

On Learning-based Control of Dynamical Systems

Application to Fluid Dynamics



université
PARIS-SACLAY

LISN
LABORATOIRE INTERDISCIPLINAIRE
DES SCIENCES DU NUMÉRIQUE

anr
agence nationale
de la recherche

Rémy Hosseinkhan-Boucher^{1,2}

PhD Defense - April 10th 2025

Advisors:

Anne Vilnat^{1,2}

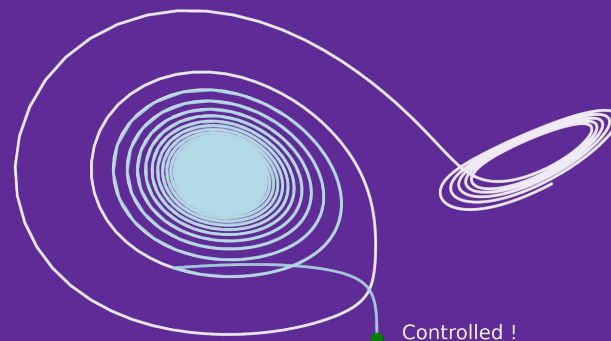
Onofrio Semeraro²

Lionel Mathelin²

¹Université Paris-Saclay

²LISN, CNRS

— Without control
— With control



Dynamical Systems

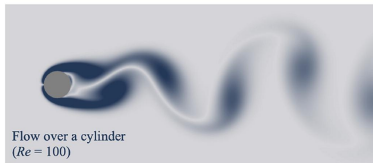
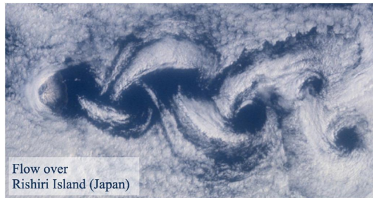
→ Termed by Poincaré in “*Les méthodes nouvelles de la mécanique céleste*” (1892)

Definition

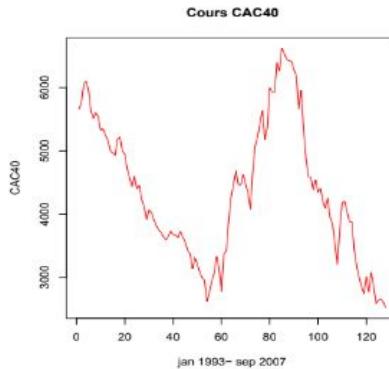
System + **time-evolution law**

Examples of fields

Fluids dynamics Financial markets Games Chemistry Astronomy Epidemiology Autonomous Driving
Natural Language Processing



Von Kármán vortex street generated by an island in Japan
Taira *et al.* “Modal Analysis of Fluid Flows: Applications and Outlook”, *AIAA Journal* (2019)



CAC40 Index between 1993 and 2007
I. Kharroubi - Gestion de Portefeuilles (2015)



Motion of a particle in the Bunimovich billiard
Credits: George Datseris, Wikipedia (2019)

(Optimal) Control

→ *Optimal control emerged post-WWII from automatic control needs in flight systems.*

Control

Force the system **initial state** → **final state**

Optimal Control

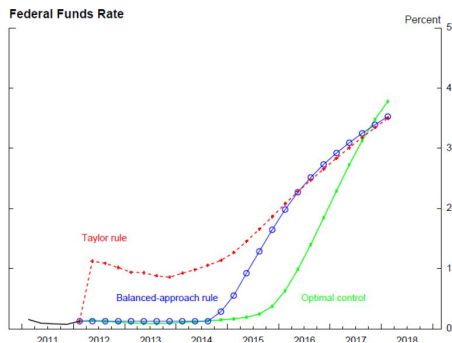
Control + **optimise** a **criterion** J

Control Policy π

Behaviour of the control input

Other Examples

$$J = \min[\text{Cash}], \min[\text{Time}], \min[\text{Energy}], \min[\text{Unemployment}], \max[\text{Score}]$$

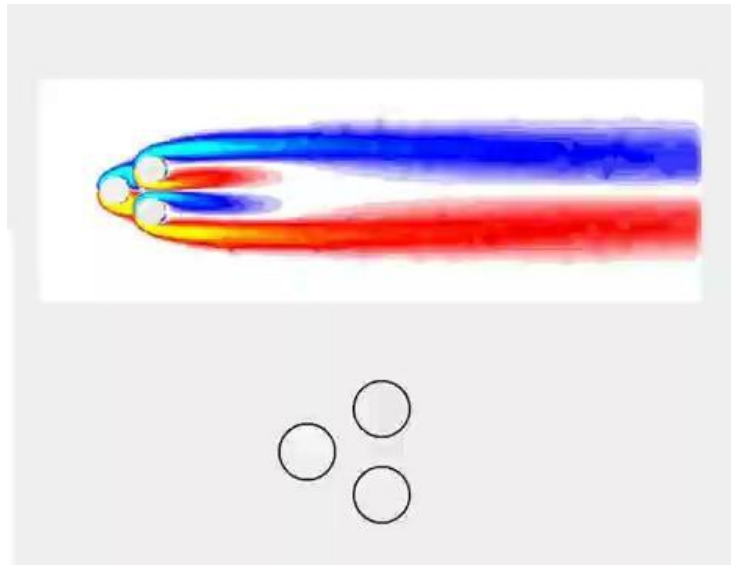


*Optimal Control of
Monetary Policy*

J. L. Yellen
“Perspectives on
Monetary Policy”
Boston Economic Club
(2012)

Application in Fluids Dynamics

$$J = \max[\text{Lift}], \min[\text{Drag}], \min[\text{Noise}]$$



Fluidic Pinball under control laws - Credits: E. Kaiser

Physical Model-based vs. Learning-based Control

Physical Model

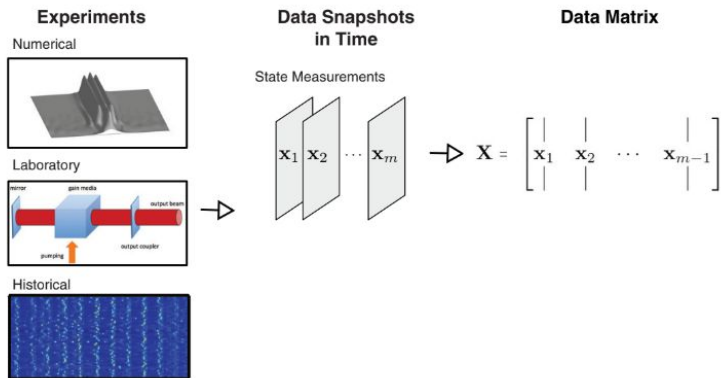
Explicit system representation (e.g. equation) → **Build** control policy π

Example (Navier-Stokes equation in Fluids Dynamics) →

$$\frac{\partial \mathbf{u}}{\partial t} + (\mathbf{u} \cdot \nabla) \mathbf{u} = -\nabla p + \frac{1}{\text{Re}} \Delta \mathbf{u},$$
$$\nabla \cdot \mathbf{u} = 0.$$

Learning-based Model

Implicit system representation from **data and statistics** → **Build** control policy π



Data collection → **system interaction**

Data-driven methods → Machine Learning Control

Data can be collected from a number of different sources

J. N. Kutz et al. - Dynamic Mode Decomposition: Data Driven Modeling Of Complex Systems, *SIAM* (2016)

Dynamical Systems Control: Challenges

Optimal Control Problem

Dynamics

$$\partial_t x(t) = f(x(t), u(t))$$

Goal

$$\min_u J(u) = \int_0^T c(x(t), u(t)) dt$$

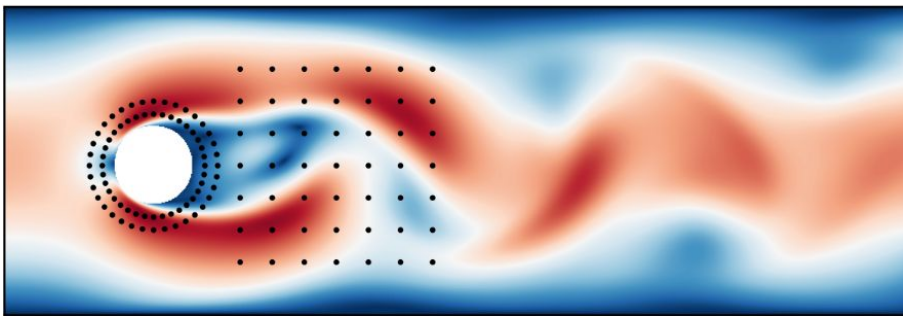
Challenges¹

- **Robustness**
- **Sample complexity**
- **Partial observability (PO) and delays**
- Controllability
- High dimensional state space
- Very large degrees of freedom (e.g. sensors & actuators config.)

Example

$f \rightarrow$ Navier-Stokes equation

Energy cost $\rightarrow c(x, u) = \|x\|^2 + \|u\|^2$



Cylinder flow drag reduction. Partial observation through sensors. Illustration from ¹.

PhD thesis Goal

- Addressing the challenges
- Combine concepts from different fields

¹J. Viquerat et al. "A review on deep reinforcement learning for fluid mechanics: An update", AIP Publishing (2022)

Presentation Outline

- Maximum Entropy: Noise Robustness
- Sampling Strategies with Semi-Markov Decision Process
- Towards Neural Controlled Delay Differential Equations
- *Academic and Scientific Involvement*
- Conclusion

Maximum Entropy: Noise Robustness

Entropy = Uncertainty Measure

How to quantify the **uncertainty** on the control law π ?

$du \rightarrow$ Infinitesimal volume (a.k.a. event) of $\mathcal{U}$

Uncertainty on du

$$I(du) = \log \left(\frac{1}{\pi(du)} \right)$$

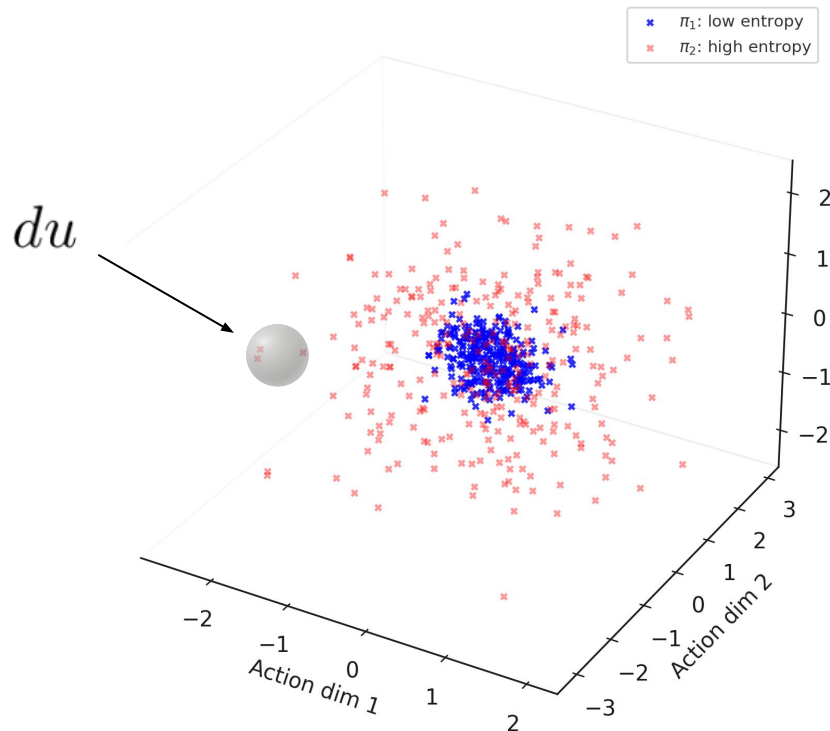
Entropy (average uncertainty)

$$\mathcal{H}(\pi) = \int_{\mathbb{R}} \log \left(\frac{1}{\pi(du)} \right) \pi(du)$$

Gaussian case

$$\pi(du) = f_{\mathcal{N}(\mu, \Sigma)}(u) du$$

$$\mathcal{H}(\pi) = \frac{1}{2} \log (2\pi e \Sigma)$$



Robustness: Maximum Entropy and Regularity

Maximum Entropy Reinforcement Learning

$$\arg \min_{\pi} \mathbb{E}^{\pi} \left[\sum_{k=0}^T \gamma^k c(X_k, U_k) - \alpha \mathcal{H}[\pi(du | X_k)] \right], \quad \alpha > 0,$$

Previous Results¹

- Better exploration
- **Robustness**
- **Loss regularisation**

Control objective

Entropy

Questions

Why does entropy **improve robustness**?

Why does entropy **regularise** the optimisation landscape?

Hypothesis

Entropy $\rightarrow$ Policy complexity

¹A. Ahmed *et al.* "Understanding Flat Minima in Neural Networks", *ICLR* (2019)

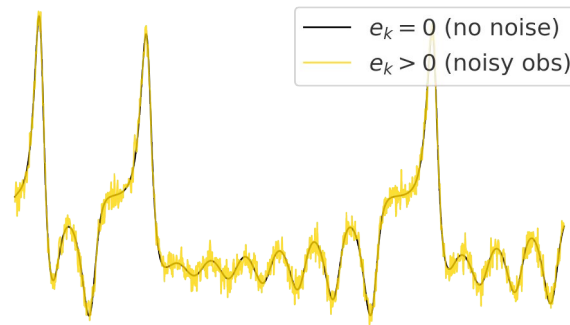
Robustness Measure

Noisy Observable

$$X_{k+1} = F(X_k, U_k)$$

$$Y_{k+1} = G(X_k) + \epsilon_k \quad \epsilon_k \sim \mathcal{N}(0, \sigma_\epsilon^2 I_d)$$

→ Feedback observation policy $\pi : \mathcal{Y} \mapsto \pi(Y_{k+1}, du)$



Noisy observation of the 1st coordinate of the Lorenz system

Rate of Excess Risk Under Noise¹

$$\epsilon \equiv 0 \longrightarrow \mathbb{P}^\pi$$

$$\epsilon \not\equiv 0 \longrightarrow \mathbb{P}^{\pi, \epsilon}$$

$$J^{\pi, \epsilon} = \mathbb{E}^{\pi, \epsilon} \left[\sum_{k=0}^T \gamma^k c(X_k, U_k) \right]$$

$$\mathcal{R}^\pi = \frac{J^{\pi, \epsilon} - J^\pi}{J^\pi}$$

→ How much the objective function deteriorates when $\epsilon \not\equiv 0$

Hypothesis

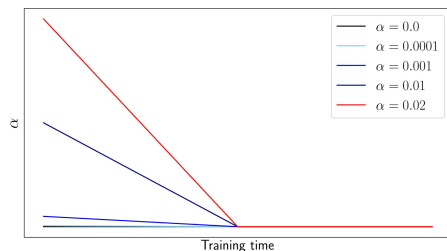
$$\begin{array}{ll} \epsilon \nearrow & \longrightarrow J^{\pi^*, \epsilon} \nearrow \\ \alpha > 0 & \longrightarrow \mathcal{R}^{\pi^\alpha} \searrow \end{array}$$

¹R. Hosseinkhan-Boucher et al. "Evidence on the Regularisation Properties of Maximum-Entropy Reinforcement Learning", *Optimization and Learning Conference* (2024)

Robustness Measure: Experiments

Experiment

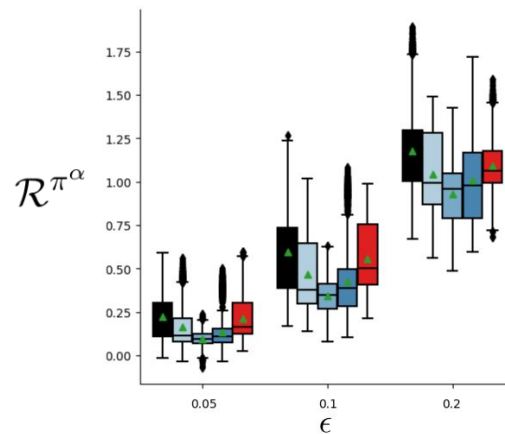
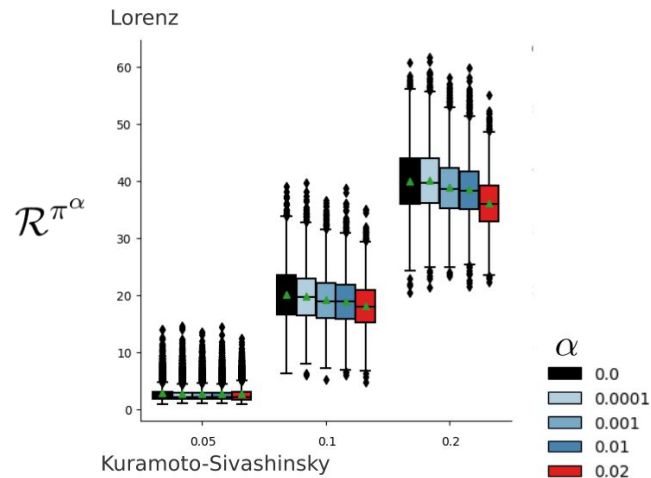
- Train 10 seeds x 5 entropy levels α



- Test over different noise intensities ϵ

Results

$$\begin{array}{lcl} \epsilon \nearrow & \longrightarrow & J^{\pi^*, \epsilon} \nearrow \\ \alpha > 0 & \longrightarrow & \mathcal{R}^{\pi^\alpha} \searrow \end{array}$$



Model Regularity

Complexity Measure¹

$$\mathcal{M}: \pi \in \Pi \rightarrow \mathbb{R}_+$$

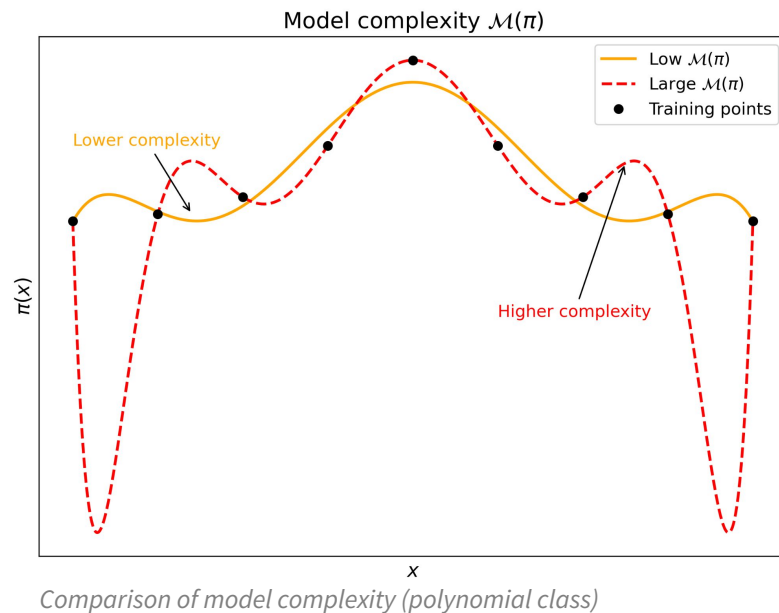
$\mathcal{M}(\pi)$ measures the **model complexity**

Robustness Measure

$$\mathcal{R}^\pi \leq \text{Bound}(\mathcal{M}(\pi))$$

Question

Which complexity measures characterise robustness to noise?



¹B. Neyshabur et al. "Exploring Generalization in Deep Learning", *NIPS* (2017)

Complexity Measure: Lipschitz Upper Bound

Policy $\rightarrow \pi_{\theta}(\cdot|x) \sim \mathcal{N}_{d_U}(\mu_{\theta}(x), \theta_{\sigma_{\pi}} I_{d_U})$

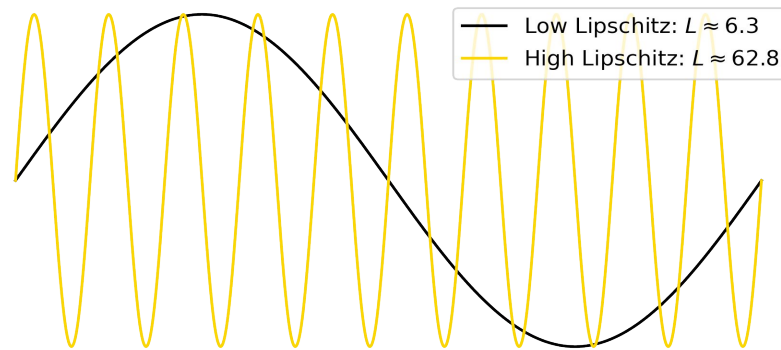
Neural Network $\rightarrow \mu_{\theta}(x) = (\sigma_l \circ \theta_l \circ \dots \circ \sigma_1 \circ \theta_1)(x)$

Lipshitz Bound

$$\text{Lips}(\mu_{\theta}) \leq \prod_{i=1}^l \underbrace{\text{Lips}(\sigma_i)}_{=1} \prod_{i=1}^l \underbrace{\text{Lips}(\theta_i)}_{=\|\theta_i\|} = \prod_{i=1}^l \|\theta_i\|$$

Lipshitz-based Complexity Measure

$$\mathcal{M}(\pi_{\theta}) = \mathcal{M}(\mu_{\theta}) = \prod_{i=1}^l \|\theta_i\|$$



Comparison of Lipschitz constant for a trigonometric class of functions

Complexity Measure: Lipschitz Upper Bound

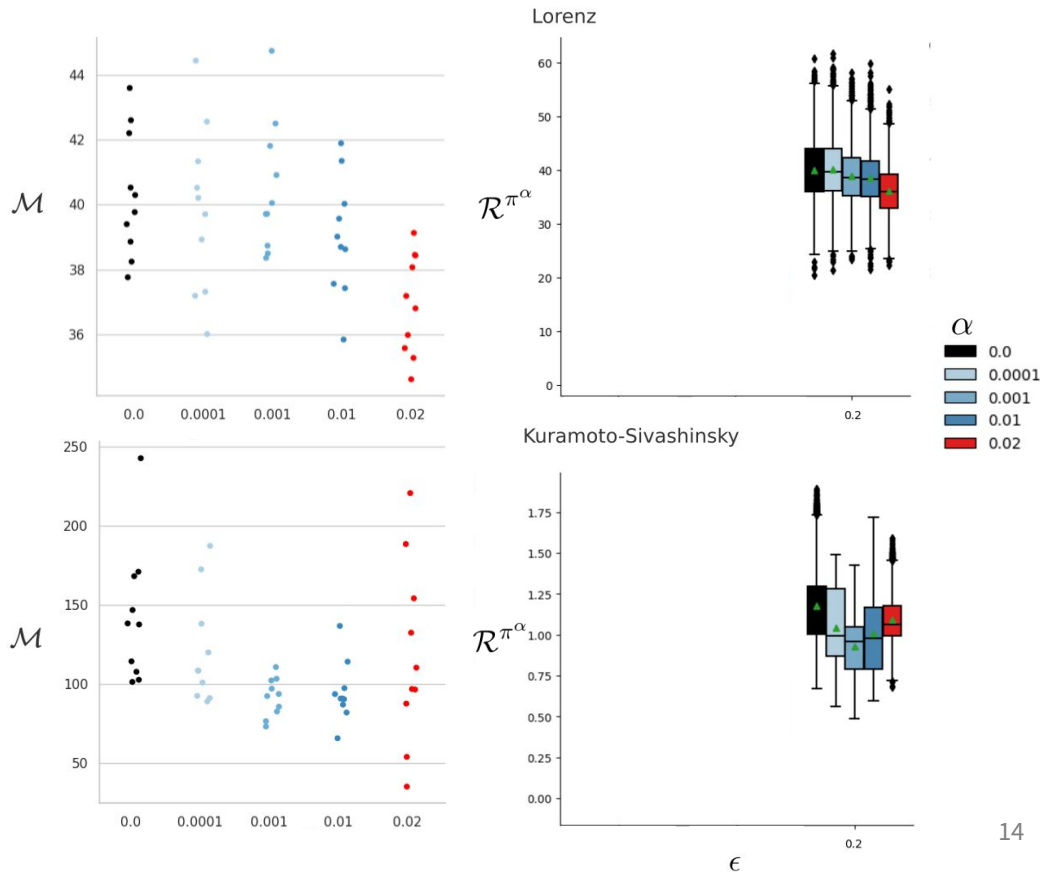
Results

$\alpha > 0 \longrightarrow \mathcal{M}(\pi_{\theta}^{\alpha}) \searrow$
(Up to a threshold for KS)

$\mathcal{R}^{\pi^{\alpha}}$ and $\mathcal{M}(\pi_{\theta})$ have **similar trend**

Lipshitz-based Complexity Measure

$$\mathcal{M}(\pi_{\theta}) = \mathcal{M}(\mu_{\theta}) = \prod_{i=1}^l \|\theta_i\|$$



Contributions and Perspectives

Contributions

- Robustness Metric: **Excess Risk Under Noise**
- Complexity Measure: **Model Regularity**
- Complexity Measure: **Optim. Landscape**

Results

Entropy → Robustness to noise

Entropy → **Regularity**

Perspectives

Mathematical analysis (theoretical Deep Reinforcement Learning)



Evidence on the Regularisation Properties of Maximum-Entropy Reinforcement Learning

Rémy Hosseinkhan Boucher^{1,2} , Onofrio Semeraro^{1,2} ,
and Lionel Mathelin^{1,2}

¹ Université Paris-Saclay, Orsay, France

{[onofrio.semeraro](mailto:onofrio.semeraro@upsaclay.fr), [lionel.mathelin](mailto:lionel.mathelin@upsaclay.fr)}@upsaclay.fr

² CNRS, Laboratoire Interdisciplinaire des Sciences du Numérique, Orsay, France
remy.hosseinkhan@upsaclay.fr

Abstract. The generalisation and robustness properties of policies learnt through Maximum-Entropy Reinforcement Learning are investigated on chaotic dynamical systems with Gaussian noise on the observable. First, the robustness under noise contamination of the agent's observation of entropy regularised policies is observed. Second, notions of statistical learning theory, such as complexity measures on the learnt model, are borrowed to explain and predict the phenomenon. Results show the existence of a relationship between entropy-regularised policy optimisation and robustness to noise, which can be described by the chosen complexity measures.

Optimization and Learning: 7th International Conference, Revised Selected Papers (2024)

Learning-Based Control: Sampling Strategies with Semi-Markov Decision Process

Model-based Control: Gaussian Process Modelling¹

Controlled Markov Chain

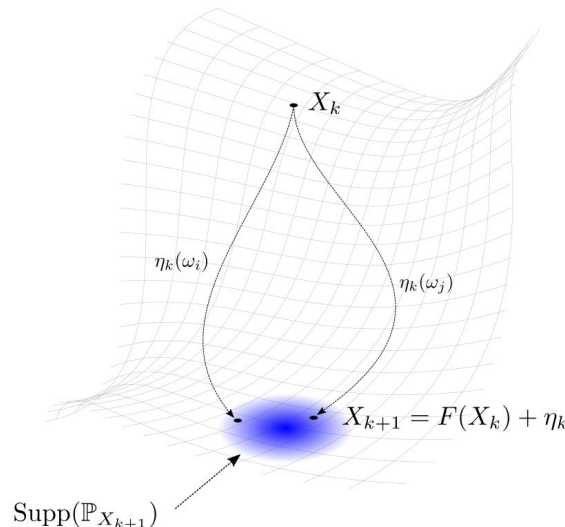
$$\mathbb{P}^\pi(dx_0 du_0 dx_1 du_1 \dots dx_T) = \mathbb{P}_{X_0}(dx_0) \pi(x_0, du_0) \mathcal{P}(dx_1 | x_0, u_0) \\ \pi(x_1 | du_1) \dots \pi(x_{T-1} | du_{T-1}) \mathcal{P}(dx_T | x_{T-1}, u_{T-1})$$

Transition Kernel $\mathcal{P}$

$$(x, u) \in \mathcal{X} \times \mathcal{U} \rightarrow \mathcal{P}(dx' | x, u) \text{ Probability on the next state } x'$$

Learning Dynamics with Gaussian (Spatial) Process¹

$$\hat{\mathcal{P}}_{\mathcal{D}}(\cdot, (x, u)) \sim \mathcal{N}(\mu_{(x,u)}, \Sigma_{(x,u), (x,u)} | \mathcal{D})$$



Learning-based Model Predictive Control

Model Predictive Control¹

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

$$s.t. \quad (u_0^*, \dots, u_{K^{\text{MPC}}}^*) = \arg \min_{(u_0, \dots, u_{K^{\text{MPC}}})} \mathbb{E}^{(u_0, \dots, u_{K^{\text{MPC}}})} \left[\sum_{k=0}^{K^{\text{MPC}}} c(X_k, u_k) \mid X_0 = x \right]$$
$$s.t. \quad X_{k+1} \sim \hat{\mathcal{P}}_{\mathcal{D}}(X_k, u_k)$$

Problem (Model Learning)

Fixed sampling budget $\rightarrow n$

$$\mathcal{D}_n = \{(x_0, u_0, x_1), \dots, (x_{n-1}, u_{n-1}, x_n)\}$$

Learn $\hat{\mathcal{P}}_{\mathcal{D}} \simeq \mathcal{P}$

Question

Which **online sampling strategy**?

¹L. Grune, J. Pannek - Nonlinear Model Predictive Control, Springer (2011)

Entropy Map

How to quantify the uncertainty on $X_{k+1} \sim \mathbb{P}_{X_{k+1}}?$

Infinitesimal volume element of $\mathcal{X} \rightarrow dx$

Uncertainty on dx

$$I(dx) = \log\left(\frac{1}{\mathbb{P}_{X_{k+1}}(dx)}\right)$$

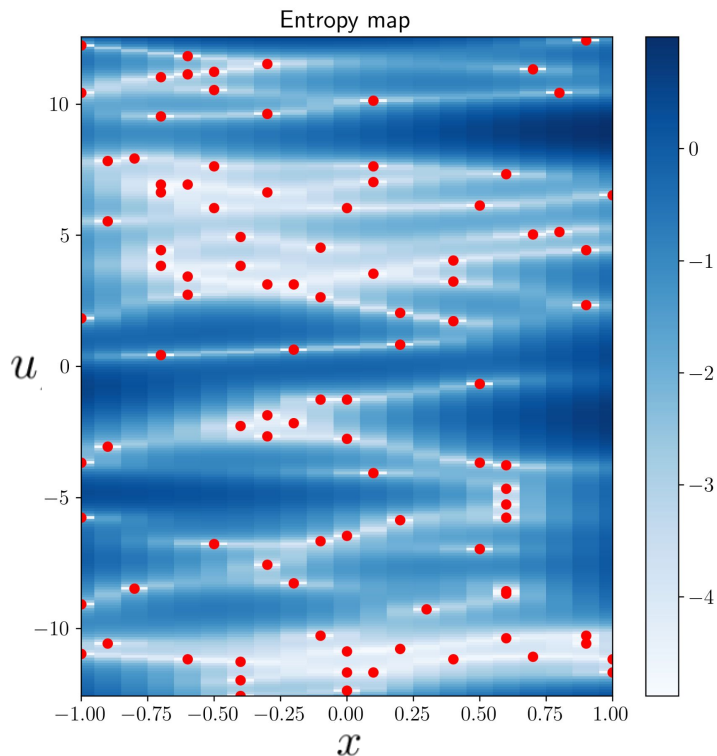
Entropy (average uncertainty)

$$\mathcal{H}(\mathbb{P}_{X_{k+1}}) = \int_{\mathbb{R}} \log \frac{1}{\mathbb{P}_{X_{k+1}}(dx)} \mathbb{P}_{X_{k+1}}(dx)$$

Gaussian case

$$\mathbb{P}_{X_{k+1}}(dx) = f_{\mathcal{N}(\mu, \Sigma)}(x)dx$$

$$\mathcal{H}(\mathbb{P}_{X_{k+1}}) = \frac{1}{2} \log(2\pi e \Sigma)$$



Red dots are observation contained in $\mathcal{D}_n$

$$\hat{\mathcal{P}}_{\mathcal{D}}(\cdot, (x, u)) \sim \mathcal{N}(\mu_{(x,u)}, \Sigma_{(x,u), (x,u)} \mid \mathcal{D})$$

Expected Information Gain¹

Dataset Construction (Sampling)

How to select the next data point (x, u) ?

$$\mathcal{D}_{n+1} = \mathcal{D}_n \cup (x, u, X_{n+1})$$

Process trajectory $\rightarrow H_T = (X_0, U_0, \dots, U_T, X_T)$

Optimal trajectory under $\hat{\mathcal{P}}_{\mathcal{D}} \rightarrow \hat{H}_T^*$

Expected Information Gain (EIG) on the Optimal Trajectory²

$$\text{EIG}(x, u) = \mathcal{H}[\hat{H}_T^* \mid \mathcal{D}_n] - \mathbb{E}_{\mathbb{P}_{X_{n+1} \mid \mathcal{D}_n, X_n=x, U_n=u}} \left[\mathcal{H}[\hat{H}_T^* \mid \underbrace{\mathcal{D}_n, X_n=x, U_n=u, X_{n+1}}_{\mathcal{D}_{n+1}}] \right]$$

By symmetry $\rightarrow$ Uncertainty on X_{n+1}

$$\text{EIG}_n(x, u) = \mathcal{H}[X_{n+1} \mid \mathcal{D}_n, X_n=x, U_n=u] - \mathbb{E}_{\mathbb{P}_{\hat{H}_T^* \mid \mathcal{D}_n}} \left[\mathcal{H}[X_{n+1} \mid \mathcal{D}_n, X_n=x, U_n=u, \hat{H}_T^*] \right]$$

¹D. V. Lindley "On a measure of the Information Provided by an Experiment", *Chapel Hill and Berkley meetings of the Institute of Mathematical Statistics* (1955)

²V. Mehta et al. "An Experimental Design Perspective on Model-Based Reinforcement Learning", *ICLR* (2022)

Decision Epochs: Temporal Abstraction with Options¹

Question

Data $\mathcal{D}_n = \{(x_0, u_0), \dots, (x_n, u_n)\}$ is **collected online along the dynamics**

Which **online sampling strategy**?

Hypothesis

Exploit the **auto-correlation** of (X_{n+1}, U_{n+1}) from $\mathcal{D}_n$

Temporal Abstraction

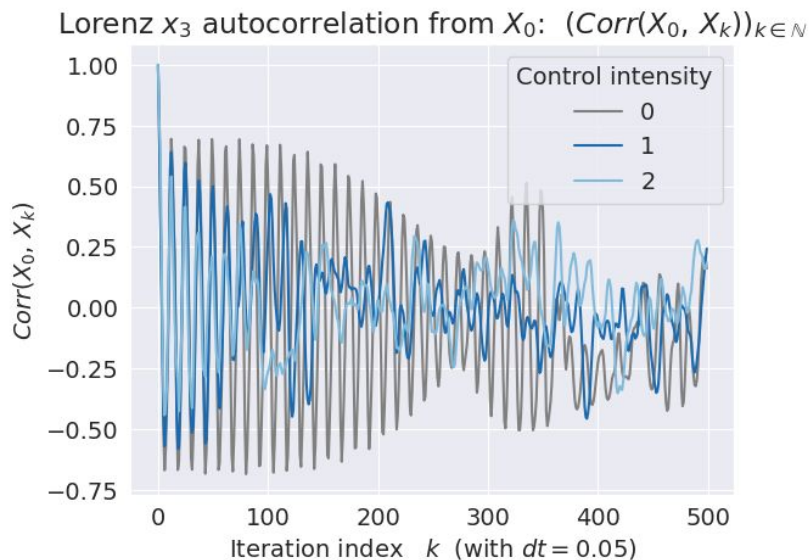
Irregular (Optional) Decision Epochs $\rightarrow (\kappa_j)_{j \in \mathbb{N}}$

Semi-Markov Decision Process $\rightarrow (X_{\kappa_j})_{j \in \mathbb{N}}$

Interdecision delay $\rightarrow \tau \in \mathcal{T}$

Constant control during $\tau \in \mathcal{T}$

New control space $\rightarrow \mathcal{U} \times \mathcal{T}$



¹R. S. Sutton *et al.* "Between MDPs and semi-MDPs: A Framework for Temporal Abstraction in Reinforcement Learning", *NIPS* (1999)

Semi-Markov Expected Information Gain

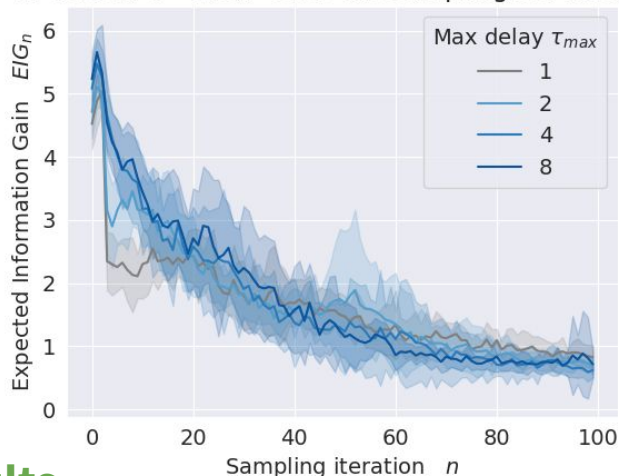
New criterion¹

$$\text{EIG}^{\text{New}}(x, (u, \tau)) = \mathcal{H}[X_{\kappa_n + \tau + 1} \mid \mathcal{D}_n, X_{\kappa_n + \tau} = x, U_{\kappa_n + \tau} = u, \kappa_n] \\ - \mathbb{E}_{\mathbb{P}_{\hat{H}_T^* \mid \mathcal{D}_n}} \left[\mathcal{H}[X_{\kappa_n + \tau + 1} \mid \mathcal{D}_n, X_{\kappa_n + \tau} = x, U_{\kappa_n + \tau} = u, \hat{H}_T^*, \kappa_n] \right]$$

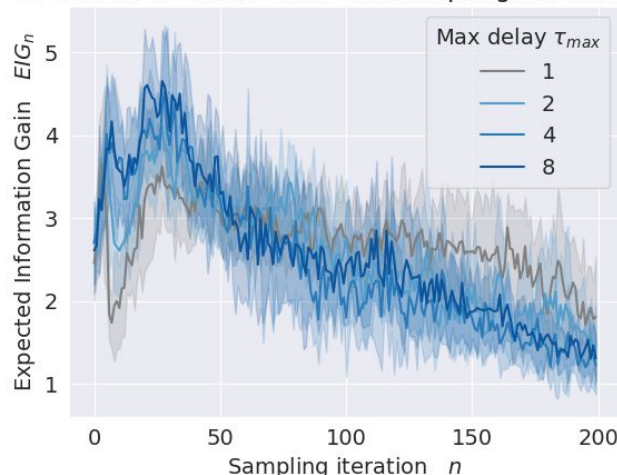


Query the **interdecision time** to maximise information!

Evolution of EIG_n over the sampling iterations n



Evolution of EIG_n over the sampling iterations n



Results

Temporal Abstraction $\rightarrow$ Information Gain

¹R. Hosseinkhan-Boucher et al. "Increasing Information for Model Predictive Control with Semi-Markov Decision Processes", *L4DC* (2024)

Contributions and Perspectives

Contributions

- Temporal Abstraction (Options framework)
- Extension of Information-based acquisition funct.

Results

Temporal Abstraction + EIG $\rightarrow$ Information gain
Information gain $\rightarrow$ Control performances

Perspectives

Mathematical analysis (theoretical Gaussian Process MPC)

Proceedings of Machine Learning Research vol 242:1400–1414, 2024 6th Annual Conference on Learning for Dynamics and Control

Increasing Information for Model Predictive Control with Semi-Markov Decision Processes

Rémy Hosseinkhan-Boucher*
Stella Douka*
Onofrio Semeraro
Lionel Mathelin

REMY.HOSSEINKHAN@UNIVERSITE-PARIS-SACLAY.FR
DOUKA.STYLIANI@UNIVERSITE-PARIS-SACLAY.FR
ONOFRIO.SEMERARO@UNIVERSITE-PARIS-SACLAY.FR
LIONEL.MATHELIN@LISN.UPSACLAY.FR

Université Paris-Saclay, Orsay, France
CNRS, Laboratoire Interdisciplinaire des Sciences du Numérique, Orsay,
France

Editors: A. Abate, M. Cannon K. Margellos, A. Papachristodoulou

Abstract

Recent works in Learning-Based Model Predictive Control of dynamical systems show impressive sample complexity performances using criteria from Information Theory to accelerate the learning procedure. However, the sequential exploration opportunities are limited by the system local state, restraining the amount of information of the observations from the current exploration trajectory. This article resolves this limitation by introducing temporal abstraction through the framework of Semi-Markov Decision Processes. The framework increases the total information of the gathered data for a fixed sampling budget, thus reducing the sample complexity.

Keywords: Expected Information Gain; Temporal Abstraction; Sample Complexity

Proceedings of the 6th Annual Learning for Dynamics & Control Conference, PMLR (2024)

Towards Neural Controlled Delay Differential Equations for Model Based Control

Learning Model-based Continuous-time Control¹

Model Learning $\partial_t x(t) = f(x(t), u(t)) \xrightarrow{\text{Neural net. } f_{\theta}}$ $\partial_t x(t) = f_{\theta}(x(t), u(t)) \quad \theta \in \Theta$

Advantages

- Temporal abstraction
- Irregularly sampled data $\rightarrow$ **Robustness**
- Model-based $\rightarrow$ **Sample efficient**

Limitations

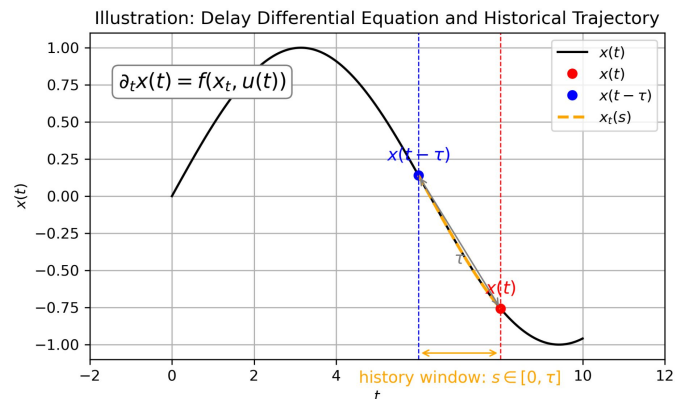
- **No partially-observed system**

$$\partial_t y(t) = g(x(t), u(t))$$

- **No delay handling**

$$\partial_t x(t) = f(x_t, u(t))$$

$$x_t : [0, \tau] \rightarrow \mathcal{X} \quad s \mapsto x(t-s) \quad \tau > 0$$



¹C. Yildiz et al. "Continuous-Time Model-Based Reinforcement Learning", ICML (2021)

Learning Model-based Continuous-time Control¹

Model Learning $\partial_t x(t) = f(x(t), u(t)) \xrightarrow[\text{Neural net. } f_{\theta}]{} \partial_t x(t) = f_{\theta}(x(t), u(t)) \quad \theta \in \Theta$

Advantages

- Temporal abstraction
- Irregularly sampled data → **Robustness**
- Model-based → **Sample efficient**

Limitations

- **No partially-observed** system

$$\partial_t y(t) = g(x(t), u(t))$$

- No **delay** handling

$$\partial_t x(t) = f(x_t, u(t))$$

$$x_t : [0, \tau] \rightarrow \mathcal{X} \quad s \mapsto x(t-s) \quad \tau > 0$$

Not Markovian

Dynamic Programming → Difficult

Delays and Partial Observability: Embedding

Solution → Markov Property in Larger Space

$H_{t+1} = (H_t, Y_{t+1}, U_{t+1}) \rightarrow$ Always Markov but **increasing dimension**

Information State^{1,2}

$$I_{t+1} = \phi(I_t, Y_{t+1}, U_{t+1})$$

$$\mathbb{P}(\cdot | H_t) = \mathbb{P}(\cdot | I_t) \rightarrow \text{Sufficient information}$$

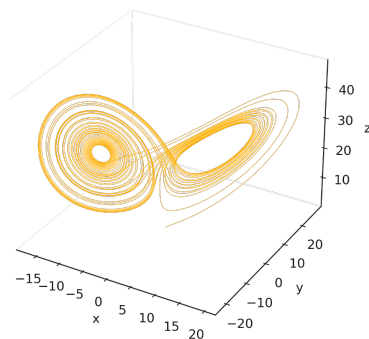
Takens theorem³

Theorem 1. Let M be a compact manifold of dimension m . For pairs (ϕ, y) , $\phi: M \rightarrow M$ a smooth diffeomorphism and $y: M \rightarrow \mathbb{R}$ a smooth function, it is a generic property that the map $\Phi_{(\phi, y)}: M \rightarrow \mathbb{R}^{2m+1}$, defined by

$$\Phi_{(\phi, y)}(x) = (y(x), y(\phi(x)), \dots, y(\phi^{2m}(x)))$$

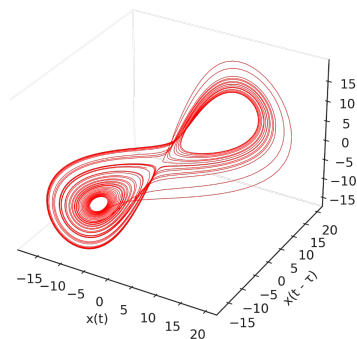
is an embedding; by "smooth" we mean at least C^2 .

Lorenz Attractor



$$(x^1(t), x^2(t), x^3(t))$$

Takens Reconstructed Space



$$(x^1(t), x^1(t - \tau), x^1(t - 2\tau))$$

$\exists$ plongement (bijection)

Phase space $\longrightarrow$ Delayed observables

¹O. Sigaud et al. - Processus décisionnels de Markov en intelligence artificielle, *Lavoisier* (2008)

²J. Subramanian et al. "Approximate Information State for Approximate Planning and Reinforcement Learning in Partially Observed Systems", *JMLR* (2022)

³F. Takens "Detecting Strange Attractors in Turbulence", *Dynamical Systems and Turbulence, Proceeding of a Warwick Symposium* (1981)

Learning Neural Delay Differential Equations (NDDE)

→ **implicit Information State**

Delay Differential Equations (DDE)¹

$$\partial_t y(t) = g(y_t, u(t)) \quad \text{with} \quad y_t : [0, \tau] \rightarrow \mathcal{Y}, \quad y_t(s) = y(t - s)$$

Particular Case

$$\partial_t y(t) = g(y(t), y(t - \tau), u(t))$$

Method → Neural Delay Differential Equations^{2,3}

$$\partial_t y(t) = g_\theta(y(t), y(t - \tau), u(t)) \quad \theta \in \Theta \quad \tau \in \mathbb{R}_+ \quad (\tau \text{ is learnable})$$

Optimisation → **backpropagation** through **DDE solver**

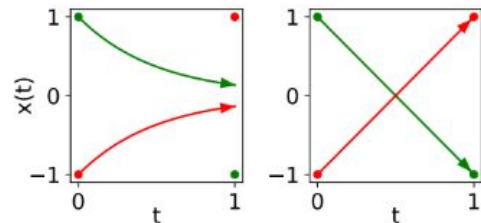
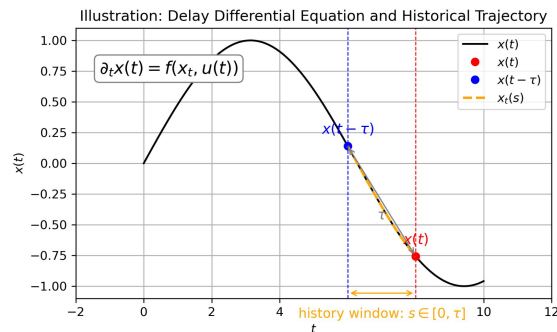


Figure 2: (Right) Two continuous trajectories generated by the DDEs are intersected, mapping -1 (resp., 1) to 1 (resp., -1), while (Left) the ODEs cannot represent such mapping.

¹J. K. Hale - Functional Differential Equations, Springer (1971)

²Q. Zhu et al. "Neural Delay Differential Equation", ICLR (2021)

³T. Monsel et al. "Time and State Dependent Neural Delay Differential Equations", PMLR - ECAI Workshop (2024)

Learning NDDE: Ablation Study

Experiment

Compare 4 neural architectures

$$\text{NODE} \rightarrow \partial_t y(t) = g_{\theta}(y(t))$$

$$\text{NCDE} \rightarrow \partial_t y(t) = g_{\theta}(y(t), u(t))$$

$$\text{NDDE} \rightarrow \partial_t y(t) = g_{\theta}(y(t), y(t - \tau))$$

$$\text{NCDDE} \rightarrow \partial_t y(t) = g_{\theta}(y(t), y(t - \tau), u(t))$$

System: Van der Pol Oscillator

$$\partial_t x^1(t) = x^2(t) + u^1(t)$$

$$\partial_t x^2(t) = \epsilon_{\text{VDP}}(1 - (x^1(t))^2)x^2(t) - x^1(t) + u^2(t)$$

$$y(t) = x(t - \tau)$$

Collect dataset of trajectories $(y_t)_{t \in I}$
with **noisy control** $u(t)$
and $\tau \in \{0, 0.1\}$

Compare 4 neural architectures $\rightarrow$

Generate predictive data (time series) $\rightarrow (y_t^{\theta})_{t \in I}$

$$\text{Minimise over } \theta \in \Theta \rightarrow \mathcal{L} = \mathbb{E}^{\mathbb{P}_{(y_t)_{t \in I}}} \left\| (y_t)_{t \in I} - (y_t^{\theta})_{t \in I} \right\|_{L^2}^2$$

Hypothesis

$u(t)$ -dependant models $\rightarrow \mathcal{L} \searrow$

$y(t - \tau)$ -dependant models fit better $\rightarrow \mathcal{L} \searrow$

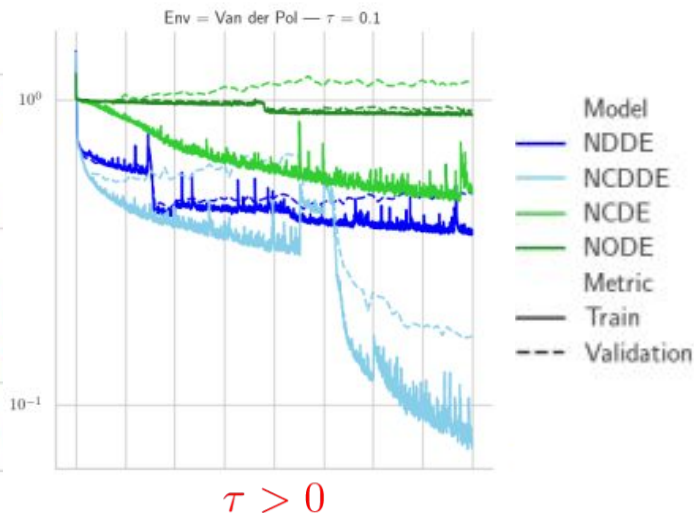
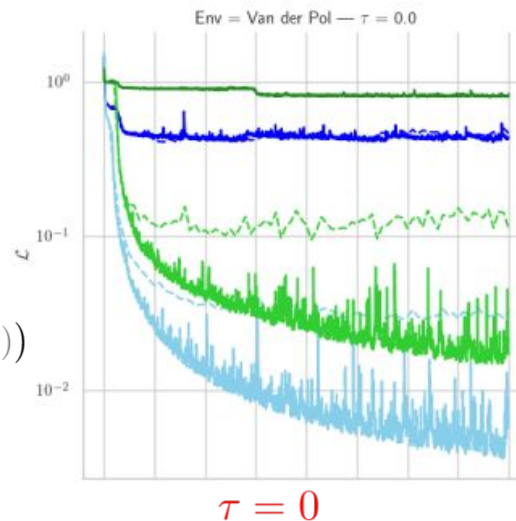
Learning NDDE: Ablation Study - Results

NODE $\rightarrow \partial_t y(t) = g_{\theta}(y(t))$

NCDE $\rightarrow \partial_t y(t) = g_{\theta}(y(t), u(t))$

NDDE $\rightarrow \partial_t y(t) = g_{\theta}(y(t), y(t - \tau))$

NCDDE $\rightarrow \partial_t y(t) = g_{\theta}(y(t), y(t - \tau), u(t))$



Results

$u(t)$ -dependant models $\rightarrow \mathcal{L} \searrow$

$y(t - \tau)$ -dependant models fit better $\rightarrow \mathcal{L} \searrow$

$y(t - \tau)$ -dependant models also improve when $\tau = 0$

Drawback

Increased computational complexity

Contributions and Perspectives

Contributions

- Functional Differential Equation framework in Continuous-time model-based control
- Links with Information States and Dynamical Systems theory
- Controlled Neural Differential Equations

Results

Delay Differential Equation → **Better regression performances**

Perspectives

Continuous-time Dynamic Programming in **Infinite Dimensional Spaces?**

Academic and Scientific Involvement

Academic and Scientific Involvement

Internship supervision

- *Information-driven learning-based MPC (with S. Douka)*
- *Learning-based Functional Dynamic Programming (with E. Pradeleix)*

Teaching

- C++ (Université Paris-Saclay)
- *Data Science Project* (CentraleSupélec)
- *Advanced Deep Learning* (ENS Paris-Saclay MVA)

Reviews

- *Transaction in Automatic and Control*
- *Journal of Fluid Mechanics*
- *European Workshop in Reinforcement Learning*

Contribution in open-source packages

- `control_dde` (with T. Monzel)
- `stable_baselines3`
- `hydrogym`
- `torchdde`

General Conclusion



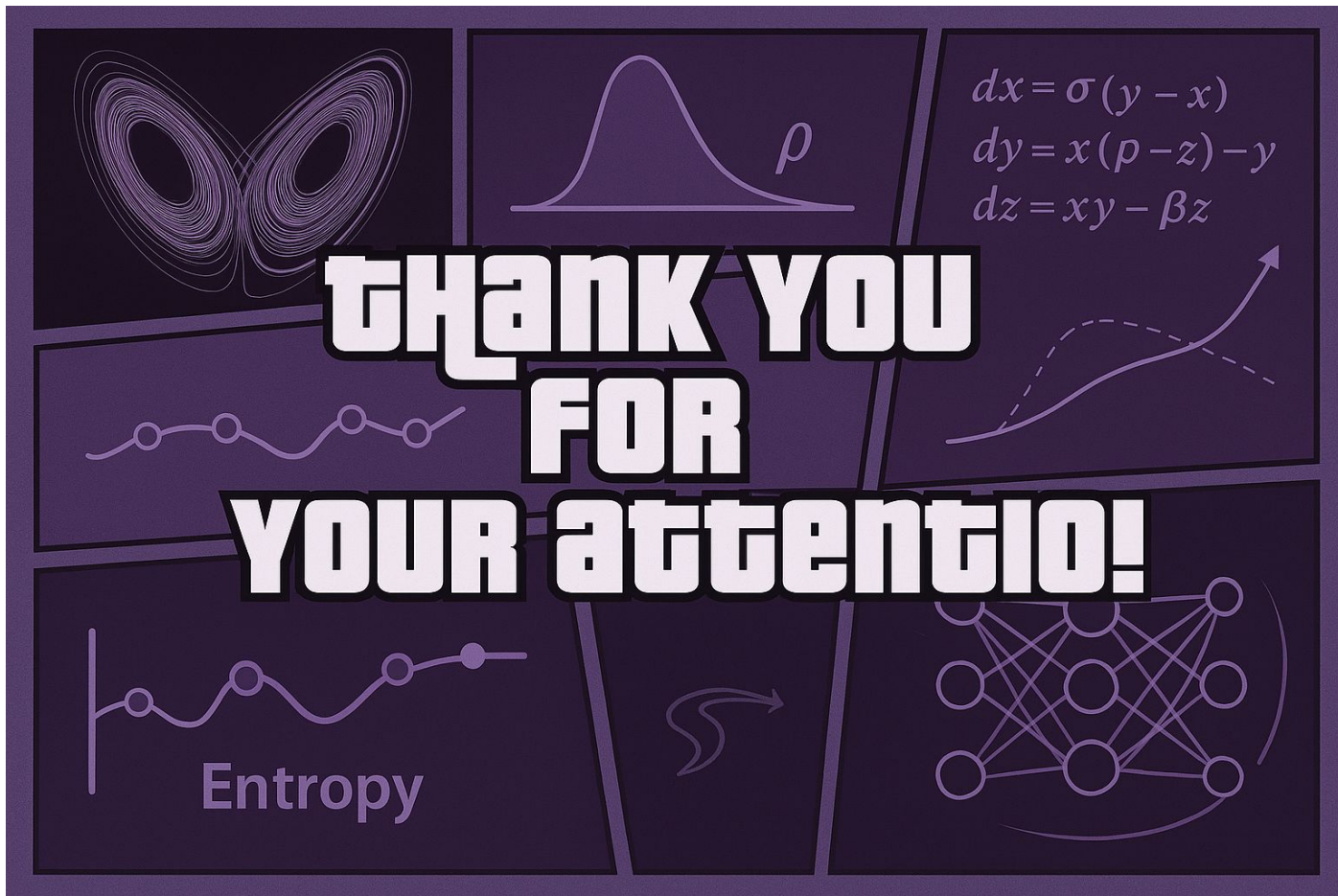
Interdisciplinary Research



Contribution to Learning-based Control Challenges



Brings concepts from various fields and a novel viewpoint



Learning-based Control of Dynamical Systems: Challenges

Complexity Measure: Conditional Fisher Information

Hessian and Fisher Information

$$\text{Objective Hessian} \rightarrow \nabla_{\theta}^2 J^{\pi_{\theta}} = \mathbb{E}^{\pi_{\theta}} \left[\sum_{h,i,j=0}^T c(X_h, U_h) (\nabla_{\theta} \log \pi_{\theta}(U_i | X_j) \nabla_{\theta} \log \pi_{\theta}(U_j | X_j)^T + \nabla_{\theta}^2 [\log \pi_{\theta}(U_i | X_j)]) \right]$$

$$\text{Fisher Information} \rightarrow \mathcal{I}(\theta) = -\mathbb{E}^{X \sim \rho, U \sim \pi_{\theta}(\cdot | X)} [\nabla_{\theta}^2 \log \pi_{\theta}(U | X)]$$

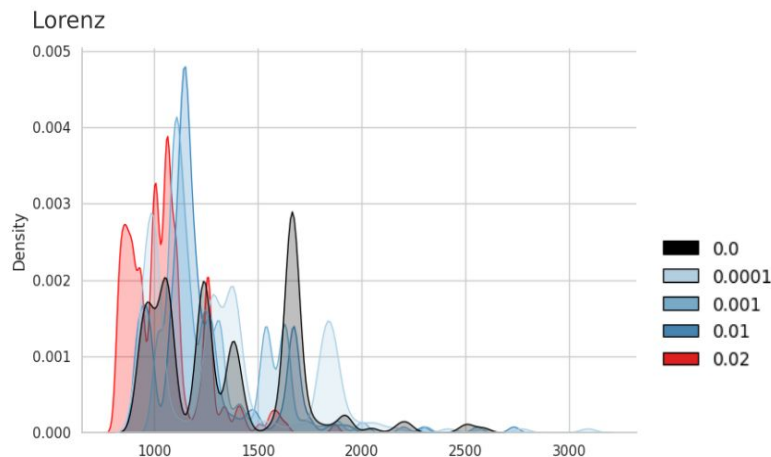
Fisher Information Complexity Measure

$$\mathcal{M}(\pi_{\theta}) = -\text{Tr} \left(\mathbb{E}^{X \sim \rho^{\pi_{\theta}}, U \sim \pi_{\theta}(\cdot | X)} [\nabla_{\theta_{\mu}}^2 \log \pi_{\theta}(U | X)] \right)$$

Results

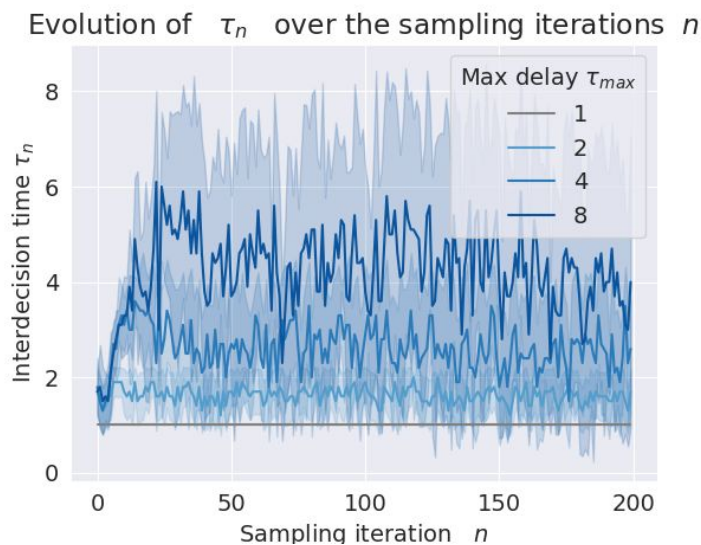
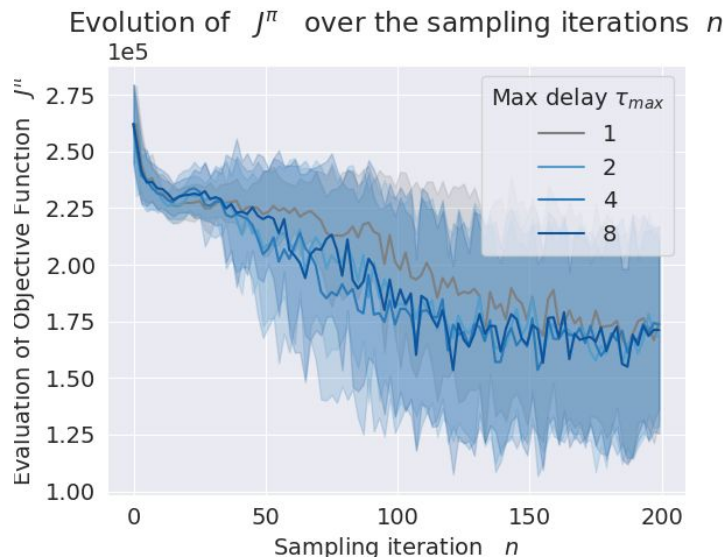
$\alpha = 0 \rightarrow$ Right Leptokurtic distribution

$\alpha > 0 \rightarrow$ Relatively less extreme values



Density of $\mathcal{M}(\pi_{\theta}(\cdot | X)) = -\text{Tr} \left(\mathbb{E}^{U \sim \pi_{\theta}(\cdot | X)} [\nabla_{\theta_{\mu}}^2 \log \pi_{\theta}(U | X)] \right)$ where $X \sim \rho$

Objective Function and Inter-decision time



Results

Temporal Abstraction → Information Gain → Control performances

Placeholder

Placeholder

Placeholder