
Characterizing control between interacting subsystems with deep Jacobian estimation

Adam J. Eisen

Brain and Cognitive Sciences
MIT
Cambridge, MA 02139
eisenaj@mit.edu

Mitchell Ostrow

Brain and Cognitive Sciences
MIT
Cambridge, MA 02139
ostrow@mit.edu

Sarthak Chandra*

Brain and Cognitive Sciences
MIT
Cambridge, MA 02139
sarthak.chandra@icts.res.in

Leo Kozachkov†

IBM Thomas J. Watson Research Center
IBM Research
Yorktown Heights, NY 10598
leokoz8@brown.edu

Earl K. Miller

Brain and Cognitive Sciences
MIT
Cambridge, MA 02139
ekmiller@mit.edu

Ila R. Fiete

Brain and Cognitive Sciences
MIT
Cambridge, MA 02139
fiete@mit.edu

Abstract

Biological function arises through the dynamical interactions of multiple subsystems, including those between brain areas, within gene regulatory networks, and more. A common approach to understanding these systems is to model the dynamics of each subsystem and characterize communication between them. An alternative approach is through the lens of control theory: how the subsystems control one another. This approach involves inferring the directionality, strength, and contextual modulation of control between subsystems. However, methods for understanding subsystem control are typically linear and cannot adequately describe the rich contextual effects enabled by nonlinear complex systems. To bridge this gap, we devise a data-driven nonlinear control-theoretic framework to characterize subsystem interactions via the Jacobian of the dynamics. We address the challenge of learning Jacobians from time-series data by proposing the JacobianODE, a deep learning method that leverages properties of the Jacobian to directly estimate it for arbitrary dynamical systems from data alone. We show that JacobianODE models outperform existing Jacobian estimation methods on challenging systems, including high-dimensional chaos. Applying our approach to a multi-area recurrent neural network (RNN) trained on a working memory selection task, we show that the “sensory” area gains greater control over the “cognitive” area over learning. Furthermore, we leverage the JacobianODE to directly control the trained RNN, enabling precise manipulation of its behavior. Our work lays the foundation for a theoretically grounded and data-driven understanding of interactions among biological subsystems.

*This author is now at The International Centre for Theoretical Sciences.

†This author is now at Brown University.

1 Introduction

Complex systems are ubiquitous in nature. These systems exhibit a wide range of behavior and function, in large part through the dynamic interaction of multiple component subsystems within them. One approach to understanding such complex systems is to build detailed models of their underlying dynamics. An alternative and simpler yet powerful approach is offered by control theory, focusing instead on how subsystems influence and regulate one another, and how they can be controlled.

Control theory thus offers a complementary approach to both understanding and manipulating biological systems. The theory describes how inputs must be coordinated with system dynamics to achieve desired behaviors, and can be applied across domains ranging from robotics to biology (Figure 1A). The brain coordinates neural activity across multiple interconnected brain areas, dynamically modulating which regions receive information from which others depending on need and context [1, 2]. Interareal interactions play central roles in cognition and consciousness [3–7], in selective attention [8–16], decision making [17, 18], working memory [19, 20], feature binding [21, 22], motor control [23–27], and learning and memory [28–32].

A common approach to characterizing interareal interactions is to quantify *communication* between them, using methods such as reduced-rank regression to define “communication subspaces”. These subspaces determine low-dimensional projections that maximally align high-dimensional states of the input area with high-dimensional states of the target area [33–35]. However, effective *control* not only involves alignment of high-variance input states with high-variance target states, but also appropriate alignment of the inputs with the *dynamics* of the target area. Given connected subsystems A and B, an identical signal from B will have dramatically different control effects on A, depending on whether the signal aligns with stable or unstable directions of A’s dynamics: projections onto more unstable eigenvectors can much more readily drive the system to novel states. (For more detail see Appendix C.1.)

Accordingly, recent work in neuroscience has espoused control-theoretic perspectives on interareal interactions (Figure 1B) [24, 36–42]. The dominant approach has involved linear control [43–53]. However, the inherently nonlinear dynamics of the brain enable richer contextual control than possible to fully model with linear systems, necessitating a nonlinear control approach.

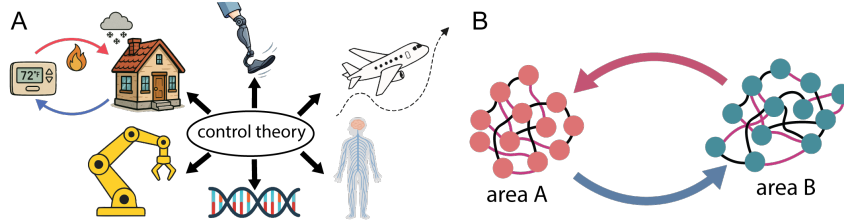


Figure 1: **Schematic overview of control-theoretic framework applied to neural interactions.** (A) Control theory generalizes across diverse systems. (B) Illustration of interareal control, highlighting how neural activity in one area directly influences dynamics in another.

One approach to extend control-theoretic analyses to nonlinear systems is by linearizing the nonlinear dynamics through Taylor expansion, which involves the Jacobian. This converts the nonlinear system into a linear state- and time-dependent one [40, 54–56], allowing for simple control. Jacobian linearization for control is straightforward with access to analytical expressions for the non-linear system, but it becomes non-trivial in purely data-driven scenarios. Estimating the Jacobian involves conjunctively inferring both a function and its derivative, yet good function approximation need not yield good approximations of its derivatives (see Section 4 and Appendix C.2).

This paper introduces several key contributions:

- **Robust data-driven Jacobian estimation.** We present JacobianODE, a deep learning-based method for estimating Jacobians from noisy trajectories in high-dimensional dynamical systems. We demonstrate the validity of this approach in several sample systems that include high-dimensional chaos.
- **A framework to characterize control between interacting subsystems via data-driven Jacobian estimation.** Harnessing our Jacobian estimation method, we devise a rigorous,

data-driven approach for nonlinear control-theoretic analysis of how paired interacting systems, including brain areas, drive and regulate each other across different contexts.

- **Data-driven inference of control dynamics in trained recurrent neural networks.** We apply our data-driven framework to a recurrent neural network (RNN) trained on a working memory task. We show that, purely from data, we can identify key control-theoretic interactions between the areas, and that these interactions crucially evolve over the course of learning to produce the desired behavior.
- **Demonstration of accurate control of rich interacting high-dimensional coupled dynamical subsystems.** We demonstrate high accuracy in a challenging high-dimensional data-driven nonlinear control task, enabling precise control of the behavior of the RNN.

Overall, our work lays the foundation for data-driven control-theoretic analyses in complex high-dimensional nonlinear systems, including the brain.

2 Related work

Interareal communication A wide range of tools have been developed and harnessed to study interareal communication in neural data [34, 57]. This includes, but is not limited to, methods based on reduced rank regression [33, 35], recurrent neural network models of neural dynamics [40, 58–63], Gaussian process factor analysis [64, 65], canonical correlation analysis [66–68], convergent cross mapping [69–71], switching dynamical systems [72, 73], granger causality [74, 75], dynamic causal mapping [76], point process models [77], and machine learning methods [78, 79].

Nonlinear controllability Classical results assess the controllability of nonlinear control systems via the Lie theory [80–84]. Another approach to nonlinear network controllability is based on attractor strength [85]. Liu et al. [54] and Parkes et al. [55] note that the large literature of linear controllability analyses could be extended locally to nonlinear systems with an appropriate linearization method.

Neural network controllability Linear structural network controllability has been implicated across a wide range of contexts, tasks, and neuropsychiatric conditions [36, 42, 44–47, 49, 50]. This approach has been extended to functional brain networks [51–53]. Recent work has characterized the subspaces of neural activity that are most feedback controllable as opposed to feedforward controllable using linear methods [43]. Other approaches to neural control analyses consider the identification of structural driver nodes [86], and input novelty [87].

Data-driven Jacobian estimation Besides approaches estimating Jacobians through weighted linear regressions [88, 89], some methods have used direct parameterization via neural networks to learn Jacobians of general functions [90, 91]. These approaches inform our method but do not explicitly address dynamical systems. Applying path-integral-based Jacobian estimation to dynamical systems is challenging, as the target function (the system’s time derivative) is typically unobserved. Beik-Mohammadi et al. [92] utilized this idea in dynamical systems to learn contracting latent dynamics from demonstrations.

3 JacobianODE: learning Jacobians from data

3.1 Jacobian linearization

We consider nonlinear dynamical systems in $\mathbb{R}^n$, defined by $\dot{\mathbf{x}}(t) = \mathbf{f}(\mathbf{x}(t))$. The Jacobian of the dynamics is a matrix-valued function $\mathbf{J}_{\mathbf{f}} : \mathbb{R}^n \rightarrow \mathbb{R}^{n \times n}$ (henceforth, $\mathbf{J}$) given by

$$\mathbf{J}(\mathbf{x}(t)) = \frac{\partial}{\partial \mathbf{x}} \mathbf{f}(\mathbf{x}(t)) = \begin{bmatrix} \frac{\partial f_1}{\partial x_1} & \frac{\partial f_1}{\partial x_2} & \cdots & \frac{\partial f_1}{\partial x_n} \\ \frac{\partial f_2}{\partial x_1} & \frac{\partial f_2}{\partial x_2} & \cdots & \frac{\partial f_2}{\partial x_n} \\ \vdots & \vdots & \ddots & \vdots \\ \frac{\partial f_n}{\partial x_1} & \frac{\partial f_n}{\partial x_2} & \cdots & \frac{\partial f_n}{\partial x_n} \end{bmatrix}. \quad (1)$$

At each time t , the Jacobian defines a linear subspace relating input and output changes, capturing how perturbations to the system will propagate. This recasts nonlinear dynamics as linear time-varying dynamics in the tangent space locally along trajectories (formally, $\delta \dot{\mathbf{x}}(t) = \mathbf{J}_{\mathbf{f}}(\mathbf{x}(t)) \delta \mathbf{x}(t)$, see Figure 2 left, also see Lohmiller and Slotine [93] for a discussion in the context of virtual displacements).

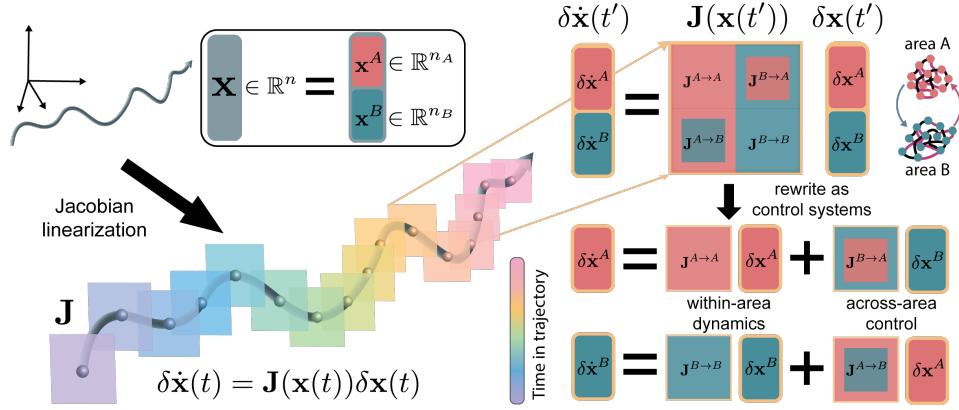


Figure 2: **Analytical framework for pairwise interacting subsystem control.** Trajectory dynamics (left, top) are locally linearized via Jacobians (left, bottom), explicitly separating within-area (diagonal blocks) and interareal (off-diagonal blocks) dynamics (right). These separated dynamics can be used to construct interaction-specific control systems.

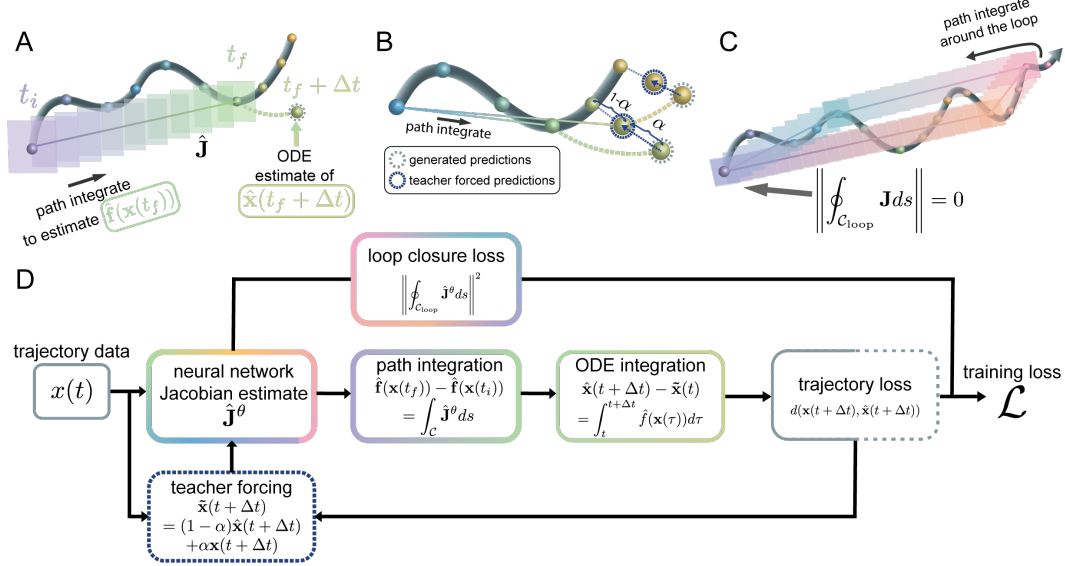


Figure 3: **Jacobian estimation with JacobianODE models.** (A) Path integration of the Jacobian predicts future states. (B) Generalized teacher forcing stabilizes trajectory predictions during training. (C) Loop-closure constraints enforce consistency of Jacobian estimates. (D) Training pipeline, combining neural Jacobian estimation, path integration, teacher forcing, and self-supervised loop-closure loss.

3.2 Parameterizing differential equations via the Jacobian

We now turn to the problem of how to estimate the Jacobian $\mathbf{J}$ from data. We assume that we have access only to observed trajectories of the system, of the form $\mathbf{x}^{(j)}(t_0^{(j)} + k\Delta t)$, $k = 1, 2, \dots$, $j = 1, 2, \dots$ where Δt is a fixed sampling interval, j indexes the trajectory, and $t_0^{(j)}$ is a trajectory-specific start time. Crucially, we do not assume access to the function $\mathbf{f}$. Our method estimates the Jacobian directly via a neural network. To do this, we parameterize a neural network function $\hat{\mathbf{J}}^\theta$ with learnable parameters θ that is then trained to approximate $\mathbf{J}$.

Path integration Following Lorraine and Hossain [91] and Beik-Mohammadi et al. [92], we exploit the fact that the path integral of the Jacobian is path independent in the construction of the loss function. This is because the rows of the Jacobian are conservative vector fields (i.e., they are the

gradients of scalar functions). For an intuitive picture, consider the work done by the force of gravity as you climb a mountain. Regardless of the path you take to climb, the resulting work is dependent only on the start and end points of the path. Formally, for the time derivative function $\mathbf{f}$ we have that

$$\mathbf{f}(\mathbf{x}(t_f)) - \mathbf{f}(\mathbf{x}(t_i)) = \int_{\mathcal{C}} \mathbf{J} ds = \int_{t_i}^{t_f} \mathbf{J}(\mathbf{c}(r)) \mathbf{c}'(r) dr, \quad (2)$$

where $\mathcal{C}$ is a piecewise smooth curve in $\mathbb{R}^n$ and $\mathbf{c} : [t_i, t_f] \rightarrow \mathcal{C}$ is a parameterization of $\mathcal{C}$ with $\mathbf{c}(t_i) = \mathbf{x}(t_i)$ and $\mathbf{c}(t_f) = \mathbf{x}(t_f)$ (Figure 3A). Given estimates of $\mathbf{f}(\mathbf{x}(t_i))$ and the Jacobian $\mathbf{J}$, we can then use Equation 2 to approximate $\mathbf{f}(\mathbf{x}(t_f))$ at any time t_f . For these integrals, we use a line between the endpoints as a simple choice of path (Figure 3A). Then, to generate an estimate $\hat{\mathbf{x}}(t_f + \Delta t)$ of the next step, we can use a standard ordinary differential equation (ODE) integrator (e.g., Euler, fourth-order Runge–Kutta, etc.) to integrate the estimated time derivative $\hat{\mathbf{f}}(\mathbf{x}(t))$ (see Appendix A for more detail). To avoid the need to represent $\mathbf{f}$ directly (thereby enabling all gradients to backpropagate through the Jacobian network) we note that we parameterize an estimate of $\mathbf{f}(\mathbf{x}(t_i))$ in Equation 2 in terms of the Jacobian (see Appendix A.1).

3.3 Loss functions

Trajectory reconstruction loss Given an observed trajectory $\mathbf{x}(t_0 + k\Delta t), k = 0, \dots, T - 1$ of length T , we can compute the trajectory reconstruction loss, $\mathcal{L}_{\text{traj}}(\theta; \mathbf{x})$, between the true trajectory and the estimated trajectory using an appropriate distance measure d (e.g., mean squared error, see Appendix A.2 for more detail on generating predictions and the trajectory prediction loss).

Generalized Teacher Forcing To avoid trajectory divergence in chaotic or noisy systems, we employ Generalized Teacher Forcing, generating recursive predictions partially guided by true states (Figure 3B, and Appendix D.8.1) [94].

Loop closure loss The Jacobian captures how perturbations to the system will propagate along *any* direction in state space. Estimating it purely from dynamics constrains only the direction of the flow, leaving the full solution underdetermined. To address this, we again exploit the fact that each row of the Jacobian is a conservative vector field. Specifically, we note that for any piecewise smooth loop $\mathcal{C}_{\text{loop}}$, we have $\oint_{\mathcal{C}_{\text{loop}}} \mathbf{J} ds = 0$ (see Figure 3C). Thus, by integrating along loops that contain directions *orthogonal* to the system’s dynamics (and penalizing the deviation from zero), we encourage the estimated Jacobians to capture information about other directions in state space (see Appendices A.3 and D.8.2 for full technical details). To ensure broad coverage of tangent space directions, we form loops from concatenations of line integrals between randomly selected data points. This strategy samples diverse directions from the tangent space while remaining easy to compute. The resulting self-supervised loss term, $\mathcal{L}_{\text{loop}}(\theta; \mathbf{x})$, builds on the loss introduced by Iyer et al. [95]. It constrains $\hat{\mathbf{J}}^\theta$ to satisfy both the dynamics and conservativity. This improves Jacobian estimation accuracy significantly (see Appendix C.4 for ablation studies).

Training loss We therefore minimize the following loss function with respect to the parameters θ :

$$\mathcal{L}(\theta; \mathbf{x}) = \mathcal{L}_{\text{traj}}(\theta; \mathbf{x}) + \lambda_{\text{loop}} \mathcal{L}_{\text{loop}}(\theta; \mathbf{x}) \quad (3)$$

where λ_{loop} controls the relative weighting of the loop closure loss $\mathcal{L}_{\text{loop}}(\theta; \mathbf{x})$ compared to the trajectory prediction, and is a hyperparameter of the learning procedure (Figure 3D).

4 Jacobian estimation in dynamical systems

Data To evaluate the quality of the Jacobian estimation procedure, we apply our approach to several example systems for which the dynamics are known. For this analysis, we used the Van der Pol oscillator [96], Lorenz system [97], and the Lorenz 96 system across three different system sizes (12, 32, and 64 dimensional) [98]. All systems were simulated using the `dysts` package, which samples dynamical systems with respect to the characteristic timescale τ of their Fourier spectrum [99, 100]. All training data consisted of 26 trajectories of 12 periods, sampled at 100 time steps per τ . All models were trained on a 10 time-step prediction task with teacher forcing.

To evaluate the performance of the methods in the presence of noise, we trained the models on data with 1%, 5%, and 10% Gaussian observation noise added i.i.d over time, where the percentage is defined via the ratio of the euclidean norm of the noise to the mean euclidean norm of the data.

JacobianODE model For the JacobianODE framework, loop closure loss weights were chosen via line search (Appendix D.8.8), where the neural network $\mathbf{J}^\theta$ was taken to be a four-layer multilayer perceptron (MLP) with hidden layer sizes 256, 1024, 2048 and 2048. Path integration was performed using the trapezoid method from the `torchquad` package, with each integral discretized into 20 steps [101]. ODE integration was performed using the fourth-order Runge–Kutta (RK4) method from the `torchdiffeq` package [102]. The JacobianODE models used 15 observed points to generate the initial estimate of $\mathbf{f}$ (see Section 3.2 and Appendices A.1 and A.2). All models were built in PyTorch [103]. Full implementation details are provided in appendices A and D.

Baselines We chose two different Jacobian estimation procedures for comparison. The first was a neural ordinary differential equation (NeuralODE) model trained to reproduce the dynamics [102]. The NeuralODE was implemented as a four-layer MLP with hidden layers of the same size as the one used for the JacobianODE model. Jacobians were computed via automatic differentiation. NeuralODEs were regularized via a penalty on the Frobenius norm of the estimated Jacobians to prevent the model from learning unnecessarily large negative eigenvalues (see Appendix D.3) [104–106]. We also employed a baseline that estimates Jacobian via a weighted linear regression, which computes locally linear models at each point in the space [88, 89] (see Appendix D.4).

Table 1: Mean Frobenius norm error on Jacobian estimation, $\langle \|\mathbf{J} - \hat{\mathbf{J}}\|_F \rangle$, for each system and noise level. Errors are reported as mean $\pm$ standard deviation, with mean and standard deviation computed over five random initializations of the model architectures.

Project	Training noise	JacobianODE	NeuralODE	Weighted Linear
VanDerPol (2 dim)	1%	0.7 $\pm$ 0.1	1.0 $\pm$ 0.3	6.11
	5%	0.72 $\pm$ 0.05	0.72 $\pm$ 0.08	6.09
	10%	1.35 $\pm$ 0.06	2.2 $\pm$ 0.2	6.08
Lorenz (3 dim)	1%	3.3 $\pm$ 0.2	8.7 $\pm$ 0.3	21.94
	5%	5.1 $\pm$ 0.9	26.0 $\pm$ 1.5	21.90
	10%	6.4 $\pm$ 0.1	26.7 $\pm$ 0.9	21.84
Lorenz 96 (12 dim)	1%	1.2 $\pm$ 0.2	4.8 $\pm$ 0.2	28.67
	5%	2.7 $\pm$ 0.2	5.9 $\pm$ 0.2	28.64
	10%	4.6 $\pm$ 0.1	6.1 $\pm$ 0.1	28.56
Lorenz 96 (32 dim)	1%	8.7 $\pm$ 0.2	16.8 $\pm$ 0.4	47.13
	5%	7.8 $\pm$ 0.2	17.7 $\pm$ 0.6	46.96
	10%	13.45 $\pm$ 0.09	19.5 $\pm$ 0.4	47.03
Lorenz 96 (64 dim)	1%	30.9 $\pm$ 0.5	45.5 $\pm$ 0.1	66.39
	5%	34.0 $\pm$ 0.2	45.7 $\pm$ 0.1	66.26
	10%	36.0 $\pm$ 0.2	46.0 $\pm$ 0.2	66.29
Task-trained RNN	1%	188.5 $\pm$ 7.1	294.02 $\pm$ 0.03	301.46
	5%	166.8 $\pm$ 3.6	294.357 $\pm$ 0.003	297.63
	10%	180.4 $\pm$ 0.5	294.328 $\pm$ 0.004	296.63

Performance We tested the approaches on Jacobian estimation on held-out trajectories without noise. The JacobianODE method outperforms the baseline methods in terms of mean Frobenius norm error for virtually all systems (Table 1). This was also true when considering the spectral matrix 2-norm (see Table S3 in Appendix C.3).

We plot performance in Figure 4. While both the JacobianODEs and the NeuralODEs reproduce the observed dynamics (Figure 4A-D), the JacobianODE learns a more accurate estimate of the Jacobian (Figure 4 E,F). In particular, looking at Lyapunov spectra learned by the models, we note that the JacobianODE exceeds the other methods in estimating the full Lyapunov exponent spectrum, indicating a better overall representation of how perturbations along different directions will evolve (Figure 4G,H).

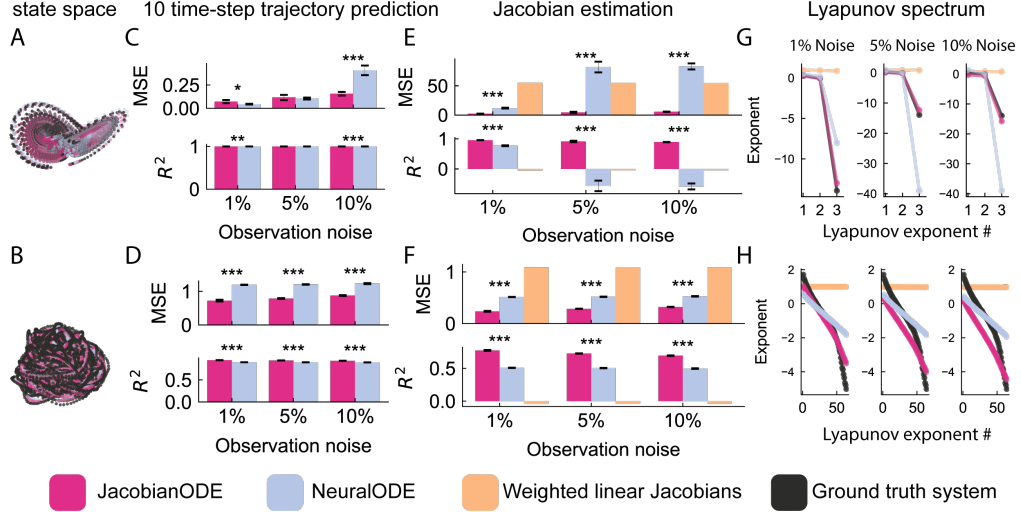


Figure 4: **JacobianODE surpasses benchmark models on chaotic dynamical systems.** Error bars indicate standard deviation, with statistics computed over five different model initializations. (A,B) State-space trajectories for (A) Lorenz and (B) 64-dimensional Lorenz 96 systems, with 10 time-step predictions. Spectral embedding was used to reduce the Lorenz 96 data to three dimensions. (C,D) Accuracy of 10 time-step trajectory predictions at varying noise levels. (E,F) Comparison of Jacobian estimation accuracy, quantified by mean squared error (MSE) and R^2 . (G,H) Estimated Lyapunov spectra averaged over initializations.

5 Control-theoretic analyses in a task-trained RNN

5.1 Characterizing control between subsystems with Jacobian linearization

Consider neural data recorded from two areas, A and B. Concatenating their data into a state vector $\mathbf{x} \in \mathbb{R}^n$, composed of $\mathbf{x}^A \in \mathbb{R}^{n_A}$ and $\mathbf{x}^B \in \mathbb{R}^{n_B}$, and assuming dynamics governed by $\mathbf{f}$, we linearize around a reference trajectory ($\delta \dot{\mathbf{x}}(t) = \mathbf{J}(\mathbf{x}(t)) \delta \mathbf{x}(t)$). Splitting the Jacobian into block matrices yields:

$$\begin{aligned} \delta \dot{\mathbf{x}}^A(t) &= \mathbf{J}^{A \rightarrow A}(\mathbf{x}(t)) \delta \mathbf{x}^A + \mathbf{J}^{B \rightarrow A}(\mathbf{x}(t)) \delta \mathbf{x}^B \\ \delta \dot{\mathbf{x}}^B(t) &= \mathbf{J}^{A \rightarrow B}(\mathbf{x}(t)) \delta \mathbf{x}^A + \mathbf{J}^{B \rightarrow B}(\mathbf{x}(t)) \delta \mathbf{x}^B \end{aligned} \quad (4)$$

where diagonal blocks $\mathbf{J}^{A \rightarrow A} \in \mathbb{R}^{n_A \times n_A}$ and $\mathbf{J}^{B \rightarrow B} \in \mathbb{R}^{n_B \times n_B}$ represent within-area dynamics, and off-diagonal blocks $\mathbf{J}^{B \rightarrow A} \in \mathbb{R}^{n_A \times n_B}$, $\mathbf{J}^{A \rightarrow B} \in \mathbb{R}^{n_B \times n_A}$ represent interareal interactions (Figure 2, right). Explicit separation of each area’s control dynamics quantifies the direct influence each exerts on the other, and readily generalizes beyond two areas (see Appendix B.3).

Since Jacobians are time-dependent, we obtain a linear time-varying representation of control dynamics along the trajectory. This enables computation of time-varying reachability ease, capturing how readily each area drives the other toward novel states [107–109]. Reachability is quantified via the reachability Gramian, a matrix defining a local metric in tangent space. For the above control system capturing the influence of area B on area A, the time-varying reachability Gramian on the interval $[t_0, t_1]$ is defined as

$$\mathbf{W}_r(t_0, t_1) \triangleq \int_{t_0}^{t_1} \Phi(t_1, \tau) \mathbf{B}(\tau) \mathbf{B}^T(\tau) \Phi^T(t_1, \tau) d\tau \quad (5)$$

where Φ (computed from $\mathbf{J}^{A \rightarrow A}$) denotes the state-transition matrix of the intrinsic dynamics of subsystem A without any input (i.e., $\delta \mathbf{x}^A(t) = \Phi(t, t_0) \delta \mathbf{x}^A(t_0)$), and $\mathbf{B}(\tau) = \mathbf{J}^{B \rightarrow A}(\mathbf{x}(\tau))$ (see Appendix B.1) [107]. The Gramian $\mathbf{W}_r(t_0, t_1)$ is symmetric and positive semidefinite for every $t_1 > t_0$ [107]. Each eigenvalue of the reachability Gramian quantifies how easily the target system can be driven along its corresponding eigenvector. Thus, the trace of the reachability Gramian reflects average ease of reaching new states, and its minimum eigenvalue reflects the ease along the most challenging direction of control.

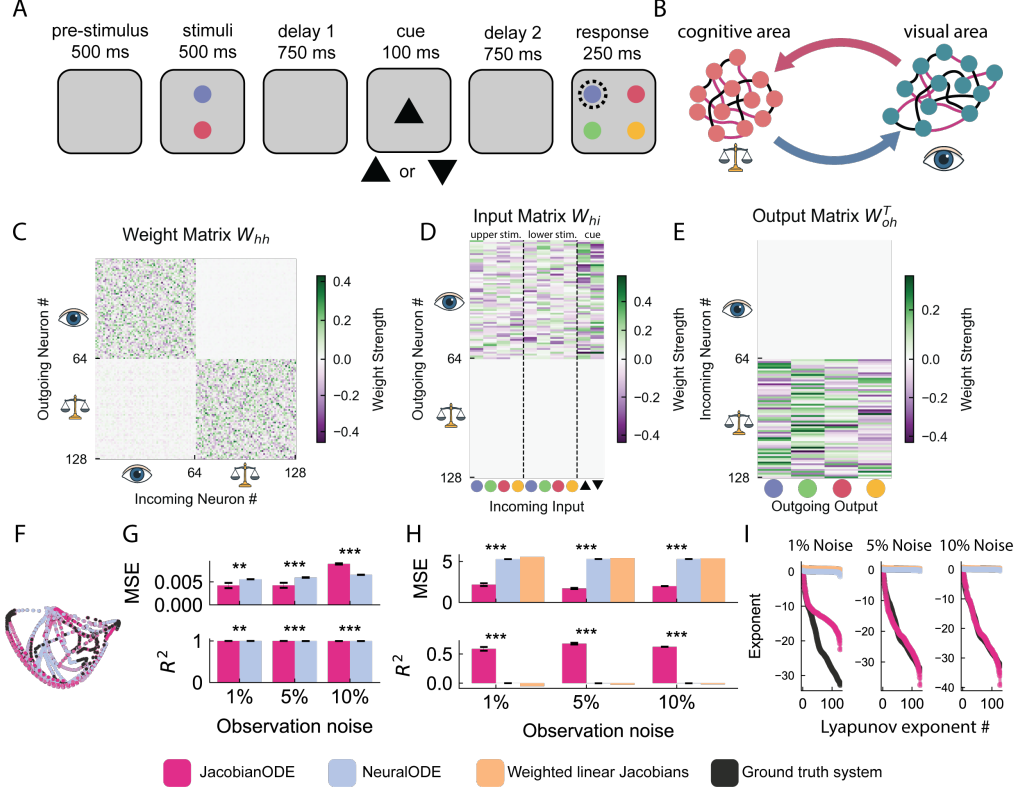


Figure 5: **JacobianODE accurately infers trained RNN Jacobians.** Error bars indicate standard deviation, with statistics computed over five different model initializations. (A) Task schematic, involving stimuli presentation, delays, cueing, and response. (B) RNN architecture with distinct visual and cognitive areas interacting. (C–E) Structure of trained weight matrices: (C) recurrent weights within RNN, (D) input connectivity pattern, and (E) output connectivity. (F) State-space trajectories from the trained RNN, with 10 time-step predictions. Spectral embedding was used to reduce the data to three dimensions. (G,H) JacobianODE performance evaluated against other models for trajectory prediction (G) and Jacobian estimation accuracy (H) at multiple noise levels. (I) Estimated Lyapunov spectra averaged over initializations.

5.2 Estimating the Jacobian of a task-trained RNN

Task To demonstrate how JacobianODE models could be used in neuroscience, we performed a control-theoretic analysis of a task-trained RNN. We used a working memory selection task from Panichello and Buschman [110] (Figure 5A). On each trial, the network is presented with two of four possible "colors", denoted by one-hot vectors. After a delay, the network is presented with a cue indicating which of the colors to select. The network is then asked to reach a state of sustained activation that corresponds to the selected color.

RNN model To perform this task, we trained a 128-dimensional continuous-time RNN. The RNN had hidden dynamics defined by

$$\begin{aligned}\tau \dot{\mathbf{h}}(t) &= -\mathbf{h} + \mathbf{W}_{hh}\sigma(\mathbf{h}(t)) + \mathbf{W}_{hi}\mathbf{u}(t) + \mathbf{b} \\ \mathbf{o}(t) &= \mathbf{W}_{oh}\mathbf{h}(t)\end{aligned}\tag{6}$$

where $\mathbf{W}_{hh}$ defines the internal dynamics, $\mathbf{W}_{hi}$ maps the input into the hidden state, $\mathbf{W}_{oh}$ maps the hidden state to a four-dimensional output $\mathbf{o}(t)$, $\mathbf{b}$ is a bias term, and σ is the exponential linear unit activation. The RNN had two 64-neuron areas: a "visual" area (which received sensory input) and a "cognitive" (which output the RNN's color choice) (Figure 5B). To encourage multi-area structure, we initialized the within-area weights with greater connectivity strength than the across-area weights (Figure 5C). Since input comes only to the visual area and output only from the cognitive area, the two areas are forced to interact to solve the task (Figure 5D,E). The RNN solves this task with 100% accuracy.

Jacobian reconstruction quality We trained JacobianODEs on RNN trajectories from the post-cue delay and response portion of the trials. We used the same baselines as in Section 4. JacobianODEs are not only robust to noise, but can also benefit from it (for multiple systems in Table 1, 5% training noise improves estimation). Noise encourages the model to explore how perturbations around the observed trajectory evolve, which is crucial for learning accurate Jacobians in high-dimensional systems. Thus (for both the JacobianODE and NeuralODE) we add a small amount of additional noise to the data during learning. Although the models perform similarly on trajectory reconstruction (Figure 5F,G), on Jacobian estimation, JacobianODEs drastically outperform both baseline models (Figure 5H,I, and Table 1).

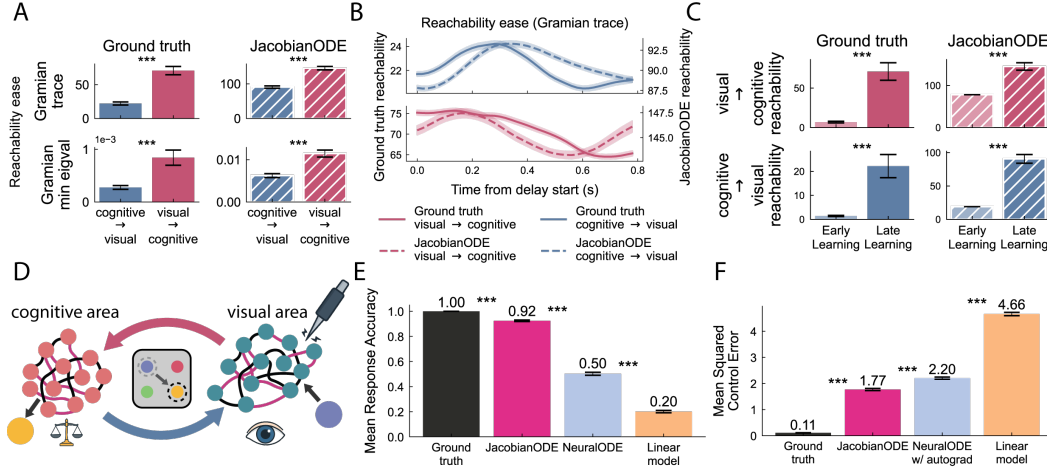


Figure 6: **JacobianODE reveals differential interareal reachability.** Error bars indicate standard deviation for (A) and (C), and standard error for (B), (E), and (F), with statistics computed over trajectories. (A) Comparison of interareal reachability between ground truth and JacobianODE estimates. (B) Temporal evolution of reachability (Gramian trace) throughout the delay period. (C) Comparison of reachability in early and late learning. (D) Schematic illustrating targeted control of cognitive-area via visual-area stimulation. (E) Mean accuracy of targeted responses using JacobianODE versus alternative methods. (F) MSE for multiple ILQR-based methods.

5.3 Reachability in the task-trained RNN across contexts

Next, we used the Jacobians learned by the JacobianODE (trained on 5% noise) to evaluate reachability control in the RNN, and compared it to evaluations using ground truth Jacobians. All analyses were performed using the 10 time-step reachability Gramian. We first found that it was on average easier for the visual area to drive the cognitive area, both when considering overall ease (Gramian trace) and worst-case ease (Gramian minimum eigenvalue) (Figure 6A, larger values indicate greater reachability). We next considered how reachability ease varied throughout the delay period. We found that the visual area could drive the cognitive area more easily at the beginning of the delay, with ease decreasing into the middle of the delay period (Figure 6B, bottom). The cognitive area was able to drive the visual area more easily slightly later in the delay period (Figure 6B, top). Finally, we considered whether reachability changes over the course of learning. We found that both directions of reachability increased after learning (Figure 6C). The JacobianODE reproduced all results accurately (Figure 6A-C, comparison with ground truth). Our analysis reveals that reachability is crucial in the RNN’s ability to perform the working memory task, with the visual area’s ability to drive the cognitive area shortly after the cue being especially important.

5.4 Controlling the task-trained RNN

To further validate our approach, we used JacobianODE-learned Jacobians to control the task-trained RNN (Figure 6D). Given a trial in which the network was cued to respond with a particular color, we tested if we could induce a specific incorrect response by input to the visual area alone. To do so, we implemented Iterative Linear Quadratic Regulator (ILQR) control, which relies on knowledge of the Jacobian [111, 112]. The controller guided the network towards the mean hidden state of

training trials corresponding to the desired incorrect color. We defined accuracy as the percentage of time points during which the RNN output the desired color. We found that the JacobianODE widely outperformed the baseline models on this task, achieving an accuracy nearing that of the ground truth system (Figure 6E). The JacobianODE was furthermore able to achieve the lowest mean squared error on the desired control trajectory (Figure 6F). This illustrates that while both the JacobianODE and the NeuralODE can learn the dynamics, only the JacobianODE learns a representation of the Jacobian that is sufficient for control.

6 Discussion

Extended Jacobian estimation A natural extension of JacobianODE models would add inductive biases for particular classes of dynamics. For example, one could parameterize a negative definite matrix and thus ensure contracting dynamics [92]. Other extensions could include an L1 penalty to encourage sparsity, as well as the inclusion of known dynamic structure (e.g., hierarchical structure, low-rank structure, etc.). In neuroscience, connectomic constraints could be incorporated to capture interactions within a neural circuit.

Limitations A future challenge for JacobianODE models is partially observed dynamics. Recent work has identified that it is possible to learn latent embeddings that approximately recover the true state from partial observation [113–117]. Jacobian-based dynamics learning has been performed in a latent space [92, 118], however it is unclear whether this translates to accurate Jacobian estimation, given the challenges related to automatic differentiation and Jacobian estimation presented here. We also note that reachability estimates depend sensitively on several factors: the alignment of cross-subsystem interactions and within-subsystem dynamics, the eigenvectors and eigenvalues of within-subsystem Jacobians, and the way activity propagates within each subsystem (see Appendix C.1). While our method reliably captures broad trends in reachability over time, fine-grained, time-point-specific comparisons should be interpreted with caution. Finally, although JacobianODE models scale well to moderately high-dimensional systems, their performance in systems that are orders of magnitude larger than those considered here (e.g., recordings of thousands of neurons via calcium imaging) remains to be tested. In many practical settings, this may not be a limitation: neural representations during tasks often exhibit intrinsic dimensionalities comparable to those studied here, enabling JacobianODE models to operate in reduced dimensionality (see Appendix C.7). Notably, most models converged in under 45 minutes on a single GPU, suggesting favorable scaling (see Table S7). Future work should explore the viability of the method in higher dimensions, and assess whether dimensionality reduction strategies (such as latent state models or low-rank Jacobian approximations) can further improve scalability and accuracy.

Broader impact on dynamical systems We demonstrated in this work that it is possible to learn high accuracy Jacobians of arbitrary dynamical systems from data. While in this work we focus on control-theoretic analyses, we emphasize that JacobianODE models can be trained in any dynamical setting. Jacobians are widely applicable, and are essential in stability analysis [80, 119–123], contraction theory [40, 80, 93, 124–126], and nonlinear control [56, 111, 112, 127–129], among many other settings.

Broader impact on control-theoretic analyses The control-theoretic analysis presented here is performed with respect to a particular reference trajectory. It therefore describes the local controllability properties along the reference trajectory, rather than global properties for the overall system [56]. This locality is desirable in biological settings due to the context-sensitive nature of the relevant processes. Applications of this approach could involve a comparison of the functioning of the autonomic nervous system, or gene regulatory networks, between control and disease states, given the association of these processes with disease [130–138]. Future work should investigate whether a more global Jacobian-based control analysis (e.g., contraction theory) is achievable in a data-driven manner via JacobianODE models. This would be useful especially in engineering settings where global controllability is desirable, such as robotics and prosthetic limbs.

Acknowledgments and Disclosure of Funding

The authors would like to thank Andrew Kirjner, Laureline Logiaco, Federico Claudi, Lakshmi Narasimhan Govindarajan, Jaedong Hwang, and Akhilan Boopathy for providing stimulating discussion about this work. The authors also acknowledge the MIT Office of Research Computing and Data

for providing high-performance computing resources that have contributed to the research results reported within this paper. I.R.F. acknowledges funding from The National Science Foundation Computer and Information Science and Engineering Directorate, The Simons Collaboration on the Global Brain, and The McGovern Institute at MIT. E.K.M. acknowledges funding from ONR MURI N00014-23-1-2768, NIMH 1R01MH131715-01, NIMH R01MH11559, The Simons Center for the Social Brain, The Freedom Together Foundation, and The Picower Institute for Learning and Memory. L.K. acknowledges funding from the Goldstone Postdoctoral Fellowship at IBM Research. M.O. acknowledges funding from NSF GRFP 2141064.

References

- [1] Laura Lee Colgin, Tobias Denninger, Marianne Fyhn, Torkel Hafting, Tora Bonnevie, Ole Jensen, May-Britt Moser, and Edvard I. Moser. Frequency of gamma oscillations routes flow of information in the hippocampus. *Nature*, 462(7271):353–357, November 2009. ISSN 1476-4687. doi: 10.1038/nature08573. URL <https://doi.org/10.1038/nature08573>.
- [2] Shencong Ni, Brendan Harris, and Pulin Gong. Distributed and dynamical communication: a mechanism for flexible cortico-cortical interactions and its functional roles in visual attention. *Commun. Biol.*, 7(1):550, 8 May 2024.
- [3] F Crick and C Koch. Towards a neurobiological theory of consciousness. *Seminars in the Neurosciences*, 2:263–275, 1990.
- [4] Andreas K Engel and Pascal Fries. Chapter 3 - neuronal oscillations, coherence, and consciousness. In Steven Laureys, Olivia Gosseries, and Giulio Tononi, editors, *The Neurology of Consciousness (Second Edition)*, pages 49–60. Academic Press, San Diego, 1 January 2016.
- [5] Evan Thompson and Francisco J Varela. Radical embodiment: neural dynamics and consciousness. *Trends Cogn. Sci.*, 5(10):418–425, 1 October 2001.
- [6] Lawrence M Ward. Synchronous neural oscillations and cognitive processes. *Trends Cogn. Sci.*, 7(12):553–559, December 2003.
- [7] Markus Siegel, Tobias H Donner, and Andreas K Engel. Spectral fingerprints of large-scale neuronal interactions. *Nat. Rev. Neurosci.*, 13(2):121–134, 11 January 2012.
- [8] Satu Palva and J Matias Palva. New vistas for alpha-frequency band oscillations. *Trends Neurosci.*, 30(4):150–158, April 2007.
- [9] Timothy J Buschman and Earl K Miller. Top-down versus bottom-up control of attention in the prefrontal and posterior parietal cortices. *Science*, 315(5820):1860–1862, 30 March 2007.
- [10] Georgia G Gregoriou, Stephen J Gotts, Huihui Zhou, and Robert Desimone. High-frequency, long-range coupling between prefrontal and visual cortex during attention. *Science*, 324(5931):1207–1210, 29 May 2009.
- [11] Georgia G Gregoriou, Stephen J Gotts, and Robert Desimone. Cell-type-specific synchronization of neural activity in FEF with V4 during attention. *Neuron*, 73(3):581–594, 9 February 2012.
- [12] E Salinas and T J Sejnowski. Correlated neuronal activity and the flow of neural information. *Nat. Rev. Neurosci.*, 2(8):539–550, August 2001.
- [13] P Fries, J H Reynolds, A E Rorie, and R Desimone. Modulation of oscillatory neuronal synchronization by selective visual attention. *Science*, 291(5508):1560–1563, 23 February 2001.
- [14] Markus Siegel, Tobias H Donner, Robert Oostenveld, Pascal Fries, and Andreas K Engel. Neuronal synchronization along the dorsal visual pathway reflects the focus of spatial attention. *Neuron*, 60(4):709–719, 26 November 2008.
- [15] Jochem van Kempen, Marc A Gieselmann, Michael Boyd, Nicholas A Steinmetz, Tirin Moore, Tatiana A Engel, and Alexander Thiele. Top-down coordination of local cortical state during selective attention. *Neuron*, 109(5):894–904.e8, 3 March 2021.

- [16] C M Gray, P König, A K Engel, and W Singer. Oscillatory responses in cat visual cortex exhibit inter-columnar synchronization which reflects global stimulus properties. *Nature*, 338(6213):334–337, 23 March 1989.
- [17] Markus Siegel, Andreas K Engel, and Tobias H Donner. Cortical network dynamics of perceptual decision-making in the human brain. *Front. Hum. Neurosci.*, 5:21, 28 February 2011.
- [18] Markus Siegel, Timothy J Buschman, and Earl K Miller. Cortical information flow during flexible sensorimotor decisions. *Science*, 348(6241):1352–1355, 19 June 2015.
- [19] Scott L Brincat, Jacob A Donoghue, Meredith K Mahnke, Simon Kornblith, Mikael Lundqvist, and Earl K Miller. Interhemispheric transfer of working memories. *Neuron*, 109(6):1055–1066.e4, 17 March 2021.
- [20] R F Salazar, N M Dotson, S L Bressler, and C M Gray. Content-specific fronto-parietal synchronization during visual working memory. *Science*, 338(6110):1097–1100, 23 November 2012.
- [21] A K Engel and W Singer. Temporal binding and the neural correlates of sensory awareness. *Trends Cogn. Sci.*, 5(1):16–25, 1 January 2001.
- [22] W Singer and C M Gray. Visual feature integration and the temporal correlation hypothesis. *Annu. Rev. Neurosci.*, 18:555–586, 1995.
- [23] V N Murthy and E E Fetz. Oscillatory activity in sensorimotor cortex of awake monkeys: synchronization of local field potentials and relation to behavior. *J. Neurophysiol.*, 76(6):3949–3967, December 1996.
- [24] Laureline Logiaco, L F Abbott, and Sean Escola. Thalamic control of cortical dynamics in a model of flexible motor sequencing. *Cell Rep.*, 35(9):109090, 1 June 2021.
- [25] Fritzie I Arce-McShane, Callum F Ross, Kazutaka Takahashi, Barry J Sessle, and Nicholas G Hatsopoulos. Primary motor and sensory cortical areas communicate via spatiotemporally coordinated networks at multiple frequencies. *Proc. Natl. Acad. Sci. U. S. A.*, 113(18):5083–5088, 3 May 2016.
- [26] Matthew G Perich, Juan A Gallego, and Lee E Miller. A neural population mechanism for rapid learning. *Neuron*, 100(4):964–976.e7, 21 November 2018.
- [27] Matthew T Kaufman, Mark M Churchland, Stephen I Ryu, and Krishna V Shenoy. Cortical activity in the null space: permitting preparation without movement. *Nat. Neurosci.*, 17(3):440–448, March 2014.
- [28] Scott L Brincat and Earl K Miller. Frequency-specific hippocampal-prefrontal interactions during associative learning. *Nat. Neurosci.*, 18(4):576–581, April 2015.
- [29] Matthew W Jones and Matthew A Wilson. Theta rhythms coordinate hippocampal-prefrontal interactions in a spatial memory task. *PLoS Biol.*, 3(12):e402, December 2005.
- [30] A G Siapas and M A Wilson. Coordinated interactions between hippocampal ripples and cortical spindles during slow-wave sleep. *Neuron*, 21(5):1123–1128, November 1998.
- [31] Antonio Fernández-Ruiz, Azahara Oliva, Marisol Soula, Florbela Rocha-Almeida, Gergo A Nagy, Gonzalo Martín-Vazquez, and György Buzsáki. Gamma rhythm communication between entorhinal cortex and dentate gyrus neuronal assemblies. *Science*, 372(6537), 2 April 2021.
- [32] Evan G Antzoulatos and Earl K Miller. Increases in functional connectivity between prefrontal cortex and striatum during category learning. *Neuron*, 83(1):216–225, 2 July 2014.
- [33] João D Semedo, Amin Zandvakili, Christian K Machens, Byron M Yu, and Adam Kohn. Cortical areas interact through a communication subspace. *Neuron*, 102(1):249–259.e4, 3 April 2019.

- [34] João D Semedo, Evren Gokcen, Christian K Machens, Adam Kohn, and Byron M Yu. Statistical methods for dissecting interactions between brain areas. *Curr. Opin. Neurobiol.*, 65:59–69, December 2020.
- [35] Camden J MacDowell, Alexandra Libby, Caroline I Jahn, Sina Tafazoli, and Timothy J Buschman. Multiplexed subspaces route neural activity across brain-wide networks. *bioRxiv*, 12 February 2023.
- [36] Shi Gu, Fabio Pasqualetti, Matthew Cieslak, Qawi K Telesford, Alfred B Yu, Ari E Kahn, John D Medaglia, Jean M Vettel, Michael B Miller, Scott T Grafton, and Danielle S Bassett. Controllability of structural brain networks. *Nat. Commun.*, 6:8414, 1 October 2015.
- [37] Ta-Chu Kao and Guillaume Hennequin. Neuroscience out of control: control-theoretic perspectives on neural circuit dynamics. *Curr. Opin. Neurobiol.*, 58:122–129, October 2019.
- [38] Danielle S Bassett and Olaf Sporns. Network neuroscience. *Nat. Neurosci.*, 20(3):353–364, 23 February 2017.
- [39] Ta-Chu Kao, Mahdieh S Sadabadi, and Guillaume Hennequin. Optimal anticipatory control as a theory of motor preparation: A thalamo-cortical circuit model. *Neuron*, 109(9):1567–1581.e12, 5 May 2021.
- [40] L Kozachkov, M Ennis, and J Slotine. RNNs of RNNs: Recursive construction of stable assemblies of recurrent neural networks. *Adv. Neural Inf. Process. Syst.*, 2022.
- [41] Marine Schimel, Ta-Chu Kao, Kristopher T Jensen, and Guillaume Hennequin. iLQR-VAE : control-based learning of input-driven dynamics with applications to neural data. In *International Conference on Learning Representations*, 6 October 2021.
- [42] Sarah Feldt Muldoon, Fabio Pasqualetti, Shi Gu, Matthew Cieslak, Scott T. Grafton, Jean M. Vettel, and Danielle S. Bassett. Stimulation-based control of dynamic brain networks. *PLOS Computational Biology*, 12(9):1–23, 09 2016. doi: 10.1371/journal.pcbi.1005076. URL <https://doi.org/10.1371/journal.pcbi.1005076>.
- [43] Kristofer Bouchard and Ankit Kumar. Feedback controllability is a normative theory of neural population dynamics. *Research Square*, 29 March 2024.
- [44] Urs Braun, Anais Harneit, Giulio Pergola, Tommaso Menara, Axel Schäfer, Richard F Betzel, Zhenxiang Zang, Janina I Schweiger, Xiaolong Zhang, Kristina Schwarz, Junfang Chen, Giuseppe Blasi, Alessandro Bertolino, Daniel Durstewitz, Fabio Pasqualetti, Emanuel Schwarz, Andreas Meyer-Lindenberg, Danielle S Bassett, and Heike Tost. Brain network dynamics during working memory are modulated by dopamine and diminished in schizophrenia. *Nat. Commun.*, 12(1):3478, 9 June 2021.
- [45] Dale Zhou, Yoona Kang, Danielle Cosme, Mia Jovanova, Xiaosong He, Arun Mahadevan, Jeesung Ahn, Ovidiu Stanoiu, Julia K Brynildsen, Nicole Cooper, Eli J Cornblath, Linden Parkes, Peter J Mucha, Kevin N Ochsner, David M Lydon-Staley, Emily B Falk, and Dani S Bassett. Mindful attention promotes control of brain network dynamics for self-regulation and discontinues the past from the present. *Proc. Natl. Acad. Sci. U. S. A.*, 120(2):e2201074119, 10 January 2023.
- [46] Daniela Zöller, Corrado Sandini, Marie Schaer, Stephan Eliez, Danielle S Bassett, and Dimitri Van De Ville. Structural control energy of resting-state functional brain states reveals less cost-effective brain dynamics in psychosis vulnerability. *Hum. Brain Mapp.*, 42(7):2181–2200, May 2021.
- [47] Xiaosong He, Lorenzo Caciagli, Linden Parkes, Jennifer Stiso, Teresa M Karrer, Jason Z Kim, Zhixin Lu, Tommaso Menara, Fabio Pasqualetti, Michael R Sperling, Joseph I Tracy, and Dani S Bassett. Uncovering the biological basis of control energy: Structural and metabolic correlates of energy inefficiency in temporal lobe epilepsy. *Sci. Adv.*, 8(45):eabn2293, 11 November 2022.

- [48] S Parker Singleton, Andrea I Luppi, Robin L Carhart-Harris, Josephine Cruzat, Leor Roseman, David J Nutt, Gustavo Deco, Morten L Kringelbach, Emmanuel A Stamatakis, and Amy Kuceyeski. Receptor-informed network control theory links LSD and psilocybin to a flattening of the brain’s control energy landscape. *Nat. Commun.*, 13(1):1–13, 3 October 2022.
- [49] John D Medaglia, Denise Y Harvey, Nicole White, Apoorva Kelkar, Jared Zimmerman, Danielle S Bassett, and Roy H Hamilton. Network controllability in the inferior frontal gyrus relates to controlled language variability and susceptibility to TMS. *J. Neurosci.*, 38(28): 6399–6410, 11 July 2018.
- [50] Zhoukang Wu, Liangjiecheng Huang, Min Wang, and Xiaosong He. Development of the brain network control theory and its implications. *Psychoradiology*, 4:kkæ028, 14 December 2024.
- [51] Weidong Cai, Srikanth Ryali, Ramkrishna Pasumarthy, Viswanath Talasila, and Vinod Menon. Dynamic causal brain circuits during working memory and their functional controllability. *Nat. Commun.*, 12(1):3314, 29 June 2021.
- [52] Ali Moradi Amani, Amirhessam Tahmassebi, Andreas Stadlbauer, Uwe Meyer-Baese, Vincent Noblet, Frederic Blanc, Hagen Malberg, and Anke Meyer-Baese. Controllability of functional and structural brain networks. *Complexity*, 2024(1):7402894, 1 January 2024.
- [53] Qian Li, Li Yao, Wanfang You, Jiang Liu, Shikuang Deng, Bin Li, Lekai Luo, Youjin Zhao, Yuxia Wang, Yaxuan Wang, Qian Zhang, Fenghua Long, John A Sweeney, Shi Gu, Fei Li, and Qiyong Gong. Controllability of functional brain networks and its clinical significance in first-episode schizophrenia. *Schizophr. Bull.*, 49(3):659–668, 3 May 2023.
- [54] Yang-Yu Liu, Jean-Jacques Slotine, and Albert-László Barabási. Controllability of complex networks. *Nature*, 473(7346):167–173, 12 May 2011.
- [55] Linden Parkes, Jason Z Kim, Jennifer Stiso, Julia K Brynildsen, Matthew Cieslak, Sydney Covitz, Raquel E Gur, Ruben C Gur, Fabio Pasqualetti, Russell T Shinohara, Dale Zhou, Theodore D Satterthwaite, and Dani S Bassett. A network control theory pipeline for studying the dynamics of the structural connectome. *Nat. Protoc.*, 19(12):3721–3749, 29 December 2024.
- [56] David R. Tyner and Andrew D. Lewis. Geometric jacobian linearization and LQR theory. *J. Geom. Mech.*, 2(4):397–440, 2010.
- [57] Robert E Kass, Heejong Bong, Motolani Olarinre, Qi Xin, and Konrad N Urban. Identification of interacting neural populations: methods and statistical considerations. *J. Neurophysiol.*, 130(3):475–496, 1 September 2023.
- [58] Matthew G Perich, Charlotte Arlt, Sofia Soares, Megan E Young, Clayton P Mosher, Juri Minxha, Eugene Carter, Ueli Rutishauser, Peter H Rudebeck, Christopher D Harvey, and Kanaka Rajan. Inferring brain-wide interactions using data-constrained recurrent neural network models. *bioRxiv*, page 2020.12.18.423348, 11 March 2021.
- [59] Matthew G Perich and Kanaka Rajan. Rethinking brain-wide interactions through multi-region ‘network of networks’ models. *Curr. Opin. Neurobiol.*, 65:146–151, December 2020.
- [60] Aaron S Andalman, Vanessa M Burns, Matthew Lovett-Barron, Michael Broxton, Ben Poole, Samuel J Yang, Logan Grosenick, Talia N Lerner, Ritchie Chen, Tyler Benster, Philippe Mourrain, Marc Levoy, Kanaka Rajan, and Karl Deisseroth. Neuronal dynamics regulating brain and behavioral state transitions. *Cell*, 177(4):970–985.e20, 2 May 2019.
- [61] Lucas Pinto, Kanaka Rajan, Brian DePasquale, Stephan Y Thiberge, David W Tank, and Carlos D Brody. Task-dependent changes in the large-scale dynamics and necessity of cortical regions. *Neuron*, 104(4):810–824.e9, 20 November 2019.
- [62] Michael Kleinman, Chandramouli Chandrasekaran, and Jonathan Kao. A mechanistic multi-area recurrent network model of decision-making. *Adv. Neural Inf. Process. Syst.*, 34:23152–23165, 6 December 2021.

- [63] Joao Barbosa, Rémi Proville, Chris C Rodgers, Michael R DeWeese, Srdjan Ostojic, and Yves Boubenec. Early selection of task-relevant features through population gating. *Nat. Commun.*, 14(1):6837, 27 October 2023.
- [64] Evren Gokcen, Anna I Jasper, João D Semedo, Amin Zandvakili, Adam Kohn, Christian K Machens, and Byron M Yu. Disentangling the flow of signals between populations of neurons. *Nature Computational Science*, 2(8):512–525, 18 August 2022.
- [65] Evren Gokcen, Anna Ivic Jasper, Alison Xu, Adam Kohn, Christian K Machens, and Byron M Yu. Uncovering motifs of concurrent signaling across multiple neuronal populations. In *Thirty-seventh Conference on Neural Information Processing Systems*, 2 November 2023.
- [66] Sadegh Ebrahimi, Jérôme Lecoq, Oleg Rumyantsev, Tugce Tasci, Yanping Zhang, Cristina Irimia, Jane Li, Surya Ganguli, and Mark J Schnitzer. Emergent reliability in sensory cortical coding and inter-area communication. *Nature*, 605(7911):713–721, May 2022.
- [67] Jordan Rodu, Natalie Klein, Scott L Brincat, Earl K Miller, and Robert E Kass. Detecting multivariate cross-correlation between brain regions. *J. Neurophysiol.*, 120(4):1962–1972, 1 October 2018.
- [68] João D Semedo, Anna I Jasper, Amin Zandvakili, Aravind Krishna, Amir Aschner, Christian K Machens, Adam Kohn, and Byron M Yu. Feedforward and feedback interactions between visual cortical areas use different population activity patterns. *Nat. Commun.*, 13(1):1099, 1 March 2022.
- [69] Hao Ye, Ethan R Deyle, Luis J Gilarranz, and George Sugihara. Distinguishing time-delayed causal interactions using convergent cross mapping. *Sci. Rep.*, 5:14750, 5 October 2015.
- [70] Satoshi Tajima, Toru Yanagawa, Naotaka Fujii, and Taro Toyoizumi. Untangling brain-wide dynamics in consciousness by cross-embedding. *PLoS Comput. Biol.*, 11(11):e1004537, November 2015.
- [71] George Sugihara, Robert May, Hao Ye, Chih-Hao Hsieh, Ethan Deyle, Michael Fogarty, and Stephan Munch. Detecting causality in complex ecosystems. *Science*, 338(6106):496–500, 26 October 2012.
- [72] Joshua I Glaser, Matthew Whiteway, John P Cunningham, Liam Paninski, and Scott W Linderman. Recurrent switching dynamical systems models for multiple interacting neural populations. In *Advances in Neural Information Processing Systems 33 (NeurIPS 2020)*, page 2020.10.21.349282, 22 October 2020.
- [73] Orren Karniol-Tambour, David M Zoltowski, E Mika Diamanti, Lucas Pinto, David W Tank, Carlos D Brody, and Jonathan W Pillow. Modeling communication and switching nonlinear dynamics in multi-region neural activity. *bioRxiv*, page 2022.09.13.507841, 15 September 2022.
- [74] Steven L Bressler and Anil K Seth. Wiener-granger causality: a well established methodology. *Neuroimage*, 58(2):323–329, 15 September 2011.
- [75] Anil K Seth, Adam B Barrett, and Lionel Barnett. Granger causality analysis in neuroscience and neuroimaging. *J. Neurosci.*, 35(8):3293–3297, 25 February 2015.
- [76] K J Friston, L Harrison, and W Penny. Dynamic causal modelling. *Neuroimage*, 19(4):1273–1302, August 2003.
- [77] Yu Chen, Hannah Douglas, Bryan J Medina, Motolani Olarinre, Joshua H Siegle, and Robert E Kass. Population burst propagation across interacting areas of the brain. *J. Neurophysiol.*, 128(6):1578–1592, 1 December 2022.
- [78] Josuan Calderon and Gordon J Berman. Inferring the time-varying coupling of dynamical systems with temporal convolutional autoencoders. *eLife*, 12 November 2024.
- [79] Ziyu Lu, Anika Tabassum, Shruti Kulkarni, Lu Mi, J Nathan Kutz, Eric Shea-Brown, and Seung-Hwan Lim. Attention for causal relationship discovery from biological neural dynamics. In *NeurIPS 2023 Workshop on Causal Representation Learning*, 12 November 2023.

- [80] Jean-Jacques E Slotine and Weiping Li. *Applied Nonlinear Control*. Prentice-Hall, 1991.
- [81] R W Brockett. Nonlinear systems and differential geometry. *Proc. IEEE Inst. Electr. Electron. Eng.*, 64(1):61–72, 1976.
- [82] G W Haynes and H Hermes. Nonlinear controllability via lie theory. *SIAM J. Control*, 8(4): 450–460, November 1970.
- [83] Cesar O Aguilar and Bahman Ghahsifard. Necessary conditions for controllability of nonlinear networked control systems. In *2014 American Control Conference*, pages 5379–5383. IEEE, June 2014.
- [84] Andrew J Whalen, Sean N Brennan, Timothy D Sauer, and Steven J Schiff. Observability and controllability of nonlinear networks: The role of symmetry. *Phys. Rev. X*, 5(1):011005, 23 January 2015.
- [85] Le-Zhi Wang, Ri-Qi Su, Zi-Gang Huang, Xiao Wang, Wen-Xu Wang, Celso Grebogi, and Ying-Cheng Lai. A geometrical approach to control and controllability of nonlinear dynamical networks. *Nat. Commun.*, 7:11323, 14 April 2016.
- [86] Yang Tang, Huijun Gao, Wei Zou, and Jürgen Kurths. Identifying controlling nodes in neuronal networks in different scales. *PLoS One*, 7(7):e41375, 27 July 2012.
- [87] Gautam Kumar, Delsin Menolascino, and Shinung Ching. Input novelty as a control metric for time varying linear systems. *arXiv [math.OC]*, 21 November 2014.
- [88] Ethan R Deyale, Robert M May, Stephan B Munch, and George Sugihara. Tracking and forecasting ecosystem interactions in real time. *Proc. Biol. Sci.*, 283(1822), 13 January 2016.
- [89] Tosif Ahamed, Antonio C Costa, and Greg J Stephens. Capturing the continuous complexity of behaviour in caenorhabditis elegans. *Nat. Phys.*, 17(2):275–283, 5 October 2020.
- [90] Frédéric Latrémolière, Sadananda Narayanappa, and Petr Vojtěchovský. Estimating the jacobian matrix of an unknown multivariate function from sample values by means of a neural network. *arXiv [cs.LG]*, April 2022.
- [91] Jonathan Lorraine and Safwan Hossain. Jacnet: Learning functions with structured jacobians. *arXiv preprint arXiv:2408.13237*, 2024.
- [92] Hadi Beik-Mohammadi, Søren Hauberg, Georgios Arvanitidis, Nadia Figueroa, Gerhard Neumann, and Leonel Roza. Neural contractive dynamical systems. *arXiv [cs.RO]*, 17 January 2024.
- [93] Winfried Lohmiller and Jean-Jacques E Slotine. On contraction analysis for non-linear systems. *Automatica*, 34(6):683–696, 1 June 1998.
- [94] Florian Hess, Zahra Monfared, Manuel Brenner, and Daniel Durstewitz. Generalized teacher forcing for learning chaotic dynamics. *arXiv [cs.LG]*, 7 June 2023.
- [95] Abhiram Iyer, Sarthak Chandra, Sugandha Sharma, and I Fiete. Flexible mapping of abstract domains by grid cells via self-supervised extraction and projection of generalized velocity signals. *Neural Inf Process Syst*, 37:85441–85466, 2024.
- [96] Balth van der Pol. LXXXVIII. on “relaxation-oscillations”. *The London, Edinburgh, and Dublin Philosophical Magazine and Journal of Science*, 2(11):978–992, 1 November 1926.
- [97] Edward N Lorenz. Deterministic nonperiodic flow. *J. Atmos. Sci.*, 20(2):130–141, 1 March 1963.
- [98] E N Lorenz. Predictability: a problem partly solved. In *ECMWF Seminar on Predictability, 4-8 September 1995*. European Centre for Medium-Range Weather Forecasts, 1996.
- [99] William Gilpin. Chaos as an interpretable benchmark for forecasting and data-driven modelling. In J Vanschoren and S Yeung, editors, *Proceedings of the Neural Information Processing Systems Track on Datasets and Benchmarks*, volume 1, 2021.

- [100] William Gilpin. Model scale versus domain knowledge in statistical forecasting of chaotic systems. *Phys. Rev. Res.*, 5(4):043252, 15 December 2023.
- [101] Pablo Gómez, Håvard Hem Toftevaag, and Gabriele Meoni. torchquad: Numerical Integration in Arbitrary Dimensions with PyTorch. *Journal of Open Source Software*, 64(6), 2021. doi: 10.21105/joss.03439.
- [102] Ricky T Q Chen, Yulia Rubanova, Jesse Bettencourt, and David K Duvenaud. Neural ordinary differential equations. In S Bengio, H Wallach, H Larochelle, K Grauman, N Cesa-Bianchi, and R Garnett, editors, *Advances in Neural Information Processing Systems*, volume 31. Curran Associates, Inc., 2018.
- [103] Adam Paszke, Sam Gross, Francisco Massa, Adam Lerer, James Bradbury, Gregory Chanan, Trevor Killeen, Zeming Lin, N Gimelshein, L Antiga, Alban Desmaison, Andreas Köpf, E Yang, Zach DeVito, Martin Raison, Alykhan Tejani, Sasank Chilamkurthy, Benoit Steiner, Lu Fang, Junjie Bai, and Soumith Chintala. PyTorch: An imperative style, high-performance deep learning library. *Adv. Neural Inf. Process. Syst.*, abs/1912.01703, 3 December 2019.
- [104] Judy Hoffman, Daniel A Roberts, and Sho Yaida. Robust learning with jacobian regularization. *arXiv [stat.ML]*, 7 August 2019.
- [105] Alexander Wikner, Joseph Harvey, Michelle Girvan, Brian R Hunt, Andrew Pomerance, Thomas Antonsen, and Edward Ott. Stabilizing machine learning prediction of dynamics: Novel noise-inspired regularization tested with reservoir computing. *Neural Netw.*, 170: 94–110, 1 February 2024.
- [106] Steffen Schneider, Rodrigo González Laiz, Anastasiia Filippova, Markus Frey, and Mackenzie Weygandt Mathis. Time-series attribution maps with regularized contrastive learning. *arXiv [stat.ML]*, 17 February 2025.
- [107] Panos J Antsaklis and Anthony N Michel. *Linear Systems*. Birkhäuser Boston, 2006.
- [108] Yu Kawano and Jacquélien M A Scherpen. Empirical differential gramians for nonlinear model reduction. *Automatica (Oxf.)*, 127(109534):109534, May 2021.
- [109] Gustav Lindmark and Claudio Altafini. Minimum energy control for complex networks. *Sci. Rep.*, 8(1):3188, 16 February 2018.
- [110] Matthew F Panichello and Timothy J Buschman. Shared mechanisms underlie the control of working memory and attention. *Nature*, 31 March 2021.
- [111] Weiwei Li and Emanuel Todorov. Iterative linear quadratic regulator design for nonlinear biological movement systems. *ICINCO 2004, Proceedings of the First International Conference on Informatics in Control, Automation and Robotics, Setúbal, Portugal, August 25-28, 2004*, 1: 222–229, 1 January 2004.
- [112] Yuval Tassa, Tom Erez, and Emanuel Todorov. Synthesis and stabilization of complex behaviors through online trajectory optimization. In *2012 IEEE/RSJ International Conference on Intelligent Robots and Systems*, pages 4906–4913. IEEE, October 2012.
- [113] S Ouala, D Nguyen, L Drumetz, B Chapron, A Pascual, F Collard, L Gaultier, and R Fablet. Learning latent dynamics for partially observed chaotic systems. *Chaos*, 30(10):103121, 20 October 2020.
- [114] William Gilpin. Deep reconstruction of strange attractors from time series. In *Proceedings of the 34th International Conference on Neural Information Processing Systems*, number Article 18 in NIPS’20, pages 204–216, Red Hook, NY, USA, 6 December 2020. Curran Associates Inc.
- [115] Peter Y Lu, Joan Ariño Bernad, and Marin Soljačić. Discovering sparse interpretable dynamics from partial observations. *Commun. Phys.*, 5(1):1–7, 12 August 2022.
- [116] Elise Özalp, Georgios Margazoglou, and Luca Magri. Reconstruction, forecasting, and stability of chaotic dynamics from partial data. *Chaos*, 33(9):093107, 1 September 2023.

- [117] Charles D Young and Michael D Graham. Deep learning delay coordinate dynamics for chaotic attractors from partial observable data. *Phys. Rev. E.*, 107(3-1):034215, 30 March 2023.
- [118] Sean Jaffe, Alexander Davydov, Deniz Lapsekili, Ambuj Singh, and Francesco Bullo. Learning neural contracting dynamics: Extended linearization and global guarantees. *arXiv [cs.LG]*, 12 February 2024.
- [119] Steven H Strogatz. *Nonlinear Dynamics and Chaos: With Applications to Physics, Biology, Chemistry, and Engineering*. Avalon Publishing, 29 July 2014.
- [120] Rudiger U Seydel. *Practical bifurcation and stability analysis*. Interdisciplinary applied mathematics. Springer, New York, NY, 3 edition, 10 December 2009.
- [121] Luca Dieci, Robert D Russell, and Erik S Van Vleck. On the computation of lyapunov exponents for continuous dynamical systems. *SIAM J. Numer. Anal.*, 34(1):402–423, 1997.
- [122] F Christiansen and H H Rugh. Computing lyapunov spectra with continuous gram - schmidt orthonormalization. *Nonlinearity*, 10(5):1063, 1 September 1997.
- [123] Vincent Duchaine and Clement M Gosselin. Investigation of human-robot interaction stability using lyapunov theory. In *2008 IEEE International Conference on Robotics and Automation*, pages 2189–2194. IEEE, May 2008.
- [124] Ian R Manchester and Jean-Jacques E Slotine. Control contraction metrics: Convex and intrinsic criteria for nonlinear feedback design. *IEEE Trans. Automat. Contr.*, 62(6):3046–3053, June 2017.
- [125] Patrick M Wensing and Jean-Jacques Slotine. Beyond convexity-contraction and global convergence of gradient descent. *PLoS One*, 15(8):e0236661, 4 August 2020.
- [126] Leo Kozachkov, Mikael Lundqvist, Jean Jacques Slotine, and Earl K Miller. Achieving stable dynamics in neural circuits. *PLoS Comput. Biol.*, 16(8):1–15, 2020.
- [127] N A M Hootsmans and S Dubowsky. Large motion control of mobile manipulators including vehicle suspension characteristics. In *Proceedings. 1991 IEEE International Conference on Robotics and Automation*, pages 2336–2341 vol.3. IEEE Comput. Soc. Press, 2002.
- [128] S Ali A Moosavian and Evangelos Papadopoulos. Modified transpose jacobian control of robotic systems. *Automatica (Oxf.)*, 43(7):1226–1233, 1 July 2007.
- [129] Krzysztof Tchoń, Adam Ratajczak, and Ida Góral. Lagrangian jacobian inverse for nonholonomic robotic systems. *Nonlinear Dyn.*, 82(4):1923–1932, December 2015.
- [130] Chevelle Brudey, Jeanie Park, Jan Wiaderkiewicz, Ihori Kobayashi, Thomas A Mellman, and Paul J Marvar. Autonomic and inflammatory consequences of posttraumatic stress disorder and the link to cardiovascular disease. *Am. J. Physiol. Regul. Integr. Comp. Physiol.*, 309(4):R315–21, 15 August 2015.
- [131] Gail A Alvares, Daniel S Quintana, Ian B Hickie, and Adam J Guastella. Autonomic nervous system dysfunction in psychiatric disorders and the impact of psychotropic medications: a systematic review and meta-analysis. *J. Psychiatry Neurosci.*, 41(2):89–104, March 2016.
- [132] Jeffrey J Goldberger, Rishi Arora, Una Buckley, and Kalyanam Shivkumar. Autonomic nervous system dysfunction: JACC focus seminar. *J. Am. Coll. Cardiol.*, 73(10):1189–1206, 19 March 2019.
- [133] Li Xiong and Thomas W H Leung. Autonomic dysfunction in neurological disorders. *Aging (Albany NY)*, 11(7):1903–1904, 9 April 2019.
- [134] Albert-László Barabási, Natali Gulbahce, and Joseph Loscalzo. Network medicine: a network-based approach to human disease. *Nat. Rev. Genet.*, 12(1):56–68, January 2011.

- [135] Matthew T Maurano, Richard Humbert, Eric Rynes, Robert E Thurman, Eric Haugen, Hao Wang, Alex P Reynolds, Richard Sandstrom, Hongzhu Qu, Jennifer Brody, Anthony Shafer, Fidencio Neri, Kristen Lee, Tanya Kutayavin, Sandra Stehling-Sun, Audra K Johnson, Theresa K Canfield, Erika Giste, Morgan Diegel, Daniel Bates, R Scott Hansen, Shane Neph, Peter J Sabo, Shelly Heimfeld, Antony Raubitschek, Steven Ziegler, Chris Cotsapas, Nona Sotoodehnia, Ian Glass, Shamil R Sunyaev, Rajinder Kaul, and John A Stamatoyannopoulos. Systematic localization of common disease-associated variation in regulatory DNA. *Science*, 337(6099): 1190–1195, 7 September 2012.
- [136] Piyush B Madhamshettiwar, Stefan R Maetschke, Melissa J Davis, Antonio Reverter, and Mark A Ragan. Gene regulatory network inference: evaluation and application to ovarian cancer allows the prioritization of drug targets. *Genome Med.*, 4(5):41, 1 May 2012.
- [137] Tong Ihn Lee and Richard A Young. Transcriptional regulation and its misregulation in disease. *Cell*, 152(6):1237–1251, 14 March 2013.
- [138] Paula Unger Avila, Tsimafei Padvitski, Ana Carolina Leote, He Chen, Julio Saez-Rodriguez, Martin Kann, and Andreas Beyer. Gene regulatory networks in disease and ageing. *Nat. Rev. Nephrol.*, 20(9):616–633, 12 September 2024.
- [139] Jaideep Pathak, Zhixin Lu, Brian R Hunt, Michelle Girvan, and Edward Ott. Using machine learning to replicate chaotic attractors and calculate lyapunov exponents from data. *Chaos*, 27(12):121102, December 2017.
- [140] Jaideep Pathak, Brian Hunt, Michelle Girvan, Zhixin Lu, and Edward Ott. Model-free prediction of large spatiotemporally chaotic systems from data: A reservoir computing approach. *Phys. Rev. Lett.*, 120(2):024102, 12 January 2018.
- [141] Zhixin Lu, Jaideep Pathak, Brian Hunt, Michelle Girvan, Roger Brockett, and Edward Ott. Reservoir observers: Model-free inference of unmeasured variables in chaotic systems. *Chaos*, 27(4):041102, 5 April 2017.
- [142] Amitava Banerjee, Jaideep Pathak, Rajarshi Roy, Juan G Restrepo, and Edward Ott. Using machine learning to assess short term causal dependence and infer network links. *Chaos*, 29(12):121104, 26 December 2019.
- [143] Mattia Rigotti, Omri Barak, Melissa R Warden, Xiao-Jing Wang, Nathaniel D Daw, Earl K Miller, and Stefano Fusi. The importance of mixed selectivity in complex cognitive tasks. *Nature*, 497(7451):585–590, 30 May 2013.

A JacobianODE technical details

A.1 Parameterizing the time derivative via the Jacobian

Recall that, for the time derivative function $\mathbf{f}$, we have that the path integral is independent of the choice of path:

$$\mathbf{f}(\mathbf{x}(t_f)) - \mathbf{f}(\mathbf{x}(t_i)) = \int_{\mathcal{C}} \mathbf{J} ds = \int_{t_i}^{t_f} \mathbf{J}(\mathbf{c}(r)) \mathbf{c}'(r) dr, \quad (7)$$

where $\mathcal{C}$ is a piecewise smooth curve in $\mathbb{R}^n$ and $\mathbf{c} : [t_i, t_f] \rightarrow \mathcal{C}$ is a parameterization of $\mathcal{C}$ with $\mathbf{c}(t_i) = \mathbf{x}(t_i)$ and $\mathbf{c}(t_f) = \mathbf{x}(t_f)$ (also see equation 2). Path integration notably provides only differences between $\mathbf{f}$ at distinct points. Thus, knowing the value of $\mathbf{f}$ at one point is needed for trajectory reconstruction, which we do not assume. To address this, we note it is possible to represent $\mathbf{f}$ solely through the Jacobian at a point $\mathbf{x}(t)$ as

$$\mathbf{f}(\mathbf{x}(t)) = \frac{G(t_0, t; \mathbf{J}, \mathbf{c}) + \mathbf{x}(t) - \mathbf{x}(t_0)}{t - t_0}, \quad (8)$$

where

$$G(t_0, t; \mathbf{J}, \mathbf{c}) = \int_{t_0}^t \int_s^t \mathbf{J}(\mathbf{c}_{s,t}(r)) \mathbf{c}'_{s,t}(r) dr ds \quad (9)$$

and we have abbreviated $\mathbf{c}_{s,t}(r) = \mathbf{c}(r; s, t, \mathbf{x}(s), \mathbf{x}(t))$, a piecewise smooth curve on $[s, t]$ parameterized by r and beginning and ending at $\mathbf{x}(s)$ and $\mathbf{x}(t)$ respectively. Intuitively, we can circumvent the need to know $\mathbf{f}$ by recognizing that integrating $\mathbf{f}$ between time points will produce the difference in the system states at those time points, which are known to us. With this formalism, we avoid the need to represent $\mathbf{f}$ directly, thereby enabling all gradients to backpropagate through the Jacobian network. The proof of this formalism is presented below.

Proposition 1 (Jacobian-parameterized ODEs). *Let $\dot{\mathbf{x}}(t) = \mathbf{f}(\mathbf{x}(t))$ and let $\mathbf{J}_{\mathbf{f}}(\mathbf{x}(t)) = \mathbf{J}(\mathbf{x}(t)) = \frac{\partial}{\partial \mathbf{x}} \mathbf{f}(\mathbf{x}(t))$. Then given times t_0, t we can express $\mathbf{f}$ parameterized by the Jacobian as*

$$\mathbf{f}(\mathbf{x}(t)) = \frac{G(t_0, t; \mathbf{J}, \mathbf{c}) + \mathbf{x}(t) - \mathbf{x}(t_0)}{t - t_0},$$

where

$$G(t_0, t; \mathbf{J}, \mathbf{c}) = \int_{t_0}^t \int_s^t \mathbf{J}(\mathbf{c}_{s,t}(r)) \mathbf{c}'_{s,t}(r) dr ds$$

and we have abbreviated $\mathbf{c}_{s,t}(r) = \mathbf{c}(r; s, t, \mathbf{x}(s), \mathbf{x}(t))$, a piecewise smooth curve on $[s, t]$ parameterized by r and beginning and ending at $\mathbf{x}(s)$ and $\mathbf{x}(t)$ respectively.

Proof. For given times t, t_0 , and s , from the fundamental theorem of calculus, we have that

$$\mathbf{x}(t) - \mathbf{x}(t_0) = \int_{t_0}^t \mathbf{f}(\mathbf{x}(s)) ds$$

and

$$\mathbf{f}(\mathbf{x}(t)) - \mathbf{f}(\mathbf{x}(s)) = \int_s^t \mathbf{J}(\mathbf{c}_{s,t}(r)) \mathbf{c}'_{s,t}(r) dr$$

Letting

$$H(s, t) = \int_s^t \mathbf{J}(\mathbf{c}_{s,t}(r)) \mathbf{c}'_{s,t}(r) dr = \mathbf{f}(\mathbf{x}(t)) - \mathbf{f}(\mathbf{x}(s))$$

We then have that

$$\begin{aligned}
\int_{t_0}^t H(s, t) ds &= \int_{t_0}^t \mathbf{f}(\mathbf{x}(t)) - \mathbf{f}(\mathbf{x}(s)) ds \\
&= \int_{t_0}^t \mathbf{f}(\mathbf{x}(t)) ds - \int_{t_0}^t \mathbf{f}(\mathbf{x}(s)) ds \\
&= (t - t_0)\mathbf{f}(\mathbf{x}(t)) - (\mathbf{x}(t) - \mathbf{x}(t_0))
\end{aligned}$$

Letting now

$$G(t_0, t; \mathbf{J}, \mathbf{c}) = \int_{t_0}^t H(s, t) ds = \int_{t_0}^t \int_s^t \mathbf{J}(\mathbf{c}_{s,t}(r)) \mathbf{c}'_{s,t}(r) dr ds$$

We can see that

$$G(t_0, t; \mathbf{J}, \mathbf{c}) = (t - t_0)\mathbf{f}(\mathbf{x}(t)) - (\mathbf{x}(t) - \mathbf{x}(t_0))$$

and thus

$$\mathbf{f}(\mathbf{x}(t)) = \frac{G(t_0, t; \mathbf{J}, \mathbf{c}) + \mathbf{x}(t) - \mathbf{x}(t_0)}{t - t_0}$$

□

A.2 Path integrating to generate predictions

Given an initial observed trajectory $\mathbf{x}(t_0 + k\Delta t)$, $k = 0, \dots, b$ of length at least two ($b \geq 1$), we compute an estimate of $\hat{\mathbf{f}}(\mathbf{x}(t_0 + b\Delta t))$ of $\mathbf{f}(\mathbf{x}(t_0 + b\Delta t))$ by replacing $\mathbf{J}$ with $\hat{\mathbf{J}}^\theta$ in equation 8. Specifically, we compute an estimate of $\hat{\mathbf{f}}(\mathbf{x}(t_b))$ of $\mathbf{f}(\mathbf{x}(t_b))$ as

$$\hat{\mathbf{f}}(\mathbf{x}(t_b)) = \frac{G(t_0, t_b; \hat{\mathbf{J}}^\theta, \mathbf{c}) + \mathbf{x}(t_b) - \mathbf{x}(t_0)}{t_b - t_0}, \quad (10)$$

In practice, we compute this estimate by constructing a cubic spline on the initial observed trajectory using 15 points (i.e., $b = 14$). Given that the computation of G involves a double integral - one over states and over time - using a spline is computationally advantageous. This is because the path integral (and all intermediate steps) can be quickly computed along the full spline, with the result of each intermediate step along the path then being summed to approximate the time integral. The integral is computed by interpolating 4 points for every gap between observed points, resulting in a discretization of 58 points along the spline. Integrals are computed using the trapezoid method from `torchquad` [101].

Once we have constructed our estimate of $\hat{\mathbf{f}}(\mathbf{x}(t_b))$ we can estimate $\mathbf{f}$ at any other point $\mathbf{x}(t)$ as

$$\hat{\mathbf{f}}(\mathbf{x}(t)) = \begin{cases} H(t_b, t) + \hat{\mathbf{f}}(\mathbf{x}(t_b)), & \text{if } t_b < t \\ \hat{\mathbf{f}}(\mathbf{x}(t_b)) - H(t, t_b), & \text{if } t_b > t \\ \mathbf{f}(\mathbf{x}(t_b)) & \text{if } t_b = t \end{cases}$$

where by convention we integrate forwards in time and H is the path integral defined above, with $\hat{\mathbf{J}}$ in place of $\mathbf{J}$. In practice, for the integration path $\mathbf{c}(r; s, t, \mathbf{x}(s), \mathbf{x}(t))$, we construct a line from $\mathbf{x}(s)$ to $\mathbf{x}(t)$ as

$$\mathbf{c}(r; s, t, \mathbf{x}(s), \mathbf{x}(t)) = \left(1 - \frac{r-s}{t-s}\right) \mathbf{x}(s) + \frac{r-s}{t-s} \mathbf{x}(t)$$

to maintain the interpretability of having r in the range $[s, t]$ however it can be easily seen that setting $r' = \frac{r-s}{t-s}$ we recover the familiar line

$$\mathbf{c}(r') = (1 - r')\mathbf{x}(s) + r'\mathbf{x}(t)$$

with r' taking values on $[0, 1]$. Line integrals are computed with 20 discretization steps using the trapezoid method from `torchquad` [101].

Using $\hat{\mathbf{f}}(\mathbf{x}(t))$, we can then generate predictions as

$$\hat{\mathbf{x}}(t + \Delta t) = \mathbf{x}(t) + \int_t^{t+\Delta t} \hat{\mathbf{f}}(\mathbf{x}(\tau)) d\tau \quad (11)$$

where the integral can be computed by a standard ODE solver (we used the RK4 method from `torchdiffeq` with default values for the relative and absolute tolerance).

We can then compute the trajectory reconstruction loss as

$$\mathcal{L}_{\text{traj}}(\theta; \mathbf{x}) = \sum_{k=b+1}^{T-1} d(\mathbf{x}(t_0 + k\Delta t), \hat{\mathbf{x}}(t_0 + k\Delta t)) \quad (12)$$

where d is a distance measure between trajectories (e.g. mean squared error).

A.3 Loop closure loss

While the space of Jacobian functions $\hat{\mathbf{J}}^\theta$ that solve the dynamics problem may be quite large, we are interested only in functions $\hat{\mathbf{J}}^\theta$ that *both* solve the dynamics problem and generate matrices with rows that are conservative vector fields. To effectively restrict our optimization to this space, we employ a self-supervised loss building on the loss introduced by Iyer et al. [95]. This loss imposes a conservative constraint on the rows of the Jacobian matrix. Specifically, for any piecewise smooth curve $\mathcal{C}_{\text{loop}}$ starting and ending at the same state $\mathbf{x}_0$, we have that $\left\| \int_{\mathcal{C}_{\text{loop}}} \mathbf{J} ds \right\|_2 = 0$.

Thus, given n_{loops} sets of any L (not necessarily sequential) points $\{\mathbf{x}^{(l)}(t_1), \mathbf{x}^{(l)}(t_2), \dots, \mathbf{x}^{(l)}(t_L)\}$ we can form a loop $\mathcal{C}_{\text{loop}}^{(l)}$ consisting of the sequence lines from $\mathbf{x}^{(l)}(t_i)$ to $\mathbf{x}^{(l)}(t_{i+1})$, $i = 1, \dots, L-1$, followed by a line from $\mathbf{x}^{(l)}(t_L)$ to $\mathbf{x}^{(l)}(t_1)$. We then define the following regularization loss to enforce this conservation constraint:

$$\mathcal{L}_{\text{loop}}(\theta; \mathbf{x}) = \frac{1}{n_{\text{loops}}} \sum_{l=1}^{n_{\text{loops}}} \frac{1}{n} \left\| \int_{\mathcal{C}_{\text{loop}}^{(l)}} \hat{\mathbf{J}}^\theta ds \right\|_2^2 \quad (13)$$

We note that while other choices for loops would have been possible (for instance, forward and backward passes along observed trajectories, or arbitrary circles beginning and ending at the same point), the approach presented here has several advantages (as discussed in Section 3.3). Namely, it balances computational tractability with uniform sampling of the directions in the tangent space. Enforcing that all the points that comprise the loop are on the data manifold ensures that the loops are integrated in directions that are as informative as possible for estimating the Jacobian of the given dynamical system.

B Control-theoretic analysis details

B.1 Gramian computation

We begin with equation 4 from Section 5, in which we separate the locally-linearized dynamics in the tangent space into separate pairwise inter-subsystem control interactions. These pairwise control interactions are in general of the form

$$\delta \dot{\mathbf{x}}^A(t) = \mathbf{A}(t) \delta \mathbf{x}^A(t) + \mathbf{B}(t) \delta \mathbf{x}^B(t)$$

where $\mathbf{A}$ is the within-subsystem Jacobian, $\mathbf{B}$ is the across subsystem Jacobian and $\mathbf{x}^A, \mathbf{x}^B$ are the states of subsystems A and B, respectively. For a given control system, the time-varying reachability Gramian on the interval $[t_0, t_1]$ is defined as

$$\mathbf{W}_r(t_0, t_1) \triangleq \int_{t_0}^{t_1} \Phi(t_1, \tau) \mathbf{B}(\tau) \mathbf{B}^T(\tau) \Phi^T(t_1, \tau) d\tau$$

where Φ denotes the state-transition matrix of the intrinsic dynamics of subsystem A without any input (i.e., $\delta \mathbf{x}^A(t) = \Phi(t, t_0) \delta \mathbf{x}^A(t_0)$) [107]. Differentiating with respect to t , we obtain $\delta \dot{\mathbf{x}}^A(t) = \frac{\partial}{\partial t} \Phi(t, t_0) \delta \mathbf{x}^A(t_0)$. Noting also that, in the absence of input from subsystem B, we have

$$\delta \dot{\mathbf{x}}^A(t) = \mathbf{A}(t) \delta \mathbf{x}^A(t) = \mathbf{A}(t) \Phi(t, t_0) \delta \mathbf{x}^A(t_0)$$

Thus setting the equations equal to each other and canceling $\delta \mathbf{x}^A(t_0)$ from both sides we obtain

$$\frac{\partial}{\partial t} \Phi(t, t_0) = \mathbf{A}(t) \Phi(t, t_0)$$

Note that $\Phi(t_0, t_0) = \mathbf{I}$. Now, letting $\Gamma(t, \tau) = \Phi(t, \tau) \mathbf{B}(\tau) \mathbf{B}^T(\tau) \Phi^T(t, \tau)$ and differentiating the Gramian expression with respect to the second argument (and using the Leibniz integral rule) yields

$$\begin{aligned} \frac{\partial}{\partial t} \mathbf{W}_r(t_0, t) &= \frac{\partial}{\partial t} \int_{t_0}^t \Gamma(t, \tau) d\tau \\ &= \Gamma(t, t) \frac{\partial}{\partial t}(t) - \Gamma(t, t_0) \frac{\partial}{\partial t}(t_0) + \int_{t_0}^t \frac{\partial}{\partial t} \Gamma(t, \tau) d\tau \\ &= \mathbf{I} \mathbf{B}(t) \mathbf{B}^T(t) \mathbf{I}(1) - 0 + \int_{t_0}^t \frac{\partial}{\partial t} \Gamma(t, \tau) d\tau \end{aligned}$$

and observing

$$\begin{aligned} \frac{\partial}{\partial t} \Gamma(t, \tau) &= \left(\frac{\partial}{\partial t} \Phi(t, t_0) \right) \mathbf{B}(\tau) \mathbf{B}^T(\tau) \Phi^T(t, \tau) + \Phi(t, t_0) \mathbf{B}(\tau) \mathbf{B}^T(\tau) \left(\frac{\partial}{\partial t} \Phi^T(t, \tau) \right)^T \\ &= \mathbf{A}(t) \Phi(t, t_0) \mathbf{B}(\tau) \mathbf{B}^T(\tau) \Phi^T(t, \tau) + \Phi(t, t_0) \mathbf{B}(\tau) \mathbf{B}^T(\tau) \Phi^T(t, \tau) \mathbf{A}^T(t) \\ &= \mathbf{A}(t) \Gamma(t, \tau) + \Gamma(t, \tau) \mathbf{A}^T(t) \end{aligned}$$

we can continue

$$\begin{aligned} \frac{\partial}{\partial t} \mathbf{W}_r(t_0, t) &= \mathbf{B}(t) \mathbf{B}^T(t) + \mathbf{A}(t) \int_{t_0}^t \Gamma(t, \tau) d\tau + \left(\int_{t_0}^t \Gamma(t, \tau) d\tau \right) \mathbf{A}^T(t) \\ &= \mathbf{B}(t) \mathbf{B}^T(t) + \mathbf{A}(t) \mathbf{W}_r(t_0, t) + \mathbf{W}_r(t_0, t) \mathbf{A}^T(t) \end{aligned}$$

which illustrates that the reachability Gramian can be solved for using an ODE integrator (with initial condition $\mathbf{W}_r(t_0, t_0) = \mathbf{0}$) [56]. In practice, to compute the reachability Gramians using the trained JacobianODE models, we fit a cubic spline $\mathbf{c}(t)$ to the reference trajectory $\mathbf{x}(t)$ and compute $\mathbf{J}(t) = \mathbf{J}(\mathbf{c}(t))$. We can then parse the Jacobian matrix into its component submatrices and compute the Gramians accordingly.

The reachability Gramian is a symmetric positive semidefinite matrix [107]. The optimal cost of driving the system from state $\delta \mathbf{x}_0$ to state $\delta \mathbf{x}_1$ on time interval $[t_0, t_1]$ can be computed as

$$(\delta \mathbf{x}_1 - \Phi(t_1, t_0) \delta \mathbf{x}_0)^T \mathbf{W}_r^{-1}(t_0, t_1) (\delta \mathbf{x}_1 - \Phi(t_1, t_0) \delta \mathbf{x}_0)$$

Thus each eigenvalue of $\mathbf{W}_r$ reflects the ease of control along the corresponding eigenvector [107, 109]. Eigenvectors with larger corresponding eigenvalues will have smaller inverse eigenvalues, and thus scale the cost down along those directions when computing the cost as above.

B.2 Extension to non-autonomous dynamics

While we deal with autonomous systems in this work, we note that this construction can easily be extended to include non-autonomous dynamical systems, of the form $\dot{\mathbf{x}}(t) = \mathbf{f}(\mathbf{x}, t)$ by constructing an augmented state $\tilde{\mathbf{x}} \in \mathbb{R}^{n+1}$ which is simply the concatenation of $\mathbf{x}$ with t . This yields the autonomous dynamics $\dot{\tilde{\mathbf{x}}} = \tilde{\mathbf{f}}(\tilde{\mathbf{x}})$, where $\tilde{\mathbf{f}} : \mathbb{R}^{n+1} \rightarrow \mathbb{R}^{n+1}$ is the concatenation of $\mathbf{f}$ with a function that maps all states to 1.

B.3 Extension to more than two subsystems

Consider a nonlinear dynamical system in $\mathbb{R}^n$, defined by $\dot{\mathbf{x}}(t) = \mathbf{f}(\mathbf{x}(t))$. Suppose now that the system is composed of K subsystems, where the time evolution of subsystem k is given by $\mathbf{x}^{(k)}(t) \in \mathbb{R}^{n_k}$. $\mathbf{x}(t)$ is then comprised of a concatenation of the $\mathbf{x}^{(k)}(t)$, with $\sum_{k=1}^K n_k = n$. Given a particular reference trajectory, $\mathbf{x}(t)$, with $\mathbf{J}$ as the Jacobian of $\mathbf{f}$, then the tangent space dynamics around the reference trajectory for subsystem $\alpha \in \{1, \dots, K\}$ are given by

$$\delta \dot{\mathbf{x}}^{(\alpha)}(t) = \sum_{k=1}^K \mathbf{J}^{k \rightarrow \alpha}(\mathbf{x}(t)) \delta \mathbf{x}^{(k)}(t)$$

where $\mathbf{J}^{k \rightarrow \alpha}(\mathbf{x}(t)) \in \mathbb{R}^{n_\alpha \times n_k}$ is the submatrix of $\mathbf{J}(\mathbf{x}(t))$ in which the columns correspond to subsystem k and the rows correspond to subsystem α . Now, for a given subsystem $\beta \in \{1, \dots, K\}$ with $\beta \neq \alpha$, we wish to analyze the ease with which β can control α locally around the reference trajectory, without intervention from other subsystems. Discounting the interventions from other subsystems equates to setting $\delta \mathbf{x}^{(k)}(t) = 0$ for $k \neq \alpha, \beta$, leaving the expression

$$\delta \dot{\mathbf{x}}^{(\alpha)}(t) = \mathbf{J}^{\alpha \rightarrow \alpha}(\mathbf{x}(t)) \delta \mathbf{x}^{(\alpha)}(t) + \mathbf{J}^{\beta \rightarrow \alpha}(\mathbf{x}(t)) \delta \mathbf{x}^{(\beta)}(t)$$

which quantifies the influence β can exert over α in the absence of perturbations from any other subsystem. Using this representation, the reachability Gramian can be computed as described above. Performing this procedure for all pairs of subsystems $\alpha, \beta \in \{1, \dots, K\}$ thus characterizes all pairwise control relationships between subsystems along the reference trajectory.

C Supplementary results

C.1 Control and communication capture different phenomena

In the Introduction (Section 1) we note that measuring communication between two systems is different than measuring control. To illustrate this, consider two brain areas, A and B, whose dynamics are internally linear. The areas are connected only through a linear feedforward interaction term from area B to area A (Figure S1A). Concretely, we consider the dynamics:

$$\begin{aligned} \dot{\mathbf{x}}^A(t) &= \mathbf{J}^{A \rightarrow A} \mathbf{x}^A(t) + \mathbf{J}^{B \rightarrow A} \mathbf{x}^B(t) \\ \dot{\mathbf{x}}^B(t) &= \mathbf{J}^{B \rightarrow B} \mathbf{x}^B(t) \end{aligned}$$

Here, the $\mathbf{J}$ matrices are time-invariant. When considering communication between brain areas, one might aim to find the subspaces in which communication between B and A occurs, as well as the messages passed [33, 73]. In this construction, the subspace in which area B communicates with area A is explicitly given by $\mathbf{J}^{B \rightarrow A}$, and thus the message, or input, from area B to area A at any time t is simply given by $\mathbf{J}^{B \rightarrow A} \mathbf{x}^B(t)$. We pick the dimensions of each region to be $n_A = n_B = 4$ and let the eigenvectors of the (negative definite matrix) $\mathbf{J}^{A \rightarrow A}$ be given by $\mathbf{v}_i, i = 1, 2, 3, 4$. The eigenvectors are numbered in order of decreasing real part of their corresponding eigenvalue. Now, suppose we construct the interaction matrix $\mathbf{J}^{B \rightarrow A}$ in two different ways: If we let (1) $\mathbf{J}_1^{B \rightarrow A} = \mathbf{v}_1 \mathbf{1}^T$ (Figure S1B, left), then the signal from B is projected onto the most stable mode of region A, whereas if we let (2) $\mathbf{J}_2^{B \rightarrow A} = \mathbf{v}_4 \mathbf{1}^T$ (Figure S1B, right), the signal is projected onto the least stable mode. While interaction 1 and interaction 2 communicated messages with identical magnitudes, interaction 2 led to much lower cost reachability control (Figure S1C). This illustrates that control depends not only on the directions along which the areas can communicate, but also on how *aligned* the communication is with the target area's dynamics.

C.2 Derivative estimation is not implied by function estimation

An alternative approach to directly estimating the Jacobians would be to learn an approximation $\hat{\mathbf{f}}$ of the function $\mathbf{f}$, and then approximate the Jacobian via automatic differentiation (i.e., an estimate $\hat{\mathbf{J}} = \frac{\partial}{\partial \mathbf{x}} \hat{\mathbf{f}}$, as with the NeuralODEs). While this can be effective in certain scenarios, it is not generally the case that approximating a function well will yield a good approximation of its derivative. To illustrate this, we recall an example from Latrémoière et al. [90], in which functions

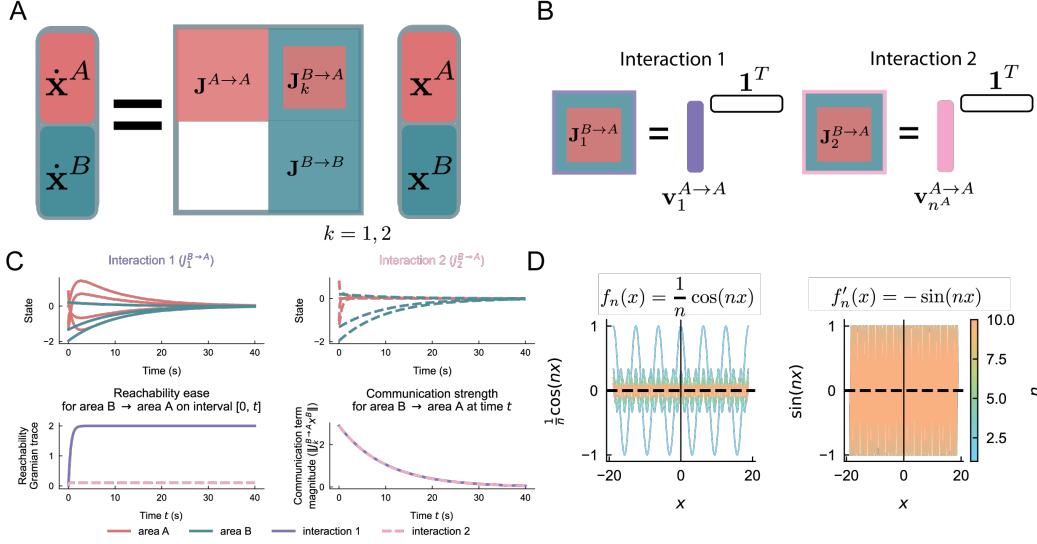


Figure S1: Communication versus control in linear systems and the challenge of Jacobian estimation. (A) Setup of two linearly connected brain areas, A and B. (B) Interaction matrices projecting signals from area B onto either the most stable (Interaction 1) or least stable (Interaction 2) eigenvectors of area A. (C) Although both interactions communicate identical signal magnitudes, Interaction 2 provides significantly enhanced reachability due to alignment with unstable modes of area A dynamics. (D) An illustrative example demonstrating that accurate approximation of a function (left panel) does not guarantee accurate approximation of its derivative (right panel), emphasizing the necessity of directly estimating the Jacobian.

$f_n = \frac{1}{n} \cos(nx)$ are used to approximate the function $f(x) = 0$ (Figure S1D). While these approximations improve with increasing n (i.e., $\lim_{n \rightarrow \infty} f_n = f$), this is not the case for the derivative ($\lim_{n \rightarrow \infty} f'_n = -\sin(nx) \neq f'$). This demonstrates the necessity of learning $\mathbf{J}$ directly (rather than first approximating $\mathbf{f}$), which we also demonstrate empirically. In the context of machine learning, this setting could be interpreted as overfitting [90]. As long as function estimates match at the specific points in the training set, how the function fluctuates between these points is not constrained to match the true function. For this reason, we included the Frobenius norm Jacobian regularization in our implementation of the NeuralODEs (Appendix D.3).

C.3 Full benchmark dynamical systems results

We here present the full results for all considered example dynamical systems. 10 time-step trajectory predictions along with Jacobian estimation and estimated Lyapunov spectra are displayed in Figure S2, along with MSE and R^2 in Table S1 (trajectory prediction) and Table S2 (Jacobian estimation). Jacobian estimation errors using the 2-norm are presented in Table S3.

C.4 Ablation studies

To determine the value of the different components of the JacobianODE learning framework, we performed several ablation studies. We chose to evaluate on the Lorenz system and the task-trained RNNs, as these together provide two common settings of chaos and stability, as well as low- and high-dimensional dynamics.

Ablating the Jacobian-parameterized ODEs The JacobianODE framework constructs an estimate of the initial time derivative $\mathbf{f}$ via a double integral of the Jacobian as described in Section 3.2. To briefly recall, Jacobian path integration is described as

$$\mathbf{f}(\mathbf{x}(t_f)) - \mathbf{f}(\mathbf{x}(t_i)) = \int_C \mathbf{J} ds = \int_{t_i}^{t_f} \mathbf{J}(\mathbf{c}(r)) \mathbf{c}'(r) dr, \quad (14)$$

When we do not have access to the initial derivative and base point $\mathbf{f}(\mathbf{x}(t_i)), \mathbf{x}(t_i)$, we use the following formulation

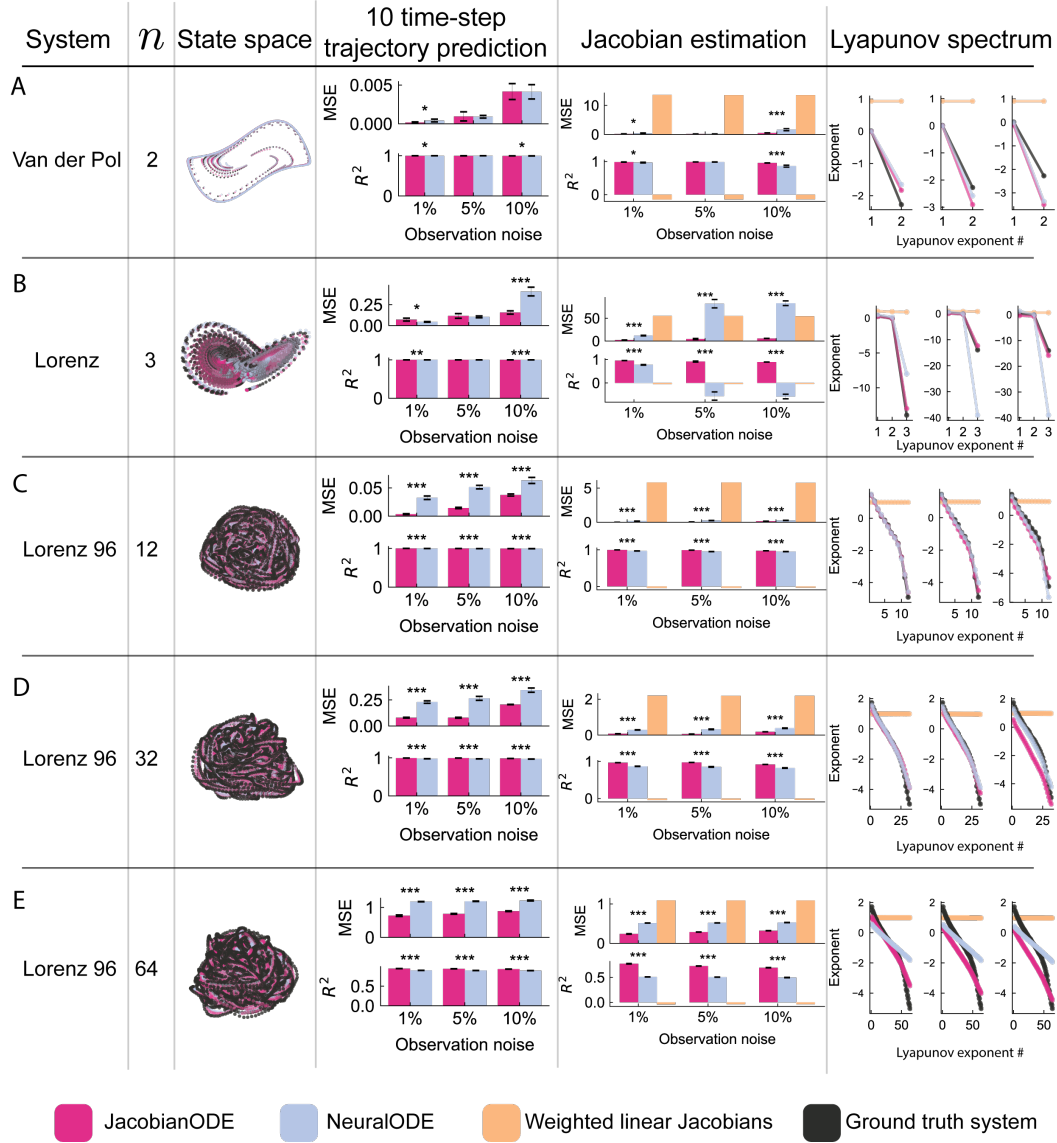


Figure S2: Full dynamical systems prediction results. State space representations (using spectral embedding for systems that have more than 3 dimensions), 10 time-step trajectory predictions, Jacobian estimation, and Lyapunov spectrum estimates for each of (A) the Van der Pol oscillator, (B) the Lorenz system, and the Lorenz 96 system with (C) 12, (D) 32, and (E) 64 dimensions.

Table S1: Trajectory prediction metrics (MSE and R^2) for each system and training noise. Errors are mean $\pm$ standard deviation across five random initializations of the model architectures.

Project	Training noise	MSE		R^2	
		JacobianODE	NeuralODE	JacobianODE	NeuralODE
VanDerPol (2 dim)	1%	0.00017 $\pm$ 0.00007	0.0004 $\pm$ 0.0002	0.99987 $\pm$ 0.00005	0.9996 $\pm$ 0.0002
	5%	0.0010 $\pm$ 0.0006	0.0009 $\pm$ 0.0002	0.9995 $\pm$ 0.0003	0.99949 $\pm$ 0.00009
	10%	0.004 $\pm$ 0.001	0.0041 $\pm$ 0.0009	0.9979 $\pm$ 0.0005	0.9969 $\pm$ 0.0003
Lorenz (3 dim)	1%	0.07 $\pm$ 0.02	0.045 $\pm$ 0.004	0.99979 $\pm$ 0.00004	0.99986 $\pm$ 0.00001
	5%	0.12 $\pm$ 0.03	0.10 $\pm$ 0.01	0.9996 $\pm$ 0.0001	0.99958 $\pm$ 0.00004
	10%	0.16 $\pm$ 0.02	0.41 $\pm$ 0.05	0.99928 $\pm$ 0.00008	0.9981 $\pm$ 0.0003
Lorenz 96 (12 dim)	1%	0.0035 $\pm$ 0.0007	0.033 $\pm$ 0.003	0.99969 $\pm$ 0.00007	0.9969 $\pm$ 0.0003
	5%	0.014 $\pm$ 0.001	0.051 $\pm$ 0.003	0.9985 $\pm$ 0.0001	0.9955 $\pm$ 0.0003
	10%	0.038 $\pm$ 0.002	0.063 $\pm$ 0.005	0.9960 $\pm$ 0.0002	0.9940 $\pm$ 0.0007
Lorenz 96 (32 dim)	1%	0.079 $\pm$ 0.004	0.23 $\pm$ 0.01	0.9939 $\pm$ 0.0003	0.977 $\pm$ 0.002
	5%	0.079 $\pm$ 0.004	0.26 $\pm$ 0.02	0.9939 $\pm$ 0.0003	0.974 $\pm$ 0.002
	10%	0.205 $\pm$ 0.001	0.34 $\pm$ 0.02	0.9837 $\pm$ 0.0001	0.968 $\pm$ 0.002
Lorenz 96 (64 dim)	1%	0.72 $\pm$ 0.03	1.194 $\pm$ 0.005	0.946 $\pm$ 0.002	0.8997 $\pm$ 0.0007
	5%	0.79 $\pm$ 0.01	1.205 $\pm$ 0.007	0.9400 $\pm$ 0.0010	0.8986 $\pm$ 0.0010
	10%	0.87 $\pm$ 0.01	1.23 $\pm$ 0.01	0.9332 $\pm$ 0.0009	0.896 $\pm$ 0.001
Task-trained RNN	1%	0.0042 $\pm$ 0.0006	0.00558 $\pm$ 0.00002	0.9981 $\pm$ 0.0003	0.99762 $\pm$ 0.00001
	5%	0.0043 $\pm$ 0.0006	0.00596 $\pm$ 0.00005	0.9981 $\pm$ 0.0002	0.99747 $\pm$ 0.00002
	10%	0.0089 $\pm$ 0.0001	0.00655 $\pm$ 0.00003	0.99591 $\pm$ 0.00004	0.99721 $\pm$ 0.00001

Table S2: Jacobian estimation metrics (MSE and R^2) for each system and training noise. Errors are mean $\pm$ standard deviation across five random initializations of the model architectures.

Project	Training noise	MSE		R^2	
		JacobianODE	NeuralODE	JacobianODE	NeuralODE
VanDerPol (2 dim)	1%	0.18 $\pm$ 0.06	0.4 $\pm$ 0.1	0.985 $\pm$ 0.005	0.97 $\pm$ 0.01
	5%	0.16 $\pm$ 0.02	0.17 $\pm$ 0.03	0.987 $\pm$ 0.002	0.985 $\pm$ 0.002
	10%	0.50 $\pm$ 0.04	1.6 $\pm$ 0.3	0.958 $\pm$ 0.003	0.86 $\pm$ 0.03
Lorenz (3 dim)	1%	2.4 $\pm$ 0.1	12.0 $\pm$ 0.7	0.954 $\pm$ 0.002	0.77 $\pm$ 0.01
	5%	4.6 $\pm$ 1.2	82.0 $\pm$ 9.2	0.91 $\pm$ 0.02	-0.6 $\pm$ 0.2
	10%	5.7 $\pm$ 0.2	83.0 $\pm$ 5.2	0.891 $\pm$ 0.004	-0.58 $\pm$ 0.10
Lorenz 96 (12 dim)	1%	0.012 $\pm$ 0.003	0.18 $\pm$ 0.01	0.9979 $\pm$ 0.0006	0.969 $\pm$ 0.002
	5%	0.053 $\pm$ 0.006	0.27 $\pm$ 0.02	0.991 $\pm$ 0.001	0.953 $\pm$ 0.003
	10%	0.153 $\pm$ 0.008	0.27 $\pm$ 0.01	0.973 $\pm$ 0.001	0.951 $\pm$ 0.002
Lorenz 96 (32 dim)	1%	0.075 $\pm$ 0.004	0.28 $\pm$ 0.01	0.965 $\pm$ 0.002	0.866 $\pm$ 0.006
	5%	0.060 $\pm$ 0.003	0.32 $\pm$ 0.02	0.972 $\pm$ 0.001	0.850 $\pm$ 0.009
	10%	0.180 $\pm$ 0.003	0.38 $\pm$ 0.02	0.915 $\pm$ 0.001	0.820 $\pm$ 0.008
Lorenz 96 (64 dim)	1%	0.238 $\pm$ 0.007	0.515 $\pm$ 0.002	0.773 $\pm$ 0.007	0.509 $\pm$ 0.002
	5%	0.286 $\pm$ 0.003	0.519 $\pm$ 0.003	0.727 $\pm$ 0.003	0.505 $\pm$ 0.003
	10%	0.321 $\pm$ 0.004	0.527 $\pm$ 0.005	0.694 $\pm$ 0.003	0.497 $\pm$ 0.005
Task-trained RNN	1%	2.2 $\pm$ 0.2	5.2777 $\pm$ 0.0009	0.59 $\pm$ 0.03	-0.0008 $\pm$ 0.0002
	5%	1.70 $\pm$ 0.08	5.2900 $\pm$ 0.0001	0.68 $\pm$ 0.01	-0.00315 $\pm$ 0.00002
	10%	1.99 $\pm$ 0.01	5.2889 $\pm$ 0.0001	0.623 $\pm$ 0.002	-0.00295 $\pm$ 0.00002

$$\mathbf{f}(\mathbf{x}(t)) = \frac{G(t_0, t; \mathbf{J}, \mathbf{c}) + \mathbf{x}(t) - \mathbf{x}(t_0)}{t - t_0}, \quad (15)$$

where

$$G(t_0, t; \mathbf{J}, \mathbf{c}) = \int_{t_0}^t \int_s^t \mathbf{J}(\mathbf{c}_{s,t}(r)) \mathbf{c}'_{s,t}(r) dr ds \quad (16)$$

Alternatively, one could use Equation 14 and learn $\mathbf{f}(\mathbf{x}(t_i))$ and $\mathbf{x}(t_i)$ as learnable parameters ($\mathbf{x}(t_i)$ is needed as the first point of the path, i.e. $\mathbf{c}(t_i) = \mathbf{x}(t_i)$ inside the integral) [92]. Here, the path integral

Table S3: Mean 2-norm error on Jacobian estimation, $\langle \|\mathbf{J} - \hat{\mathbf{J}}\|_2 \rangle$, for each system and noise level. Errors are reported as mean $\pm$ standard deviation, with mean and standard deviation computed over five random initializations of the model architectures.

Project	Training noise	JacobianODE	NeuralODE	Weighted Linear
VanDerPol (2 dim)	1%	0.7 ± 0.1	1.0 ± 0.3	6.05
	5%	0.71 ± 0.05	0.71 ± 0.08	6.03
	10%	1.31 ± 0.05	2.2 ± 0.2	6.02
Lorenz (3 dim)	1%	3.2 ± 0.2	8.6 ± 0.3	17.06
	5%	4.9 ± 0.9	25.9 ± 1.5	17.02
	10%	6.0 ± 0.1	26.4 ± 0.9	16.95
Lorenz 96 (12 dim)	1%	0.8 ± 0.1	3.5 ± 0.2	14.94
	5%	1.75 ± 0.09	4.1 ± 0.1	14.93
	10%	2.91 ± 0.09	4.2 ± 0.1	14.92
Lorenz 96 (32 dim)	1%	3.75 ± 0.08	7.8 ± 0.1	16.29
	5%	3.10 ± 0.09	8.1 ± 0.2	16.26
	10%	4.89 ± 0.05	8.6 ± 0.1	16.28
Lorenz 96 (64 dim)	1%	8.4 ± 0.1	12.63 ± 0.02	16.84
	5%	9.04 ± 0.07	12.66 ± 0.03	16.81
	10%	9.42 ± 0.06	12.72 ± 0.05	16.82
Task-trained RNN	1%	30.4 ± 0.4	38.565 ± 0.006	39.29
	5%	29.4 ± 0.9	38.5611 ± 0.0004	38.91
	10%	36.0 ± 0.1	38.5600 ± 0.0004	38.70

between $\mathbf{x}(t_i)$ and $\mathbf{x}(t_f)$ is a linear interpolation, as before. These ablated models were trained on 10 time-step prediction, as with the original JacobianODEs. For these models, we completed a full sweep over the loop closure loss weight λ_{loop} in order to determine the best hyperparameter.

Ablating teacher forcing We trained models without any teacher-forcing. That is, models were able to generate only one-step predictions, without any recursive predictions. Again we did a full hyperparameter sweep to pick λ_{loop} .

Ablating loop closure loss We ablated the loop closure loss in two ways. The first was to set $\lambda_{\text{loop}} = 0$ to illustrate what would happen if there were no constraints placed on the learned Jacobians. The second was to instead use the Jacobian Frobenius norm regularization that was used for the NeuralODEs (details are in Appendix D.3). We did a full sweep to pick λ_{jac} , the Frobenius norm regularization weight.

Table S4: Mean Frobenius norm error $\langle \|\mathbf{J} - \hat{\mathbf{J}}\|_F \rangle$ for different model ablations on the Lorenz system with 10% observation noise. Errors are reported as mean $\pm$ standard deviation, with statistics computed over 8 test trajectories, each consisting of 1200 points.

Model variant	Frobenius norm error
JacobianODE (original)	6.46 ± 1.50
With learned base derivative point	7.87 ± 1.14
No teacher forcing	11.59 ± 1.30
No loop closure	58.22 ± 2.00
With Jacobian penalty instead of loop closure	9.59 ± 2.45

Ablation results The performance of the ablated models on Jacobian estimation in the Lorenz system are presented in Table S4. The original JacobianODE outperforms all ablated models, indicating that all components of the JacobianODE training framework improve the model’s performance in this setting. Ablating the Jacobian-parameterized initial derivative estimate resulted in a slight decrease in the estimation loss. This is potentially because the network could offload some of the

responsibility for generating correct trajectory predictions onto the estimated base point $\mathbf{x}(t_i)$ and derivative estimate $\hat{\mathbf{f}}(\mathbf{x}(t_i))$, slightly reducing the necessity of estimating correct Jacobians. Ablating the teacher forcing annealing predictably led to worse Jacobian estimation, as the network no longer has to consider how errors will propagate along the trajectory. The most dramatic increase in error was with the ablation of the loop closure loss. Without this important regularization, the learned Jacobians reproduced the dynamics but were not constrained to be conservative, resulting in poor Jacobian estimation. The inclusion of the Frobenius penalty on the Jacobians mitigated this, although it did not encourage accurate Jacobian estimation to the same degree as the loop closure loss.

Table S5: Mean Frobenius norm error $\langle \|\mathbf{J} - \hat{\mathbf{J}}\|_F \rangle$ for different model ablations on the task-trained RNN with 10% observation noise. Errors are reported as mean $\pm$ standard deviation, with statistics computed over 409 test trajectories, each consisting of the 49 points from the second delay and response epochs.

Model variant	Frobenius norm error
JacobianODE (original)	180.13 ± 1.55
With learned base derivative point	186.28 ± 1.17
No teacher forcing	187.72 ± 1.63
No loop closure	313.31 ± 29.88
With Jacobian penalty instead of loop closure	163.69 ± 4.22

We then tested the ablated the models on Jacobian estimation in the task-trained RNN, with results presented in Table S5. Again, ablating the Jacobian-parameterized derivative estimates, teacher forcing, and loop closure resulted in worse Jacobian estimation. Interestingly, in this setting, the inclusion of a penalty on the Frobenius norm of the Jacobians outperformed the use of the loop closure loss. This could potentially be because the loop closure loss is more difficult to drive to zero in high dimensional systems, or because the loop closure loss is more important in chaotic systems like the Lorenz system considered above. Future work should consider in what contexts each kind of regularization is most beneficial to JacobianODE models.

C.5 NeuralODEs achieve improved performance at the cost of increased inference time

In the main paper, we implemented both the JacobianODEs and the NeuralODEs as four-layer MLPs, with the four layers having sizes of 256, 1024, 2048, and 2048 respectively. This was done for the fairest architectural comparison between the models, to ensure that both models had the same representational capacity when generating their respective outputs. However, there are many architectural changes that we could make to this setup that impact performance. We hypothesized based on the discussion in Appendix C.2 that increasing the hidden layer size of the NeuralODEs would improve Jacobian estimation, as larger models have been known to learn smoother representations. Furthermore, we wondered whether including residual blocks in place of the standard MLP implementation would improve Jacobian estimation.

To test this, we implemented the NeuralODEs as four-layer residual networks and tested three different sizes of hidden layer: 1024, 2048, and 4096. Results are in Table S6. For nearly all models, these changes yielded only marginal improvements over the original NeuralODE model. Only the model with 4096-dimensional hidden layers under 10% training noise achieves a performance near the *original* JacobianODEs. However, the performance limit is still below that of the JacobianODEs, even with a large increase in the representational capacity of the model. It is furthermore of note that the NeuralODEs were able to significantly improve performance only in the high noise setting. This suggests that high noise is necessary for the model to be forced to learn the response of the system to perturbation. In contrast, JacobianODEs perform similarly across all noise levels, indicating a stronger inductive bias to learn the response of the system to perturbation.

While the increased hidden layer size seems to induce smoother hidden representations (as expected), it comes with increasing computational cost. Recall that the NeuralODE models directly parameterize the dynamics as $\dot{\mathbf{x}} = \hat{\mathbf{f}}^\theta(\mathbf{x})$. The Jacobian at state $\mathbf{x}$ is computed by directly backpropagating through the Neural ODE $\hat{\mathbf{f}}^\theta$, thereby significantly increasing compute. On the other hand, the JacobianODEs only require a forward pass. We evaluated the Jacobian inference time of JacobianODE

Table S6: Mean Frobenius norm error $\langle \|\mathbf{J} - \hat{\mathbf{J}}\|_F \rangle$ for different model types on the task-trained RNN data across observation noise levels. Errors are reported as mean $\pm$ standard deviation, with statistics computed over 409 test trajectories.

Model type	Training noise 5%	Training noise 10%	Learnable parameters
JacobianODE (original)	174.1 $\pm$ 2.4	180.1 $\pm$ 1.5	4.02e+07
NeuralODE (original)	294.0 $\pm$ 2.4	293.9 $\pm$ 2.4	6.85e+06
NeuralODE (1024 dim. hidden layers)	291.6 $\pm$ 2.3	289.8 $\pm$ 2.3	3.41e+06
NeuralODE (2048 dim. hidden layers)	288.5 $\pm$ 2.3	275.0 $\pm$ 2.7	1.31e+07
NeuralODE (4096 dim. hidden layers)	275.8 $\pm$ 2.4	184.2 $\pm$ 5.6	5.14e+07

and NeuralODE models with four hidden layers of the same size on an H100 GPU. Each model was timed on inferring the Jacobians of 100 batches of the 128-dimensional task-trained RNN, with each batch containing 16 sequences of length 25. Timings were repeated ten times for each model. Results are plotted in Figure S3.

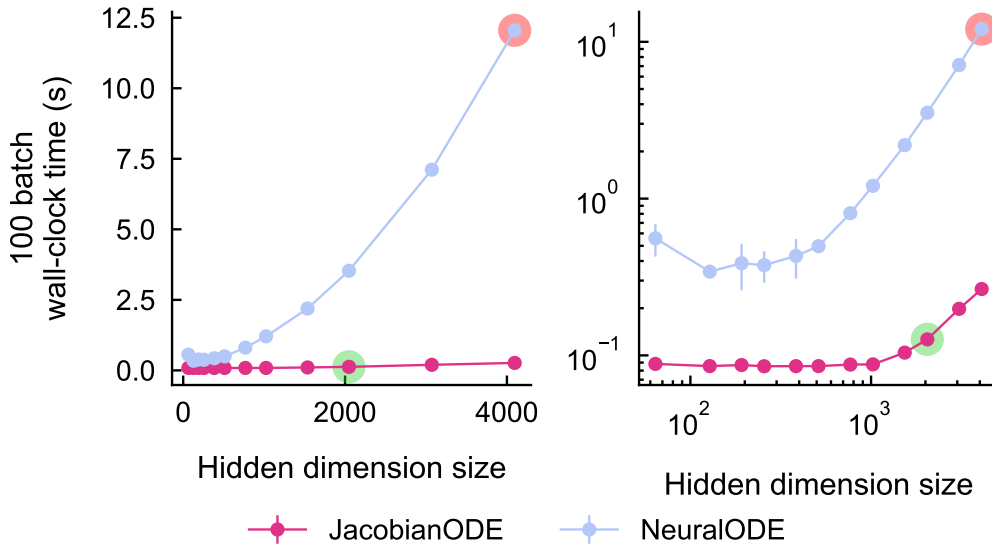


Figure S3: **JacobianODEs achieve highly efficient Jacobian inference.** Jacobian inference times computed over ten repetitions of 100 batches (error bars indicate mean $\pm$ standard deviation). Red circles indicate the inference time corresponding to the largest hidden layer dimension of the highest performing NeuralODE model in Table S6. Green circles indicate the inference time corresponding to the largest hidden layer dimension of the highest performing JacobianODE model in Table S6. Each plot illustrates the same data but with different x and y scaling.

The JacobianODE achieves much faster inference times than the NeuralODE – approximately two orders of magnitude faster at large hidden dimension sizes. Furthermore, as shown in Table S6, the JacobianODE with a maximum hidden layer size of 2048 outperforms the NeuralODE with a maximum hidden layer size of 4096 on Jacobian estimation, and does so with orders of magnitude faster inference (Figure S3, green and red circles, Table S6). This suggests that while both architectures appear to have inference times that scale approximately exponentially, the JacobianODE achieves more favorable scaling across every hidden layer size we tested. Our analysis therefore illustrates that it is possible to improve the NeuralODEs’ Jacobian estimation with larger models, but the inference time scaling renders these models ill-equipped for important settings such as real-time control and the analysis of high-volume neural data.

C.6 Comparison to echo state networks

Echo state networks (ESNs, or reservoir networks) are tools for time-series forecasting and analysis. ESNs are typically implemented as RNNs with fixed, sparse, hidden connectivity. The output weights

of the ESNs can then be trained via linear regression to reproduce specific temporal patterns. ESNs have been shown to be able to forecast the Lorenz 63 system for a particular choice of hyperparameters [139]. We note that Pathak et al. [139] implements the same method as Pathak et al. [140]. We implemented the ESNs exact hyperparameters reported in Pathak et. al. 2017 for the Lorenz system, which involved a reservoir of 300 nodes. Our implementation of the Lorenz system was simulated with a sampling interval of $\tilde{0.015}$ s, very similar to the 0.02 s used in the paper. Thus, for consistency of comparison with JacobianODEs, we fit ESN models to data sampled every 0.015 s.

To enable Jacobian estimation, we changed one small implementation detail from Pathak et al. [139]. Specifically, we trained the network output weights to predict the third variable z using only the network state $\mathbf{r}(t)$, as opposed to the vector $\tilde{\mathbf{r}}(t)$ for which the node state is squared ($\tilde{r}_i(t) = r_i^2(t)$) for half of the nodes i . This is because squaring the reservoir state complicates Jacobian computation, and was stated by Pathak et al. [139] and Lu et al. [141] to only be necessary for symmetry breaking if only the z variable of the Lorenz system were provided. We trained the ESNs on the same volume of data as the JacobianODEs and NeuralODEs, with 1% observation noise. We swept the regularization parameter β to ensure a good fit. We furthermore generated predictions using a 150 time-step lead-up to ensure the ESNs converged to a good internal representation. The ESNs predicted the Lorenz system well, achieving an MSE of 0.03 on the 10-step prediction task from our paper (on test data). This was very similar to the reported 10-step prediction MSEs of 0.07 and 0.04 for the JacobianODEs and NeuralODEs, respectively.

In Pathak et al. [139], the Lyapunov exponents of the system were computed using the Jacobian of the *interior* dynamics of the ESN, as the recursion necessary to compute the Jacobian can be formulated uniquely only in regard to this state representation. Thus, each of these Jacobians would be a 300×300 matrix. While this suffices to compute the Lyapunov exponents, it does not suffice for our purposes of obtaining Jacobian estimates of the dynamics of the *original* system. For this purpose, it has been shown that it is possible to obtain an estimate of the Jacobian of the original dynamical system (i.e., not the echo state space) using ESNs [142]. However, as Banerjee et al. [142] noted, it is necessary to invert the non-square ESN output matrix in order to obtain these Jacobians, thus they are non-uniquely determined, and are used in the paper "only for causality estimation purposes". Nevertheless, we computed ESN-estimated Jacobians on the 1% observation noise test data. We used a large (7,200 time-step) lead-up and ensured coherence to the trajectory by feeding in the true state as input. The Frobenius norm error over five random seeds was 85.8 ± 3.1 , much larger than the reported errors of 3.3 ± 0.2 and 8.7 ± 0.3 for the JacobianODEs and NeuralODEs respectively. Thus, similar to the NeuralODE, while ESNs can achieve comparable prediction on chaotic systems, they do so without enabling an accurate reconstruction of the Jacobian of the original dynamical system.

C.7 Analyzing the dimensionality of Jacobian manifolds

RNN activity is often confined to low-dimensional manifolds, potentially simplifying Jacobian estimation. To address the extent to which this influences our results, we performed the following analyses.

Jacobian rank: We first note that both the RNN hidden weight matrix and nearly all Jacobians had full rank (128; $\sim 0.1\%$ had rank 127). To assess the *effective* rank, we computed the participation ratio (PR), a singular value-based measure of intrinsic dimensionality. The RNN weight matrix had a PR of 123, and the Jacobians had a mean PR of ~ 106 (std 0.7).

Intrinsic-data dimensionality: Another possibility is that the data itself may be low-dimensional. Since the Jacobian is a function of the trajectory data, the Jacobian manifold’s intrinsic dimension is at most that of the trajectory manifold. To estimate the intrinsic dimensions, we pooled time points from a large sample of trajectories. For RNN state vectors, we computed PR directly; for Jacobians, we flattened them into n^2 -dimensional vectors before computing PR. We found the RNN trajectory PR to be ~ 13.5 , and the Jacobian PR to be ~ 6 . Given the 128D extrinsic space, this confirms both RNN states and Jacobians lie on lower-dimensional manifolds.

Notably, prior work estimated the dimensionality of prefrontal cortex delay activity at 24 using ~ 4000 neurons, falling to ~ 6 when subsampled to 100 neurons [143]. These values were interpreted as high-dimensional. Thus, the dimensionalities we observe are consistent with biological systems performing complex tasks.

Furthermore, in the 64-dimensional Lorenz 96 system, the trajectory PR was ~ 47.5 , and the Jacobians had a PR of ~ 50 , which we believe to be high-dimensional enough to demonstrate the effectiveness of our method on complex systems.

We note that in larger systems, approximating the Jacobian as low rank may be more efficient. Rather than computing each of the outputs directly, we can constrain the Jacobian to be low-rank by designing the output as a low-rank combination of vectors (e.g., it learns to output a low-rank singular value decomposition).

D Experimental details

D.1 Dynamical systems data

We used the following dynamical systems for the testing and validation of Jacobian estimation.

Van der Pol oscillator We implement the classic Van der Pol oscillator [96]. The system is governed by the following equations

$$\begin{aligned}\dot{x} &= y \\ \dot{y} &= \mu(1 - x^2)y - x\end{aligned}$$

We pick $\mu = 2$ in our implementation.

Lorenz We implement the Lorenz system introduced by Lorenz [97] as

$$\begin{aligned}\dot{x} &= \sigma(y - x) \\ \dot{y} &= x(\rho - z) - y \\ \dot{z} &= xy - \beta z\end{aligned}$$

with the typical choices for parameters ($\sigma = 10, \rho = 28, \beta = 8/3$).

Lorenz 96 We implement the Lorenz 96 system introduced by Lorenz [98] and defined by

$$\dot{x}_i = (x_{(i+1) \pmod N} - x_{(i-2) \pmod N})x_{(i-1) \pmod N} - x_i \pmod N + F$$

with $F = 8$ and $N \in \{12, 32, 64\}$.

Simulation We use the `dysts` package to simulate all dynamical systems [99, 100]. The characteristic timescale of their Fourier spectrum τ is selected and the systems are sampled with respect to τ . For all systems, the training data consisted of 26 trajectories of 12 periods, sampled at 100 time steps per τ . The validation data consisted of 6 trajectories of 12 periods sampled at 100 time steps per τ . The test data consisted of 8 trajectories of 12 periods sampled at 100 time steps per τ . Trajectories were initialized using a random normal distribution with standard deviation 0.2. The simulation algorithm used was Radau. Batches were constructed by moving a sliding window along the signal. The sequence length was selected such that the generated predictions would generate 10 novel time points (i.e., 11 time steps for the NeuralODE, and 25 time steps for the JacobianODE, due to the 15 time steps used to estimate the initial time derivative $\mathbf{f}$).

Noise We define $P\%$ observation noise with $P = 100p$ in the following way. Let $A_{\text{signal}} = \mathbb{E} [\|x(t)\|_2^2]$ be the expected squared norm of the signal with $\mathbf{x}(t) \in \mathbb{R}^n$. Then consider a noise signal $\eta(t) \in \mathbb{R}^n$ where each component $\eta_i(t) \sim \mathcal{N}(0, \frac{1}{\sqrt{n}}p\sqrt{A_{\text{signal}}})$. Then

$$\mathbb{E} [\|\eta(t)\|_2^2] = \mathbb{E} \left[\sum_{i=1}^n \eta_i^2(t) \right] = \sum_{i=1}^n \mathbb{E} [\eta_i^2(t)] = \sum_{i=1}^n \frac{1}{n} p^2 A_{\text{signal}} = p^2 A_{\text{signal}}$$

and thus the noise percent is

$$\sqrt{\frac{\mathbb{E} [\|\eta(t)\|_2^2]}{\mathbb{E} [\|x(t)\|_2^2]}} = \sqrt{\frac{p^2 A_{\text{signal}}}{A_{\text{signal}}}} = p$$

D.2 Task-trained RNN

The task used to train the RNN was exactly as defined in Section 5. The hidden dimensionality of the RNN was 128, and the input dimensionality was 10, where the first four dimensions represented the one-hot encoded "upper" color, the second four dimensions represented the one-hot encoded "lower" color, and the last two dimensions represented the one-hot encoded cue. The RNN used for the task had hidden dynamics defined by

$$\begin{aligned}\tau \dot{\mathbf{h}}(t) &= -\mathbf{h} + \mathbf{W}_{hh}\sigma(\mathbf{h}(t)) + \mathbf{W}_{hi}\mathbf{u}(t) + \mathbf{b} \\ \mathbf{o}(t) &= \mathbf{W}_{oh}\mathbf{h}(t)\end{aligned}$$

with $\tau = 50$ ms, which for the purposes of training was discretized with Euler integration with a time step of $\Delta t = 20$ ms. $\mathbf{W}_{hh}$ is the 128x128 dimensional matrix that defines the internal dynamics, $\mathbf{W}_{hi}$ is the 128 x 10 dimensional matrix that maps the input into the hidden state, $\mathbf{W}_{oh}$ is the 4x128 dimensional output matrix that maps the hidden state to a four-dimensional output $\mathbf{o}(t)$, and $\mathbf{b}$ is a static bias term. σ was taken to be an exponential linear unit activation with $\alpha = 1$. The RNN hidden state was split into two "areas" each with 64 dimensions. The input matrix $\mathbf{W}_{hi}$ was masked during training so that inputs could only flow into the first 64 dimensions – the "visual" area. The same procedure was performed for the output matrix $\mathbf{W}_{oh}$, except the mask was such that outputs could stem only from the second group of 64 dimensions – the "cognitive" area. The within-area subblocks of the matrix $\mathbf{W}_{hh}$ were first initialized such that the real part of the eigenvalues were randomly distributed on the interval $[-0.1, 0]$ and the imaginary part of the eigenvalues were randomly distributed on the interval $[0, 2\pi]$. The eigenvectors were random orthonormal matrices. We then computed the matrix exponential of this matrix. The across-area weights were first initialized to be random normal, then divided by the 2-norm of the resulting matrix and multiplied by 0.05. The input and output matrices were initialized as random normal and then scaled by the 2-norm of the resulting matrix. The static bias $\mathbf{b}$ was initialized at $\mathbf{0}$. After initialization, all weights could be altered unimpeded (except for the masks). Notably, the inputs $\mathbf{u}(t)$ were present only during the stimulus and cue presentation epochs – otherwise the network evolved autonomously. The loss was computed via cross entropy loss on the RNN outputs during the response period (the final 250 ms of the trial).

For the training data, we generated 4096 random trials, and used 80% for training and the remainder for validation. The batch size used was 32. Training was performed for 40 epochs. The learning rate was 0.0005. For use with the Jacobian estimation models, data was batched and used for training exactly as was done with the other dynamical systems data (see Appendix D.1). Observation noise was also computed in the same way.

D.3 NeuralODE details

NeuralODE models directly estimate the time derivative $\mathbf{f}$ with a neural-network parameterized function $\hat{\mathbf{f}}^\theta$. Then the Jacobians can be computed as $\hat{\mathbf{J}} = \frac{\partial}{\partial \mathbf{x}} \hat{\mathbf{f}}^\theta$.

The NeuralODEs were implemented as described in Section 4 and ODE integration was done exactly as for the JacobianODE using the `torchdiffeq` package with the RK4 method [102]. To regularize the NeuralODE we implemented a Frobenius norm penalty on the estimated Jacobians, i.e.

$$\mathcal{L}_{\text{jac}} = \lambda_{\text{jac}} \langle \|\hat{\mathbf{J}}(\mathbf{x}(t))\|_F \rangle$$

where $\hat{\mathbf{J}}$ is the estimated Jacobian computed via automatic differentiation and λ_{jac} is a hyperparameter that controls the relative weighting of the Jacobian penalty [104–106]. As mentioned in the main text, this penalty prevents the model from learning unnecessarily large eigenvalues and encourages better Jacobian estimation.

D.4 Weighted linear Jacobian details

We implemented a baseline Jacobian estimation method using weighted linear regression models as described in Deyle et al. [88]. Given a reference point $\mathbf{x}(t^*)$ at which the (discrete) Jacobian will be

computed, all other points are weighted according to

$$w_k = \exp \frac{-\theta \|\mathbf{x}(t_k) - \mathbf{x}(t^*)\|}{\bar{d}}$$

where

$$\bar{d} = \sum_{i=1}^T \|\mathbf{x}(t_i) - \mathbf{x}(t^*)\|$$

is the average distance from $\mathbf{x}(t^*)$ to all other points. We then perform a linear regression using the weighted points (and a bias term), the result of which is an estimate of the discrete Jacobian at $\mathbf{x}(t^*)$, which can be converted to continuous time by subtracting the identity matrix and dividing by the sampling time step (i.e., $\hat{\mathbf{J}} = \frac{\hat{\mathbf{J}}_{\text{discrete}} - \mathbf{I}}{\Delta t}$, where $\hat{\mathbf{J}}_{\text{discrete}}$ is the discrete Jacobian). The parameter θ tunes how strongly the regression is weighted towards local points. To pick θ , we sweep over values range from 0 to 10, and pick the value that yields the best one-step prediction according to

$$\mathbf{x}(t^* + 2\Delta t) = \mathbf{x}(t^* + \Delta t) + e^{\hat{\mathbf{J}}\Delta t}(\mathbf{x}(t^* + \Delta t) - \mathbf{x}(t^*))$$

where $\hat{\mathbf{J}}$ is the estimated Jacobian. This form of prediction has been previously used to learn Jacobians in machine learning settings [90]. To test the method, we pick θ based on data with observation noise at a particular noise level, then add in the denoised data to the data pool in order to compute regressions and estimate Jacobians at the true points.

D.5 Model details

All models were implemented as four-layer MLPS, with the four layers having sizes of 256, 1024, 2048, and 2048 respectively. All models used a sigmoid linear unit activation. JacobianODE models output to the dimension n^2 which was then reshaped into the matrix of the appropriate dimension. NeuralODEs output to the dimension n .

D.6 Lyapunov spectrum computation

To compute the Lyapunov spectrum, we employ a QR based algorithm [121, 122]. We discretize the Jacobians using the matrix exponential (i.e., $\hat{\mathbf{J}}_{\text{discrete}} = e^{\hat{\mathbf{J}}\Delta t}$) and then propagate a bundle of small vectors through the Jacobians, using QR to ensure the perturbations remain bounded.

D.7 Iterative Linear Quadratic Regulator (ILQR)

We implement the standard algorithm for ILQR, the details of which can be found in Li and Todorov [111] and Tassa et al. [112]. In brief, the ILQR algorithm linearizes the system dynamics around a nominal trajectory using the Jacobian, and then iteratively optimizes the control sequence using forward and backward passes to minimize the total control cost. The state cost matrix $\mathbf{Q}$ was a diagonal matrix with 1.0 along the diagonal. The final state cost matrix $\mathbf{Q}_f$ was a diagonal matrix with 1.0 along the diagonal. The control cost matrix $\mathbf{R}$ was a diagonal matrix with 0.01 along the diagonal. The control matrix was a 128×128 matrix in which the 64×64 block corresponding to the first 64 neurons (the "visual" area) was the 64-dimensional identity matrix. The control algorithm was seeded with only the initial state of the test trajectory with 5% noise. The control sequence was initialized random normal with standard deviation 0.001 and mean 0. The ILQR algorithm was run for a max of 100 iterations. The regularization was initialized at 1.0, with a minimum of 1×10^{-6} and a maximum of 1×10^{10} . Δ_0 was set to 2, as in Tassa et al. [112]. If the backward pass failed 20 times in a row, the optimization was stopped. The list of values for the line search parameter α was 1.1^{-k^2} for $k \in 0, \dots, 9$ (see Tassa et al. [112]). The linear model used for the linear baseline was computed via linear regression.

D.8 Training details

All models were implemented in PyTorch. The batch size used was 16. Gradients were accumulated for 4 batches. Training epochs were limited to 500 shuffled batches. Validation epochs were limited to 100 randomly chosen batches. Testing used all testing data. Training was run for a maximum of 1000 epochs, 3 hours, or until the early stopping was activated (see Appendix D.8.6), whichever came first.

D.8.1 Generalized Teacher Forcing

The Jacobian can be best learned when training predictions are generated recursively (i.e., replacing $\mathbf{x}(t)$ by $\hat{\mathbf{x}}(t)$). However, in chaotic systems, and/or systems with measurement noise (as considered here), this could lead to catastrophic divergence of the predicted trajectory from the true trajectory during training. We therefore employ Generalized Teacher Forcing when training all models [94]. Generalized Teacher Forcing prevents catastrophic divergence by forcing the generated predictions along the line from the prediction to the true state. Specifically, for a given predicted state $\hat{\mathbf{x}}(t)$ and true state $\mathbf{x}(t)$, the teacher forced state is

$$\tilde{\mathbf{x}}(t) = (1 - \alpha)\hat{\mathbf{x}}(t) + \alpha\mathbf{x}(t)$$

with $\alpha \in [0, 1]$. This effectively forces the predictions along a line from the prediction to the true state, by an amount with proportion α . $\alpha = 1$ corresponds to fully-forced prediction (i.e., one-step prediction) and $\alpha = 0$ corresponds to completely unforced prediction (i.e., autonomous prediction). Hess et al. [94] suggested that a good estimate of α is

$$\alpha = \max \left(\max_p \left[1 - \frac{1}{\|\mathcal{G}(\mathbf{J}_{T:2}^{(p)})\|} \right], 0 \right) \quad (17)$$

where $\mathbf{J}_{T:2}^{(p)}$ are the Jacobians of the modeled dynamics computed at data-constrained states, p indicates the batch or sequence index, and

$$\|\mathcal{G}(\mathbf{J}_{T:2}^{(p)})\| = \left\| \left(\prod_{k=0}^{T-2} \mathbf{J}_{T-k} \right)^{\frac{1}{T-1}} \right\|$$

effectively computes the discrete maximum Lyapunov exponent. In our implementation, we compute $\|\mathcal{G}(\mathbf{J}_{T:2}^{(p)})\|$ using a QR-decomposition-based Lyapunov computation algorithm [121]. As the Jacobian of the dynamics is necessary to compute this quantity, the JacobianODEs enjoy an advantage over other models in that the Jacobians are directly output by the model, and do not have to be computed via differentiating the model itself.

We furthermore employ a slightly modified version of the suggested annealing process in Hess et al. [94], which sets $\alpha_0 = 1$ and updates α_n as

$$\alpha_n = \gamma\alpha_{n-1} + (1 - \gamma)\alpha$$

where α is computed according to equation 17. Following the suggested hyperparameters, we set $\gamma = 0.999$ and update α_n every 5 batches. Once the teacher forced state $\tilde{\mathbf{x}}(t)$ is computed, it can simply replace $\mathbf{x}(t)$ in equation 11 to generate predictions.

D.8.2 Loop closure loss

We implemented a loop closure loss as discussed in Section 3.3 and Appendix A.3. For each loop, we used 20 randomly chosen points from the batch. For each batch, we constructed the same number of loops as there were batches. Path integrals were discretized in 20 steps and computed using the trapezoid method from `torchquad` [101].

D.8.3 Validation loss

All models were validated on 10 time-step prediction task with teacher forcing parameter $\alpha = 0$ (i.e., autonomous prediction).

D.8.4 Learning rate scheduling

For all models, the learning rate was annealed in accordance with teacher forcing annealing. Given an initial and final learning rates η_i and η_f we compute the effective learning rate as

$$\eta = \eta_f + \sigma(\alpha_n)(\eta_i - \eta_f)$$

where α_n is the current value of the teacher forcing parameter and

$$\sigma(\alpha_n) = \frac{\alpha_n}{\alpha_n + (1 - \alpha_n)e^{-k\alpha_n}}$$

$\sigma(\alpha_n)$ is a scaling function with $\sigma(1) = 1$ and $\sigma(0) = 0$ and for which the shape of the scaling is controlled by the parameter k . For positive values of k , the scaling is super-linear, and for negative values of k it is sub-linear. We use $k = 1$, ensuring that the learning rate does not decrease too quickly at the start of learning. We set $\eta_i = 10^{-4}$ and $\eta_f = 10^{-6}$ for all models.

D.8.5 Optimizer and weight decay

All models were trained with PyTorch’s AdamW optimizer with the learning rate as described above, and weight decay parameter 10^{-4} . All other parameters were default ($\beta_1 = 0.9$, $\beta_2 = 0.999$, $\epsilon = 10^{-8}$). We also used gradient norm clipping, with a clipping value of 1.0.

D.8.6 Early stopping

For all models, we implemented an early stopping scheme that halted the training if the validation loss improved by less than 1% for two epochs in a row.

D.8.7 Added noise during learning

For the models trained on the task-trained RNN dynamics, we added 5% Gaussian i.i.d. noise (defined relative to the norm of the training data with observation noise already added). Noise was sampled for each batch and added prior to the trajectory generation step of the learning process. Additional noise was not added for the loop closure computation.

D.8.8 Hyperparameter selection

For the JacobianODEs, the primary hyperparameter to select is the loop closure loss λ_{loop} . To select this hyperparameter, we trained JacobianODE models with $\lambda_{\text{loop}} \in [0, 10^{-6}, 10^{-5}, 10^{-4}, 10^{-3}, 10^{-2}, 10^{-1}, 1, 10]$. For each run, the epoch with the lowest trajectory validation loss ($\mathcal{L}_{\text{traj}}$) is kept. Then, for this model, we compute the one-step prediction error on validation data, the validation loop closure loss ($\mathcal{L}_{\text{loop}}$), and the percentage of Jacobian eigenvalues on all validation data that have a decay rate faster than the sampling rate $\frac{1}{\Delta t}$. We exclude any models that meet any of the following criteria:

1. **One-step prediction error greater than the persistence baseline.** The persistence baseline is computed as the mean error between each time step $k\Delta t$ and the subsequent time step $(k+1)\Delta t$ across the dataset, and constitutes a sanity check for whether a model is capturing meaningful information about the dynamics.
2. **Loop closure loss greater than $\sqrt{n}$,** where n is the system dimension (see Appendix D.11 for the derivation of this bound). As discussed in the main text, we are interested in Jacobians that not only solve the trajectory prediction problem, but that also are constructed so that the rows of the matrix are approximately conservative vector fields.
3. **More than 0.1% of the Jacobian eigenvalues have a decay rate faster than the sampling rate $\frac{1}{\Delta t}$.** Since large negative eigenvalues do not impact trajectory prediction, the models may erroneously learn Jacobians with large negative eigenvalues. If the decay rate of these eigenvalues is faster than the sampling rate, we can infer that the eigenvalues are not aligned with the observed data.

If none of the models that meet criterion (2) meet criterion (1), we discount criterion (2), as this suggests that a loop closure loss below $\sqrt{n}$ bound is too strict to obtain good prediction on this system. Additionally, if none of the models that meet criterion (3) meet criterion (1), we discount criterion (1), as this suggests that noise is very high in the data, which leads to both high one-step prediction error, and large negative eigenvalues to compensate for the perturbations introduced by the noise. Of the remaining models, we select the one with the lowest trajectory validation loss $\mathcal{L}_{\text{traj}}$.

For the NeuralODEs, we needed to select the hyperparameter λ_{jac} , which regularized the mean frobenius norm of the Jacobians computed through automatic differentiation. Again, to select this hyperparameter, we trained JacobianODE models with $\lambda_{\text{jac}} \in [0, 10^{-6}, 10^{-5}, 10^{-4}, 10^{-3}, 10^{-2}, 10^{-1}, 1, 10]$. We followed exactly the above procedure with the exception of criterion (2), which was deemed unnecessary, as computing the Jacobians implicitly via a gradient of the model (using automatic differentiation) ensures that the rows of the matrix are conservative.

All other hyperparameters (model size, learning rate, length of initial trajectory, number of discretization steps, etc.) were fixed for all systems. Given the wide range of systems and behaviors and dimensionalities that the JacobianODEs are capable of capturing, this indicates that the method is robust given a reasonable choice of these hyperparameters.

D.8.9 Model hyperparameters and training details

Below are presented the details of all Jacobian estimation models considered in the main paper.

Table S7: Hyperparameters used and model details for each system and noise level. Training time is reported in seconds.

System	Noise	Model	Loop closure weight	Jacobian penalty	Training time (s)	Final epoch	Learning rate	Min learning rate	Weight decay	Learnable parameters
VanDerPol (2 dim)	1%	JacobianODE	0.0010	0	882.139	15	0.0001	1.00e-06	0.0001	6.568e+06
	1%	NeuralODE	0	1.00e-06	949.217	21	0.0001	1.00e-06	0.0001	6.564e+06
	5%	JacobianODE	0.0001	0	692.980	11	0.0001	1.00e-06	0.0001	6.568e+06
	5%	NeuralODE	0	0	249.451	6	0.0001	1.00e-06	0.0001	6.564e+06
	10%	JacobianODE	0.010	0	678.794	10	0.0001	1.00e-06	0.0001	6.568e+06
	10%	NeuralODE	0	0.0010	713.221	16	0.0001	1.00e-06	0.0001	6.564e+06
Lorenz (3 dim)	1%	JacobianODE	0.0010	0	972.268	16	0.0001	1.00e-06	0.0001	6.578e+06
	1%	NeuralODE	0	0.0010	889.819	18	0.0001	1.00e-06	0.0001	6.566e+06
	5%	JacobianODE	0.010	0	1.147e+03	18	0.0001	1.00e-06	0.0001	6.578e+06
	5%	NeuralODE	0	0.0010	680.651	15	0.0001	1.00e-06	0.0001	6.566e+06
	10%	JacobianODE	0.010	0	388.346	6	0.0001	1.00e-06	0.0001	6.578e+06
	10%	NeuralODE	0	0.010	421.488	8	0.0001	1.00e-06	0.0001	6.566e+06
Lorenz 96 (12 dim)	1%	JacobianODE	0.0010	0	1.817e+03	31	0.0001	1.00e-06	0.0001	6.857e+06
	1%	NeuralODE	0	1.00e-06	1.892e+03	32	0.0001	1.00e-06	0.0001	6.587e+06
	5%	JacobianODE	0.0010	0	716.408	12	0.0001	1.00e-06	0.0001	6.857e+06
	5%	NeuralODE	0	0.0010	1.147e+03	19	0.0001	1.00e-06	0.0001	6.587e+06
	10%	JacobianODE	0.010	0	1.213e+03	18	0.0001	1.00e-06	0.0001	6.857e+06
	10%	NeuralODE	0	0.0001	851.803	14	0.0001	1.00e-06	0.0001	6.587e+06
Lorenz 96 (32 dim)	1%	JacobianODE	0.010	0	5.160e+03	85	0.0001	1.00e-06	0.0001	8.665e+06
	1%	NeuralODE	0	0.0010	4.204e+03	43	0.0001	1.00e-06	0.0001	6.633e+06
	5%	JacobianODE	0.0010	0	1.341e+03	22	0.0001	1.00e-06	0.0001	8.665e+06
	5%	NeuralODE	0	0.0010	3.855e+03	39	0.0001	1.00e-06	0.0001	6.633e+06
	10%	JacobianODE	0.0001	0	868.262	14	0.0001	1.00e-06	0.0001	8.665e+06
	10%	NeuralODE	0	0.0010	2.695e+03	28	0.0001	1.00e-06	0.0001	6.633e+06
Lorenz 96 (64 dim)	1%	JacobianODE	0.0010	0	2.749e+03	40	0.0001	1.00e-06	0.0001	1.497e+07
	1%	NeuralODE	0	0.010	4.801e+03	30	0.0001	1.00e-06	0.0001	6.706e+06
	5%	JacobianODE	0.0001	0	1.748e+03	27	0.0001	1.00e-06	0.0001	1.497e+07
	5%	NeuralODE	0	0.010	4.879e+03	30	0.0001	1.00e-06	0.0001	6.706e+06
	10%	JacobianODE	0.0010	0	1.701e+03	25	0.0001	1.00e-06	0.0001	1.497e+07
	10%	NeuralODE	0	0.010	4.237e+03	26	0.0001	1.00e-06	0.0001	6.706e+06
Task-trained RNN	1%	JacobianODE	0.0001	0	2.496e+03	32	0.0001	1.00e-06	0.0001	4.016e+07
	1%	NeuralODE	0	0	9.777e+03	38	0.0001	1.00e-06	0.0001	6.854e+06
	5%	JacobianODE	0.0001	0	2.352e+03	29	0.0001	1.00e-06	0.0001	4.016e+07
	5%	NeuralODE	0	1.00e-05	7.770e+03	27	0.0001	1.00e-06	0.0001	6.854e+06
	10%	JacobianODE	0.010	0	2.172e+03	27	0.0001	1.00e-06	0.0001	4.016e+07
	10%	NeuralODE	0	1.00e-05	5.260e+03	18	0.0001	1.00e-06	0.0001	6.854e+06

D.9 Information about computing resources and efficiency

All models were able to be trained on a single H100 GPU, with 80 GB of memory.

Jacobian inference times Jacobian inference times for the JacobianODE and NeuralODE models are discussed in Appendix C.5. As discussed, models were implemented with four hidden layers of the same size, and tested on 100 batches of the 128-dimensional task-trained RNN data, with each batch consisting of 16 sequences of length 25. Timings were repeated ten times for each model (see Figure S3 for details).

Training time Total training times for each of the chosen models are presented in Table S7. Furthermore, we include the training time (including backward pass) for 100 batches (with 16 sequences per batch), using 10 time-step prediction, in Table S8.

D.10 Statistical details

All statistics were computed using `scipy`. For the comparison between JacobianODE and NeuralODE trajectory and Jacobian predictions, as well as the comparison of Gramian traces and minimum eigenvalues, we used a two-sample t-test. For the comparison of ILQR control accuracies and errors, we used a Wilcoxon signed-rank test.

Table S8: Trajectory training time (seconds) for each system and noise level. Training used 100 batches, with 16 sequences per batch, as well as 10 time-step prediction.

Model	Lorenz (3 dim)	VanDerPol (2 dim)	Lorenz 96 (12 dim)	Lorenz 96 (32 dim)	Lorenz 96 (64 dim)	Task-trained RNN
1% noise						
JacobianODE	15.772	13.737	11.604	16.469	16.214	23.855
NeuralODE	8.215	8.403	11.852	12.075	18.192	14.955
5% noise						
JacobianODE	12.422	16.772	10.655	13.121	12.949	23.197
NeuralODE	6.191	5.841	13.501	8.021	21.307	30.209
10% noise						
JacobianODE	12.590	15.719	18.447	10.506	18.900	22.482
NeuralODE	11.083	9.875	15.269	8.008	17.826	33.074

D.11 Derivation of loop closure loss bound

We consider the loop closure loss as defined in 3.3. We are interested in estimating a bound on the error

$$\mathbb{E} \left[\frac{1}{n} \left\| \int_{\mathcal{C}_{\text{loop}}^{(l)}} \mathbf{J} ds \right\|_2^2 \right]$$

where n is the system dimension. While in theory this quantity should be 0, in practice due to numerical estimation error, it will not be. First recall, that

$$\int_{\mathcal{C}_{\text{loop}}^{(l)}} \mathbf{J} ds = \sum_{i=1}^L \int_{\mathbf{x}(t_i \pmod{L})}^{\mathbf{x}(t_{i+1} \pmod{L})} \mathbf{J}(\mathbf{c}(r)) \mathbf{c}'(r) dr$$

where L is the number of loop points and $\mathbf{c}$ is a line from $\mathbf{x}(t_i \pmod{L})$ to $\mathbf{x}(t_{i+1} \pmod{L})$. We assume that

$$\mathbb{E} \left[\left\| \int_{\mathbf{x}(t_i \pmod{L})}^{\mathbf{x}(t_{i+1} \pmod{L})} \mathbf{J}(\mathbf{c}(r)) \mathbf{c}'(r) dr \right\|_2^2 \right] = \mathbb{E} \left[\left\| \int_{\mathbf{x}(t_j \pmod{L})}^{\mathbf{x}(t_{j+1} \pmod{L})} \mathbf{J}(\mathbf{c}(r)) \mathbf{c}'(r) dr \right\|_2^2 \right]$$

$\forall i, j \in [1, \dots, L]$, which is justified as the numerical error accrued will likely be similar along different lines for the same system. Thus

$$\begin{aligned} \mathbb{E} \left[\frac{1}{n} \left\| \int_{\mathcal{C}_{\text{loop}}^{(l)}} \mathbf{J} ds \right\|_2^2 \right] &= \mathbb{E} \left[\frac{1}{n} \left\| \sum_{i=1}^L \int_{\mathbf{x}(t_i \pmod{L})}^{\mathbf{x}(t_{i+1} \pmod{L})} \mathbf{J}(\mathbf{c}(r)) \mathbf{c}'(r) dr \right\|_2^2 \right] \\ &\leq \frac{1}{n} \mathbb{E} \left[\sum_{i=1}^L \left\| \int_{\mathbf{x}(t_i \pmod{L})}^{\mathbf{x}(t_{i+1} \pmod{L})} \mathbf{J}(\mathbf{c}(r)) \mathbf{c}'(r) dr \right\|_2^2 \right] \\ &= \frac{1}{n} \sum_{i=1}^L \mathbb{E} \left[\left\| \int_{\mathbf{x}(t_i \pmod{L})}^{\mathbf{x}(t_{i+1} \pmod{L})} \mathbf{J}(\mathbf{c}(r)) \mathbf{c}'(r) dr \right\|_2^2 \right] \\ &= \frac{1}{n} \sum_{i=1}^L \mathbb{E} \left[\left\| \int_{\mathbf{x}(t_0)}^{\mathbf{x}(t_1)} \mathbf{J}(\mathbf{c}(r)) \mathbf{c}'(r) dr \right\|_2^2 \right] \\ &= \frac{L}{n} \mathbb{E} \left[\left\| \int_{\mathbf{x}(t_0)}^{\mathbf{x}(t_1)} \mathbf{J}(\mathbf{c}(r)) \mathbf{c}'(r) dr \right\|_2^2 \right] \end{aligned}$$

where we have used the fact that the errors are assumed to be equivalent for each of the line segments comprising the overall loop path. Recall now that for the trapezoid integration rule, the error E in

integrating $\int_a^b f(x)dx$ can be computed as $E = -\frac{(b-a)^3}{12M^2} f''(u)$ for some $u \in [a, b]$. Thus the squared error is bounded as

$$E^2 \leq \left\| \max_{u \in [a, b]} \frac{(b-a)^3}{12M^2} f''(u) \right\|_2^2$$

In our case, when path integrating lines, we are effectively integrating from $a = 0$ to $b = 1$. Furthermore, let $\mathbf{g}(r) = \mathbf{J}(\mathbf{c}(r))\mathbf{c}'(r)$, the integrand of our line integrations. We assume that $\mathbb{E} \left[\left\| \max_r \mathbf{g}''(r) \right\|_2^2 \right] \propto n \cdot \sqrt{n}$ where the first n comes from the fact that the norm involves summing over the n components of the vector $\mathbf{g}''(r)$ and the second $\sqrt{n}$ involves an assumption that loop closure loss will be more difficult to compute accurately in higher dimensions, though this will be more pronounced as dimensionality initially starts increasing. Thus $\mathbb{E} \left[\left\| \max_r \mathbf{g}''(r) \right\|_2^2 \right] = kn \cdot \sqrt{n}$ for some $k \in \mathbb{R}^+$. Now, continuing on,

$$\begin{aligned} \mathbb{E} \left[\frac{1}{n} \left\| \int_{\mathcal{C}_{\text{loop}}^{(l)}} \mathbf{J} ds \right\|_2^2 \right] &\leq \frac{L}{n} \mathbb{E} \left[\left\| \int_{\mathbf{x}(t_0)}^{\mathbf{x}(t_1)} \mathbf{J}(\mathbf{c}(r))\mathbf{c}'(r) dr \right\|_2^2 \right] \\ &\leq \frac{L}{n} \mathbb{E} \left[\left\| \frac{1^3}{12M^2} \max_r \mathbf{g}''(r) \right\|_2^2 \right] \\ &= \frac{L}{12^2 n M^4} \mathbb{E} \left[\left\| \max_r \mathbf{g}''(r) \right\|_2^2 \right] \\ &= \frac{Lkn \cdot \sqrt{n}}{12^2 n M^4} \\ &= \frac{Lk\sqrt{n}}{12^2 M^4} \end{aligned}$$

Finally, assuming that the number of discretization steps M was chosen to be large enough such that $M^4 \geq \frac{1}{12^2} Lk$ we finally obtain

$$\mathbb{E} \left[\frac{1}{n} \left\| \int_{\mathcal{C}_{\text{loop}}^{(l)}} \mathbf{J} ds \right\|_2^2 \right] \leq \sqrt{n}$$

NeurIPS Paper Checklist

1. Claims

Question: Do the main claims made in the abstract and introduction accurately reflect the paper's contributions and scope?

Answer: [\[Yes\]](#)

Justification: The abstract and introduction claim the the paper contributes robust data-driven Jacobian estimation, a framework to characterize control between interacting subsystems, inference of control dynamics in trained recurrent neural networks, and accurate control of the trained network. All claims are substantiated in detail through a description of the presented method and analytical framework, as well as comprehensive testing. See Sections 3, 4, and 5.

Guidelines:

- The answer NA means that the abstract and introduction do not include the claims made in the paper.
- The abstract and/or introduction should clearly state the claims made, including the contributions made in the paper and important assumptions and limitations. A No or NA answer to this question will not be perceived well by the reviewers.
- The claims made should match theoretical and experimental results, and reflect how much the results can be expected to generalize to other settings.
- It is fine to include aspirational goals as motivation as long as it is clear that these goals are not attained by the paper.

2. Limitations

Question: Does the paper discuss the limitations of the work performed by the authors?

Answer: [\[Yes\]](#)

Justification: The primary limitation of the present approach is that it was not tested on partially observed system dynamics, a limitation which is discussed in the Discussion (Section 6). The method was tested on a wide range of datasets (Section 4) and computational efficiency is discussed in Appendix D.9.

Guidelines:

- The answer NA means that the paper has no limitation while the answer No means that the paper has limitations, but those are not discussed in the paper.
- The authors are encouraged to create a separate "Limitations" section in their paper.
- The paper should point out any strong assumptions and how robust the results are to violations of these assumptions (e.g., independence assumptions, noiseless settings, model well-specification, asymptotic approximations only holding locally). The authors should reflect on how these assumptions might be violated in practice and what the implications would be.
- The authors should reflect on the scope of the claims made, e.g., if the approach was only tested on a few datasets or with a few runs. In general, empirical results often depend on implicit assumptions, which should be articulated.
- The authors should reflect on the factors that influence the performance of the approach. For example, a facial recognition algorithm may perform poorly when image resolution is low or images are taken in low lighting. Or a speech-to-text system might not be used reliably to provide closed captions for online lectures because it fails to handle technical jargon.
- The authors should discuss the computational efficiency of the proposed algorithms and how they scale with dataset size.
- If applicable, the authors should discuss possible limitations of their approach to address problems of privacy and fairness.
- While the authors might fear that complete honesty about limitations might be used by reviewers as grounds for rejection, a worse outcome might be that reviewers discover limitations that aren't acknowledged in the paper. The authors should use their best judgment and recognize that individual actions in favor of transparency play an important role in developing norms that preserve the integrity of the community. Reviewers will be specifically instructed to not penalize honesty concerning limitations.

3. Theory assumptions and proofs

Question: For each theoretical result, does the paper provide the full set of assumptions and a complete (and correct) proof?

Answer: [Yes]

Justification: A complete and correct proof of the presented theoretical result, that the time derivative $\mathbf{f}$ can be represented in terms of its Jacobian, is provided in Appendix A.1.

Guidelines:

- The answer NA means that the paper does not include theoretical results.
- All the theorems, formulas, and proofs in the paper should be numbered and cross-referenced.
- All assumptions should be clearly stated or referenced in the statement of any theorems.
- The proofs can either appear in the main paper or the supplemental material, but if they appear in the supplemental material, the authors are encouraged to provide a short proof sketch to provide intuition.
- Inversely, any informal proof provided in the core of the paper should be complemented by formal proofs provided in appendix or supplemental material.
- Theorems and Lemmas that the proof relies upon should be properly referenced.

4. Experimental result reproducibility

Question: Does the paper fully disclose all the information needed to reproduce the main experimental results of the paper to the extent that it affects the main claims and/or conclusions of the paper (regardless of whether the code and data are provided or not)?

Answer: [Yes]

Justification: The paper provides full information about the presented methods and experiments, with sufficient detail that all results could be reproduced.

Guidelines:

- The answer NA means that the paper does not include experiments.
- If the paper includes experiments, a No answer to this question will not be perceived well by the reviewers: Making the paper reproducible is important, regardless of whether the code and data are provided or not.
- If the contribution is a dataset and/or model, the authors should describe the steps taken to make their results reproducible or verifiable.
- Depending on the contribution, reproducibility can be accomplished in various ways. For example, if the contribution is a novel architecture, describing the architecture fully might suffice, or if the contribution is a specific model and empirical evaluation, it may be necessary to either make it possible for others to replicate the model with the same dataset, or provide access to the model. In general, releasing code and data is often one good way to accomplish this, but reproducibility can also be provided via detailed instructions for how to replicate the results, access to a hosted model (e.g., in the case of a large language model), releasing of a model checkpoint, or other means that are appropriate to the research performed.
- While NeurIPS does not require releasing code, the conference does require all submissions to provide some reasonable avenue for reproducibility, which may depend on the nature of the contribution. For example
 - (a) If the contribution is primarily a new algorithm, the paper should make it clear how to reproduce that algorithm.
 - (b) If the contribution is primarily a new model architecture, the paper should describe the architecture clearly and fully.
 - (c) If the contribution is a new model (e.g., a large language model), then there should either be a way to access this model for reproducing the results or a way to reproduce the model (e.g., with an open-source dataset or instructions for how to construct the dataset).
 - (d) We recognize that reproducibility may be tricky in some cases, in which case authors are welcome to describe the particular way they provide for reproducibility. In the case of closed-source models, it may be that access to the model is limited in some way (e.g., to registered users), but it should be possible for other researchers to have some path to reproducing or verifying the results.

5. Open access to data and code

Question: Does the paper provide open access to the data and code, with sufficient instructions to faithfully reproduce the main experimental results, as described in supplemental material?

Answer: [Yes]

Justification: As part of the submission, we provide an anonymized link to download the analyzed data, as well as all code implementing the models, baselines, control analyses, and data generation. If the paper is accepted, the full code will be open sourced.

Guidelines:

- The answer NA means that paper does not include experiments requiring code.
- Please see the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- While we encourage the release of code and data, we understand that this might not be possible, so “No” is an acceptable answer. Papers cannot be rejected simply for not including code, unless this is central to the contribution (e.g., for a new open-source benchmark).
- The instructions should contain the exact command and environment needed to run to reproduce the results. See the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- The authors should provide instructions on data access and preparation, including how to access the raw data, preprocessed data, intermediate data, and generated data, etc.
- The authors should provide scripts to reproduce all experimental results for the new proposed method and baselines. If only a subset of experiments are reproducible, they should state which ones are omitted from the script and why.
- At submission time, to preserve anonymity, the authors should release anonymized versions (if applicable).
- Providing as much information as possible in supplemental material (appended to the paper) is recommended, but including URLs to data and code is permitted.

6. Experimental setting/details

Question: Does the paper specify all the training and test details (e.g., data splits, hyperparameters, how they were chosen, type of optimizer, etc.) necessary to understand the results?

Answer: [Yes]

Justification: The paper specifies all training and test details, including data splits, hyperparameter selection, and training details such as optimizers.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The experimental setting should be presented in the core of the paper to a level of detail that is necessary to appreciate the results and make sense of them.
- The full details can be provided either with the code, in appendix, or as supplemental material.

7. Experiment statistical significance

Question: Does the paper report error bars suitably and correctly defined or other appropriate information about the statistical significance of the experiments?

Answer: [Yes]

Justification: The paper includes error bars and statistical information. Additionally, the factors of variability that the errors are capturing are clearly stated.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The authors should answer "Yes" if the results are accompanied by error bars, confidence intervals, or statistical significance tests, at least for the experiments that support the main claims of the paper.

- The factors of variability that the error bars are capturing should be clearly stated (for example, train/test split, initialization, random drawing of some parameter, or overall run with given experimental conditions).
- The method for calculating the error bars should be explained (closed form formula, call to a library function, bootstrap, etc.)
- The assumptions made should be given (e.g., Normally distributed errors).
- It should be clear whether the error bar is the standard deviation or the standard error of the mean.
- It is OK to report 1-sigma error bars, but one should state it. The authors should preferably report a 2-sigma error bar than state that they have a 96% CI, if the hypothesis of Normality of errors is not verified.
- For asymmetric distributions, the authors should be careful not to show in tables or figures symmetric error bars that would yield results that are out of range (e.g. negative error rates).
- If error bars are reported in tables or plots, The authors should explain in the text how they were calculated and reference the corresponding figures or tables in the text.

8. Experiments compute resources

Question: For each experiment, does the paper provide sufficient information on the computer resources (type of compute workers, memory, time of execution) needed to reproduce the experiments?

Answer: [Yes]

Justification: We provide information about compute resources needed to reproduce the experiments in Appendix D.9.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The paper should indicate the type of compute workers CPU or GPU, internal cluster, or cloud provider, including relevant memory and storage.
- The paper should provide the amount of compute required for each of the individual experimental runs as well as estimate the total compute.
- The paper should disclose whether the full research project required more compute than the experiments reported in the paper (e.g., preliminary or failed experiments that didn't make it into the paper).

9. Code of ethics

Question: Does the research conducted in the paper conform, in every respect, with the NeurIPS Code of Ethics <https://neurips.cc/public/EthicsGuidelines>?

Answer: [Yes]

Justification: Ethical considerations have been taken into account and the research conforms fully to the NeurIPS code of ethics.

Guidelines:

- The answer NA means that the authors have not reviewed the NeurIPS Code of Ethics.
- If the authors answer No, they should explain the special circumstances that require a deviation from the Code of Ethics.
- The authors should make sure to preserve anonymity (e.g., if there is a special consideration due to laws or regulations in their jurisdiction).

10. Broader impacts

Question: Does the paper discuss both potential positive societal impacts and negative societal impacts of the work performed?

Answer: [Yes]

Justification: Yes, the paper discusses societal impacts, including potential applications to modeling biological dysfunction in disease. There are no expected negative impacts.

Guidelines:

- The answer NA means that there is no societal impact of the work performed.
- If the authors answer NA or No, they should explain why their work has no societal impact or why the paper does not address societal impact.

- Examples of negative societal impacts include potential malicious or unintended uses (e.g., disinformation, generating fake profiles, surveillance), fairness considerations (e.g., deployment of technologies that could make decisions that unfairly impact specific groups), privacy considerations, and security considerations.
- The conference expects that many papers will be foundational research and not tied to particular applications, let alone deployments. However, if there is a direct path to any negative applications, the authors should point it out. For example, it is legitimate to point out that an improvement in the quality of generative models could be used to generate deepfakes for disinformation. On the other hand, it is not needed to point out that a generic algorithm for optimizing neural networks could enable people to train models that generate Deepfakes faster.
- The authors should consider possible harms that could arise when the technology is being used as intended and functioning correctly, harms that could arise when the technology is being used as intended but gives incorrect results, and harms following from (intentional or unintentional) misuse of the technology.
- If there are negative societal impacts, the authors could also discuss possible mitigation strategies (e.g., gated release of models, providing defenses in addition to attacks, mechanisms for monitoring misuse, mechanisms to monitor how a system learns from feedback over time, improving the efficiency and accessibility of ML).

11. Safeguards

Question: Does the paper describe safeguards that have been put in place for responsible release of data or models that have a high risk for misuse (e.g., pretrained language models, image generators, or scraped datasets)?

Answer: [NA]

Justification: This paper contributes an analytical framework for characterizing nonlinear control in interacting subsystems, and poses no risk of misuse. We do not make use of any language models or image generators, and do not use any scraped datasets. We therefore do not describe safeguards.

Guidelines:

- The answer NA means that the paper poses no such risks.
- Released models that have a high risk for misuse or dual-use should be released with necessary safeguards to allow for controlled use of the model, for example by requiring that users adhere to usage guidelines or restrictions to access the model or implementing safety filters.
- Datasets that have been scraped from the Internet could pose safety risks. The authors should describe how they avoided releasing unsafe images.
- We recognize that providing effective safeguards is challenging, and many papers do not require this, but we encourage authors to take this into account and make a best faith effort.

12. Licenses for existing assets

Question: Are the creators or original owners of assets (e.g., code, data, models), used in the paper, properly credited and are the license and terms of use explicitly mentioned and properly respected?

Answer: [Yes]

Justification: Yes, we implement existing methods in our paper and cite the original creators appropriately.

Guidelines:

- The answer NA means that the paper does not use existing assets.
- The authors should cite the original paper that produced the code package or dataset.
- The authors should state which version of the asset is used and, if possible, include a URL.
- The name of the license (e.g., CC-BY 4.0) should be included for each asset.
- For scraped data from a particular source (e.g., website), the copyright and terms of service of that source should be provided.

- If assets are released, the license, copyright information, and terms of use in the package should be provided. For popular datasets, paperswithcode.com/datasets has curated licenses for some datasets. Their licensing guide can help determine the license of a dataset.
- For existing datasets that are re-packaged, both the original license and the license of the derived asset (if it has changed) should be provided.
- If this information is not available online, the authors are encouraged to reach out to the asset's creators.

13. New assets

Question: Are new assets introduced in the paper well documented and is the documentation provided alongside the assets?

Answer: [\[Yes\]](#)

Justification: We provide well-documented code for generating the models and implementing the datasets used in this paper.

Guidelines:

- The answer NA means that the paper does not release new assets.
- Researchers should communicate the details of the dataset/code/model as part of their submissions via structured templates. This includes details about training, license, limitations, etc.
- The paper should discuss whether and how consent was obtained from people whose asset is used.
- At submission time, remember to anonymize your assets (if applicable). You can either create an anonymized URL or include an anonymized zip file.

14. Crowdsourcing and research with human subjects

Question: For crowdsourcing experiments and research with human subjects, does the paper include the full text of instructions given to participants and screenshots, if applicable, as well as details about compensation (if any)?

Answer: [\[NA\]](#)

Justification: This paper does not involve crowdsourcing nor research with human subjects.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Including this information in the supplemental material is fine, but if the main contribution of the paper involves human subjects, then as much detail as possible should be included in the main paper.
- According to the NeurIPS Code of Ethics, workers involved in data collection, curation, or other labor should be paid at least the minimum wage in the country of the data collector.

15. Institutional review board (IRB) approvals or equivalent for research with human subjects

Question: Does the paper describe potential risks incurred by study participants, whether such risks were disclosed to the subjects, and whether Institutional Review Board (IRB) approvals (or an equivalent approval/review based on the requirements of your country or institution) were obtained?

Answer: [\[NA\]](#)

Justification: This paper does not involve crowdsourcing nor research with human subjects.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Depending on the country in which research is conducted, IRB approval (or equivalent) may be required for any human subjects research. If you obtained IRB approval, you should clearly state this in the paper.

- We recognize that the procedures for this may vary significantly between institutions and locations, and we expect authors to adhere to the NeurIPS Code of Ethics and the guidelines for their institution.
- For initial submissions, do not include any information that would break anonymity (if applicable), such as the institution conducting the review.

16. **Declaration of LLM usage**

Question: Does the paper describe the usage of LLMs if it is an important, original, or non-standard component of the core methods in this research? Note that if the LLM is used only for writing, editing, or formatting purposes and does not impact the core methodology, scientific rigorousness, or originality of the research, declaration is not required.

Answer: [NA]

Justification: The core method development in this research does not involve LLMs as any important, original, or non-standard components.

Guidelines:

- The answer NA means that the core method development in this research does not involve LLMs as any important, original, or non-standard components.
- Please refer to our LLM policy (<https://neurips.cc/Conferences/2025/LLM>) for what should or should not be described.