

Stochastic Optimal Control as Approximate Input Inference

Joe Watson, Hany Abdulsamad, Jan Peters[†]

Department of Computer Science, Technische Universität Darmstadt, Germany

[†]Robot Learning Group, Max Planck Institute for Intelligent Systems, Tübingen, Germany

{watson, abdulamad, peters}@ias.informatik.tu-darmstadt.de

Abstract: Optimal control of stochastic nonlinear dynamical systems is a major challenge in the domain of robot learning. Given the intractability of the global control problem, state-of-the-art algorithms focus on approximate sequential optimization techniques, that heavily rely on heuristics for regularization in order to achieve stable convergence. By building upon the duality between inference and control, we develop the view of Optimal Control as Input Estimation, devising a probabilistic stochastic optimal control formulation that iteratively infers the optimal input distributions by minimizing an upper bound of the control cost. Inference is performed through Expectation Maximization and message passing on a probabilistic graphical model of the dynamical system, and time-varying linear Gaussian feedback controllers are extracted from the joint state-action distribution. This perspective incorporates uncertainty quantification, effective initialization through priors, and the principled regularization inherent to the Bayesian treatment. Moreover, it can be shown that for deterministic linearized systems, our framework derives the maximum entropy linear quadratic optimal control law. We provide a complete and detailed derivation of our probabilistic approach and highlight its advantages in comparison to other deterministic and probabilistic solvers.

Keywords: Stochastic Optimal Control, Approximate Inference

1 Introduction

Trajectory Optimization for nonlinear dynamical systems is among the most fundamental paradigms in the field of robotics. It has proven itself to be a cornerstone for both low- and high-level planning techniques [1, 2]. A popular tool for devising such planning schemes is Optimal Control [3, 4], which frames the search for the best sequence of inputs into a dynamical system as the optimization of the state-action trajectory. While Optimal Control has had great success both in theory and application, mainly represented by Sequential Quadratic Programming (SQP) techniques [5], it is known to struggle with stochastic environments due to its feedforward nature. Meanwhile, a popular tool for dealing with uncertainty is Bayesian statistics [6], which in part uses the notion of random variables to describe model uncertainty. The process of determining the characteristics of this uncertainty is known as inference, and this too is often framed as an optimization problem. Control-as-inference [7, 8, 9, 10] is a body of research combining these two paradigms, with the proposition that the principled mechanisms of inference will bring the benefits of faster convergence, more principled regularization and the addition of uncertainty quantification [11].

In this work we present Input Inference for Control (I2C), a new perspective on control-as-inference. By moving away from the typical Optimal Control formulation, while preserving the underlying operations, recursive Bayesian inference can be applied to the inputs to manner that optimizes a control objective. This builds on previous work that performs recursive approximate inference of the state trajectory [12] and exact input inference for linear systems [13]. Consider the fundamental task of control: to find the sequence of actions that generate a desired trajectory. From an inference perspective, we would call this problem Input Estimation, where the ‘desired’ observed trajectory in this case is a set of measurements. In Optimal Control, as the desired trajectory is not expected to be fully achieved, the notion of a cost function is used to describe the desired deviation of the observed trajectory. Statistically this deviation would be framed as a ‘disturbance’ and described

by a probability distribution. As likelihoods are often the optimization objective of an inference problem, by comparing the likelihood of this formulation to typical control cost functions informs our choice of disturbance noise in order to achieve equivalence. In this work, we focus on the well-established duality between Gaussian noise and quadratic penalties [6].

By making the linear Gaussian assumption on both our dynamics and observation models, inference can be performed in closed-form using message passing, and we show that this input inference reduces to the Linear Quadratic Regulator (LQR) solution in the deterministic case. Moreover, the inference is in fact performing the same Discrete Algebraic Ricatti equation (DARE) computation [13]. Additionally, making the inference approximate through local linearizations, we extend the scheme to nonlinear dynamical systems and arrive at a procedure akin to the popular trajectory optimization of Differential Dynamic Programming (DDP) [14] and variants (e.g. iLQR [15], eLQR[16], GPS [17]). While these methods require explicit regularization, bounds and heuristics to maintain steady convergence, the behaviour of our scheme is governed primarily by the choice of priors, and the regularization only required to account for the log-likelihood approximation. The use of Bayesian inference also results in self-regularized exploration, as the covariance of each input is a measure of confidence / robustness. Moreover, by examining the conditional distributions between the resultant posterior state-action distribution, we arrive at (Bayes) optimal time-varying linear (Gaussian) controllers, as in LQR [13]. We show that the covariance of these controllers naturally exhibits the maximum entropy characteristic, achieved without explicit incorporation of a policy entropy term in the objective as done previously.

The contributions of this work are as follows:

A **control-as-inference formulation (I2C)** that posits optimal control as input estimation for a dynamical system, such that the optimization objective is separated from the priors over the controls. This allows for Bayesian inference of the controls, rather than fixing them for exploration.

A **practical realisation** through approximate Expectation Maximisation, performing inference via linearized Gaussian message passing in the E-Step and hyperparameter optimization in the M-Step. Compared to previous methods, I2C has more principled regularization, relying primarily on the priors rather than heuristic methods such as line search, smoothing and annealing.

2 Input Inference for Control

Given a stochastic discrete-time fully-observed nonlinear dynamical system, $\mathbf{x}_{t+1} \sim \mathbf{f}(\mathbf{x}_t, \mathbf{u}_t)$ with state $\mathbf{x} \in \mathbb{R}^{d_x}$ and input $\mathbf{u} \in \mathbb{R}^{d_u}$, we wish to find the optimal control inputs $\mathbf{u}_{0:T}^*$ over time horizon T that minimizes the cost function $C(\mathbf{x}, \mathbf{u})$ for moving from an initial state $\mathbf{x}_0$ to goal state $\mathbf{x}_g$.

Our proposed method reframes optimal control as inference of the inputs of the dynamical system. This can be achieved with access to a dynamics model and by incorporating the cost function into the likelihood in an affine manner through an ‘observation model’ $p(\mathbf{z}_t | \mathbf{x}_t, \mathbf{u}_t)$ of our optimization variables $\mathbf{z} \in \mathbb{R}^{d_z}$, such that $\alpha C(\mathbf{x}, \mathbf{u}) + \beta = \log p(\mathbf{z}_t | \mathbf{x}_t, \mathbf{u}_t)$. By maximizing this likelihood

$$\max_{\mathbf{u}_{0:T}, \boldsymbol{\theta}} p(\mathbf{z}_{0:T}, \mathbf{x}_{0:T}, \mathbf{u}_{0:T}, \boldsymbol{\theta}) = p(\mathbf{x}_0) \prod_{t=0}^{T-1} p(\mathbf{x}_{t+1} | \mathbf{x}_t, \mathbf{u}_t) \prod_{t=0}^T p(\mathbf{z}_t | \mathbf{x}_t, \mathbf{u}_t, \boldsymbol{\theta}) p(\mathbf{u}_t | \mathbf{x}_t), \quad (1)$$

both the control cost (observation likelihood) and trajectory likelihood are jointly optimized, generating an estimated optimal state-action joint distribution $p(\mathbf{x}, \mathbf{u})$. From this, the conditional distribution $p(\mathbf{u} | \mathbf{x})$ can be found and used as a policy. The likelihood acts as an unconstrained control cost function by incorporating the constraint of the dynamical system, present in typical Optimal Control formulations, as an additional likelihood. This makes sense for stochastic systems, where the dynamical system can no longer be treated as a deterministic constraint. The likelihood also depends on hyperparameters $\boldsymbol{\theta}$, which can be optimized via the marginal likelihood.

2.1 The Linear Gaussian Assumption

By applying the linear Gaussian assumption to the models and their respective uncertainties, Equation 1 can not only be tackled in a tractable manner, but also compared to LQR control (Section A). Firstly, we can express the conditionals as linear state-space models,

$$\text{Dynamics: } p(\mathbf{x}_{t+1} | \mathbf{x}_t, \mathbf{u}_t) : \quad \mathbf{x}_{t+1} = \mathbf{A}_t \mathbf{x}_t + \mathbf{B}_t \mathbf{u}_t + \mathbf{a}_t + \boldsymbol{\eta}_t, \quad \boldsymbol{\eta}_t \sim \mathcal{N}(\mathbf{0}, \boldsymbol{\Sigma}_{\boldsymbol{\eta}_t}), \quad (2)$$

$$\text{Cost: } p(\mathbf{z}_t | \mathbf{x}_t, \mathbf{u}_t) : \quad \mathbf{z}_t = \mathbf{E}_t \mathbf{x}_t + \mathbf{F}_t \mathbf{u}_t + \mathbf{e}_t + \boldsymbol{\xi}_t, \quad \boldsymbol{\xi}_t \sim \mathcal{N}(\mathbf{0}, \boldsymbol{\Sigma}_{\boldsymbol{\xi}_t}). \quad (3)$$

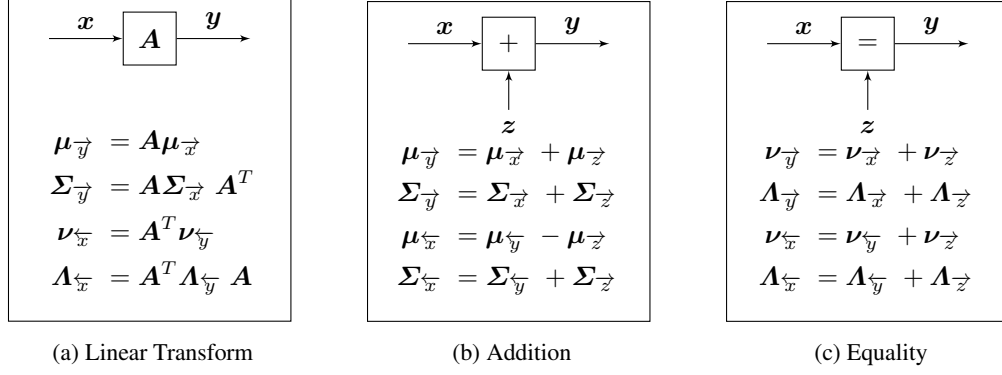


Figure 1: Linear Gaussian Message Passing rules for elementary state-space operations [18], with the mean (μ), covariance (Σ), precision ($\Lambda = \Sigma^{-1}$) and scaled mean ($\nu = \Lambda\mu$), which describe the moment and information (or canonical) form of the Normal distribution respectively.

Secondly, the log-likelihood is transformed into a convex function (Equation 4) which is quadratic in the optimization variables x , u and z [19].

$$\begin{aligned}
-\mathcal{L}(\theta) = & \frac{1}{2} \sum_{t=0}^{T-1} \log |\Sigma_{\eta_t}| + \frac{1}{2} \sum_{t=0}^T (z_t - E_t x_t - F_t u_t - e_t)^\top \Sigma_{\xi}^{-1} (z_t - E_t x_t - F_t u_t - e_t) \\
& + \frac{T}{2} \log |\Sigma_{\xi}| + \frac{1}{2} \sum_{t=0}^{T-1} (x_{t+1} - A_t x_t - B_t u_t - a_t)^\top \Sigma_{\eta_t}^{-1} (x_{t+1} - A_t x_t - B_t u_t - a_t) + \dots \quad (4)
\end{aligned}$$

In i2C, the ‘measurement’ of z represents the desired state-action trajectory. Therefore to transform the log-likelihood of z to a quadratic control cost, the precision of the ‘observation noise’ ξ is $\Sigma_{\xi}^{-1} = \Lambda_{\xi} = \alpha \Theta$, where Θ represents the weights of the cost function and α accounts for its scale invariance. For the standard LQ problem (Section A), $z_t = [x_g \ u_g]^\top$ and $\Theta = \text{diag}(\mathbf{Q}, \mathbf{R})$. Our hyperparameters θ include α , the scale factor, along with the priors over the inputs u . In Equation 4, α acts as the scale factor of the LQ cost against the other terms in the likelihood. Typically for multi-objective cost functions this scaling must be user-defined, but as it has a probabilistic interpretation here, it can be iteratively estimated during inference. As α scales the given control cost Θ such that it can be used as the observation noise precision Λ_{ξ} , it can be estimated based on the current estimated state-action trajectory deviation about the goal. This inference is carried out using the Expectation Maximization (EM) algorithm [20], treating α as a latent variable.

Expectation Step

The E-Step, estimating the state-action trajectory, can be performed in a tractable manner through linear Gaussian message passing. For model-based signal processing on linear Gaussian state space models [18, 21, 22], expressing inference problems as Forney-style factor graphs enables the construction of message-passing algorithms by following straightforward rules (see Figure 1). For cycle-free graphs, the messages can be expressed in closed-form. The forward messages (i.e. $\vec{x}$) represents the priors, while the backward messages (i.e. $\overleftarrow{x}$) represent likelihood functions (up to a scale factor). The updated belief is the posterior of an edge, which are the product of the edge’s forward and backward message:

$$\Sigma_x = (\Lambda_{\vec{x}} + \Lambda_{\overleftarrow{x}})^{-1}, \quad \mu_x = \Sigma_x (\nu_{\vec{x}} + \nu_{\overleftarrow{x}}). \quad (5)$$

In i2C, the backward messages perform optimal control, so the posterior states and controls represent a regularized update of the estimated optimal state-action trajectory.

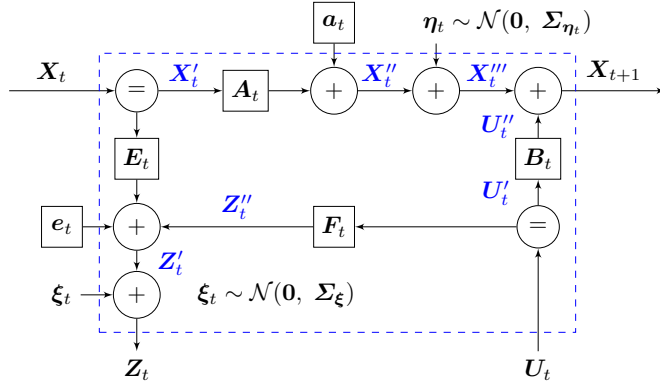


Figure 2: Forney factor graph of the linear Gaussian dynamical system used by i2C. Blue terms are intermediate variables used in the message derivations (Section B).

The message-passing on the graph of Figure 2 performs the same inference as Kalman filtering and smoothing [23], with the addition that the inputs are also uncertain¹. Additionally, the inference starts with an ‘innovation’ (observation) of x_0 in order to evaluate (x_t, u_t) rather than (x_{t+1}, u_t) , but this is a minor discrepancy as the subsequent prediction and innovation steps are the same. The forward and backward messages are derived in Sections B.1-B.2. While the message-passing form is more verbose than the standard Kalman filtering and smoothing equations, they allow us to appreciate how this framework performs optimal control [13]. From Equation 4 with the LQ-equivalent z_t and Θ , it is clear that the negative log-likelihood acts as an upper bound on the LQ cost, as it incorporates the trajectory likelihood, which depends on the system’s stochasticity and uncertainty in controls. Therefore, as the EM algorithm maximizes the log-likelihood, it in turn minimizes the LQ cost, performing Bayesian optimal control. The further connections between i2C and LQ control are discussed in Section 2.1.1 and B.5.

Maximisation Step

To update α , the scale factor between the LQ cost Θ and the estimated Λ_ξ must be found. This is derived by maximizing the expected log-likelihood via the derivative:

$$-2 \frac{\partial}{\partial \alpha} \mathbb{E}[\mathcal{L}(\alpha)] = \frac{\partial}{\partial \alpha} (\text{tr}\{\Sigma_\xi^{-1} \hat{\Sigma}_\xi\} + T \log |\Sigma_\xi|) = -\text{tr}\{\Theta \hat{\Sigma}_\xi\} + T d_z \alpha^{-1} = 0, \quad (6)$$

$$\text{where } \hat{\Sigma}_\xi = \sum_{t=0}^T [(z_t - E_t \mu_{x_t} - F_t \mu_{u_t})(z_t - E_t \mu_{x_t} - F_t \mu_{u_t})^\top + E_t \Sigma_{x_t} E_t^\top + F_t \Sigma_{u_t} F_t^\top].$$

In practice this means that over EM iterations, as the state-action trajectory moves towards the goal, Λ_ξ and therefore α steadily increases. This in turn results in the control cost term increasing in significance in the log-likelihood (Equation 4). The resulting annealing effect aids in stabilizing the optimization. This effect bears a resemblance to curriculum learning [24], where the task (e.g. cost function) increases in difficulty as the performance improves, as a strategy for learning complex tasks effectively.

Linear Gaussian Controller

For finite horizon LQ control, it can be shown that a time-varying linear controller is the optimal policy. Here we show that this is true for the inference setting as well. By examining the conditional distribution between the marginalized posteriors of x and u at each timestep, a time-varying linear Gaussian controller (Equation 7-9) can be derived from the messages (see Section B.4). For a time-varying linear Gaussian controller of the form $u_t \sim \mathcal{N}(K_t x_t + k_t, \Sigma_{k_t})$, i2C computes the parameters as

$$K_t = -\Sigma_{u_t} B_t \Gamma_{t+1} \Lambda_{\bar{x}_{t+1}} \Psi_{t+1} A_t, \quad (7)$$

$$k_t = \Sigma_{u_t} (\nu_{\bar{u}_t} + F_t^\top (\Sigma_\xi + E_t \Sigma_{\bar{x}_t} E_t^\top)^{-1} (z_t - E_t \mu_{\bar{x}_t} - e_t) + B_t^\top (\Gamma_{t+1} \nu_{\bar{x}_{t+1}} + (I - \Gamma_{t+1}) \nu_{\bar{x}_t} - \Gamma_{t+1} \Lambda_{\bar{x}_{t+1}} \Psi_{t+1} a_t)), \quad (8)$$

$$\Sigma_{k_t} = \Sigma_{u_t} = (\Lambda_{\bar{u}_t} + F_t^\top (\Sigma_\xi + E_t \Sigma_{\bar{x}_t} E_t^\top)^{-1} F_t + B_t^\top \Gamma_{t+1} \Lambda_{\bar{x}_{t+1}} B_t)^{-1}. \quad (9)$$

In Section 2.1.1, it is shown how the expressions for the controller resemble the corresponding expressions for LQ control. Moreover, the discrepancy between the i2C and LQ controllers can be interpreted as uncertainty-derived regulation. In the i2C controller two additional (dimensionless)

Data: $T, \alpha, \delta_\alpha, f(x, u), g(x, u)$
 $\mu_{\bar{x}_0}, \Sigma_{\bar{x}_0}, \mu_{\bar{u}_t}, \Sigma_{\bar{u}_t}$ for $t = 0:T$
Result: K_t, k_t, Σ_{k_t} for $t = 0:T$
while not converged do
 // E-Step
 for $i \leftarrow 0$ **to** $T - 1$ **do**
 Compute $\mu_{\bar{x}_{t+1}}, \Sigma_{\bar{x}_{t+1}}$ from
 forward messages (Equation 18-29),
 updating A_t, a_t, B_t, E_t, e_t and F_t
 end
 for $i \leftarrow T$ **to** 1 **do**
 Compute $\mu_{x_t}, \Sigma_{x_t}, \mu_{u_t}, \Sigma_{u_t}$ from
 backward messages and
 marginalisation (Equation 32-45)
 end
 // M-Step
 Update α with reg. (Equation 6, 10)
 Update priors, $\mu_{\bar{u}} = \mu_u, \Sigma_{\bar{u}} = \Sigma_u$
end
 // Controller
 Compute linear Gaussian controller
 K_t, k_t, Σ_{k_t} for $t = 0:T$ from messages
 (Equation 7-9)

Algorithm 1: EM for Linear Gaussian i2C

¹If the input is incorporated into the state, the two procedures become identical, however the joint dynamics then become degenerate due to the independence of the inputs

LQR	Riccati Backward Message	Message-derived Controller
Q	$E_t^\top (\Sigma_\xi + F_t \Sigma_{\vec{u}_t} F_t^\top)^{-1} E_t$	$E_t^\top (\Sigma_\xi + F_t \Sigma_{\vec{u}_t} F_t^\top)^{-1} E_t$
R	$F_t^\top (\Sigma_\xi + E_t \Sigma_{\vec{x}_t} E_t^\top)^{-1} F_t$	$F_t^\top (\Sigma_\xi + E_t \Sigma_{\vec{x}_t} E_t^\top)^{-1} F_t$
P_t	$\Lambda_{\vec{x}_t}^{\leftarrow}$	$\Gamma_t \Lambda_{\vec{x}_t}^{\leftarrow}$
p_t	$-\nu_{\vec{x}_t}^{\leftarrow}$	$-\Gamma_t \nu_{\vec{x}_t}^{\leftarrow}$

Table 1: Due to the formulation of i2C, the precision of the observation noise is proportional to the LQ cost function weights. Additionally, due to the linear Gaussian assumption, we can show that the precision and scaled-mean of the backward messages of the state belief correspond to the value function parameters in LQR. These equivalences are explained further in Section B.5.

terms appear, Γ and Ψ (see Section B.4), which are functions of $\Sigma_{\vec{x}_{t+1}}^{\leftarrow}$, $\Sigma_{\vec{x}_t}^{\leftarrow}$ and $\Sigma_{\vec{u}_t}^{\leftarrow}$. As process uncertainty increases, Γ acts to ‘turn off’ the optimal control terms of the controller and rely on the priors. Meanwhile, Ψ represents the confidence in the controller, which counteracts the attenuating effects of Γ given sufficient control certainty. These findings parallel the ‘turn-off phenomenon’ observed in Dual Control [25, 26] and Bayesian Reinforcement Learning [27], where actions are attenuated under uncertainty. This behaviour is important for settings such as probabilistic Model-based Reinforcement Learning [28], where localised regions of uncertainty can indicate modelling error, and such errors can lead to detrimental policy updates. Attenuating the policy updates in these regions between model learning iterations would mitigate this pitfall.

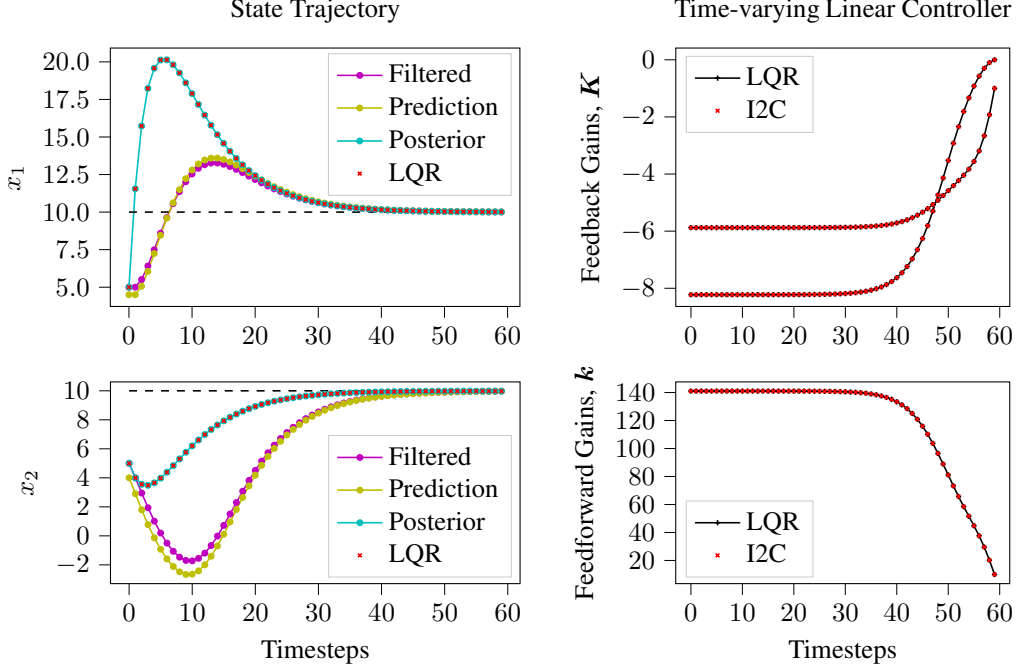
2.1.1 Connections to Finite Horizon Maximum Entropy LQR

To understand how this framework performs optimal control, we look at the backward messages of the probabilistic graphical model described in Figure 2 with a control perspective [13]. By looking at the backward messages of the state $\vec{x}_t$ in Section B.3, the backwards evolution of the precision (Equation 52) and scaled-mean (Equation 57) can be seen to have a similar Ricatti form to the quadratic value function parameters for LQ control (Equation 16-17). Extending this analysis to find the linear Gaussian controllers from the conditional distributions, we see that some of the equivalent terms have the additional uncertainty-weighted scalar term Γ . Table 1 details the correspondence.

The control covariance, Equation 9, can be seen to resemble that of a Maximum Entropy controller. In Control Theory and Reinforcement Learning, the entropy of a policy can be interpreted as a metric for robustness, so a maximum entropy objective has been added to cost functions as regularization [29]. Augmenting the LQ cost function with the entropy of the control inputs, the covariance of the input at each timestep can be shown to be $\Sigma_t = (R + B^\top P_{t+1} B)^{-1}$ (using LQ notation, see Section A) [30]. Comparing this to Equation 9 and Table 1, it can be seen that this maximum entropy control is calculated by the backward message, and combined with the prior (forward message) to construct the posterior. This fusion is important as the prior can be used to regularize exploration during inference, which is essential for mitigating the effects of linearizing the dynamics during approximate inference of nonlinear systems (see Section 2.2). This smoothing mechanism has previously been added explicitly or via constraints on the trajectory update during optimization.

2.2 Nonlinear i2C through Approximate Inference

The linear analysis conducted here can be naturally extended to nonlinear dynamical systems through linearization, taking the Jacobian of the dynamics and observation models about the current state-action trajectory. This approach has been applied to both state estimation (i.e. Extended Kalman Smoothing) and optimal control (i.e. DDP). From a probabilistic perspective, this linearization renders the inference approximate. As a consequence, careful consideration of the priors and additional regularization is required, as the act of linearizing imposes a requirement of *local* improvement during inference. Placing small priors on u ensures that the Bayesian posterior remains close to the prior, and was found to be critical for systems that were highly nonlinear or with low sampling frequencies. As in Extended Kalman Filtering and other inference schemes for nonlinear systems [31], the dynamics are linearized in the forward pass. This linearization-based approximate inference can be viewed as Gauss-Newton optimization [32], making it closely related to approximate trajectory optimization algorithms such as iLQR. Additionally, it was found that the α update during the M-Step must be restricted to ensure the state-action distribution did not change signif-



(a) Comparing state trajectories

(b) Comparing the parameters of the linear controller

Figure 3: Demonstrating how I2C generalizes the Dynamic Programming Finite Horizon LQR solution. This is achieved when the controls have a large prior and the certainty in the target observation is high. Note ‘Filtered’ and ‘Prediction’ correspond to $\mu_{\vec{x}_t'}$ and $\mu_{\vec{x}_{t+1}}$ in Figure 2 respectively.

icantly between iterations. By looking at a bound δ_ξ on the KL divergence between $\mathbf{Z}$ updates (Equation 10), this in fact can be applied as a bound δ_α on the update ratio. As the expression

$$D_{\text{KL}}(\mathbf{Z}^i || \mathbf{Z}^{i+1}) = \frac{1}{2} \left[\log \frac{|\Sigma_\xi^{i+1}|}{|\Sigma_\xi^i|} + \text{tr}\{\mathbf{A}_\xi^{i+1} \Sigma_\xi^i\} - d_z \right] = \frac{1}{2} \left[\log \frac{\alpha^{i+1}}{\alpha^i} + d_z \frac{\alpha^i}{\alpha^{i+1}} - d_z \right] \leq \delta_\xi \quad (10)$$

is monotonic increasing in the ratio α^{i+1}/α^i . From the perspective of approximate EM, the regularized M-Step is motivated by mitigating the adverse effect of the linearization assumption on the likelihood estimate. [33].

3 Experimental Results

An empirical evaluation is presented, first to highlight the equivalence of I2C to the LQR solution and second to compare I2C to state-of-the-art algorithms on nonlinear dynamical systems².

3.1 Equivalence with finite-horizon LQR by Dynamic Programming

In Section 2.1, the LQR problem was used to motivate the linear Gaussian assumption for I2C. In Section B.5 it is shown how, under specific settings, the message passing expressions reduce to those found when solving the LQR problem via Dynamic Programming. Figure 3 illustrates this numerically, for an LQR problem described in Section C.1.

3.2 Evaluation on nonlinear trajectory optimization tasks

To evaluate the viability of I2C for nonlinear trajectory optimization, its performance on three standard control tasks were compared to similar baseline methods. iLQR and GPS are two popular algorithms that use local linearization for time-varying controllers and have demonstrated strong performance on complex control problems. iLQR is deterministic, so here it is used as a baseline for

²The code is available at <https://github.com/JoeMWatson/input-inference-for-control>

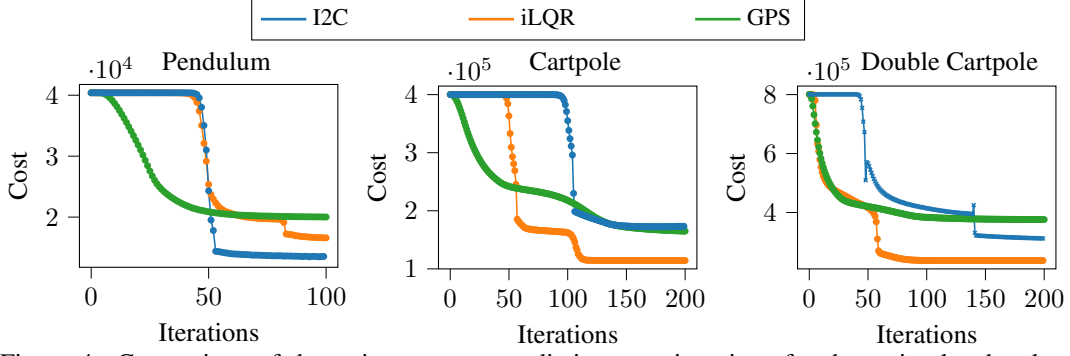


Figure 4: Comparison of the trajectory cost prediction over iterations for three simulated tasks during trajectory optimization. For all algorithms, the dynamics are linearized once per iteration. For experimental details see Section C.2.

the ignoring uncertainty in stochastic control problems. While GPS was motivated to train Neural Network policies, here we use its time-varying linear controllers, viewing it as Maximum Entropy iLQG. In order to perform the linearization required for approximate inference (and the baseline approaches), the test environments were implemented using the Autograd library [34]. We test on three classical problems of increasing complexity in state-action-observation dimensionality (d_x, d_u, d_z): Pendulum (2, 1, 4), Cartpole (4, 1, 6) and Double Cartpole (6, 1, 9) swing-up. Both Cartpole domains are also underactuated, which presents a significant planning challenge. All environments also have constrained actuation, which introduces both a nonlinearity and increased sensitivity to disturbances. Experimental details and additional trajectory plots are included in Section C.2.

Figure 4 shows that I2C is capable of performing effective trajectory optimization. The EM aspect of the algorithm results in a significant portion of the time is used ‘warming up’ the priors, which are set to be small in order to carry out steady exploration, rather than optimizing the control cost. iLQR performs superior trajectory optimization, both in rate and final cost. However, actuation constraints were found to lead to suboptimal convergence (in the Pendulum task, Figure 5), and the optimized controllers were comparatively highly aggressive. GPS performed steadier optimization due to the KL bound and exploration in the forward pass. In Table 2, the optimized (deterministic) controllers were evaluated on the stochastic environment. I2C performs the most consistently, operating close to its predicted cost for each task. GPS and iLQR, with more aggressive controllers and trajectories, both suffered reduced performance when evaluated on the simulated systems. We attribute this to the high-risk strategy of operating at the actuation limits, when also subjected to disturbances, especially as time-varying control strategies are inherently very brittle to any deviation in trajectory.

4 Related Work

Optimal control of nonlinear dynamical systems through iterative linearization originated from Differential Dynamic Programming (DDP) [14]. A drawback of DDP is the need for the computationally expensive second-order approximation of the dynamics. In the framework of Iterative LQR (iLQR) [35] and its stochastic extension iLQG [36], this requirement is dropped. Both algorithms perform only first order approximations, making them akin to a regularized Gauss-Newton method.

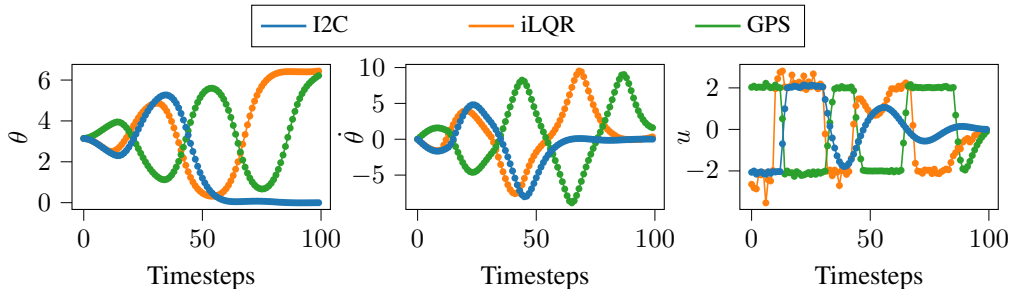


Figure 5: Comparison of the state-action trajectories of I2C, iLQR and GPS on the Pendulum swing-up task after convergence.

Environment	Algorithm	Predicted Cost	Evaluated Cost
Pendulum	i2C	1.35×10^4	$1.37 \times 10^4 \pm 3.82$
	iLQR	1.66×10^3	$1.11 \times 10^5 \pm 20.38$
	GPS	2.00×10^4	$7.01 \times 10^4 \pm 30.96$
Cartpole	i2C	1.73×10^5	$1.74 \times 10^5 \pm 0.14$
	iLQR	1.14×10^5	$1.76 \times 10^7 \pm 88.63$
	GPS	1.65×10^5	$2.94 \times 10^6 \pm 17.60$
Double Cartpole	i2C	3.12×10^5	$3.21 \times 10^5 \pm 1.79$
	iLQR	2.37×10^5	$1.76 \times 10^7 \pm 5.27 \times 10^5$
	GPS	3.76×10^5	$2.94 \times 10^6 \pm 44.39$

Table 2: Evaluating the optimized deterministic controller of each algorithm on the simulated stochastic environments. Predicted Cost refers to the converged value from Figure 4, Evaluated Cost shows the mean and standard deviation after 100 trials.

All former methods however lack a principled forward pass and instead rely on a line-search approach to find a suitable regularization, that counteracts the greediness of their local approximations. Extended LQR (eLQR) [16] and its stochastic extension seLQR [37] address this issue and perform a forward pass based on the ‘cost-to-come’, that has similarities to Kalman filtering. A more elegant solution to the problem of regularization is proposed in Guided Policy Search (GPS) [38, 39], where the Stochastic Optimal Control problem is formulated with a KL bound on the change of trajectory distributions. GPS derives a Maximum Entropy iLQG as a means to train neural network policies.

The connection between optimal control and inference, also known as the estimation-control duality and Kalman duality [4, 40] was initially noted by Kalman [41], while working on the Kalman Filter and Optimal Control. Probabilistic Control Design [42, 43, 44] derives a probabilistic variant of LQR through a KL divergence minimization, also noting the connection between the LQR cost weight matrices and the precisions of multivariate normal distributions. Furthermore, the similarity between LQR and Kalman Smoothed trajectories has previously been utilized for the ERTS controller [45]. However, this work uses standard smoothing in the state and does not derive a corresponding controller, relying on an approximate inverse dynamics model instead.

Inference has been applied to reinforcement learning for discrete environments [7], through maximizing the likelihood of a discrete latent optimality variable. AICO [12] applies this approach to the continuous LQR setting, with the state cost defining the optimality probability, and the action weight defining the precision of the action prior, which is treated like a disturbance for exploration. As with i2C, the backward messages were found to share similarities with the DAREs of LQR, however unlike i2C, the input priors are fixed. AICO was generalised to Posterior Policy Iteration (PPI) [46, 47] in which a risk-tuned linear Gaussian controller was obtained from the inferred value function. The idea of controls as a random diffusion process is shared by Todorov [48], along with Path Integral (PI) Control (KL Control for discrete environments) [9, 49, 50] that takes advantage of Feynman-Kac lemma to approximately solve the continuous-time Hamilton-Jacobi-Bellman equation using stochastic processes. PI methods iteratively compute local improvements to the controls, allowing them to be used to train parametric policies or for model predictive control.

5 Conclusion

In this work we have introduced Input Inference for Control (i2C), a novel control-as-inference formulation, by casting optimal control as Bayesian inference over the inputs. Through making the linear Gaussian assumption, we arrived at a tractable approximate EM algorithm with the use of message passing for approximate inference, and are able to draw connections with linear quadratic optimal control, Kalman filtering and Kalman smoothing through examination of the messages. Compared to prior work, this scheme employs natural regularization through the mechanisms of Bayesian inference, offering a more principled approach than currently established deterministic solvers. Moreover, our approach improves previous probabilistic approaches by naturally incorporating and optimizing over actions, enabling us to retrieve time-variant feedback controllers. Future avenues of research include the analysis of different approximate inference techniques, such as Monte Carlo, variational methods, and numerical quadrature, and the investigation of the trade-off between accuracy of inference, computational cost and benefit to control optimization.

Acknowledgments

The authors would like to thank Michael Lutter and Julen Urain for valuable feedback on the draft. This project has received funding from the European Unions Horizon 2020 research and innovation programme under grant agreement No. 640554 (SKILLS4ROBOTS).

References

- [1] Y. Tassa, T. Erez, and E. Todorov. Synthesis and stabilization of complex behaviors through on-line trajectory optimization. In *2012 IEEE/RSJ International Conference on Intelligent Robots and Systems*. IEEE, 2012.
- [2] M. Toussaint, K. Allen, K. A. Smith, and J. B. Tenenbaum. Differentiable physics and stable modes for tool-use and manipulation planning. In *Robotics: Science and Systems*, 2018.
- [3] D. E. Kirk. *Optimal control theory: an introduction*. Courier Corporation, 2012.
- [4] R. F. Stengel. *Stochastic optimal control: Theory and application*. John Wiley & Sons, Inc., 1986.
- [5] A. E. Bryson. *Applied optimal control: Optimization, estimation and control*. Routledge, 2018.
- [6] C. M. Bishop. *Pattern Recognition and Machine Learning (Information Science and Statistics)*. Springer-Verlag, 2006.
- [7] H. Attias. Planning by probabilistic inference. In *Proc. of the 9th Int. Workshop on Artificial Intelligence and Statistics*, 2003.
- [8] M. Toussaint and A. Storkey. Probabilistic inference for solving discrete and continuous state markov decision processes. In *Proceedings of the 23rd international conference on Machine learning*. ACM, 2006.
- [9] H. J. Kappen, V. Gómez, and M. Opper. Optimal control as a graphical model inference problem. In *Proceedings of the Twenty-Third International Conference on Automated Planning and Scheduling, ICAPS 2013*, 2013.
- [10] S. Levine. Reinforcement learning and control as probabilistic inference: Tutorial and review. *arXiv preprint arXiv:1805.00909*, 2018.
- [11] P. Hennig, M. A. Osborne, and M. Girolami. Probabilistic numerics and uncertainty in computations. *Proceedings of the Royal Society A: Mathematical, Physical and Engineering Sciences*, 2015.
- [12] M. Toussaint. Robot trajectory optimization using approximate inference. In *Proceedings of the 26th annual international conference on machine learning*. ACM, 2009.
- [13] C. Hoffmann and P. Rostalski. Linear optimal control on factor graphs - a message passing perspective -. *IFAC (International Federation of Automatic Control)*, 2017.
- [14] D. H. Jacobson and D. Q. Mayne. *Differential dynamic programming*. 1970.
- [15] W. Li and E. Todorov. Iterative linear quadratic regulator design for nonlinear biological movement systems. In *ICINCO (1)*, 2004.
- [16] J. van den Berg. Extended LQR: Locally-optimal feedback control for systems with non-linear dynamics and non-quadratic cost. In *Robotics Research*. Springer, 2016.
- [17] S. Levine and V. Koltun. Guided policy search. In *Proceedings of the 30th International Conference on Machine Learning, ICML 2013*, 2013.
- [18] H.-A. Loeliger, J. Dauwels, J. Hu, S. Korl, L. Ping, and F. R. Kschischang. The factor graph approach to model-based signal processing. *Proceedings of the IEEE*, 2007.
- [19] Z. Ghahramani and G. E. Hinton. Parameter estimation for linear dynamical systems. Technical report, 1996.
- [20] A. P. Dempster, N. M. Laird, and D. B. Rubin. Maximum likelihood from incomplete data via the em algorithm. *Journal of the Royal Statistical Society: Series B (Methodological)*, 1977.
- [21] L. Bruderer. *Input estimation and dynamical system identification: New algorithms and results*. PhD thesis, ETH Zurich, 2015.
- [22] H.-A. Loeliger, L. Bruderer, H. Malmberg, F. Wadehn, and N. Zalmi. On sparsity by nuv-em, gaussian message passing, and Kalman smoothing. In *2016 Information Theory and Applications Workshop (ITA)*. IEEE, 2016.
- [23] B. D. Anderson and J. B. Moore. *Optimal filtering*. Courier Corporation, 2012.
- [24] Y. Bengio, J. Louradour, R. Collobert, and J. Weston. Curriculum learning. In *Proceedings of the 26th annual international conference on machine learning*. ACM, 2009.

- [25] M. Aoki. *Optimization of stochastic systems: topics in discrete-time systems*, volume 32. Academic Press, 1967.
- [26] Y. Bar-Shalom. Stochastic dynamic programming: Caution and probing. *IEEE Transactions on Automatic Control*, 1981.
- [27] E. D. Klenke and P. Hennig. Dual control for approximate bayesian reinforcement learning. *Journal of Machine Learning Research*, 2016.
- [28] M. P. Deisenroth, G. Neumann, J. Peters, et al. A survey on policy search for robotics. *Foundations and Trends® in Robotics*, 2013.
- [29] B. D. Ziebart. Modeling purposeful adaptive behavior with the principle of maximum causal entropy. 2010.
- [30] S. Levine and V. Koltun. Variational policy search via trajectory optimization. In *Advances in neural information processing systems*, 2013.
- [31] Z. Ghahramani and S. T. Roweis. Learning nonlinear dynamical systems using an em algorithm. In *Advances in neural information processing systems*, 1999.
- [32] B. M. Bell. The iterated Kalman smoother as a Gauss–Newton method. *SIAM Journal on Optimization*, 1994.
- [33] X. Yi and C. Caramanis. Regularized em algorithms: A unified framework and statistical guarantees. In *Advances in Neural Information Processing Systems*, 2015.
- [34] D. Maclaurin, D. Duvenaud, and R. P. Adams. Autograd: Effortless gradients in numpy. In *ICML 2015 AutoML Workshop*, volume 238, 2015.
- [35] W. Li and E. Todorov. Iterative linear quadratic regulator design for nonlinear biological movement systems. In *ICINCO 2004, Proceedings of the First International Conference on Informatics in Control, Automation and Robotics, 2004*, 2004.
- [36] E. Todorov and W. Li. A generalized iterative LQG method for locally-optimal feedback control of constrained nonlinear stochastic systems. In *Proceedings of the 2005, American Control Conference, 2005*. IEEE, 2005.
- [37] W. Sun, J. van den Berg, and R. Alterovitz. Stochastic extended LQR for optimization-based motion planning under uncertainty. *IEEE Trans. Automation Science and Engineering*, 2016.
- [38] S. Levine. *Motor skill learning with local trajectory methods*. PhD thesis, Stanford University, 2014.
- [39] S. Levine and P. Abbeel. Learning neural network policies with guided policy search under unknown dynamics. In *Advances in Neural Information Processing Systems*, 2014.
- [40] E. Todorov. General duality between optimal control and estimation. In *2008 47th IEEE Conference on Decision and Control*. IEEE, 2008.
- [41] R. E. Kalman. A new approach to linear filtering and prediction problems. *Journal of basic Engineering*, 1960.
- [42] M. Kárný. Towards fully probabilistic control design. *Automatica*, 1996.
- [43] M. Kárný and T. V. Guy. Fully probabilistic control design. *Systems & Control Letters*, 2006.
- [44] J. Šindelář, I. Vajda, and M. Kárný. Stochastic control optimal in the kullback sense. *Kybernetika*, 2008.
- [45] M. Zima, L. Armesto, V. Girbés, A. Sala, and V. Šmídl. Extended rauch-tung-striebeel controller. In *52nd IEEE Conference on Decision and Control*. IEEE, 2013.
- [46] K. Rawlik, M. Toussaint, and S. Vijayakumar. On stochastic optimal control and reinforcement learning by approximate inference. In *Twenty-Third International Joint Conference on Artificial Intelligence*, 2013.
- [47] K. C. Rawlik. On probabilistic inference approaches to stochastic optimal control. 2013.
- [48] E. Todorov. Linearly-solvable markov decision problems. In *Advances in neural information processing systems*, 2007.
- [49] Y. Pan, E. Theodorou, and M. Kontitsis. Sample efficient path integral control under uncertainty. In *Advances in Neural Information Processing Systems 28: Annual Conference on Neural Information Processing Systems 2015*, 2015.
- [50] V. Gómez, H. J. Kappen, J. Peters, and G. Neumann. Policy search for path integral control. In *Machine Learning and Knowledge Discovery in Databases - European Conference, ECML PKDD 2014*, 2014.
- [51] K. B. Petersen, M. S. Pedersen, et al. The matrix cookbook. *Technical University of Denmark*, 2008.

A The Dynamic Programming solution to the Linear Quadratic Regulator

Given a linear system, we wish to find a control sequence $\mathbf{u}_{0:T}^*$ that minimizes a quadratic cost function over a finite time horizon T for a goal state $\mathbf{x}_g$ and input $\mathbf{u}_g$:

$$\begin{aligned} \min_{\mathbf{u}_{0:T}} & \left[(\mathbf{x}_T - \mathbf{x}_g)^\top \mathbf{Q}_f (\mathbf{x}_T - \mathbf{x}_g) + \sum_{t=0}^{T-1} (\mathbf{x}_t - \mathbf{x}_g)^\top \mathbf{Q} (\mathbf{x}_t - \mathbf{x}_g) + (\mathbf{u}_t - \mathbf{u}_g)^\top \mathbf{R} (\mathbf{u}_t - \mathbf{u}_g) \right] \\ \text{s.t. } & \mathbf{x}_{t+1} = \mathbf{A}\mathbf{x}_t + \mathbf{a} + \mathbf{B}\mathbf{u}_t \end{aligned} \quad (11)$$

Solving this method via Dynamic Programming, we can construct a quadratic value function backwards through time to find the optimal control at each timestep, which we can calculate for Equation 11 using Bellman's Principle of Optimality.

Starting with $\mathbf{P}_T = \mathbf{Q}_f$, $\mathbf{p}_T = -\mathbf{x}_g^\top \mathbf{Q}$, $p_T = 0$.

$$V_t(\mathbf{x}) = \mathbf{x}^\top \mathbf{P}_t \mathbf{x} + 2\mathbf{x}^\top \mathbf{p}_t + p_t \quad (12)$$

$$= \min_{\mathbf{u}} [(\mathbf{x}_t - \mathbf{x}_g)^\top \mathbf{Q} (\mathbf{x}_t - \mathbf{x}_g) + (\mathbf{u}_t - \mathbf{u}_g)^\top \mathbf{R} (\mathbf{u}_t - \mathbf{u}_g) + V_{t+1}(\mathbf{x}_{t+1})] \quad (13)$$

The optimal input can be found to be linear in state,

$$\mathbf{u}_t^* = -(\mathbf{R} + \mathbf{B}^\top \mathbf{P}_{t+1} \mathbf{B})^{-1} (\mathbf{B}^\top \mathbf{P}_{t+1} (\mathbf{A}\mathbf{x}_t + \mathbf{a}) + \mathbf{B}^\top \mathbf{p}_{t+1} - \mathbf{R}\mathbf{u}_g) \quad (14)$$

$$= \mathbf{K}_t \mathbf{x}_t + \mathbf{k}_t \quad (15)$$

The parameters of the value functions follow the recursive form, i.e.

$$\mathbf{P}_t = \mathbf{Q} + \mathbf{A}^\top \mathbf{P}_{t+1} \mathbf{A} + \mathbf{A}^\top \mathbf{P}_{t+1} \mathbf{B} (\mathbf{R} + \mathbf{B}^\top \mathbf{P}_{t+1} \mathbf{B})^{-1} \mathbf{B}^\top \mathbf{P}_{t+1} \mathbf{A} \quad (16)$$

$$\begin{aligned} \mathbf{p}_t &= \mathbf{A}^\top (\mathbf{P}_{t+1} \mathbf{a} + \mathbf{p}_{t+1} - \mathbf{P}_{t+1} \mathbf{B} (\mathbf{R} + \mathbf{B}^\top \mathbf{P}_{t+1} \mathbf{B})^{-1} (\mathbf{B}^\top \mathbf{P}_{t+1} \mathbf{a} + \mathbf{B}^\top \mathbf{p}_{t+1} - \mathbf{R}\mathbf{u}_g)) \\ &\quad - \mathbf{Q}\mathbf{x}_g \end{aligned} \quad (17)$$

B Derivation of I2C Linear Gaussian Messages

All messages are derived following the graphical model in Figure 2. Note the figure includes the intermediate variables (denoted with primes), used to add clarity to the derivations.

B.1 Forward Messages

The forward message are very close to those of Kalman Filtering, except the inputs are also observed and so have their own innovation step.

The innovation and propagation of the input into the system dynamics:

$$\nu_{\vec{u}_t'} = \nu_{\vec{u}_t} + \mathbf{F}_t^\top (\Sigma_{\vec{x}_t} + \mathbf{E}_t \Sigma_{\vec{x}_t} \mathbf{E}_t^\top)^{-1} (z_t - \mathbf{E}_t \mu_{\vec{x}_t} - e_t) \quad (18)$$

$$\Lambda_{\vec{u}_t'} = \Lambda_{\vec{u}_t} + \mathbf{F}_t^\top (\Sigma_{\vec{x}_t} + \mathbf{E}_t \Sigma_{\vec{x}_t} \mathbf{E}_t^\top)^{-1} \mathbf{F}_t \quad (19)$$

$$\mu_{\vec{u}_t''} = \mathbf{B}_t \mu_{\vec{u}_t'} \quad (20)$$

$$\Sigma_{\vec{u}_t''} = \mathbf{B}_t \Sigma_{\vec{u}_t'} \mathbf{B}_t^\top \quad (21)$$

The innovation and propagation of the state, incorporating the input:

$$\nu_{\vec{x}_t'} = \nu_{\vec{x}_t} + \mathbf{E}_t^\top (\Sigma_{\vec{x}_t} + \mathbf{F}_t \Sigma_{\vec{u}_t} \mathbf{F}_t^\top)^{-1} (z_t - \mathbf{F}_t \mu_{\vec{u}_t} - e_t) \quad (22)$$

$$\Lambda_{\vec{x}_t'} = \Lambda_{\vec{x}_t} + \mathbf{E}_t^\top (\Sigma_{\vec{x}_t} + \mathbf{F}_t \Sigma_{\vec{u}_t} \mathbf{F}_t^\top)^{-1} \mathbf{E}_t \quad (23)$$

$$\mu_{\vec{x}_t''} = \mathbf{A}_t \mu_{\vec{x}_t'} + \mathbf{a}_t \quad (24)$$

$$\Sigma_{\vec{x}_t''} = \mathbf{A}_t \Sigma_{\vec{x}_t'} \mathbf{A}_t^\top \quad (25)$$

$$\mu_{\vec{x}_t'''} = \mu_{\vec{x}_t''} \quad (26)$$

$$\Sigma_{\vec{x}_t'''} = \Sigma_{\vec{x}_t''} + \Sigma_{\eta_t} \quad (27)$$

$$\mu_{\vec{x}_{t+1}} = \mu_{\vec{x}_t'''} + \mu_{\vec{u}_t''} \quad (28)$$

$$\Sigma_{\vec{x}_{t+1}} = \Sigma_{\vec{x}_t'''} + \Sigma_{\vec{u}_t''} \quad (29)$$

B.2 Backward Messages

The most efficient means of constructing the backward messages for marginalisation is to make use of the ‘auxiliary’ form (see [18, 21]), which has several useful properties for message propagation. In particular they are invariant to the Addition factor, so the various offsets are automatically considered.

$$\tilde{\Lambda}_x = (\Sigma_{\vec{x}} + \Sigma_{\leftarrow x})^{-1} = \Lambda_{\vec{x}} - \Lambda_{\vec{x}} \Sigma_x \Lambda_{\vec{x}} \quad (30)$$

$$\tilde{\nu}_x = \tilde{\Lambda}_x (\mu_{\vec{x}_t} - \mu_{\leftarrow x_t}) = \nu_{\vec{x}_t} - \Lambda_{\vec{x}_t} \mu_{x_t} \quad (31)$$

Like the marginal they are a fusion of the forward and backward message, but in the ‘dual’ form.

For initialising the backward pass there are two approaches. One is to follow the idea of the terminal cost from Equation 11, where for example $P_T = Q_f = Q$, $p_T = 0$, $\Lambda_{\leftarrow x_T} = \Lambda_\xi$, $\nu_{\leftarrow x_T} = 0$. The marginals can then be constructed following Equation 5. Any Q_f can be used, so long as $\Lambda_{\leftarrow x_T}$ is constructed with an α following Section 2.1. In practice, it was found crucial to tune up this terminal cost to ensure the target state is reached with a responsive controller (as many target states lay at unstable equilibria). However Q_f the represents another (multi-dimensional) hyperparameter to tune. Using the probabilistic perspective, instead choose Σ_{x_T} such that the prior $\Sigma_{\vec{x}_T}$ has been reduced by a scale factor κ . While this deviates from the previous quadratic cost formulation into an adaptive cost function, it was found to be both simple and effective when tackling difficult domains. The adaptation becomes an important quality for nonlinear problems where the initial dynamics are stable and the target state dynamics are unstable. This scheme acts to tune up the terminal cost as the dynamics become more unstable, which causes the state uncertainty to grow at a greater rate, which in turn acts to increase the responsiveness of the controller. As we also wish to keep the prior and posterior trajectories tight during optimization (to ensure the linearization assumption is valid), we set $\mu_{x_T} = \mu_{\vec{x}_T}$. In the experiments of Section 3.2, $Q_f = Q$.

Starting with $\Sigma_{x_T} = \Sigma_{\vec{x}_T}$, $\mu_{x_T} = \mu_{\vec{x}_T}$.

Construct the auxiliary for x_{t+1} ,

$$\tilde{\Lambda}_{x_{t+1}} = \Lambda_{\vec{x}_{t+1}} - \Lambda_{\vec{x}_{t+1}} \Sigma_{x_{t+1}} \Lambda_{\vec{x}_{t+1}} \quad (32)$$

$$\tilde{\nu}_{x_{t+1}} = \nu_{\vec{x}_{t+1}} - \Lambda_{\vec{x}_{t+1}} \mu_{x_{t+1}} \quad (33)$$

The auxiliary is invariant across an addition operation, so

$$\tilde{\Lambda}_{x_t''} = \tilde{\Lambda}_{x_t'''} = \tilde{\Lambda}_{x_{t+1}} \quad (34)$$

$$\tilde{\nu}_{x_t''} = \tilde{\nu}_{x_t'''} = \tilde{\nu}_{x_{t+1}} \quad (35)$$

Propagate the state belief backwards through system dynamics,

$$\tilde{\Lambda}_{x_t'} = A_t^\top \tilde{\Lambda}_{x_t''} A_t \quad (36)$$

$$\tilde{\nu}_{x_t'} = A_t^\top \tilde{\nu}_{x_t''} \quad (37)$$

Marginalized variables are invariant across the Equality node, so marginalize x_t at x_t' ,

$$\Sigma_{x_t} = \Sigma_{x_t'} = \Sigma_{\vec{x}_t'} - \Sigma_{\vec{x}_t'} \tilde{\Lambda}_{x_t'} \Sigma_{\vec{x}_t'} \quad (38)$$

$$\mu_{x_t} = \mu_{x_t'} = \mu_{\vec{x}_t'} - \Sigma_{\vec{x}_t'} \tilde{\nu}_{x_t'} \quad (39)$$

To find μ_{u_t} , note that due to the addition operation, the auxiliary of u_t'' is equal to that of x_t''' ,

$$\tilde{\Lambda}_{u_t''} = \tilde{\Lambda}_{x_t'''} \quad (40)$$

$$\tilde{\nu}_{u_t''} = \tilde{\nu}_{x_t'''} \quad (41)$$

$$\tilde{\Lambda}_{u_t'} = B_t^\top \tilde{\Lambda}_{u_t''} B_t \quad (42)$$

$$\tilde{\nu}_{u_t'} = B_t^\top \tilde{\nu}_{u_t''} \quad (43)$$

$$\Sigma_{u_t} = \Sigma_{u_t'} = \Sigma_{\vec{u}_t'} - \Sigma_{\vec{u}_t'} \tilde{\Lambda}_{u_t'} \Sigma_{\vec{u}_t'} \quad (44)$$

$$\mu_{u_t} = \mu_{u_t'} = \mu_{\vec{u}_t'} - \Sigma_{\vec{u}_t'} \tilde{\nu}_{u_t'} \quad (45)$$

B.3 Riccati Backward Messages for Control

To understand the relation to optimal control, the backward messages must be represented recursively as a Discrete Algebraic Ricatti Equation.

Recursion of the precision

$$\Lambda_{\underline{x}_t} = \Lambda_{\underline{x}_t} + \mathbf{E}_t^\top \Lambda_{\underline{z}_t} \mathbf{E}_t \quad (46)$$

$$= \mathbf{A}_t^\top \Lambda_{\underline{x}_t} \mathbf{A}_t + \mathbf{E}_t^\top \Lambda_{\underline{z}_t} \mathbf{E}_t \quad (47)$$

$$\Lambda_{\underline{z}_t} = (\Sigma_\xi + \mathbf{F}_t \Sigma_{\underline{u}_t} \mathbf{F}_t^\top)^{-1} \quad (48)$$

Using the matrix inversion identity $(\mathbf{A}^{-1} + \mathbf{B})^{-1} = \mathbf{A} - \mathbf{A}(\mathbf{A} + \mathbf{B}^{-1})\mathbf{A}$ [51],

$$\Lambda_{\underline{x}_t} = (\Sigma_{\eta_t} + \Sigma_{\underline{u}_t} + \Sigma_{\underline{x}_{t+1}})^{-1} \quad (49)$$

$$= \Lambda_{\underline{x}_{t+1}} - \Lambda_{\underline{x}_{t+1}} ((\Sigma_{\eta_t} + \Sigma_{\underline{u}_t})^{-1} + \Lambda_{\underline{x}_{t+1}})^{-1} \Lambda_{\underline{x}_{t+1}} \quad (50)$$

$$\Sigma_{\underline{u}_t} = \mathbf{B}_t (\Lambda_{\underline{u}_t} + \mathbf{F}_t^\top (\Sigma_\xi + \mathbf{E}_t \Sigma_{\underline{x}_t} \mathbf{E}_t^\top)^{-1} \mathbf{F}_t)^{-1} \mathbf{B}_t^\top \quad (51)$$

So the recursion in full is

$$\begin{aligned} \Lambda_{\underline{x}_t} &= \mathbf{E}_t^\top (\Sigma_\xi + \mathbf{F}_t \Sigma_{\underline{u}_t} \mathbf{F}_t^\top)^{-1} \mathbf{E}_t + \mathbf{A}_t^\top \Lambda_{\underline{x}_{t+1}} \mathbf{A}_t - \mathbf{A}_t^\top \Lambda_{\underline{x}_{t+1}} \\ &\quad ((\Sigma_{\eta_t} + \mathbf{B}_t (\Lambda_{\underline{u}_t} + \mathbf{F}_t^\top (\Sigma_\xi + \mathbf{E}_t \Sigma_{\underline{x}_t} \mathbf{E}_t^\top)^{-1} \mathbf{F}_t)^{-1} \mathbf{B}_t^\top)^{-1} + \Lambda_{\underline{x}_{t+1}})^{-1} \Lambda_{\underline{x}_{t+1}} \mathbf{A}_t \end{aligned} \quad (52)$$

Recursion of the scaled-mean

$$\nu_{\underline{x}_t} = \nu_{\underline{x}_t} + \mathbf{E}_t^\top \Lambda_{\underline{z}_t} (z_t - \mathbf{F}_t \mu_{\underline{u}_t} - e_t) \quad (53)$$

$$\nu_{\underline{x}_t} = \mathbf{A}_t^\top \Lambda_{\underline{x}_t} (\mu_{\underline{x}_t} - \mathbf{a}_t) = \mathbf{A}_t^\top \Lambda_{\underline{x}_t} (\Sigma_{\underline{x}_{t+1}} \nu_{\underline{x}_{t+1}} - \mu_{\underline{u}_t} - \mathbf{a}_t) \quad (54)$$

Substituting Equation 18-21 into Equation 54 (using $\Lambda_{\underline{z}_t}$ for brevity)

$$\nu_{\underline{x}_t} = \mathbf{A}_t^\top \Lambda_{\underline{x}_t} (\Sigma_{\underline{x}_{t+1}} \nu_{\underline{x}_{t+1}} - \mathbf{a}_t - \mathbf{B}_t \Sigma_{\underline{u}_t} (\nu_{\underline{u}_t} + \mathbf{F}_t^\top \Lambda_{\underline{z}_t} (z_t - \mathbf{E}_t \mu_{\underline{x}_t} - e_t))) \quad (55)$$

Substituting $\Lambda_{\underline{x}_t}$ through Equation 50,

$$\begin{aligned} \nu_{\underline{x}_t} &= \mathbf{A}_t^\top (\Lambda_{\underline{x}_{t+1}} - \Lambda_{\underline{x}_{t+1}} ((\Sigma_{\eta_t} + \Sigma_{\underline{u}_t})^{-1} + \Lambda_{\underline{x}_{t+1}})^{-1} \Lambda_{\underline{x}_{t+1}}) \\ &\quad (\Sigma_{\underline{x}_{t+1}} \nu_{\underline{x}_{t+1}} - \mathbf{a}_t - (\mathbf{B}_t (\Lambda_{\underline{u}_t} + \mathbf{F}_t^\top \Lambda_{\underline{z}_t} \mathbf{F}_t)^{-1} (\nu_{\underline{u}_t} + \mathbf{F}_t^\top \Lambda_{\underline{z}_t} z_t))) \end{aligned} \quad (56)$$

So the full recursion is,

$$\begin{aligned} \nu_{\underline{x}_t} &= \mathbf{A}_t^\top (\mathbf{I} - \Lambda_{\underline{x}_{t+1}} ((\Sigma_{\eta_t} + \mathbf{B}_t (\Lambda_{\underline{u}_t} + \mathbf{F}_t^\top (\Sigma_\xi + \mathbf{E}_t \Sigma_{\underline{x}_t} \mathbf{E}_t^\top)^{-1} \mathbf{F}_t)^{-1} \mathbf{B}_t^\top)^{-1} + \Lambda_{\underline{x}_{t+1}})^{-1}) \\ &\quad (\nu_{\underline{x}_{t+1}} - \Lambda_{\underline{x}_{t+1}} \mathbf{a}_t - \Lambda_{\underline{x}_{t+1}} (\mathbf{B}_t (\Lambda_{\underline{u}_t} + \mathbf{F}_t^\top (\Sigma_\xi + \mathbf{E}_t \Sigma_{\underline{x}_t} \mathbf{E}_t^\top)^{-1} \mathbf{F}_t)^{-1} \\ &\quad (\nu_{\underline{u}_t} + \mathbf{F}_t^\top (\Sigma_\xi + \mathbf{E}_t \Sigma_{\underline{x}_t} \mathbf{E}_t^\top)^{-1} (z_t - \mathbf{E}_t \mu_{\underline{x}_t} - e_t)))) \\ &\quad + \mathbf{E}_t^\top (\Sigma_\xi + \mathbf{F}_t \Sigma_{\underline{u}_t} \mathbf{F}_t^\top)^{-1} (z_t - \mathbf{F}_t \mu_{\underline{u}_t} - e_t) \end{aligned} \quad (57)$$

B.4 Linear Gaussian Controller

To extract the linear Gaussian controllers, we find the conditional distribution between $\mathbf{u}_t$ and $\mathbf{x}_t$.

The input estimate is marginalized by fusing the forward and backward message:

$$\mu_{\mathbf{u}_t} = \Sigma_{\mathbf{u}_t} (\nu_{\underline{u}_t} + \nu_{\underline{u}_t}) \quad (58)$$

$$\nu_{\underline{u}_t} = \mathbf{F}_t^\top (\Sigma_\xi + \mathbf{E}_t^\top \Sigma_{\underline{x}_t} \mathbf{E}_t)^{-1} (z_t - \mathbf{E}_t \mu_{\underline{x}_t} - e_t) + \mathbf{B}_t^\top \nu_{\underline{u}_t} \quad (59)$$

$$= \mathbf{F}_t^\top (\Sigma_\xi + \mathbf{E}_t^\top \Sigma_{\underline{x}_t} \mathbf{E}_t)^{-1} (z_t - \mathbf{E}_t \mu_{\underline{x}_t} - e_t) + \mathbf{B}_t^\top \Lambda_{\underline{u}_t} \mu_{\underline{u}_t} \quad (60)$$

$$\begin{aligned} &= \mathbf{F}_t^\top (\Sigma_\xi + \mathbf{E}_t^\top \Sigma_{\underline{x}_t} \mathbf{E}_t)^{-1} (z_t - \mathbf{E}_t \mu_{\underline{x}_t} - e_t) \\ &\quad + \mathbf{B}_t^\top \Lambda_{\underline{u}_t} (\mu_{\underline{x}_{t+1}} - \mu_{\underline{x}_t}) \end{aligned} \quad (61)$$

Eventually we need an expression in terms of the marginal x_t , so we need to be able to express Eq. (61) in terms of μ_{x_t} . Taking the marginalisation rule from Equation 5,

$$\mu_{x_t}''' = A\mu_{x_t} + a_t = \Sigma_{x_t}'''(\nu_{x_t}''' + \nu_{x_t}''') \quad (62)$$

$$\mu_{x_t}''' = \Sigma_{x_t}'''(\Lambda_{x_t}''' \mu_{x_t}''' - \nu_{x_t}''') = \Sigma_{x_t}'''(\Lambda_{x_t}''' \mu_{x_t}''' - \nu_{x_t}''') \quad (63)$$

$$\mu_{x_{t+1}} - \mu_{x_t}''' = \mu_{x_{t+1}} + \Sigma_{x_t}''' \nu_{x_t}''' - \Sigma_{x_t}''' \Lambda_{x_t}''' \mu_{x_t}''' \quad (64)$$

To find Λ_{u_t}'' , we can use the matrix inversion identity $(A^{-1} + B^{-1})^{-1} = B(A+B)^{-1}A = A(A+B)^{-1}B$ [51]

$$\Lambda_{u_t}'' = (\Sigma_{x_{t+1}} + \Sigma_{x_t}''')^{-1} \quad (65)$$

$$= \Lambda_{x_t}'''(\Lambda_{x_{t+1}} + \Lambda_{x_t}''')^{-1} \Lambda_{x_{t+1}} \quad (66)$$

$$= \Lambda_{x_{t+1}}(\Lambda_{x_{t+1}} + \Lambda_{x_t}''')^{-1} \Lambda_{x_t}''' \quad (67)$$

We introduce dimensionless term Γ for brevity, and will discuss interpretations of it in subsequent sections,

$$\Gamma_{t+1} = \Lambda_{x_t}'''(\Lambda_{x_{t+1}} + \Lambda_{x_t}''')^{-1} \quad (68)$$

$$I - \Gamma_{t+1} = \Lambda_{x_{t+1}}(\Lambda_{x_{t+1}} + \Lambda_{x_t}''')^{-1} \quad (69)$$

This allows us to express the last term of Equation 61 as,

$$\Lambda_{u_t}''(\mu_{x_{t+1}} - \mu_{x_t}''') = (\Gamma_{t+1} \nu_{x_{t+1}} + (I - \Gamma_{t+1}) \nu_{x_t}''') - \Gamma_{t+1} \Lambda_{x_{t+1}} \Sigma_{x_t}''' \Lambda_{x_t}''' \mu_{x_t}''' \quad (70)$$

where

$$\nu_{x_t}''' = \Lambda_{x_t}'''(\mu_{x_{t+1}} - \mu_{u_t}''') \quad (71)$$

To develop the last term of Equation 70, recall the marginalisation rule for Λ (Equation 5),

$$\Sigma_{x_t}''' \Lambda_{x_t}''' \mu_{x_t}''' = \Sigma_{x_t}'''(\Lambda_{x_t}''' + \Lambda_{x_t}''') \mu_{x_t}''' \quad (72)$$

To understand this better, it is best to expand Λ_{x_t}''' ,

$$\Lambda_{x_t}''' = (\Sigma_{x_{t+1}} + \Sigma_{u_t}''')^{-1} = \Lambda_{u_t}''(\Lambda_{x_{t+1}} + \Lambda_{u_t}''')^{-1} \Lambda_{x_{t+1}} \quad (73)$$

Applying this to Equation 72 and introducing Ψ (another dimensionless scaling term)

$$\Sigma_{x_t}'''(\Lambda_{x_t}''' + \Lambda_{x_t}''') \mu_{x_t}''' = \Sigma_{x_t}'''(\Lambda_{x_t}''' + \Lambda_{u_t}''')(\Lambda_{x_{t+1}} + \Lambda_{u_t}''')^{-1} \Lambda_{x_{t+1}} \mu_{x_t}''' \quad (74)$$

$$= \Psi_{t+1} \mu_{x_t}''' \quad (75)$$

$$\text{where } \Psi_{t+1} = \Sigma_{x_t}'''(\Lambda_{x_t}''' + \Lambda_{u_t}''')(\Lambda_{x_{t+1}} + \Lambda_{u_t}''')^{-1} \Lambda_{x_{t+1}} \quad (76)$$

To summarize:

$$\begin{aligned} \mu_{u_t} &= \Sigma_{u_t}(\nu_{u_t} + F_t^\top(\Sigma_\xi + E_t^\top \Sigma_{x_t} E_t)^{-1}(z_t - E_t \mu_{x_t} - e_t) \\ &\quad + B_t^\top(\Gamma_{t+1} \nu_{x_{t+1}} + (I - \Gamma_{t+1}) \nu_{x_t}''') \\ &\quad - \Gamma_{t+1} \Lambda_{x_{t+1}} \Psi_{t+1}(A\mu_{x_t} + a)) \end{aligned} \quad (77)$$

To find Σ_{u_t} ,

$$\Sigma_{u_t} = (\Lambda_{u_t}'' + \Lambda_{u_t}')^{-1} \quad (78)$$

$$\Sigma_{u_t} = (\Lambda_{u_t}'' + F_t^\top \Lambda_\xi F_t + B_t^\top(\Sigma_{x_{t+1}} + \Sigma_{x_t}''')^{-1} B_t)^{-1} \quad (79)$$

Using the matrix inversion identity $(A^{-1} + B^{-1})^{-1} = B(A+B)^{-1}A$ [51],

$$\begin{aligned} \Sigma_{u_t} &= (\Lambda_{u_t}'' + F_t^\top(\Sigma_\xi + E_t^\top \Sigma_{x_t} E_t)^{-1} F_t \\ &\quad + B_t^\top \Lambda_{x_t}'''(\Lambda_{x_t}''' + \Lambda_{x_{t+1}})^{-1} \Lambda_{x_{t+1}} B_t)^{-1} \end{aligned} \quad (80)$$

$$\Sigma_{u_t} = (\Lambda_{u_t}'' + F_t^\top(\Sigma_\xi + E_t^\top \Sigma_{x_t} E_t)^{-1} F_t + B_t^\top \Gamma_{t+1} \Lambda_{x_{t+1}} B_t)^{-1} \quad (81)$$

Interpreting the scale matrices Γ and Ψ

Deriving the controller lead to the emergence of two scale matrices Γ and Ψ , representing matrix fractions of the forward messages (uncertainty) and backward messages (optimality). To interpret

the meaning of these terms, there are four scenarios that are important to consider: high process uncertainty ($\Lambda_{\vec{x}_t} \rightarrow 0$), low process uncertainty ($\Lambda_{\vec{x}_t} \rightarrow \infty$), high input prior uncertainty ($\Lambda_{\vec{u}_t} \rightarrow 0$) and low input prior uncertainty ($\Lambda_{\vec{u}_t} \rightarrow \infty$). Note, ‘process uncertainty’ includes accumulated uncertainty from previous timesteps, including that from the input priors. The input prior described above is specific to that timestep.

1. **High process uncertainty, high input prior uncertainty**

Here $\Gamma_{t+1} \rightarrow 0$, $\Psi_{t+1} \rightarrow 0$ and $\Lambda_{\vec{x}_t} \rightarrow 0$. Therefore the controller becomes cut off from the backward messages (i.e. any sense of optimality) and becomes a weighted average of its prior and goal state.

2. **High process uncertainty, low input prior uncertainty**

Here $\Gamma_{t+1} \rightarrow 0$, $\Psi_{t+1} \rightarrow \Gamma_{t+1}^{-1}$ and $\Lambda_{\vec{x}_t} \rightarrow \Lambda_{\vec{x}_{t+1}}$. Despite the system uncertainty, the controller confidence reactivates the control terms by cancelling out Γ .

3. **Low process uncertainty, high input prior uncertainty**

Here $\Gamma_{t+1} \rightarrow I$, $\Psi_{t+1} \rightarrow I$ and $\Lambda_{\vec{x}_t} \rightarrow 0$. This is the equivalent LQR setting, assuming the deterministic controller is used (see Section B.5).

4. **Low process uncertainty, low input prior uncertainty**

Here $\Gamma_{t+1} \rightarrow I$, $\Psi_{t+1} \rightarrow I$ and $\Lambda_{\vec{x}_t} \rightarrow \Lambda_{\vec{x}_{t+1}}$. As above this is similar to the LQR setting, however now the controller update will be closer to its prior.

B.5 Equivalence to the Dynamic Programming LQR Solution

First, remembering that the backwards message correspond to likelihoods, the log-likelihood of a Gaussian distribution is,

$$\log \mathcal{N}(x; \mu, \Sigma) = (x - \mu)^\top \Sigma^{-1} (x - \mu) + \text{constant} \quad (82)$$

$$= x^\top \Sigma^{-1} x - 2x^\top \Sigma^{-1} \mu + \text{constant} \quad (83)$$

$$= x^\top \Lambda x - 2x^\top \nu + \text{constant} \quad (84)$$

By comparing this to the LQR value function in Equation 12, the equivalence between Λ , P , $-\nu$ and p outlined in Table 1 may be appreciated.

To arrive at the recursive LQR expressions outlined in Section A we must consider linear models, deterministic dynamics and infinitely broad priors, which requires $\Sigma_v \rightarrow 0$ and $\Lambda_{\vec{u}_t} \rightarrow 0$. Additionally, from the formulation outlined in Section 2.1, $E^\top \Lambda_\xi E = \alpha Q$ and $F_t^\top \Lambda_\xi F_t = \alpha R$. To recover the LQR result, we require the observation likelihood to dominate, which occurs for suitable large α . Moreover, given large input priors, we require α such that $(\Sigma_\xi + F_t \Sigma_{\vec{u}_t} F_t^\top)^{-1} \approx \Lambda_\xi$, therefore $\alpha \rightarrow \infty$ as $\Lambda_{\vec{u}_t} \rightarrow 0$. As the value function parameters scale linearly with the cost function and the controller is invariant to the scale, we can omit α from the analysis for brevity.

Recursion of the precision

Applying these conditions to the $\Lambda_{\vec{x}_t}$ recursion in Equation 52,

$$\Lambda_{\vec{x}_t} = Q + A^\top \Lambda_{\vec{x}_{t+1}} A - A^\top \Lambda_{\vec{x}_{t+1}} ((B_t R^{-1} B_t^\top)^{-1} + \Lambda_{\vec{x}_{t+1}})^{-1} \Lambda_{\vec{x}_{t+1}} A \quad (85)$$

$$= Q + A^\top \Lambda_{\vec{x}_{t+1}} A - A^\top \Lambda_{\vec{x}_{t+1}} B (R + B^\top \Lambda_{\vec{x}_{t+1}} B)^{-1} B^\top \Lambda_{\vec{x}_{t+1}} A \quad (86)$$

Equation 85 to 86 is achieved using the identity $(A + B)^{-1} = B^{-1} (B^{-1} + A^{-1})^{-1} A^{-1}$ [51],

$$((B R^{-1} B^\top)^{-1} + \Lambda_{\vec{x}_{t+1}})^{-1} = \Sigma_{\vec{x}_{t+1}} (B R^{-1} B^\top + \Sigma_{\vec{x}_{t+1}})^{-1} B R^{-1} B^\top \quad (87)$$

along with $(A + J^\top B J)^{-1} J^\top B = A^{-1} J^\top (B^{-1} + J A J^\top)^{-1}$ [51]

$$= B (R + B^\top \Lambda_{\vec{x}_{t+1}} B)^{-1} B^\top \quad (88)$$

Recursion of the scaled-mean

Applying the conditions to the $\nu_{\bar{x}_t}$ recursion in Equation 57, along with the identity tricks used above,

$$\begin{aligned} \nu_{\bar{x}_t} &= A^\top (I - \Lambda_{\bar{x}_{t+1}} ((BR^{-1}B^\top)^{-1} + \Lambda_{\bar{x}_{t+1}})^{-1}) (\nu_{\bar{x}_{t+1}} - \Lambda_{\bar{x}_{t+1}} (BR^{-1}(Ru_g)) - \Lambda_{\bar{x}_{t+1}} a) \\ &\quad + Qx_g \end{aligned} \quad (89)$$

$$\begin{aligned} &= A^\top (I - \Lambda_{\bar{x}_{t+1}} B(R + B^\top \Lambda_{\bar{x}_{t+1}} B^\top)^{-1} B^\top) (\nu_{\bar{x}_{t+1}} - \Lambda_{\bar{x}_{t+1}} (BR^{-1}(Ru_g)) - \Lambda_{\bar{x}_{t+1}} a) \\ &\quad + Qx_g \end{aligned} \quad (90)$$

$$\begin{aligned} &= A^\top (\nu_{\bar{x}_{t+1}} - \Lambda_{\bar{x}_{t+1}} a - \Lambda_{\bar{x}_{t+1}} Bu_g - \Lambda_{\bar{x}_{t+1}} B(R + B^\top \Lambda_{\bar{x}_{t+1}} B^\top)^{-1} B^\top \\ &\quad (\nu_{\bar{x}_{t+1}} - \Lambda_{\bar{x}_{t+1}} Bu_g - \Lambda_{\bar{x}_{t+1}} a) + Qx_g \end{aligned} \quad (91)$$

Here, there is a discrepancy in the u_g terms, but this can be rectified through adding $Ru_g - Ru_g$ and rearranging,

$$\begin{aligned} &= A^\top (\nu_{\bar{x}_{t+1}} - \Lambda_{\bar{x}_{t+1}} a - \Lambda_{\bar{x}_{t+1}} Bu_g - \Lambda_{\bar{x}_{t+1}} B(R + B^\top \Lambda_{\bar{x}_{t+1}} B^\top)^{-1} \\ &\quad (B^\top \nu_{\bar{x}_{t+1}} - B^\top \Lambda_{\bar{x}_{t+1}} Bu_g + Ru_g - Ru_g - B^\top \Lambda_{\bar{x}_{t+1}} a) + Qx_g \end{aligned} \quad (92)$$

$-(R + B^\top \Lambda_{\bar{x}_{t+1}} B^\top)u_g$ can be taken outside to cancel out the existing term there, so only one u_g term remains,

$$\begin{aligned} &= A^\top (\nu_{\bar{x}_{t+1}} - \Lambda_{\bar{x}_{t+1}} a - \Lambda_{\bar{x}_{t+1}} Bu_g + \Lambda_{\bar{x}_{t+1}} Bu_g - \Lambda_{\bar{x}_{t+1}} B(R + B^\top \Lambda_{\bar{x}_{t+1}} B^\top)^{-1} \\ &\quad (B^\top \nu_{\bar{x}_{t+1}} + Ru_g - B^\top \Lambda_{\bar{x}_{t+1}} a) + Qx_g \end{aligned} \quad (93)$$

$$\begin{aligned} &= A^\top (\nu_{\bar{x}_{t+1}} - \Lambda_{\bar{x}_{t+1}} a - \Lambda_{\bar{x}_{t+1}} B(R + B^\top \Lambda_{\bar{x}_{t+1}} B^\top)^{-1} \\ &\quad (B^\top \nu_{\bar{x}_{t+1}} - B^\top \Lambda_{\bar{x}_{t+1}} a + Ru_g) + Qx_g \end{aligned} \quad (94)$$

Recall that p_t is equivalent to $-\nu_t$, so all non- ν terms should have the opposite sign to those in Eq. 17.

The linear Gaussian controller

As mentioned above, for the LQR conditions $\Gamma_{t+1} \rightarrow I$, $\Psi_{t+1} \rightarrow I$ and $\Lambda_{\bar{x}_t}''' \rightarrow 0$.

Applying the conditions to the controller,

$$\Sigma_{u_t} = (R + B_t^\top \Lambda_{\bar{x}_{t+1}} B_t)^{-1}, \quad (95)$$

$$\mu_{u_t} = -\Sigma_{u_t} (-Ru_g + B_t^\top (-\nu_{\bar{x}_{t+1}} + \Lambda_{\bar{x}_{t+1}} (A\mu_{x_t} + a))), \quad (96)$$

remembering that $\nu_{\bar{x}_{t+1}}$ is the opposite sign to p_{t+1} .

Expanding on the case of highly uncertainty, where $\Gamma_t \rightarrow 0$, here the stochastic controller is independent of the backward messages (and therefore any notion of optimality). Therefore it would depend purely on a weighted combination of its prior and goal:

$$\Sigma_{u_t} = (\Lambda_{\bar{u}_t} + \alpha R)^{-1}, \quad (97)$$

$$\mu_{u_t} = (\Lambda_{\bar{u}_t} + \alpha R)^{-1} (\nu_{\bar{u}_t} + \alpha Ru_g), \quad (98)$$

where the stationary distribution would be $\Sigma_{u_t} \rightarrow 0$ and $\mu_{u_t} \rightarrow 0$.

B.6 Hyperparameter Sensitivity

For nonlinear tasks, the crucial hyperparameters for i2C are the initial input priors $\Sigma_{\bar{u}}$ and the update limit (motivated as a KL bound) of α , δ_α . The role of $\Sigma_{\bar{u}_t}$ is to facilitate exploration, but too much uncertainty in the trajectory leads to the linearization assumption becoming invalid during inference. This failure mode manifests as the posterior inputs becoming inaccurate, therefore leading to the subsequent prior trajectory deviating from the previous posterior trajectory. This means that the predicted performance of the controller diverges from the true performance when evaluated on the actual system. Therefore, for fast and successful convergence, $\Sigma_{\bar{u}}$ depends not only on the expected input range, but should also be tuned based on the inherent uncertainty of the system and nonlinearity of the dynamics.

Even after tuning $\Sigma_{\bar{u}}$, the approximate inference can fail after aggressive updates to α in the M-Step, due to the approximate nature of the log-likelihood evaluation with the linearization assumption. A KL bound, simplified to a bound δ_α on the update ratio, smooths the optimization by limiting

Environment	z	z_g	$\Theta = \text{diag}(Q, R)$	u_{limit}	Σ_η
Pendulum	$[\sin \theta, \cos \theta, \dot{\theta}, u]^\top$	$[0, 1, 0, 0]^\top$	$\text{diag}(1, 100, 1, 1)$	$[-2, 2]$	$\text{diag}(\epsilon_1, \epsilon_3)$
Cartpole	$[x, \sin \theta, \cos \theta, \dot{x}, \dot{\theta}, u]^\top$	$[0, 0, 1, 0, 0, 0]^\top$	$\text{diag}(1, 1, 100, 1, 1, 1)$	$[-5, 5]$	$\text{diag}(\epsilon_1, \epsilon_1, \epsilon_2, \epsilon_2)$
Double Cartpole	$[x, \sin \theta_1, \cos \theta_1, \sin \theta_2, \cos \theta_2, \dot{x}, \dot{\theta}_1, \dot{\theta}_2, u]^\top$	$[0, 0, 1, 0, 1, 0, 0, 0, 0]^\top$	$\text{diag}(1, 1, 100, 1, 100, 1, 1, 1, 1)$	$[-10, 10]$	$\text{diag}(\epsilon_1, \epsilon_1, \epsilon_1, \epsilon_2, \epsilon_2, \epsilon_2)$

Table 3: Environment parameters of the nonlinear tasks. $\epsilon_1 = 1\text{e-}12$, $\epsilon_2 = 1\text{e-}6$ and $\epsilon_3 = 1\text{e-}3$.

aggressive updates. For tuning, δ_α should smooth out any large updates to α while limiting the impact to the rate of convergence. Note that as α is increasing, it is numerically easier to work with α^{-1} , which tends to zero, so the limiting is implemented in practice as $\delta_{\alpha^{-1}}$ acting on α^i/α^{i+1} . Figure 6 demonstrates the behaviour of the hyperparameters for the Cartpole swing-up task.

C Experimental Details

C.1 Equivalence with finite-horizon LQR by Dynamic Programming

$$\mathbf{x}_{t+1} = \begin{bmatrix} 1.1 & 0 \\ 0.1 & 1.1 \end{bmatrix} \mathbf{x}_t + \begin{bmatrix} 0.1 \\ 0 \end{bmatrix} \mathbf{u}_t + \begin{bmatrix} -1 \\ -2 \end{bmatrix} \quad (99)$$

$$Q = \begin{bmatrix} 10 & 0 \\ 0 & 10 \end{bmatrix}, \quad R = [1], \quad \mathbf{x}_g = \begin{bmatrix} 10 \\ 10 \end{bmatrix}, \quad \mathbf{u}_g = [0], \quad \alpha = 1\text{e}5, \quad \Sigma_{\vec{u}_t} = [100] \quad (100)$$

C.2 Evaluation on nonlinear trajectory optimization tasks

Both iLQR and GPS required the cost function in Table 3 to be scaled in order to have good numerics. In Table 5 and 6 we refer to this has α (as it performs the same role as the i2C parameter). Additionally, iLQR and GPS were able to optimize without a random initialization. In order to compare with i2C, which initializes by design with fixed priors, the random initialisation was set to have a smallest amplitude that allowed optimization to take place. All algorithms achieved faster converged with random initialisation, however such ‘warm start’ strategies were not the focus of this work, instead we wished to focus on i2C strength in deterministic initialisation. For these experiments we use the terminal cost $Q_f = Q$.

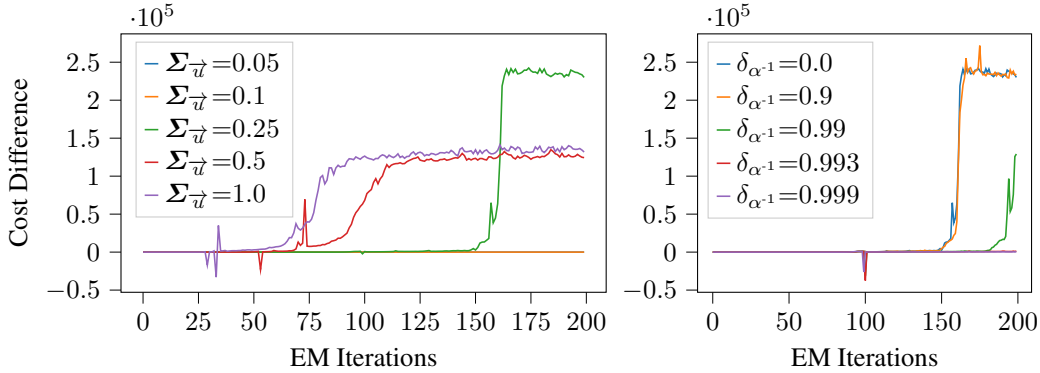


Figure 6: Difference between evaluated cost and predicted cost over EM iterations across hyperparameters for the Cartpole swing-up task.

Environment	$\Sigma_{\vec{u}}$ (init.)	α (init.)	$\delta_{\alpha^{-1}}$
Pendulum	[0.2]	1/100	0.99
Cartpole	[0.25]	1/67	0.993
Double Cartpole	[0.04]	1/90	0.9995

Table 4: i2C parameters for the nonlinear trajectory optimization tasks.

Environment	λ range	$\lambda_{\text{multiplier}}$	σ_k (init.)	α
Pendulum	[1 - 1e-9]	1.002	1e-2	1e-4
Cartpole	[1 - 1e-7]	1.001	1e-2	1e-3
Double Cartpole	[1 - 1e-7]	1.001	1e-2	1e-3

Table 5: iLQR parameters for the nonlinear trajectory optimization tasks.

Environment	Σ_{Explore}	KL bound	σ_k (init.)	α
Pendulum	[2.0]	0.07	1e-2	1e-4
Cartpole	[1.25]	1.0	1e-1	1e-3
Double Cartpole	[5.0]	0.75	1e-1	1e-3

Table 6: GPS parameters for the nonlinear trajectory optimization tasks

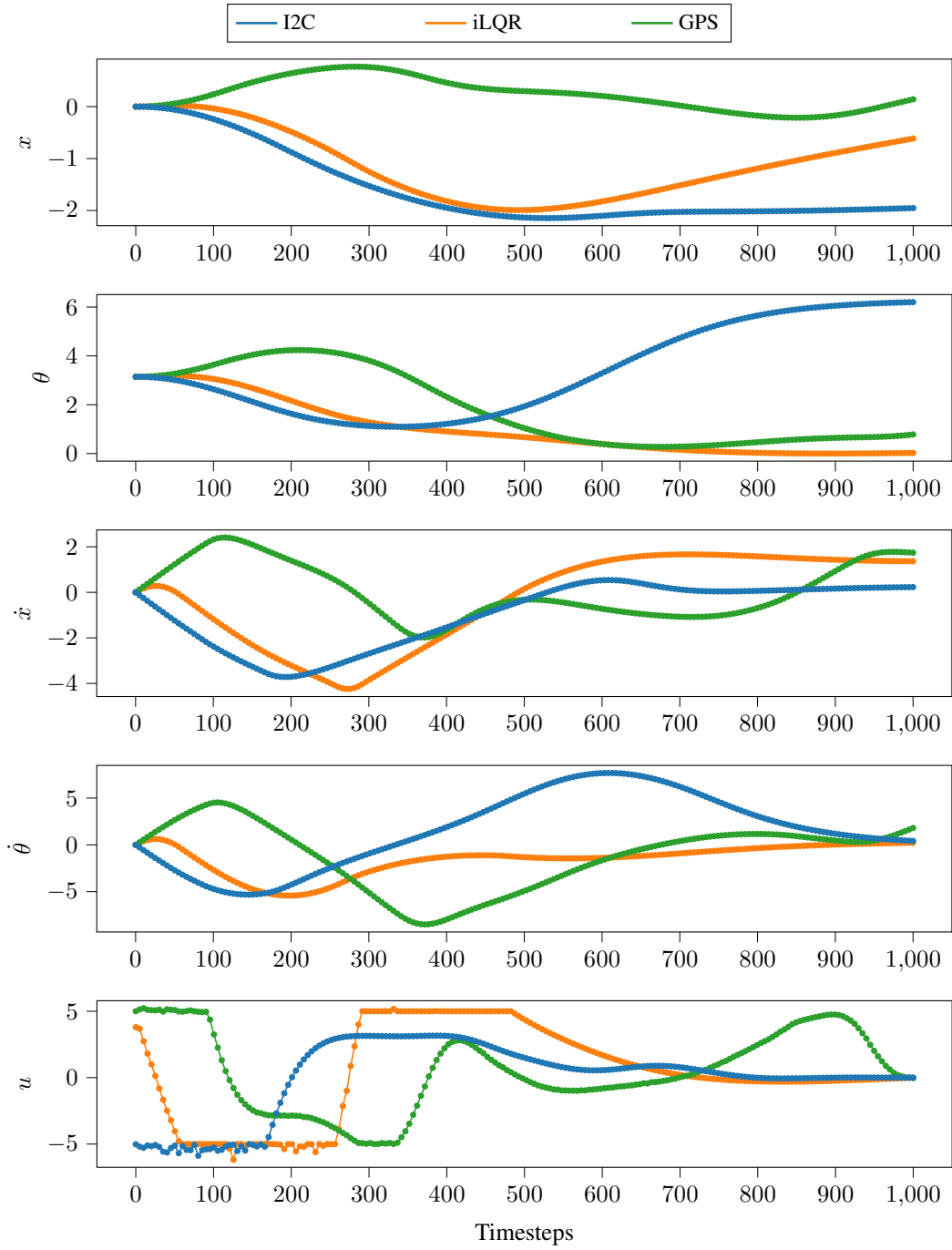


Figure 7: Comparison of the state-action trajectories of I2C, iLQR and GPS on the Cartpole swing-up task after convergence.

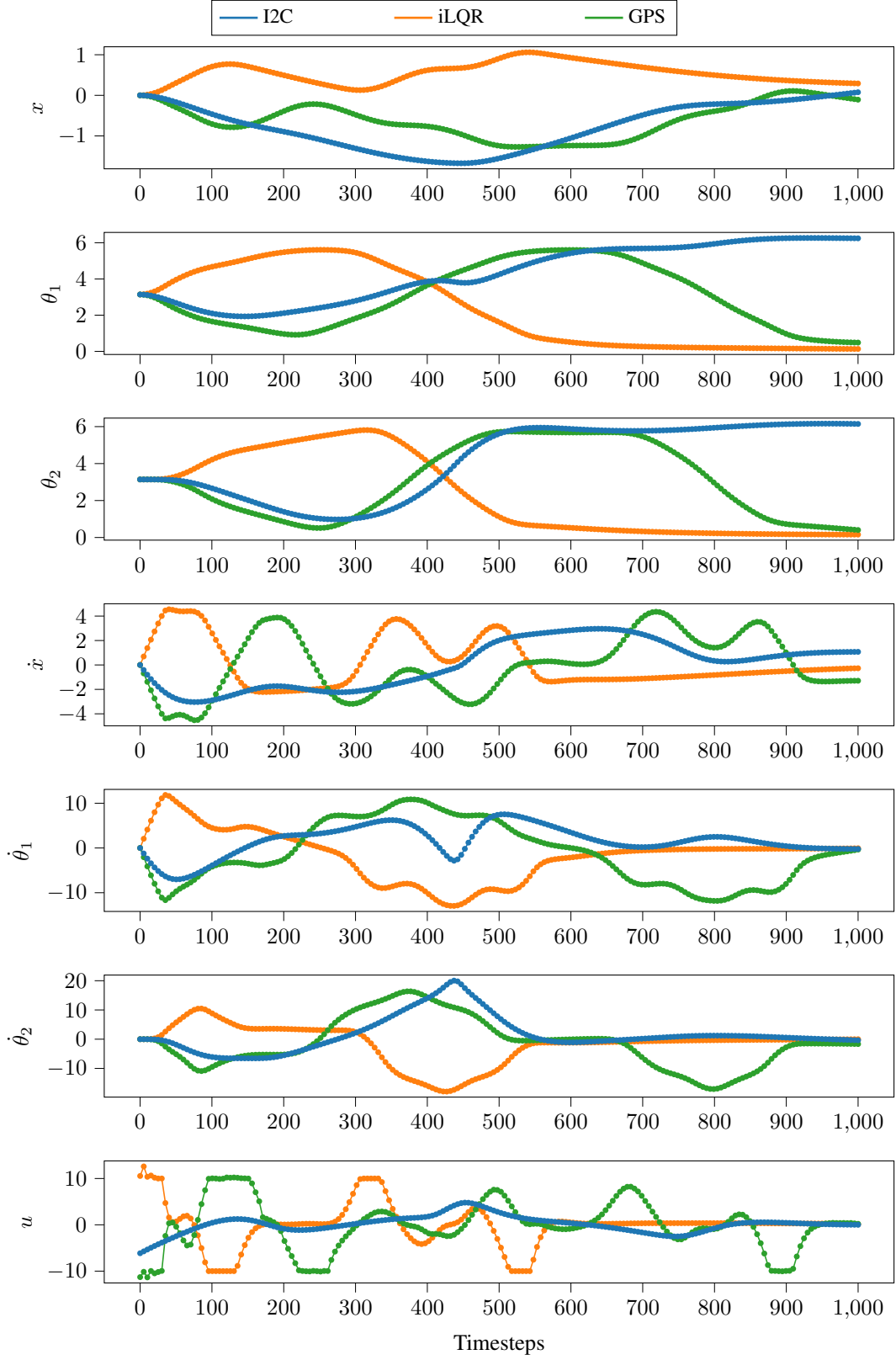


Figure 8: Comparison of the state-action trajectories of I2C, iLQR and GPS on the Double Cartpole swing-up task after convergence.