

Model-plant mismatch compensation using reinforcement learning

Koryakovskiy, Ivan; Kudruss, Manuel; Vallery, Heike; Babuska, Robert; Caarls, Wouter

DOI

[10.1109/LRA.2018.2800106](https://doi.org/10.1109/LRA.2018.2800106)

Publication date

2018

Document Version

Accepted author manuscript

Published in

IEEE Robotics and Automation Letters

Citation (APA)

Koryakovskiy, I., Kudruss, M., Vallery, H., Babuska, R., & Caarls, W. (2018). Model-plant mismatch compensation using reinforcement learning. *IEEE Robotics and Automation Letters*, 3(3), 2471 - 2477. <https://doi.org/10.1109/LRA.2018.2800106>

Important note

To cite this publication, please use the final published version (if applicable).
Please check the document version above.

Copyright

Other than for strictly personal use, it is not permitted to download, forward or distribute the text or part of it, without the consent of the author(s) and/or copyright holder(s), unless the work is under an open content license such as Creative Commons.

Takedown policy

Please contact us and provide details if you believe this document breaches copyrights.
We will remove access to the work immediately and investigate your claim.

Model-plant Mismatch Compensation Using Reinforcement Learning

Ivan Koryakovskiy¹, Manuel Kudruss², Heike Vallery¹, Robert Babuška¹, Wouter Caarls³

Abstract—Learning-based approaches are suitable for the control of systems with unknown dynamics. However, learning from scratch involves many trials with exploratory actions until a good control policy is discovered. Real robots usually cannot withstand the exploratory actions and suffer damage. This problem can be circumvented by combining learning with model-based control. In this article, we employ a nominal model-predictive controller that is impeded by the presence of an unknown model-plant mismatch. To compensate for the mismatch, we propose two approaches of combining reinforcement learning with the nominal controller. The first approach learns a compensatory control action which minimizes the same performance measure as is minimized by the nominal controller. The second approach learns a compensatory signal from a difference of a transition predicted by the internal model and an actual transition. We compare the approaches on a robot attached to the ground and performing a setpoint reaching task in simulations. We implement the better approach on the real robot and demonstrate successful learning results.

Index Terms—Learning and Adaptive Systems; Humanoid Robots

I. INTRODUCTION

MECHANICALLY and electronically, robotics have advanced to the point where cognitive abilities have become the main limiting factor. While robots can flawlessly execute a set of commands to achieve a task, these commands are mostly encoded or tuned by hand. Reinforcement learning (RL) allows to find an optimal sequence of commands without any prior assumption about the world. However, the application of pure learning to real systems is very limited due to intrinsically damaging exploratory policies. For example, when learning from scratch the robot Leo depicted in Figure 1 can withstand only five minutes of operation due to large and rapidly changing motor torques and frequent falls [1].

Usually, the dynamics of physical systems are known, but various uncertainties do not allow achieving optimal performance with model-based control methods [2]. Whereas for the estimation of parametric uncertainties moving horizon

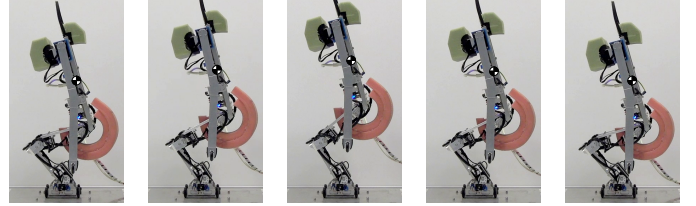


Fig. 1: Robot Leo performs up and down motions. Root point is shown by the circle with black and white sectors.

estimation techniques [3] can often be employed, for structural uncertainties, such as backlash, Coulomb friction or wear and tear, this is not easily possible. Nevertheless, model-based methods can predict the evolution of the system for a short horizon, thus enabling the implementation of safety barriers to limit risky exploration in dangerous state space regions.

By safety, we mean the prevention of actions that cause damage to the system. For example, in a bipedal robot, safety is particularly related to the robot not falling. Falls result in impact forces applied to the limbs and gearboxes, and some robots cannot withstand even a single fall. In model-based control, it is natural to constrain the angles and velocities to remain within an admissible range, and to enforce static stability constraints such that the center of mass projection is some distance away from the support polygon border. These constraints help prevent falls, but a momentary violation of them does not necessarily result in the one. In RL, it is possible to consider angle and velocity constraints by means of negative rewards. However, to learn avoiding such constraints, they need to be violated multiple times in different robot configurations. Random exploration exacerbates the problem and can lead to a very large number of falls.

Therefore, we propose to combine RL and nonlinear model predictive control (NMPC) in one framework that allows RL to gather the required experience without damaging a many-degree-of-freedom system. The experience is used by RL to compensate the difference between the internal model of the system and the real one. Any model-based nominal controller is suitable, but the choice of NMPC is particularly motivated by the complexity of the robot.

This paper proposes two different approaches shown in Figure 2. Similarly to [4], the first approach learns a compensatory control action, but instead of a proportional-derivative (PD) controller, we use NMPC, which introduces an additional optimization problem. Since both NMPC and RL optimize similar performance measures, the obtained policy is optimal with respect to the real system. The second approach learns a compensatory signal from the difference of transitions predicted by the internal model and the actual transition. In this case,

Manuscript received: September, 10, 2017; Revised December, 9, 2017; Accepted January, 12, 2018.

This paper was recommended for publication by Editor Dongheui Lee upon evaluation of the Associate Editor and Reviewers' comments. This work was supported by the European project (KoroiBot) FP7-ICT-2013-10/611909 and ROBOTIS which provided motors for Leo. The authors would like to thank A.V. Proskurnikov for his insightful comments.

¹Robotics Institute, TU Delft, Netherlands {i.koryakovskiy, h.vallery, r.babuska}@tudelft.nl

²Interdisciplinary Center for Scientific Computing, Heidelberg University, Germany manuel.kudruss@iwr.uni-heidelberg.de

³Department of Electrical Engineering, Pontifical Catholic University of Rio de Janeiro, Brazil wouter@caarls.org

Digital Object Identifier (DOI): see top of this page.

RL uses a different optimization goal, which does not divert NMPC from reaching its objective. As a result, the model-plant mismatch is eliminated by forcing the real system to behave as if it has no uncertainties.

We conduct simulated and real experiments with Leo and demonstrate the advantage of our proposal in the presence of temperature- and torque-dependent Coulomb friction.

II. RELATED WORK

From a control theoretic viewpoint, our approaches can be compared to adaptive internal model control (IMC) [5]. The implementation requires an explicit model of the plant to be used as part of the controller. However, in adaptive IMC, the structure of the unknown system is determined offline, while its parameters can be inferred by online parameter estimation [3]. A particular shortcoming is that the structure needs to be identified precisely; otherwise a model-plant mismatch remains. The proposed approaches require neither precise identification of the structure nor of the parameters.

As a learning controller, we employ model-free on-policy deterministic policy gradient (DPG) [6]. In principle, any model-free RL algorithm such as [7]–[10] can be used. However, lack of safety measures and sample complexity of the algorithms limits their application to real systems.

Learning the forward model of the system demonstrates the lowest number of interactions with it [11], [12]. Learning the inverse model [13], [14] assumes that the model can connect successive states prescribed by the nominal controller.

When the approximate model of the system is available, it is possible to pre-train the initial policy, which can speed up learning. The two-step sequential approach is proposed in [15]. First, an iterative linear-quadratic-Gaussian algorithm is used to design an initial policy. Then, the policy is refined using the PI^2 algorithm on the real system. Another approach is to iteratively learn the difference model of the measured state and the state obtained on the approximate model and adopt this difference model for improving the policy [16], [17]. Finally, the authors in [18] use an ensemble of slightly perturbed model parameters to learn a robust policy.

Learning involving off-line planning or human-expert demonstrations [1], [19]–[22] constrains the problem space, thus reducing hardware damage. This option requires either a hand-coded suboptimal policy or a skilled human operator.

For a bipedal robot where any failure can be catastrophic, the above methods are not suited even given a good starting policy, because it is likely that RL will result in at least several failures during subsequent policy improvement episodes.

It is possible to guarantee safe learning when one can either predict repercussions of bad actions [23] or has a backup policy to lead the system back to safe states [24], [25]. In this article, we do not guarantee safe learning, though the proposed approaches in practice can be safe.

Our contribution is twofold. First, we propose approaches that can in principle compensate any type of uncertainty which preserves the Markov property, without a time-consuming structure identification process and expert-designed models of friction, backlash, etc. The approaches can be implemented

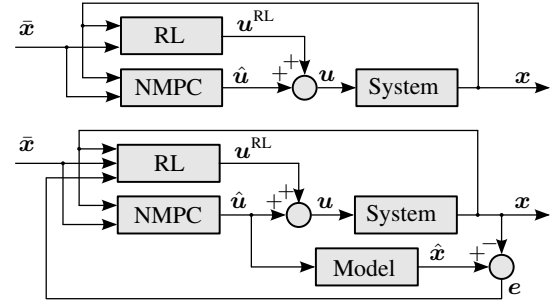


Fig. 2: Compensatory action learning (*top*). Model-plant mismatch learning (*bottom*).

on top of an existing model-based controller which makes it easier to integrate into the various fields of engineering, such as robotics, chemistry or computer science. Second, we demonstrate successful learning results on a real robot.

III. BACKGROUND

A. Problem statement

Consider the nonlinear time-invariant system in the form of

$$\dot{x}(t) = f(x(t), u(t), \rho) \quad (1)$$

where $x(t) \in \mathbb{R}^{n_x}$ the system state at time t , $u(t) \in \mathbb{R}^{n_u}$ is the control vector of joint motor voltages applied to the system at time t , and ρ is an unknown structural uncertainty. The presence of uncertainty causes the model-plant mismatch e which is formulated as the difference between the real system state x and the simulated state of the model $\hat{x}$, see Figure 2. We do not make any assumption on how the uncertainty enters the equations. Thus it represents the general concept of mismatch.

B. Nonlinear model predictive control

The nominal feedback is provided by NMPC, a closed-loop control strategy in which the control action is computed from the current system state by solving an open-loop optimal control problem on a finite prediction horizon $[0, T]$ online,

$$\min_{x(\cdot), u(\cdot)} \int_0^T L(x(t), u(t)) dt \quad (2a)$$

$$\text{s.t. } \dot{x}(t) = f(x(t), u(t), 0), \quad (2b)$$

$$\begin{aligned} x(0) &= x_0, \\ g(x(t), u(t)) &\geq 0. \end{aligned} \quad (2c)$$

Here we strive to find a control trajectory $u(t)$ such that an objective function composed of a Lagrange term $L : \mathbb{R}^{n_x} \times \mathbb{R}^{n_u} \rightarrow \mathbb{R}$ is minimized. The state trajectory $x(t)$ is characterized by the dynamic system (1). In (2b), we assume the idealized model, $\rho = 0$, because nothing is known about the uncertainties in the real system. In addition, we impose mixed state-control path constraints $g : \mathbb{R}^{n_x} \times \mathbb{R}^{n_u} \rightarrow \mathbb{R}^{n_g}$ such as constraints on joint angles ensuring the static stability of the robot together with constraints formulating limits on the maximum motor input voltage.

In the simulated experiment, we use a nominal NMPC scheme that is based on direct multiple shooting [26]. Controls u are approximated as piecewise constant functions. The cost

of some discretized trajectory $\mathbf{x}_0, \mathbf{u}_0, \mathbf{x}_1, \mathbf{u}_1, \dots$ obtained using policy $\mathbf{u}_k = \boldsymbol{\pi}(\mathbf{x}_k)$ is denoted as $\mathcal{L} = \sum_k L(\mathbf{x}_k, \mathbf{u}_k)$.

To achieve real-time control on the robot, we implement a parallelized version of the NMPC scheme [27], where one controller provides a fast feedback by efficiently reusing control problem linearizations of the last iteration, while the second controller prepares the next nonlinear step.

C. Reinforcement learning

RL is a trial-and-error method which does not require an explicitly given model, and can naturally adapt to uncertainties in the real system [28]. RL assumes the system is stochastic, and thus it maximizes the expected discounted return

$$G_k^\gamma = \mathbb{E} \left\{ \sum_{i=0}^{\infty} \gamma^i r(\mathbf{x}_{k+i}, \mathbf{u}_{k+i}, \mathbf{x}_{k+i+1}) \right\} \quad (3)$$

$$r(\mathbf{x}_k, \mathbf{u}_k, \mathbf{x}_{k+1}) = \begin{cases} -L(\mathbf{x}_{k+1}, \mathbf{u}_k) & \text{if } g(\mathbf{x}_{k+1}, \mathbf{u}_k) \geq 0 \\ R_a & \text{otherwise.} \end{cases} \quad (4)$$

Here $r(\mathbf{x}_k, \mathbf{u}_k, \mathbf{x}_{k+1})$ is the scalar reward given for a transition from $\mathbf{x}_k$ to $\mathbf{x}_{k+1}$ caused by the control signal $\mathbf{u}_k = \boldsymbol{\pi}(\mathbf{x}_k) + \mathbf{n}$, $\mathbf{n} \sim \mathcal{N}$ chosen from some policy $\boldsymbol{\pi}$. Discount rate $\gamma \in [0, 1]$ is required for integrability of the infinite sum. Its role is similar to NMPC horizon T . Constraints (2c) are established by means of the large negative reward R_a . Subsequently, the episode is terminated, and the system is restarted in state $\mathbf{x}_0$. Usually, RL requires at least several repetitions to estimate the return (3) correctly. These repetitions are obtained by adding exploration noise $\mathcal{N}$ to control signals $\mathbf{u}_k$ at every time step. The outcome of repetitions is not known in advance and therefore may be damaging to the system.

An important aspect of learning is the preservation of the Markov property, which assumes that the next state $\mathbf{x}_{k+1}$ depends only on the current state $\mathbf{x}_k$ and action $\mathbf{u}_k$, but not on previous states or actions [28].

We solve problem (3) using DPG with linear function approximation and compatible features, chosen for its ability to optimize continuous control policies and fast convergence.

IV. PROPOSED COMBINATION APPROACHES

A. Compensatory action learning (CAL)

In the proposed combination schemes shown in Figure 2, we use $\hat{\mathbf{u}}$ notation for the output of the NMPC controller and $\mathbf{u}^{\text{RL}}$ notation for the output of the RL controller.

CAL approach learns a compensatory control action added to the control input computed by nominal NMPC. For learning, we use the NMPC-inspired reward (4), which establishes similar optimization goals for both controllers. Due to small differences in formulation and function approximations in RL, the obtained policy might be suboptimal compared to NMPC.

B. Model-plant mismatch learning (MPML)

In the MPML approach, RL maximizes return (3) where the reward is given by

$$r(\mathbf{x}_k, \mathbf{u}_k, \mathbf{x}_{k+1}) = -\|\mathbf{e}_{k+1}\|_2 = -\|\mathbf{x}_{k+1} - \hat{\mathbf{x}}_{k+1}\|_2. \quad (5)$$

Cumulative model-plant mismatch along some trajectory is calculated as the negative value of the undiscounted return $\mathcal{E} = \sum_k \|\mathbf{e}_k\|_2$.

In the following, we prove two theorems. The first one explains the behavior of the system when the model-plant mismatch is minimized by RL. The subsequent corollary considers the case when the cumulative return ($\gamma > 0$) is useful for discovering a better control policy. The second theorem specifies conditions under which the system retains the Markov property. If the property is preserved, then RL will not diverge, and the mismatch can be minimized. In this case, the performance depends on the RL learning capability.

Theorem 1. *The outcome of the control policy approaches the outcome of the optimal policy with respect to the idealized model iff the model-plant mismatch $\mathbf{e}_k \rightarrow 0$ when $k \rightarrow \infty$.*

Proof. Writing the mismatch as $\mathbf{e}_k = \mathbf{x}_k - \hat{\mathbf{x}}_k \rightarrow 0$ results in $\mathbf{x}_k \rightarrow \hat{\mathbf{x}}_k$. Assuming the nominal controller can reach the setpoint on the model, $\hat{\mathbf{x}}_k \rightarrow \bar{\mathbf{x}}_k$, implies that the system state will also approach the setpoint, $\mathbf{x}_k \rightarrow \bar{\mathbf{x}}_k$. Since the mismatch $\mathbf{e}_k$ is minimized in every point of the reference trajectory $\bar{\mathbf{x}}_k$, we arrive at the proof of the theorem. The same logic holds for the reverse. $\square$

Corollary 1. *If there is a time step k such that $\mathbf{e}_k \neq 0 \forall \mathbf{u}_k$, then $\min_{\boldsymbol{\pi}}|_{\gamma=0} \mathcal{E} \geq \min_{\boldsymbol{\pi}}|_{\gamma>0} \mathcal{E}$. When the mismatch is eliminated along the reference trajectory, strict equality holds.*

The corollary is based on the RL result that larger γ improves the quality of the policy [29]. However, if there exists a control action which achieves zero mismatch, then maximizing immediate rewards ($\gamma = 0$) is desirable because the problem becomes computationally easier.

In the following theorem, we assume that the system (1) can be discretized as $\mathbf{x}_{k+1} = f(\mathbf{x}_k, \mathbf{u}_k, \boldsymbol{\rho})$ and the setpoint $\bar{\mathbf{x}}$ can be included into the state $\mathbf{x}$ for simplicity.

Theorem 2. *The system controlled by the nominal controller is Markov w.r.t. RL if (a) the system itself is Markov w.r.t. RL, $\mathbf{x}_{k+1} = f(\mathbf{x}_k, \mathbf{u}_k, \boldsymbol{\rho})$; (b) the internal model is Markov w.r.t. the nominal controller, $\hat{\mathbf{x}}_{k+1} = f(\mathbf{x}_k, \hat{\mathbf{u}}_k, 0)$; and (c) the nominal controller response $\hat{\mathbf{u}}_k \equiv \hat{\mathbf{u}}(\mathbf{x}_k) + \mathbf{m}$, $\mathbf{m} \sim \mathcal{M}$ is stochastic with some stationary distribution $\mathcal{M}$.*

Proof. First, by looking at the bottom diagram of Figure 2 we write the condition (a) and show that the distribution of states $\mathbf{x}_{k+1}$ is defined by the current state $\mathbf{x}_k$

$$\begin{aligned} \mathbf{x}_{k+1} &= f(\mathbf{x}_k, \hat{\mathbf{u}}_k + \mathbf{u}_k^{\text{RL}}, \boldsymbol{\rho}) \\ &= f(\mathbf{x}_k, \hat{\mathbf{u}}(\mathbf{x}_k) + \boldsymbol{\pi}(\mathbf{x}_k) + \mathbf{m} + \mathbf{n}, \boldsymbol{\rho}). \end{aligned}$$

Next, we show that the reward is also defined by $\mathbf{x}_k$.

$$\begin{aligned} r(\mathbf{x}_k, \mathbf{u}_k, \mathbf{x}_{k+1}) &= -\|\mathbf{x}_{k+1} - \hat{\mathbf{x}}_{k+1}\|_2 \\ &= -\|\mathbf{x}_{k+1} - f(\mathbf{x}_k, \hat{\mathbf{u}}(\mathbf{x}_k) + \mathbf{m}, 0)\|_2 \end{aligned}$$

The expected return averages the sum of discounted rewards over the distribution of states and controls. Since both the dynamics and the return are predictable from the current state $\mathbf{x}_k$, we conclude that the system controlled by the nominal controller is Markov w.r.t. RL. $\square$

Note that the real system does not have to be Markov with respect to NMPC, which means that any uncertainty ρ can be compensated. We use this observation in the real example where ρ depends on motor temperature which does not enter the model, but is used in RL as an extra state variable.

In the special case of an affine system (1) w.r.t. controls, $\mathbf{x}_{k+1} = f^x(\mathbf{x}_k, \rho) + f^u(\mathbf{x}_k, \rho)\mathbf{u}_k$, a perfectly learned RL control, i.e., $e = 0$, is explicitly given by

$$\mathbf{u}_k^{\text{RL}} = (f^u(\mathbf{x}_k, \rho)^\top f^u(\mathbf{x}_k, \rho))^{-1} f^u(\mathbf{x}_k, \rho)^\top [f^x(\mathbf{x}_k, 0) - f^x(\mathbf{x}_k, \rho) + (f^u(\mathbf{x}_k, 0) - f^u(\mathbf{x}_k, \rho)) \hat{\mathbf{u}}_k],$$

where $f^x(\mathbf{x}_k, \rho) \in \mathbb{R}^{n_x \times 1}$ and $f^u(\mathbf{x}_k, \rho) \in \mathbb{R}^{n_u \times n_x}$ are terms independent of $\mathbf{u}_k$. As expected, the control $\mathbf{u}^{\text{RL}}$ captures the model-plant mismatch caused by ρ .

V. EXPERIMENTS

A. Robot Leo

Robot Leo is depicted in Figure 1. We focus on the task of reaching upper and lower setpoints which together realize a squatting motion. For this purpose, the robot is fixed to the ground plate below its feet. Falling situation is recognized when absolute torso angle becomes larger than 57.3° . Leo has seven degrees of freedom driven by Dynamixel XM430 servo motors, three in each leg at ankle, knee and hip and one motor in the shoulder. Gearboxes are subject to Coulomb friction dependent on motor temperature and torque.

The robot state $\mathbf{x} = (\phi, \dot{\phi}, p, \tau_{\text{knee}})^\top$ is defined as a vector of all but shoulder joint angles ϕ , corresponding angular velocities $\dot{\phi}$, setpoint height $p \in \{0.28 \text{ m}, 0.35 \text{ m}\}$ and mean temperature of knee motors τ_{knee} .

Exploiting the symmetry of Leo, we apply the same control voltages to both legs. The shoulder is actuated using a PD controller. The setpoints are switched over when the robot root point appears to be within $\pm 0.01 \text{ m}$ away from it. In the simulated experiment, we add shoulder angle and velocity to the state, and the shoulder voltage is also learned. The robot is initialized in a setpoint chosen randomly at the beginning of every episode. The idealized model is made such that no friction in joints is present. For the realistic model, we add Coulomb friction $\mathbf{u}_{\text{fr}} = -0.2 \tanh(2000\dot{\phi})$ in all joints.

The control delay of $13.0 \pm 1.7 \text{ ms}$ comprises measurement, computation and actuation delays. A sampling period of 33.3 ms is chosen to be larger than the control delay.

B. Objective function and constraints

The NMPC objective function is defined by (2a) with

$$L(\mathbf{x}, \mathbf{u}) = 0.05 (h(\phi) - p)^2 + 0.10 (x_c(\phi) - \bar{x}_c)^2 + 0.05 (\text{pose}(\phi) - 0.3)^2 + 0.003 \dot{\phi}^\top \dot{\phi}.$$

Here, the first term accounts for the vertical distance $h(\phi)$ to a setpoint p , the second term maintains the horizontal position of the center of mass $x_c(\phi)$ close to predefined value $\bar{x}_c$, the third term containing $\text{pose}(\phi) = \phi_{\text{ankle}} + \phi_{\text{knee}} + \phi_{\text{hip}}$ is used as a regularization term improving the stability of the robot, and the last term favors small velocities.

We formulate static stability as a constraint $\mathbf{g}(\mathbf{x}, \mathbf{u}) = (x_t - x_c(\phi), x_c(\phi) - x_h)^\top$, where x_t, x_h denote the position of the tip and the heel of robot feet. Additionally, robot angles and controls are subject to constraints

$$\begin{bmatrix} -1.57 \\ -2.53 \\ -0.61 \end{bmatrix} \leq \phi_i \leq \begin{bmatrix} 1.45 \\ -0.02 \\ 2.53 \end{bmatrix} \quad \{|u_i|, |\dot{u}_i|, |u_i^{\text{RL}}|\} \leq 10.0 \text{ V} \quad i \in \{\text{ankle, knee, hip}\}.$$

For MPML approach the reward (5) is calculated based on the joint angles $r(\mathbf{x}_k, \mathbf{u}_k, \mathbf{x}_{k+1}) = -\|\phi_{k+1} - \hat{\phi}_{k+1}\|_2$.

C. Parameters

The time horizon T for NMPC optimization is selected to be 1 s. One learning or testing episode lasts for 15 s. Advantage, critic and actor learning rates are chosen to be 0.01, 0.10 and 0.01, respectively. Additional parameters include discount rate $\gamma = 0.97$, and an eligibility trace decay rate of 0.65. We rely on NMPC to avoid falls of the robot, therefore negative reward R_a is not used. Exploration is achieved by Ornstein-Uhlenbeck noise $\Delta \mathbf{u}_{k+1} = 0.5 \Delta \mathbf{u}_k + \mathcal{N}(0, \sigma)$ with $\sigma = 0.005$. For the real experiment, we select a higher advantage learning rate and increase exploration.

D. Evaluation

For quantitative assessment, we evaluate objective (2a) separately for reaching upper and lower setpoints $\mathcal{L}^{\{\text{u,l}\}} = \sum L(\mathbf{x}, \mathbf{u})$. Second, we evaluate the minimization of model-plant mismatch $\mathcal{E}^{\{\text{u,l}\}} = \sum \|e\|_2$ for reaching both setpoints separately. Third, we calculate root mean squared error (RMSE) between transitions obtained by both approaches and NMPC executed on the idealized model. Finally, to experimentally demonstrate safety barriers imposed by NMPC, we calculate the cumulative number of falls and violation of NMPC constraints at multiple levels of exploration noise σ for two proposed approaches and DPG.

For qualitative assessment, we calculate the number of squats the robot performs during the testing episode. This measure should be accounted only as a learning progress indicator since it is not included in the optimization objective.

E. Simulation results

In Figure 3 and Table I, we demonstrate a significant difference in the performance of standalone NMPC on the idealized and realistic models. On the idealized model, NMPC realizes three squats. On the realistic model, NMPC can reach neither upper nor lower switching points, which results in the inability to squat and high costs $\mathcal{L}$. This result motivates the need for an adaptive component in the controller.

To compare the performance of the proposed approaches to the baseline performance of NMPC, we perform learning for 10^6 time steps. The time was enough for CAL and MPML to converge, while DPG required about hundred times more steps. Thus, its results were excluded from the comparison.

We notice that on the idealized model the performance of both approaches becomes slightly worse than the baseline performance of NMPC. For CAL $\gamma = 0.97$, we observe deviation

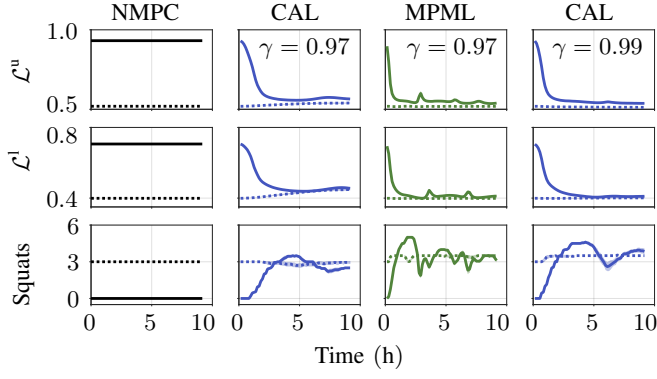


Fig. 3: Learning to reach upper (*top row*) and lower setpoints (*middle row*) in simulation. Number of squats is given in the *bottom row*. Dotted and solid lines show learning on the idealized and realistic models, respectively. Means with upper and lower 95% confidence limits are shown for 10 runs.

TABLE I: Final performance of methods. Significant width of the confidence interval is shown in brackets.

Method	$\mathcal{L}^u$ $\times 10$	$\mathcal{L}^l$ $\times 10$	Number of squats	RMSE $\times 10^4$
Idealized model				
NMPC	5.2	4.0	3.0	0.0
CAL $\gamma = 0.97$	5.4	4.4	3.0(0.1)	6.4(1.5)
MPML	5.2	4.0	3.5	1.7(0.1)
CAL $\gamma = 0.99$	5.1	4.0	3.5	11.8(4.0)
Realistic model				
NMPC	9.3	7.1	0	20.1
CAL $\gamma = 0.97$	5.6	4.6	2.5	10.4(0.5)
MPML	5.4	4.1	3.2(0.1)	4.3(1.2)
CAL $\gamma = 0.99$	5.4	4.1	3.9(0.2)	13.2(1.2)
Real robot				
NMPC 38°C	22.0(2.1)	—	0	—
MPML	7.9(2.9)	4.4(1.8)	3.3(1.1)	94.9(26.5)

from the optimal policy which is seen in the increase of $\mathcal{L}$ cost. Yet, the approach is able to keep the number of squats close to the baseline value. For MPML, costs $\mathcal{L}$ do not change, but the number of squats increases by 0.5 which indicates that the approach reaches the upper setpoint right before the episode ends. Deviation of the learned trajectory from the idealized one is captured by RMSE which is nonzero for both approaches. For the CAL and MPML approaches the mean RMSE is 68.6 % and 91.7 % below the reference of 20.4 ± 0.2 which is RMSE of NMPC trajectory obtained on the realistic model.

Results of the realistic model experiment show that both approaches improve the performance of NMPC. The decrease of the $\mathcal{L}$ cost is at least 35.2 % and 41.9 % for CAL $\gamma = 0.97$ and MPML approaches, respectively. In terms of squats, both learning approaches overshoot the NMPC baseline of 3 squats and then continuously reduce the number towards the baseline. RMSE increases comparing to the idealized model experiment, but still remains significantly below the reference value.

To find the reason of CAL $\gamma = 0.97$ performance decrease on the idealized model, we test the approach with a discount rate of $\gamma = 0.99$. Increasing γ leads to a longer planning horizon which makes RL return (3) more similar to NMPC objective (2a). It turns out that CAL $\gamma = 0.99$ obtains lower $\mathcal{L}$ costs not only comparing to CAL $\gamma = 0.97$ but also comparing

to the baseline NMPC. We believe the reason of this is due to the early switching of setpoints described above. While CAL $\gamma = 0.99$ can learn this fact, NMPC is not aware of it.

In Figure 4, we plot the number of falls and NMPC constraint violations accumulated over 10^6 time steps. Here, we prematurely stopped DPG for the sake of results comparability. The proposed approaches are almost identical. Both approaches prevent the robot from falling, while constraints get violated at $\sigma > 0.1$. A different picture is seen in DPG results. The smallest number of falls and constraint violations is achieved for the value of $\sigma = 0.02$. Smaller σ reduces the learning pace, while larger values increase chances of fall.

MPML learns almost twice as fast as CAL and does not exhibit deviating behavior, which are the main reasons for testing the approach on the real robot.

F. Results on the real robot

Results of standalone NMPC on Leo are shown in Figure 5. While on the idealized model NMPC successfully reaches switching points, on the real robot the controller is not able to do so. The reason is due to Coulomb friction in gearboxes, which depends on motor temperature and the applied torque. Modeling these effects is possible but requires a precise identification of the underlying physical processes.

To circumvent this problem, we apply the proposed MPML approach. In Figure 5, results of three independent runs are shown. MPML successfully realizes squatting by learning the compensation signal. The variation of motor temperature leads to noticeable differences in the trajectories. The trajectory obtained in the 3rd run is less noisy and squatting is faster than the one achieved in 1st and 2nd runs. In particular, the gradient of the downwards motion in the 3rd run is very similar to the idealized NMPC run, except for the later part where the slow approach towards the setpoint diminishes.

Variation of motor temperature also leads to differences in the learning progress, see Figure 6. 1st and 2nd runs require substantially longer time before the squatting cycle is observed. This is due to the increase of motor temperature above 40.0°C which requires additional exploration of the state space. Nevertheless, all runs successfully attain a stable squatting cycle after 7.25 h.

Model-plant mismatch $\mathcal{E}^u$ and $\mathcal{E}^l$ is minimized to about 0.5 and 0.8 for reaching upper and lower setpoints, respectively. As it was expected, minimization of model-plant mismatch leads to minimization of the nominal controller objective (2a) shown by plots $\mathcal{L}^u$ and $\mathcal{L}^l$. The smallest final costs are incurred by the 3rd run because it was stuck the least due to temperature fluctuations, while the largest costs are incurred by the 1st run which was stuck the most.

RMSE after learning is calculated in Table I. RMSE of MPML trajectories is much higher than for the realistic model, and it also exhibits more variability. Unfortunately, it is not possible to obtain the reference RMSE value of the real robot.

Figure 7 shows the MPML knee control signal and the RL compensation component of it. For reference, NMPC control on the idealized model is also shown. MPML controls are very oscillatory comparing to NMPC. Nonetheless, the robot

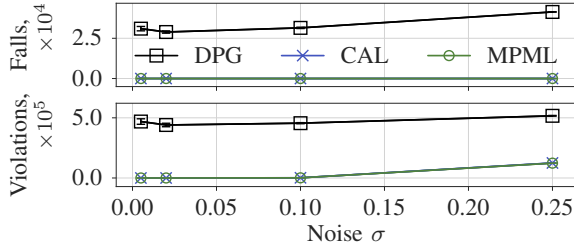


Fig. 4: The number of falls (*top*) and NMPC constraint violations (*bottom*) as functions of σ . Means with upper and lower 95% confidence limits are shown for 10 runs.

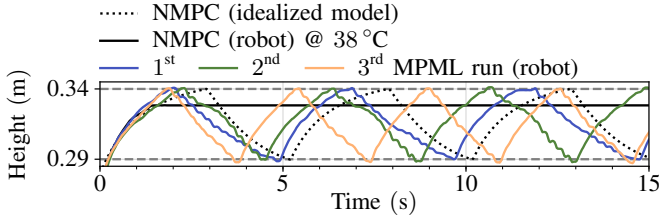


Fig. 5: Robot root point trajectories obtained after learning.

neither fell down, nor were its motors damaged, which is a significant result, cf. [1]. As is expected with Coulomb friction compensation, RL learned to apply positive and negative controls for the upward and downward motions, respectively.

VI. DISCUSSION

Even though the final CAL policy is optimal w.r.t. the real system, the policy is suboptimal w.r.t. the objective of the nominal controller (2a). This result can be explained by the fact that RL and NMPC objectives are not exactly same. While NMPC optimizes the undiscounted cost up to horizon T , RL optimizes γ -discounted reward on the infinite horizon. All in all, RL views the system and the nominal controller as a hybrid entity and the obtained policy becomes optimal w.r.t. the RL objective. This observation also explains the inability of $CAL^{\gamma=0.97}$ to reach the optimal performance on the realistic model, even though it can significantly improve the performance of the nominal controller. The longer prediction horizon used by $CAL^{\gamma=0.99}$ attains a better performance.

Another problem of CAL is the slow convergence which is probably caused by the fact that the reward constructed from the quadratic objective function of the nominal controller results in small gradients [2]. This hypothesis is supported by the fact that DPG with a quadratic cost function learns the task extremely slowly. To mitigate this, the RL cost function can be modified. The downside of it can be the difficulty of predicting the outcome of such modification, e.g. robot velocity may change drastically.

The MPML approach is free from these complications. However, it should be emphasized that MPML optimizes policy w.r.t. the internal model, that is RL forces the system to behave like the idealized model. In principle, a large mismatch may pose a problem because the obtained policy will be less optimal w.r.t. the real system and control constraints may prevent the necessary compensation to be applied. However, in our experiments, this is not a problem. Minimization of the model-plant mismatch $\mathcal{E}$ closely follows minimization of the nominal

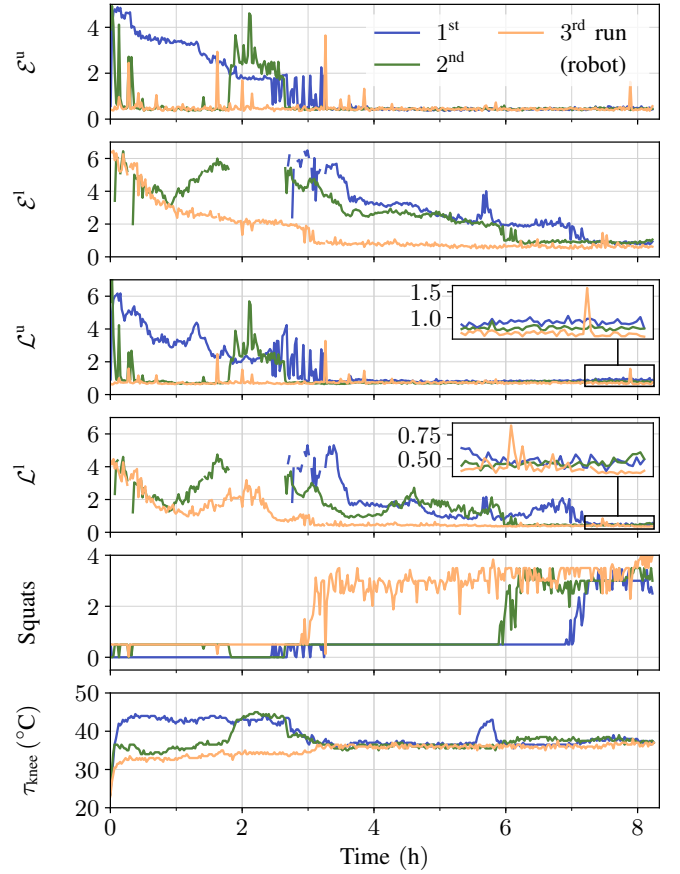


Fig. 6: Model-plant mismatch (*upper two*), nominal controller objective (*middle two*), number of squats and the mean temperature of knee motors obtained during three real learning experiments. Measurements of $\mathcal{E}^l$ and $\mathcal{L}^l$ are missing when upper switching point is not reached.

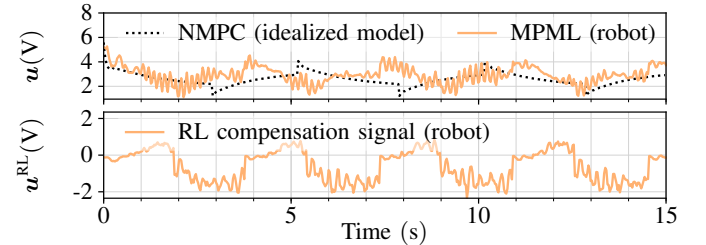


Fig. 7: Knee control signal of NMPC applied to the idealized model and of MPML applied to the real robot after learning (*top*). Compensation signal learned by MPML (*bottom*).

controller cost $\mathcal{L}$. MPML successfully learns to compensate the unknown Coulomb friction as well as its dependency on motor temperature and torque.

The MPML approach obtains the lowest RMSE. This does not come as a surprise, as MPML directly minimizes the mismatch by the specifically constructed reward function. CAL also minimizes RMSE, even though its primary goal is not defined in terms of such minimization. Arguably, in order for NMPC to successfully complete the task, the realistic model should resemble the idealized one which is achieved by learning with RL in both approaches.

For both proposed approaches, a little deviation caused by RL

exploration leads to an immediate setback reaction from NMPC. Our simulated experiment reveals a wide range of admissible exploration noise σ for which the number of NMPC constraint violations and robot falls is zero. This result demonstrates the role of NMPC which provides safety barriers to constrain RL exploratory actions near dangerous state space regions. However, there are disadvantages. First, the formulated task demands deliberate control learning which is difficult in the presence of Coulomb friction. If the robot starts moving after a slight overshoot caused by RL exploration, this immediately causes the decrease of friction (Stribeck effect), and at the next sampling moment, the system displacement appears to be too large. NMPC counteracts, so that resulting trajectories appear to be oscillatory. The other reason of oscillations is due to the large control delay. Given the results, it is hard to assess the role of NMPC counter-reaction in the oscillatory trajectories, but we expect that reduction of sampling time and control delay will reduce oscillations. Second, NMPC can drive the system very close to constraint boundaries. For some systems violation of constraints can be very critical, however, this is not true in our case.

We note that the success of model-plant mismatch compensation depends on the learning capabilities of RL on the hybrid system mentioned above. Whether there is a decrease or increase of computational complexity of that system against the original system remains an open problem.

It is common to compensate for a steady-state error in a task completion by adding an integral term to the objective that is tuned by experimental data. Learning the actual model-plant mismatch with MPML goes far beyond cost tuning. It allows to predict the outcome of executed actions since the learned trajectory is expected to be optimal w.r.t. the idealized model.

VII. CONCLUSION

We proposed two learning approaches to compensate model-plant mismatch. Our simulation results demonstrated the feasibility of both approaches. We implemented the better one on a real robot affected by torque and temperature dependent friction and autonomously learned a squatting task. Trying to achieve a similar performance with the standalone nominal controller would require tedious identification of the law of such a dependency. During learning, the robot did not fall.

Several avenues can be explored in future. First, one may reduce learning time by using different function approximators or RL algorithms. A wide range of recent model-free RL can be utilized in a straightforward way. Second, it is important to validate MPML on more challenging tasks. It is also a relatively straightforward implementation if the nominal controller is already set up, and RL actions are bounded.

REFERENCES

- [1] E. Schuitema, "Reinforcement learning on autonomous humanoid robots," Ph.D. dissertation, TU Delft, 2012.
- [2] I. Koryakovskiy, M. Kudruss, R. Babuska, W. Caarls, C. Kirches, K. Mombaur, J. P. Schlöder, and H. Vallery, "Benchmarking model-free and model-based optimal control," *Robotics and Autonomous Systems*, vol. 92, pp. 81–90, 2017.
- [3] P. Kühl, M. Diehl, T. Kraus, J. P. Schlöder, and H. G. Bock, "A real-time algorithm for moving horizon state and parameter estimation," *Computers & Chemical Engineering*, vol. 35, no. 1, pp. 71–83, 2011.
- [4] Y. E. Bayiz and R. Babuska, "Nonlinear disturbance compensation and reference tracking via reinforcement learning with fuzzy approximators," *IFAC Proceedings Volumes*, vol. 47, no. 3, pp. 5393–5398, 2014.
- [5] A. Datta and L. Xing, "The theory and design of adaptive internal model control schemes," in *American Control Conf.*, vol. 6, 1998, pp. 3677–3684.
- [6] D. Silver, G. Lever, N. Heess, T. Degris, D. Wierstra, and M. Riedmiller, "Deterministic policy gradient algorithms," in *Int. Conf. on Machine Learning*, 2014, pp. 387–395.
- [7] S. M. Kakade, "A natural policy gradient," in *Advances in Neural Information Processing Systems*, 2002, pp. 1531–1538.
- [8] J. Schulman, S. Levine, P. Moritz, M. I. Jordan, and P. Abbeel, "Trust region policy optimization," in *Int. Conf. on Machine Learning*, 2015.
- [9] J. Schulman, P. Moritz, S. Levine, M. Jordan, and P. Abbeel, "High-dimensional continuous control using generalized advantage estimation," in *Int. Conf. on Learning Representations*, 2016.
- [10] J. Schulman, F. Wolski, P. Dhariwal, A. Radford, and O. Klimov, "Proximal policy optimization algorithms," *CoRR*, vol. abs/1707.06347, 2017.
- [11] S. Kamthe and M. P. Deisenroth, "Data-efficient reinforcement learning with probabilistic model predictive control," *CoRR*, vol. abs/1706.06491, 2017.
- [12] S. Gu, T. Lillicrap, I. Sutskever, and S. Levine, "Continuous deep Q-learning with model-based acceleration," in *Int. Conf. on Machine Learning*, 2016, pp. 2829–2838.
- [13] M. Kawato, "Feedback-error-learning neural network for supervised motor learning," *Advanced neural computers*, vol. 6, no. 3, pp. 365–372, 1990.
- [14] P. Christiano, Z. Shah, I. Mordatch, J. Schneider, T. Blackwell, J. Tobin, P. Abbeel, and W. Zaremba, "Transfer from simulation to real world through learning deep inverse dynamics model," *CoRR*, vol. abs/1610.03518, 2016.
- [15] F. Farshidian, M. Neunert, and J. Buchli, "Learning of closed-loop motion control," in *IEEE/RSJ Int. Conf. on Intelligent Robots and Systems*, 2014, pp. 1441–1446.
- [16] S. Ha and K. Yamane, "Reducing hardware experiments for model learning and policy optimization," in *IEEE Int. Conf. on Robotics and Automation*, 2015, pp. 2620–2626.
- [17] M. Saveriano, Y. Yin, P. Falco, and D. Lee, "Data-efficient control policy search using residual dynamics learning," in *IEEE/RSJ Int. Conf. on Intelligent Robots and Systems*, 2017.
- [18] A. Rajeswaran, S. Ghotra, B. Ravindran, and S. Levine, "EPOpt: learning robust neural network policies using model ensembles," in *Int. Conf. on Learning Representations*, 2017.
- [19] P. Abbeel, A. Coates, and A. Y. Ng, "Autonomous helicopter aerobatics through apprenticeship learning," *Int. J. of Robotics Research*, vol. 29, no. 13, pp. 1608–1639, 2010.
- [20] J. Peters, K. Mlling, and Y. Altun, "Relative entropy policy search," in *AAAI Conf. on Artificial Intelligence*, 2010, pp. 1607–1612.
- [21] S. Levine and V. Koltun, "Guided policy search," in *Int. Conf. on Machine Learning*, 2013.
- [22] T. Zhang, G. Kahn, S. Levine, and P. Abbeel, "Learning deep control policies for autonomous aerial vehicles with MPC-guided policy search," in *IEEE Int. Conf. on Robotics and Automation*, 2016, pp. 528–535.
- [23] T. M. Moldovan and P. Abbeel, "Safe exploration in Markov decision processes," in *Int. Conf. on Machine Learning*, 2012, pp. 1711–1718.
- [24] A. Hans, D. Schneegass, A. M. Schafer, and S. Udluft, "Safe exploration for reinforcement learning," in *European Symposium on Artificial Neural Networks*, 2008.
- [25] J. F. Fisac, A. K. Akametalu, M. N. Zeilinger, S. Kaynama, J. H. Gillula, and C. J. Tomlin, "A general safety framework for learning-based control in uncertain robotic systems," *CoRR*, vol. abs/1705.01292, 2017.
- [26] H. G. Bock, M. Diehl, E. A. Kostina, and J. P. Schlöder, "Constrained optimal feedback control of systems governed by large differential algebraic equations," in *Real-Time PDE-Constrained Optimization*, L. Biegler, O. Ghattas, M. Heinkenschloss, D. Keyes, and B. van Bloemen Waanders, Eds. SIAM, 2007, ch. 1, pp. 3–24.
- [27] M. Kudruss, I. Koryakovskiy, H. Vallery, K. Mombaur, and C. Kirches, "Combining multi-level real-time iterations of nonlinear model predictive control to realize squatting motions on Leo," *Optimization Online*, Tech. Rep., 2018.
- [28] R. S. Sutton and A. G. Barto, *Reinforcement Learning: An Introduction*. MIT Press, 1998.
- [29] N. Jiang, A. Kulesza, S. Singh, and R. Lewis, "The dependence of effective planning horizon on model accuracy," in *Int. Conf. on Autonomous Agents and Multiagent Systems*, 2015.