

# Synthesis of Model Predictive Control and Reinforcement Learning

– Overview over Synthesis of MPC and RL –

Dirk Reinhardt, Jasper Hoffmann

Fall School on Model Predictive Control and Reinforcement Learning  
Freiburg, 6-10 October 2025

universität freiburg



NTNU

Norwegian University of  
Science and Technology

- 🇳🇴 Postdoc at Cybernetics Department, NTNU, Norway
- 💻 Working on combining MPC and RL for real-world systems
- ⚡ Interested in applications in energy systems
- 🎓 PhD in Cybernetics, worked on Optimal Control of Unmanned Aerial Vehicles (acados inside)





## Reinforcement Learning

- ▶ Machine-Learning Community: *Use learning to mitigate complexities*
- ▶ **Philosophy:** Trial-and-error learning to build policies directly from experience (implicit)
- ▶ Offline Solution (explicit function)
- ▶ Simulation first
- ▶ Roots in Computer Science



## Model Predictive Control

- ▶ Control Community: *Use feedback to mitigate uncertainties*
- ▶ **Philosophy:** Model-based optimization over engineering models for planning (explicit)
- ▶ Online Solution (implicit function)
- ▶ Theory first
- ▶ Roots in Process Control

**But both want to solve sequential decision-making problems!**

[Video: Miki et al. 2022] [Video: Romero et al. 2024]

## Motivation:

- ▶ MPC and RL are complementary approaches for solving MDPs
- ▶ Nearly orthogonal advantages: weaknesses of one are strengths of the other
- ▶ Growing research interest in combination methods

## Goals of this presentation:

- ▶ Illuminate differences, similarities, and fundamentals enabling combinations
- ▶ Categorize existing work based on actor-critic RL framework

## Synthesis of Model Predictive Control and Reinforcement Learning: Survey and Classification\*

Rudolf Reiter<sup>a,\*,</sup>, Jasper Hoffmann<sup>b,\*,</sup>, Dirk Reinhardt<sup>d</sup>, Florian Messerer<sup>a</sup>, Katrin Baumgärtner<sup>a</sup>, Shambhura Sawant<sup>d</sup>,  
Joschka Bödecker<sup>a</sup>, Moritz Diehl<sup>a,c</sup>, Sebastian Gros<sup>d</sup>

<sup>a</sup>Department of Microsystems Engineering, University of Freiburg, Freiburg im Breisgau, 79110, Baden-Württemberg, Germany

<sup>b</sup>Department of Computer Science, University of Freiburg, Freiburg im Breisgau, 79110, Baden-Württemberg, Germany

<sup>c</sup>Department of Mathematics, University of Freiburg, Freiburg im Breisgau, 79104, Baden-Württemberg, Germany

<sup>d</sup>Department of Engineering Cybernetics, Norwegian University of Science and Technology, Trondheim, 7034, Norway



Comparing MPC and RL

Inference Architectures

- Parameterized

- Parallel

- Hierarchical

- Integrated

- Algorithmic

Taxonomy of Combination Approaches

- MPC as an Expert Actor

- MPC within the Deployed Policy

  - Aligned Learning

  - Closed-Loop Optimal Learning

- MPC as the Critic

- MPC for Preprocessing/Postprocessing

Current Challenges

Summary

## Comparing MPC and RL

### Inference Architectures

- Parameterized
- Parallel
- Hierarchical
- Integrated
- Algorithmic

### Taxonomy of Combination Approaches

- MPC as an Expert Actor
- MPC within the Deployed Policy
  - Aligned Learning
  - Closed-Loop Optimal Learning
- MPC as the Critic
- MPC for Preprocessing/Postprocessing

### Current Challenges

### Summary

# Comparison of Practical Properties



| Property                 | MPC                                       | RL                                 |
|--------------------------|-------------------------------------------|------------------------------------|
| state-space              | model specific ✗                          | (quite) arbitrary ✓                |
| model requirements       | differentiability/<br>online simulation ✗ | offline simulation ✓               |
| uncertainty              | guarantees with<br>known uncertainty ✓    | probabilistic<br>guarantees ✓      |
| stability                | strong theory ✓                           | minor theory ✗                     |
| constraint handling      | inherent ✓                                | hard to achieve ✗                  |
| online computation time  | high ✓                                    | low ✗                              |
| offline computation time | low ✗                                     | high ✓                             |
| adaptability             | inherent ✓                                | needs retraining ✗                 |
| generalization           | inherent ✓                                | poor when<br>out-of-distribution ✗ |

Comparison of practical properties between MPC and RL.  
The evaluation is simplified and conceptual. Exemptions may exist.

## MPC Problem

$$\begin{aligned} \min_z \quad & \sum_{k=0}^{N-1} l^{\text{MPC}}(x_k, u_k) + \bar{V}^{\text{MPC}}(x_N) \\ \text{s.t.} \quad & x_0 = s, \\ & x_{k+1} = f^{\text{MPC}}(x_k, u_k), \\ & 0 \leq h^{\text{MPC}}(x_k, u_k), \\ & 0 \leq h_N^{\text{MPC}}(x_N). \end{aligned}$$

**Change in dynamics?**

Modify  $f^{\text{MPC}}$   $\rightarrow$  new optimal policy

## RL Policy (DDPG)

$$\begin{aligned} w &\leftarrow w + \frac{\alpha_w}{B} \sum_{i=1}^B \delta_i \nabla_w Q_w(S_i, A_i), \\ \theta &\leftarrow \theta + \frac{\alpha_\theta}{B} \sum_{i=1}^B \nabla_\theta \mu_\theta(S_i) \nabla_a Q_w(S_i, a)|_{a=\mu_\theta(S_i)}, \\ \delta_i &:= l(S_i, A_i) + \gamma Q_{\bar{w}}(S_i^+, A_i^+) - Q_w(S_i, A_i), \\ &\text{with } (S_i, A_i, S_i^+) \sim \mathcal{D}^{\text{buffer}}, \quad A_i^+ = \mu_\theta(S_i^+), \end{aligned}$$

**Change in dynamics?**

Collect new data  $\mathcal{D}^{\text{buffer}}$   $\rightarrow$  retrain policy

## MPC Problem

$$\begin{aligned} \min_z \quad & \sum_{k=0}^{N-1} l^{\text{MPC}}(x_k, u_k) + \bar{V}^{\text{MPC}}(x_N) \\ \text{s.t.} \quad & x_0 = s, \\ & x_{k+1} = f^{\text{MPC}}(x_k, u_k), \\ & 0 \leq h^{\text{MPC}}(x_k, u_k), \\ & 0 \leq h_N^{\text{MPC}}(x_N). \end{aligned}$$

### Constraints?

Directly handled through  $h^{\text{MPC}}$  and  $h_N^{\text{MPC}}$

## RL Policy (DDPG)

$$\begin{aligned} w &\leftarrow w + \frac{\alpha_w}{B} \sum_{i=1}^B \delta_i \nabla_w Q_w(S_i, A_i), \\ \theta &\leftarrow \theta + \frac{\alpha_\theta}{B} \sum_{i=1}^B \nabla_\theta \mu_\theta(S_i) \nabla_a Q_w(S_i, a)|_{a=\mu_\theta(S_i)}, \\ \delta_i &:= l(S_i, A_i) + \gamma Q_{\bar{w}}(S_i^+, A_i^+) - Q_w(S_i, A_i), \\ &\text{with } (S_i, A_i, S_i^+) \sim \mathcal{D}^{\text{buffer}}, \quad A_i^+ = \mu_\theta(S_i^+), \end{aligned}$$

### Constraints?

Hard to handle formally  $\rightarrow$  use penalties in cost  $l$  or restrict action space of policy  $\mu_\theta$ .

## MPC Problem

$$\begin{aligned} \min_z \quad & \sum_{k=0}^{N-1} l^{\text{MPC}}(x_k, u_k) + \bar{V}^{\text{MPC}}(x_N) \\ \text{s.t.} \quad & x_0 = s, \\ & x_{k+1} = f^{\text{MPC}}(x_k, u_k), \\ & 0 \leq h^{\text{MPC}}(x_k, u_k), \\ & 0 \leq h_N^{\text{MPC}}(x_N). \end{aligned}$$

### Inference time?

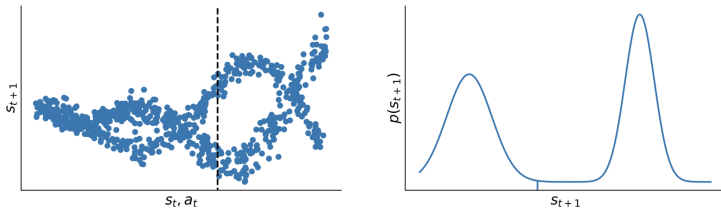
Often high due to **online optimization** → may be prohibitive for fast systems

## RL Policy (DDPG)

$$\begin{aligned} w &\leftarrow w + \frac{\alpha_w}{B} \sum_{i=1}^B \delta_i \nabla_w Q_w(S_i, A_i), \\ \theta &\leftarrow \theta + \frac{\alpha_\theta}{B} \sum_{i=1}^B \nabla_\theta \mu_\theta(S_i) \nabla_a Q_w(S_i, a)|_{a=\mu_\theta(S_i)}, \\ \delta_i &:= l(S_i, A_i) + \gamma Q_{\bar{w}}(S_i^+, A_i^+) - Q_w(S_i, A_i), \\ &\text{with } (S_i, A_i, S_i^+) \sim \mathcal{D}^{\text{buffer}}, \quad A_i^+ = \mu_\theta(S_i^+), \end{aligned}$$

### Inference time?

Low, only need forward pass of  $\mu_\theta$  → suitable for fast systems



[Image: Moerland et al. 2022]

- ▶ **Definition:** Environment is inherently stochastic - deterministic future state prediction impossible even with perfect model knowledge
- ▶ **RL approaches:**
  - ▶ **Natural fit:** Stochasticity inherent in MDP framework
  - ▶ **Direct training:** Policies trained on real-world environment to account for actual uncertainty
  - ▶ **Challenge:** Rare events (large but infrequent deviations) make value estimation difficult

**Challenge:** Training environment often differs from deployment environment

## **RL approaches:**

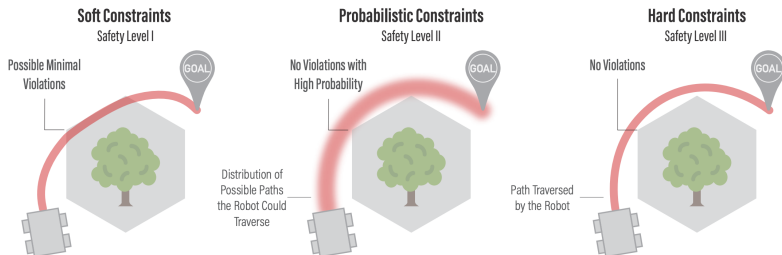
- ▶ Domain randomization
- ▶ Robust RL
- ▶ Meta RL - training on adaptive distribution of models

**RL assumption:** Real-world environment lies within training model distribution

## **MPC approaches:**

- ▶ Stochastic MPC - explicitly accounts for model mismatch
- ▶ Robust MPC - handles correlated predictive errors across time
- ▶ Distributionally robust MPC - robustifies against distribution mismatches

**MPC advantage:** Inherent robustness - performs well even without explicit model mismatch consideration



[Image: Brunke et al. 2023]

- ▶ **MPC:** Constraint satisfaction and stability are key strengths:
  - ▶ **Challenge:** Hard to handle probabilistic constraints with (unknown) uncertainty
  - ▶ **Common MPC approaches:** Robust MPC, tube-based MPC
- ▶ **RL:** Increasingly focused on safety guarantees using constrained MDPs:
  - ▶ **Challenge:** Safe RL policies often become too conservative or may still violate constraints.
  - ▶ **Common RL approaches:** Safe action sets, Lagrange formulation, trust-region optimization, control-barrier functions, MPC based safety filter

Comparing MPC and RL

Inference Architectures

- Parameterized

- Parallel

- Hierarchical

- Integrated

- Algorithmic

Taxonomy of Combination Approaches

- MPC as an Expert Actor

- MPC within the Deployed Policy

  - Aligned Learning

  - Closed-Loop Optimal Learning

- MPC as the Critic

- MPC for Preprocessing/Postprocessing

Current Challenges

Summary

## Parameterized OCP with parameters $\phi$ :

### OCP:

$$\begin{aligned} \min_z \quad & L_\phi(z) \\ \text{s.t.} \quad & x_0 = s, \\ & x_{k+1} = f_\phi^{\text{MPC}}(x_k, u_k), \quad 0 \leq k < N, \\ & 0 \leq h_\phi^{\text{MPC}}(x_k, u_k), \quad 0 \leq k < N, \\ & 0 \leq \tilde{h}_\phi^{\text{MPC}}(x_N), \end{aligned}$$

### Objective:

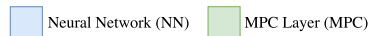
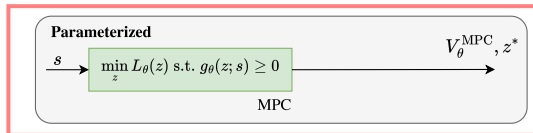
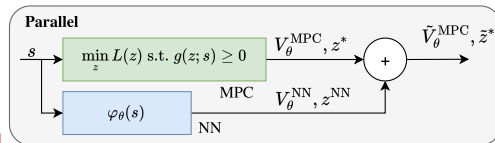
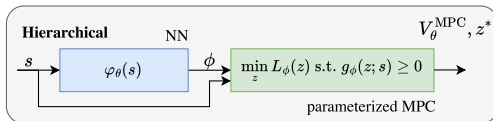
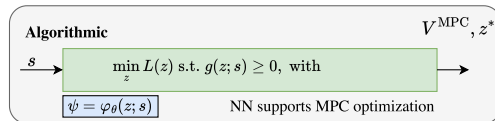
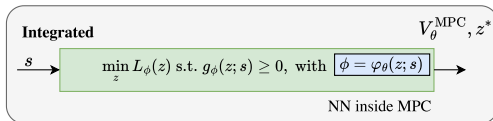
$$L_\phi(z) := \bar{V}_\phi^{\text{MPC}}(x_N) + \sum_{k=0}^{N-1} l_\phi^{\text{MPC}}(x_k, u_k).$$

### Short notation:

$$\min_z L_\phi(z) \quad \text{s.t.} \quad g_\phi(z; s) \geq 0.$$

**Parameters**  $\phi$  can affect: objective  $L_\phi$ , dynamics  $f_\phi^{\text{MPC}}$ , and constraints  $h_\phi^{\text{MPC}}$ .

**Note:**  $g_\phi$  encodes both dynamics and constraints.



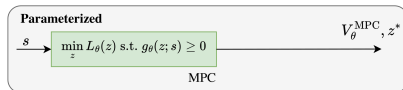
**Architecture:** Parameters  $\theta$  are constant and independent of decision variables  $z$  or state  $s$

## Key characteristics:

- ▶ Parameters learned offline, fixed during deployment
- ▶ Optimization-friendly MPC formulation
- ▶ No evaluation of highly nonlinear functions

## Advantages:

- ✓ Preserves MPC problem structure
- ✓ No NN evaluation during inference
- ✓ Suitable for standard MPC solvers
- ✓ Lower computational complexity
- ✗ Less flexibility than state-dependent architectures



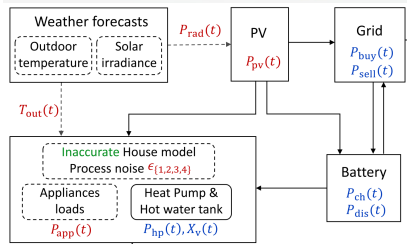
## Examples:

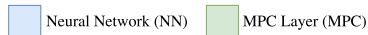
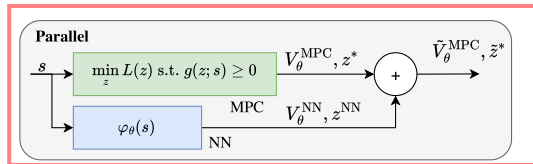
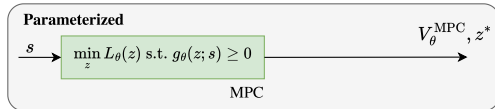
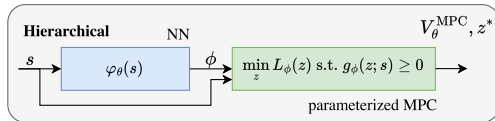
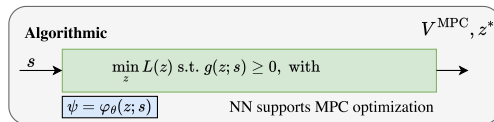
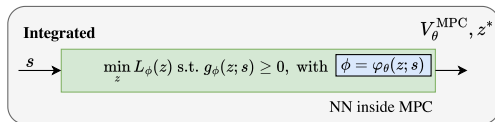
- ▶ Simplified models [Cai et al. 2023]
- ▶ Cost function tuning

## A Learning-Based Model Predictive Control Strategy for Home Energy Management Systems

WENQI CAI<sup>1</sup>, SHAMBHURAJ SAWANT<sup>1</sup>, DIRK REINHARDT<sup>1</sup>,  
SOROUSH RASTEGARPOUR<sup>2</sup>, AND SÉBASTIEN GROS<sup>1</sup>

- ▶ Build a simplified MPC scheme and parameterize model, cost constraints
- ▶ Learn parameters via RL
- ▶ Maintains MPC structure, lower inference time compared to full model





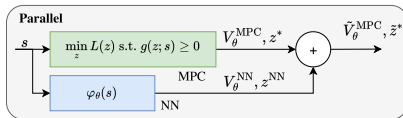
**Architecture:** MPC problem usually not parameterized; NN evaluated in parallel to correct MPC output

## Key characteristics:

- ▶ MPC problem typically not parameterized
- ▶ NN evaluated in parallel to MPC optimization
- ▶ NN output corrects optimal solution or value function
- ▶ MPC and NN operate independently

## Trade-offs:

- ▶ No differentiation through optimization ✓
- ▶ Potentially unsafe actions due to perturbation ✗



## Examples:

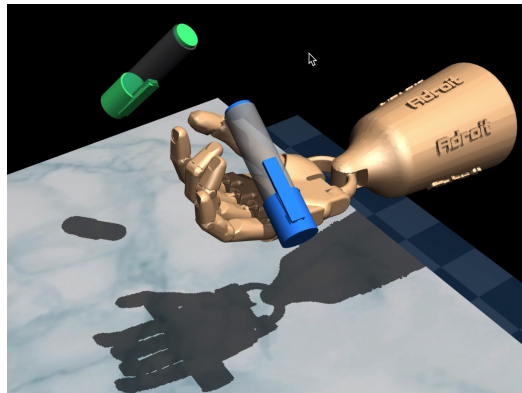
- ▶ Value function approximation  
[Bhardwaj et al. 2020]

## BLENDING MPC & VALUE FUNCTION APPROXIMATION FOR EFFICIENT REINFORCEMENT LEARNING

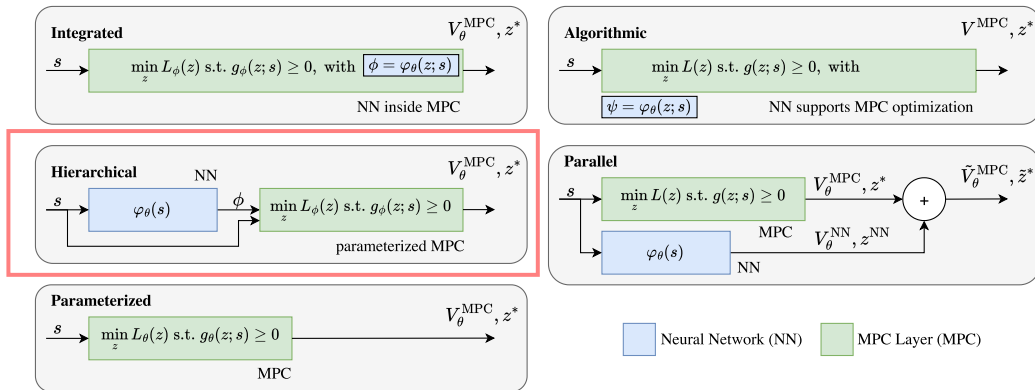
Mohak Bhardwaj<sup>1</sup>   Sanjiban Choudhury<sup>2</sup>   Byron Boots<sup>1</sup>

<sup>1</sup> University of Washington   <sup>2</sup> Aurora Innovation Inc.

- ▶ Use MPC for local value estimates
- ▶ NN provides corrections
- ▶ Blend both estimates via weighting factor  $\lambda$
- ▶ Improves data efficiency and adds structure



[Image: Bhardwaj et al. 2020]



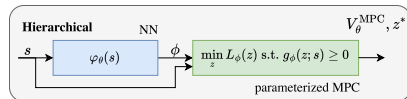
**Architecture:** NN provides input to MPC, parameter  $\phi$  depends on state via  $\phi = \varphi_{\theta}(s)$

## Key characteristics:

- ▶ NN evaluated *before* solving optimization problem
- ▶ Function  $\varphi_{\theta}(s)$  evaluated independently of MPC
- ▶ Nonlinearity does not affect optimization structure
- ▶ Parameter changes may affect convergence

## Tradeoffs:

- ✓ Optimization not harder to solve
- ✓ No NN inside optimization
- ✓ Standard MPC solvers work
- ✗ MPC theory may not hold

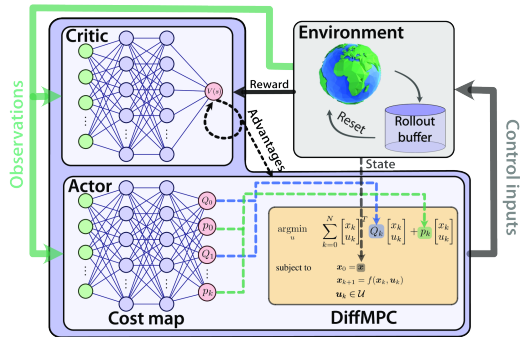
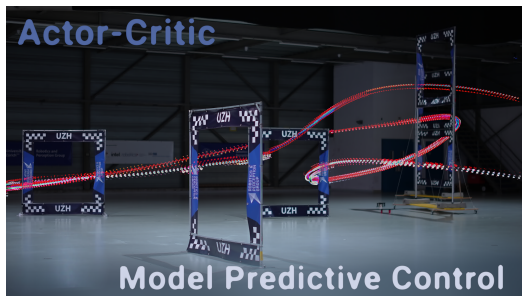


## Examples:

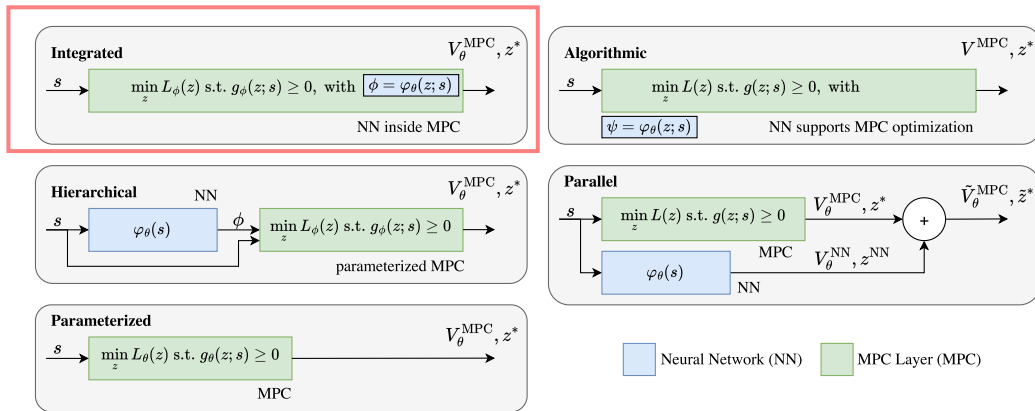
- ▶ Reference trajectory [Brito et al. 2021] [Reiter et al. 2023]
- ▶ Cost modification [Romero et al. 2024]
- ▶ Constraint modification [Jacquet et al. 2024]

## Actor-Critic Model Predictive Control

Angel Romero, Yunlong Song, Davide Scaramuzza



[Image: Romero et al. 2024]



**Architecture:** Parameters depend on both state  $s$  and optimization variables  $z$  via

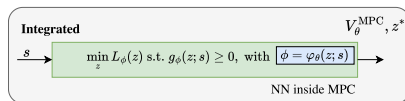
$$\phi = \varphi_{\theta}(s, z)$$

## Key characteristics:

- ▶ NN is part of the optimization problem itself
- ▶ Cannot be evaluated separately from optimization algorithm
- ▶ NN encoders/decoders infer states, rewards, transition models
- ▶ Outputs value function  $V_{\theta}^{\text{MPC}}$  or decision variables  $z^*$

## Tradeoffs:

- ✗ Highly nonlinear, possibly nonconvex
- ✗ Computationally expensive
- ✗ Requires  $\nabla_z \varphi_{\theta}(z; s)$



## Examples:

- ▶ Latent space MPC [Lowrey et al. 2019]
- ▶ NN transition models [Salzmann et al. 2023][Adhau et al. 2024]
- ▶ Cost function [Seel et al. 2022]

## Idea:

- ▶ Learn NN dynamics model from data
- ▶ Use model for MPC planning

## Anyone wants to try?



[i4casadi]



[i4acados]

IEEE ROBOTICS AND AUTOMATION LETTERS, VOL. 8, NO. 4, APRIL 2023

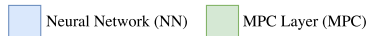
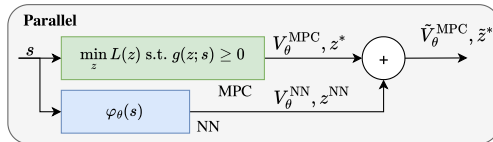
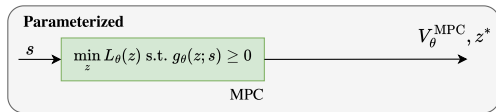
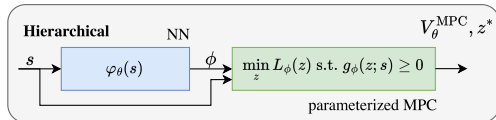
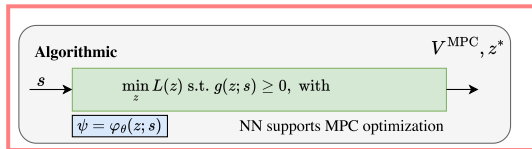
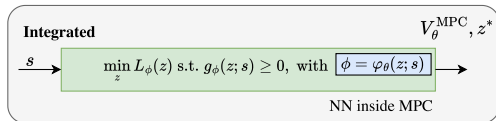
2307

## Real-Time Neural MPC: Deep Learning Model Predictive Control for Quadrotors and Agile Robotic Platforms

Tim Salzmann<sup>✉</sup>, Elia Kaufmann<sup>✉</sup>, *Member, IEEE*, Jon Arrizabalaga<sup>✉</sup>, *Graduate Student Member, IEEE*,  
Marco Pavone<sup>✉</sup>, *Member, IEEE*, Davide Scaramuzza<sup>✉</sup>, *Senior Member, IEEE*, and Markus Ryll<sup>✉</sup>, *Member, IEEE*



[Image: Salzmann et al. 2023]



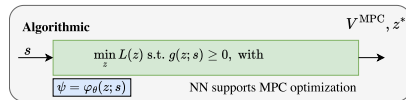
**Architecture:** RL policy guides the MPC optimization algorithm without modifying the optimization problem itself

## Key characteristics:

- ▶ RL policy alters hyperparameters  $\psi$  of the solver
- ▶ Optimization problem remains unchanged
- ▶ Examples of hyperparameters:
  - ▶ Initial trajectory guess
  - ▶ Trajectory samples (MPPI)
  - ▶ Active-set predictions
  - ▶ MPC recompute flag
  - ▶ Horizon length

## Tradeoffs:

- ✓ Improve online computation time
- ✓ Can help escape local minima
- ✗ Increased complexity



## Examples:

- ▶ TD-MPC2 [Hansen et al. 2024]
- ▶ Warm-starting [Grandesso et al. 2023]

## RL Components:

- ▶ Joint-embedding prediction
- ▶ Reward prediction
- ▶ TD learning (Q-function)
- ▶ Long-term credit assignment

## MPC Components:

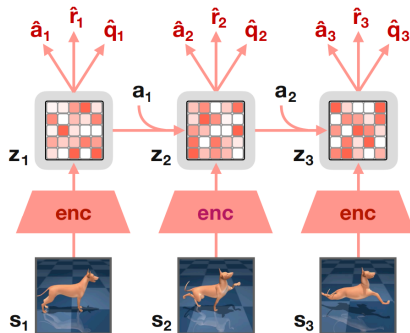
- ▶ Local trajectory optimization
- ▶ Planning in latent space
- ▶ Policy prior guidance

## Key Innovation:

Bootstraps beyond planning horizon using learned Q-function

## TD-MPC<sup>2</sup>: Scalable, Robust World Models for Continuous Control

Nicklas Hansen\*, Hao Su<sup>†</sup>, Xiaolong Wang<sup>†</sup>  
\*University of California San Diego, <sup>†</sup>Equal advising  
{nihansen, haosu, xiw012}@ucsd.edu



[Image: Hansen et al. 2023]

Comparing MPC and RL

Inference Architectures

Parameterized

Parallel

Hierarchical

Integrated

Algorithmic

Taxonomy of Combination Approaches

MPC as an Expert Actor

MPC within the Deployed Policy

Aligned Learning

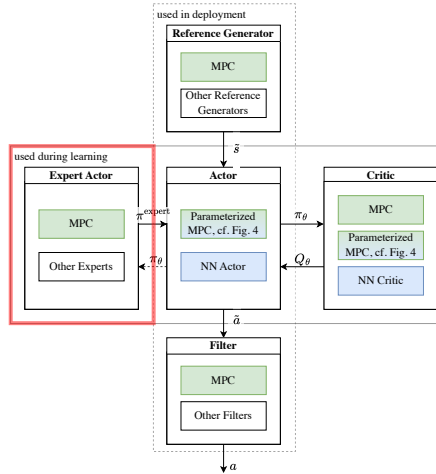
Closed-Loop Optimal Learning

MPC as the Critic

MPC for Preprocessing/Postprocessing

Current Challenges

Summary



## MPC used to generate training data for RL policy

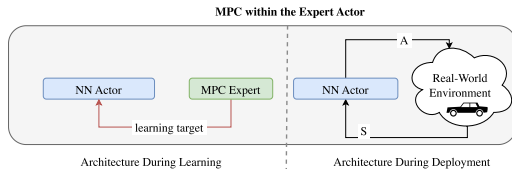
**Purpose:** Use MPC to provide expert trajectories for imitation learning

### Characteristics:

- ▶ MPC generates state-action pairs
- ▶ RL policy learns to mimic MPC behavior
- ▶ Can use behavioral cloning or DAgger
- ▶ See also [Karg and Lucia 2020]

### Benefits:

- ✓ Faster inference than MPC
- ✓ Leverages domain knowledge
- ✓ Sample-efficient learning



### Examples:

- ▶ Autonomous driving [Hoffmann et al. 2024]

### Challenge:

- ✗ Limited to MPC performance

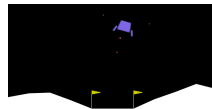
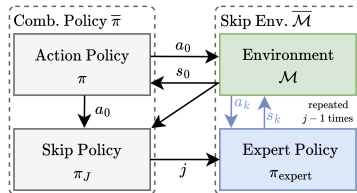
## Learning When to Trust the Expert for Guided Exploration in RL

Felix Schulz \*  
Neurorobotics Lab  
University of Freiburg  
Germany  
felix.schulz@student.hpi.de

Jasper Hoffmann \*  
Neurorobotics Lab  
University of Freiburg  
Germany  
hoffmaja@informatik.uni-freiburg.de

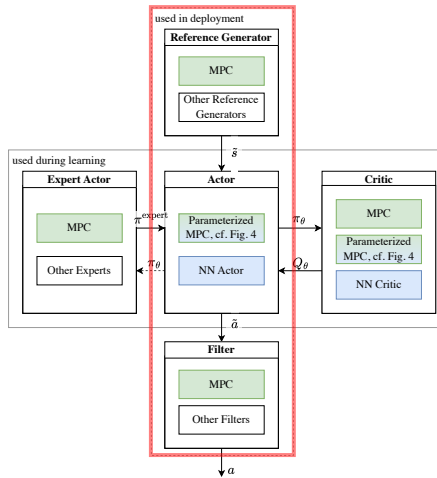
Yuan Zhang  
Neurorobotics Lab  
University of Freiburg  
Germany  
yzhang@informatik.uni-freiburg.de

Joschka Bödecker  
Neurorobotics Lab  
University of Freiburg  
Germany  
jboedeck@informatik.uni-freiburg.de

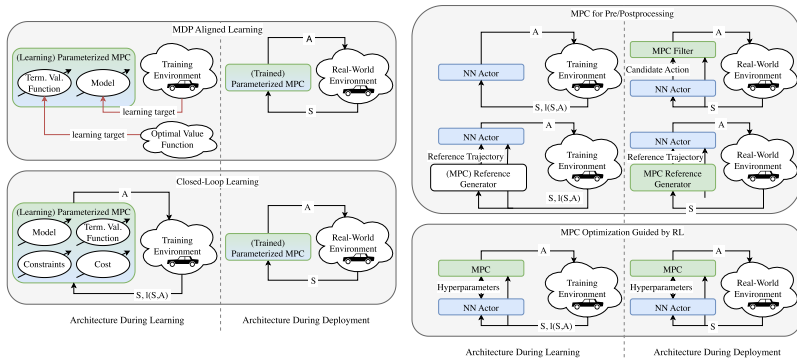


[Image: Schulz and Hoffmann et al. 2024]

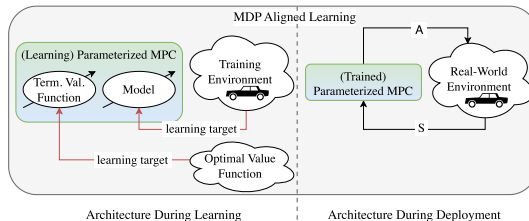
- ▶ Use an expert actor for exploration during training
- ▶ A learned high-level skip policy  $\pi_J$  decides at each for how many steps it should either follow a learned action policy  $\pi$  or the MPC expert  $\pi_{\text{expert}}$
- ▶ Reduce usage of expert during training



# MPC within the Deployed policy



**In the following, we will focus on Aligned and Closed-Loop Learning!**



**Paradigm:** Learn individual components of the MPC controller to better approximate the true environment. This aligns the controller's internal understanding with reality.

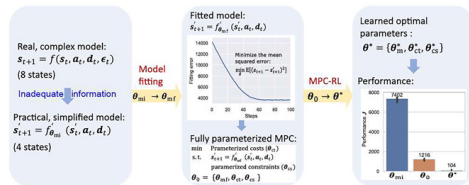
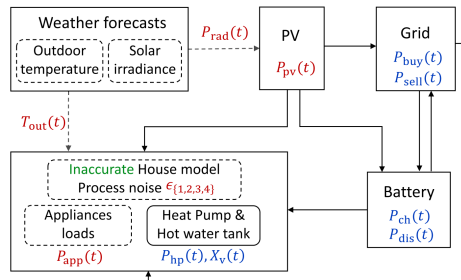
## Components aligned with MDP:

- ⚙️ Approximate the system's dynamics by learning:  $f_{\theta}^{\text{MPC}} \approx P$
- 🧭 Approximate the infinite-horizon cost with RL:  $\bar{V}_{\theta}^{\text{MPC}} \approx V^*$

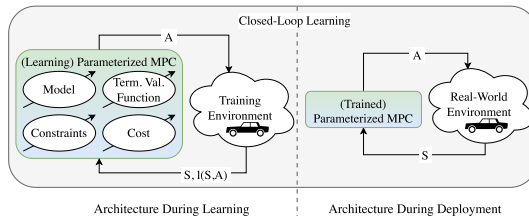
## A Learning-Based Model Predictive Control Strategy for Home Energy Management Systems

WENQI CAI<sup>1</sup>, SHAMBHURAJ SAWANT<sup>1</sup>, DIRK REINHARDT<sup>1</sup>,  
SOROUSH RASTEGARPOUR<sup>2</sup>, AND SÉBASTIEN GROS<sup>1</sup>

- Build a simplified MPC scheme and parameterize model, cost constraints
- Learn parameters via RL
- Maintains MPC structure, lower inference time compared to full model



[Image: Cai et al. 2023]



## Paradigm:

🏆 The focus is on what **works best**, not what is most **accurate**. The internal model or costs may become misaligned.

## How is it implemented?

📈 Learn MPC parameters  $\phi$  to directly optimize the final closed-loop performance  $J(\phi)$ .

🔗 **Differentiable MPC:** Treat the MPC as a layer in a policy network. Gradients are backpropagated *through* the optimization process to update parameters.

🏠 **MPC as part of the Environment:** An outer-loop RL agent learns to choose the best MPC parameters  $\phi$  as its "action", treating the MPC-controlled system as a black box.

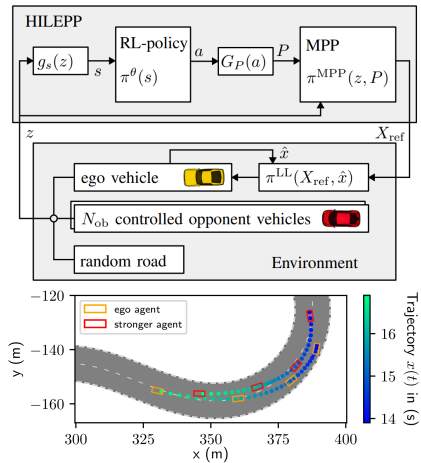
# Closed-loop Optimal Learning: Example

2023 European Control Conference (ECC)  
June 13-16, 2023, Bucharest, Romania

## A Hierarchical Approach for Strategic Motion Planning in Autonomous Racing

Rudolf Reiter<sup>1</sup>, Jasper Hoffmann<sup>2</sup>, Joschka Boedecker<sup>2</sup> and Moritz Diehl<sup>1,3</sup>

- ▶ Car-racing environment
- ▶ RL learns to set reference trajectories
- ▶ MPC tracks learned reference trajectories



[Image: Reiter et al. 2023]

## Aligned Learning

 Approximate the true Markov Decision Process

**Model Learning:** Supervised learning on transition data, e.g.,

$$\min_{\theta} \mathcal{L}(S^+ - f_{\theta}^{\text{MPC}}(S, A))$$

**Terminal Value Learning:** RL-style updates towards

$$V_{\theta}^{\text{MPC}} \approx V^*$$

### Tradeoffs:

- ✓ Interpretable model, cost, constraints
- ✓ Division of design into manageable subtasks
- ✓ Safe, generalizes better
- ✗ Sub-optimal closed-loop performance
- ✗ May require complex models
- ✗ Restrictive model requirements (differentiability, smoothness)

## Closed-Loop Learning

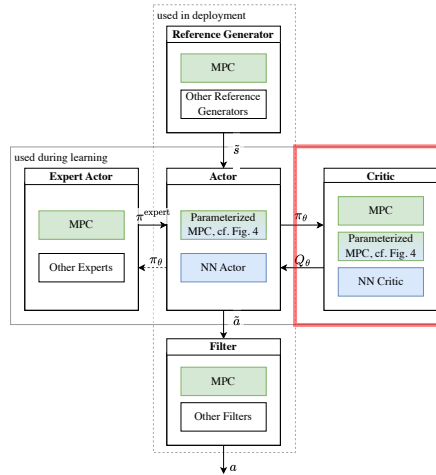
 Approximate the optimal policy

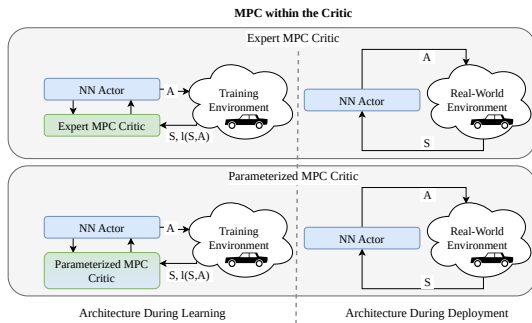
**Policy Learning:** Directly optimize closed-loop performance, towards

$$\mu_{\theta}^{\text{MPC}} \approx \mu^*$$

### Tradeoffs:

- ✓ Direct optimization of actual objective
- ✓ Can correct for model errors through learning
- ✓ Flexibility in model choice
- ✗ Limited generalization to new scenarios
- ✗ Not representing true system
- ✗ Reduced adaptability when requirements change





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

s.t.  $x_0 = s,$   
 $u_0 = a,$   
 $x_{k+1} = f^{\text{MPC}}(x_k, u_k), 0 \leq k < N,$   
 $0 \leq h^{\text{MPC}}(x_k, u_k), 1 \leq k < N,$   
 $0 \leq h_N^{\text{MPC}}(x_N).$

- We can use MPC to derive a Q-function  $Q^{\text{MPC}}$ , which can be used as a critic!
- We constrain the initial control  $u_0$  to be  $a$ .



**Instead of using a fixed MPC critic, we can further learn to improve the critic:**

- ▶ In [Bhardwaj et al. 2020], the author combine sampling-based MPC with a learned value function using a parallel architecture solving multiple robot manipulation tasks. A simplified version of the critic is given by:

$$Q(s, a) = (1 - \lambda) \underbrace{Q_w(s, a)}_{\text{learned critic}} + \lambda \underbrace{Q^{\text{MPC}}(s, a)}_{\text{model-based}},$$

where  $Q_w$  is a learned critic and  $Q^{\text{MPC}}$  is the MPC-based Q-function.

- ▶ The parameter  $\lambda$  balances the two Q-functions and is decayed over time.

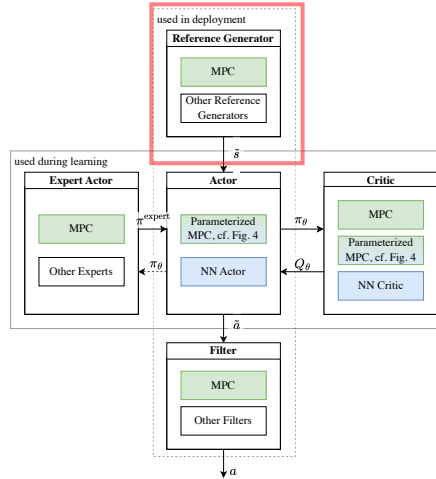
- ▶ [Anand et al. 2023] proposed a simple actor-critic scheme, where a single MPC scheme is used for both actor and critic.
- ▶ **Idea:** Close-to-optimal MPC parameters  $\phi$  allow approximation of the optimal Q-function  $Q^*$ .

## Local Q-function approximation:

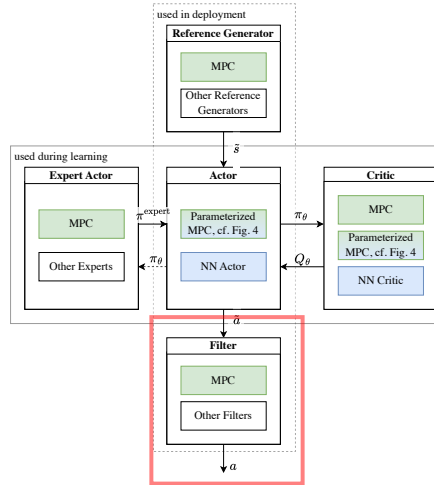
$$Q_w(s, a) = Q_\phi^{\text{MPC}}(s, a) + \nabla_\phi Q_\phi^{\text{MPC}}(s, a)^\top w$$

## MPC Q-Function

$$\begin{aligned} Q_\phi^{\text{MPC}}(s, a) = & \min_z \sum_{k=0}^{N-1} l_\phi^{\text{MPC}}(x_k, u_k) + \bar{V}_\phi^{\text{MPC}}(x_N) \\ \text{s.t.} \quad & x_0 = s, \\ & u_0 = a, \\ & x_{k+1} = f_\phi^{\text{MPC}}(x_k, u_k), \quad , \\ & 0 \leq h_\phi^{\text{MPC}}(x_k, u_k), \quad , \\ & 0 \leq h_{N_\phi}^{\text{MPC}}(x_N). \end{aligned}$$



## 46



## MPC used in deployed policy, but not during training

### Postprocessing: Safety Filter

**Purpose:** Ensure safety w.r.t. model and constraints

#### Characteristics:

- ▶ MPC filters RL policy output
- ▶ Projects actions to safe set

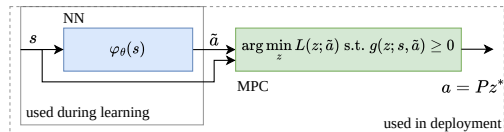
#### Benefit:

- ✓ Provides safety guarantees w.r.t. model and constraints

#### Issue:

- ✗ Policy unaware of filter during training
- ✗ May lead to suboptimal actions

#### MPC for postprocessing

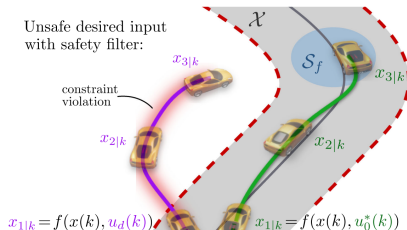


IEEE ROBOTICS AND AUTOMATION LETTERS, VOL. 6, NO. 4, OCTOBER 2021

#### A Predictive Safety Filter for Learning-Based Racing Control

Ben Tearle<sup>1</sup>, Kim P. Wabersich<sup>2</sup>, Andrea Carron<sup>3</sup>, and Melanie N. Zeilinger<sup>1</sup>, Member, IEEE

Unsafe desired input  
with safety filter:



Comparing MPC and RL

Inference Architectures

- Parameterized

- Parallel

- Hierarchical

- Integrated

- Algorithmic

Taxonomy of Combination Approaches

- MPC as an Expert Actor

- MPC within the Deployed Policy

  - Aligned Learning

  - Closed-Loop Optimal Learning

- MPC as the Critic

- MPC for Preprocessing/Postprocessing

Current Challenges

Summary

-  **Computational efficiency:** Scaling to high-dimensional offline RL with massive datasets remains challenging
-  **Risk-aware learning:** Moving beyond robust/stochastic MPC toward sophisticated risk metrics (e.g., mission completion probabilities)
-  **Expanding frameworks:** Generalization from MPC to broader planning (combinatorial optimization, MILP, stochastic multi-stage programming)
-  **Multi-agent systems:** Coordination, communication, and distributed decision-making largely unresolved
-  **Application domains:** Need for more practical demonstrations across diverse sectors (energy, healthcare, logistics)



## Key features:

- ⚙️ Hierarchical NN-MPC architecture implementation
- 🦆 Seamless integration of acados into PyTorch as differentiable layer
- 📈 Supports differentiation of control/state trajectories and Q-function w.r.t. parameters
- 🚀 Efficient multithreaded implementation for batched inputs
- 🧰 Modular codebase
- ↺️ Focus on closed-loop learning
- ⚠️ Missing a differentiable convex programming layer
- 🔗 <https://github.com/leap-c/leap-c>



**Having heard this lecture, you should have a better understanding of...**

-  Identify strengths and weaknesses of MPC and RL
-  Recognize potential benefits of combining both approaches
-  Distinguish between parameterized MPC architectures:
  - ▶ Integrated, hierarchical, parallel, parameterized
-  Design or use MPC-RL algorithms based on:
  -  MPC as expert actor
  -  MPC within deployed policy
  -  MPC as critic
  -  MPC as reference generator
  -  MPC as a safety filter

**Thank you for your attention!**