



Delft University of Technology

DeepKoCo

Efficient latent planning with a task-relevant Koopman representation

van der Heijden, Bas; Ferranti, Laura; Kober, Jens; Babuska, Robert

DOI

[10.1109/IROS51168.2021.9636408](https://doi.org/10.1109/IROS51168.2021.9636408)

Publication date

2021

Document Version

Final published version

Published in

Proceedings of the IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS 2021)

Citation (APA)

van der Heijden, B., Ferranti, L., Kober, J., & Babuska, R. (2021). DeepKoCo: Efficient latent planning with a task-relevant Koopman representation. In *Proceedings of the IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS 2021)* (pp. 183-189). IEEE.
<https://doi.org/10.1109/IROS51168.2021.9636408>

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.

Green Open Access added to TU Delft Institutional Repository

'You share, we take care!' - Taverne project

<https://www.openaccess.nl/en/you-share-we-take-care>

Otherwise as indicated in the copyright section: the publisher is the copyright holder of this work and the author uses the Dutch legislation to make this work public.

DeepKoCo: Efficient latent planning with a task-relevant Koopman representation

Bas van der Heijden¹, Laura Ferranti¹, Jens Kober¹, Robert Babuška¹

Abstract—This paper presents DeepKoCo, a novel model-based agent that learns a latent Koopman representation from images. This representation allows DeepKoCo to plan efficiently using linear control methods, such as linear model predictive control. Compared to traditional agents, DeepKoCo learns task-relevant dynamics, thanks to the use of a tailored lossy autoencoder network that allows DeepKoCo to learn latent dynamics that reconstruct and predict only observed costs, rather than all observed dynamics. As our results show, DeepKoCo achieves a similar final performance as traditional model-free methods on complex control tasks, while being considerably more robust to distractor dynamics, making the proposed agent more amenable for real-life applications.

Index Terms—Model-based reinforcement learning, Koopman theory, model-predictive control

I. INTRODUCTION

From self-driving cars to vision-based robotic manipulation, emerging technologies are characterized by visual measurements of strongly nonlinear physical systems. Unlike in highly controlled lab environments where any measured change is likely relevant, cameras in real-world settings are notorious for mainly capturing task-irrelevant information, such as, the movement of other robots outside of a manipulator's workspace or cloud movements captured by the cameras of self-driving cars.

While Deep Reinforcement Learning (DRL) algorithms can learn to perform various tasks using raw images, they will require an enormous number of trials. Prior methods mitigate this by encoding the raw images into a lower-dimensional representation that allows for faster learning. However, these methods can be easily distracted by irrelevant dynamics [1]. This motivates data-driven methodologies that learn low-dimensional latent dynamics that are task-relevant and useful for control.

In the learning of latent dynamics for control, there is a trade-off between having an accurate dynamic model and one that is suitable for control. On one hand, latent dynamic models based on neural networks (NN) can provide accurate predictions over long horizons. On the other hand, their inherent nonlinearity renders them incompatible with efficient planning algorithms. Alternatively, one can choose to approximate the latent dynamics with a more restricted function approximation class to favor the use of efficient planning algorithms. In this respect, a promising strategy

is represented by the Koopman framework [2]. Loosely speaking, this framework allows one to map observations with nonlinear dynamics to a latent space where the *global dynamics* of the autonomous system are approximately linear (Koopman representation). This enables the use of powerful linear optimal control techniques in the latent space [2].

While the Koopman framework is promising, existing methods have fundamental limitations that must be addressed to fully exploit the benefits of this method for control applications. First, methods that identify Koopman representations from data were designed for prediction and estimation. These methods were later adapted for control. These adaptations, however, lead to limiting assumptions on the underlying dynamics, such as assuming the Koopman representation to be linear in the states and actions [3], [4], [5], [6], [7]. Second, these methods are *task agnostic*, that is, the models represent all dynamics they observe, whether they are relevant to the task or not. This focuses the majority of their model capacity on potentially task-irrelevant dynamics.

Therefore, we introduce *Deep Koopman Control* (DeepKoCo), that is, a model-based agent that learns a latent Koopman representation from raw pixel images and achieves its goal through planning in this latent space. The representation is (i) robust to task-irrelevant dynamics and (ii) compatible with efficient planning algorithms. We propose a lossy autoencoder network that reconstructs and predicts observed costs, rather than all observed dynamics, which leads to a representation that is task-relevant. The latent-dynamics model can represent continuously differentiable nonlinear systems and does not require knowledge of the underlying environment dynamics or cost function. We demonstrate the success of our approach on two continuous control tasks and show that our method is more robust to irrelevant dynamics than state-of-the-art approaches, that is, DeepMDP [8] and Deep Bisimulation for Control (DBC) [1].

II. RELATED WORK

1) *Koopman control*: Koopman theory has been used to control various nonlinear systems with linear control techniques, both in simulation [2], [9], [6] and in real-world robotic applications [4], [5]. Herein, [10], [2], [4] used a linear quadratic regulator (LQR), while [9], [3], [5], [6], [7] applied linear model predictive control (MPC). [9], [3], [5], [7] used data-driven methods that were derived from the Extended Dynamic Mode Decomposition (EDMD) [11] to find the Koopman representation. In contrast, [4], [2], [10] require prior knowledge of the system dynamics to hand-craft parts of the lifting function. Similar to [6], we rely on deep

*This work was supported by the European Union's H2020 project Open Deep Learning Toolkit for Robotics (OpenDR) under grant agreement No 871449.

¹Cognitive Robotics at the Faculty of 3mE, Delft University of Technology, The Netherlands. d.s.vanderheijden@tudelft.nl

learning to derive the Koopman representation for control. However, we do not assume the Koopman representation to be (bi-)linear in the states and actions and we show how our representation can be used to control systems that violate this assumption. Compared to existing methods, we propose an agent that learns the representation online in a reinforcement learning setting using high-dimensional observations that contain irrelevant dynamics.

2) *Latent planning*: Extensive work has been conducted to learn latent dynamics from images and use them to plan suitable actions [12], [13], [14]. [14] proposes a model-based agent that uses NNs for the latent dynamics and cost model. To find suitable action sequences, however, their method requires a significant computational budget to evaluate many candidate sequences. Alternatively, [12], [13] propose locally linear dynamic models, which allowed them to efficiently plan for actions using LQR. However, their cost function was defined in the latent space and required observations of the goal to be available. In contrast to our approach, all aforementioned methods are trained towards full observation reconstruction, which focuses the majority of their model capacity on potentially task-irrelevant dynamics.

3) *Relevant representation learning*: [8], [1] filter task-irrelevant dynamics by minimizing an auxiliary bisimulation loss. Similar to our approach, they propose learning latent dynamics and predicting costs. Their method, however, is limited to minimizing a single-step prediction loss, while we incorporate multi-step predictions. This optimizes our model towards accurate long-term predictions. [15], [16] also proposed training a dynamics model towards predicting the future sum of costs given an action sequence. However, their method focused on discrete control variables, while we focus on continuous ones.

III. PRELIMINARIES

This section briefly introduces the Koopman framework for autonomous and controlled nonlinear systems. A detailed description can be found in [2]. This framework is fundamental to the design of our latent model and control strategy.

A. Koopman eigenfunctions for autonomous systems

Consider the following autonomous nonlinear system $\dot{\mathbf{o}} = F(\mathbf{o})$, where the observations $\mathbf{o} \in \mathbb{R}^N$ evolve according to the smooth continuous-time dynamics $F(\mathbf{o})$. For such a system, there exists a lifting function $g(\cdot) : \mathbb{R}^N \rightarrow \mathbb{R}^n$ that maps the observations to a latent space where the dynamics are linear, that is,

$$\frac{d}{dt}g(\mathbf{o}) = \mathcal{K} \circ g(\mathbf{o}), \quad (1)$$

where $\mathcal{K}$ is the infinitesimal operator generator of Koopman operators K . In theory, K is infinite dimensional (i.e., $n \rightarrow \infty$), but a finite-dimensional matrix representation can be obtained by restricting it to an invariant subspace. Any set of eigenfunctions of the Koopman operator spans such a subspace. Identifying these eigenfunctions [9], [17] provides a set of intrinsic coordinates that enable global

linear representations of the underlying nonlinear system. A Koopman eigenfunction satisfies

$$\frac{d}{dt}\phi(\mathbf{o}) = K\phi(\mathbf{o}) = \lambda\phi(\mathbf{o}), \quad (2)$$

where $\lambda \in \mathbb{C}$ is the continuous-time eigenvalue corresponding to eigenfunction $\phi(\mathbf{o})$.

B. Koopman eigenfunctions for controlled systems

For controlled nonlinear system with action $\mathbf{a} \in \mathbb{R}^m$ and smooth continuous-time dynamics $\dot{\mathbf{o}} = \tilde{F}(\mathbf{o}, \mathbf{a})$, we follow the procedure in [2]. Given the eigenfunction $\phi(\mathbf{o}, \mathbf{a})$ augmented with $\mathbf{a}$ for the controlled system, we can take its time derivative and apply the chain rule with respect to $\mathbf{o}$ and $\mathbf{a}$, leading to

$$\frac{d}{dt}\phi(\mathbf{o}, \mathbf{a}) = \underbrace{\nabla_{\mathbf{o}}\phi(\mathbf{o}, \mathbf{a})\tilde{F}(\mathbf{o}, \mathbf{a})}_{\lambda\phi(\mathbf{o}, \mathbf{a})} + \nabla_{\mathbf{a}}\phi(\mathbf{o}, \mathbf{a})\dot{\mathbf{a}}, \quad (3)$$

where λ is now the eigenvalue that corresponds to eigenfunction $\phi(\mathbf{o}, \mathbf{a})$. Since $\dot{\mathbf{a}}$ can be chosen arbitrarily, we could set it to zero and instead interpret each action as a parameter of the Koopman eigenfunctions. Thus, for any given choice of parameter $\mathbf{a}$ the standard relationship in (2) is recovered in the presence of actions. A local approximation of the Koopman representation is obtained when $\dot{\mathbf{a}}$ is nonzero.

C. Identifying Koopman eigenfunctions from data

To facilitate eigenfunction identification with discrete data, (3) can be discretized with a procedure similar to [17]. The eigenvalues $\lambda_{\pm} = \mu \pm i\omega$ are used to parametrize block-diagonal $\Lambda = \text{diag}(J^{[1]}, J^{[2]}, \dots, J^{[P]}) \in \mathbb{R}^{2P \times 2P}$. For all P pairs of complex eigenvalues, the discrete-time operator Λ has a Jordan real block of the form

$$J(\mu, \omega) = e^{\mu\Delta t} \begin{bmatrix} \cos(\omega\Delta t) & -\sin(\omega\Delta t) \\ \sin(\omega\Delta t) & \cos(\omega\Delta t) \end{bmatrix}, \quad (4)$$

with sampling time Δt . The “forward Euler method” provides a discrete approximation of the control matrix, so that (3) can be discretized as

$$\varphi(\mathbf{o}_{k+1}, \mathbf{a}_{k+1}) = \Lambda\varphi(\mathbf{o}_k, \mathbf{a}_k) + \underbrace{\nabla_{\mathbf{a}_k}\varphi(\mathbf{o}_k, \mathbf{a}_k)}_{B_{\varphi_k}} \underbrace{\dot{\mathbf{a}}_k\Delta t}_{\Delta\mathbf{a}_k}. \quad (5)$$

Herein, the stacked vector $\varphi = (\phi^{[1]}, \phi^{[2]}, \dots, \phi^{[P]})$ comprises a set of P eigenfunctions with $\phi^{[j]} \in \mathbb{R}^2$ associated with complex eigenvalue pair $\lambda_{\pm}^{[j]}$ and Jordan block $J^{[j]}$. Subscript k corresponds to discretized snapshots in time. If we view the action increment $\Delta\mathbf{a}_k = \mathbf{a}_{k+1} - \mathbf{a}_k$ in (5) as the controlled input instead, we obtain a discrete *control-affine* Koopman eigenfunction formulation with *linear autonomous* dynamics for the original *non-control-affine nonlinear* system. In the next section, we show that (5) plays a central role in our latent model.

IV. LEARNING TASK-RELEVANT KOOPMAN EIGENFUNCTIONS

For efficient planning in the latent space, we propose to learn a latent dynamics model that uses Koopman eigenfunctions as its latent state. This section describes this model and how the Koopman eigenfunctions can be identified robustly, that is, in a way that the identified eigenfunctions remain unaffected by task-irrelevant dynamics that are expected to contaminate the observations.

A. Koopman latent model

We propose a *lossy* autoencoder that builds on the deep autoencoder in [17]. Compared to [17], our autoencoder enables control. To train the latent model, we provide the training objective that is to be minimized given a buffer D that contains observed sequences $\{t_k\}_{k=0}^T$ of a Markov decision process with tuples $t_k = (\mathbf{o}_k, \mathbf{a}_k, \Delta\mathbf{a}_k, c_k)$, where c_k are observed scalar costs. The proposed latent model is illustrated in Fig. 1 with more details below on the individual components of the architecture.

1) *Encoder*: The encoder φ is the approximate eigenfunction that maps an observation-action pair $(\mathbf{o}_k, \mathbf{a}_k)$ to the latent state $\mathbf{s}_k$. The encoder φ is parametrized by a neural network, defined as

$$\mathbf{s}_k = \varphi(\mathbf{o}_k, \mathbf{a}_k). \quad (6)$$

2) *Latent dynamics*: The latent state $\mathbf{s}_k$ approximates a Koopman eigenfunction, so the autonomous time evolution in the latent space is linear and dictated by Λ . Note here that $\mathbf{a}_k$ is part of the *augmented* latent state and we view the action increment $\Delta\mathbf{a}_k$ as the controlled variable that is determined by the policy. This leads to the dynamics model in (7), which we derived from (5). The Koopman operator Λ is parametrized by P complex conjugate eigenvalue pairs $\lambda_{\pm}^{[j]}$. We do not assume the latent dynamics to be linear in the control. Instead, the influence of $\Delta\mathbf{a}_k$ on the latent state varies and depends on the partial derivative of the encoder with respect to the action, i.e., the state-dependent matrix $B_{\varphi_k} = \nabla_{\mathbf{a}_k} \varphi(\mathbf{o}_k, \mathbf{a}_k) \in \mathbb{R}^{2P \times m}$.

$$\begin{bmatrix} \mathbf{s}_{k+1} \\ \mathbf{a}_{k+1} \end{bmatrix} = \begin{bmatrix} \Lambda & 0 \\ 0 & I \end{bmatrix} \begin{bmatrix} \mathbf{s}_k \\ \mathbf{a}_k \end{bmatrix} + \begin{bmatrix} B_{\varphi_k} \\ I \end{bmatrix} \Delta\mathbf{a}_k. \quad (7)$$

3) *Cost model*: The environment contains a cost function that produces a scalar cost observation c_k at every time-step. For planning in the latent space, we require a cost model $\hat{c}_k$ as a function of the latent state. This cost approximates the observed cost (i.e., $\hat{c}_k \approx c_k$). We adopt a latent state-dependent quadratic cost model to facilitate the use of fast planning algorithms (Section V). The entries of $C_{\mathbf{s}_k} \in \mathbb{R}^{1 \times 2P}$ are determined by a function $\psi(\mathbf{s}_k)$ that is parametrized by a neural network. The weights of ψ are initially unknown and must be learned together with the rest of the latent model. We assume that the cost of applying action $\mathbf{a}_k$ is known *a priori* and defined by matrix R . This leads to the cost model

$$\hat{c}_k = \|C_{\mathbf{s}_k} \mathbf{s}_k\|_2^2 + \mathbf{a}_k^\top R \mathbf{a}_k. \quad (8)$$

4) *Policy*: The action increment $\Delta\mathbf{a}_k$ is the controlled variable that is sampled from a probability distribution π , conditioned on the augmented latent state (i.e., $\Delta\mathbf{a}_k \sim \pi(\Delta\mathbf{a}_k | \mathbf{s}_k, \mathbf{a}_k)$). Even though the model is deterministic, we define the policy to be stochastic to allow for stochastic exploration. The policy will be specified further in Section V.

$$\Delta\mathbf{a}_k \sim \pi(\Delta\mathbf{a}_k | \mathbf{s}_k, \mathbf{a}_k) \quad (9)$$

5) *Decoder*: After learning the latent model we intend to plan over it, which involves a multi-step prediction. Given only the encoder (6), dynamics model (7), policy (9) and the current $(\mathbf{o}_k, \mathbf{a}_k)$ -pair, we would be limited to single-step predictions of $\mathbf{s}_{k+1}$ at run-time, because multi-step predictions $\mathbf{s}_{k+i}$ with $i > 1$ would require knowledge of future observations $\mathbf{o}_{k+i}$ to evaluate $B_{\varphi_{k+i}}$. Therefore, to make multi-step predictions, we introduce a decoder φ^{-1} (parametrized by a NN) in (10) that uses predicted latent states $\mathbf{s}_{k+i}$ to construct pseudo-observations $\hat{\mathbf{o}}_{k+i}$ that produce the same partial derivative as the true observation (i.e., $\nabla_{\mathbf{a}_k} \varphi(\varphi^{-1}(\mathbf{s}_k), \mathbf{a}_k) \approx \nabla_{\mathbf{a}_k} \varphi(\mathbf{o}_k, \mathbf{a}_k)$). Future values $\mathbf{a}_{k+i}$ do not pose a problem, because they can be inferred from the policy (9) and dynamics model (7).

$$\hat{\mathbf{o}}_k = \varphi^{-1}(\mathbf{s}_k). \quad (10)$$

6) *Image processor*: When the observations are raw pixel images $\mathbf{p}_k$, not all relevant information can be inferred from a single observation. To restore the Markov property, we pass the last d consecutive pixel images through a convolutional neural network Ω in (11), stack the output into a single vector, and consider that to be the observation $\mathbf{o}_k$ instead. In that case, the observed sequences consist of tuples $t_k = (\mathbf{p}_k, \dots, \mathbf{p}_{k-d+1}, \mathbf{a}_k, \Delta\mathbf{a}_k, c_k)$.

$$\mathbf{o}_k = \Omega(\mathbf{p}_k, \dots, \mathbf{p}_{k-d+1}), \quad (11)$$

B. Learning the latent model

Our latent model should have linear dynamics and be predictive of observed costs. These two high-level requirements lead to the following three losses which are minimized during training.

1) *Linear dynamics*: To ensure that the latent state is a valid Koopman eigenfunction, we regularize the time evolution in the latent space to be linear by using the following loss,

$$\text{Linear loss: } \mathcal{L}_{\text{lin}} = \frac{1}{T} \sum_{k=0}^{T-1} \|\varphi(\mathbf{o}_{k+1}, \mathbf{a}_{k+1}) - \mathbf{s}_{k+1}\|_{\text{MSE}}, \quad (12)$$

where $\mathbf{s}_{k+1}$ is obtained by rolling out a latent trajectory as illustrated in Fig. 1.

2) *Cost prediction*: We want the latent representation to contain all necessary information to solve the task. If we would naively apply an autoencoder that predicts future observations, we focus the majority of the model capacity on potentially task-irrelevant dynamics contained in the observations. To learn a latent representation that *only* encodes relevant information, we propose to use a lossy autoencoder

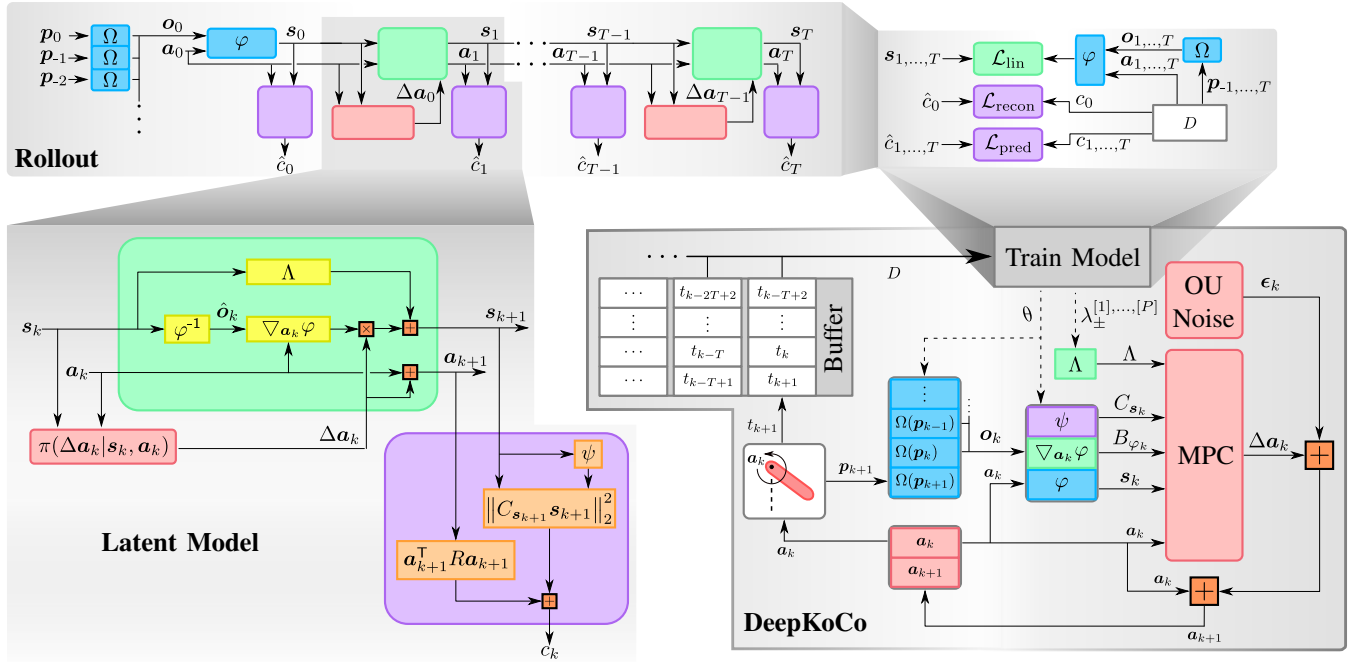


Fig. 1: **Latent Model** The proposed network architecture of the *latent model*, consisting of the dynamic model, cost model, and policy (depicted in green, purple, and red, respectively). **Rollout** A multi-step ahead prediction with the *latent model*. Note that we only encode an observation at the first time-step (blue boxes), after which we remain in the latent space. **DeepKoCo** The training procedure that corresponds to Algorithm 1.

that is predictive of current and future costs instead. Such a representation would allow an agent to predict the cost evolution of various action sequences and choose the sequence that minimizes the predicted cumulative cost, which is essentially equivalent to solving the task. Because we only penalize inaccurate cost predictions, the encoder is not incentivized to encode task-irrelevant dynamics into the latent representation as they are not predictive of the cost. This leads to the task-relevant identification of the lifting function φ . Cost prediction accuracy is achieved by using the following two losses,

$$\text{Reconstruction loss: } \mathcal{L}_{\text{recon}} = \|c_0 - \hat{c}_0\|_{\text{MSE}} \quad (13)$$

$$\text{Prediction loss: } \mathcal{L}_{\text{pred}} = \frac{1}{T} \sum_{k=1}^T \|c_k - \hat{c}_k\|_{\text{MSE}} \quad (14)$$

3) *Training objective*: We minimize the losses in (12), (13), and (14), corresponding to linear dynamics regularization and cost prediction, together with an L_2 -regularization loss $\mathcal{L}_{\text{reg}}$ on the trainable variables (excluding neural network biases). This leads to the following training objective,

$$\min_{\theta, \lambda_{\pm}^{[1], \dots, [P]}} \mathcal{L}_{\text{lin}} + \alpha_1 (\mathcal{L}_{\text{recon}} + \mathcal{L}_{\text{pred}}) + \alpha_2 \mathcal{L}_{\text{reg}}, \quad (15)$$

where θ is the collection of all the trainable variables that parametrize the encoder φ , decoder φ^{-1} , cost model ψ , and convolutional network Ω (in case of image observations). Weights α_1, α_2 are hyperparameters. The model is trained using the Adam optimizer [18] with learning rate $\tilde{\alpha}$, on batches of B sequences $\{t_k\}_{k=0}^T$ for E epochs.

V. DEEP KOOPMAN CONTROL

This section introduces the agent that uses the Koopman latent model to find the action sequence that minimizes the predicted cumulative cost. We use linear model-predictive control (LMPC) to allow the agent to adapt its plan based on new observations, meaning the agent re-plans at each step. Re-planning at each time-step can be computationally costly. In the following, we explain how to exploit and adapt our latent model and cost model to formulate a sparse and convex MPC problem that can be solved efficiently online.

The planning algorithm should achieve competitive performance, while only using a limited amount of computational resources. This motivates choosing Koopman eigenfunctions as the latent state, because the *autonomous* dynamics are linear. The dynamics are affine in the controlled variable Δa_k that is multiplied in the definition of the state space by B_{φ_k} , which depends on the latent state. Similarly, C_{s_k} requires the evaluation of the nonlinear function $\psi(s_k)$. There exist methods that can be applied in this setting, such as the State-Dependent Ricatti Equation (SDRE) method [19]. While the SDRE requires less complexity compared to sample-based nonlinear MPC (e.g. CEM [20]), it remains computationally demanding as it also requires the derivative of ψ with respect to s_k at every step of the planning horizon.

Our goal is to reduce the online complexity of our planning strategy, while also dealing with input constraints. Hence, we trade-off some prediction accuracy (due to the mismatch between the latent model and the MPC prediction model) to simplify the online planning strategy by using linear MPC.

We propose to evaluate the state-dependent matrices C_{s_0} and B_{φ_0} at time-step $k = 0$ (obtained from our latent model) and keep them both fixed for the rest of the LMPC horizon. This assumes that the variation of B_{φ_k} and C_{s_k} is limited over the prediction horizon (compared to (7) and (8)). Nevertheless, thanks to this simplification we can rely on LMPC for planning that can be solved efficiently. Specifically, once we evaluate s_0 , C_{s_0} , and B_{φ_0} , the computational cost of solving the MPC problem in the *dense form* [21] scales linearly with the latent state dimension due to the diagonal structure of Λ . As Section VI details, this simplification allows our method to achieve competitive final performance, while only requiring a single evaluation of the NNs Ω , φ , and ψ . This significantly decreases the computational cost at run-time compared to sample-based nonlinear MPC (e.g., CEM [20]) that would require many evaluations of the NNs at every time-step. In contrast to LQR, LMPC can explicitly deal with actuator saturation by incorporating constraints on $\mathbf{a}$. The proposed planning strategy based on LMPC is defined as follows:

$$\begin{aligned} \min_{\Delta \mathbf{a}^{[0], \dots, [H-1]}} \quad & \sum_{k=1}^H \|C_{s_0} \mathbf{s}_k\|_2^2 + \mathbf{a}_k^\top R \mathbf{a}_k + \Delta \mathbf{a}_k^\top \tilde{R} \Delta \mathbf{a}_k, \quad (16) \\ \text{s.t.} \quad & \begin{bmatrix} \mathbf{s}_{k+1} \\ \mathbf{a}_{k+1} \end{bmatrix} = \begin{bmatrix} \Lambda & 0 \\ 0 & I \end{bmatrix} \begin{bmatrix} \mathbf{s}_k \\ \mathbf{a}_k \end{bmatrix} + \begin{bmatrix} B_{\varphi_0} \\ I \end{bmatrix} \Delta \mathbf{a}_k, \\ & \mathbf{a}^{\min} \leq \mathbf{a}_k \leq \mathbf{a}^{\max}, \text{ for } k = 1, \dots, H, \end{aligned}$$

where H is the prediction horizon. Positive-definite matrix $\tilde{R}$ penalizes the use of $\Delta \mathbf{a}_k$ and is required to make the problem well-conditioned. Its use does introduce a discrepancy between the approximate cost model (8) and the cumulative cost ultimately minimized by the agent (16). Therefore, the elements in $\tilde{R}$ are kept as low as possible.

To align the representation learning objective (15) with the linear MPC objective (16), we also fix the state-dependent terms C_{s_0} and B_{φ_0} at time-step $k = 0$ in the evaluation of the cost prediction loss (14) and linear loss (12). Note that this does not mean that C_{s_k} and B_{φ_k} are constant in the latent model (7), (8). The matrices remain state-dependent, but their variation is limited over the sequence length T . In general, we choose the sequence length to be equal to the prediction horizon. Hence, we learn a representation that provides local linear models that are particularly accurate around $\mathbf{s}_k$ in the direction of the (goal-directed) trajectories gathered during training.

To gather a rich set of episodes to learn the Koopman latent model, we add colored noise to the actions commanded by the agent's linear MPC policy, that is, $\mathbf{a}_{k+1} = \mathbf{a}_k + \Delta \mathbf{a}_k + \epsilon_k$. This adds a stochastic exploration component to the policy. We use an Ornstein-Uhlenbeck (OU) process to generate the additive colored noise with decay rate λ^{ou} . The variance $\sigma^{2, \text{ou}}$ is linearly annealed from $\sigma_{\text{init}}^{2, \text{ou}} \rightarrow 0$ over $1, \dots, N^{\text{ou}}$ episodes, that is, after N^{ou} episodes the policy becomes deterministic.

An overview of the proposed method is shown in Algorithm 1. First, we initialize all model parameters. Then, we construct the Koopman operator Λ with (4) and gather N

Algorithm 1: Deep Koopman Control (DeepKoCo)

Input: Model parameters: P, d
Policy parameters: $\zeta = \{H, R, \tilde{R}\}$
Noise parameters: $\lambda^{\text{ou}}, \sigma_{\text{init}}^{2, \text{ou}}, N^{\text{ou}}$
Train parameters: N, L, T, E, B
Optimization parameters: $\xi = \{\alpha_1, \alpha_2, \tilde{\alpha}\}$

Output: Eigenvalues $\lambda_{\pm}^{[1], \dots, [P]}$
Trained networks $\varphi, \varphi^{-1}, \psi, \Omega$

```

1  $\theta, \lambda_{\pm}^{[1], \dots, [P]} \leftarrow \text{InitializeModel}(P, R)$ 
2 while not converged do
3    $\Lambda \leftarrow \text{GetKoopmanOperator}(\lambda_{\pm}^{[1], \dots, [P]})$ 
4   for episode  $l = 1, \dots, N$  do
5      $\mathbf{p}_0, \dots, \mathbf{p}_{1-d} \leftarrow \text{ResetEnvironment}()$ 
6      $\mathbf{a}_0 \leftarrow 0$ 
7     for time-step  $k = 0, \dots, L$  do
8        $\mathbf{o}_k \leftarrow \text{ProcessImages}(\mathbf{p}_k, \dots, \mathbf{p}_{k-d+1}, \theta)$ 
9        $\mathbf{s}_k, B_{\varphi_k}, C_{s_k} \leftarrow \text{LatentModel}(\mathbf{o}_k, \mathbf{a}_k, \theta)$ 
10       $\Delta \mathbf{a}_k \leftarrow \text{LMPC}(\mathbf{s}_k, \mathbf{a}_k, B_{\varphi_k}, C_{s_k}, \Lambda, \zeta)$ 
11       $\mathbf{a}_{k+1} \leftarrow \mathbf{a}_k + \Delta \mathbf{a}_k + \text{Noise}(\lambda^{\text{ou}}, \sigma^{2, \text{ou}})$ 
12       $\mathbf{p}_{k+1}, c_k \leftarrow \text{ApplyAction}(\mathbf{a}_k)$ 
13     $D \leftarrow D \cup \text{CreateSequences}(T, \{t_k\}_{k=0}^L)$ 
14   $\theta, \lambda_{\pm}^{[1], \dots, [P]} \leftarrow \text{TrainModel}(D, \theta, \lambda_{\pm}^{[1], \dots, [P]}, \xi, E, B)$ 

```

episodes of experience. Each time-step, we process the image observations with (11), evaluate the latent model (6), (7), and (8), and use it to find $\Delta \mathbf{a}_k$ with (16). Noise is added to the action increment before it is applied to the environment. We fill the experience buffer D with N episodes, split into sequences of length T , and train on them for E epochs with (15). This is repeated until convergence.

VI. RESULTS

We evaluate *DeepKoCo* on two continuous control tasks, namely OpenAI's pendulum swing-up task and a manipulator task. The manipulator task is similar to OpenAI's reacher task where the two joints of a two-link arm are torque controlled and the Euclidean distance between the arm's end-point and a target must be minimized. However, we increase the difficulty by allowing the target to move at a constant angular velocity and radius around the arm's center-joint. Given that the angular velocity and radius vary randomly over episodes, the manipulator must learn to track arbitrary circular trajectories at different speeds. The dynamics for the manipulator can be formulated as $\mathbf{x}_{k+1} = F(\mathbf{x}_k) + B(\mathbf{x}_k)\mathbf{a}_k$, where $\mathbf{x}$ is the original nonlinear state. Such dynamics do not necessarily admit a Koopman representation that is (bi-)linear in the states and actions, as is often assumed in literature [6], [4].

To investigate the effect of distractor dynamics, we test each task in two different scenarios. In the first scenario, only relevant dynamics are observed, while in the second one we purposely contaminate the observations with distractor dynamics. The agent is either provided with a concatenated state observation containing the state measurements of both the relevant system and distractors (while not knowing which

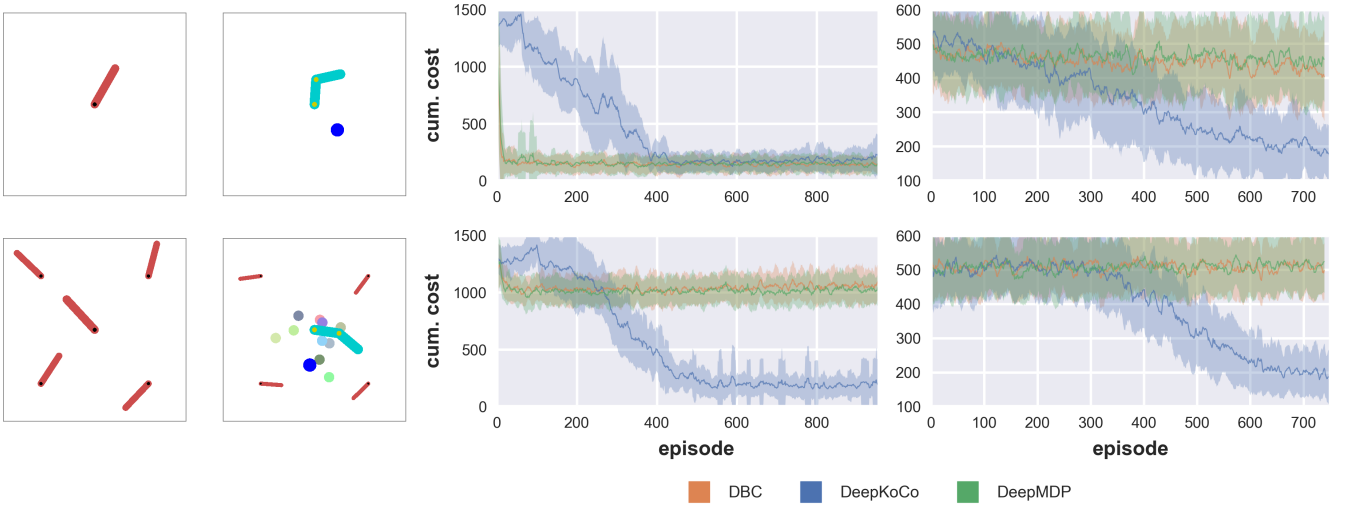


Fig. 2: **Left** Typical setups for the two tasks in a clean scenario (first row) and distractor scenario (second row). In all setups, the center system is controlled by the agent. In the manipulator task, the moving target is a blue ball. **Right** Learning curves when using state observations. The grid-location of each figure corresponds to the grid-location of each setup on the left. The mean cumulative cost over the last 10 episodes (line) with one standard deviation (shaded area) over 5 random seed runs are shown.

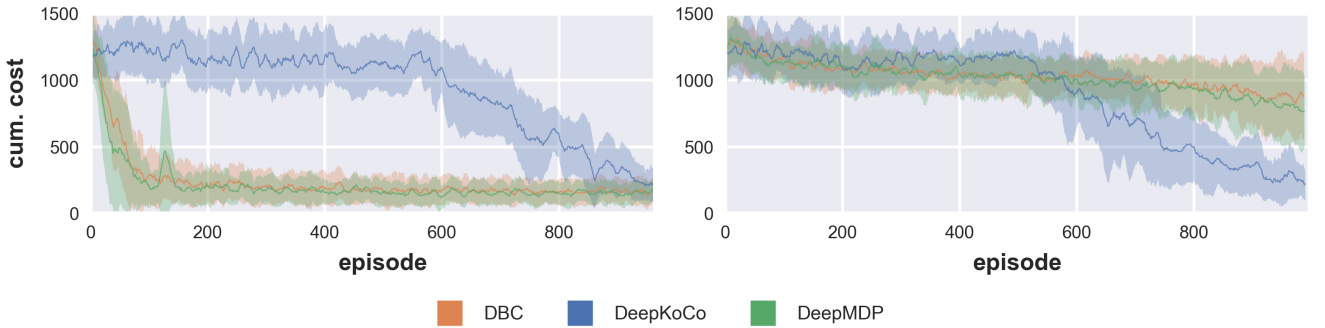


Fig. 3: Learning curves when using images as observations in the pendulum task. The mean cumulative cost over the last 10 episodes (line) with one standard deviation (shaded area) over 5 random seed runs are shown. **Left** Clean scenario. **Right** Distractor scenario.

states are the relevant ones) or image observations with all systems in-frame (refer to Fig. 2 for the setup). The state observation dimension from the clean to the distractor scenario increases from 3 to 15 for the pendulum and from 10 to 50 for the manipulator. A video of the simulations using DeepKoCo accompanies the paper [22].

1) *Baselines*: We compare with two baselines that both combine model-free reinforcement learning with an auxiliary bisimulation loss to be robust against task-irrelevant dynamics: (i) Deep Bisimulation for Control (DBC) [1], and (ii) DeepMDP [8]. In case of state observations, we replace their convolutional encoder, with our fully connected encoder.

2) *Hyperparameters*: We use the same set of hyperparameters throughout all experiments, except for the number of complex eigenvalue pairs P , that is, $P = 10$ and $P = 30$ in the swing-up and manipulator task, respectively to cope with the complexity of the scenarios. As policy parameters, we use $H = 15$, $R = 0.001$, $\tilde{R} = 0.01$, $\lambda^{\text{OU}} = 0.85$, $\sigma_{\text{init}}^{2,\text{ou}} = 0.85$. Initially, we fill the experience buffer D with $N = 90$ episodes, split into sequences of length $T = 15$, and train

on them for $E = 100$ epochs. Then, we continuously add the sequences of $N = 20$ episodes to the buffer and train on the complete buffer for $E = 3$ epochs. As optimization parameters, we use $\alpha_1 = 10$, $\alpha_2 = 10^{-14}$, $\tilde{\alpha} = 0.001$. In case of image observations, we stack the last $d = 3$ images, downsample them to $3 \times 64 \times 64$ pixels before passing them through the convolutional NN defined in [23]. The networks $\varphi, \varphi^{-1}, \psi$ are 2-layered fully connected NNs with 90, 90, and 70 units per layer, respectively. The layers use ReLU activation and are followed by a linear layer. The number of complex eigenvalue pairs P , planning horizon H , and action increment cost $\tilde{R}$ are the most important parameters to tune.

3) *Clean scenario*: Concerning the pendulum task, the baselines converge more quickly to the final performance compared to DeepKoCo, as the top-left graph of Fig. 2 shows. Nevertheless, we do consistently achieve a similar final performance. The slower convergence can be explained by the added noise, required for exploration, that is only fully annealed after 400 episodes. We believe the convergence rate can be significantly improved by performing a parameter

search together with a faster annealing rate, but we do not expect to be able to match the baselines in this ideal scenario. Note that, despite the apparent *simplicity* of the application scenario, finding an accurate Koopman representation for the pendulum system is challenging, because it exhibits a continuous eigenvalue spectrum and has multiple (unstable) fixed points [17]. Concerning the manipulator task, both baselines were not able to solve the manipulator task with a moving target as the top-right graph of Fig. 2 shows, while they were able to learn in case the target was fixed (results omitted due to space limitations). This shows that learning to track arbitrary circular references is significantly harder than regulating towards a fixed goal. Despite the increased difficulty, DeepKoCo learns to track the moving target. Finally, note that the manipulator task shows that the proposed method can deal with a multi-dimensional action-space and non-quadratic cost functions, that is, the Euclidean norm (the square root of an inner product).

4) *Distractor scenario*: In the more realistic scenario, both baselines fail to learn anything. In contrast, our approach is able to reach the same final performance as in the clean scenario, in a comparable amount of episodes, as the bottom row of Fig. 2 shows. This result can be explained by noticing that our multi-step cost prediction (in our loss function (14)) provides a stronger learning signal for the agent to distinguish relevant from irrelevant dynamics. For the manipulator task, there is a tracking error caused by the trade-off of using fixed B_φ and C_s along the MPC prediction horizon for efficiency. While our latent model presented in Section IV supports state-dependent matrices, we decided to keep them fixed in the control design for efficiency.

5) *Image observations*: Fig. 3 shows the results for the pendulum task when images are used instead of state observations. In both clean and distractor scenarios, our approach is able to reach a similar final performance compared to using state observations. As expected, the baselines struggle to learn in the distractor scenario. This supports our statement that our approach learns a task-relevant Koopman representation from high-dimensional observations. We plan to test the manipulator task with images both in simulation and in real-world experiments.

VII. CONCLUSIONS AND FUTURE WORK

We presented a model-based agent that uses the Koopman framework to learn a latent Koopman representation from images. DeepKoCo can find Koopman representations that (i) enable efficient linear control, (ii) are robust to distractor dynamics. Thanks to these features, DeepKoCo outperforms the baselines (two state-of-the-art model-free agents) in the presence of distractions. As part of our future work, we will extend our deterministic latent model with stochastic components to deal with partial observability and aleatoric uncertainty. Furthermore, we will extend our cost model to deal with sparse rewards, as they are often easier to provide.

REFERENCES

- [1] A. Zhang, R. McAllister, R. Calandra, Y. Gal, and S. Levine, "Learning invariant representations for reinforcement learning without reconstruction," *arXiv preprint arXiv:2006.10742*, 2020.
- [2] E. Kaiser, J. N. Kutz, and S. L. Brunton, "Data-driven discovery of Koopman eigenfunctions for control," in *APS Division of Fluid Dynamics Meeting Abstracts*, ser. APS Meeting Abstracts, Nov. 2017, p. M27.006.
- [3] H. Arbabi, M. Korda, and I. Mezić, "A data-driven koopman model predictive control framework for nonlinear partial differential equations," *IEEE Conference on Decision and Control (CDC)*, pp. 6409–6414, 2018.
- [4] G. Mamakoukas, M. Castano, X. Tan, and T. Murphey, "Local koopman operators for data-driven control of robotic systems," in *Proceedings of Robotics: Science and Systems (RSS)*, 2019.
- [5] D. Bruder, B. Gillespie, C. D. Remy, and R. Vasudevan, "Modeling and control of soft robots using the koopman operator and model predictive control," in *Proceedings of Robotics: Science and Systems (RSS)*, 2019.
- [6] Y. Li, H. He, J. Wu, D. Katabi, and A. Torralba, "Learning compositional koopman operators for model-based control," in *International Conference on Learning Representations (ICLR)*, 2020.
- [7] M. Korda and I. Mezić, "Optimal construction of koopman eigenfunctions for prediction and control," *IEEE Transactions on Automatic Control*, 2020.
- [8] C. Gelada, S. Kumar, J. Buckman, O. Nachum, and M. G. Bellemare, "DeepMDP: Learning continuous latent space models for representation learning," *36th International Conference on Machine Learning (ICML)*, pp. 3802–3826, 2019.
- [9] M. Korda and I. Mezić, "Linear predictors for nonlinear dynamical systems: Koopman operator meets model predictive control," *Automatica*, vol. 93, pp. 149–160, 2018.
- [10] S. L. Brunton, B. W. Brunton, J. L. Proctor, and J. N. Kutz, "Koopman invariant subspaces and finite linear representations of nonlinear dynamical systems for control," *PloS one*, vol. 11, no. 2, 2016.
- [11] M. O. Williams, I. G. Kevrekidis, and C. W. Rowley, "A Data-Driven Approximation of the Koopman Operator: Extending Dynamic Mode Decomposition," *Journal of Nonlinear Science*, vol. 25, no. 6, pp. 1307–1346, 2015.
- [12] M. Watter, J. Springenberg, J. Boedecker, and M. Riedmiller, "Embed to control: A locally linear latent dynamics model for control from raw images," in *Advances in neural information processing systems (NIPS)*, 2015, pp. 2746–2754.
- [13] E. Banijamali, R. Shu, H. Bui, A. Ghodsi *et al.*, "Robust locally-linear controllable embedding," in *International Conference on Artificial Intelligence and Statistics*. PMLR, 2018, pp. 1751–1759.
- [14] D. Hafner, T. Lillicrap, I. Fischer, R. Villegas, D. Ha, H. Lee, and J. Davidson, "Learning latent dynamics for planning from pixels," in *International Conference on Machine Learning (ICML)*. PMLR, 2019, pp. 2555–2565.
- [15] J. Oh, S. Singh, and H. Lee, "Value prediction network," in *Advances in Neural Information Processing Systems (NIPS)*, 2017, pp. 6118–6128.
- [16] J. Schrittwieser, I. Antonoglou, T. Hubert, K. Simonyan, L. Sifre, S. Schmitt, A. Guez, E. Lockhart, D. Hassabis, T. Graepel, T. Lillicrap, and D. Silver, "Mastering Atari, Go, Chess and Shogi by Planning with a Learned Model," *arXiv preprint arXiv:1911.08265*, pp. 1–21, 2019.
- [17] B. Lusch, J. N. Kutz, and S. L. Brunton, "Deep learning for universal linear embeddings of nonlinear dynamics," *Nature Communications*, vol. 9, no. 1, pp. 1–10, dec 2018.
- [18] D. P. Kingma and J. Ba, "Adam: A method for stochastic optimization," *arXiv preprint arXiv:1412.6980*, 2014.
- [19] I. Chang and J. Bentsman, "Constrained discrete-time state-dependent riccati equation technique: A model predictive control approach," in *52nd IEEE Conference on Decision and Control (CDC)*. IEEE, 2013, pp. 5125–5130.
- [20] I. Szita and A. Lőrincz, "Learning tetris using the noisy cross-entropy method," *Neural computation*, vol. 18, no. 12, pp. 2936–2941, 2006.
- [21] J. M. Maciejowski, *Predictive control: with constraints*. Pearson education, 2002.
- [22] B. van der Heijden, L. Ferranti, J. Kober, and R. Babuška, "Supplementary video material of DeepKoCo simulations," <https://youtu.be/751OuQyHBMQ>, July 2021.
- [23] D. Ha and J. Schmidhuber, "World Models," *arXiv preprint arXiv:1803.10122*, 2018.