

# Reinforcement Learning over Model Predictive Control

– Why does it work? –

Dirk Reinhardt

Department of Engineering Cybernetics, NTNU, Norway

Predict, Control, Learn: Combining Model Predictive Control and Reinforcement Learning  
Tutorial Session at ECC 2026, Reykjavík, July 2026

**universität freiburg**

 **NTNU** | Norwegian University of  
Science and Technology



**University of  
Zurich**<sup>UZH</sup>



# Tutorial Roadmap

| <i>Talk</i>                                                        | <i>Presenter</i>      | <i>Duration</i> |
|--------------------------------------------------------------------|-----------------------|-----------------|
| ✓ Markov Decision Processes – a unifying perspective on MPC and RL | Jasper Hoffmann       | 20min           |
| ✓ Synthesis of Model Predictive Control and Reinforcement Learning | Rudolf Reiter         | 30min           |
| ➤ <b>Reinforcement Learning over MPC: Why does it work?</b>        | <b>Dirk Reinhardt</b> | <b>30min</b>    |
| Differentiable NMPC with acados                                    | Katrin Baumgärtner    | 20min           |
| MPCRL with leap-c                                                  | Jasper Hoffmann       | 20min           |

At the end of each talk there will be time for questions!



Tutorial webpage for slides  
and more!

Background: MPC, MDPs, and Model Fitting

MPC as a Model of the MDP Solution

Reinforcement Learning Tunes the MPC

Background: MPC, MDPs, and Model Fitting

MPC as a Model of the MDP Solution

Reinforcement Learning Tunes the MPC

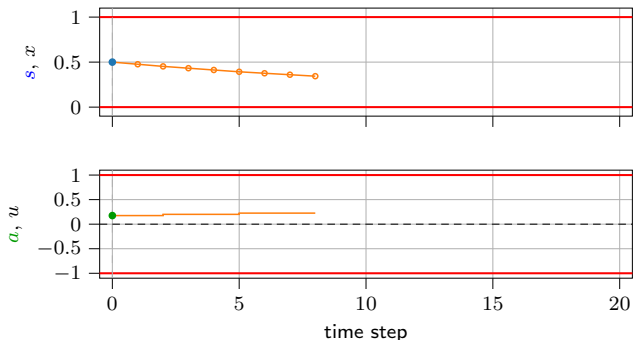
**MPC:** at **current state**  $s$

$$\min_{x, u} \quad \bar{V}(x_N) + \sum_{k=0}^{N-1} l(x_k, u_k)$$

$$\text{s.t.} \quad x_{k+1} = f(x_k, u_k)$$

$$h(x_k, u_k) \geq 0, \quad x_0 = s$$

apply **action**  $a = u_0^*$ , then repeat



MPC: at **current state**  $s$

$$\min_{x, u} \quad \bar{V}(x_N) + \sum_{k=0}^{N-1} l(x_k, u_k)$$

$$\text{s.t.} \quad x_{k+1} = f(x_k, u_k)$$

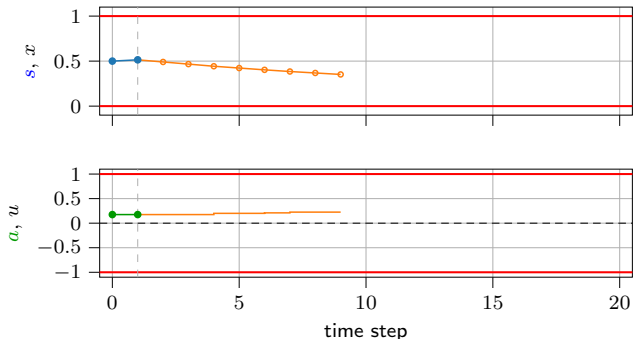
$$h(x_k, u_k) \geq 0, \quad x_0 = s$$

apply **action**  $a = u_0^*$ , then repeat

MPC

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

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



MPC: at **current state**  $s$

$$\min_{x, u} \quad \bar{V}(x_N) + \sum_{k=0}^{N-1} l(x_k, u_k)$$

$$\text{s.t.} \quad x_{k+1} = f(x_k, u_k)$$

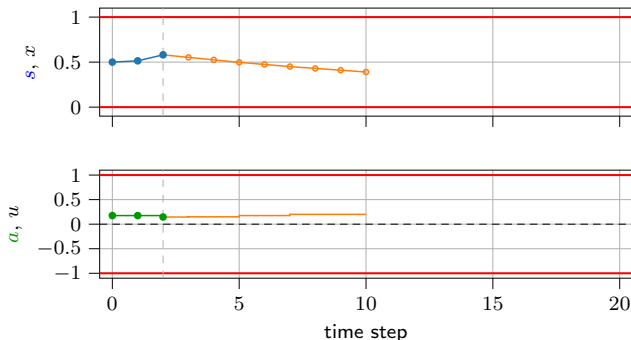
$$h(x_k, u_k) \geq 0, \quad x_0 = s$$

apply **action**  $a = u_0^*$ , then repeat

MPC

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

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



MPC: at **current state**  $s$

$$\min_{x, u} \quad \bar{V}(x_N) + \sum_{k=0}^{N-1} l(x_k, u_k)$$

$$\text{s.t.} \quad x_{k+1} = f(x_k, u_k)$$

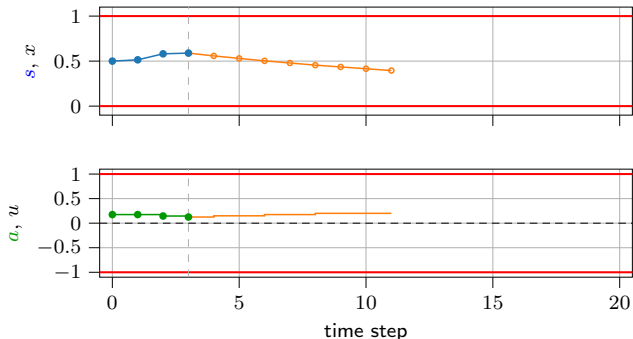
$$h(x_k, u_k) \geq 0, \quad x_0 = s$$

apply **action**  $a = u_0^*$ , then repeat

MPC

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

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





MPC: at **current state**  $s$

$$\min_{x, u} \quad \bar{V}(x_N) + \sum_{k=0}^{N-1} l(x_k, u_k)$$

$$\text{s.t.} \quad x_{k+1} = f(x_k, u_k)$$

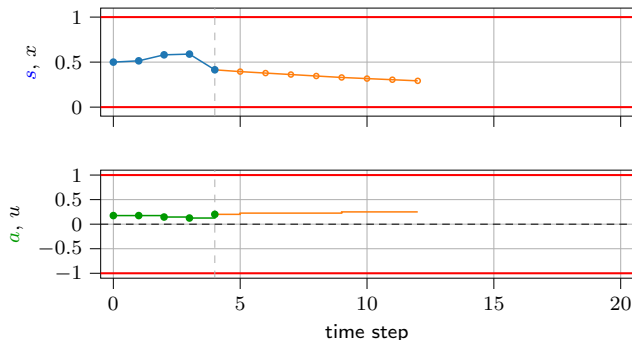
$$h(x_k, u_k) \geq 0, \quad x_0 = s$$

apply **action**  $a = u_0^*$ , then repeat

MPC

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

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



MPC: at **current state**  $s$

$$\min_{x, u} \quad \bar{V}(x_N) + \sum_{k=0}^{N-1} l(x_k, u_k)$$

$$\text{s.t.} \quad x_{k+1} = f(x_k, u_k)$$

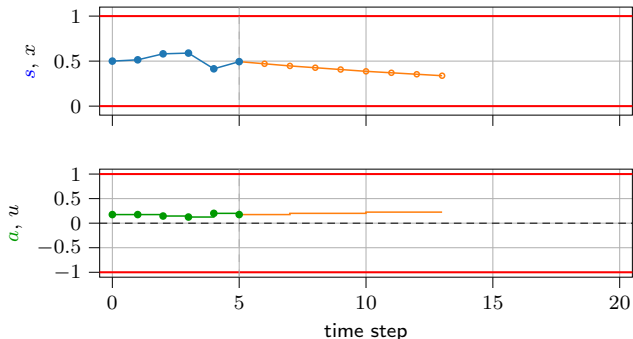
$$h(x_k, u_k) \geq 0, \quad x_0 = s$$

apply **action**  $a = u_0^*$ , then repeat

MPC

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

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



MPC: at **current state**  $s$

$$\min_{x, u} \quad \bar{V}(x_N) + \sum_{k=0}^{N-1} l(x_k, u_k)$$

$$\text{s.t.} \quad x_{k+1} = f(x_k, u_k)$$

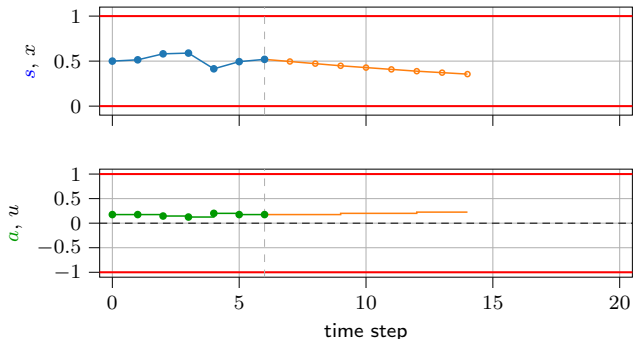
$$h(x_k, u_k) \geq 0, \quad x_0 = s$$

apply **action**  $a = u_0^*$ , then repeat

MPC

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

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



MPC: at **current state**  $s$

$$\min_{x, u} \quad \bar{V}(x_N) + \sum_{k=0}^{N-1} l(x_k, u_k)$$

$$\text{s.t.} \quad x_{k+1} = f(x_k, u_k)$$

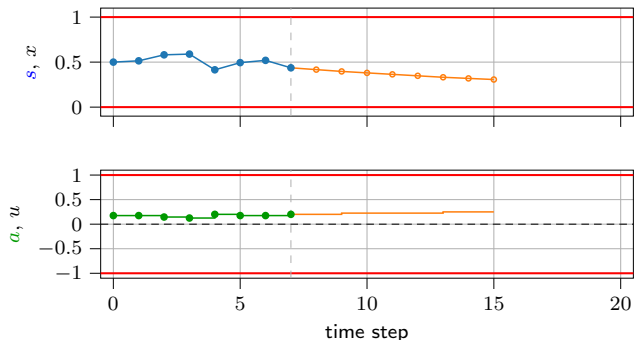
$$h(x_k, u_k) \geq 0, \quad x_0 = s$$

apply **action**  $a = u_0^*$ , then repeat

MPC

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

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



MPC: at **current state**  $s$

$$\min_{x, u} \quad \bar{V}(x_N) + \sum_{k=0}^{N-1} l(x_k, u_k)$$

$$\text{s.t.} \quad x_{k+1} = f(x_k, u_k)$$

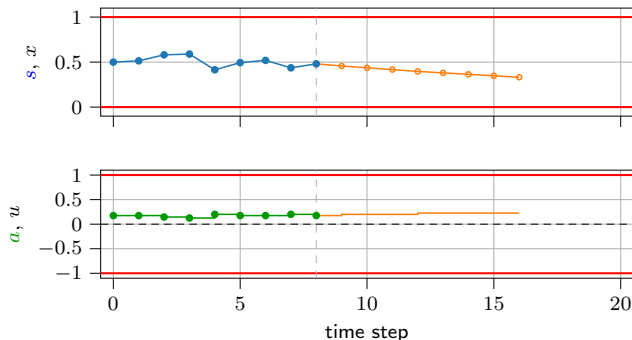
$$h(x_k, u_k) \geq 0, \quad x_0 = s$$

apply **action**  $a = u_0^*$ , then repeat

MPC

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

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



MPC: at **current state**  $s$

$$\min_{x, u} \quad \bar{V}(x_N) + \sum_{k=0}^{N-1} l(x_k, u_k)$$

$$\text{s.t.} \quad x_{k+1} = f(x_k, u_k)$$

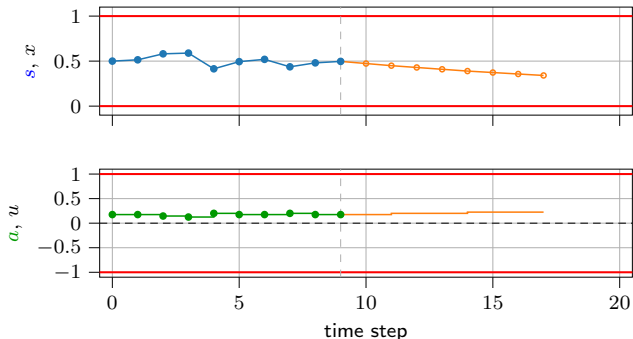
$$h(x_k, u_k) \geq 0, \quad x_0 = s$$

apply **action**  $a = u_0^*$ , then repeat

MPC

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

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



MPC: at **current state**  $s$

$$\min_{x, u} \quad \bar{V}(x_N) + \sum_{k=0}^{N-1} l(x_k, u_k)$$

$$\text{s.t.} \quad x_{k+1} = f(x_k, u_k)$$

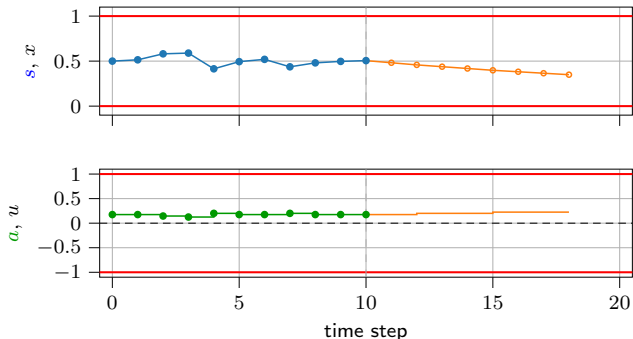
$$h(x_k, u_k) \geq 0, \quad x_0 = s$$

apply **action**  $a = u_0^*$ , then repeat

MPC

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

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



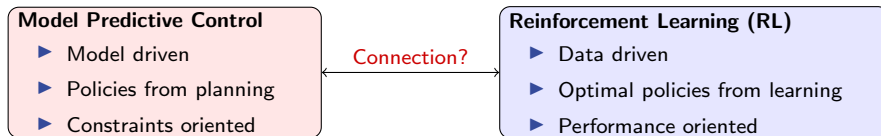
### Model Predictive Control

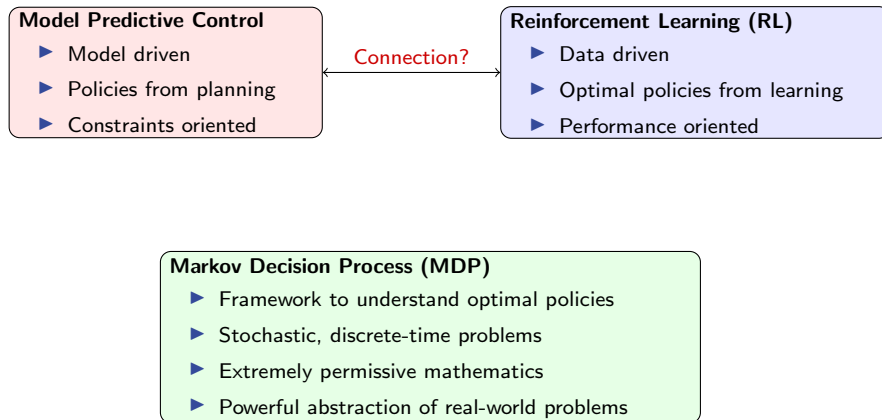
- ▶ Model driven
- ▶ Policies from planning
- ▶ Constraints oriented

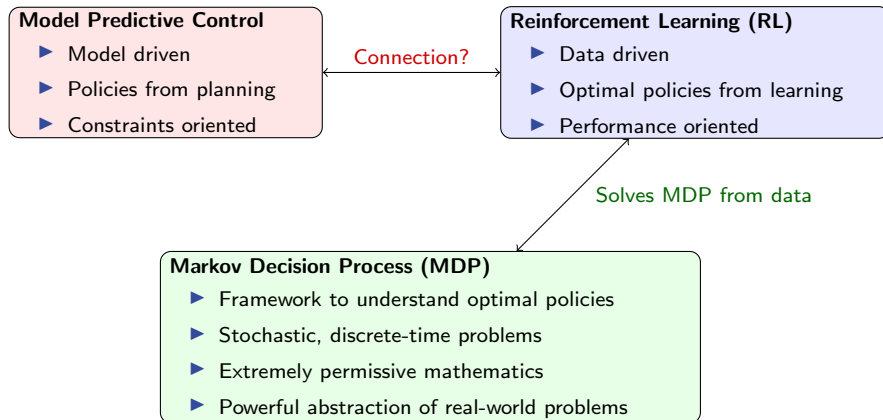
### Reinforcement Learning (RL)

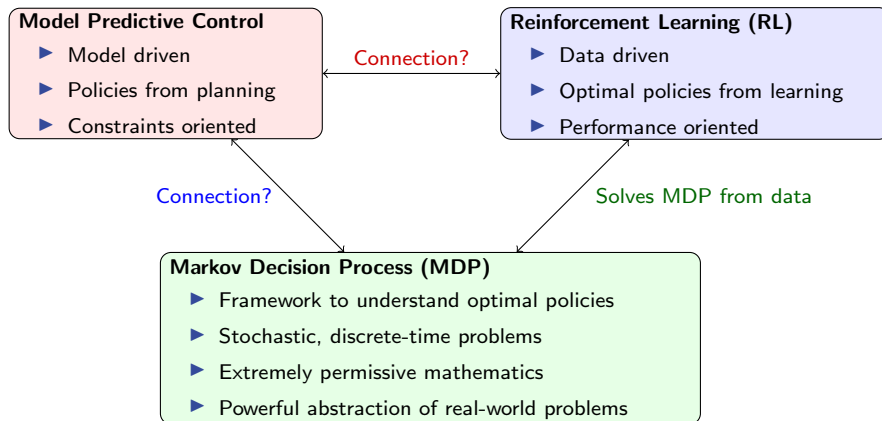
- ▶ Data driven
- ▶ Optimal policies from learning
- ▶ Performance oriented

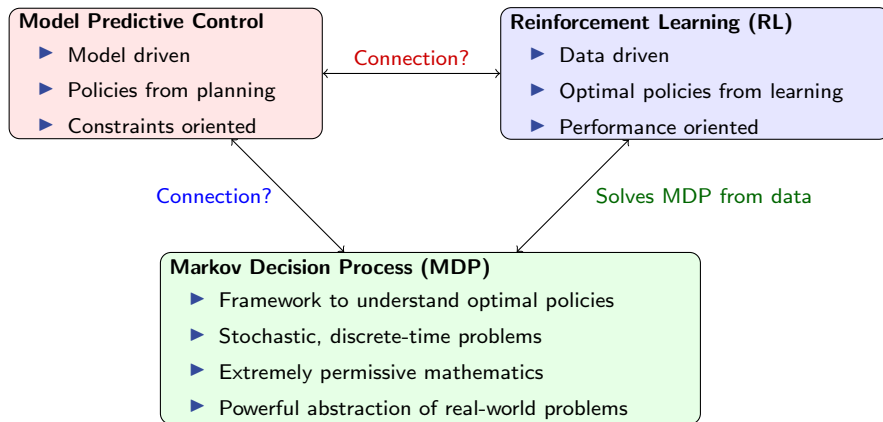




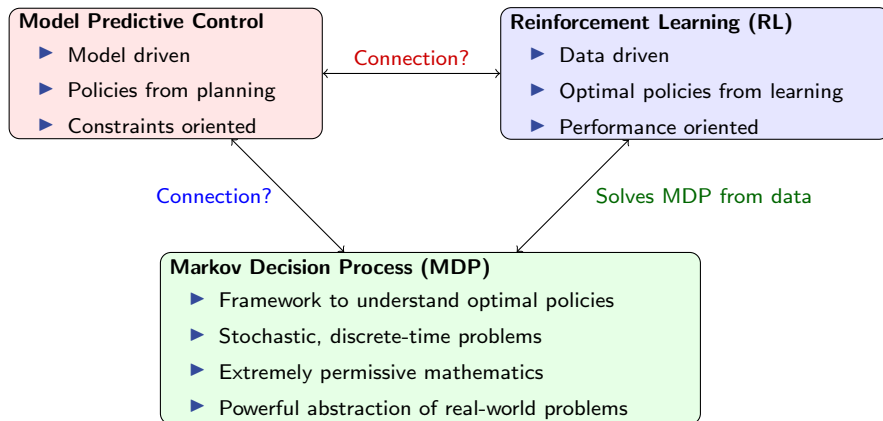








**We have connected MPC and RL from an “implementation” point of view**



**We have connected MPC and RL from an “implementation” point of view  
The theory behind is about the connection between MPC and MDPs**

**Stochastic state transitions** (real world)

$$s^+ \sim P(\cdot | s, a)$$

**Cost function** (instant performance)  $L(s, a) \in \mathbb{R}$

**Stochastic state transitions** (real world)

$$s^+ \sim P(\cdot | s, a)$$

**Cost function** (instant performance)  $L(s, a) \in \mathbb{R}$

► **Policy**

$$a = \pi(s), \quad \text{or} \quad A \sim \pi(\cdot | s)$$



**Stochastic state transitions** (real world)

$$s^+ \sim P(\cdot | s, a)$$

**Cost function** (instant performance)  $L(s, a) \in \mathbb{R}$

► **Policy**

$$a = \pi(s), \quad \text{or} \quad A \sim \pi(\cdot | s)$$

► **Closed-loop performance**

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

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

## Stochastic state transitions (real world)

$$s^+ \sim P(\cdot | s, a)$$

### ► Policy

$$a = \pi(s), \quad \text{or} \quad A \sim \pi(\cdot | s)$$

### ► Closed-loop performance

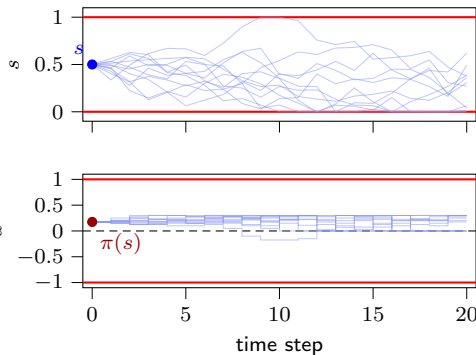
$$J(\pi) = \mathbb{E} \left[ \sum_{k=0}^{\infty} \gamma^k l(s_k, a_k) \mid \pi \right]$$

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

### ► Optimal policy: $\pi^*$ from

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

## Cost function (instant performance) $L(s, a) \in \mathbb{R}$



## Stochastic state transitions (real world)

$$s^+ \sim P(\cdot | s, a)$$

### ► Policy

$$a = \pi(s), \quad \text{or} \quad A \sim \pi(\cdot | s)$$

### ► Closed-loop performance

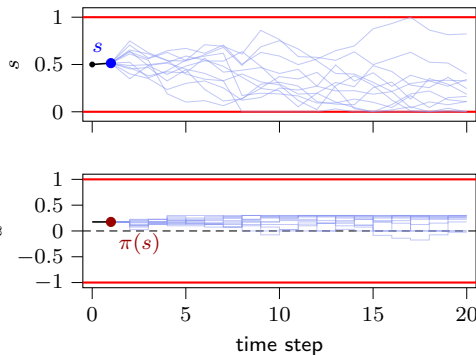
$$J(\pi) = \mathbb{E} \left[ \sum_{k=0}^{\infty} \gamma^k l(s_k, a_k) \mid \pi \right]$$

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

### ► Optimal policy: $\pi^*$ from

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

## Cost function (instant performance) $L(s, a) \in \mathbb{R}$



## Stochastic state transitions (real world)

$$s^+ \sim P(\cdot | s, a)$$

### ► Policy

$$a = \pi(s), \quad \text{or} \quad A \sim \pi(\cdot | s)$$

### ► Closed-loop performance

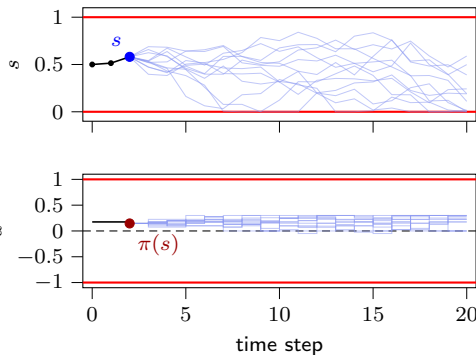
$$J(\pi) = \mathbb{E} \left[ \sum_{k=0}^{\infty} \gamma^k l(s_k, a_k) \mid \pi \right]$$

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

### ► Optimal policy: $\pi^*$ from

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

## Cost function (instant performance) $L(s, a) \in \mathbb{R}$



## Stochastic state transitions (real world)

$$s^+ \sim P(\cdot | s, a)$$

### ► Policy

$$a = \pi(s), \quad \text{or} \quad A \sim \pi(\cdot | s)$$

### ► Closed-loop performance

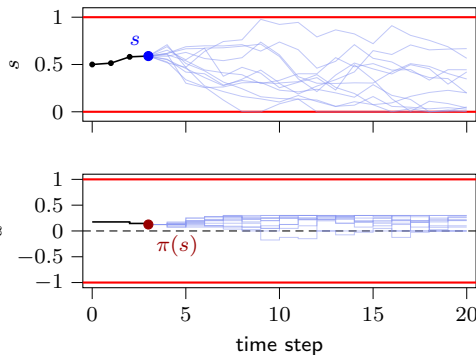
$$J(\pi) = \mathbb{E} \left[ \sum_{k=0}^{\infty} \gamma^k l(s_k, a_k) \mid \pi \right]$$

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

### ► Optimal policy: $\pi^*$ from

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

## Cost function (instant performance) $L(s, a) \in \mathbb{R}$



## Stochastic state transitions (real world)

$$s^+ \sim P(\cdot | s, a)$$

### ► Policy

$$a = \pi(s), \quad \text{or} \quad A \sim \pi(\cdot | s)$$

### ► Closed-loop performance

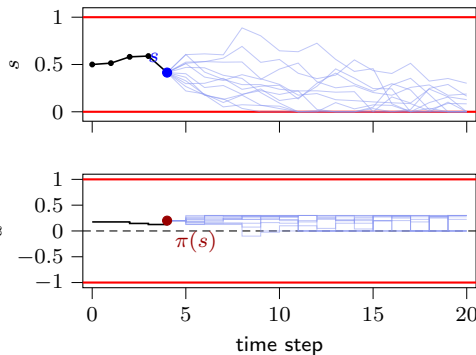
$$J(\pi) = \mathbb{E} \left[ \sum_{k=0}^{\infty} \gamma^k l(s_k, a_k) \mid \pi \right]$$

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

### ► Optimal policy: $\pi^*$ from

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

## Cost function (instant performance) $L(s, a) \in \mathbb{R}$



## Stochastic state transitions (real world)

$$s^+ \sim P(\cdot | s, a)$$

### ► Policy

$$a = \pi(s), \quad \text{or} \quad A \sim \pi(\cdot | s)$$

### ► Closed-loop performance

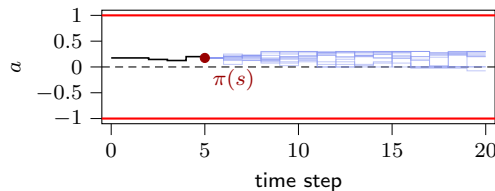
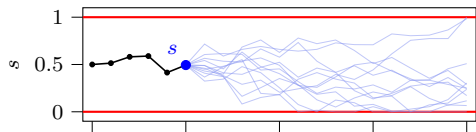
$$J(\pi) = \mathbb{E} \left[ \sum_{k=0}^{\infty} \gamma^k l(s_k, a_k) \mid \pi \right]$$

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

### ► Optimal policy: $\pi^*$ from

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

## Cost function (instant performance) $L(s, a) \in \mathbb{R}$



## Stochastic state transitions (real world)

$$s^+ \sim P(\cdot | s, a)$$

### ► Policy

$$a = \pi(s), \quad \text{or} \quad A \sim \pi(\cdot | s)$$

### ► Closed-loop performance

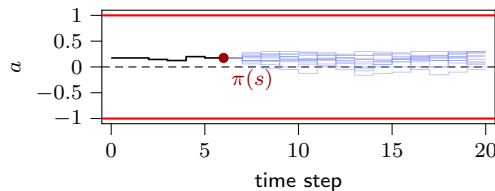
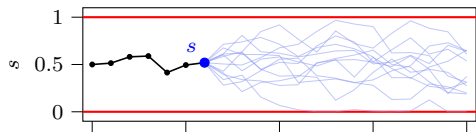
$$J(\pi) = \mathbb{E} \left[ \sum_{k=0}^{\infty} \gamma^k l(s_k, a_k) \mid \pi \right]$$

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

### ► Optimal policy: $\pi^*$ from

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

## Cost function (instant performance) $L(s, a) \in \mathbb{R}$





## Stochastic state transitions (real world)

$$s^+ \sim P(\cdot | s, a)$$

### ► Policy

$$a = \pi(s), \quad \text{or} \quad A \sim \pi(\cdot | s)$$

### ► Closed-loop performance

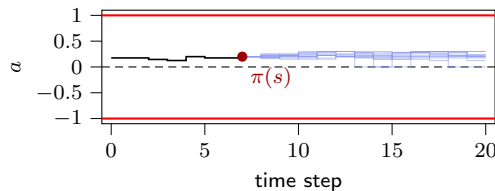
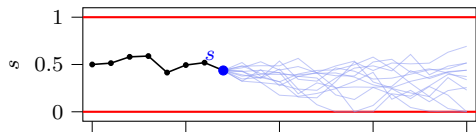
$$J(\pi) = \mathbb{E} \left[ \sum_{k=0}^{\infty} \gamma^k l(s_k, a_k) \mid \pi \right]$$

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

### ► Optimal policy: $\pi^*$ from

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

## Cost function (instant performance) $L(s, a) \in \mathbb{R}$



## Stochastic state transitions (real world)

$$s^+ \sim P(\cdot | s, a)$$

### ► Policy

$$a = \pi(s), \quad \text{or} \quad A \sim \pi(\cdot | s)$$

### ► Closed-loop performance

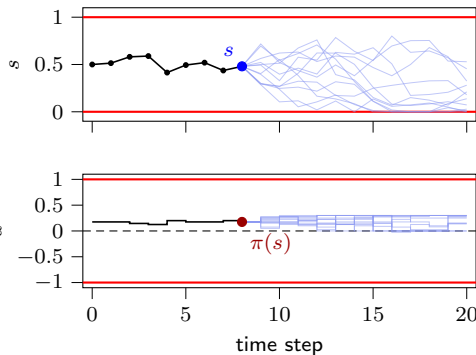
$$J(\pi) = \mathbb{E} \left[ \sum_{k=0}^{\infty} \gamma^k l(s_k, a_k) \mid \pi \right]$$

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

### ► Optimal policy: $\pi^*$ from

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

## Cost function (instant performance) $L(s, a) \in \mathbb{R}$



## Stochastic state transitions (real world)

$$s^+ \sim P(\cdot | s, a)$$

### ► Policy

$$a = \pi(s), \quad \text{or} \quad A \sim \pi(\cdot | s)$$

### ► Closed-loop performance

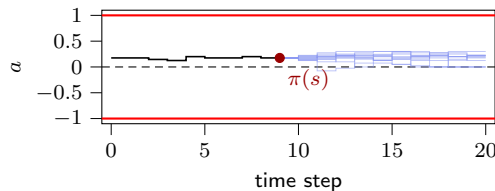
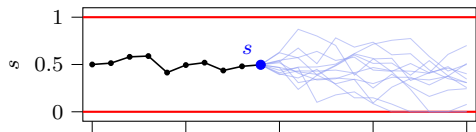
$$J(\pi) = \mathbb{E} \left[ \sum_{k=0}^{\infty} \gamma^k l(s_k, a_k) \mid \pi \right]$$

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

### ► Optimal policy: $\pi^*$ from

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

## Cost function (instant performance) $L(s, a) \in \mathbb{R}$



## Stochastic state transitions (real world)

$$s^+ \sim P(\cdot | s, a)$$

### ► Policy

$$a = \pi(s), \quad \text{or} \quad A \sim \pi(\cdot | s)$$

### ► Closed-loop performance

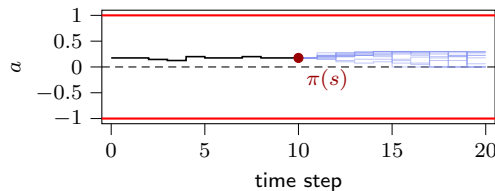
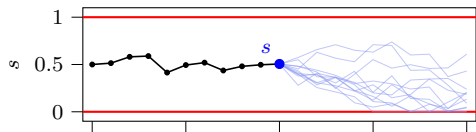
$$J(\pi) = \mathbb{E} \left[ \sum_{k=0}^{\infty} \gamma^k l(s_k, a_k) \mid \pi \right]$$

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

### ► Optimal policy: $\pi^*$ from

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

## Cost function (instant performance) $L(s, a) \in \mathbb{R}$



**Stochastic state transitions** (real world)

$$s^+ \sim P(\cdot | s, a)$$

► **Policy**

$$a = \pi(s), \quad \text{or} \quad A \sim \pi(\cdot | s)$$

► **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:**  $\pi^*$  from

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

**Cost function** (instant performance)  $L(s, a) \in \mathbb{R}$

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

$$l(s, a) = \begin{cases} L(s, a) & \text{if } h(s, a) \geq 0 \\ \infty & \text{if } h(s, a) < 0 \end{cases}$$

**Stochastic state transitions** (real world)

$$s^+ \sim P(\cdot | s, a)$$

► **Policy**

$$a = \pi(s), \quad \text{or} \quad A \sim \pi(\cdot | s)$$

► **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:**  $\pi^*$  from

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

**Cost function** (instant performance)  $L(s, a) \in \mathbb{R}$

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

$$l(s, a) = \begin{cases} L(s, a) & \text{if } h(s, a) \geq 0 \\ \infty & \text{if } h(s, a) < 0 \end{cases}$$

Solution of an MDP is described by “Bellman” equations (more in a bit), but solving them is very challenging

**Stochastic state transitions** (real world)

$$s^+ \sim P(\cdot | s, a)$$

► **Policy**

$$a = \pi(s), \quad \text{or} \quad A \sim \pi(\cdot | s)$$

► **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:**  $\pi^*$  from

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

**Cost function** (instant performance)  $L(s, a) \in \mathbb{R}$

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

$$l(s, a) = \begin{cases} L(s, a) & \text{if } h(s, a) \geq 0 \\ \infty & \text{if } h(s, a) < 0 \end{cases}$$

Solution of an MDP is described by “Bellman” equations (more in a bit), but solving them is very challenging

**A useful viewpoint:** when closed-loop performance ( $l$ ) matters and real-world is uncertain, MDPs are what we want to solve. MPC is a heuristic to deliver “good enough” solutions, at low computational cost.

# MPC as an approximation of the MDP problem

Steps:

**Original MDP problem**

$$V^*(s) = \min_{\pi} \mathbb{E} \left[ \sum_{k=0}^{\infty} \gamma^k l(S_k, A_k) \right],$$

with  $S_0 = s, S_{k+1} \sim P(\cdot \mid S_k, A_k), A_k \sim \pi(\cdot \mid S_k).$



# MPC as an approximation of the MDP problem

## Approximation 1: finite horizon

$$\min_{\pi} \mathbb{E} \left[ \gamma^N \bar{V}^{\pi}(S_N) + \sum_{k=0}^{N-1} \gamma^k l(S_k, A_k) \right],$$

with  $S_0 = s$ ,  $S_{k+1} \sim P(\cdot \mid S_k, A_k)$ ,  $A_k \sim \pi(\cdot \mid S_k)$ .

## Steps:

1.  $\infty$  horizon  $\rightarrow$  horizon  $N$  and  $\bar{V}$

# MPC as an approximation of the MDP problem

## Approximation 2: nominal predictions

$$\min_{\pi} \mathbb{E} \left[ \gamma^N \bar{V}^{\pi}(S_N) + \sum_{k=0}^{N-1} \gamma^k l(S_k, A_k) \right],$$

with  $S_0 = s$ ,  $x_0 = s$ ,  $x_{k+1} = f(x_k, u_k)$ ,  $A_k \sim \pi(\cdot | S_k)$ .

## Steps:

1.  $\infty$  horizon  $\rightarrow$  horizon  $N$  and  $\bar{V}$
2. stochastic model  $\rightarrow$  nominal dynamics

# MPC as an approximation of the MDP problem

## Approximation 3: open-loop action sequence

$$\min_{u_0, \dots, u_{N-1}} \gamma^N \bar{V}(x_N) + \sum_{k=0}^{N-1} \gamma^k l(x_k, u_k),$$

with  $x_0 = s, x_{k+1} = f(x_k, u_k)$ .

## Steps:

1.  $\infty$  horizon  $\rightarrow$  horizon  $N$  and  $\bar{V}$
2. stochastic model  $\rightarrow$  nominal dynamics
3. policy  $\pi \rightarrow$  open-loop sequence

# MPC as an approximation of the MDP problem

## Unconstrained finite-horizon OCP

$$\min_{u_0, \dots, u_{N-1}} \gamma^N \bar{V}(x_N) + \sum_{k=0}^{N-1} \gamma^k l(x_k, u_k),$$

with  $x_0 = s, x_{k+1} = f(x_k, u_k).$

## Steps:

1.  $\infty$  horizon  $\rightarrow$  horizon  $N$  and  $\bar{V}$
2. stochastic model  $\rightarrow$  nominal dynamics
3. policy  $\pi \rightarrow$  open-loop sequence

# MPC as an approximation of the MDP problem

**Deterministic MPC:**

$$\begin{aligned} \min_{x,u} \quad & \gamma^N \bar{V}(x_N) + \sum_{k=0}^{N-1} \gamma^k l(x_k, u_k) \\ \text{s.t} \quad & x_{k+1} = f(x_k, u_k) \\ & x_0 = s \end{aligned}$$

**Steps:**

1.  $\infty$  horizon  $\rightarrow$  horizon  $N$  and  $\bar{V}$
2. stochastic model  $\rightarrow$  nominal dynamics
3. policy  $\pi \rightarrow$  open-loop sequence

# MPC as an approximation of the MDP problem

## Deterministic MPC:

$$\begin{aligned} \min_{x,u} \quad & \gamma^N \bar{V}(x_N) + \sum_{k=0}^{N-1} \gamma^k l(x_k, u_k) \\ \text{s.t.} \quad & x_{k+1} = f(x_k, u_k) \\ & x_0 = s \end{aligned}$$

## Steps:

1.  $\infty$  horizon  $\rightarrow$  horizon  $N$  and  $\bar{V}$
2. stochastic model  $\rightarrow$  nominal dynamics
3. policy  $\pi \rightarrow$  open-loop sequence

MPC solves this problem online, applies the first action, observes the next state, and repeats.

## Deterministic MPC:

$$\begin{aligned} \min_{x,u} \quad & \gamma^N \bar{V}(x_N) + \sum_{k=0}^{N-1} \gamma^k l(x_k, u_k) \\ \text{s.t.} \quad & x_{k+1} = f(x_k, u_k) \\ & x_0 = s \end{aligned}$$

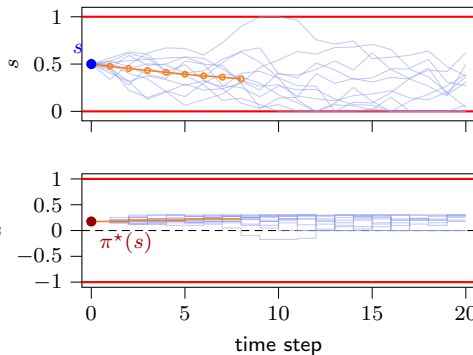
MPC defines a policy:

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

How does  $\pi^{\text{MPC}}$  relate to  $\pi^*$ ?

## MDP

$$\pi^* = \arg \min_{\pi} \mathbb{E} \left[ \sum_{k=0}^{\infty} \gamma^k l(S_k, A_k) \right]$$



## Deterministic MPC:

$$\begin{aligned} \min_{x,u} \quad & \gamma^N \bar{V}(x_N) + \sum_{k=0}^{N-1} \gamma^k l(x_k, u_k) \\ \text{s.t.} \quad & x_{k+1} = f(x_k, u_k) \\ & x_0 = s \end{aligned}$$

MPC defines a policy:

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

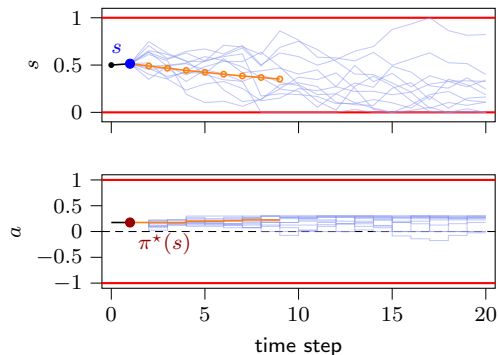
**How does  $\pi^{\text{MPC}}$  relate to  $\pi^*$ ?**

## No reason to match:

- ▶ MPC model approximates stochasticity, often deterministic
- ▶ Planning  $\neq$  policy for stochastic problems

## MDP

$$\pi^* = \arg \min_{\pi} \mathbb{E} \left[ \sum_{k=0}^{\infty} \gamma^k l(S_k, A_k) \right]$$





## Deterministic MPC:

$$\begin{aligned} \min_{x,u} \quad & \gamma^N \bar{V}(x_N) + \sum_{k=0}^{N-1} \gamma^k l(x_k, u_k) \\ \text{s.t.} \quad & x_{k+1} = f(x_k, u_k) \\ & x_0 = s \end{aligned}$$

MPC defines a policy:

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

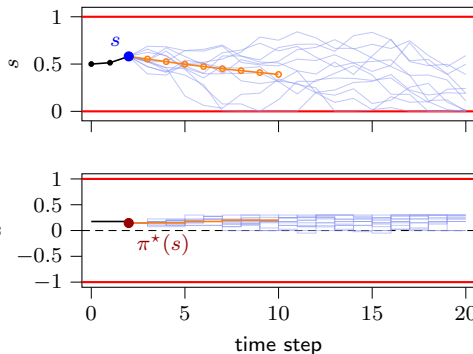
How does  $\pi^{\text{MPC}}$  relate to  $\pi^*$ ?

## No reason to match:

- ▶ MPC model approximates stochasticity, often deterministic
- ▶ Planning  $\neq$  policy for stochastic problems

## MDP

$$\pi^* = \arg \min_{\pi} \mathbb{E} \left[ \sum_{k=0}^{\infty} \gamma^k l(S_k, A_k) \right]$$



## Deterministic MPC:

$$\begin{aligned} \min_{x,u} \quad & \gamma^N \bar{V}(x_N) + \sum_{k=0}^{N-1} \gamma^k l(x_k, u_k) \\ \text{s.t.} \quad & x_{k+1} = f(x_k, u_k) \\ & x_0 = s \end{aligned}$$

MPC defines a policy:

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

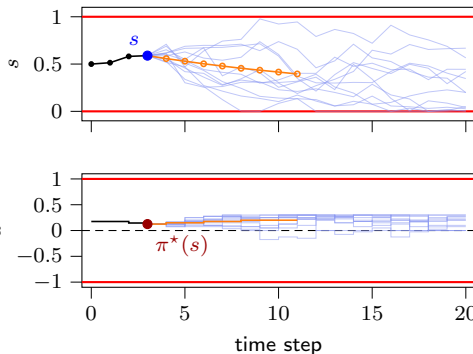
How does  $\pi^{\text{MPC}}$  relate to  $\pi^*$ ?

## No reason to match:

- ▶ MPC model approximates stochasticity, often deterministic
- ▶ Planning  $\neq$  policy for stochastic problems

## MDP

$$\pi^* = \arg \min_{\pi} \mathbb{E} \left[ \sum_{k=0}^{\infty} \gamma^k l(S_k, A_k) \right]$$



## Deterministic MPC:

$$\begin{aligned} \min_{x,u} \quad & \gamma^N \bar{V}(x_N) + \sum_{k=0}^{N-1} \gamma^k l(x_k, u_k) \\ \text{s.t.} \quad & x_{k+1} = f(x_k, u_k) \\ & x_0 = s \end{aligned}$$

MPC defines a policy:

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

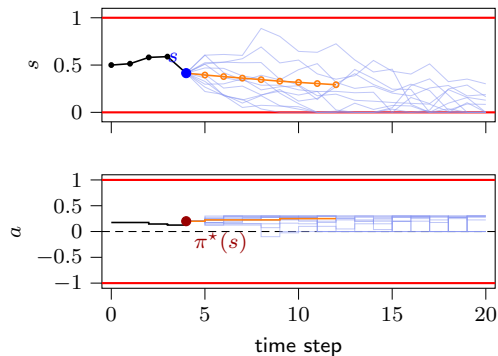
**How does  $\pi^{\text{MPC}}$  relate to  $\pi^*$ ?**

## No reason to match:

- ▶ MPC model approximates stochasticity, often deterministic
- ▶ Planning  $\neq$  policy for stochastic problems

## MDP

$$\pi^* = \arg \min_{\pi} \mathbb{E} \left[ \sum_{k=0}^{\infty} \gamma^k l(S_k, A_k) \right]$$



## Deterministic MPC:

$$\begin{aligned} \min_{x,u} \quad & \gamma^N \bar{V}(x_N) + \sum_{k=0}^{N-1} \gamma^k l(x_k, u_k) \\ \text{s.t.} \quad & x_{k+1} = f(x_k, u_k) \\ & x_0 = s \end{aligned}$$

MPC defines a policy:

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

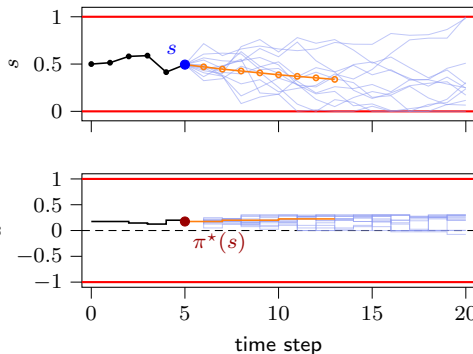
How does  $\pi^{\text{MPC}}$  relate to  $\pi^*$ ?

## No reason to match:

- ▶ MPC model approximates stochasticity, often deterministic
- ▶ Planning  $\neq$  policy for stochastic problems

## MDP

$$\pi^* = \arg \min_{\pi} \mathbb{E} \left[ \sum_{k=0}^{\infty} \gamma^k l(S_k, A_k) \right]$$



## Deterministic MPC:

$$\begin{aligned} \min_{x,u} \quad & \gamma^N \bar{V}(x_N) + \sum_{k=0}^{N-1} \gamma^k l(x_k, u_k) \\ \text{s.t.} \quad & x_{k+1} = f(x_k, u_k) \\ & x_0 = s \end{aligned}$$

MPC defines a policy:

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

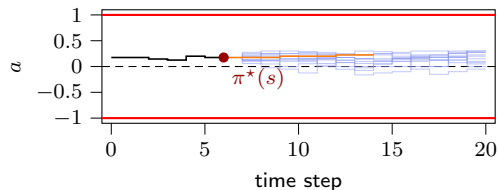
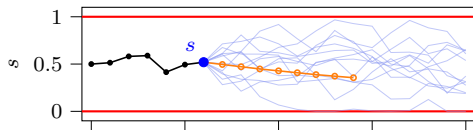
**How does  $\pi^{\text{MPC}}$  relate to  $\pi^*$ ?**

## No reason to match:

- ▶ MPC model approximates stochasticity, often deterministic
- ▶ Planning  $\neq$  policy for stochastic problems

## MDP

$$\pi^* = \arg \min_{\pi} \mathbb{E} \left[ \sum_{k=0}^{\infty} \gamma^k l(S_k, A_k) \right]$$



## Deterministic MPC:

$$\begin{aligned} \min_{x,u} \quad & \gamma^N \bar{V}(x_N) + \sum_{k=0}^{N-1} \gamma^k l(x_k, u_k) \\ \text{s.t.} \quad & x_{k+1} = f(x_k, u_k) \\ & x_0 = s \end{aligned}$$

MPC defines a policy:

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

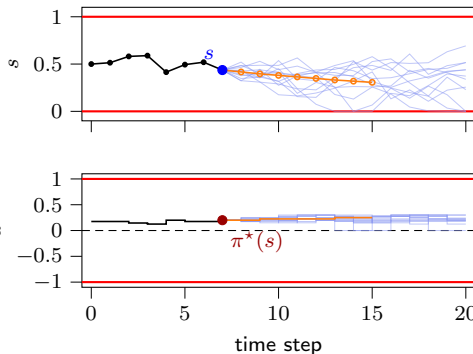
How does  $\pi^{\text{MPC}}$  relate to  $\pi^*$ ?

## No reason to match:

- ▶ MPC model approximates stochasticity, often deterministic
- ▶ Planning  $\neq$  policy for stochastic problems

## MDP

$$\pi^* = \arg \min_{\pi} \mathbb{E} \left[ \sum_{k=0}^{\infty} \gamma^k l(S_k, A_k) \right]$$



## Deterministic MPC:

$$\begin{aligned} \min_{x,u} \quad & \gamma^N \bar{V}(x_N) + \sum_{k=0}^{N-1} \gamma^k l(x_k, u_k) \\ \text{s.t.} \quad & x_{k+1} = f(x_k, u_k) \\ & x_0 = s \end{aligned}$$

MPC defines a policy:

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

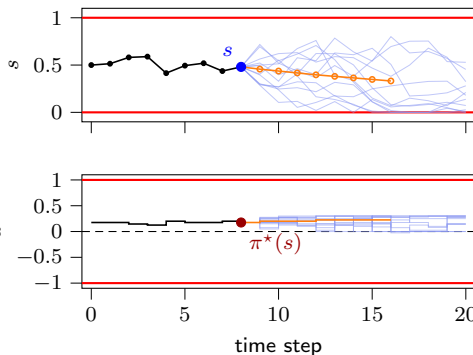
How does  $\pi^{\text{MPC}}$  relate to  $\pi^*$ ?

## No reason to match:

- ▶ MPC model approximates stochasticity, often deterministic
- ▶ Planning  $\neq$  policy for stochastic problems

## MDP

$$\pi^* = \arg \min_{\pi} \mathbb{E} \left[ \sum_{k=0}^{\infty} \gamma^k l(S_k, A_k) \right]$$



## Deterministic MPC:

$$\begin{aligned} \min_{x,u} \quad & \gamma^N \bar{V}(x_N) + \sum_{k=0}^{N-1} \gamma^k l(x_k, u_k) \\ \text{s.t.} \quad & x_{k+1} = f(x_k, u_k) \\ & x_0 = s \end{aligned}$$

MPC defines a policy:

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

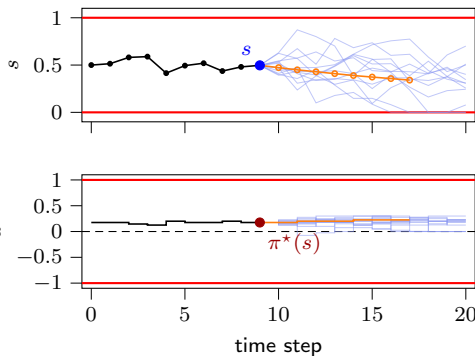
How does  $\pi^{\text{MPC}}$  relate to  $\pi^*$ ?

## No reason to match:

- ▶ MPC model approximates stochasticity, often deterministic
- ▶ Planning  $\neq$  policy for stochastic problems

## MDP

$$\pi^* = \arg \min_{\pi} \mathbb{E} \left[ \sum_{k=0}^{\infty} \gamma^k l(S_k, A_k) \right]$$





## Deterministic MPC:

$$\begin{aligned} \min_{x,u} \quad & \gamma^N \bar{V}(x_N) + \sum_{k=0}^{N-1} \gamma^k l(x_k, u_k) \\ \text{s.t.} \quad & x_{k+1} = f(x_k, u_k) \\ & x_0 = s \end{aligned}$$

MPC defines a policy:

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

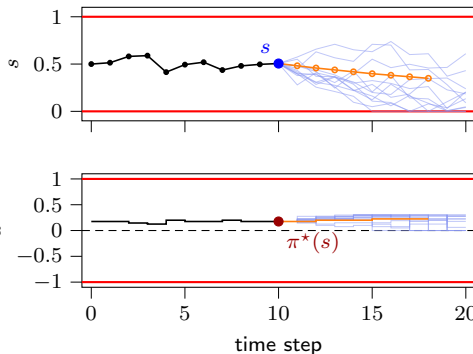
How does  $\pi^{\text{MPC}}$  relate to  $\pi^*$ ?

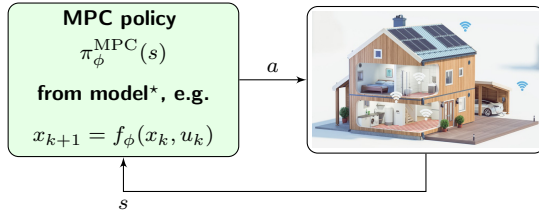
## No reason to match:

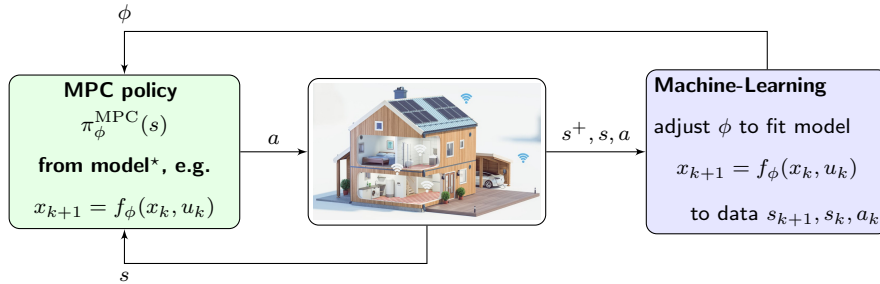
- ▶ MPC model approximates stochasticity, often deterministic
- ▶ Planning  $\neq$  policy for stochastic problems

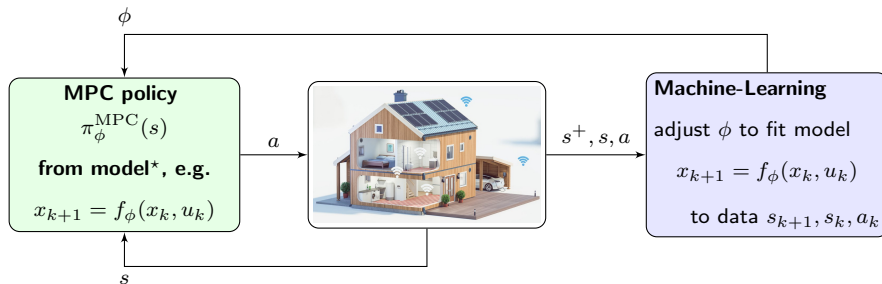
## MDP

$$\pi^* = \arg \min_{\pi} \mathbb{E} \left[ \sum_{k=0}^{\infty} \gamma^k l(S_k, A_k) \right]$$



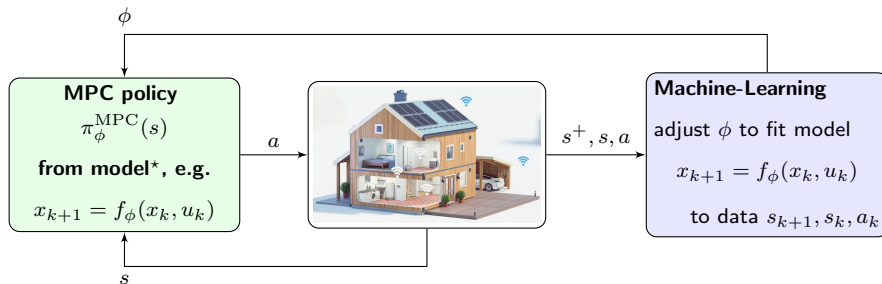






“Machine-Learning” in-the-loop  $f_{\phi}$  from

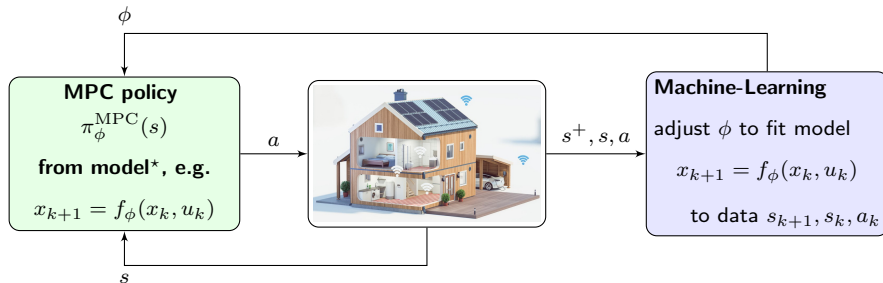
- **Physics-based:** first principles + SYSID
- **Neural Network:** DNN, LSTM, TFT, ...
- **Statistical:** GP, Kernel ...



“Machine-Learning” in-the-loop  $f_{\phi}$  from

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

*\* any consistent predictor: extends to input-output, multi-step, Koopman, etc...*



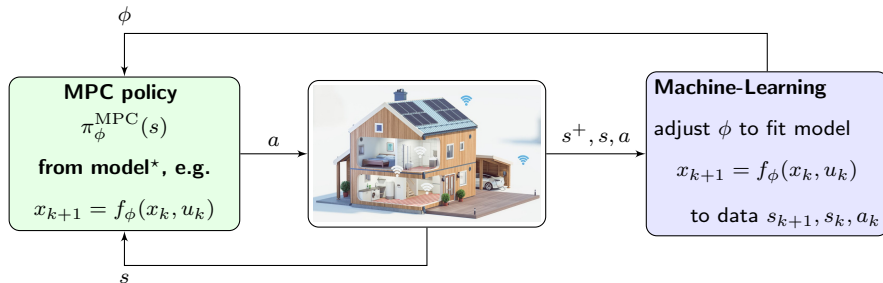
“Machine-Learning” in-the-loop  $f_{\phi}$  from

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

*\* any consistent predictor: extends to input-output, multi-step, Koopman, etc...*

## Limitations of this paradigm

- Open-loop prediction accuracy  $\xrightarrow{\text{hopefully}}$  Closed-loop performance
- Ignores that MPC is a policy



“Machine-Learning” in-the-loop  $f_{\phi}$  from

- Physics-based: first principles + SYSID
- Neural Network: DNN, LSTM, TFT, ...
- Statistical: GP, Kernel ...

**RLMPC focuses on “breaking” these limitations**  
**RL plays a key role**

*\* any consistent predictor: extends to input-output, multi-step, Koopman, etc...*

**Limitations of this paradigm**

- Open-loop prediction accuracy  $\xrightarrow{\text{hopefully}}$  Closed-loop performance
- Ignores that MPC is a policy

Background: MPC, MDPs, and Model Fitting

MPC as a Model of the MDP Solution

Reinforcement Learning Tunes the MPC



### MDP

$$\pi^* = \arg \min_{\pi} \mathbb{E} \left[ \sum_{k=0}^{\infty} \gamma^k l(S_k, A_k) \right]$$

### Bellman equations

$$V^*(s) = \min_a l(s, a) + \gamma \mathbb{E} [V^*(S^+) \mid s, a]$$

minimizer  $a$  for  $s$  gives the optimal policy

### Bellman equations

$$V^*(s) = \min_a l(s, a) + \gamma \mathbb{E} [V^*(S^+) \mid s, a]$$

minimizer  $a$  for  $s$  gives the optimal policy

### MDP

$$\pi^* = \arg \min_{\pi} \mathbb{E} \left[ \sum_{k=0}^{\infty} \gamma^k l(S_k, A_k) \right]$$

### A more expressive form

$$Q^*(s, a) = \mathbb{E} \left[ \sum_{k=0}^{\infty} \gamma^k l(S_k, A_k) \mid \begin{array}{l} S_0 = s \\ A_0 = a \\ A_{k>0} = \pi^*(S_k) \end{array} \right]$$

**Remarks:** the action-value function  $Q^*$  is

- ▶ the optimal “cost-to-go” from state  $s$  if we force the first action to be  $a$ , and follow the optimal policy after that

## Bellman equations

### MDP

$$\pi^* = \arg \min_{\pi} \mathbb{E} \left[ \sum_{k=0}^{\infty} \gamma^k l(S_k, A_k) \right]$$

$$V^*(s) = \min_a l(s, a) + \gamma \mathbb{E} [V^*(S^+) | s, a]$$

minimizer  $a$  for  $s$  gives the optimal policy

### A more expressive form

$$Q^*(s, a) = \mathbb{E} \left[ \sum_{k=0}^{\infty} \gamma^k l(S_k, A_k) \middle| \begin{array}{l} S_0 = s \\ A_0 = a \\ A_{k>0} = \pi^*(S_k) \end{array} \right]$$

$$Q^*(s, a) = l(s, a) + \gamma \mathbb{E} [V^*(S^+) | s, a]$$

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

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

**Remarks:** the action-value function  $Q^*$  is

- ▶ the optimal “cost-to-go” from state  $s$  if we force the first action to be  $a$ , and follow the optimal policy after that

## Bellman equations

### MDP

$$\pi^* = \arg \min_{\pi} \mathbb{E} \left[ \sum_{k=0}^{\infty} \gamma^k l(S_k, A_k) \right]$$

$$V^*(s) = \min_a l(s, a) + \gamma \mathbb{E} [V^*(S^+) | s, a]$$

minimizer  $a$  for  $s$  gives the optimal policy

### A more expressive form

$$Q^*(s, a) = \mathbb{E} \left[ \sum_{k=0}^{\infty} \gamma^k l(S_k, A_k) \middle| \begin{array}{l} S_0 = s \\ A_0 = a \\ A_{k>0} = \pi^*(S_k) \end{array} \right]$$

$$Q^*(s, a) = l(s, a) + \gamma \mathbb{E} [V^*(S^+) | s, a]$$

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

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

**Remarks:** the action-value function  $Q^*$  is

- ▶ the optimal “cost-to-go” from state  $s$  if we force the first action to be  $a$ , and follow the optimal policy after that
- ▶ the “whole” MDP solution

## Bellman equations

### MDP

$$\pi^* = \arg \min_{\pi} \mathbb{E} \left[ \sum_{k=0}^{\infty} \gamma^k l(S_k, A_k) \right]$$

$$V^*(s) = \min_a l(s, a) + \gamma \mathbb{E} [V^*(S^+) | s, a]$$

minimizer  $a$  for  $s$  gives the optimal policy

### A more expressive form

$$Q^*(s, a) = \mathbb{E} \left[ \sum_{k=0}^{\infty} \gamma^k l(S_k, A_k) \middle| \begin{array}{l} S_0 = s \\ A_0 = a \\ A_{k>0} = \pi^*(S_k) \end{array} \right]$$

$$Q^*(s, a) = l(s, a) + \gamma \mathbb{E} [V^*(S^+) | s, a]$$

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

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

**Remarks:** the action-value function  $Q^*$  is

- ▶ the optimal “cost-to-go” from state  $s$  if we force the first action to be  $a$ , and follow the optimal policy after that
- ▶ the “whole” MDP solution
- ▶ very expensive to compute if  $s, a$  are not “trivially small-dimensional”

## Bellman equations

### MDP

$$\pi^* = \arg \min_{\pi} \mathbb{E} \left[ \sum_{k=0}^{\infty} \gamma^k l(S_k, A_k) \right]$$

$$V^*(s) = \min_a l(s, a) + \gamma \mathbb{E} [V^*(S^+) | s, a]$$

minimizer  $a$  for  $s$  gives the optimal policy

### A more expressive form

$$Q^*(s, a) = \mathbb{E} \left[ \sum_{k=0}^{\infty} \gamma^k l(S_k, A_k) \middle| \begin{array}{l} S_0 = s \\ A_0 = a \\ A_{k>0} = \pi^*(S_k) \end{array} \right]$$

$$Q^*(s, a) = l(s, a) + \gamma \mathbb{E} [V^*(S^+) | s, a]$$

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

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

**Remarks:** the action-value function  $Q^*$  is

- ▶ the optimal “cost-to-go” from state  $s$  if we force the first action to be  $a$ , and follow the optimal policy after that
- ▶ the “whole” MDP solution
- ▶ very expensive to compute if  $s, a$  are not “trivially small-dimensional”
- ▶ a “cornerstone” in Reinforcement Learning

## MPC as a model of the MDP solution?

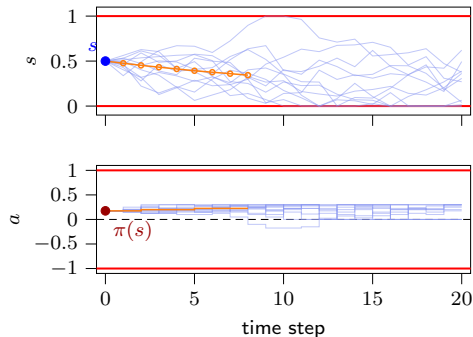
**MDP** optimal policy

$$\pi^* = \arg \min_{\pi} \mathbb{E} \left[ \sum_{k=0}^{\infty} \gamma^k l(S_k, A_k) \right]$$

is “solved” by  $Q^*(s, a)$

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

$$\begin{aligned} \min_{x, u} \quad & \bar{V}(x_N) + \sum_{k=0}^{N-1} \gamma^k l(x_k, u_k) \\ \text{s.t} \quad & x_{k+1} = f(x_k, u_k), \quad x_0 = s \\ & h(x_k, u_k) \geq 0 \end{aligned}$$



## MPC as a model of the MDP solution?

**MDP** optimal policy

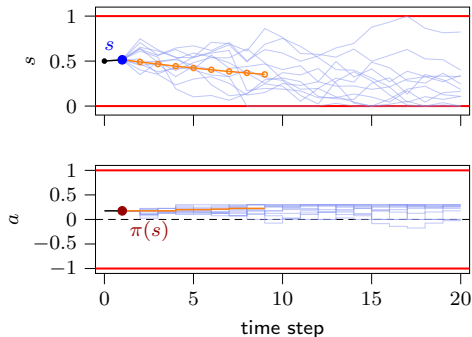
$$\pi^* = \arg \min_{\pi} \mathbb{E} \left[ \sum_{k=0}^{\infty} \gamma^k l(S_k, A_k) \right]$$

is “solved” by  $Q^*(s, a)$

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

$$\begin{aligned} \min_{x, u} \quad & \bar{V}(x_N) + \sum_{k=0}^{N-1} \gamma^k l(x_k, u_k) \\ \text{s.t} \quad & x_{k+1} = f(x_k, u_k), \quad x_0 = s \\ & h(x_k, u_k) \geq 0 \end{aligned}$$

- ▶ Solving MDP  $\equiv$  build  $Q^*$  model
- ▶ MPC can model a value function...
- ▶ Can it model an action-value function?





## MPC as a model of the MDP solution?

**MDP** optimal policy

$$\pi^* = \arg \min_{\pi} \mathbb{E} \left[ \sum_{k=0}^{\infty} \gamma^k l(S_k, A_k) \right]$$

is “solved” by  $Q^*(s, a)$

$$V^{\text{MPC}}(s) := \min_{x, u} \bar{V}(x_N) + \sum_{k=0}^{N-1} \gamma^k l(x_k, u_k)$$

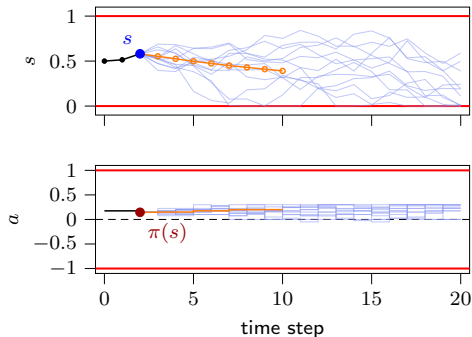
s.t  $x_{k+1} = f(x_k, u_k), \quad x_0 = s$   
 $h(x_k, u_k) \geq 0$

- ▶ Solving MDP  $\equiv$  build  $Q^*$  model
- ▶ MPC can model a value function...
- ▶ Can it model an action-value function?

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

$$\min_{x, u} \bar{V}(x_N) + \sum_{k=0}^{N-1} \gamma^k l(x_k, u_k)$$

s.t  $x_{k+1} = f(x_k, u_k), \quad x_0 = s$   
 $h(x_k, u_k) \geq 0$



## MDP optimal policy

$$\pi^* = \arg \min_{\pi} \mathbb{E} \left[ \sum_{k=0}^{\infty} \gamma^k l(S_k, A_k) \right]$$

is “solved” by  $Q^*(s, a)$

$$V^{\text{MPC}}(s) := \min_{x, u} \bar{V}(x_N) + \sum_{k=0}^{N-1} \gamma^k l(x_k, u_k)$$

s.t.  $x_{k+1} = f(x_k, u_k), \quad x_0 = s$   
 $h(x_k, u_k) \geq 0$

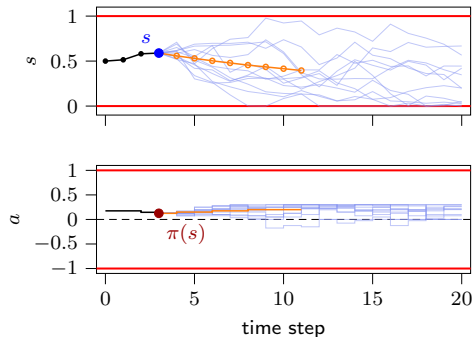
$$Q^{\text{MPC}}(s, a) := \min_{x, u} \bar{V}(x_N) + \sum_{k=0}^{N-1} \gamma^k l(x_k, u_k)$$

s.t.  $x_{k+1} = f(x_k, u_k)$   
 $h(x_k, u_k) \geq 0$   
 $x_0 = s, \quad u_0 = a$

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

$$\min_{x, u} \bar{V}(x_N) + \sum_{k=0}^{N-1} \gamma^k l(x_k, u_k)$$

s.t.  $x_{k+1} = f(x_k, u_k), \quad x_0 = s$   
 $h(x_k, u_k) \geq 0$



**MDP optimal policy**

$$\pi^* = \arg \min_{\pi} \mathbb{E} \left[ \sum_{k=0}^{\infty} \gamma^k l(S_k, A_k) \right]$$

is “solved” by  $Q^*(s, a)$

$$V^{\text{MPC}}(s) := \min_{x, u} \bar{V}(x_N) + \sum_{k=0}^{N-1} \gamma^k l(x_k, u_k)$$

s.t.  $x_{k+1} = f(x_k, u_k), \quad x_0 = s$   
 $h(x_k, u_k) \geq 0$

$$Q^{\text{MPC}}(s, a) := \min_{x, u} \bar{V}(x_N) + \sum_{k=0}^{N-1} \gamma^k l(x_k, u_k)$$

s.t.  $x_{k+1} = f(x_k, u_k)$   
 $h(x_k, u_k) \geq 0$   
 $x_0 = s, \quad u_0 = a$

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

$$\min_{x, u} \bar{V}(x_N) + \sum_{k=0}^{N-1} \gamma^k l(x_k, u_k)$$

s.t.  $x_{k+1} = f(x_k, u_k), \quad x_0 = s$   
 $h(x_k, u_k) \geq 0$

**MPC is consistent:**

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

**MDP optimal policy**

$$\pi^* = \arg \min_{\pi} \mathbb{E} \left[ \sum_{k=0}^{\infty} \gamma^k l(S_k, A_k) \right]$$

is “solved” by  $Q^*(s, a)$

$$V^{\text{MPC}}(s) := \min_{x, u} \bar{V}(x_N) + \sum_{k=0}^{N-1} \gamma^k l(x_k, u_k)$$

s.t.  $x_{k+1} = f(x_k, u_k), \quad x_0 = s$   
 $h(x_k, u_k) \geq 0$

$$Q^{\text{MPC}}(s, a) := \min_{x, u} \bar{V}(x_N) + \sum_{k=0}^{N-1} \gamma^k l(x_k, u_k)$$

s.t.  $x_{k+1} = f(x_k, u_k)$   
 $h(x_k, u_k) \geq 0$   
 $x_0 = s, \quad u_0 = a$

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

$$\min_{x, u} \quad \bar{V}(x_N) + \sum_{k=0}^{N-1} \gamma^k l(x_k, u_k)$$

s.t.  $x_{k+1} = f(x_k, u_k), \quad x_0 = s$   
 $h(x_k, u_k) \geq 0$

**MPC is consistent:**

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

**MPC is a complete MDP model if:**

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

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

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

**MDP optimal policy**

$$\pi^* = \arg \min_{\pi} \mathbb{E} \left[ \sum_{k=0}^{\infty} \gamma^k l(S_k, A_k) \right]$$

is “solved” by  $Q^*(s, a)$

$$V^{\text{MPC}}(s) := \min_{x, u} \bar{V}(x_N) + \sum_{k=0}^{N-1} \gamma^k l(x_k, u_k)$$

$$\begin{aligned} \text{s.t. } & x_{k+1} = f(x_k, u_k), \quad x_0 = s \\ & h(x_k, u_k) \geq 0 \end{aligned}$$

$$Q^{\text{MPC}}(s, a) := \min_{x, u} \bar{V}(x_N) + \sum_{k=0}^{N-1} \gamma^k l(x_k, u_k)$$

$$\begin{aligned} \text{s.t. } & x_{k+1} = f(x_k, u_k) \\ & h(x_k, u_k) \geq 0 \\ & x_0 = s, \quad u_0 = a \end{aligned}$$

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

$$\begin{aligned} \min_{x, u} \quad & \bar{V}(x_N) + \sum_{k=0}^{N-1} \gamma^k l(x_k, u_k) \\ \text{s.t. } \quad & x_{k+1} = f(x_k, u_k), \quad x_0 = s \\ & h(x_k, u_k) \geq 0 \end{aligned}$$

**Punchline:** see MPC as a  $Q^*$  model i.e. what we want to achieve is

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

or at least

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

$\neq$  fitting  $f$  to the real world

## How to model $Q^*$ with MPC?

---

**Classical** view...

**MPC:** at current state  $s$  solve

$$\min_{x,u} \quad \bar{V}(x_N) + \sum_{k=0}^{N-1} l(x_k, u_k)$$

$$\text{s.t.} \quad x_{k+1} = f_{\phi}(x_k, u_k)$$

$$h(x_k, u_k) \geq 0$$

$$x_0 = s$$

gives policy  $\pi_{\phi}^{\text{MPC}}(s) = u_0^*$

Find  $\phi$  such that prediction “fits”  
the data

## Shift 1: focus on performance instead of fitting

- ▶ **from:**  $f_\phi$  is a model for the system dynamics
- ▶ **to:** MPC is a model of the MDP solution

**Classical** view...

**MPC:** at current state  $s$  solve

$$\min_{x,u} \quad \bar{V}(x_N) + \sum_{k=0}^{N-1} l(x_k, u_k)$$

$$\text{s.t.} \quad x_{k+1} = f_\phi(x_k, u_k)$$

$$h(x_k, u_k) \geq 0$$

$$x_0 = s$$

gives policy  $\pi_\phi^{\text{MPC}}(s) = u_0^*$

Find  $\phi$  such that prediction “fits”  
the data



## Shift 1: focus on performance instead of fitting

- ▶ **from:**  $f_\phi$  is a model for the system dynamics
- ▶ **to:** MPC is a model of the MDP solution

**Classical** view...

**MPC:** at current state  $s$  solve

$$\min_{x,u} \quad \bar{V}(x_N) + \sum_{k=0}^{N-1} l(x_k, u_k)$$

$$\text{s.t.} \quad x_{k+1} = f_\phi(x_k, u_k)$$

$$h(x_k, u_k) \geq 0$$

$$x_0 = s$$

gives policy  $\pi_\phi^{\text{MPC}}(s) = u_0^*$

**Shift to...**

**Find  $\phi$  that “fits MPC to optimality” according to the data**  
( $Q^{\text{MPC}} \rightarrow Q^*$  or at least  $\pi^{\text{MPC}} \rightarrow \pi^*$ )

**Find  $\phi$  such that prediction “fits” the data**

## Shift 1: focus on performance instead of fitting

- ▶ **from:**  $f_\phi$  is a model for the system dynamics
- ▶ **to:** MPC is a model of the MDP solution

**Classical** view...

**MPC:** at current state  $s$  solve

$$\min_{x,u} \quad \bar{V}(x_N) + \sum_{k=0}^{N-1} l(x_k, u_k)$$

$$\text{s.t.} \quad x_{k+1} = f_\phi(x_k, u_k)$$

$$h(x_k, u_k) \geq 0$$

$$x_0 = s$$

gives policy  $\pi_\phi^{\text{MPC}}(s) = u_0^*$

**Find  $\phi$  such that prediction “fits” the data**

**Shift to...**

**Find  $\phi$  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!\*

\* Theorem 1 in Anand et al., "Optimality conditions for model predictive control: Rethinking predictive model design." Automatica (2026).

## Shift 1: focus on performance instead of fitting

- ▶ **from:**  $f_\phi$  is a model for the system dynamics
- ▶ **to:** MPC is a model of the MDP solution

**Classical** view...

**MPC:** at current state  $s$  solve

$$\min_{x,u} \quad \bar{V}(x_N) + \sum_{k=0}^{N-1} l(x_k, u_k)$$

$$\text{s.t.} \quad x_{k+1} = f_\phi(x_k, u_k)$$

$$h(x_k, u_k) \geq 0$$

$$x_0 = s$$

gives policy  $\pi_\phi^{\text{MPC}}(s) = u_0^*$

**Find  $\phi$  such that prediction “fits” the data**

**Shift to...**

**Find  $\phi$  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 this places “high demands” on  $f_\phi$

Can we do better?

\* Theorem 1 in Anand et al., “Optimality conditions for model predictive control: Rethinking predictive model design.” Automatica (2026).

## Shift 1: focus on performance instead of fitting

- ▶ **from:**  $f_\phi$  is a model for the system dynamics
- ▶ **to:** MPC is a model of the MDP solution

**Classical** view...

**MPC:** at current state  $s$  solve

$$\min_{x,u} \quad \bar{V}(x_N) + \sum_{k=0}^{N-1} l(x_k, u_k)$$

$$\text{s.t.} \quad x_{k+1} = f_\phi(x_k, u_k)$$

$$h(x_k, u_k) \geq 0$$

$$x_0 = s$$

gives policy  $\pi_\phi^{\text{MPC}}(s) = u_0^*$

**Find  $\phi$  such that prediction “fits” the data**

**Shift to...**

**Find  $\phi$  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 this places “high demands” on  $f_\phi$

Can we do better? **Yes!**

\* Theorem 1 in Anand et al., “Optimality conditions for model predictive control: Rethinking predictive model design.” Automatica (2026).

## Shift 1: focus on performance instead of fitting

- ▶ **from:**  $f_\phi$  is a model for the system dynamics
- ▶ **to:** MPC is a model of the MDP solution

**Classical** view...

**MPC:** at current state  $s$  solve

$$\min_{x,u} \quad \bar{V}(x_N) + \sum_{k=0}^{N-1} l(x_k, u_k)$$

$$\text{s.t.} \quad x_{k+1} = f_\phi(x_k, u_k)$$

$$h(x_k, u_k) \geq 0$$

$$x_0 = s$$

gives policy  $\pi_\phi^{\text{MPC}}(s) = u_0^*$

Find  $\phi$  such that prediction “fits”  
the data

**Shift to...**

Find  $\phi$  that “fits MPC to optimality” according to the data  
( $Q^{\text{MPC}} \rightarrow Q^*$  or at least  $\pi^{\text{MPC}} \rightarrow \pi^*$ )

## Shift 2: “holistic” parametrization

$$\min_{x,u} \quad \bar{V}_\phi(x_N) + \sum_{k=0}^{N-1} l_\phi(x_k, u_k)$$

$$\text{s.t.} \quad x_{k+1} = f_\phi(x_k, u_k), \quad x_0 = s$$

$$h_\phi(x_k, u_k) \geq 0$$

real-world performance still assessed via  $l$  and  $h$

## Shift 1: focus on performance instead of fitting

- ▶ **from:**  $f_\phi$  is a model for the system dynamics
- ▶ **to:** MPC is a model of the MDP solution

**Classical** view...

**MPC:** at current state  $s$  solve

$$\min_{x,u} \quad \bar{V}(x_N) + \sum_{k=0}^{N-1} l(x_k, u_k)$$

$$\text{s.t.} \quad x_{k+1} = f_\phi(x_k, u_k)$$

$$h(x_k, u_k) \geq 0$$

$$x_0 = s$$

gives policy  $\pi_\phi^{\text{MPC}}(s) = u_0^*$

Find  $\phi$  such that prediction “fits”  
the data

**Shift to...**

Find  $\phi$  that “fits MPC to optimality” according to the data  
( $Q^{\text{MPC}} \rightarrow Q^*$  or at least  $\pi^{\text{MPC}} \rightarrow \pi^*$ )

## Shift 2: “holistic” parametrization

$$\min_{x,u} \quad \bar{V}_\phi(x_N) + \sum_{k=0}^{N-1} l_\phi(x_k, u_k)$$

$$\text{s.t.} \quad x_{k+1} = f_\phi(x_k, u_k), \quad x_0 = s$$

$$h_\phi(x_k, u_k) \geq 0$$

real-world performance still assessed via  $l$  and  $h$

*Intuitively ok... but does that make sense?*

### Full MPC parametrization:

$$\min_{x,u} \quad \bar{V}_{\phi}(x_N) + \sum_{k=0}^{N-1} l_{\phi}(x_k, u_k)$$

$$\text{s.t.} \quad x_{k+1} = f_{\phi}(x_k, u_k)$$

$$h_{\phi}(x_k, u_k) \geq 0$$

$$x_0 = s$$

gives policy  $\pi_{\phi}^{\text{MPC}}(s) = u_0^*$

### Full MPC parametrization:

$$\min_{x,u} \quad \bar{V}_{\phi}(x_N) + \sum_{k=0}^{N-1} l_{\phi}(x_k, u_k)$$

$$\text{s.t.} \quad x_{k+1} = f_{\phi}(x_k, u_k)$$

$$h_{\phi}(x_k, u_k) \geq 0$$

$$x_0 = s$$

gives policy  $\pi_{\phi}^{\text{MPC}}(s) = u_0^*$

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

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

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

even if the model  $f_{\phi}$  *cannot* describe the real system accurately

\* Theorem 1 from Gros et al., “Data-driven economic NMPC using reinforcement learning,” IEEE Transactions on Automatic Control 65.2 (2019): 636–648.



### Full MPC parametrization:

$$\min_{x,u} \quad \bar{V}_{\phi}(x_N) + \sum_{k=0}^{N-1} l_{\phi}(x_k, u_k)$$

$$\text{s.t.} \quad x_{k+1} = f_{\phi}(x_k, u_k)$$

$$h_{\phi}(x_k, u_k) \geq 0$$

$$x_0 = s$$

gives policy  $\pi_{\phi}^{\text{MPC}}(s) = u_0^*$

**Theorem<sup>\*</sup>**: under some technical conditions and for a “rich” parametrization of the MPC, there is a  $\phi$  such that

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

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

even if the model  $f_{\phi}$  *cannot* describe the real system accurately

### Remarks

- **Generic**: works for robust MPC, stochastic MPC, economic MPC, Multi-Step PC, etc.

<sup>\*</sup> Theorem 1 from Gros et al., “Data-driven economic NMPC using reinforcement learning,” IEEE Transactions on Automatic Control 65.2 (2019): 636–648.

### Full MPC parametrization:

$$\min_{x,u} \quad \bar{V}_{\phi}(x_N) + \sum_{k=0}^{N-1} l_{\phi}(x_k, u_k)$$

$$\text{s.t.} \quad x_{k+1} = f_{\phi}(x_k, u_k)$$

$$h_{\phi}(x_k, u_k) \geq 0$$

$$x_0 = s$$

gives policy  $\pi_{\phi}^{\text{MPC}}(s) = u_0^*$

**Theorem<sup>\*</sup>**: under some technical conditions and for a “rich” parametrization of the MPC, there is a  $\phi$  such that

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

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

even if the model  $f_{\phi}$  *cannot* describe the real system accurately

### Remarks

- ▶ **Generic**: works for robust MPC, stochastic MPC, economic MPC, Multi-Step PC, etc.
- ▶ **Sanity check**: technical conditions are mild but forbid  $f_{\phi}$  to lead to non-finite  $V^*$  under the resulting closed-loop policy

<sup>\*</sup> Theorem 1 from Gros et al., “Data-driven economic NMPC using reinforcement learning,” IEEE Transactions on Automatic Control 65.2 (2019): 636–648.

### Full MPC parametrization:

$$\min_{x,u} \quad \bar{V}_{\phi}(x_N) + \sum_{k=0}^{N-1} l_{\phi}(x_k, u_k)$$

$$\text{s.t.} \quad x_{k+1} = f_{\phi}(x_k, u_k)$$

$$h_{\phi}(x_k, u_k) \geq 0$$

$$x_0 = s$$

gives policy  $\pi_{\phi}^{\text{MPC}}(s) = u_0^*$

**Theorem<sup>\*</sup>**: under some technical conditions and for a “rich” parametrization of the MPC, there is a  $\phi$  such that

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

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

even if the model  $f_{\phi}$  *cannot* describe the real system accurately

### Remarks

- ▶ **Generic**: works for robust MPC, stochastic MPC, economic MPC, Multi-Step PC, etc.
- ▶ **Sanity check**: technical conditions are mild but forbid  $f_{\phi}$  to lead to non-finite  $V^*$  under the resulting closed-loop policy

<sup>\*</sup> Theorem 1 from Gros et al., “Data-driven economic NMPC using reinforcement learning,” IEEE Transactions on Automatic Control 65.2 (2019): 636–648.

### Full MPC parametrization:

$$\min_{x,u} \quad \bar{V}_{\phi}(x_N) + \sum_{k=0}^{N-1} l_{\phi}(x_k, u_k)$$

$$\text{s.t.} \quad x_{k+1} = f_{\phi}(x_k, u_k)$$

$$h_{\phi}(x_k, u_k) \geq 0$$

$$x_0 = s$$

gives policy  $\pi_{\phi}^{\text{MPC}}(s) = u_0^*$

**Full MPC parametrization enables modeling  $Q^*$ ,  $\pi^*$  without “torturing”  $f_{\phi}$**

**Theorem<sup>\*</sup>:** under some technical conditions and for a “rich” parametrization of the MPC, there is a  $\phi$  such that

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

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

even if the model  $f_{\phi}$  *cannot* describe the real system accurately

### Remarks

- ▶ **Generic:** works for robust MPC, stochastic MPC, economic MPC, Multi-Step PC, etc.
- ▶ **Sanity check:** technical conditions are mild but forbid  $f_{\phi}$  to lead to non-finite  $V^*$  under the resulting closed-loop policy

<sup>\*</sup> Theorem 1 from Gros et al., “Data-driven economic NMPC using reinforcement learning,” IEEE Transactions on Automatic Control 65.2 (2019): 636–648.

Background: MPC, MDPs, and Model Fitting

MPC as a Model of the MDP Solution

Reinforcement Learning Tunes the MPC

**Policy**  $\pi_{\phi}^{\text{MPC}}(s) = u_0^*$  from

$$\begin{aligned} \min_{x,u} \quad & \bar{V}_{\phi}(x_N) + \sum_{k=0}^{N-1} l_{\phi}(x_k, u_k) \\ \text{s.t.} \quad & x_{k+1} = f_{\phi}(x_k, u_k) \\ & h_{\phi}(x_k, u_k) \geq 0, \quad x_0 = s \end{aligned}$$

**We want**

$$\min_{\phi} J\left(\pi_{\phi}^{\text{MPC}}\right)$$

where  $J(\cdot)$  is the closed-loop cost in the  
real-world

**What to do?**

- ▶  $\min_{\phi} J\left(\pi_{\phi}^{\text{MPC}}\right)$  using data
- ▶ Complex “chain”:  $\phi \rightarrow \pi_{\phi}^{\text{MPC}}(s) \rightarrow \text{real world} \rightarrow J\left(\pi_{\phi}^{\text{MPC}}\right)$

**Policy**  $\pi_{\phi}^{\text{MPC}}(s) = u_0^*$  from

$$\begin{aligned} \min_{x,u} \quad & \bar{V}_{\phi}(x_N) + \sum_{k=0}^{N-1} l_{\phi}(x_k, u_k) \\ \text{s.t.} \quad & x_{k+1} = f_{\phi}(x_k, u_k) \\ & h_{\phi}(x_k, u_k) \geq 0, \quad x_0 = s \end{aligned}$$

**We want**

$$\min_{\phi} J(\pi_{\phi}^{\text{MPC}})$$

where  $J(\cdot)$  is the closed-loop cost in the  
real-world

**What to do?**

- ▶  $\min_{\phi} J(\pi_{\phi}^{\text{MPC}})$  using data
- ▶ Complex “chain”:  $\phi \rightarrow \pi_{\phi}^{\text{MPC}}(s) \rightarrow \text{real world} \rightarrow J(\pi_{\phi}^{\text{MPC}})$

### Reinforcement Learning

**Toolbox to solve MDPs from data**  
**This is not (necessarily) about DNNs**

**Policy**  $\pi_{\phi}^{\text{MPC}}(s) = u_0^*$  from

$$\begin{aligned} \min_{x,u} \quad & \bar{V}_{\phi}(x_N) + \sum_{k=0}^{N-1} l_{\phi}(x_k, u_k) \\ \text{s.t.} \quad & x_{k+1} = f_{\phi}(x_k, u_k) \\ & h_{\phi}(x_k, u_k) \geq 0, \quad x_0 = s \end{aligned}$$

**We want**

$$\min_{\phi} J(\pi_{\phi}^{\text{MPC}})$$

where  $J(\cdot)$  is the closed-loop cost in the real-world

**What to do?**

- ▶  $\min_{\phi} J(\pi_{\phi}^{\text{MPC}})$  using data
- ▶ Complex “chain”:  $\phi \rightarrow \pi_{\phi}^{\text{MPC}}(s) \rightarrow \text{real world} \rightarrow J(\pi_{\phi}^{\text{MPC}})$

### Reinforcement Learning

**Toolbox to solve MDPs from data**  
**This is not (necessarily) about DNNs**

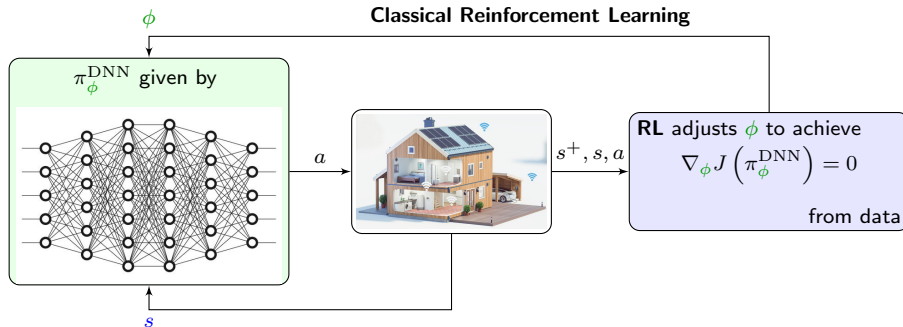
**For MPC:** tools to find best  $\phi$ , e.g.

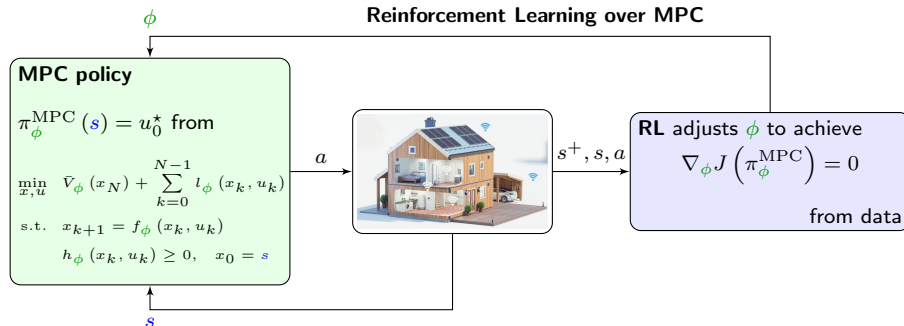
- ▶ **Policy Gradient:** estimations of

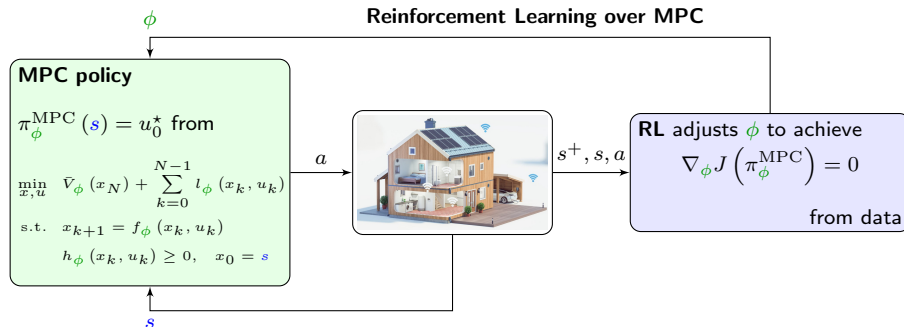
$$\nabla_{\phi} J(\pi_{\phi}^{\text{MPC}}), \quad \text{possibly} \quad \nabla_{\phi}^2 J(\pi_{\phi}^{\text{MPC}})$$

- ▶ **Q-learning:** directly seek  $Q^{\text{MPC}} \rightarrow Q^*$



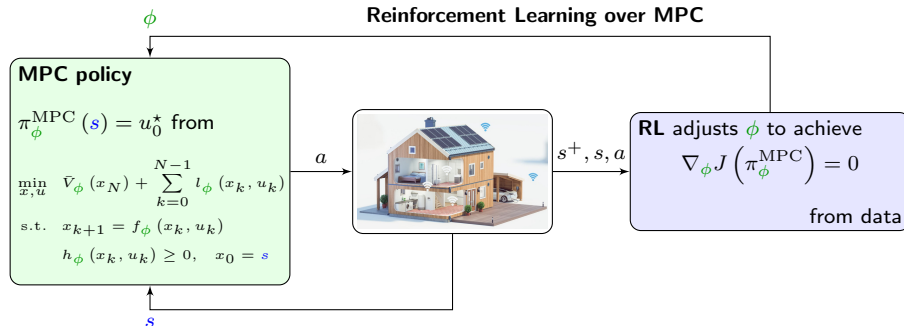






## Why is it useful?

- ✓ Tune MPC for real-world performance
- ✓ 25 years of guarantees & methods
- ✓ Learning does not start from scratch
- ✓ Well accepted in the industry
- ✓ Repeated Planning provides explainability
- ✓ Great at processing forecast information



## Why is it useful?

- ✓ Tune MPC for real-world performance
- ✓ 25 years of guarantees & methods
- ✓ Learning does not start from scratch
- ✓ Well accepted in the industry
- ✓ Repeated Planning provides explainability
- ✓ Great at processing forecast information

## Why can it be difficult?

- ✗ MPC is computationally more expensive than a DNN
- ✗ Deployment on GPUs is lagging
- ✗ Non-convexity can get in the way...
- ✗ RL methods are tailored to DNN, not MPC. There is still a gap here.

### Proposed paradigm

Policy  $\pi_{\phi}^{\text{MPC}}(s) = u_0^*$  from

$$\begin{aligned} \min_{x,u} \quad & \bar{V}_{\phi}(x_N) + \sum_{k=0}^{N-1} l_{\phi}(x_k, u_k) \\ \text{s.t.} \quad & x_{k+1} = f_{\phi}(x_k, u_k) \\ & h_{\phi}(x_k, u_k) \geq 0, \quad x_0 = s \end{aligned}$$

A step back: **model adjustment?**

Policy  $\pi_{\phi}^{\text{MPC}}(s) = u_0^*$  from

$$\begin{aligned} \min_{x,u} \quad & \bar{V}(x_N) + \sum_{k=0}^{N-1} l(x_k, u_k) \\ \text{s.t.} \quad & x_{k+1} = f_{\phi}(x_k, u_k) \\ & h(x_k, u_k) \geq 0, \quad x_0 = s \end{aligned}$$

Adjust  $\phi$  according to (e.g.)

1.  $\min_{\phi} \mathbb{E} \left[ \|f_{\phi}(S, A) - S^+\|^2 \right]$  vs.
2.  $\min_{\phi} J \left( \pi_{\phi}^{\text{MPC}} \right)$

... is that different?

A step back: **model adjustment?**

Policy  $\pi_{\phi}^{\text{MPC}}(s) = u_0^*$  from

$$\begin{aligned} \min_{x,u} \quad & \bar{V}(x_N) + \sum_{k=0}^{N-1} l(x_k, u_k) \\ \text{s.t.} \quad & x_{k+1} = f_{\phi}(x_k, u_k) \\ & h(x_k, u_k) \geq 0, \quad x_0 = s \end{aligned}$$

Adjust  $\phi$  according to (e.g.)

1.  $\min_{\phi} \mathbb{E} \left[ \|f_{\phi}(S, A) - S^+\|^2 \right]$  vs.
2.  $\min_{\phi} J \left( \pi_{\phi}^{\text{MPC}} \right)$

... 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

A step back: **model adjustment?**

Policy  $\pi_{\phi}^{\text{MPC}}(s) = u_0^*$  from

$$\begin{aligned} \min_{x,u} \quad & \bar{V}(x_N) + \sum_{k=0}^{N-1} l(x_k, u_k) \\ \text{s.t.} \quad & x_{k+1} = f_{\phi}(x_k, u_k) \\ & h(x_k, u_k) \geq 0, \quad x_0 = s \end{aligned}$$

Adjust  $\phi$  according to (e.g.)

1.  $\min_{\phi} \mathbb{E} \left[ \|f_{\phi}(S, A) - S^+\|^2 \right]$  vs.
2.  $\min_{\phi} J \left( \pi_{\phi}^{\text{MPC}} \right)$

... 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

**Formal answers:**\* 1. and 2. are

- ▶ the same for LQR
- ▶ *locally* the same for smooth MDPs with a “narrow” optimal steady state distribution away from the constraints



A step back: **model adjustment?**

Policy  $\pi_{\phi}^{\text{MPC}}(s) = u_0^*$  from

$$\begin{aligned} \min_{x,u} \quad & \bar{V}(x_N) + \sum_{k=0}^{N-1} l(x_k, u_k) \\ \text{s.t.} \quad & x_{k+1} = f_{\phi}(x_k, u_k) \\ & h(x_k, u_k) \geq 0, \quad x_0 = s \end{aligned}$$

Adjust  $\phi$  according to (e.g.)

1.  $\min_{\phi} \mathbb{E} \left[ \|f_{\phi}(S, A) - S^+ \|^2 \right]$  vs.
2.  $\min_{\phi} J \left( \pi_{\phi}^{\text{MPC}} \right)$

... 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

**Formal answers:**\* 1. and 2. are

- ▶ the same for LQR
- ▶ *locally* the same for smooth MDPs with a “narrow” optimal steady state distribution away from the constraints

**Practical answers:**

- ▶ RL for tracking(-like) MPC to a fixed reference is disappointing
- ▶ Otherwise the gain of performance from RL is often good

### Take-home messages

- ▶ The **MDP** is the common language linking **MPC** and **RL**
- ▶ Read MPC as a **structured model of  $Q^*, \pi^*$**  – via a holistic parametrization  $(\bar{V}_\phi, l_\phi, f_\phi, h_\phi)$ , not as fitting  $f$  to reality – exact even with an inexact model
- ▶ **RL** tunes  $\phi$  to  $\min_\phi J(\pi_\phi^{\text{MPC}})$  from data (policy gradient, Q-learning)
- ▶ Keep MPC's structure, constraints & guarantees; gain data-driven real-world performance

### Next in this tutorial:

- ▶ Differentiable NMPC with acados – Katrin Baumgärtner
- ▶ MPCRL with leap-c – Jasper Hoffmann

**Thank you! Questions?**



Slides & more