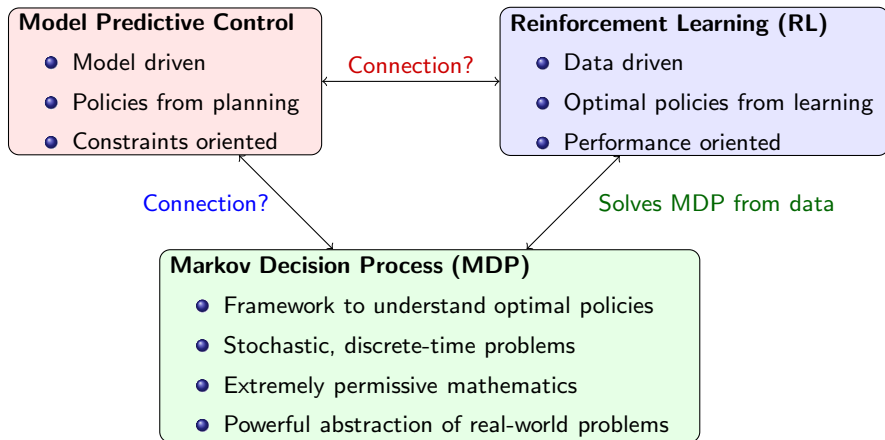
Markov Decision Process (MDP)

- Framework to understand optimal policies
- Stochastic, discrete-time problems
- Extremely permissive mathematics
- Powerful abstraction of real-world problems

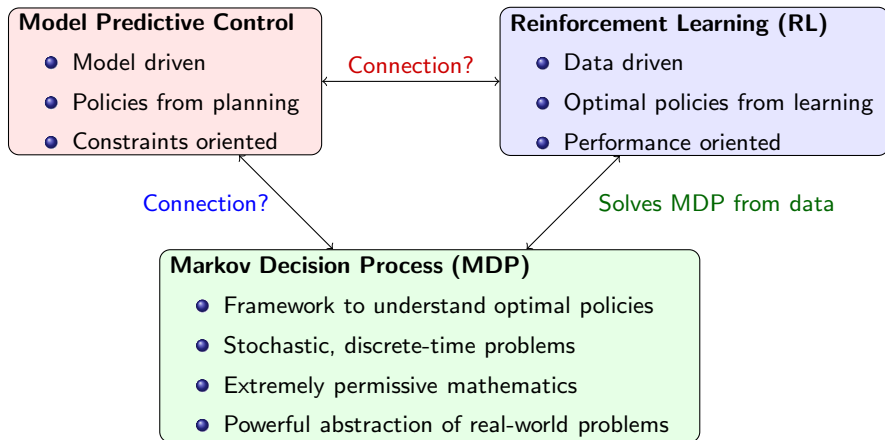
Theoretical Framework to connect RL and MPC



Theoretical Framework to connect RL and MPC

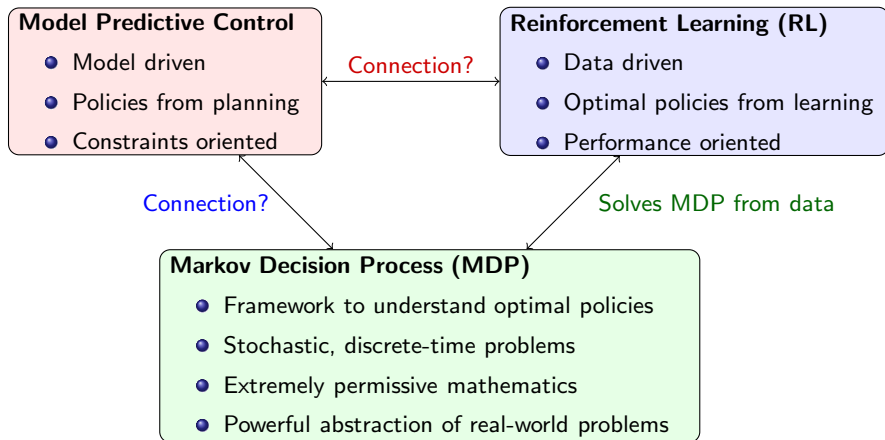


Theoretical Framework to connect RL and MPC



We have connected MPC and RL from an “implementation” point of view

Theoretical Framework to connect RL and MPC



**We have connected MPC and RL from an “implementation” point of view
But understanding what we are doing is about connecting MPC to MDPs!!**

Markov Decision Processes (MDP)

Stochastic state transitions (real world)

$$\mathbf{s}, \mathbf{a} \rightarrow \mathbf{s}_+$$

(state-action $\rightarrow$ next state)

Cost function (instant performance)

$$L(\mathbf{s}, \mathbf{a}) \in \mathbb{R}$$

A (fairly) general way of describing optimal control

Markov Decision Processes (MDP)

Stochastic state transitions (real world)

$$\mathbf{s}, \mathbf{a} \rightarrow \mathbf{s}_+$$

(state-action $\rightarrow$ next state)

Cost function (instant performance)

$$L(\mathbf{s}, \mathbf{a}) \in \mathbb{R}$$

- **Policy**

$$\mathbf{a} = \boldsymbol{\pi}(\mathbf{s})$$

is how we act on the system

A (fairly) general way of describing optimal control

Markov Decision Processes (MDP)

Stochastic state transitions (real world)

$$\mathbf{s}, \mathbf{a} \rightarrow \mathbf{s}_+$$

(state-action $\rightarrow$ next state)

Cost function (instant performance)

$$L(\mathbf{s}, \mathbf{a}) \in \mathbb{R}$$

- **Policy**

$$\mathbf{a} = \pi(\mathbf{s})$$

is how we act on the system

- **Closed-loop performance**

$$J(\pi) = \mathbb{E} \left[\sum_{k=0}^{\infty} \gamma^k L(\mathbf{s}_k, \mathbf{a}_k) \mid \pi \right]$$

with discount $\gamma \in [0, 1]$

A (fairly) general way of describing optimal control

Markov Decision Processes (MDP)

Stochastic state transitions (real world)

$$s, a \rightarrow s_+$$

(state-action $\rightarrow$ next state)

- **Policy**

$$a = \pi(s)$$

is how we act on the system

- **Closed-loop performance**

$$J(\pi) = \mathbb{E} \left[\sum_{k=0}^{\infty} \gamma^k L(s_k, a_k) \mid \pi \right]$$

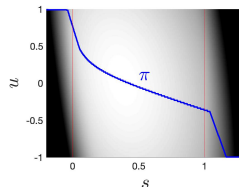
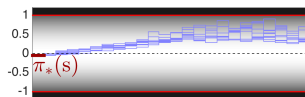
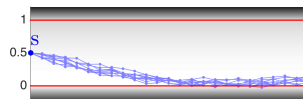
with discount $\gamma \in [0, 1]$

- **Optimal policy:** π^* from

$$\min_{\pi} J(\pi)$$

Cost function (instant performance)

$$L(s, a) \in \mathbb{R}$$



A (fairly) general way of describing optimal control

Markov Decision Processes (MDP)

Stochastic state transitions (real world)

$$s, a \rightarrow s_+$$

(state-action $\rightarrow$ next state)

- **Policy**

$$a = \pi(s)$$

is how we act on the system

- **Closed-loop performance**

$$J(\pi) = \mathbb{E} \left[\sum_{k=0}^{\infty} \gamma^k L(s_k, a_k) \mid \pi \right]$$

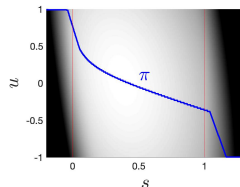
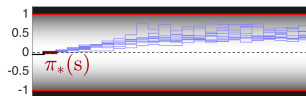
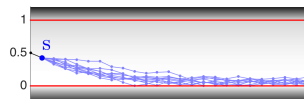
with discount $\gamma \in [0, 1]$

- **Optimal policy:** π^* from

$$\min_{\pi} J(\pi)$$

Cost function (instant performance)

$$L(s, a) \in \mathbb{R}$$



A (fairly) general way of describing optimal control

Markov Decision Processes (MDP)

Stochastic state transitions (real world)

$$s, a \rightarrow s_+$$

(state-action $\rightarrow$ next state)

- **Policy**

$$a = \pi(s)$$

is how we act on the system

- **Closed-loop performance**

$$J(\pi) = \mathbb{E} \left[\sum_{k=0}^{\infty} \gamma^k L(s_k, a_k) \mid \pi \right]$$

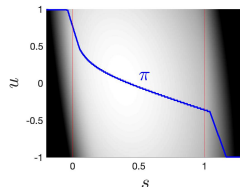
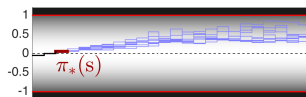
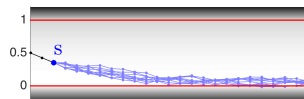
with discount $\gamma \in [0, 1]$

- **Optimal policy:** π^* from

$$\min_{\pi} J(\pi)$$

Cost function (instant performance)

$$L(s, a) \in \mathbb{R}$$



A (fairly) general way of describing optimal control

Markov Decision Processes (MDP)

Stochastic state transitions (real world)

$$s, a \rightarrow s_+$$

(state-action $\rightarrow$ next state)

- **Policy**

$$a = \pi(s)$$

is how we act on the system

- **Closed-loop performance**

$$J(\pi) = \mathbb{E} \left[\sum_{k=0}^{\infty} \gamma^k L(s_k, a_k) \mid \pi \right]$$

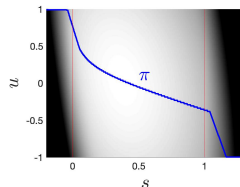
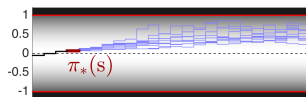
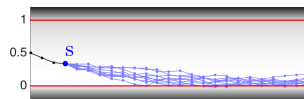
with discount $\gamma \in [0, 1]$

- **Optimal policy:** π^* from

$$\min_{\pi} J(\pi)$$

Cost function (instant performance)

$$L(s, a) \in \mathbb{R}$$



A (fairly) general way of describing optimal control

Markov Decision Processes (MDP)

Stochastic state transitions (real world)

$$s, a \rightarrow s_+$$

(state-action $\rightarrow$ next state)

- **Policy**

$$a = \pi(s)$$

is how we act on the system

- **Closed-loop performance**

$$J(\pi) = \mathbb{E} \left[\sum_{k=0}^{\infty} \gamma^k L(s_k, a_k) \mid \pi \right]$$

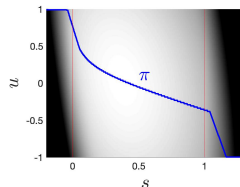
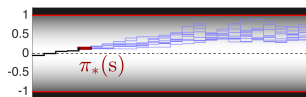
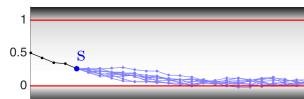
with discount $\gamma \in [0, 1]$

- **Optimal policy:** π^* from

$$\min_{\pi} J(\pi)$$

Cost function (instant performance)

$$L(s, a) \in \mathbb{R}$$



A (fairly) general way of describing optimal control

Markov Decision Processes (MDP)

Stochastic state transitions (real world)

$$s, a \rightarrow s_+$$

(state-action $\rightarrow$ next state)

- **Policy**

$$a = \pi(s)$$

is how we act on the system

- **Closed-loop performance**

$$J(\pi) = \mathbb{E} \left[\sum_{k=0}^{\infty} \gamma^k L(s_k, a_k) \mid \pi \right]$$

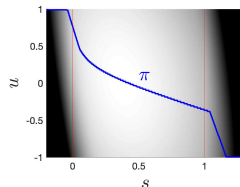
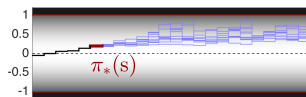
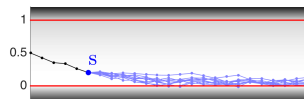
with discount $\gamma \in [0, 1]$

- **Optimal policy:** π^* from

$$\min_{\pi} J(\pi)$$

Cost function (instant performance)

$$L(s, a) \in \mathbb{R}$$



A (fairly) general way of describing optimal control

Markov Decision Processes (MDP)

Stochastic state transitions (real world)

$$s, a \rightarrow s_+$$

(state-action $\rightarrow$ next state)

- **Policy**

$$a = \pi(s)$$

is how we act on the system

- **Closed-loop performance**

$$J(\pi) = \mathbb{E} \left[\sum_{k=0}^{\infty} \gamma^k L(s_k, a_k) \mid \pi \right]$$

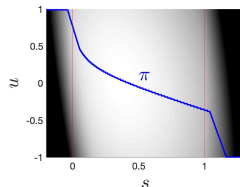
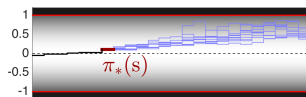
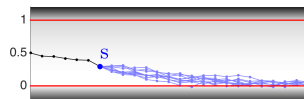
with discount $\gamma \in [0, 1]$

- **Optimal policy:** π^* from

$$\min_{\pi} J(\pi)$$

Cost function (instant performance)

$$L(s, a) \in \mathbb{R}$$



A (fairly) general way of describing optimal control

Markov Decision Processes (MDP)

Stochastic state transitions (real world)

$$s, a \rightarrow s_+$$

(state-action $\rightarrow$ next state)

- **Policy**

$$a = \pi(s)$$

is how we act on the system

- **Closed-loop performance**

$$J(\pi) = \mathbb{E} \left[\sum_{k=0}^{\infty} \gamma^k L(s_k, a_k) \mid \pi \right]$$

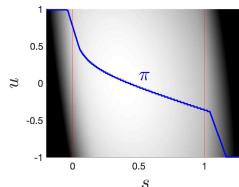
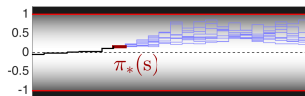
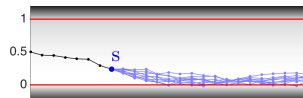
with discount $\gamma \in [0, 1]$

- **Optimal policy:** π^* from

$$\min_{\pi} J(\pi)$$

Cost function (instant performance)

$$L(s, a) \in \mathbb{R}$$



A (fairly) general way of describing optimal control

Markov Decision Processes (MDP)

Stochastic state transitions (real world)

$$s, a \rightarrow s_+$$

(state-action $\rightarrow$ next state)

- **Policy**

$$a = \pi(s)$$

is how we act on the system

- **Closed-loop performance**

$$J(\pi) = \mathbb{E} \left[\sum_{k=0}^{\infty} \gamma^k L(s_k, a_k) \mid \pi \right]$$

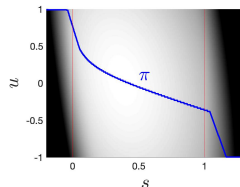
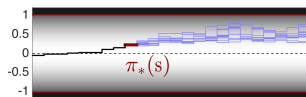
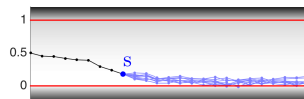
with discount $\gamma \in [0, 1]$

- **Optimal policy:** π^* from

$$\min_{\pi} J(\pi)$$

Cost function (instant performance)

$$L(s, a) \in \mathbb{R}$$



A (fairly) general way of describing optimal control

Markov Decision Processes (MDP)

Stochastic state transitions (real world)

$$s, a \rightarrow s_+$$

(state-action $\rightarrow$ next state)

- **Policy**

$$a = \pi(s)$$

is how we act on the system

- **Closed-loop performance**

$$J(\pi) = \mathbb{E} \left[\sum_{k=0}^{\infty} \gamma^k L(s_k, a_k) \mid \pi \right]$$

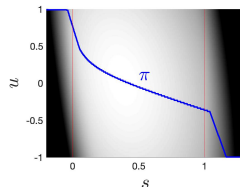
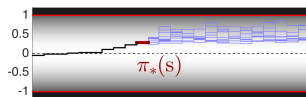
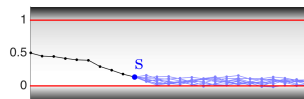
with discount $\gamma \in [0, 1]$

- **Optimal policy:** π^* from

$$\min_{\pi} J(\pi)$$

Cost function (instant performance)

$$L(s, a) \in \mathbb{R}$$



A (fairly) general way of describing optimal control

Markov Decision Processes (MDP)

Stochastic state transitions (real world)

$$s, a \rightarrow s_+$$

(state-action $\rightarrow$ next state)

- **Policy**

$$a = \pi(s)$$

is how we act on the system

- **Closed-loop performance**

$$J(\pi) = \mathbb{E} \left[\sum_{k=0}^{\infty} \gamma^k L(s_k, a_k) \mid \pi \right]$$

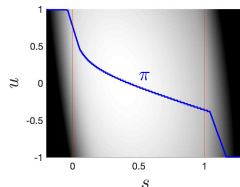
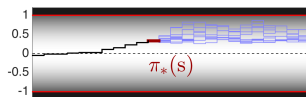
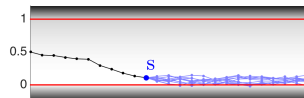
with discount $\gamma \in [0, 1]$

- **Optimal policy:** π^* from

$$\min_{\pi} J(\pi)$$

Cost function (instant performance)

$$L(s, a) \in \mathbb{R}$$



A (fairly) general way of describing optimal control

Markov Decision Processes (MDP)

Stochastic state transitions (real world)

$$s, a \rightarrow s_+$$

(state-action $\rightarrow$ next state)

- **Policy**

$$a = \pi(s)$$

is how we act on the system

- **Closed-loop performance**

$$J(\pi) = \mathbb{E} \left[\sum_{k=0}^{\infty} \gamma^k L(s_k, a_k) \mid \pi \right]$$

with discount $\gamma \in [0, 1]$

- **Optimal policy:** π^* from

$$\min_{\pi} J(\pi)$$

Cost function (instant performance)

$$L(s, a) \in \mathbb{R}$$

Impose hard constraints $h(s, a) \leq 0$?

$$L(s, a) = \begin{cases} \ell(s, a) & \text{if } h(s, a) \leq 0 \\ \infty & \text{if } h(s, a) > 0 \end{cases}$$

I will use the same view in MPC for a bit

A (fairly) general way of describing optimal control

Markov Decision Processes (MDP)

Stochastic state transitions (real world)

$$s, a \rightarrow s_+$$

(state-action $\rightarrow$ next state)

- **Policy**

$$a = \pi(s)$$

is how we act on the system

- **Closed-loop performance**

$$J(\pi) = \mathbb{E} \left[\sum_{k=0}^{\infty} \gamma^k L(s_k, a_k) \mid \pi \right]$$

with discount $\gamma \in [0, 1]$

- **Optimal policy:** π^* from

$$\min_{\pi} J(\pi)$$

Cost function (instant performance)

$$L(s, a) \in \mathbb{R}$$

Impose hard constraints $h(s, a) \leq 0$?

$$L(s, a) = \begin{cases} \ell(s, a) & \text{if } h(s, a) \leq 0 \\ \infty & \text{if } h(s, a) > 0 \end{cases}$$

I will use the same view in MPC for a bit

MDP is a go-to framework when considering general optimal control problems, useful for applications with stochastic dynamics.

A (fairly) general way of describing optimal control

Markov Decision Processes (MDP)

Stochastic state transitions (real world)

$$s, a \rightarrow s_+$$

(state-action $\rightarrow$ next state)

- **Policy**

$$a = \pi(s)$$

is how we act on the system

- **Closed-loop performance**

$$J(\pi) = \mathbb{E} \left[\sum_{k=0}^{\infty} \gamma^k L(s_k, a_k) \mid \pi \right]$$

with discount $\gamma \in [0, 1]$

- **Optimal policy:** π^* from

$$\min_{\pi} J(\pi)$$

Cost function (instant performance)

$$L(s, a) \in \mathbb{R}$$

Impose hard constraints $h(s, a) \leq 0$?

$$L(s, a) = \begin{cases} \ell(s, a) & \text{if } h(s, a) \leq 0 \\ \infty & \text{if } h(s, a) > 0 \end{cases}$$

I will use the same view in MPC for a bit

MDP is a go-to framework when considering general optimal control problems, useful for applications with stochastic dynamics.

Solution of an MDP is described by “simple” equations, but solving them is very challenging

A (fairly) general way of describing optimal control

From Policy to Repeated Planning (MPC)

Infinite horizon & discounted

$$\pi_{\infty}^* = \arg \min_{\pi} \mathbb{E} \left[\sum_{k=0}^{\infty} \gamma^k L(\mathbf{x}_k, \mathbf{a}_k) \right]$$

Policy π : state $\rightarrow$ action
belongs to a function space

*Let's transform an MDP into an MPC and
understand the approximations we make*

From Policy to Repeated Planning (MPC)

Infinite horizon & discounted

$$\pi_{\infty}^* = \arg \min_{\pi} \mathbb{E} \left[\sum_{k=0}^{\infty} \gamma^k L(\mathbf{x}_k, \mathbf{a}_k) \right]$$

Policy π : state $\rightarrow$ action
belongs to a function space

*Let's transform an MDP into an MPC and
understand the approximations we make*

Finite-horizon equivalent:

$$\pi_{0,\dots,N-1}^* = \arg \min_{\pi_{0,\dots,N-1}} \mathbb{E} \left[T(s_N) + \sum_{k=0}^{N-1} \gamma^k L(s_k, \mathbf{a}_k) \right]$$

If $T = V^*$, then $\pi_{0,\dots,N-1}^* = \pi_{\infty}^*$

From Policy to Repeated Planning (MPC)

Infinite horizon & discounted

$$\pi_{\infty}^* = \arg \min_{\pi} \mathbb{E} \left[\sum_{k=0}^{\infty} \gamma^k L(\mathbf{x}_k, \mathbf{a}_k) \right]$$

Policy π : state $\rightarrow$ action
belongs to a function space

Let's transform an MDP into an MPC and understand the approximations we make

Finite-horizon equivalent:

$$\pi_{0,\dots,N-1}^* = \arg \min_{\pi_{0,\dots,N-1}} \mathbb{E} \left[T(s_N) + \sum_{k=0}^{N-1} \gamma^k L(s_k, \mathbf{a}_k) \right]$$

If $T = V^*$, then $\pi_{0,\dots,N-1}^* = \pi_{\infty}^*$

Planning instead of a policy:

$$\min_{\mathbf{a}_{0,\dots,N-1}} \mathbb{E} \left[T(s_N) + \sum_{k=0}^{N-1} \gamma^k L(s_k, \mathbf{a}_k) \mid \mathbf{s}_0 = \mathbf{s} \right]$$

i.e. **restrict policies** to fixed $\mathbf{a}_{0,\dots,N-1}$

From Policy to Repeated Planning (MPC)

Infinite horizon & discounted

$$\pi_{\infty}^* = \arg \min_{\pi} \mathbb{E} \left[\sum_{k=0}^{\infty} \gamma^k L(\mathbf{x}_k, \mathbf{a}_k) \right]$$

Policy π : state $\rightarrow$ action
belongs to a function space

Let's transform an MDP into an MPC and understand the approximations we make

Finite-horizon equivalent:

$$\pi_{0,\dots,N-1}^* = \arg \min_{\pi_{0,\dots,N-1}} \mathbb{E} \left[T(s_N) + \sum_{k=0}^{N-1} \gamma^k L(s_k, \mathbf{a}_k) \right]$$

If $T = V^*$, then $\pi_{0,\dots,N-1}^* = \pi_{\infty}^*$

Planning instead of a policy:

$$\min_{\mathbf{a}_{0,\dots,N-1}} \mathbb{E} \left[T(s_N) + \sum_{k=0}^{N-1} \gamma^k L(s_k, \mathbf{a}_k) \mid \mathbf{s}_0 = \mathbf{s} \right]$$

i.e. **restrict policies** to fixed $\mathbf{a}_{0,\dots,N-1}$

Deterministic model, policy

$$\min_{\pi_{0,\dots,N-1}} T(s_N) + \sum_{k=0}^{N-1} \gamma^k L(s_k, \mathbf{a}_k)$$

$$\text{s.t. } s_{k+1} = \mathbf{f}(s_k, \mathbf{a}_k)$$

$$s_0 = \mathbf{s}$$

i.e. adopt **deterministic model**

From Policy to Repeated Planning (MPC)

Infinite horizon & discounted

$$\pi_{\infty}^* = \arg \min_{\pi} \mathbb{E} \left[\sum_{k=0}^{\infty} \gamma^k L(\mathbf{x}_k, \mathbf{a}_k) \right]$$

Policy π : state $\rightarrow$ action
belongs to a function space

Let's transform an MDP into an MPC and understand the approximations we make

Finite-horizon equivalent:

$$\pi_{0,\dots,N-1}^* = \arg \min_{\pi_{0,\dots,N-1}} \mathbb{E} \left[T(s_N) + \sum_{k=0}^{N-1} \gamma^k L(s_k, \mathbf{a}_k) \right]$$

If $T = V^*$, then $\pi_{0,\dots,N-1}^* = \pi_{\infty}^*$

Why attacking the problem in these ways?

Planning instead of a policy:

$$\min_{\mathbf{a}_{0,\dots,N-1}} \mathbb{E} \left[T(s_N) + \sum_{k=0}^{N-1} \gamma^k L(s_k, \mathbf{a}_k) \mid \mathbf{s}_0 = \mathbf{s} \right]$$

i.e. **restrict policies** to fixed $\mathbf{a}_{0,\dots,N-1}$

Deterministic model, planning

$$\min_{\mathbf{a}_{0,\dots,N-1}} T(s_N) + \sum_{k=0}^{N-1} \gamma^k L(s_k, \mathbf{a}_k)$$

$$\text{s.t. } \mathbf{s}_{k+1} = \mathbf{f}(\mathbf{s}_k, \mathbf{a}_k)$$

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

i.e. adopt **deterministic model**

MPC as a policy

Deterministic MPC:

$$\begin{aligned} \min_{\mathbf{u}_0, \dots, \mathbf{u}_{N-1}} \quad & T(\mathbf{x}_N) + \sum_{k=0}^{N-1} \gamma^k L(\mathbf{x}_k, \mathbf{u}_k) \\ \text{s.t.} \quad & \mathbf{x}_{k+1} = \mathbf{f}(\mathbf{x}_k, \mathbf{u}_k) \\ & \mathbf{x}_0 = \mathbf{s} \end{aligned}$$

Defines policy:

$$\pi^{\text{MPC}}(\mathbf{s}) = \mathbf{u}_0^*$$

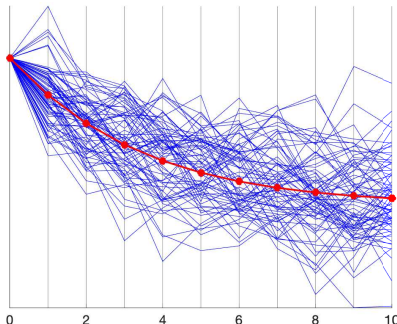
How does π^{MPC} relate to π^* ?

No reason to match:

- Planning rather than policing
- Model approximates stochasticity, often deterministic

Infinite horizon & discounted

$$\pi_{\infty}^* = \arg \min_{\pi} \mathbb{E} \left[\sum_{k=0}^{\infty} \gamma^k L(\mathbf{s}_k, \mathbf{a}_k) \right]$$



MPC as a policy

Deterministic MPC:

$$\begin{aligned} \min_{\mathbf{u}_0, \dots, \mathbf{u}_{N-1}} \quad & T(\mathbf{x}_N) + \sum_{k=0}^{N-1} \gamma^k L(\mathbf{x}_k, \mathbf{u}_k) \\ \text{s.t.} \quad & \mathbf{x}_{k+1} = \mathbf{f}(\mathbf{x}_k, \mathbf{u}_k) \\ & \mathbf{x}_0 = \mathbf{s} \end{aligned}$$

Defines policy:

$$\pi^{\text{MPC}}(\mathbf{s}) = \mathbf{u}_0^*$$

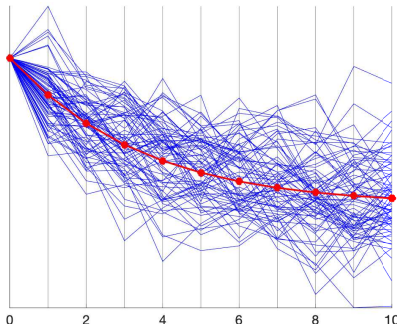
How does π^{MPC} relate to π^* ?

No reason to match:

- Planning rather than policing
- Model approximates stochasticity, often deterministic

Infinite horizon & discounted

$$\pi_\infty^* = \arg \min_{\pi} \mathbb{E} \left[\sum_{k=0}^{\infty} \gamma^k L(\mathbf{s}_k, \mathbf{a}_k) \right]$$



MPC as a policy

Deterministic MPC:

$$\begin{aligned} \min_{\mathbf{u}_0, \dots, \mathbf{u}_{N-1}} \quad & T(\mathbf{x}_N) + \sum_{k=0}^{N-1} \gamma^k L(\mathbf{x}_k, \mathbf{u}_k) \\ \text{s.t.} \quad & \mathbf{x}_{k+1} = \mathbf{f}(\mathbf{x}_k, \mathbf{u}_k) \\ & \mathbf{x}_0 = \mathbf{s} \end{aligned}$$

Defines policy:

$$\pi^{\text{MPC}}(\mathbf{s}) = \mathbf{u}_0^*$$

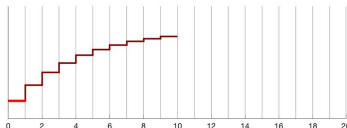
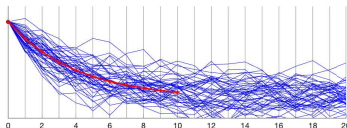
How does π^{MPC} relate to π^* ?

No reason to match:

- Planning rather than policing
- Model approximates stochasticity, often deterministic

Infinite horizon & discounted

$$\pi_\infty^* = \arg \min_{\pi} \mathbb{E} \left[\sum_{k=0}^{\infty} \gamma^k L(\mathbf{s}_k, \mathbf{a}_k) \right]$$



MPC as a policy

Deterministic MPC:

$$\begin{aligned} \min_{\mathbf{u}_0, \dots, \mathbf{u}_{N-1}} \quad & T(\mathbf{x}_N) + \sum_{k=0}^{N-1} \gamma^k L(\mathbf{x}_k, \mathbf{u}_k) \\ \text{s.t.} \quad & \mathbf{x}_{k+1} = \mathbf{f}(\mathbf{x}_k, \mathbf{u}_k) \\ & \mathbf{x}_0 = \mathbf{s} \end{aligned}$$

Defines policy:

$$\pi^{\text{MPC}}(\mathbf{s}) = \mathbf{u}_0^*$$

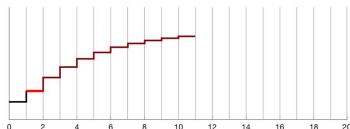
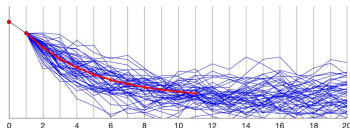
How does π^{MPC} relate to π^* ?

No reason to match:

- Planning rather than policing
- Model approximates stochasticity, often deterministic

Infinite horizon & discounted

$$\pi^* = \arg \min_{\pi} \mathbb{E} \left[\sum_{k=0}^{\infty} \gamma^k L(\mathbf{s}_k, \mathbf{a}_k) \right]$$



MPC as a policy

Deterministic MPC:

$$\begin{aligned} \min_{\mathbf{u}_0, \dots, \mathbf{u}_{N-1}} \quad & T(\mathbf{x}_N) + \sum_{k=0}^{N-1} \gamma^k L(\mathbf{x}_k, \mathbf{u}_k) \\ \text{s.t.} \quad & \mathbf{x}_{k+1} = \mathbf{f}(\mathbf{x}_k, \mathbf{u}_k) \\ & \mathbf{x}_0 = \mathbf{s} \end{aligned}$$

Defines policy:

$$\pi^{\text{MPC}}(\mathbf{s}) = \mathbf{u}_0^*$$

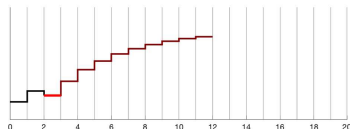
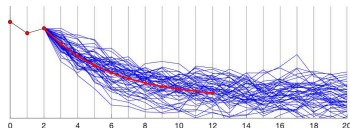
How does π^{MPC} relate to π^* ?

No reason to match:

- Planning rather than policing
- Model approximates stochasticity, often deterministic

Infinite horizon & discounted

$$\pi_\infty^* = \arg \min_{\pi} \mathbb{E} \left[\sum_{k=0}^{\infty} \gamma^k L(\mathbf{s}_k, \mathbf{a}_k) \right]$$



MPC as a policy

Deterministic MPC:

$$\begin{aligned} \min_{\mathbf{u}_0, \dots, \mathbf{u}_{N-1}} \quad & T(\mathbf{x}_N) + \sum_{k=0}^{N-1} \gamma^k L(\mathbf{x}_k, \mathbf{u}_k) \\ \text{s.t.} \quad & \mathbf{x}_{k+1} = \mathbf{f}(\mathbf{x}_k, \mathbf{u}_k) \\ & \mathbf{x}_0 = \mathbf{s} \end{aligned}$$

Defines policy:

$$\pi^{\text{MPC}}(\mathbf{s}) = \mathbf{u}_0^*$$

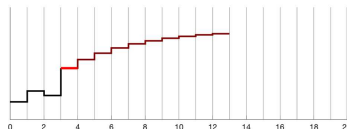
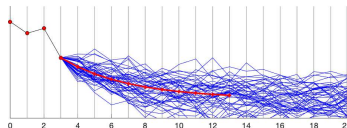
How does π^{MPC} relate to π^* ?

No reason to match:

- Planning rather than policing
- Model approximates stochasticity, often deterministic

Infinite horizon & discounted

$$\pi_\infty^* = \arg \min_{\pi} \mathbb{E} \left[\sum_{k=0}^{\infty} \gamma^k L(\mathbf{s}_k, \mathbf{a}_k) \right]$$



MPC as a policy

Deterministic MPC:

$$\begin{aligned} \min_{\mathbf{u}_0, \dots, \mathbf{u}_{N-1}} \quad & T(\mathbf{x}_N) + \sum_{k=0}^{N-1} \gamma^k L(\mathbf{x}_k, \mathbf{u}_k) \\ \text{s.t.} \quad & \mathbf{x}_{k+1} = \mathbf{f}(\mathbf{x}_k, \mathbf{u}_k) \\ & \mathbf{x}_0 = \mathbf{s} \end{aligned}$$

Defines policy:

$$\pi^{\text{MPC}}(\mathbf{s}) = \mathbf{u}_0^*$$

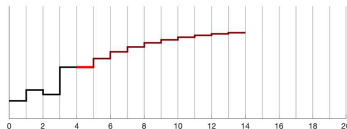
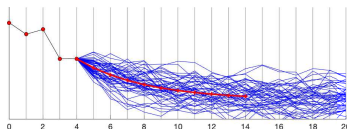
How does π^{MPC} relate to π^* ?

No reason to match:

- Planning rather than policing
- Model approximates stochasticity, often deterministic

Infinite horizon & discounted

$$\pi_\infty^* = \arg \min_{\pi} \mathbb{E} \left[\sum_{k=0}^{\infty} \gamma^k L(\mathbf{s}_k, \mathbf{a}_k) \right]$$



MPC as a policy

Deterministic MPC:

$$\begin{aligned} \min_{\mathbf{u}_0, \dots, \mathbf{u}_{N-1}} \quad & T(\mathbf{x}_N) + \sum_{k=0}^{N-1} \gamma^k L(\mathbf{x}_k, \mathbf{u}_k) \\ \text{s.t.} \quad & \mathbf{x}_{k+1} = \mathbf{f}(\mathbf{x}_k, \mathbf{u}_k) \\ & \mathbf{x}_0 = \mathbf{s} \end{aligned}$$

Defines policy:

$$\pi^{\text{MPC}}(\mathbf{s}) = \mathbf{u}_0^*$$

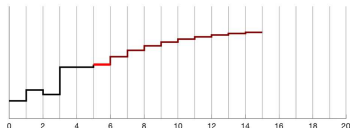
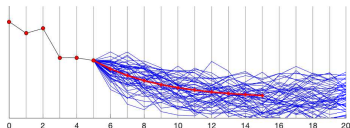
How does π^{MPC} relate to π^* ?

No reason to match:

- Planning rather than policing
- Model approximates stochasticity, often deterministic

Infinite horizon & discounted

$$\pi_\infty^* = \arg \min_{\pi} \mathbb{E} \left[\sum_{k=0}^{\infty} \gamma^k L(\mathbf{s}_k, \mathbf{a}_k) \right]$$



MPC as a policy

Deterministic MPC:

$$\begin{aligned} \min_{\mathbf{u}_0, \dots, \mathbf{u}_{N-1}} \quad & T(\mathbf{x}_N) + \sum_{k=0}^{N-1} \gamma^k L(\mathbf{x}_k, \mathbf{u}_k) \\ \text{s.t.} \quad & \mathbf{x}_{k+1} = \mathbf{f}(\mathbf{x}_k, \mathbf{u}_k) \\ & \mathbf{x}_0 = \mathbf{s} \end{aligned}$$

Defines policy:

$$\pi^{\text{MPC}}(\mathbf{s}) = \mathbf{u}_0^*$$

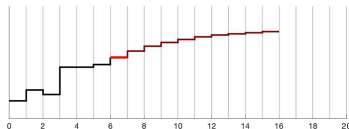
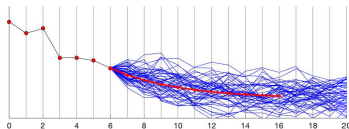
How does π^{MPC} relate to π^* ?

No reason to match:

- Planning rather than policing
- Model approximates stochasticity, often deterministic

Infinite horizon & discounted

$$\pi_\infty^* = \arg \min_{\pi} \mathbb{E} \left[\sum_{k=0}^{\infty} \gamma^k L(\mathbf{s}_k, \mathbf{a}_k) \right]$$



MPC as a policy

Deterministic MPC:

$$\begin{aligned} \min_{\mathbf{u}_0, \dots, \mathbf{u}_{N-1}} \quad & T(\mathbf{x}_N) + \sum_{k=0}^{N-1} \gamma^k L(\mathbf{x}_k, \mathbf{u}_k) \\ \text{s.t.} \quad & \mathbf{x}_{k+1} = \mathbf{f}(\mathbf{x}_k, \mathbf{u}_k) \\ & \mathbf{x}_0 = \mathbf{s} \end{aligned}$$

Defines policy:

$$\pi^{\text{MPC}}(\mathbf{s}) = \mathbf{u}_0^*$$

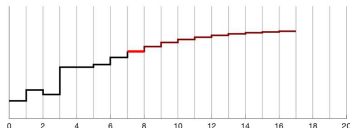
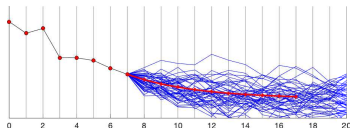
How does π^{MPC} relate to π^* ?

No reason to match:

- Planning rather than policing
- Model approximates stochasticity, often deterministic

Infinite horizon & discounted

$$\pi_\infty^* = \arg \min_{\pi} \mathbb{E} \left[\sum_{k=0}^{\infty} \gamma^k L(\mathbf{s}_k, \mathbf{a}_k) \right]$$



MPC as a policy

Deterministic MPC:

$$\begin{aligned} \min_{\mathbf{u}_0, \dots, \mathbf{u}_{N-1}} \quad & T(\mathbf{x}_N) + \sum_{k=0}^{N-1} \gamma^k L(\mathbf{x}_k, \mathbf{u}_k) \\ \text{s.t.} \quad & \mathbf{x}_{k+1} = \mathbf{f}(\mathbf{x}_k, \mathbf{u}_k) \\ & \mathbf{x}_0 = \mathbf{s} \end{aligned}$$

Defines policy:

$$\pi^{\text{MPC}}(\mathbf{s}) = \mathbf{u}_0^*$$

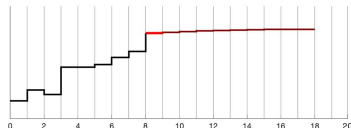
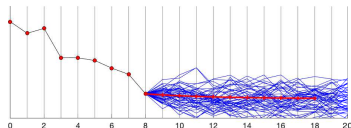
How does π^{MPC} relate to π^* ?

No reason to match:

- Planning rather than policing
- Model approximates stochasticity, often deterministic

Infinite horizon & discounted

$$\pi_\infty^* = \arg \min_{\pi} \mathbb{E} \left[\sum_{k=0}^{\infty} \gamma^k L(\mathbf{s}_k, \mathbf{a}_k) \right]$$



MPC as a policy

Deterministic MPC:

$$\begin{aligned} \min_{\mathbf{u}_0, \dots, \mathbf{u}_{N-1}} \quad & T(\mathbf{x}_N) + \sum_{k=0}^{N-1} \gamma^k L(\mathbf{x}_k, \mathbf{u}_k) \\ \text{s.t.} \quad & \mathbf{x}_{k+1} = \mathbf{f}(\mathbf{x}_k, \mathbf{u}_k) \\ & \mathbf{x}_0 = \mathbf{s} \end{aligned}$$

Defines policy:

$$\pi^{\text{MPC}}(\mathbf{s}) = \mathbf{u}_0^*$$

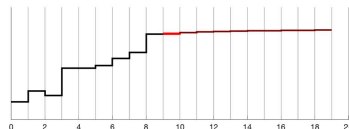
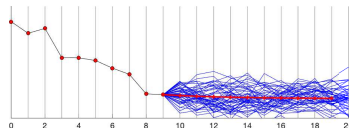
How does π^{MPC} relate to π^* ?

No reason to match:

- Planning rather than policing
- Model approximates stochasticity, often deterministic

Infinite horizon & discounted

$$\pi_\infty^* = \arg \min_{\pi} \mathbb{E} \left[\sum_{k=0}^{\infty} \gamma^k L(\mathbf{s}_k, \mathbf{a}_k) \right]$$



MPC as a policy

Deterministic MPC:

$$\begin{aligned} \min_{\mathbf{u}_0, \dots, \mathbf{u}_{N-1}} \quad & T(\mathbf{x}_N) + \sum_{k=0}^{N-1} \gamma^k L(\mathbf{x}_k, \mathbf{u}_k) \\ \text{s.t.} \quad & \mathbf{x}_{k+1} = \mathbf{f}(\mathbf{x}_k, \mathbf{u}_k) \\ & \mathbf{x}_0 = \mathbf{s} \end{aligned}$$

Defines policy:

$$\pi^{\text{MPC}}(\mathbf{s}) = \mathbf{u}_0^*$$

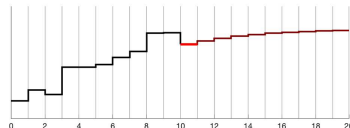
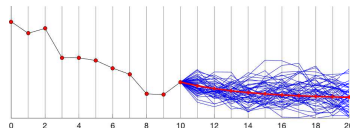
How does π^{MPC} relate to π^* ?

No reason to match:

- Planning rather than policing
- Model approximates stochasticity, often deterministic

Infinite horizon & discounted

$$\pi_{\infty}^* = \arg \min_{\pi} \mathbb{E} \left[\sum_{k=0}^{\infty} \gamma^k L(\mathbf{s}_k, \mathbf{a}_k) \right]$$



MPC as a policy

Deterministic MPC:

$$\begin{aligned} \min_{\mathbf{u}_0, \dots, \mathbf{u}_{N-1}} \quad & T(\mathbf{x}_N) + \sum_{k=0}^{N-1} \gamma^k L(\mathbf{x}_k, \mathbf{u}_k) \\ \text{s.t.} \quad & \mathbf{x}_{k+1} = \mathbf{f}(\mathbf{x}_k, \mathbf{u}_k) \\ & \mathbf{x}_0 = \mathbf{s} \end{aligned}$$

Defines policy:

$$\pi^{\text{MPC}}(\mathbf{s}) = \mathbf{u}_0^*$$

How does π^{MPC} relate to π^* ?

No reason to match:

- Planning rather than policing
- Model approximates stochasticity, often deterministic

Infinite horizon & discounted

$$\pi_{\infty}^* = \arg \min_{\pi} \mathbb{E} \left[\sum_{k=0}^{\infty} \gamma^k L(\mathbf{s}_k, \mathbf{a}_k) \right]$$



MPC as a policy

Deterministic MPC:

$$\begin{aligned} \min_{\mathbf{u}_0, \dots, \mathbf{u}_{N-1}} \quad & T(\mathbf{x}_N) + \sum_{k=0}^{N-1} \gamma^k L(\mathbf{x}_k, \mathbf{u}_k) \\ \text{s.t.} \quad & \mathbf{x}_{k+1} = \mathbf{f}(\mathbf{x}_k, \mathbf{u}_k) \\ & \mathbf{x}_0 = \mathbf{s} \end{aligned}$$

Defines policy:

$$\pi^{\text{MPC}}(\mathbf{s}) = \mathbf{u}_0^*$$

How does π^{MPC} relate to π^* ?

No reason to match:

- Planning rather than policing
- Model approximates stochasticity, often deterministic

Infinite horizon & discounted

$$\pi_{\infty}^* = \arg \min_{\pi} \mathbb{E} \left[\sum_{k=0}^{\infty} \gamma^k L(\mathbf{s}_k, \mathbf{a}_k) \right]$$



Can we clarify the relationship?

Some more context on MPC for performance

Historically MPC focuses on **constraints satisfaction & stability**. Cost is for reference tracking, not representative of the system performance.

Some more context on MPC for performance

Historically MPC focuses on **constraints satisfaction & stability**. Cost is for reference tracking, not representative of the system performance.

- “Tracking MPC”
- Classic stability theory
- Uncertainty via
 - ▶ Robust MPC
 - ▶ Stochastic MPC
- “MPC is for constraints satisfaction” (undisclosed speaker)

Some more context on MPC for performance

Historically MPC focuses on **constraints satisfaction & stability**. Cost is for reference tracking, not representative of the system performance.

More recently, focus shifted to **closed-loop performance**, e.g. energy, time, money. Cost is generic, representative of the system performance.

- “Tracking MPC”
- Classic stability theory
- Uncertainty via
 - ▶ Robust MPC
 - ▶ Stochastic MPC
- “*MPC is for constraints satisfaction*” (undisclosed speaker)

Some more context on MPC for performance

Historically MPC focuses on **constraints satisfaction & stability**. Cost is for reference tracking, not representative of the system performance.

More recently, focus shifted to **closed-loop performance**, e.g. energy, time, money. Cost is generic, representative of the system performance.

- “Tracking MPC”
- Classic stability theory
- Uncertainty via
 - ▶ Robust MPC
 - ▶ Stochastic MPC
- “MPC is for constraints satisfaction” (undisclosed speaker)

- “Economic MPC”
- Dissipativity theory
- Uncertainty via
 - ▶ Robust MPC
 - ▶ Stochastic MPC
- MPC can optimize the system performance...

Some more context on MPC for performance

Historically MPC focuses on **constraints satisfaction & stability**. Cost is for reference tracking, not representative of the system performance.

More recently, focus shifted to **closed-loop performance**, e.g. energy, time, money. Cost is generic, representative of the system performance.

- “Tracking MPC”
- Classic stability theory
- Uncertainty via
 - ▶ Robust MPC
 - ▶ Stochastic MPC
- “MPC is for constraints satisfaction” (undisclosed speaker)

- “Economic MPC”
- Dissipativity theory
- Uncertainty via
 - ▶ Robust MPC
 - ▶ Stochastic MPC
- MPC can optimize the system performance...

MPC for closed-loop performance

- is not a very old topic
- “historical robustness” of MPC does not hold in the presence of stochasticity

Some more context on MPC for performance

Historically MPC focuses on **constraints satisfaction & stability**. Cost is for reference tracking, not representative of the system performance.

More recently, focus shifted to **closed-loop performance**, e.g. energy, time, money. Cost is generic, representative of the system performance.

- “Tracking MPC”
- Classic stability theory
- Uncertainty via
 - ▶ Robust MPC
 - ▶ Stochastic MPC
- “MPC is for constraints satisfaction” (undisclosed speaker)

- “Economic MPC”
- Dissipativity theory
- Uncertainty via
 - ▶ Robust MPC
 - ▶ Stochastic MPC
- MPC can optimize the system performance...

MPC for closed-loop performance

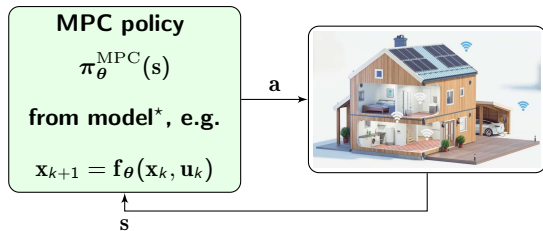
- is not a very old topic
- “historical robustness” of MPC does not hold in the presence of stochasticity

Soft claim: MPC can be used as a practical toolbox to model the solution of MDPs. This view is the “best way” for understanding what we are doing with economic MPC

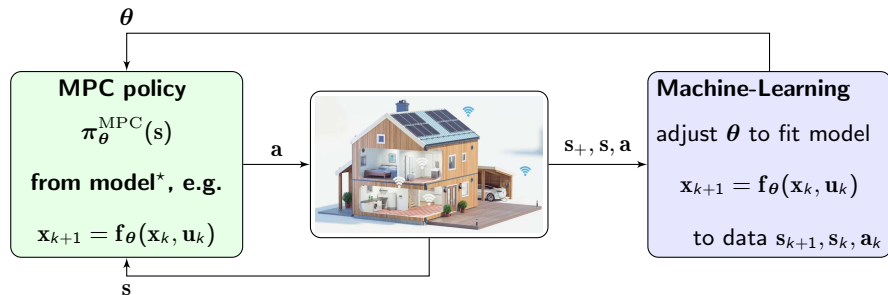
Outline

- 1 Let's rebuild some background – MPC and MDPs
- 2 MPC as a solution to MDPs
- 3 When is RL (most) beneficial for MPC?
- 4 A Deeper Look at the Theory

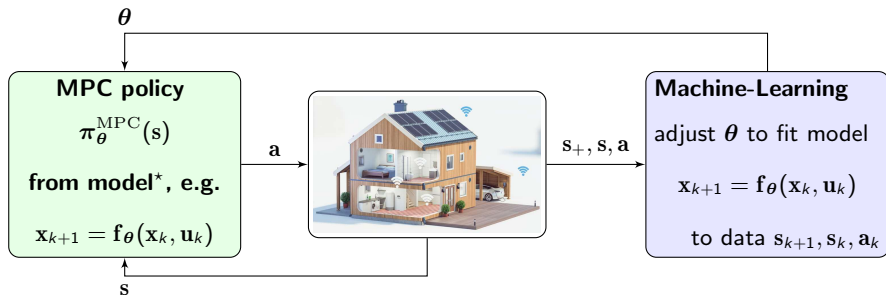
Learning for MPC - Machine Learning in-the-loop



Learning for MPC - Machine Learning in-the-loop



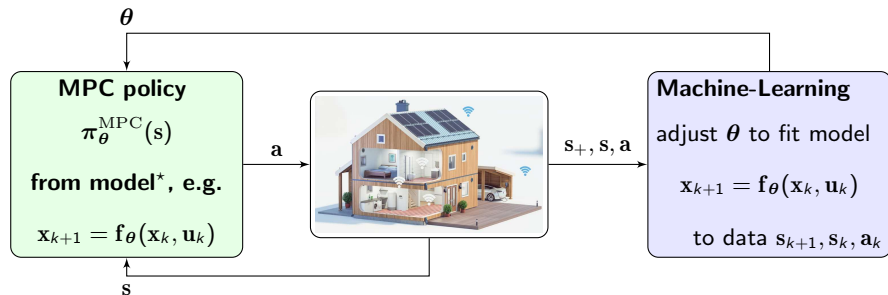
Learning for MPC - Machine Learning in-the-loop



“Machine-Learning” in-the-loop f_{θ} from

- **Physics-based:** first principles + SYSID
- **Neural Network:** DNN, LSTM, TFT, ...
- **Statistical:** GP, GPC, Probabilistic AI ...

Learning for MPC - Machine Learning in-the-loop

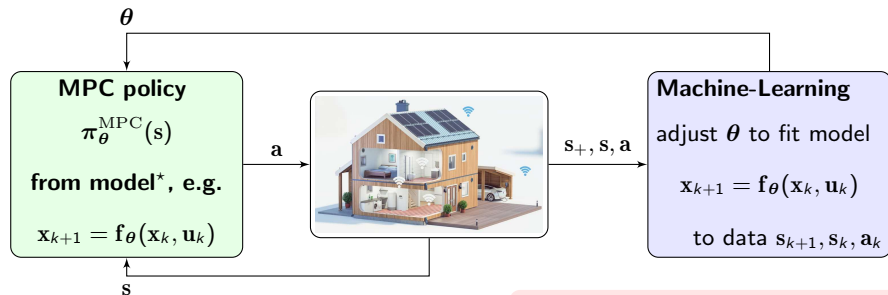


“Machine-Learning” in-the-loop f_{θ} from

- **Physics-based:** first principles + SYSID
- **Neural Network:** DNN, LSTM, TFT, ...
- **Statistical:** GP, GPC, Probabilistic AI ...

**can replace “model” by any prediction strategies:
input-output predictors, multi-step predictors, etc...*

Learning for MPC - Machine Learning in-the-loop



“Machine-Learning” in-the-loop f_{θ} from

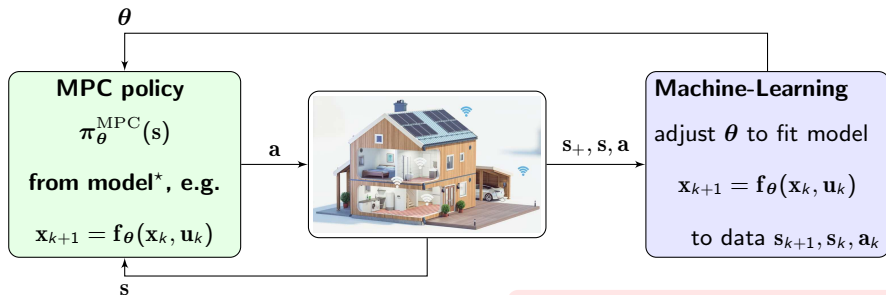
- **Physics-based:** first principles + SYSID
- **Neural Network:** DNN, LSTM, TFT, ...
- **Statistical:** GP, GPC, Probabilistic AI ...

**can replace “model” by any prediction strategies:
input-output predictors, multi-step predictors, etc...*

Paradigm

- Performance tied to prediction accuracy
- Target accuracy via ML
- Ignore that MPC is a policy

Learning for MPC - Machine Learning in-the-loop



“Machine-Learning” in-the-loop f_{θ} from

- **Physics-based:** first principles + SYSID
- **Neural Network:** DNN, LSTM, TFT, ...
- **Statistical:** GP, GPC, Probabilistic AI ...

**can replace “model” by any prediction strategies: input-output predictors, multi-step predictors, etc...*

Paradigm

- Performance tied to prediction accuracy
- Target accuracy via ML
- Ignore that MPC is a policy

We focus on “breaking” this paradigm
Learning / RL plays a key role

How can MPC model an MDP?

MDP optimal policy

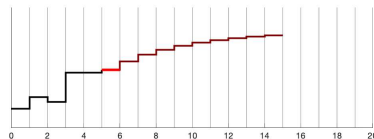
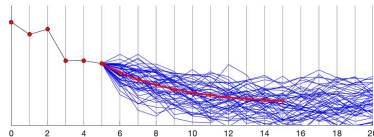
$$\pi^* = \arg \min_{\pi} \mathbb{E} \left[\sum_{k=0}^{\infty} \gamma^k L(\mathbf{x}_k, \mathbf{u}_k) \right]$$

is entirely described by $Q^*(s, \mathbf{a})$

MPC policy $\pi^{\text{MPC}}(s) = \mathbf{u}_0^*$ from

$$\min_{\mathbf{x}, \mathbf{u}} \quad T(\mathbf{x}_N) + \sum_{k=0}^{N-1} \gamma^k L(\mathbf{x}_k, \mathbf{u}_k)$$

$$\text{s.t.} \quad \mathbf{x}_{k+1} = \mathbf{f}(\mathbf{x}_k, \mathbf{u}_k), \quad \mathbf{x}_0 = \mathbf{s} \\ \mathbf{h}(\mathbf{x}_k, \mathbf{u}_k) \leq 0$$



How can MPC model an MDP?

MDP optimal policy

$$\pi^* = \arg \min_{\pi} \mathbb{E} \left[\sum_{k=0}^{\infty} \gamma^k L(\mathbf{x}_k, \mathbf{u}_k) \right]$$

is entirely described by $Q^*(s, a)$

MPC as a model of the MDP

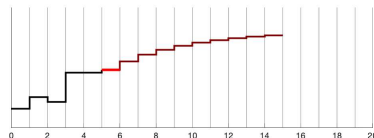
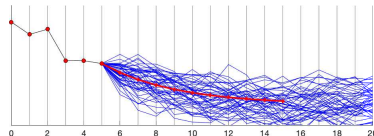
$$V^{\text{MPC}}(\mathbf{s}) := \min_{\mathbf{x}, \mathbf{u}} T(\mathbf{x}_N) + \sum_{k=0}^{N-1} \gamma^k L(\mathbf{x}_k, \mathbf{u}_k)$$
$$\text{s.t. } \mathbf{x}_{k+1} = \mathbf{f}(\mathbf{x}_k, \mathbf{u}_k), \quad \mathbf{x}_0 = \mathbf{s}$$
$$\mathbf{h}(\mathbf{x}_k, \mathbf{u}_k) \leq 0$$

MPC policy $\pi^{\text{MPC}}(s) = \mathbf{u}_0^*$ from

$$\min_{\mathbf{x}, \mathbf{u}} \quad T(\mathbf{x}_N) + \sum_{k=0}^{N-1} \gamma^k L(\mathbf{x}_k, \mathbf{u}_k)$$

$$\text{s.t. } \mathbf{x}_{k+1} = \mathbf{f}(\mathbf{x}_k, \mathbf{u}_k), \quad \mathbf{x}_0 = \mathbf{s}$$

$$\mathbf{h}(\mathbf{x}_k, \mathbf{u}_k) \leq 0$$



How can MPC model an MDP?

MDP optimal policy

$$\pi^* = \arg \min_{\pi} \mathbb{E} \left[\sum_{k=0}^{\infty} \gamma^k L(\mathbf{x}_k, \mathbf{u}_k) \right]$$

is entirely described by $Q^*(s, a)$

MPC as a model of the MDP

$$V^{\text{MPC}}(s) := \min_{\mathbf{x}, \mathbf{u}} T(\mathbf{x}_N) + \sum_{k=0}^{N-1} \gamma^k L(\mathbf{x}_k, \mathbf{u}_k)$$

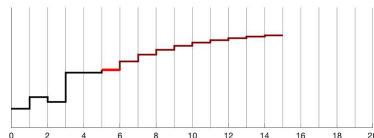
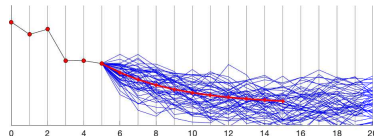
s.t. $\mathbf{x}_{k+1} = \mathbf{f}(\mathbf{x}_k, \mathbf{u}_k), \quad \mathbf{x}_0 = \mathbf{s}$
 $\mathbf{h}(\mathbf{x}_k, \mathbf{u}_k) \leq 0$

- Solving MDP $\leftarrow$ building model of Q^*
- E.g. Q-learning does that from data
- MPC can model a value function...
- Can it model an action-value function?

MPC policy $\pi^{\text{MPC}}(s) = \mathbf{u}_0^*$ from

$$\min_{\mathbf{x}, \mathbf{u}} \quad T(\mathbf{x}_N) + \sum_{k=0}^{N-1} \gamma^k L(\mathbf{x}_k, \mathbf{u}_k)$$

s.t. $\mathbf{x}_{k+1} = \mathbf{f}(\mathbf{x}_k, \mathbf{u}_k), \quad \mathbf{x}_0 = \mathbf{s}$
 $\mathbf{h}(\mathbf{x}_k, \mathbf{u}_k) \leq 0$



How can MPC model an MDP?

MDP optimal policy

$$\pi^* = \arg \min_{\pi} \mathbb{E} \left[\sum_{k=0}^{\infty} \gamma^k L(\mathbf{x}_k, \mathbf{u}_k) \right]$$

is entirely described by $Q^*(s, a)$

MPC as a model of the MDP

$$V^{\text{MPC}}(\mathbf{s}) := \min_{\mathbf{x}, \mathbf{u}} T(\mathbf{x}_N) + \sum_{k=0}^{N-1} \gamma^k L(\mathbf{x}_k, \mathbf{u}_k)$$

s.t. $\mathbf{x}_{k+1} = \mathbf{f}(\mathbf{x}_k, \mathbf{u}_k), \quad \mathbf{x}_0 = \mathbf{s}$
 $\mathbf{h}(\mathbf{x}_k, \mathbf{u}_k) \leq 0$

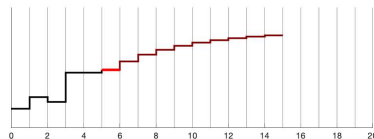
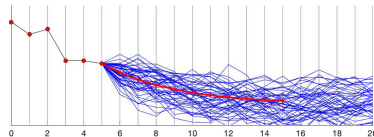
$$Q^{\text{MPC}}(\mathbf{s}, \mathbf{a}) := \min_{\mathbf{x}, \mathbf{u}} T(\mathbf{x}_N) + \sum_{k=0}^{N-1} \gamma^k L(\mathbf{x}_k, \mathbf{u}_k)$$

s.t. $\mathbf{x}_{k+1} = \mathbf{f}(\mathbf{x}_k, \mathbf{u}_k)$
 $\mathbf{h}(\mathbf{x}_k, \mathbf{u}_k) \leq 0$
 $\mathbf{x}_0 = \mathbf{s}, \quad \mathbf{u}_0 = \mathbf{a}$

MPC policy $\pi^{\text{MPC}}(\mathbf{s}) = \mathbf{u}_0^*$ from

$$\min_{\mathbf{x}, \mathbf{u}} \quad T(\mathbf{x}_N) + \sum_{k=0}^{N-1} \gamma^k L(\mathbf{x}_k, \mathbf{u}_k)$$

s.t. $\mathbf{x}_{k+1} = \mathbf{f}(\mathbf{x}_k, \mathbf{u}_k), \quad \mathbf{x}_0 = \mathbf{s}$
 $\mathbf{h}(\mathbf{x}_k, \mathbf{u}_k) \leq 0$



How can MPC model an MDP?

MDP optimal policy

$$\pi^* = \arg \min_{\pi} \mathbb{E} \left[\sum_{k=0}^{\infty} \gamma^k L(\mathbf{x}_k, \mathbf{u}_k) \right]$$

is entirely described by $Q^*(s, a)$

MPC as a model of the MDP

$$\begin{aligned} V^{\text{MPC}}(\mathbf{s}) &:= \min_{\mathbf{x}, \mathbf{u}} T(\mathbf{x}_N) + \sum_{k=0}^{N-1} \gamma^k L(\mathbf{x}_k, \mathbf{u}_k) \\ \text{s.t. } \quad &\mathbf{x}_{k+1} = \mathbf{f}(\mathbf{x}_k, \mathbf{u}_k), \quad \mathbf{x}_0 = \mathbf{s} \\ &\mathbf{h}(\mathbf{x}_k, \mathbf{u}_k) \leq 0 \end{aligned}$$

$$\begin{aligned} Q^{\text{MPC}}(\mathbf{s}, \mathbf{a}) &:= \min_{\mathbf{x}, \mathbf{u}} T(\mathbf{x}_N) + \sum_{k=0}^{N-1} \gamma^k L(\mathbf{x}_k, \mathbf{u}_k) \\ \text{s.t. } \quad &\mathbf{x}_{k+1} = \mathbf{f}(\mathbf{x}_k, \mathbf{u}_k) \\ &\mathbf{h}(\mathbf{x}_k, \mathbf{u}_k) \leq 0 \\ &\mathbf{x}_0 = \mathbf{s}, \quad \mathbf{u}_0 = \mathbf{a} \end{aligned}$$

MPC policy $\pi^{\text{MPC}}(\mathbf{s}) = \mathbf{u}_0^*$ from

$$\begin{aligned} \min_{\mathbf{x}, \mathbf{u}} \quad &T(\mathbf{x}_N) + \sum_{k=0}^{N-1} \gamma^k L(\mathbf{x}_k, \mathbf{u}_k) \\ \text{s.t. } \quad &\mathbf{x}_{k+1} = \mathbf{f}(\mathbf{x}_k, \mathbf{u}_k), \quad \mathbf{x}_0 = \mathbf{s} \\ &\mathbf{h}(\mathbf{x}_k, \mathbf{u}_k) \leq 0 \end{aligned}$$

MPC is consistent (for correct T):

$$\begin{aligned} V^{\text{MPC}}(\mathbf{s}) &= \min_{\mathbf{a}} Q^{\text{MPC}}(\mathbf{s}, \mathbf{a}) \\ \pi^{\text{MPC}}(\mathbf{s}) &= \arg \min_{\mathbf{a}} Q^{\text{MPC}}(\mathbf{s}, \mathbf{a}) \end{aligned}$$

How can MPC model an MDP?

MDP optimal policy

$$\pi^* = \arg \min_{\pi} \mathbb{E} \left[\sum_{k=0}^{\infty} \gamma^k L(\mathbf{x}_k, \mathbf{u}_k) \right]$$

is entirely described by $Q^*(s, a)$

MPC as a model of the MDP

$$\begin{aligned} V^{\text{MPC}}(s) &:= \min_{\mathbf{x}, \mathbf{u}} T(\mathbf{x}_N) + \sum_{k=0}^{N-1} \gamma^k L(\mathbf{x}_k, \mathbf{u}_k) \\ \text{s.t. } \quad &\mathbf{x}_{k+1} = \mathbf{f}(\mathbf{x}_k, \mathbf{u}_k), \quad \mathbf{x}_0 = \mathbf{s} \\ &\mathbf{h}(\mathbf{x}_k, \mathbf{u}_k) \leq 0 \end{aligned}$$

$$\begin{aligned} Q^{\text{MPC}}(\mathbf{s}, \mathbf{a}) &:= \min_{\mathbf{x}, \mathbf{u}} T(\mathbf{x}_N) + \sum_{k=0}^{N-1} \gamma^k L(\mathbf{x}_k, \mathbf{u}_k) \\ \text{s.t. } \quad &\mathbf{x}_{k+1} = \mathbf{f}(\mathbf{x}_k, \mathbf{u}_k) \\ &\mathbf{h}(\mathbf{x}_k, \mathbf{u}_k) \leq 0 \\ &\mathbf{x}_0 = \mathbf{s}, \quad \mathbf{u}_0 = \mathbf{a} \end{aligned}$$

MPC policy $\pi^{\text{MPC}}(s) = \mathbf{u}_0^*$ from

$$\begin{aligned} \min_{\mathbf{x}, \mathbf{u}} \quad &T(\mathbf{x}_N) + \sum_{k=0}^{N-1} \gamma^k L(\mathbf{x}_k, \mathbf{u}_k) \\ \text{s.t. } \quad &\mathbf{x}_{k+1} = \mathbf{f}(\mathbf{x}_k, \mathbf{u}_k), \quad \mathbf{x}_0 = \mathbf{s} \\ &\mathbf{h}(\mathbf{x}_k, \mathbf{u}_k) \leq 0 \end{aligned}$$

MPC is consistent (for correct T):

$$\begin{aligned} V^{\text{MPC}}(s) &= \min_a Q^{\text{MPC}}(s, a) \\ \pi^{\text{MPC}}(s) &= \arg \min_a Q^{\text{MPC}}(s, a) \end{aligned}$$

MPC is a complete model of MDP if:

$$Q^{\text{MPC}}(s, a) = Q^*(s, a)$$

for all s, a . Then **optimality** holds:

$$\pi^{\text{MPC}}(s) = \pi^*(s)$$

Modelling the MDP solution with MPC? Paradigm shifts...

Modelling the MDP solution with MPC? Paradigm shifts...

Shift 1: focus on performance instead of fitting

- **from:** f_θ is a model for the system dynamics
- **to:** MPC is a model of the MDP solution

Modelling the MDP solution with MPC? Paradigm shifts...

Shift 1: focus on performance instead of fitting

- **from:** f_{θ} is a model for the system dynamics
- **to:** MPC is a model of the MDP solution

Classic view...

MPC: at current state s solve

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

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

$$\mathbf{h}(\mathbf{x}_k, \mathbf{u}_k) \leq 0$$

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

gives policy $\pi_{\theta}^{\text{MPC}}(s) = \mathbf{u}_0^*$

Find θ such that prediction
“fits” the data

Modelling the MDP solution with MPC? Paradigm shifts...

Shift 1: focus on performance instead of fitting

- **from:** f_θ is a model for the system dynamics
- **to:** MPC is a model of the MDP solution

Classic view...

MPC: at current state $\mathbf{s}$ solve

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

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

$$\mathbf{h}(\mathbf{x}_k, \mathbf{u}_k) \leq 0$$

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

gives policy $\pi_\theta^{\text{MPC}}(\mathbf{s}) = \mathbf{u}_0^*$

Shift to...

**Find θ that “fits MPC to optimality”
according to the data**

$$(Q^{\text{MPC}} \rightarrow Q^* \text{ or at least } \pi^{\text{MPC}} \rightarrow \pi^*)$$

**Find θ such that prediction
“fits” the data**

Modelling the MDP solution with MPC? Paradigm shifts...

Shift 1: focus on performance instead of fitting

- **from:** f_θ is a model for the system dynamics
- **to:** MPC is a model of the MDP solution

Classic view...

MPC: at current state s solve

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

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

$$\mathbf{h}(\mathbf{x}_k, \mathbf{u}_k) \leq 0$$

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

gives policy $\pi_\theta^{\text{MPC}}(s) = \mathbf{u}_0^*$

Shift to...

**Find θ that “fits MPC to optimality”
according to the data**

$(Q^{\text{MPC}} \rightarrow Q^* \text{ or at least } \pi^{\text{MPC}} \rightarrow \pi^*)$

- $\rightarrow$ Best model for closed-loop performance
- $\neq$ Best model to fit the data!

**Find θ such that prediction
“fits” the data**

Modelling the MDP solution with MPC? Paradigm shifts...

Shift 1: focus on performance instead of fitting

- **from:** f_θ is a model for the system dynamics
- **to:** MPC is a model of the MDP solution

Classic view...

MPC: at current state s solve

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

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

$$\mathbf{h}(\mathbf{x}_k, \mathbf{u}_k) \leq 0$$

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

gives policy $\pi_\theta^{\text{MPC}}(s) = \mathbf{u}_0^*$

Find θ such that prediction
“fits” the data

Shift to...

Find θ that “fits MPC to optimality”
according to the data

($Q^{\text{MPC}} \rightarrow Q^*$ or at least $\pi^{\text{MPC}} \rightarrow \pi^*$)

- $\rightarrow$ Best model for closed-loop performance
- $\neq$ Best model to fit the data!

But $\pi^{\text{MPC}} = \pi^*$ places “high demands” on f_θ

Can we do better?

Modelling the MDP solution with MPC? Paradigm shifts...

Shift 1: focus on performance instead of fitting

- **from:** f_θ is a model for the system dynamics
- **to:** MPC is a model of the MDP solution

Classic view...

MPC: at current state s solve

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

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

$$\mathbf{h}(\mathbf{x}_k, \mathbf{u}_k) \leq 0$$

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

gives policy $\pi_\theta^{\text{MPC}}(s) = \mathbf{u}_0^*$

Find θ such that prediction
“fits” the data

Shift to...

Find θ that “fits MPC to optimality”
according to the data

($Q^{\text{MPC}} \rightarrow Q^*$ or at least $\pi^{\text{MPC}} \rightarrow \pi^*$)

- $\rightarrow$ Best model for closed-loop performance
- $\neq$ Best model to fit the data!

But $\pi^{\text{MPC}} = \pi^*$ places “high demands” on f_θ

Can we do better? **Yes!**

Modelling the MDP solution with MPC? Paradigm shifts...

Shift 1: focus on performance instead of fitting

- **from:** f_θ is a model for the system dynamics
- **to:** MPC is a model of the MDP solution

Classic view...

MPC: at current state s solve

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

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

$$\mathbf{h}(\mathbf{x}_k, \mathbf{u}_k) \leq 0$$

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

gives policy $\pi_\theta^{\text{MPC}}(s) = \mathbf{u}_0^*$

Find θ such that prediction
“fits” the data

Shift to...

Find θ that “fits MPC to optimality”
according to the data

$(Q^{\text{MPC}} \rightarrow Q^* \text{ or at least } \pi^{\text{MPC}} \rightarrow \pi^*)$

Shift 2: “holistic” parametrization

$$\min_{\mathbf{x}, \mathbf{u}} \quad T_\theta(\mathbf{x}_N) + \sum_{k=0}^{N-1} L_\theta(\mathbf{x}_k, \mathbf{u}_k)$$

$$\text{s.t.} \quad \mathbf{x}_{k+1} = \mathbf{f}_\theta(\mathbf{x}_k, \mathbf{u}_k), \quad \mathbf{x}_0 = s$$

$$\mathbf{h}_\theta(\mathbf{x}_k, \mathbf{u}_k) \leq 0$$

i.e. cost and constraints are part of the model

Modelling the MDP solution with MPC? Paradigm shifts...

Shift 1: focus on performance instead of fitting

- **from:** f_θ is a model for the system dynamics
- **to:** MPC is a model of the MDP solution

Classic view...

MPC: at current state $\mathbf{s}$ solve

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

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

$$\mathbf{h}(\mathbf{x}_k, \mathbf{u}_k) \leq 0$$

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

gives policy $\pi_\theta^{\text{MPC}}(\mathbf{s}) = \mathbf{u}_0^*$

Shift to...

**Find θ that “fits MPC to optimality”
according to the data**

$$(Q^{\text{MPC}} \rightarrow Q^* \text{ or at least } \pi^{\text{MPC}} \rightarrow \pi^*)$$

**Find θ such that prediction
“fits” the data**

“Holistic” parametrization - Is that formally justified? Yes...

Full MPC parametrization:

$$\min_{\mathbf{x}, \mathbf{u}} \quad T_{\theta}(\mathbf{x}_N) + \sum_{k=0}^{N-1} L_{\theta}(\mathbf{x}_k, \mathbf{u}_k)$$

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

$$\mathbf{h}_{\theta}(\mathbf{x}_k, \mathbf{u}_k) \leq 0$$

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

gives policy $\pi_{\theta}^{\text{MPC}}(\mathbf{s}) = \mathbf{u}_0^*$

“Holistic” parametrization - Is that formally justified? Yes...

Full MPC parametrization:

$$\min_{\mathbf{x}, \mathbf{u}} \quad T_{\theta}(\mathbf{x}_N) + \sum_{k=0}^{N-1} L_{\theta}(\mathbf{x}_k, \mathbf{u}_k)$$

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

$$\mathbf{h}_{\theta}(\mathbf{x}_k, \mathbf{u}_k) \leq 0$$

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

gives policy $\pi_{\theta}^{\text{MPC}}(\mathbf{s}) = \mathbf{u}_0^*$

Theorem: under some technical conditions and for a “rich” parametrization of the MPC, there is a θ such that

$$Q_{\theta}^{\text{MPC}}(\mathbf{s}, \mathbf{a}) = Q^*(\mathbf{s}, \mathbf{a})$$

$$\pi_{\theta}^{\text{MPC}}(\mathbf{s}) = \pi^*(\mathbf{s})$$

even if the model $\mathbf{f}_{\theta}$ *cannot* describe the real system accurately

“Holistic” parametrization - Is that formally justified? Yes...

Full MPC parametrization:

$$\min_{\mathbf{x}, \mathbf{u}} \quad T_{\theta}(\mathbf{x}_N) + \sum_{k=0}^{N-1} L_{\theta}(\mathbf{x}_k, \mathbf{u}_k)$$

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

$$\mathbf{h}_{\theta}(\mathbf{x}_k, \mathbf{u}_k) \leq 0$$

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

gives policy $\pi_{\theta}^{\text{MPC}}(\mathbf{s}) = \mathbf{u}_0^*$

Theorem: under some technical conditions and for a “rich” parametrization of the MPC, there is a θ such that

$$Q_{\theta}^{\text{MPC}}(\mathbf{s}, \mathbf{a}) = Q^*(\mathbf{s}, \mathbf{a})$$

$$\pi_{\theta}^{\text{MPC}}(\mathbf{s}) = \pi^*(\mathbf{s})$$

even if the model $\mathbf{f}_{\theta}$ *cannot* describe the real system accurately

Remarks

- **Compensate** MPC model deficiencies in the cost + constraints

“Holistic” parametrization - Is that formally justified? Yes...

Full MPC parametrization:

$$\min_{\mathbf{x}, \mathbf{u}} \quad T_{\theta}(\mathbf{x}_N) + \sum_{k=0}^{N-1} L_{\theta}(\mathbf{x}_k, \mathbf{u}_k)$$

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

$$\mathbf{h}_{\theta}(\mathbf{x}_k, \mathbf{u}_k) \leq 0$$

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

gives policy $\pi_{\theta}^{\text{MPC}}(\mathbf{s}) = \mathbf{u}_0^*$

Theorem: under some technical conditions and for a “rich” parametrization of the MPC, there is a θ such that

$$Q_{\theta}^{\text{MPC}}(\mathbf{s}, \mathbf{a}) = Q^*(\mathbf{s}, \mathbf{a})$$

$$\pi_{\theta}^{\text{MPC}}(\mathbf{s}) = \pi^*(\mathbf{s})$$

even if the model $\mathbf{f}_{\theta}$ *cannot* describe the real system accurately

Remarks

- **Compensate** MPC model deficiencies in the cost + constraints
- **Generic:** works for robust MPC, stochastic MPC, economic MPC, Multi-Step PC, etc.

“Holistic” parametrization - Is that formally justified? Yes...

Full MPC parametrization:

$$\min_{\mathbf{x}, \mathbf{u}} \quad T_{\theta}(\mathbf{x}_N) + \sum_{k=0}^{N-1} L_{\theta}(\mathbf{x}_k, \mathbf{u}_k)$$

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

$$\mathbf{h}_{\theta}(\mathbf{x}_k, \mathbf{u}_k) \leq 0$$

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

gives policy $\pi_{\theta}^{\text{MPC}}(\mathbf{s}) = \mathbf{u}_0^*$

Theorem: under some technical conditions and for a “rich” parametrization of the MPC, there is a θ such that

$$Q_{\theta}^{\text{MPC}}(\mathbf{s}, \mathbf{a}) = Q^*(\mathbf{s}, \mathbf{a})$$

$$\pi_{\theta}^{\text{MPC}}(\mathbf{s}) = \pi^*(\mathbf{s})$$

even if the model $\mathbf{f}_{\theta}$ *cannot* describe the real system accurately

Remarks

- **Compensate** MPC model deficiencies in the cost + constraints
- **Generic:** works for robust MPC, stochastic MPC, economic MPC, Multi-Step PC, etc.
- **Sanity check:** technical conditions are mild but forbid MPC model to be “very absurd”

“Holistic” parametrization - Is that formally justified? Yes...

Full MPC parametrization:

$$\min_{\mathbf{x}, \mathbf{u}} \quad T_{\theta}(\mathbf{x}_N) + \sum_{k=0}^{N-1} L_{\theta}(\mathbf{x}_k, \mathbf{u}_k)$$

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

$$\mathbf{h}_{\theta}(\mathbf{x}_k, \mathbf{u}_k) \leq 0$$

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

gives policy $\pi_{\theta}^{\text{MPC}}(\mathbf{s}) = \mathbf{u}_0^*$

MPC is a model of the optimal policy

not a policy approximation using open-loop (model-based) predictions

Theorem: under some technical conditions and for a “rich” parametrization of the MPC, there is a θ such that

$$Q_{\theta}^{\text{MPC}}(\mathbf{s}, \mathbf{a}) = Q^*(\mathbf{s}, \mathbf{a})$$

$$\pi_{\theta}^{\text{MPC}}(\mathbf{s}) = \pi^*(\mathbf{s})$$

even if the model $\mathbf{f}_{\theta}$ *cannot* describe the real system accurately

Remarks

- **Compensate** MPC model deficiencies in the cost + constraints
- **Generic:** works for robust MPC, stochastic MPC, economic MPC, Multi-Step PC, etc.
- **Sanity check:** technical conditions are mild but forbid MPC model to be “very absurd”

How to use this? Reinforcement Learning

Policy $\pi_{\theta}^{\text{MPC}}(\mathbf{s}) = \mathbf{u}_0^*$ from

$$\min_{\mathbf{x}, \mathbf{u}} \quad T_{\theta}(\mathbf{x}_N) + \sum_{k=0}^{N-1} L_{\theta}(\mathbf{x}_k, \mathbf{u}_k)$$

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

$$\mathbf{h}_{\theta}(\mathbf{x}_k, \mathbf{u}_k) \leq 0, \quad \mathbf{x}_0 = \mathbf{s}$$

- $\min_{\theta} J(\pi_{\theta}^{\text{MPC}})$ using data
- $\theta \rightarrow J(\pi_{\theta}^{\text{MPC}})$ very implicit
- $J(\cdot)$ is the real-system!

How to use this? Reinforcement Learning

Reinforcement Learning

Tools to approximate π^* from data
This is not (necessarily) about DNNs

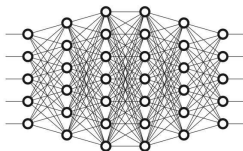
Policy $\pi_{\theta}^{\text{MPC}}(\mathbf{s}) = \mathbf{u}_0^*$ from

$$\min_{\mathbf{x}, \mathbf{u}} \quad T_{\theta}(\mathbf{x}_N) + \sum_{k=0}^{N-1} L_{\theta}(\mathbf{x}_k, \mathbf{u}_k)$$

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

$$\mathbf{h}_{\theta}(\mathbf{x}_k, \mathbf{u}_k) \leq 0, \quad \mathbf{x}_0 = \mathbf{s}$$

- $\min_{\theta} J(\pi_{\theta}^{\text{MPC}})$ using data
- $\theta \rightarrow J(\pi_{\theta}^{\text{MPC}})$ very implicit
- $J(\cdot)$ is the real-system!



How to use this? Reinforcement Learning

Policy $\pi_{\theta}^{\text{MPC}}(\mathbf{s}) = \mathbf{u}_0^*$ from

$$\min_{\mathbf{x}, \mathbf{u}} \quad T_{\theta}(\mathbf{x}_N) + \sum_{k=0}^{N-1} L_{\theta}(\mathbf{x}_k, \mathbf{u}_k)$$

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

$$\mathbf{h}_{\theta}(\mathbf{x}_k, \mathbf{u}_k) \leq 0, \quad \mathbf{x}_0 = \mathbf{s}$$

- $\min_{\theta} J(\pi_{\theta}^{\text{MPC}})$ using data
- $\theta \rightarrow J(\pi_{\theta}^{\text{MPC}})$ very implicit
- $J(\cdot)$ is the real-system!

Reinforcement Learning

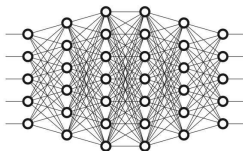
Tools to approximate π^* from data
This is not (necessarily) about DNNs

For MPC: tools to find best θ , e.g.

- **Policy Gradient:** estimations of

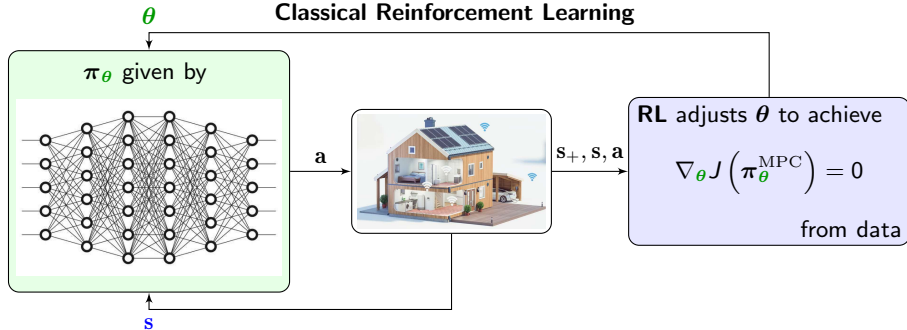
$$\nabla_{\theta} J(\pi_{\theta}^{\text{MPC}}), \quad \text{possibly} \quad \nabla_{\theta}^2 J(\pi_{\theta}^{\text{MPC}})$$

- **Q-learning:** direct “shaping” of MPC



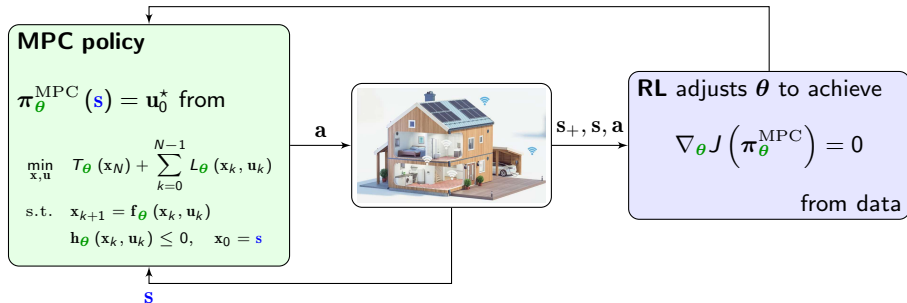
Reinforcement Learning & MPC

Classical Reinforcement Learning



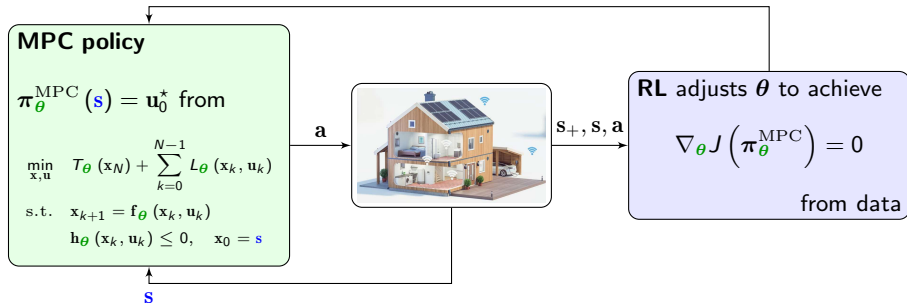
Reinforcement Learning & MPC

Reinforcement Learning over MPC



Reinforcement Learning & MPC

Reinforcement Learning over MPC

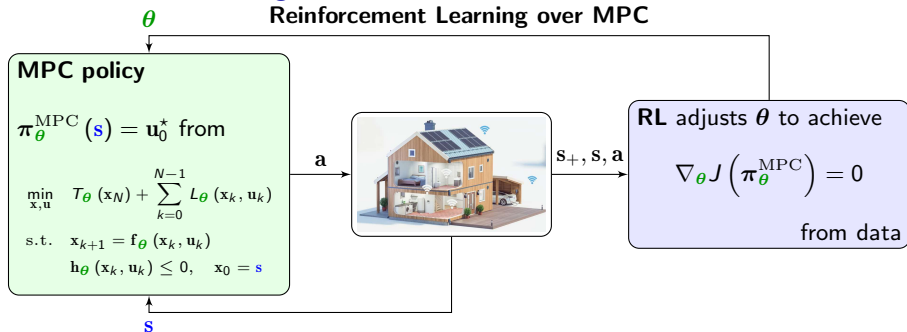


Why is it useful?

- ✓ Tunes your optimization model for real-world performance
- ✓ MPC: 25 years of results on formal guarantees & methods
- ✓ Easy to inject knowledge
- ✓ Learning does not start from scratch
- ✓ Planning provides explainability
- ✓ Software now available

Reinforcement Learning & MPC

Reinforcement Learning over MPC



Why is it useful?

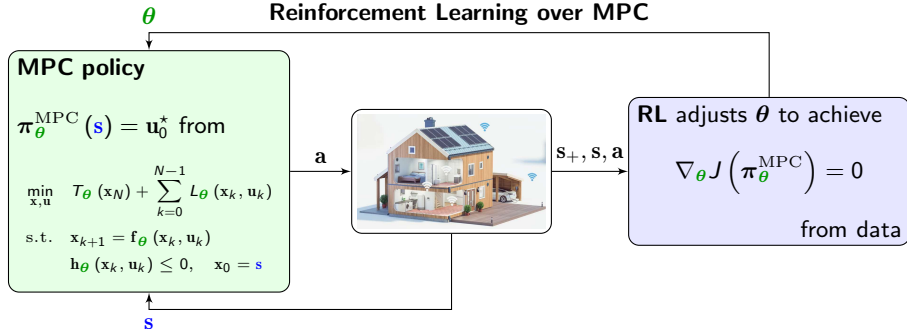
- ✓ Tunes your optimization model for real-world performance
- ✓ MPC: 25 years of results on formal guarantees & methods
- ✓ Easy to inject knowledge
- ✓ Learning does not start from scratch
- ✓ Planning provides explainability
- ✓ Software now available

Why can it be difficult?

- ✗ Optimization is computationally more expensive than a DNN
- ✗ Deployment on GPUs is lagging
- ✗ Non-convexity can get in the way...

Reinforcement Learning & MPC

Reinforcement Learning over MPC



Why is it useful?

- ✓ Tunes your optimization model for real-world performance
- ✓ MPC: 25 years of results on formal guarantees & methods
- ✓ Easy to inject knowledge
- ✓ Learning does not start from scratch
- ✓ Planning provides explainability
- ✓ Software now available

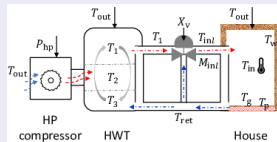
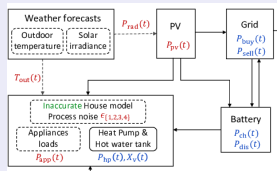
Why can it be difficult?

- ✗ Optimization is computationally more expensive than a DNN
- ✗ Deployment on GPUs is lagging
- ✗ Non-convexity can get in the way...

RL over MPC can tune your MPC for performance *beyond* what classical methods can do!

Example - Home Energy Management (simulated)

Setup



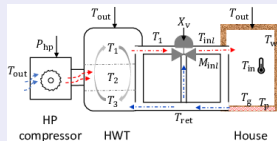
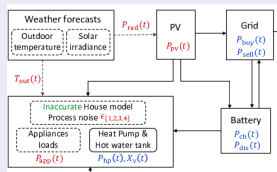
Example - Home Energy Management (simulated)

Setup



Learning process: aligned followed with closed-loop

Real world data



Example - Home Energy Management (simulated)

Setup

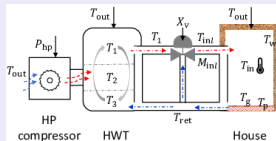
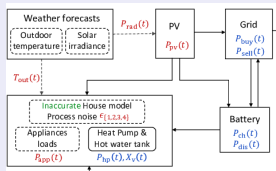


Learning process: aligned followed with closed-loop

Real world
data

Data

Classical
model fitting
to build
MPC model

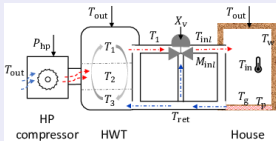
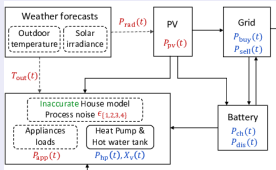
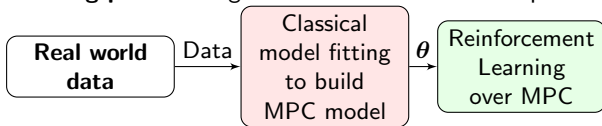


Example - Home Energy Management (simulated)

Setup

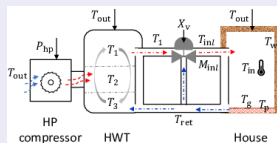
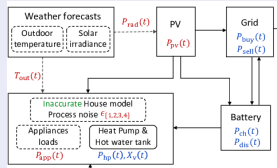


Learning process: aligned followed with closed-loop

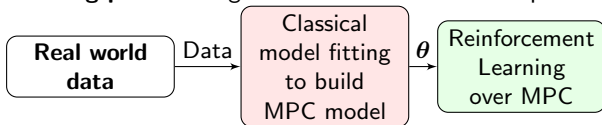


Example - Home Energy Management (simulated)

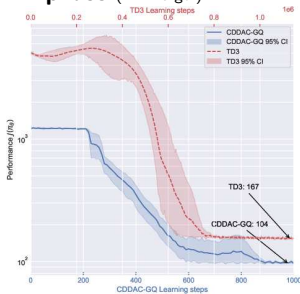
Setup



Learning process: aligned followed with closed-loop

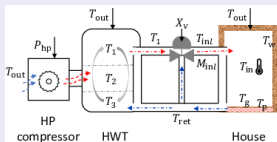
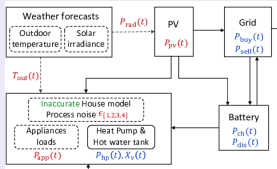


RL phase (2 RL algo.)

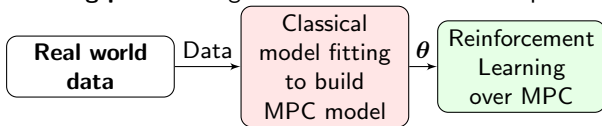


Example - Home Energy Management (simulated)

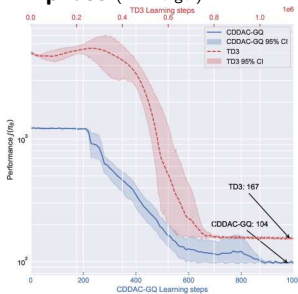
Setup



Learning process: aligned followed with closed-loop



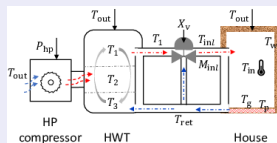
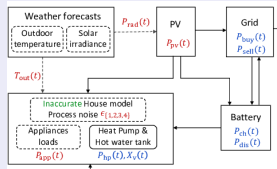
RL phase (2 RL algo.)



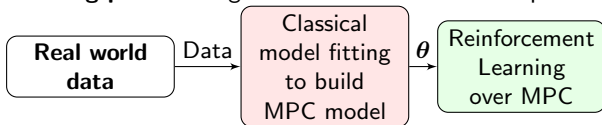
The RL step improves the MPC performance significantly from only model fitting

Example - Home Energy Management (simulated)

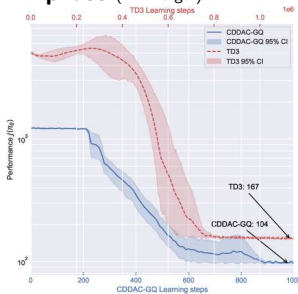
Setup



Learning process: aligned followed with closed-loop



RL phase (2 RL algo.)



The RL step improves the MPC performance significantly from only model fitting

RL over MPC can tune your MPC for performance **beyond** what classical methods can do!

Outline

- 1 Let's rebuild some background – MPC and MDPs
- 2 MPC as a solution to MDPs
- 3 When is RL (most) beneficial for MPC?
- 4 A Deeper Look at the Theory

Model fitting?

Proposed paradigm

Policy $\pi_{\theta}^{\text{MPC}}(\mathbf{s}) = \mathbf{u}_0^*$ from

$$\min_{\mathbf{x}, \mathbf{u}} \quad T_{\theta}(\mathbf{x}_N) + \sum_{k=0}^{N-1} L_{\theta}(\mathbf{x}_k, \mathbf{u}_k)$$

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

$$\mathbf{h}_{\theta}(\mathbf{x}_k, \mathbf{u}_k) \leq 0, \quad \mathbf{x}_0 = \mathbf{s}$$

Model fitting?

A step back: **model adjustment?**

Policy $\pi_{\theta}^{\text{MPC}}(\mathbf{s}) = \mathbf{u}_0^*$ from

$$\begin{aligned} \min_{\mathbf{x}, \mathbf{u}} \quad & T(\mathbf{x}_N) + \sum_{k=0}^{N-1} L(\mathbf{x}_k, \mathbf{u}_k) \\ \text{s.t.} \quad & \mathbf{x}_{k+1} = \mathbf{f}_{\theta}(\mathbf{x}_k, \mathbf{u}_k) \\ & \mathbf{h}(\mathbf{x}_k, \mathbf{u}_k) \leq 0, \quad \mathbf{x}_0 = \mathbf{s} \end{aligned}$$

Adjust θ according to (e.g.)

1. $\min_{\theta} \mathbb{E} \left[\|\mathbf{f}_{\theta}(\mathbf{s}, \mathbf{a}) - \mathbf{s}_+\|^2 \right]$ vs.
2. $\min_{\theta} J(\pi_{\theta}^{\text{MPC}})$

... is that different?

Model fitting?

A step back: **model adjustment**?

Policy $\pi_{\theta}^{\text{MPC}}(\mathbf{s}) = \mathbf{u}_0^*$ from

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

$$\text{s.t.} \quad \mathbf{x}_{k+1} = \mathbf{f}_{\theta}(\mathbf{x}_k, \mathbf{u}_k) \\ \mathbf{h}(\mathbf{x}_k, \mathbf{u}_k) \leq 0, \quad \mathbf{x}_0 = \mathbf{s}$$

Adjust θ according to (e.g.)

1. $\min_{\theta} \mathbb{E} \left[\|\mathbf{f}_{\theta}(\mathbf{s}, \mathbf{a}) - \mathbf{s}_+\|^2 \right]$ vs.
2. $\min_{\theta} J(\pi_{\theta}^{\text{MPC}})$

... is that different?

Empirical answers:

- In general it is different...
- Good to do 1. first, then 2.
- Is step 2 necessary?
 - ▶ Improves closed-loop performance
 - ▶ But gain is very case dependent (also with full parametrization)

What causes that?

Sufficient Condition for MPC Optimality (for MDPs)

MPC policy $\pi^{\text{MPC}}(\mathbf{s}) = \mathbf{u}_0^*$ from

$$\begin{aligned} \min_{\mathbf{u}, \mathbf{x}} \quad & T(\mathbf{x}_N) + \sum_{k=0}^{N-1} L(\mathbf{x}_k, \mathbf{u}_k) \\ \text{s.t.} \quad & \mathbf{x}_{k+1} = \mathbf{f}(\mathbf{x}_k, \mathbf{u}_k) \\ & \mathbf{h}(\mathbf{x}_k, \mathbf{u}_k) \leq 0, \quad \mathbf{x}_0 = \mathbf{s} \end{aligned}$$

Sufficient Condition for MPC Optimality (for MDPs)

MPC policy $\pi^{\text{MPC}}(\mathbf{s}) = \mathbf{u}_0^*$ from

$$\begin{aligned} \min_{\mathbf{u}, \mathbf{x}} \quad & T(\mathbf{x}_N) + \sum_{k=0}^{N-1} L(\mathbf{x}_k, \mathbf{u}_k) \\ \text{s.t.} \quad & \mathbf{x}_{k+1} = \mathbf{f}(\mathbf{x}_k, \mathbf{u}_k) \\ & \mathbf{h}(\mathbf{x}_k, \mathbf{u}_k) \leq 0, \quad \mathbf{x}_0 = \mathbf{s} \end{aligned}$$

Sufficient condition of optimality of MPC
relates model $\mathbf{f}$ to optimal value function and
conditional distribution $\mathbf{s}, \mathbf{a} \rightarrow \mathbf{s}_+$ of the MDP
(+ condition on T)

Sufficient Condition for MPC Optimality (for MDPs)

MPC policy $\pi^{\text{MPC}}(s) = \mathbf{u}_0^*$ from

$$\min_{\mathbf{u}, \mathbf{x}} \quad T(\mathbf{x}_N) + \sum_{k=0}^{N-1} L(\mathbf{x}_k, \mathbf{u}_k)$$
$$\text{s.t.} \quad \mathbf{x}_{k+1} = \mathbf{f}(\mathbf{x}_k, \mathbf{u}_k)$$
$$\mathbf{h}(\mathbf{x}_k, \mathbf{u}_k) \leq 0, \quad \mathbf{x}_0 = s$$

Sufficient condition of optimality of MPC
relates model $\mathbf{f}$ to optimal value function and
conditional distribution $s, \mathbf{a} \rightarrow s_+$ of the MDP
(+ condition on T)

In general

- ridge regression

$$\mathbf{f}(s, \mathbf{a}) = \mathbb{E}[s_+ | s, \mathbf{a}]$$

- max likelihood

$$\mathbf{f}(s, \mathbf{a}) = \max \rho[s_+ | s, \mathbf{a}]$$

do not satisfy the optimality
conditions

Sufficient Condition for MPC Optimality (for MDPs)

MPC policy $\pi^{\text{MPC}}(s) = u_0^*$ from

$$\min_{\mathbf{u}, \mathbf{x}} \quad T(\mathbf{x}_N) + \sum_{k=0}^{N-1} L(\mathbf{x}_k, \mathbf{u}_k)$$
$$\text{s.t.} \quad \mathbf{x}_{k+1} = \mathbf{f}(\mathbf{x}_k, \mathbf{u}_k)$$
$$\mathbf{h}(\mathbf{x}_k, \mathbf{u}_k) \leq 0, \quad \mathbf{x}_0 = s$$

Sufficient condition of optimality of MPC
relates model $\mathbf{f}$ to optimal value function and
conditional distribution $s, \mathbf{a} \rightarrow s_+$ of the MDP
(+ condition on T)

In general

- ridge regression

$$\mathbf{f}(s, \mathbf{a}) = \mathbb{E}[s_+ | s, \mathbf{a}]$$

- max likelihood

$$\mathbf{f}(s, \mathbf{a}) = \max \rho[s_+ | s, \mathbf{a}]$$

do not satisfy the optimality
conditions

Model from “classic SYSID” does not necessarily
yield the best MPC policy

Sufficient Condition for MPC Optimality (for MDPs)

MPC policy $\pi^{\text{MPC}}(s) = \mathbf{u}_0^*$ from

$$\min_{\mathbf{u}, \mathbf{x}} \quad T(\mathbf{x}_N) + \sum_{k=0}^{N-1} L(\mathbf{x}_k, \mathbf{u}_k)$$
$$\text{s.t.} \quad \mathbf{x}_{k+1} = \mathbf{f}(\mathbf{x}_k, \mathbf{u}_k)$$
$$\mathbf{h}(\mathbf{x}_k, \mathbf{u}_k) \leq 0, \quad \mathbf{x}_0 = s$$

Sufficient condition of optimality of MPC relates model $\mathbf{f}$ to optimal value function and conditional distribution $s, \mathbf{a} \rightarrow s_+$ of the MDP (+ condition on T)

In general

- ridge regression

$$\mathbf{f}(s, \mathbf{a}) = \mathbb{E}[s_+ | s, \mathbf{a}]$$

- max likelihood

$$\mathbf{f}(s, \mathbf{a}) = \max \rho[s_+ | s, \mathbf{a}]$$

do not satisfy the optimality conditions

Model from “classic SYSID” does not necessarily yield the best MPC policy

Notable exceptions: model gives $\mathbb{E}[s_+ | s, \mathbf{a}]$ and

- LQR with process noise
- By extension, locally optimal policies for
 - ▶ dissipative MDP and
 - ▶ $V^* \approx$ quadratic in positive invariant set

Smooth tracking MPC, fixed reference away from constraints, “reasonable” stochasticity

Sufficient Condition for MPC Optimality (for MDPs)

MPC policy $\pi^{\text{MPC}}(s) = \mathbf{u}_0^*$ from

$$\begin{aligned} \min_{\mathbf{u}, \mathbf{x}} \quad & T(\mathbf{x}_N) + \sum_{k=0}^{N-1} L(\mathbf{x}_k, \mathbf{u}_k) \\ \text{s.t.} \quad & \mathbf{x}_{k+1} = \mathbf{f}(\mathbf{x}_k, \mathbf{u}_k) \\ & \mathbf{h}(\mathbf{x}_k, \mathbf{u}_k) \leq 0, \quad \mathbf{x}_0 = s \end{aligned}$$

Sufficient condition of optimality of MPC
relates model $\mathbf{f}$ to optimal value function and
conditional distribution $s, \mathbf{a} \rightarrow s_+$ of the MDP
(+ condition on T)

When is RL (most) useful for MPC? *“practical” view*

- Not much if smooth tracking problem, spend most of the time close to fixed steady state, rare transients away from the constraints $\rightarrow$ small performance gain

Sufficient Condition for MPC Optimality (for MDPs)

MPC policy $\pi^{\text{MPC}}(s) = u_0^*$ from

$$\min_{u, x} \quad T(x_N) + \sum_{k=0}^{N-1} L(x_k, u_k)$$
$$\text{s.t.} \quad x_{k+1} = f(x_k, u_k)$$
$$h(x_k, u_k) \leq 0, \quad x_0 = s$$

Sufficient condition of optimality of MPC relates model f to optimal value function and conditional distribution $s, a \rightarrow s_+$ of the MDP (+ condition on T)

When is RL (most) useful for MPC? *“practical” view*

- Not much if smooth tracking problem, spend most of the time close to fixed steady state, rare transients away from the constraints $\rightarrow$ small performance gain
- Model is not rich enough to predict the (relevant) expected state transition
- Economic problem with “low dissipativity” (trajectories spread over the state space)
- Value function curvature changes a lot, problem is non-smooth
- Optimal steady-state near / at constraints
- Varying exogeneous inputs (e.g. varying prices in energy-related problems, changing reference)
- Task-based problems: racing, minimum time, termination conditions, etc.

Outline

- 1 Let's rebuild some background – MPC and MDPs
- 2 MPC as a solution to MDPs
- 3 When is RL (most) beneficial for MPC?
- 4 A Deeper Look at the Theory

Fully Parametrized MPC Can Describe the MDP Solution

Fully Parametrized MPC Can Describe the MDP Solution

The theory is about equivalences between MDPs

Fully Parametrized MPC Can Describe the MDP Solution

The theory is about equivalences between MDPs

World MDP

World MDP definition

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

Fully Parametrized MPC Can Describe the MDP Solution

The theory is about 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

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

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

Fully Parametrized MPC Can Describe the MDP Solution

The theory is about 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

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

Fully Parametrized MPC Can Describe the MDP Solution

The theory is about 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

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

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

Fully Parametrized MPC Can Describe the MDP Solution

The theory is about 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

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 $\hat{Q}^* = Q^*$

Fully Parametrized MPC Can Describe the MDP Solution

The theory is about 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

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 Equivalences

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 Equivalences

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 set of models $\hat{\rho}$ such that $\hat{Q}^* = Q^*$

More on MDP Equivalences

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 set of models $\hat{\rho}$ such that $\hat{Q}^* = Q^*$
- Conditions for model $\hat{\rho}$ “optimality” $\neq$ min of classical loss functions (except. LQR)

More on MDP Equivalences

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 set of models $\hat{\rho}$ 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 Equivalences

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 set of models $\hat{\rho}$ 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

MPC is a “weird” undiscounted Model MDP

- Deterministic model is a trivial $\hat{\rho}$, then optimal policy for Model MDP $\equiv$ planning
- Finite horizon: ok if correct terminal cost (strong assumption, but learning can tune it)
- Scenario tree: rough $\hat{\rho}$ and policy.
- Robust MPC: carries support $\hat{\rho}$, worst-case cost is a bit off...

More on MDP Equivalences

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 set of models $\hat{\rho}$ 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

More on MDP Equivalences

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 set of models $\hat{\rho}$ 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
- If V^* is continuous and support of ρ is bounded(?) and path-connected, then $\mathbf{f}(s, a) \subset \text{support of } \rho$

More on MDP Equivalences

World MDP

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

Model MDP

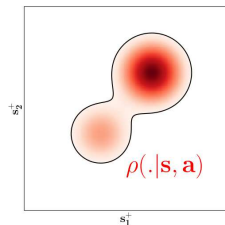
- States s and actions a
- Cost $\hat{L}(s, a)$
- Transition $s_+ \sim \hat{\rho}[\cdot | s, a]$

Theory says that

- Under some technical conditions - there is a $\hat{L}$ such that $\hat{Q}^* = Q^*$
- For $\hat{L} = L$ there is a set of models $\hat{\rho}$ 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
- If V^* is continuous and support of ρ is bounded(?) and path-connected, then $\mathbf{f}(s, a) \subset \text{support of } \rho$



More on MDP Equivalences

World MDP

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

Model MDP

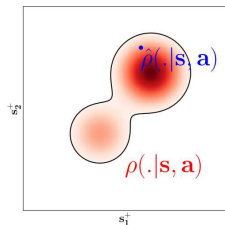
- States s and actions a
- Cost $\hat{L}(s, a)$
- Transition $s_+ \sim \hat{\rho}[\cdot | s, a]$

Theory says that

- Under some technical conditions - there is a $\hat{L}$ such that $\hat{Q}^* = Q^*$
- For $\hat{L} = L$ there is a set of models $\hat{\rho}$ 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
- If V^* is continuous and support of ρ is bounded(?) and path-connected, then $\mathbf{f}(s, a) \subset \text{support of } \rho$



More on MDP Equivalences

World MDP

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

Model MDP

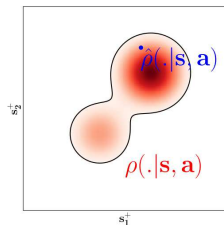
- States s and actions a
- Cost $\hat{L}(s, a)$
- Transition $s_+ \sim \hat{\rho}[\cdot | s, a]$

Theory says that

- Under some technical conditions - there is a $\hat{L}$ such that $\hat{Q}^* = Q^*$
- For $\hat{L} = L$ there is a set of models $\hat{\rho}$ 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
- If V^* is continuous and support of ρ is bounded(?) and path-connected, then $\mathbf{f}(s, a) \subset \text{support of } \rho$
- *More simply said:* we can build a “plausible” optimal deterministic model $\mathbf{f}$ (predictions have prob. > 0).



More on MDP Equivalences

World MDP

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

Model MDP

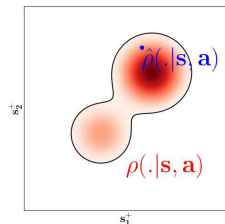
- States s and actions a
- Cost $\hat{L}(s, a)$
- Transition $s_+ \sim \hat{\rho}[\cdot | s, a]$

Theory says that

- Under some technical conditions - there is a $\hat{L}$ such that $\hat{Q}^* = Q^*$
- For $\hat{L} = L$ there is a set of models $\hat{\rho}$ 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
- If V^* is continuous and support of ρ is bounded(?) and path-connected, then $\mathbf{f}(s, a) \subset \text{support of } \rho$
- *More simply said:* we can build a “plausible” optimal deterministic model $\mathbf{f}$ (predictions have prob. > 0).
- Cost modification promotes “simplicity”.



Is there a catch?

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 $\hat{Q}^* = Q^*$

Technical assumptions? They are very technical and not intuitive at all. But not very demanding, and there is a “self-fulfilling” effect in learning, i.e. RL “learns” to satisfy them.

Is there a catch?

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 $\hat{Q}^* = Q^*$

But we are making a strong assumption (in plain sight)

Is there a catch?

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 $\hat{Q}^* = Q^*$

But we are making a strong assumption (in plain sight)

- Theory assumes that World MDP and Model MDP have the same state s
- I.e. we need to know the state of the World MDP to build a correct Model MDP
- That is often unrealistic...

Is there a catch?

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 $\hat{Q}^* = Q^*$

But we are making a strong assumption (in plain sight)

- Theory assumes that World MDP and Model MDP have the same state s
- I.e. we need to know the state of the World MDP to build a correct Model MDP
- That is often unrealistic...

The end of a nice theory? No, but it is a word of caution...

Is there a catch?

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 $\hat{Q}^* = Q^*$

Eventually the Markov state s “boils down to” the history of past observations

How to embed that in a Model MDP?

Is there a catch?

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 $\hat{Q}^* = Q^*$

Eventually the Markov state s “boils down to” the history of past observations

How to embed that in a Model MDP?

- 1 “Input-output models”: ARX, multi-step predictors (linear or not), etc.

Is there a catch?

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 $\hat{Q}^* = Q^*$

Eventually the Markov state s “boils down to” the history of past observations

How to embed that in a Model MDP?

- 1 “Input-output models”: ARX, multi-step predictors (linear or not), etc.
- 2 Latent states (a.k.a. embeddings or compressed representations) giving an AI-driven lower-dimensional version of that history

Is there a catch?

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 $\hat{Q}^* = Q^*$

Eventually the Markov state s “boils down to” the history of past observations

How to embed that in a Model MDP?

- 1 “Input-output models”: ARX, multi-step predictors (linear or not), etc.
- 2 Latent states (a.k.a. embeddings or compressed representations) giving an AI-driven lower-dimensional version of that history
- 3 State observers giving a model-based meaningful low-dimensional version of that history

Is there a catch?

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 $\hat{Q}^* = Q^*$

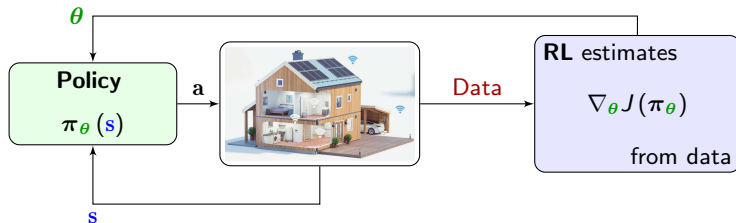
Eventually the Markov state s “boils down to” the history of past observations

How to embed that in a Model MDP?

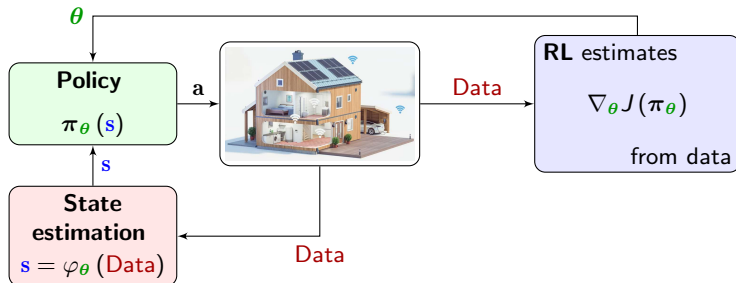
- 1 “Input-output models”: ARX, multi-step predictors (linear or not), etc.
- 2 Latent states (a.k.a. embeddings or compressed representations) giving an AI-driven lower-dimensional version of that history
- 3 State observers giving a model-based meaningful low-dimensional version of that history

Options 2 and 3 are part of the decision-making process, back-propagate through policy + state “constructor”.

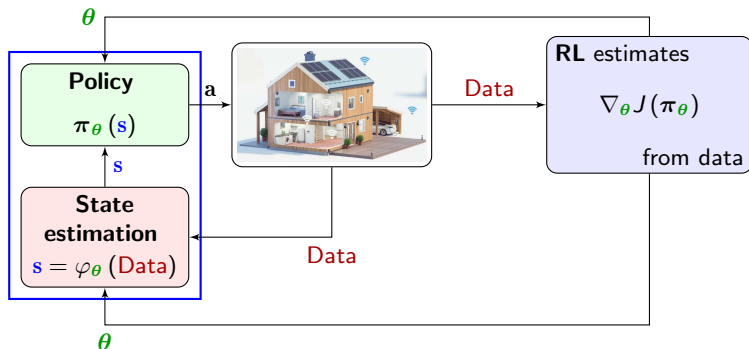
State Estimation in the Learning Loop



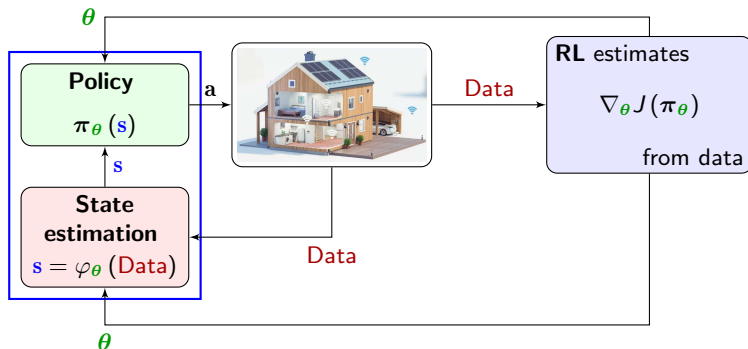
State Estimation in the Learning Loop



State Estimation in the Learning Loop



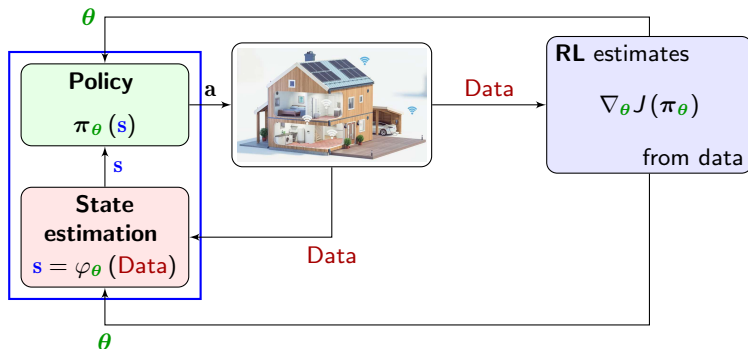
State Estimation in the Learning Loop



Remarks

- Policy is $\pi_{\theta}^D(\text{Data}) = \pi_{\theta}(\varphi_{\theta}(\text{Data}))$

State Estimation in the Learning Loop

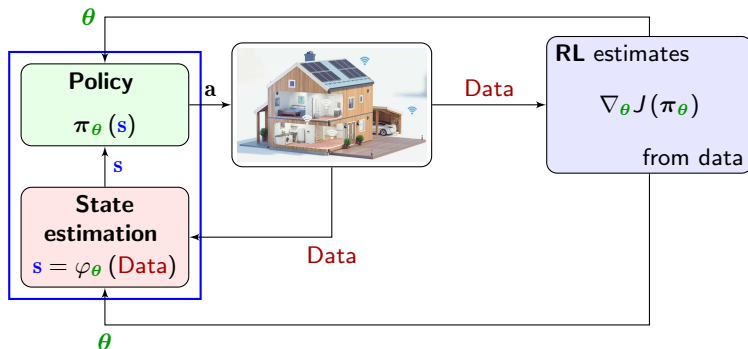


Remarks

- Policy is $\pi_{\theta}^D(\text{Data}) = \pi_{\theta}(\varphi_{\theta}(\text{Data}))$
- Gradient of the policy now is:

$$\nabla_{\theta} \pi_{\theta}^D(\text{Data}) = \nabla_{\theta} \pi_{\theta}(s) + \nabla_{\theta} \varphi_{\theta}(\text{Data}) \nabla_s \pi_{\theta}(s)$$

State Estimation in the Learning Loop



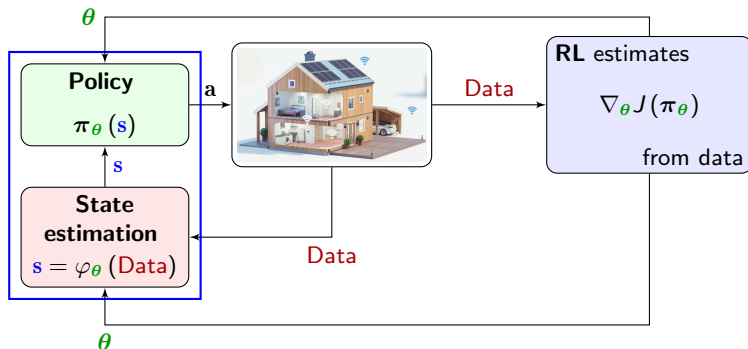
Remarks

- Policy is $\pi_{\theta}^D(\text{Data}) = \pi_{\theta}(\varphi_{\theta}(\text{Data}))$
- Gradient of the policy now is:

$$\nabla_{\theta} \pi_{\theta}^D(\text{Data}) = \nabla_{\theta} \pi_{\theta}(s) + \nabla_{\theta} \varphi_{\theta}(\text{Data}) \nabla_s \pi_{\theta}(s)$$

- Critic is looking at $\text{Data} \xrightarrow{\theta} a$, i.e. $Q^{\pi_{\theta}^D}(\text{Data}, a)$

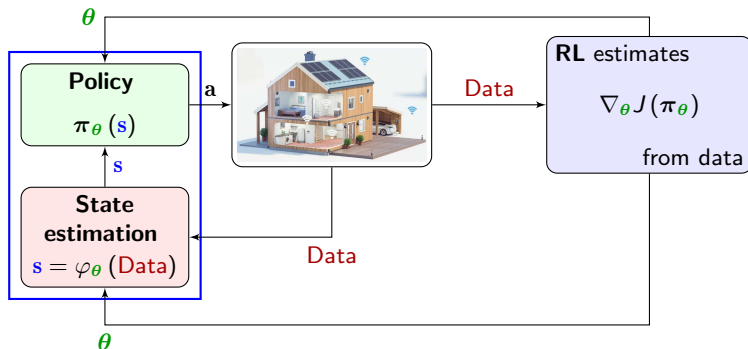
State Estimation in the Learning Loop



If $\varphi_{\theta}(\text{Data})$ is

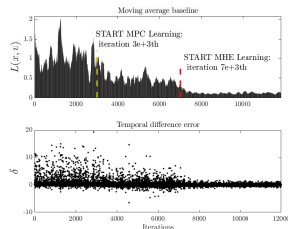
- MHE, then its gradient is computed using NLP sensitivities (same as MPC)
- DNN, then efficient sensitivity computations exist
- Bayesian inference? Belief state? Open research topic...

State Estimation in the Learning Loop



If $\varphi_\theta(\text{Data})$ is

- MHE, then its gradient is computed using NLP sensitivities (same as MPC)
- DNN, then efficient sensitivity computations exist
- Bayesian inference? Belief state? Open research topic...



Multi-Step Predictive Control

Systems with

- $\sim$ Linear dynamics
- Input-output data
- Significant stochasticity
- Modelling is difficult

Multi-Step Predictive Control

Systems with

- ~Linear dynamics
- Input-output data
- Significant stochasticity
- Modelling is difficult

Multi-step linear predictors

$$\hat{\mathbf{y}} = \Phi \begin{bmatrix} \mathbf{u} \\ \mathbf{y} \\ \mathbf{u} \end{bmatrix}$$

- Recent history of input-output sequence $\mathbf{u}, \mathbf{y}$
- Planned input sequence $\mathbf{u}$
- Predicted output sequence $\hat{\mathbf{y}}$

Multi-Step Predictive Control

SPC for $\mathbf{u}$, $\mathbf{y}$ given

$$\min_{\mathbf{u}, \hat{\mathbf{y}}} \sum_{k=0}^N L(\hat{\mathbf{y}}_k, \mathbf{u}_k)$$

$$\text{s.t. } \hat{\mathbf{y}} = \Phi \begin{bmatrix} \mathbf{u} \\ \mathbf{y} \\ \mathbf{u} \end{bmatrix}$$

$$\mathbf{h}(\hat{\mathbf{y}}_k, \mathbf{u}_k) \leq 0$$

yields policy $\pi(\mathbf{u}, \mathbf{y}) = \mathbf{u}_0^*$

Multi-step linear predictors

$$\hat{\mathbf{y}} = \Phi \begin{bmatrix} \mathbf{u} \\ \mathbf{y} \\ \mathbf{u} \end{bmatrix}$$

- Recent history of input-output sequence $\mathbf{u}, \mathbf{y}$
- Planned input sequence $\mathbf{u}$
- Predicted output sequence $\hat{\mathbf{y}}$

Multi-Step Predictive Control

SPC for $\mathbf{u}$, $\mathbf{y}$ given

$$\min_{\mathbf{u}, \hat{\mathbf{y}}} \sum_{k=0}^N L(\hat{\mathbf{y}}_k, \mathbf{u}_k)$$

$$\text{s.t. } \hat{\mathbf{y}} = \Phi \begin{bmatrix} \mathbf{u} \\ \mathbf{y} \\ \mathbf{u} \end{bmatrix}$$

$$\mathbf{h}(\hat{\mathbf{y}}_k, \mathbf{u}_k) \leq 0$$

yields policy $\pi(\mathbf{u}, \mathbf{y}) = \mathbf{u}_0^*$

Multi-step linear predictors

$$\hat{\mathbf{y}} = \Phi \begin{bmatrix} \mathbf{u} \\ \mathbf{y} \\ \mathbf{u} \end{bmatrix}$$

- Recent history of input-output sequence $\mathbf{u}, \mathbf{y}$
- Planned input sequence $\mathbf{u}$
- Predicted output sequence $\hat{\mathbf{y}}$
- Measured output sequences $\mathbf{y}$

matrix Φ can be **built from past data** $\mathcal{D}$, e.g.

$$\min_{\Phi} \sum_{i \in \mathcal{D}} \frac{1}{2} \left\| \mathbf{y}_i - \Phi \begin{bmatrix} \mathbf{u}_i \\ \mathbf{y}_i \\ \mathbf{u}_i \end{bmatrix} \right\|^2 + R(\Phi)$$

s.t. Φ causal

or alternative loss functions (e.g. quantile)

Multi-Step Predictive Control

SPC for $\mathbf{u}$, $\mathbf{y}$ given

$$\min_{\mathbf{u}, \hat{\mathbf{y}}} \sum_{k=0}^N L(\hat{\mathbf{y}}_k, \mathbf{u}_k)$$

$$\text{s.t. } \hat{\mathbf{y}} = \Phi \begin{bmatrix} \mathbf{u} \\ \mathbf{y} \\ \mathbf{u} \end{bmatrix}$$

$$\mathbf{h}(\hat{\mathbf{y}}_k, \mathbf{u}_k) \leq 0$$

yields policy $\pi(\mathbf{u}, \mathbf{y}) = \mathbf{u}_0^*$

If R promotes a dynamic system structure, regression well-posed even with limited data (and Φ high-dimensional)

But best model fit $\nRightarrow$ best policy

Multi-step linear predictors

$$\hat{\mathbf{y}} = \Phi \begin{bmatrix} \mathbf{u} \\ \mathbf{y} \\ \mathbf{u} \end{bmatrix}$$

- Recent history of input-output sequence $\mathbf{u}, \mathbf{y}$
- Planned input sequence $\mathbf{u}$
- Predicted output sequence $\hat{\mathbf{y}}$
- Measured output sequences $\mathbf{y}$

matrix Φ can be **built from past data** $\mathcal{D}$, e.g.

$$\min_{\Phi} \sum_{i \in \mathcal{D}} \frac{1}{2} \left\| \mathbf{y}_i - \Phi \begin{bmatrix} \mathbf{u}_i \\ \mathbf{y}_i \\ \mathbf{u}_i \end{bmatrix} \right\|^2 + R(\Phi)$$

s.t. Φ causal

or alternative loss functions (e.g. quantile)

Multi-Step Predictive Control

SPC for $\mathbf{u}$, $\mathbf{y}$ given

$$\min_{\mathbf{u}, \hat{\mathbf{y}}} \sum_{k=0}^N L(\hat{\mathbf{y}}_k, \mathbf{u}_k)$$

$$\text{s.t. } \hat{\mathbf{y}} = \Phi \begin{bmatrix} \mathbf{u} \\ \mathbf{y} \\ \mathbf{u} \end{bmatrix}$$

$$\mathbf{h}(\hat{\mathbf{y}}_k, \mathbf{u}_k) \leq 0$$

yields policy $\pi(\mathbf{u}, \mathbf{y}) = \mathbf{u}_0^*$

Can we close the gap via RL? **Yes!**

- RL-MPC theory applies with some twists
- State becomes $\mathbf{u}$, $\mathbf{y}$ (window of input-output)

But best model fit $\nRightarrow$ best policy

Multi-Step Predictive Control

MSPC for $\mathbf{u}$, $\mathbf{y}$ given

$$\min_{\mathbf{u}, \hat{\mathbf{y}}} \Psi_{\theta}(\mathbf{u}, \hat{\mathbf{y}}, \mathbf{u}, \mathbf{y})$$

$$\text{s.t. } \hat{\mathbf{y}} = \Phi_{\theta} \begin{bmatrix} \mathbf{u} \\ \mathbf{y} \\ \mathbf{u} \end{bmatrix}$$

$$\mathbf{H}_{\theta}(\mathbf{u}, \hat{\mathbf{y}}, \mathbf{u}, \mathbf{y}) \leq 0$$

yields policy $\pi_{\theta}(\mathbf{u}, \mathbf{y}) = \mathbf{u}_0^*$

Can we close the gap via RL? **Yes!**

- RL-MPC theory applies with some twists
- State becomes $\mathbf{u}$, $\mathbf{y}$ (window of input-output)
- Modifications in principle not localized in time

But best model fit $\nRightarrow$ best policy

Multi-Step Predictive Control

MSPC for $\mathbf{u}$, $\mathbf{y}$ given

$$\min_{\mathbf{u}, \hat{\mathbf{y}}} \Psi_{\theta}(\mathbf{u}, \hat{\mathbf{y}}, \mathbf{u}, \mathbf{y})$$

$$\text{s.t. } \hat{\mathbf{y}} = \Phi_{\theta} \begin{bmatrix} \mathbf{u} \\ \mathbf{y} \\ \mathbf{u} \end{bmatrix}$$

$$\mathbf{H}_{\theta}(\mathbf{u}, \hat{\mathbf{y}}, \mathbf{u}, \mathbf{y}) \leq 0$$

yields policy $\pi_{\theta}(\mathbf{u}, \mathbf{y}) = \mathbf{u}_0^*$

Can we close the gap via RL? **Yes!**

- RL-MPC theory applies with some twists
- State becomes $\mathbf{u}$, $\mathbf{y}$ (window of input-output)
- Modifications in principle not localized in time
- High-dimensional parameter space for RL

But best model fit $\nRightarrow$ best policy

Multi-Step Predictive Control

MSPC for $\mathbf{u}$, $\mathbf{y}$ given

$$\min_{\mathbf{u}, \hat{\mathbf{y}}} \Psi_{\theta}(\mathbf{u}, \hat{\mathbf{y}}, \mathbf{u}, \mathbf{y})$$

$$\text{s.t. } \hat{\mathbf{y}} = \Phi_{\theta} \begin{bmatrix} \mathbf{u} \\ \mathbf{y} \\ \mathbf{u} \end{bmatrix}$$

$$\mathbf{H}_{\theta}(\mathbf{u}, \hat{\mathbf{y}}, \mathbf{u}, \mathbf{y}) \leq 0$$

yields policy $\pi_{\theta}(\mathbf{u}, \mathbf{y}) = \mathbf{u}_0^*$

Can we close the gap via RL? **Yes!**

- RL-MPC theory applies with some twists
- State becomes $\mathbf{u}$, $\mathbf{y}$ (window of input-output)
- Modifications in principle not localized in time
- High-dimensional parameter space for RL

We need an add-on in Leap-C for high-speed MPC + sensitivity on this type of model structure! *For now we are down to using CVXPY.*

But best model fit $\nRightarrow$ best policy

Multi-Step Predictive Control

MSPC for $\mathbf{u}$, $\mathbf{y}$ given

$$\min_{\mathbf{u}, \hat{\mathbf{y}}} \Psi_{\theta}(\mathbf{u}, \hat{\mathbf{y}}, \mathbf{u}, \mathbf{y})$$

$$\text{s.t. } \hat{\mathbf{y}} = \Phi_{\theta} \begin{bmatrix} \mathbf{u} \\ \mathbf{y} \\ \mathbf{u} \end{bmatrix}$$

$$\mathbf{H}_{\theta}(\mathbf{u}, \hat{\mathbf{y}}, \mathbf{u}, \mathbf{y}) \leq 0$$

yields policy $\pi_{\theta}(\mathbf{u}, \mathbf{y}) = \mathbf{u}_0^*$

Can we close the gap via RL? **Yes!**

- RL-MPC theory applies with some twists
- State becomes $\mathbf{u}$, $\mathbf{y}$ (window of input-output)
- Modifications in principle not localized in time
- High-dimensional parameter space for RL

We need an add-on in Leap-C for high-speed MPC + sensitivity on this type of model structure! *For now we are down to using CVXPY.*

But best model fit $\nRightarrow$ best policy

What we have seen:

- MPC can be understood as a model of the optimal action-value function Q^* of real-world MDPs and/or of the optimal policy π^*
- MPC cost (and constraints) become part of that model
- Model that best fits the real-world does not (necessarily) yield the best policy
- RL is a toolbox to tune the MPC as a model of the MDP solution
- MPC state space should match the real world, strong assumption that can be alleviated

What we have seen:

- MPC can be understood as a model of the optimal action-value function Q^* of real-world MDPs and/or of the optimal policy π^*
- MPC cost (and constraints) become part of that model
- Model that best fits the real-world does not (necessarily) yield the best policy
- RL is a toolbox to tune the MPC as a model of the MDP solution
- MPC state space should match the real world, strong assumption that can be alleviated

What we will do next: RL over MPC

- Safe & Stable RL over MPC (*In the afternoon*)
- RL over MPC with belief states – a future prospect (*In the afternoon*)
- Beyond MPC – Model-based Decisions and AI for decisions (*Tomorrow*)

Thanks for your attention!



ResearchGate



Google Scholar