

PHD ORAL DEFENSE · NATIONAL UNIVERSITY OF SINGAPORE · 2026

Machine Learning Augmented Simulation and Analysis of Dynamical Systems

Presenter: Guo Yue

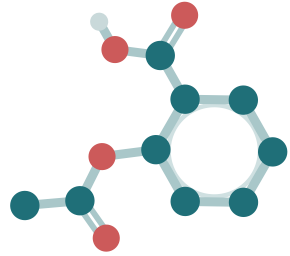
Supervisor: Li Qianxiao

Department of Mathematics

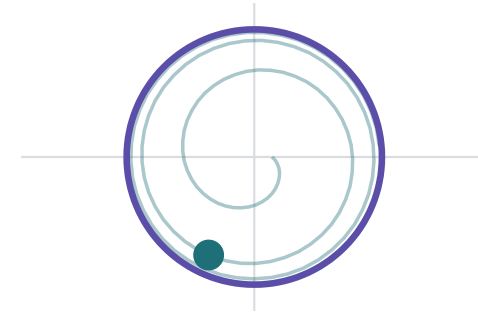


What is a dynamical system?

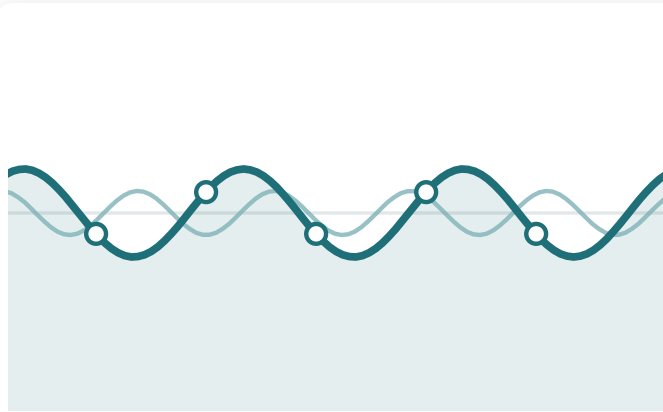
A **state that evolves in time** — four everyday examples, all in motion:



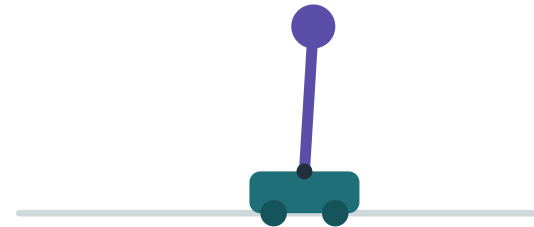
molecule · molecular dynamics



oscillator · limit cycle



wave · transport



pendulum · with control

Dynamics = choosing a computable form



molecule



oscillator



wave



pendulum · control

All are **states evolving in time**



The equation

the evolution law

$$\dot{\mathbf{y}} = \mathbf{f}(\mathbf{y})$$



The representation

coordinates

we track in



The scheme

step size & solver

Right form → the trajectory computes **accurately, stably, cheaply**;

Wrong form → it drifts or blows up.

Classically, every dial is set by hand

Each form built **from first principles, with a theorem attached:**



The equation

first principles

$$m\ddot{\mathbf{x}} = -\nabla V$$

exact · interpretable



The representation

human insight

normal modes

decoupled



The scheme

hand-designed

Euler · RK · symplectic

order p · stable

Strict and guaranteed: hand-derivable, unknown, or too expensive.

When the form is not easily hand-writable, we learn it



Only data align with trajectories

we watch the system move
but the **equation is unknown**

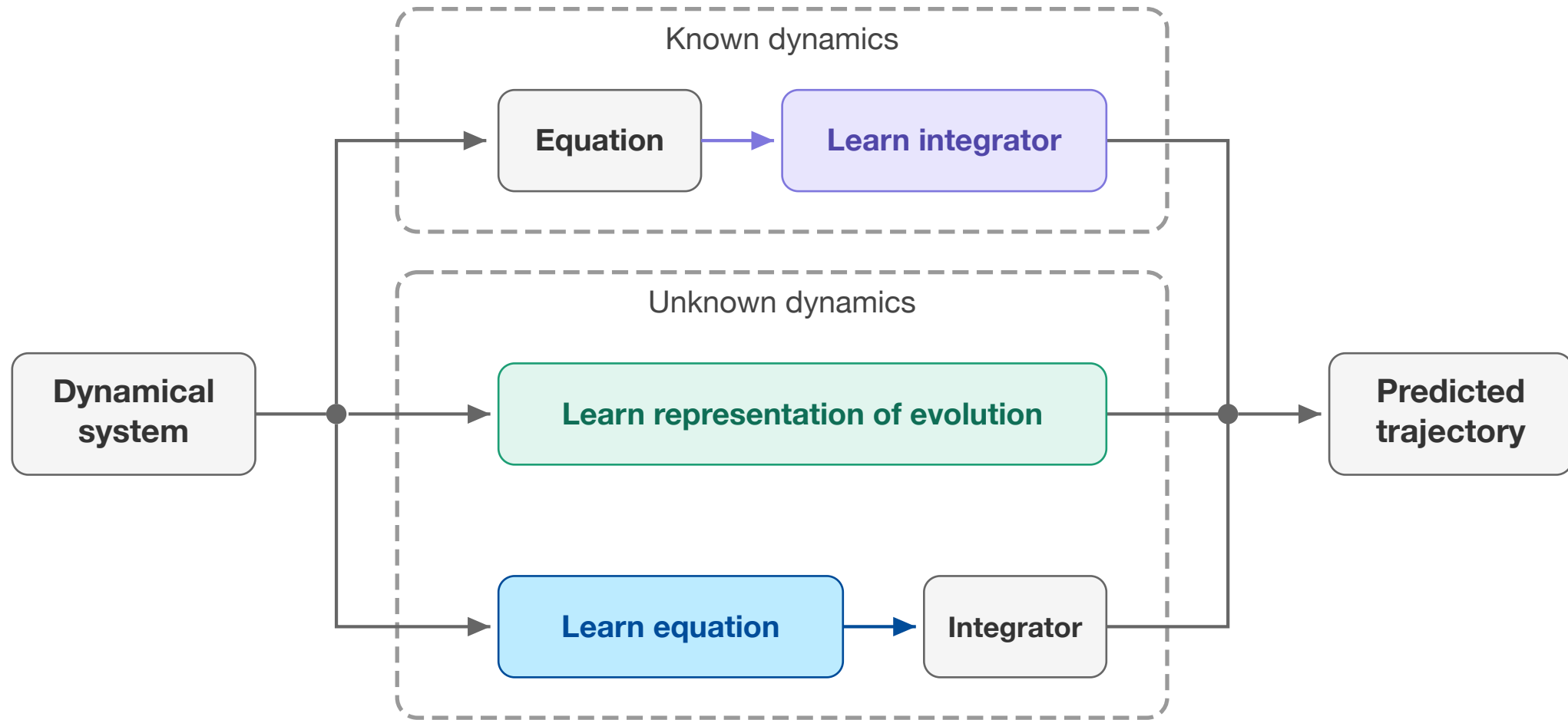


Computation costs too much
each force is a **quantum-scale calculation** (molecular dynamics)

Learning replaces a hand-designed form with
one **parameterized flexibly and fit from data**.

Bring that power into **scientific computing**, **without giving up the classical guarantees** → so **which dial do we learn?** →

Thesis viewpoint: learn the structured components



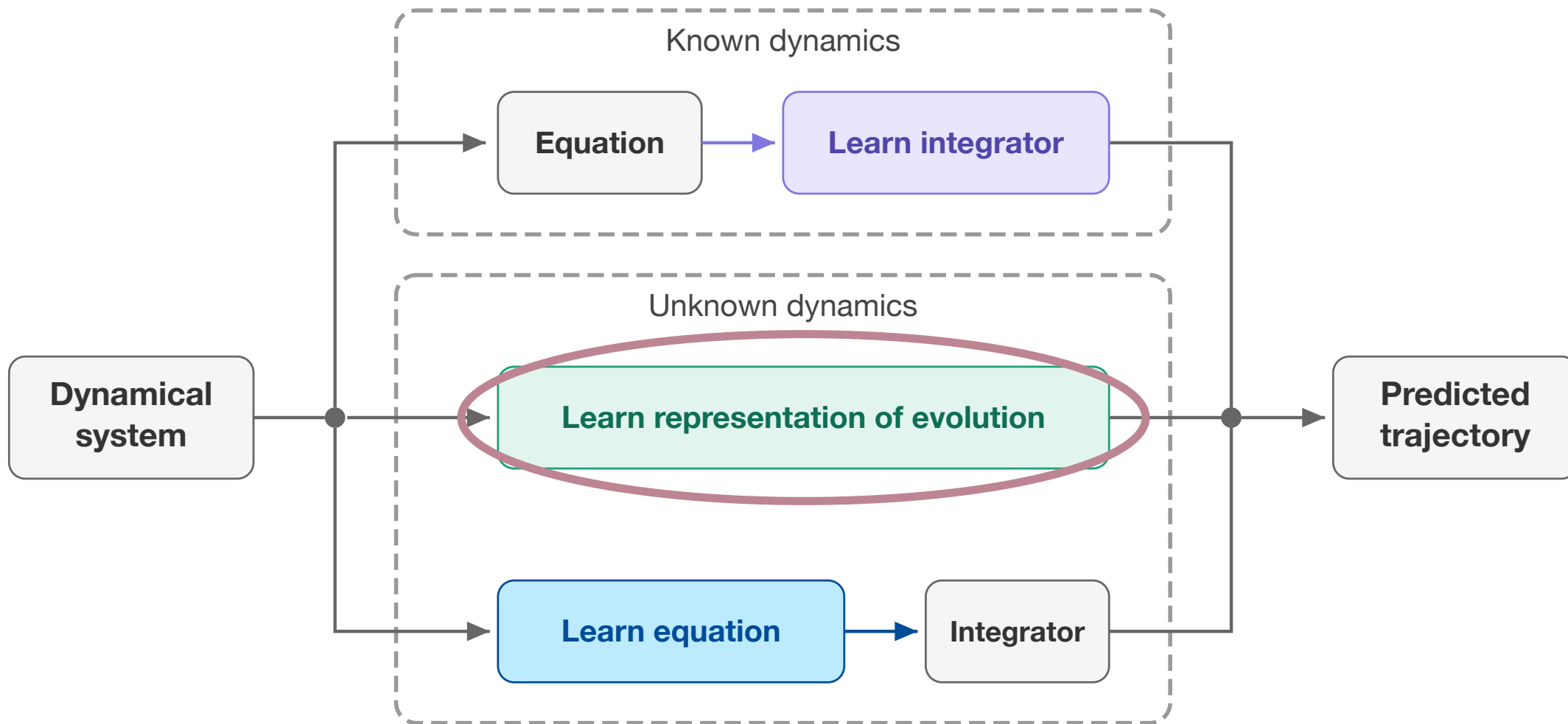
PART I · LEARNING REPRESENTATIONS

Parametric Koopman Decompositions

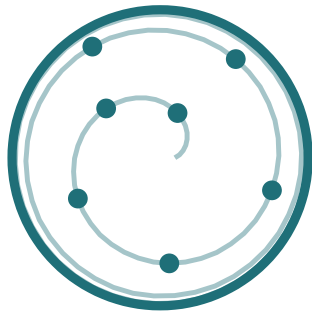
PUBLISHED AS “*Learning Parametric Koopman Decompositions for Prediction and Control*”

in **SIAM Journal on Applied Dynamical Systems** (2025)

Yue Guo, Milan Korda, Ioannis G. Kevrekidis, Qianxiao Li



Make a nonlinear system look linear PK-NN



our system: **nonlinear**, from data

learn a
representation



IF THE SYSTEM WERE LINEAR

- ||| spectral analysis
- ✓ forecast / prediction
- ✓ convex optimal control
- ✓ stability guarantees

Linear systems come with this whole toolbox **but only if the system is linear.**

Driving question: can we **learn a representation**
in which a nonlinear system becomes **linear**?

The problem: data-driven prediction & control

We have only **trajectory data** from an **unknown** transition map $\mathbf{x}_{n+1} = \mathbf{f}(\mathbf{x}_n, \mathbf{u}_n)$
— state $\mathbf{x} \in X$, parameter / control $\mathbf{u} \in U$. Two problems are posed on $\mathbf{f}$:

Forward prediction

Data: $\{\mathbf{x}_n, \mathbf{u}_n, \mathbf{x}_{n+1}\}$

Goal: given $\mathbf{x}_0, \{\mathbf{u}_i\}$, predict $\{\mathbf{x}_i\}$.

Optimal control

Data: $\{\mathbf{x}_n, \mathbf{u}_n, \mathbf{x}_{n+1}\}$

Goal: given $\mathbf{x}_0$ and a cost, find $\{\mathbf{u}_n\}$ minimizing it.

Two obstacles for data-driven modelling: **non-linearity** & **high dimensionality**.

Koopman: a finite, linear invariant subspace

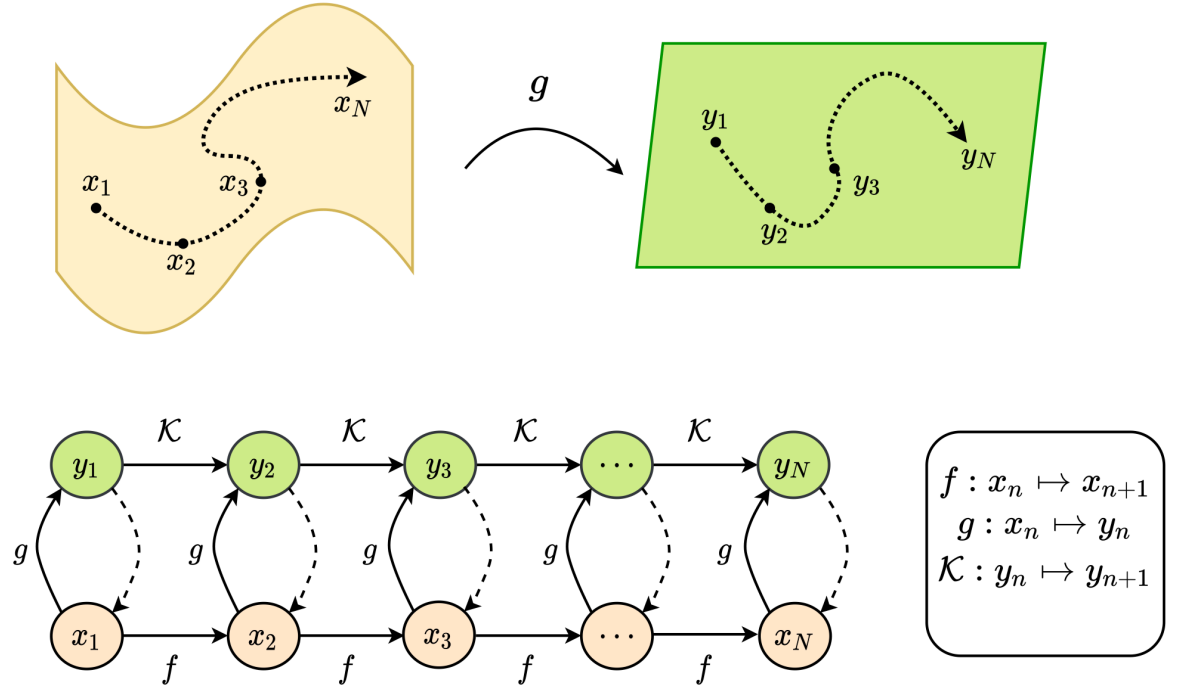
Track **observables** $g(\mathbf{x})$ instead of the state.
The **Koopman operator** $\mathcal{K}$ advances any observable **linearly**: $g(\mathbf{x}_{n+1}) = (\mathcal{K}g)(\mathbf{x}_n)$.

We keep only a chosen dictionary $\Psi = (\psi_1, \dots, \psi_m)$ spanning a finite **invariant** subspace S : $\mathcal{K} S \subset S$.

On S , $\mathcal{K}$ becomes a finite **matrix** K (basis Ψ) — the **Koopman decomposition**:

$$\Psi(\mathbf{x}_{n+1}) = K \Psi(\mathbf{x}_n)$$

finite-dim non-linear $\rightarrow$ infinite-dim
linear; DL handles high dimensions.



Nonlinear state space $\rightarrow$ flat observable space, where $\mathcal{K}$ is linear.

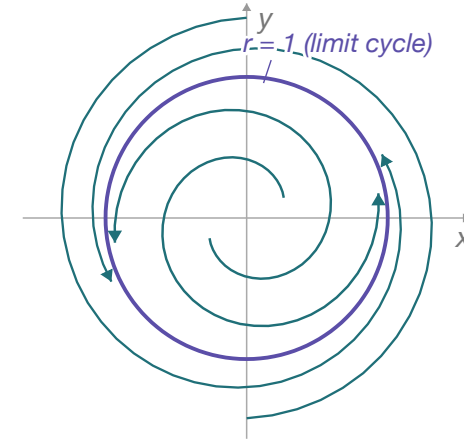
A less trivial example: a nonlinear **limit cycle**

A limit-cycle system (**Stuart–Landau**):

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

With $r^2 = x^2 + y^2$, the polar form decouples:

$$\dot{r} = r(1 - r^2), \quad \dot{\theta} = \omega$$



The amplitude relaxes **nonlinearly**, while the phase rotates uniformly.

Koopman coordinates separate decay and rotation

Observables ψ_a (amplitude), (ψ_c, ψ_s) (phase):

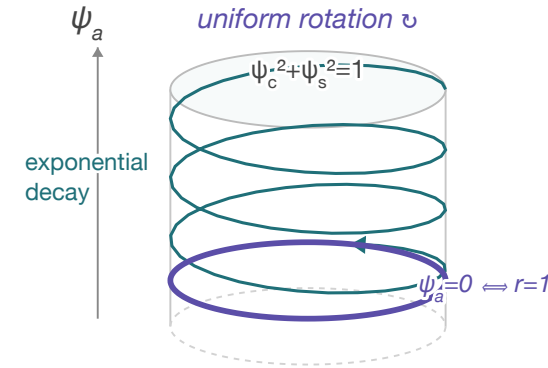
$$\psi_a = \frac{1}{r^2} - 1, \quad \psi_c = \cos \theta, \quad \psi_s = \sin \theta$$

The linear evolution on the observables:

$$\dot{\Psi} = A\Psi, \quad A = \begin{pmatrix} -2 & 0 & 0 \\ 0 & 0 & -\omega \\ 0 & \omega & 0 \end{pmatrix}$$

$$\Psi_{n+1} = K\Psi_n$$

$$K = e^{A\Delta t} = \begin{pmatrix} e^{-2\Delta t} & 0 & 0 \\ 0 & \cos \omega \Delta t & -\sin \omega \Delta t \\ 0 & \sin \omega \Delta t & \cos \omega \Delta t \end{pmatrix}$$



The lifted states still form a **two-dimensional invariant manifold**:

- No increase in intrinsic dimension.
- Linear evolution in observable space.

Finite-dimensional closure is **not generic**

Nonlinear **pendulum**:

$$\dot{\theta} = v, \quad \dot{v} = -\sin \theta$$

Start from a finite dictionary:

$$\Psi = (\theta, v, \sin \theta, \cos \theta)^\top$$

Differentiating keeps producing **new** observables: the dictionary **never closes**.

$$\frac{d}{dt} \sin \theta = v \cos \theta$$

$$\frac{d}{dt} \cos \theta = -v \sin \theta$$

$$\frac{d}{dt} (v \cos \theta) = -\sin \theta \cos \theta - v^2 \sin \theta$$

$$\frac{d}{dt} (v \sin \theta) = -\sin^2 \theta + v^2 \cos \theta$$

$$\begin{aligned} \sin \theta, \cos \theta &\rightarrow v \cos \theta, v \sin \theta \\ &\rightarrow v^2 \sin \theta, v^2 \cos \theta, \sin \theta \cos \theta, \dots \end{aligned}$$

The observable hierarchy keeps growing.

The Koopman operator is always **linear**, but its invariant observable space is generally **infinite-dimensional**.

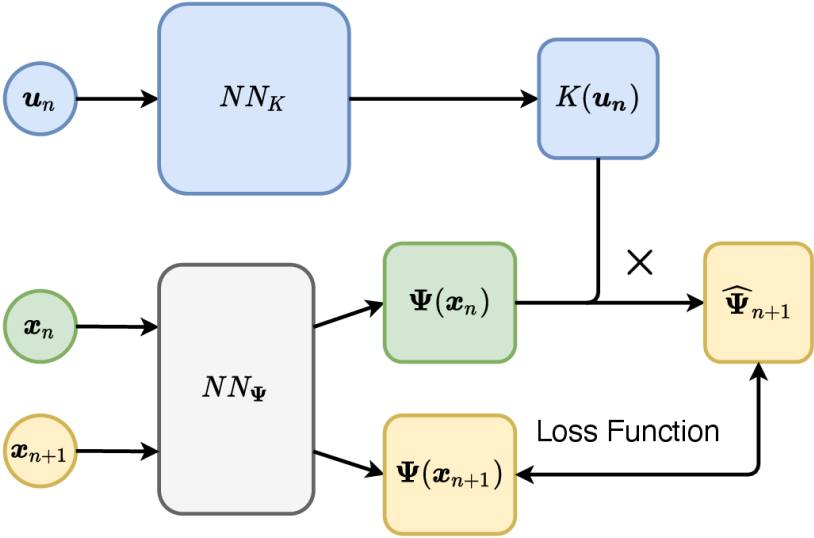
Finite-dimensional Koopman models are therefore often **approximations**, not exact closures.

PK-NN: learn Ψ and $K(u)$ jointly

For **parametric** systems $\mathcal{K}$ depends on u , PK-NN learns one shared **dictionary** Ψ and a **fully nonlinear operator** $K(u)$.

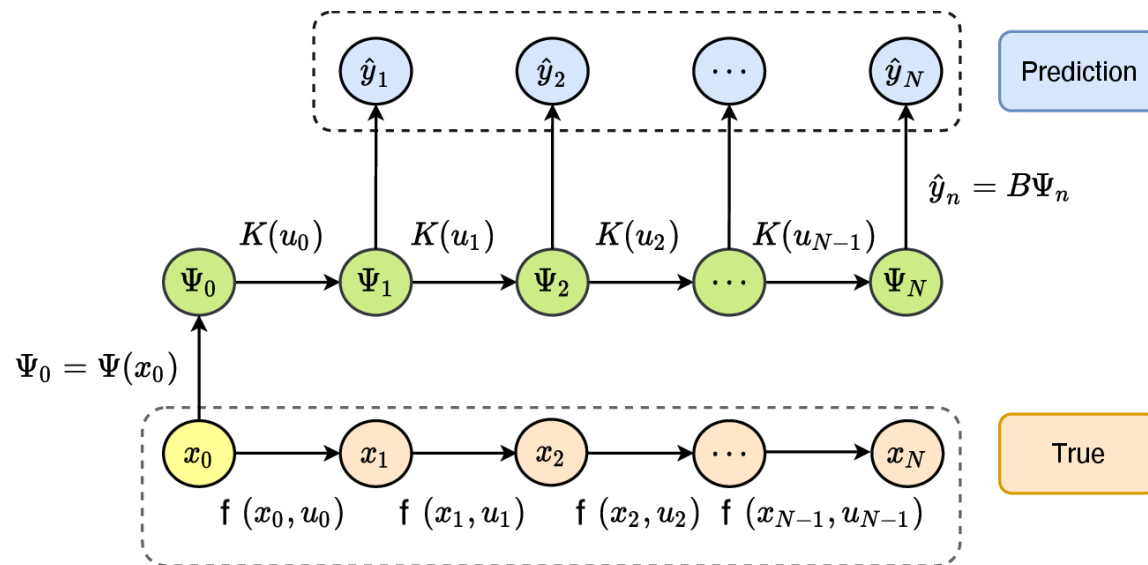
Model	Form	
Optimized DMD	$x_{n+1} = K x_n$	M0
EDMD / EDMD-DL	$\Psi(x_{n+1}) = K \Psi(x_n)$	M1
Linear control	$\Psi_{n+1} = A \Psi_n + B u_n$	M2
Bilinear	$\Psi_{n+1} = A \Psi_n + \sum_i B_i u_{n,i} \Psi_n$	M3
Fixed-basis param	$\Psi_{n+1} = \sum_i h_i(u_n) K_i \Psi_n$	M4
Ours (PK-NN)	$\Psi(x_{n+1}) = K(u_n) \Psi(x_n)$	Ours

► comparison

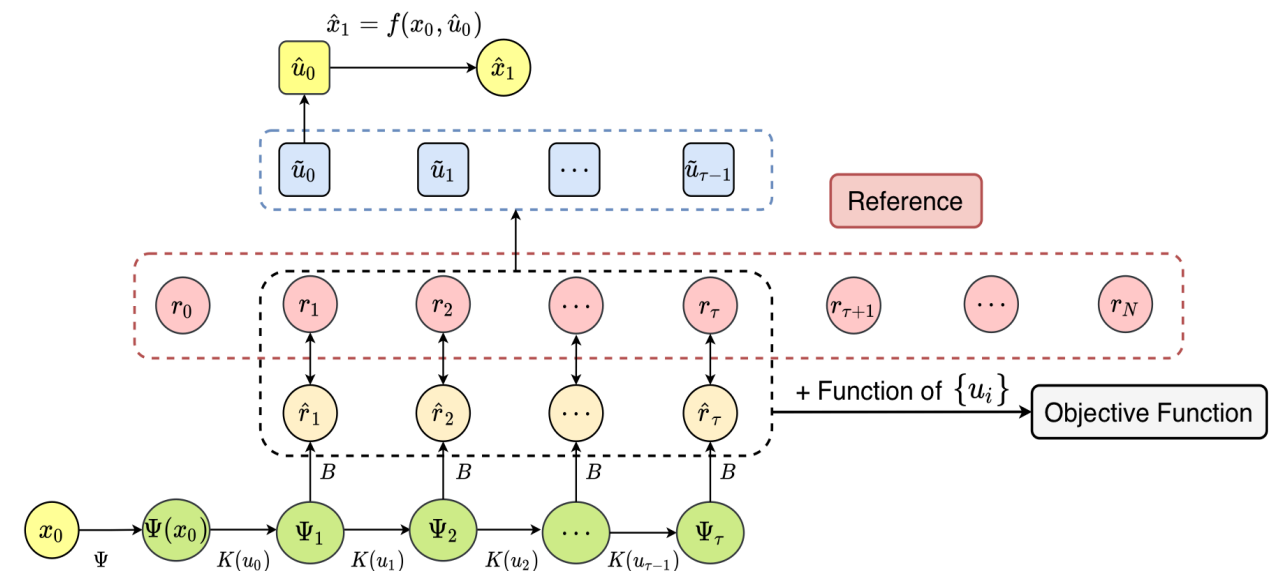


$$\min_{\theta_\Psi, \theta_K} \sum_{n,j} \left\| \Psi(x_{n+1}^{(j)}; \theta_\Psi) - K(u_n^{(j)}; \theta_K) \Psi(x_n^{(j)}; \theta_\Psi) \right\|^2$$

We solve **both**: forward prediction & optimal control

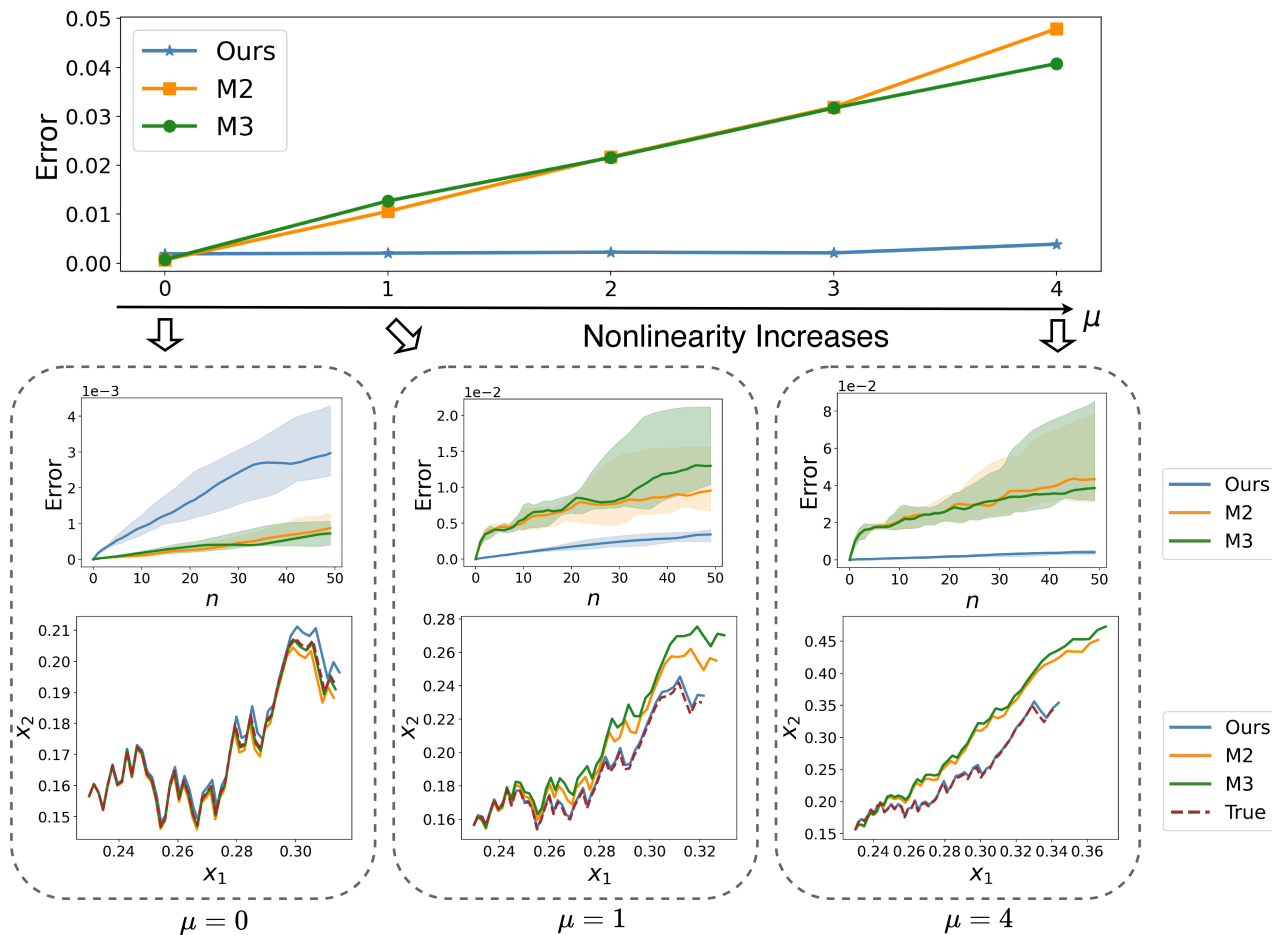


Forward prediction — lift, step with $K(u_n)$, read out $\hat{y} = B\Psi$.



Optimal control (MPC) — optimize $\{u_n\}$ to track a reference.

Prediction under strong nonlinearity in u



Van der Pol–Mathieu; μ controls how nonlinearly the control u enters:

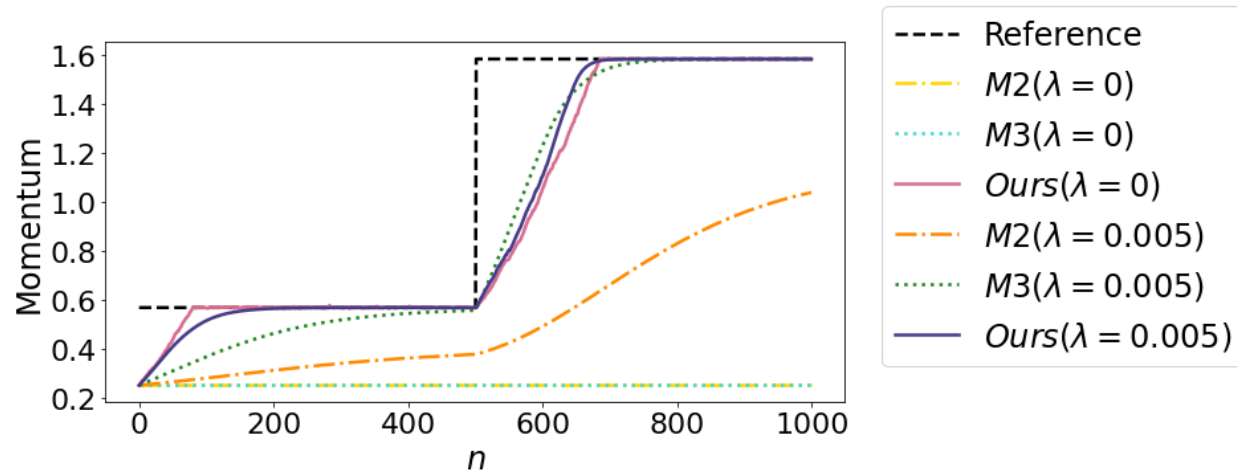
$$\dot{x}_1 = x_2$$

$$\dot{x}_2 = (2 - 2x_1^2)x_2 - (1 + 2\mu u^2 - \mu)x_1 + u$$

Advantage: the fully nonlinear $K(u)$ stays accurate as μ grows and generalizes to **unseen parameters**

Van der Pol–Mathieu — error vs. step, as nonlinearity in u increases.

Optimal control — forced KdV (MPC)



Ours (solid) follows the step reference (dashed); baselines M2 / M3 lag.

Forced KdV — control u enters nonlinearly:

$$\partial_t \eta + \eta \partial_x \eta + \partial_{xxx} \eta = \sum_i v_i(x) \sin(\pi u_i)$$

Track a reference r_n (momentum $\int \eta^2 dx$) by **Model Predictive Control** in the lifted space:

$$\min_{\{u_n\}} \sum_n |\hat{r}_n - r_n|^2 + \lambda \|u_n\|^2$$

The model reads $\hat{r}_n = \mathbf{c}^\top \Psi(\mathbf{x}_n)$ **linearly** off the lifted state $\Psi_{n+1} = K(\mathbf{u}_n) \Psi_n$.

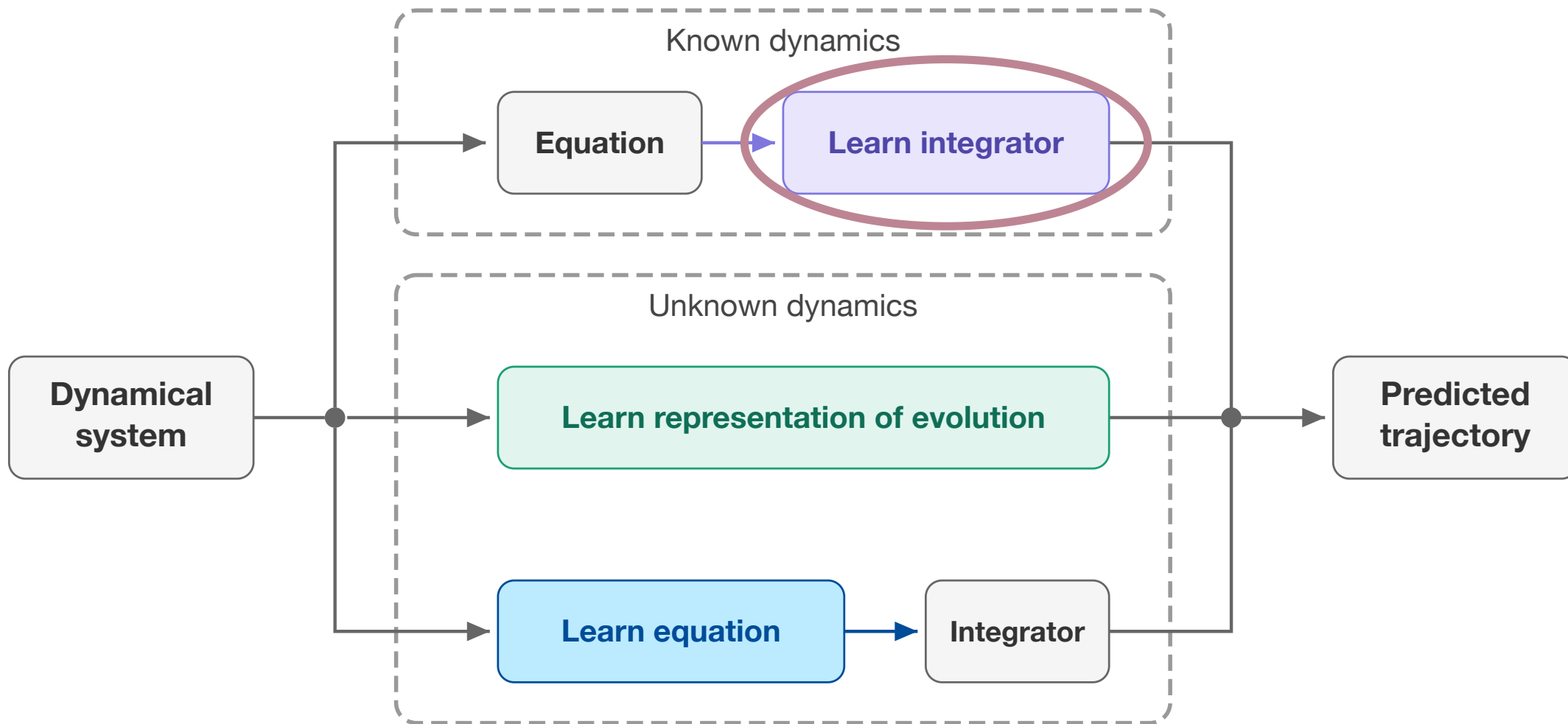
PART II · LEARNING THE NUMERICAL MECHANISM

Personalized ODE Integrators

PUBLISHED AS “*Personalized Algorithm Generation: A Case Study in Learning ODE Integrators*”

in **SIAM Journal on Scientific Computing** (2022)

Yue Guo, Felix Dietrich, Tom Bertalan, Danimir T. Doncevic, Manuel Dahmen, Ioannis G. Kevrekidis, Qianxiao Li



An integrator built for Specific system RK-NN

Classical Solver: one recipe for **all f**



A **worst-case tax**: robust
to systems you never solve

Specific family: solved **many times**

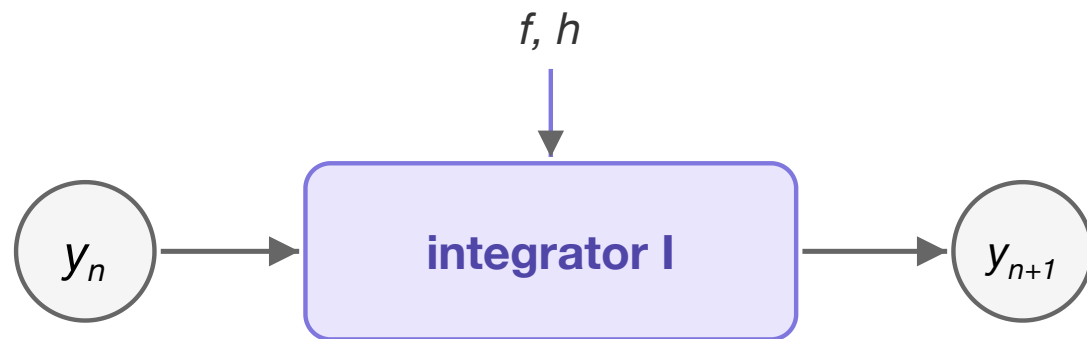


Worst-case coefficients are never tuned to it,
you pay that tax on every run

Driving question: can we build (learn) an integrator
specialized to the family we actually solve?

Personalized ODE integrators RK-NN

Known ODE $\dot{\mathbf{y}} = \mathbf{f}(\mathbf{y})$. A numerical **integrator** I advances the state one step h : $\mathbf{y}_n \rightarrow \mathbf{y}_{n+1}$.



One step, using the known f .

How should we **optimize** the integrator I ?

Classical: I tuned for the **worst case** over a large class $\mathcal{F}$:

$$\sup_{F \in \mathcal{F}} \|L(I, F)\| = \mathcal{O}(h^\alpha)$$

Ours: I optimized over the **family** μ we actually solve:

$$\mathbb{E}_\mu \|L(I, F_\mu)\| = \mathcal{O}(h^\alpha)$$

$\Rightarrow$ A **personalized** integrator

Classical RK: coefficients fixed by hand

A **Runge–Kutta** step $\hat{\mathbf{y}}_{n+1} = \mathbf{y}_n + h \sum_i b_i \mathbf{k}_i$
 , $\mathbf{k}_i = \mathbf{f}(\mathbf{y}_n + h \sum_j a_{ij} \mathbf{k}_j)$, with coefficients (a_{ij}, b_i) , $c_i = \sum_j a_{ij}$.

Taylor-expand the **exact** step in the **elementary differentials** (fixed coefficients):

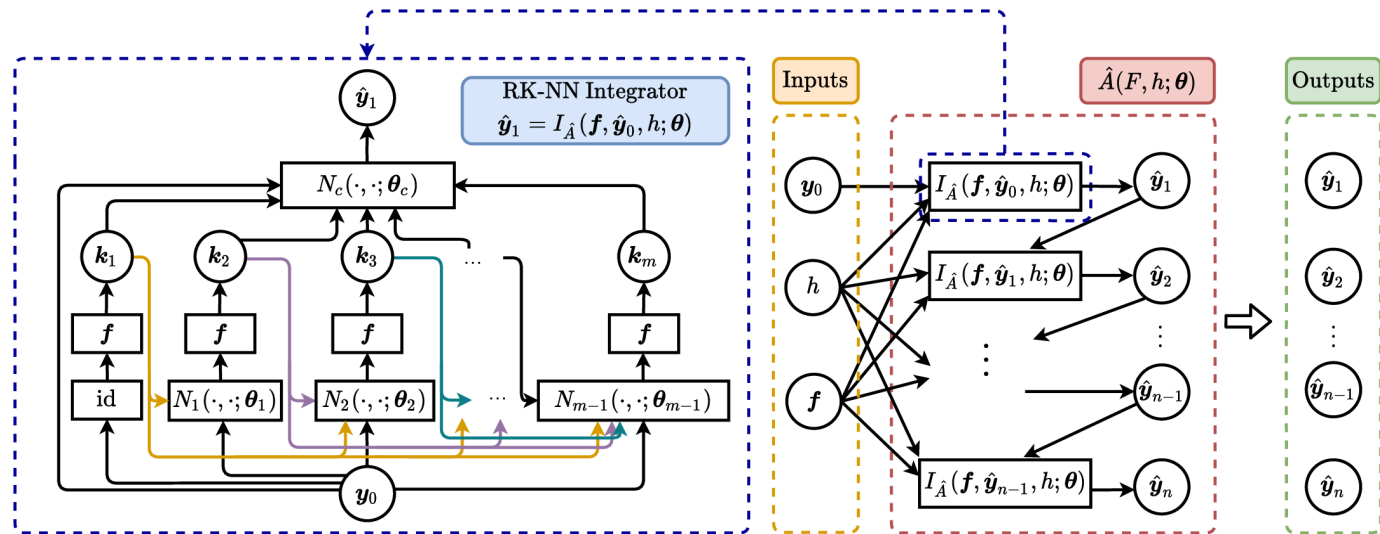
$$\mathbf{y}_{n+1} = \mathbf{y}_n + h \mathbf{f} + \frac{h^2}{2} \mathbf{f}' \mathbf{f} + \frac{h^3}{6} (\mathbf{f}'' \mathbf{f} \mathbf{f} + \mathbf{f}' \mathbf{f}' \mathbf{f}) + \cdots$$

RK expands in **the same** differentials — match each (for **all** $\mathbf{f}$) $\Rightarrow$ **one condition per term**:

term	set these two coefficients equal		$\Rightarrow$ order condition on (a_{ij}, b_i)
	in exact $\mathbf{y}_{n+1}$	in RK $\hat{\mathbf{y}}_{n+1}$	
$h \mathbf{f}$	1	$\sum_i b_i$	$\sum_i b_i = 1$
$h^2 \mathbf{f}' \mathbf{f}$	$\frac{1}{2}$	$\sum_i b_i c_i$	$\sum_i b_i c_i = \frac{1}{2}$
$h^3 \mathbf{f}'' \mathbf{f} \mathbf{f}$	$\frac{1}{6}$	$\frac{1}{2} \sum_i b_i c_i^2$	$\frac{1}{2} \sum_i b_i c_i^2 = \frac{1}{6}$
$h^3 \mathbf{f}' \mathbf{f}' \mathbf{f}$	$\frac{1}{6}$	$\sum_{ij} b_i a_{ij} c_j$	$\sum_{ij} b_i a_{ij} c_j = \frac{1}{6}$

Matched through h^3 : the step differs only by the **truncation error** $\hat{\mathbf{y}}_{n+1} - \mathbf{y}_{n+1} = O(h^4)$.

RK-NN: learn the coefficients θ



The Butcher-tableau coefficients become learnable θ inside an m -stage RK.

Data: sample the family F_μ , initial states y_0 , and steps $h \in (0.01, 0.1)$; the target $y_1 = \varphi_h(y_0)$ is the accurate one-step solution of the known f .

Train θ on the loss:

$$\mathcal{L} = \frac{\|y_1 - \hat{y}_1\|^2}{\|y_1 - \hat{y}_1^{(\text{RK})}\|^2} + \lambda \mathcal{R}$$

$$\mathcal{R} = \sum_{i=1}^{\alpha} \left\| \frac{d^i}{dh^i} \Big|_0 (y_1 - \hat{y}_1) \right\|^2$$

$\mathcal{L}$: one-step error, scaled by the classical RK error $\|y_1 - \hat{y}_1^{(\text{RK})}\|$
 $\mathcal{R}$: Taylor regularizer enforcing order α .

RK-NN keeps order 3 yet beats RK3

Two real 2-D nonlinear systems, state $\mathbf{y} = (u, v)$. A 3-stage RK-NN ($\alpha = 3$) is trained per family over $h \in (0.01, 0.1)$; outside that range it degrades — **specialized by design**.

Van der Pol

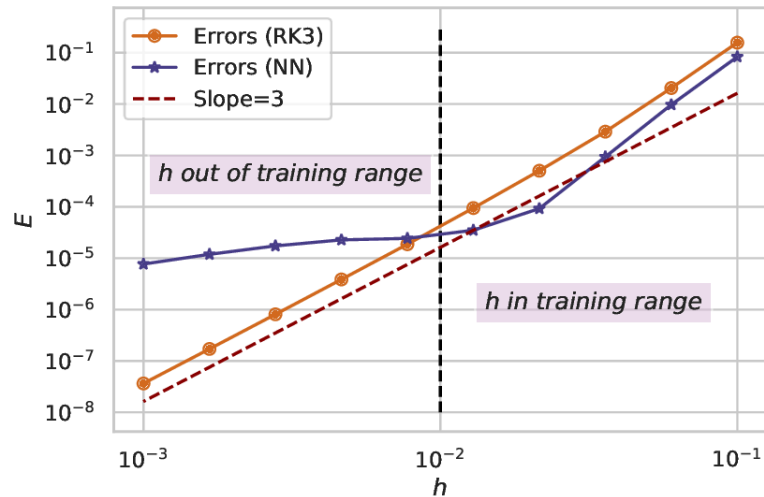
$$\dot{u} = v, \quad \dot{v} = \mu(1 - u^2)v - u$$

limit cycle; nonlinearity set by μ .

Brusselator

$$\dot{u} = A + u^2v - (B+1)u, \quad \dot{v} = Bu - u^2v$$

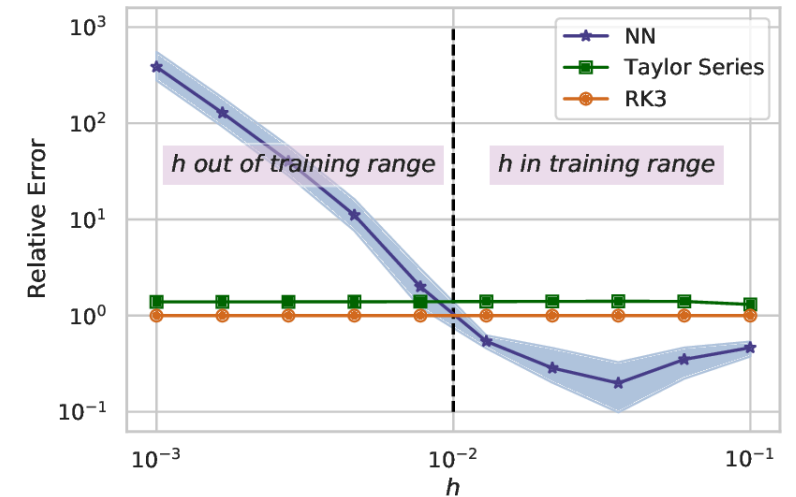
autocatalytic reaction with constants A, B .



Van der Pol — order check



Brusselator — order check



Van der Pol — relative error

Highlight: “breaking” the order barrier

2-stage scheme, $k_2 = f(y + h\theta_1 k_1)$:
$$y_{n+1} = y + h(\theta_{c1} k_1 + \theta_{c2} k_2)$$

Classical: Taylor-match → order 2:

$$\theta_{c1} + \theta_{c2} = 1, \quad 2\theta_1\theta_{c2} = 1 \quad \Rightarrow_{(\text{Heun})} \quad \theta_1=1, \quad \theta_{c1}=\theta_{c2}=\frac{1}{2}$$

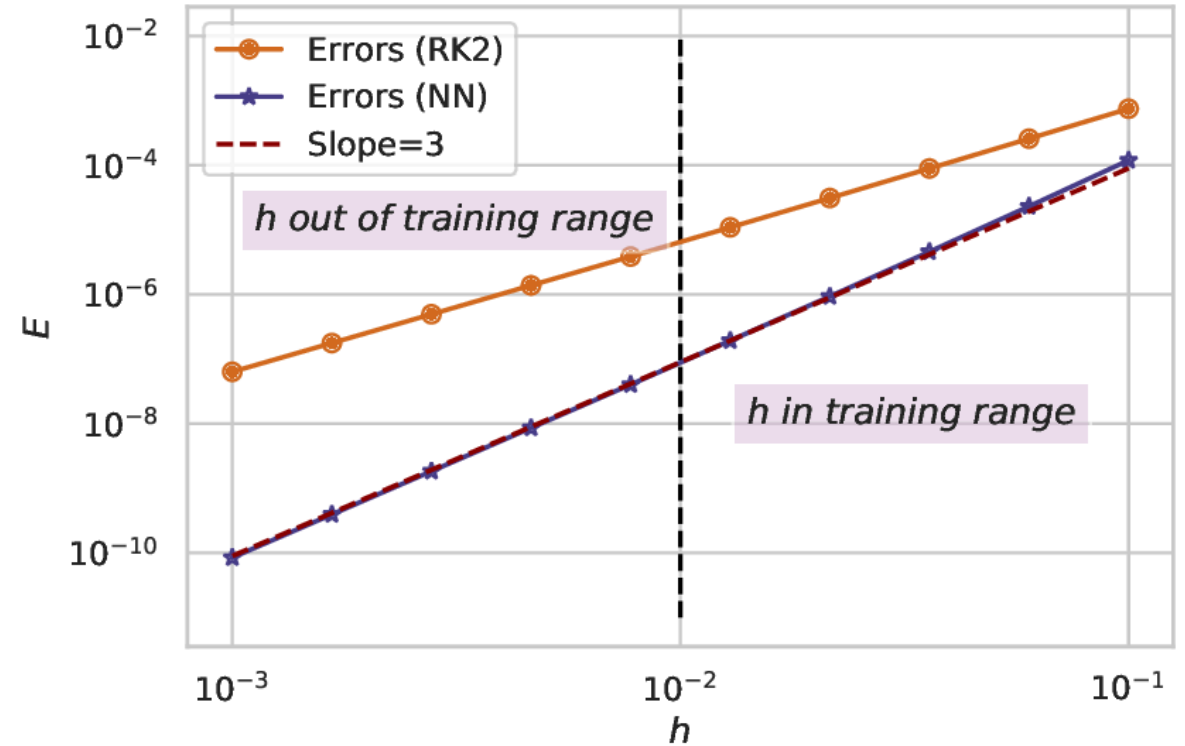
Two stages cannot meet both 3rd-order conditions → **order barrier**.

Ours: for $f=-y^2$, both 3rd-order terms are $\propto y^4$, so they merge:

$$\begin{aligned} y(t_n+h) &= y - y^2 h + y^3 h^2 - y^4 h^3 + O(h^4) \\ \hat{y}_{n+1} &= y - (\theta_{c1} + \theta_{c2}) y^2 h + 2\theta_{c2}\theta_1 y^3 h^2 \\ &\quad - \theta_{c2}\theta_1^2 y^4 h^3 + O(h^4) \end{aligned}$$

h^3 : $\theta_{c2}\theta_1^2=1$ (solvable); with

$$\theta_{c2}\theta_1 = \frac{1}{2} \Rightarrow \theta_1=2, \quad \theta_{c2}=\frac{1}{4}, \quad \theta_{c1}=\frac{3}{4}.$$



2-stage RK-NN attains slope-3 error.

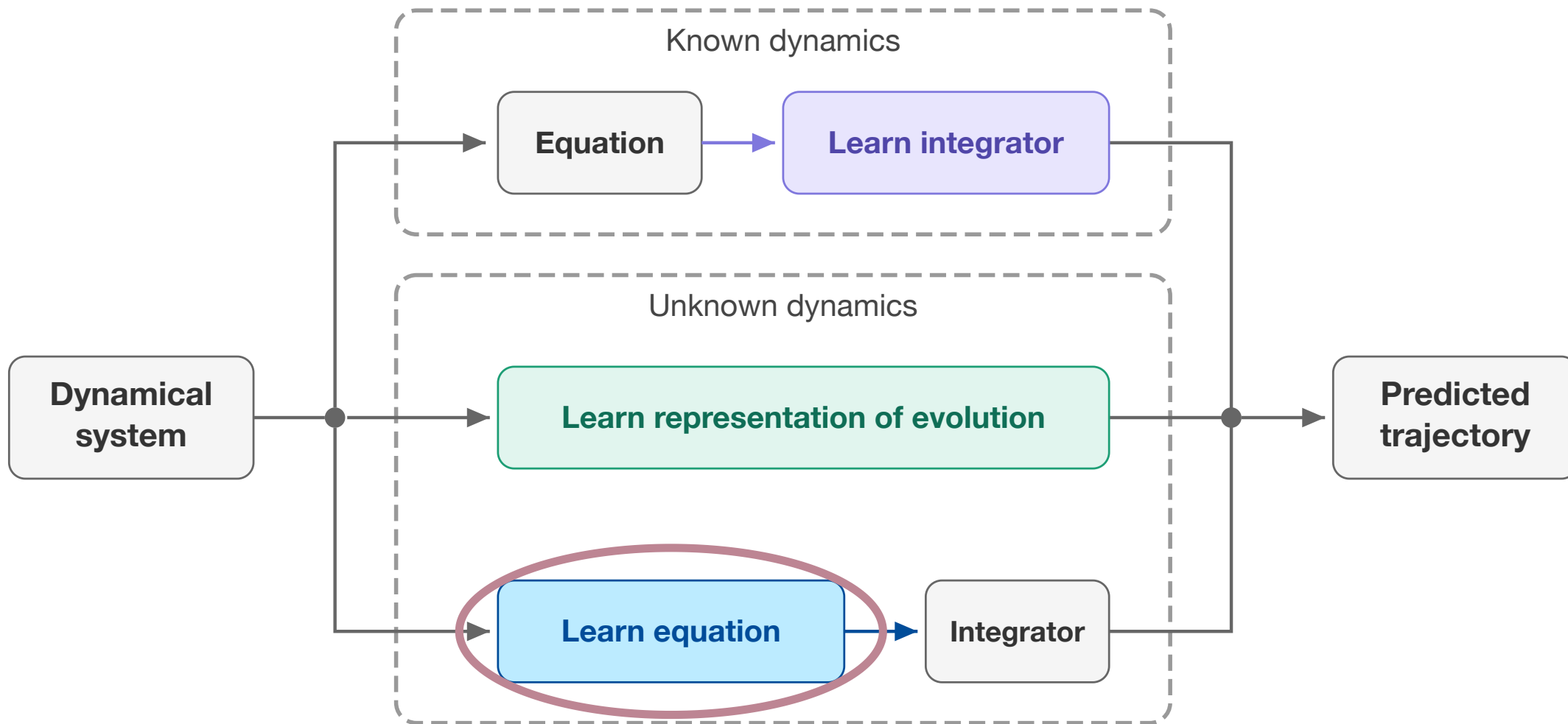
PART III · LEARNING THE CONTINUOUS DESCRIPTION

Integrator-aware Learning

MANUSCRIPT “*Integrator-aware Learning of Modified Equations*”

in process

Yue Guo, Aiqing Zhu, Haotian Jiang, Qianxiao Li.



Cheap forces, but tiny steps

ModEqLearn

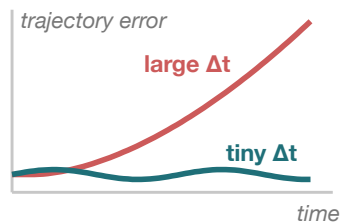
ethanol

naphthalene

aspirin

ML potentials make the **force per step** cheap

Molecular motion needs a **tiny Δt** (fs), so a long run takes lots of steps.



Take a **larger Δt** to go faster, and per-step error **accumulates** over the trajectory.

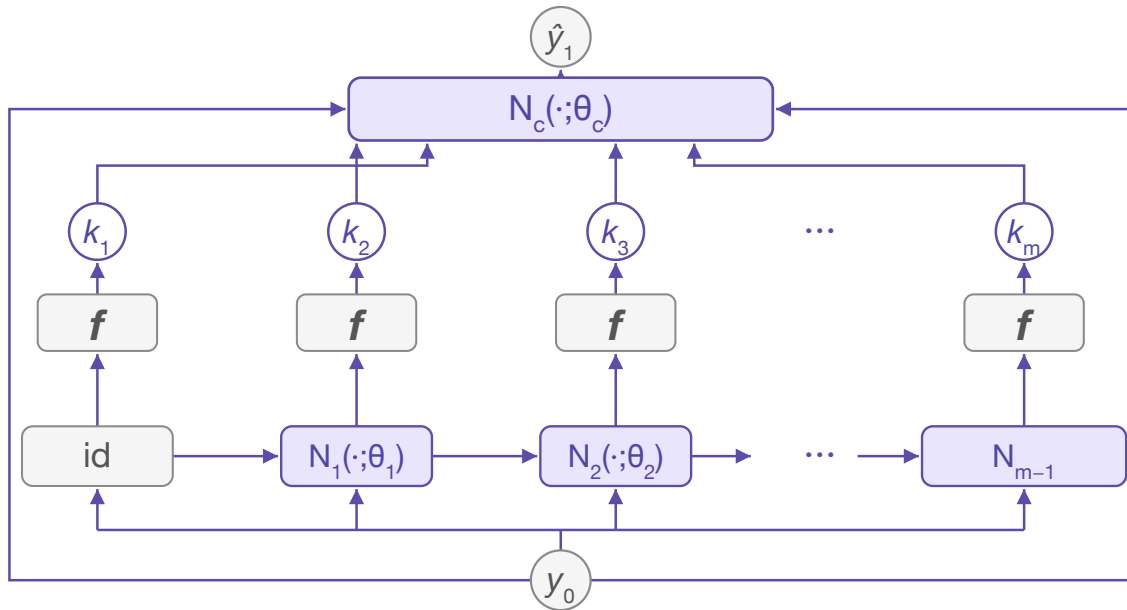
Driving question: can we learn forces accurate at a **larger Δt** ?

Key idea ①: Duality with learning the integrator

Last Project

Given a system, learn the **integrator**

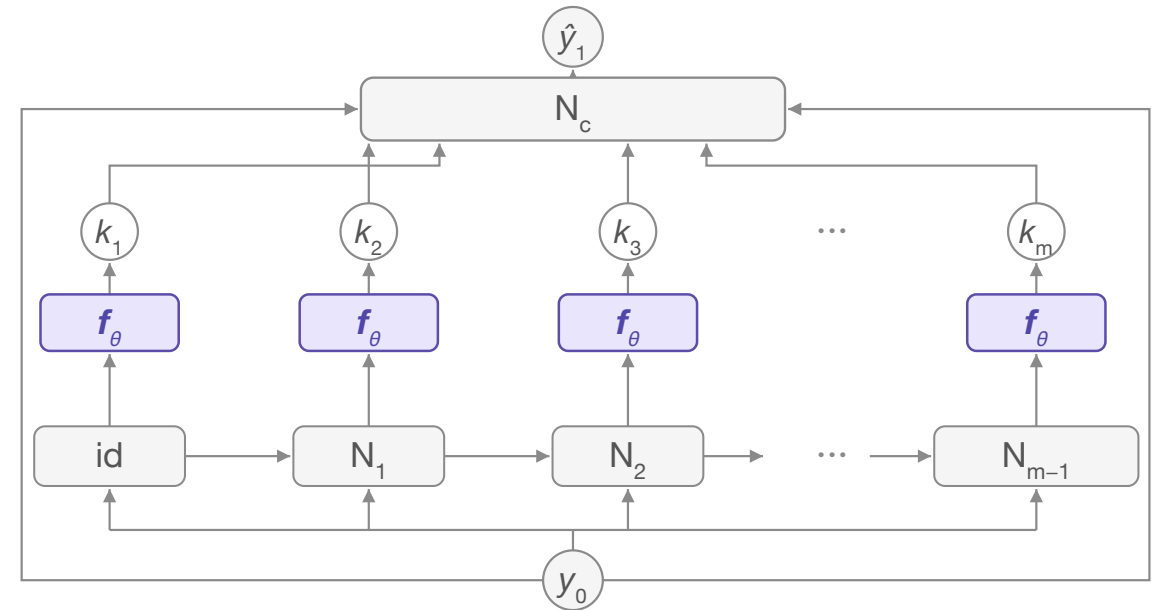
$$\min_{\Phi} \text{err}(\Phi_h, \varphi_h, \mathbf{f}_{\text{true}})$$



Current Project

Given an integrator, learn the **equation**

$$\min_{\mathbf{f}_{\theta}} \text{err}(\Phi_h, \mathbf{f}_{\theta}, \varphi_h, \mathbf{f}_{\text{true}})$$



Key idea ②: modified equations

- **Backward error analysis:** a numerical integrator exactly solves a **modified** differential equation whose field depends on the integrator **and** the step size h .

$$\Phi_{h, \mathbf{f}_{true}}(\mathbf{y}) = \varphi_{h, \mathcal{M}_h(\mathbf{f}_{true})}(\mathbf{y}), \quad \mathcal{M}_h(\mathbf{f}_{true}) = \mathbf{f}_{true} + h \mathbf{f}_{true}^{(1)} + h^2 \mathbf{f}_{true}^{(2)} + \dots$$

- **Inverse modified equation:** Given the true field $\mathbf{f}_{true}$ we want the **discrete map** to reproduce, find the field $\mathbf{f}_\theta$ to feed the integrator so that its numerical flow **exactly** matches the true flow.

$$\Phi_{h, \mathbf{f}_\theta}(\mathbf{y}) = \varphi_{h, \mathcal{M}_h(\mathbf{f}_\theta)}(\mathbf{y}) \iff \varphi_{h, \mathbf{f}_{true}}(\mathbf{y}) \approx \Phi_{h, \mathcal{M}_h^{-1}(\mathbf{f}_{true})}(\mathbf{y})$$

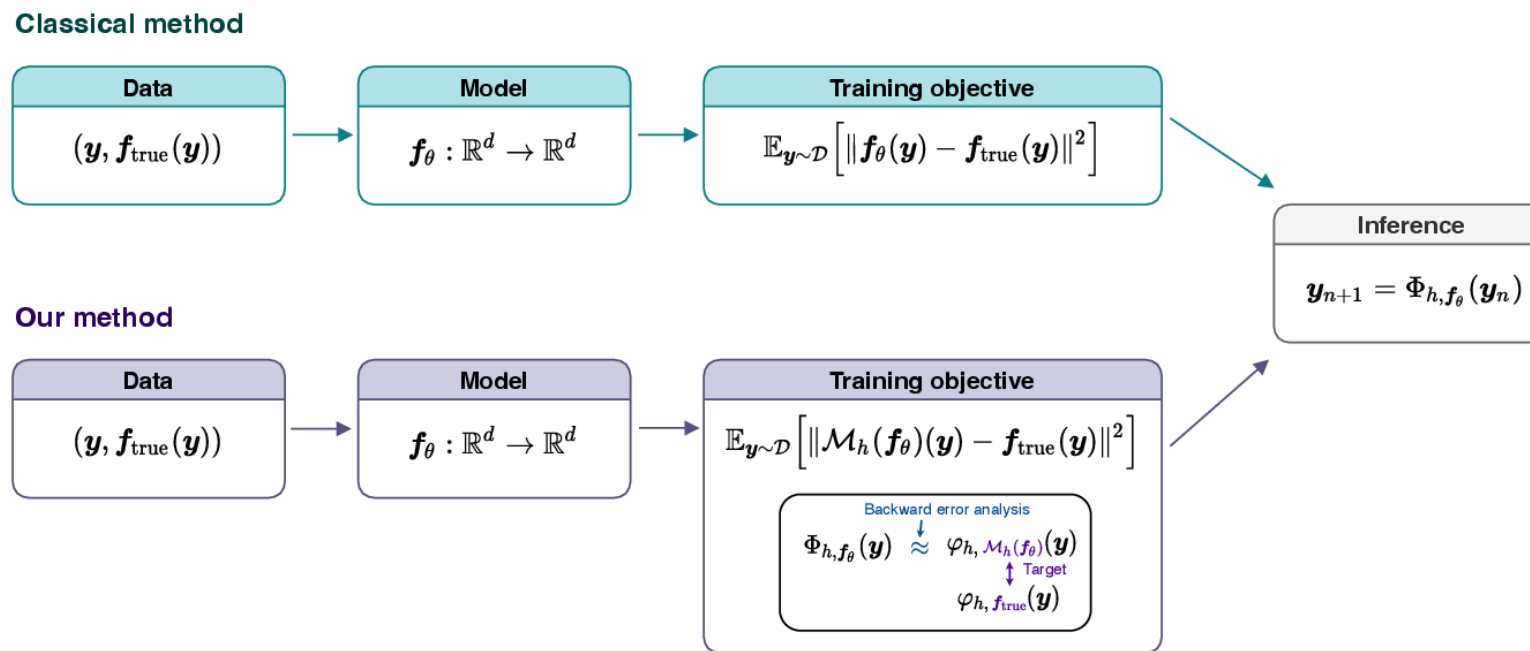
The integrator's “true” target $\mathbf{f}_\theta$ is $\mathcal{M}_h^{-1}(\mathbf{f}_{true})$, not $\mathbf{f}_{true}$.

Method: ModEqLearn

Train the **modified equation** of $\mathbf{f}_\theta$ matches $\mathbf{f}_{\text{true}}$:

$$\min_{\theta} \mathbb{E}_{\mathbf{y} \sim \mathcal{D}} \left\| \mathcal{M}_h(\mathbf{f}_\theta)(\mathbf{y}) - \mathbf{f}_{\text{true}}(\mathbf{y}) \right\|^2$$

Data: states $\mathbf{y} \sim \mathcal{D}$ sampled over the domain, each paired with the true field $\mathbf{f}_{\text{true}}(\mathbf{y})$.



Classical field matching vs. modified-equation-aware training.

Case study: molecular dynamics

Hamiltonian · leapfrog

Hamiltonian System positions $\mathbf{q}$, momenta $\mathbf{p}$ dynamics from partials of H :

$$\begin{aligned} H(\mathbf{q}, \mathbf{p}) &= \frac{1}{2} \mathbf{p}^\top \mathbf{M}^{-1} \mathbf{p} + V(\mathbf{q}) \\ \dot{\mathbf{q}} &= \partial_{\mathbf{p}} H = \mathbf{M}^{-1} \mathbf{p} \\ \dot{\mathbf{p}} &= -\partial_{\mathbf{q}} H = -\nabla V(\mathbf{q}) \end{aligned}$$

Symplectic Leapfrog (Störmer–Verlet)

$$\begin{aligned} \mathbf{p}_{n+\frac{1}{2}} &= \mathbf{p}_n - \frac{h}{2} \nabla V(\mathbf{q}_n) \\ \mathbf{q}_{n+1} &= \mathbf{q}_n + h \mathbf{M}^{-1} \mathbf{p}_{n+\frac{1}{2}} \\ \mathbf{p}_{n+1} &= \mathbf{p}_{n+\frac{1}{2}} - \frac{h}{2} \nabla V(\mathbf{q}_{n+1}) \end{aligned}$$

Molecular ML learns the potential V_θ , $H_\theta = T(\mathbf{p}) + V_\theta(\mathbf{q})$, field $\mathbf{f}_\theta = (\partial_{\mathbf{p}} H_\theta, -\partial_{\mathbf{q}} H_\theta) = (\mathbf{M}^{-1} \mathbf{p}, -\nabla V_\theta)$:

Classical MACE: pointwise fit of V_θ to energies & forces:

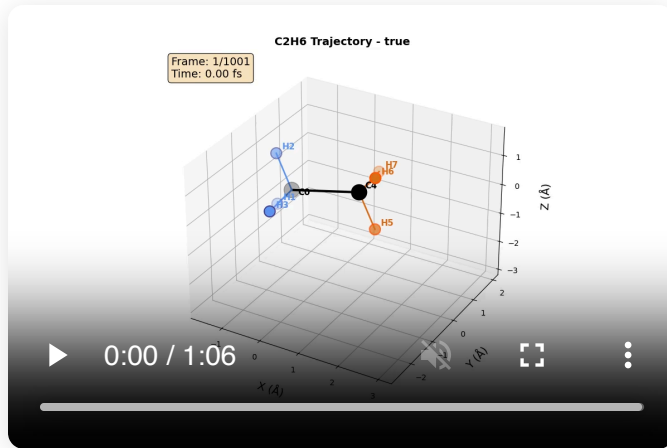
$$\min_{\theta} \sum (\|E - V_\theta\|^2 + \|\mathbf{F} + \nabla V_\theta\|^2)$$

Ours (ModEqLearn): leapfrog's **modified Hamiltonian** of H_θ :

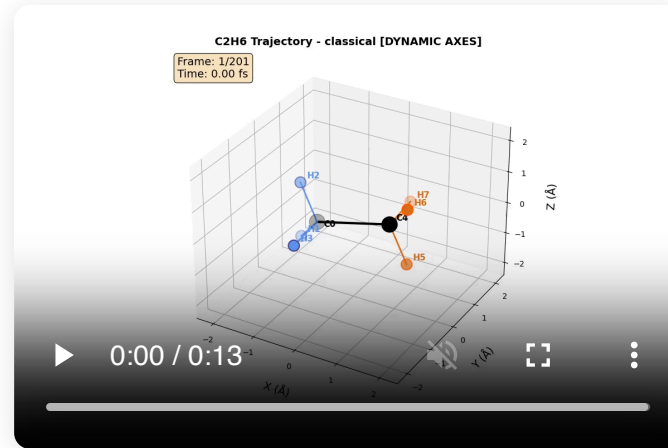
$$\begin{aligned} \tilde{H}_\theta^{(2)} &= T(\mathbf{p}) + V_\theta + \frac{h^2}{24} (2 \mathbf{v}^\top \nabla^2 V_\theta \mathbf{v} - \|\nabla V_\theta\|^2), \quad \mathbf{v} = \mathbf{M}^{-1} \mathbf{p} \\ \min_{\theta} \mathbb{E} \|\mathcal{M}_h(\mathbf{f}_\theta) - \mathbf{f}_{\text{true}}\|^2 &= \mathbb{E} \|\nabla_{(\mathbf{q}, \mathbf{p})} \tilde{H}_\theta^{(2)} - \nabla_{(\mathbf{q}, \mathbf{p})} H_{\text{true}}\|^2 \end{aligned}$$

A perfect **pointwise** V_θ may **drift** in leapfrog at large step size → **the gap our work closes.**

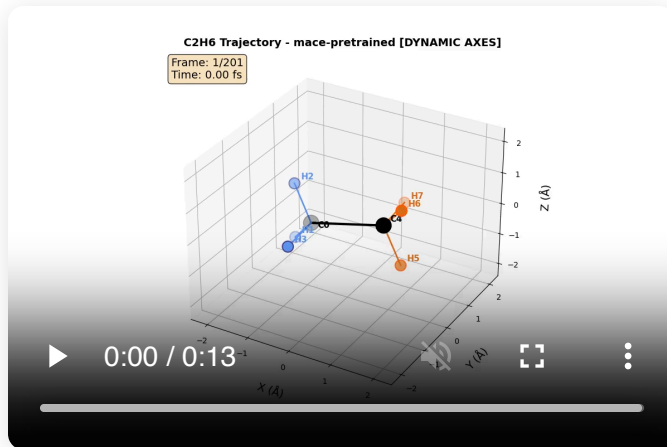
Results: Stable long-time integration



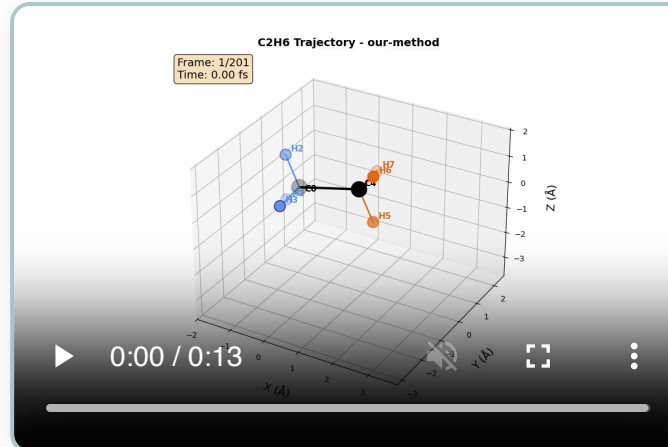
Ground truth



Classical

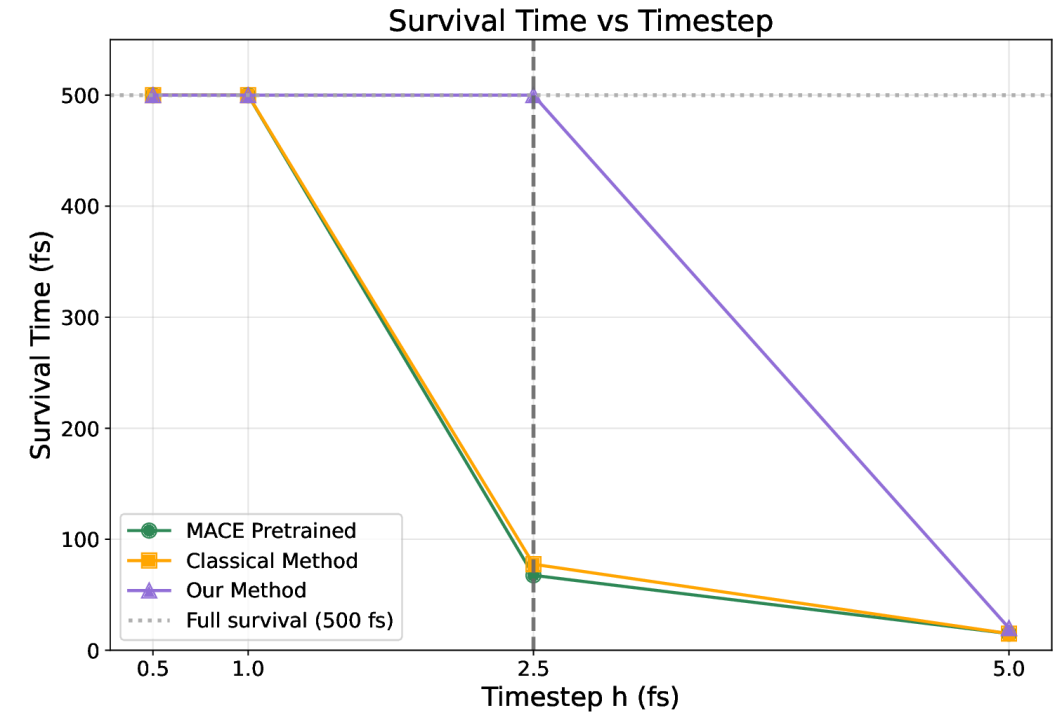


MACE pretrained



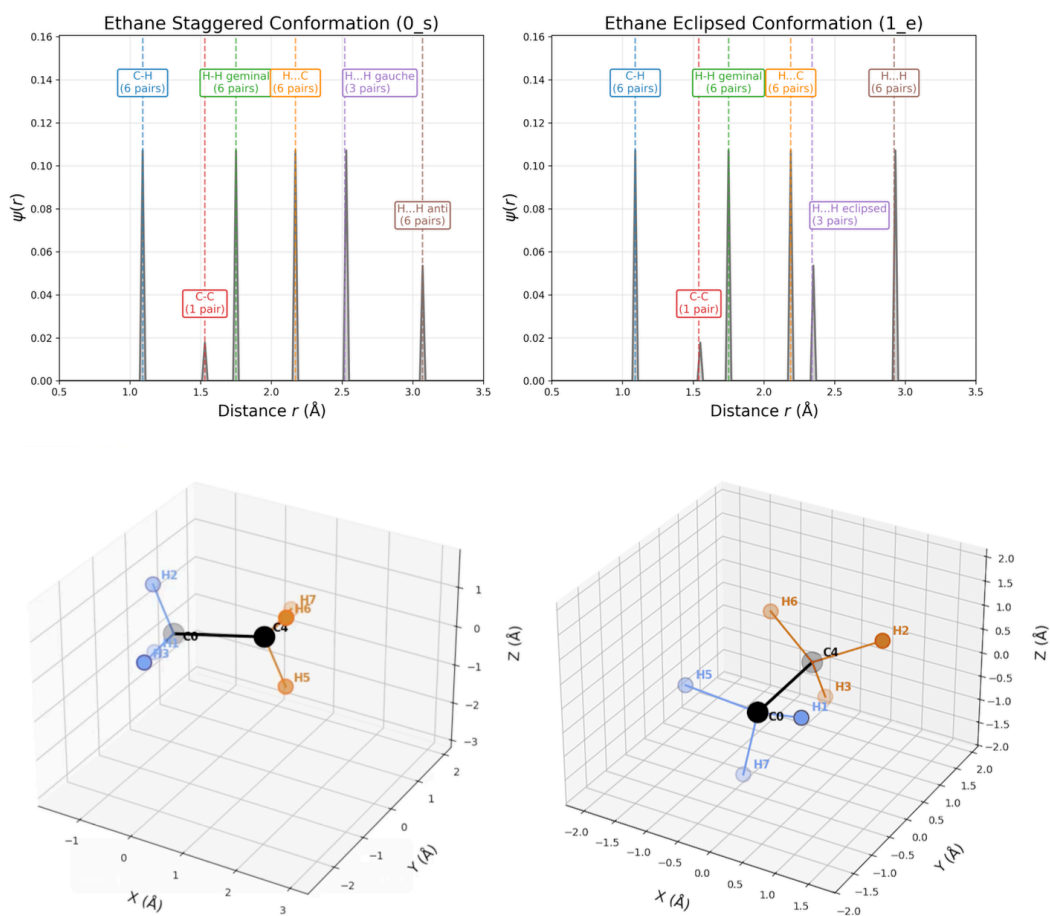
Our method

Ethane (C₂H₆): rollout trajectories at large step size.



Stable rollout survival vs. step size h .

Structural metric: bond-distance distribution



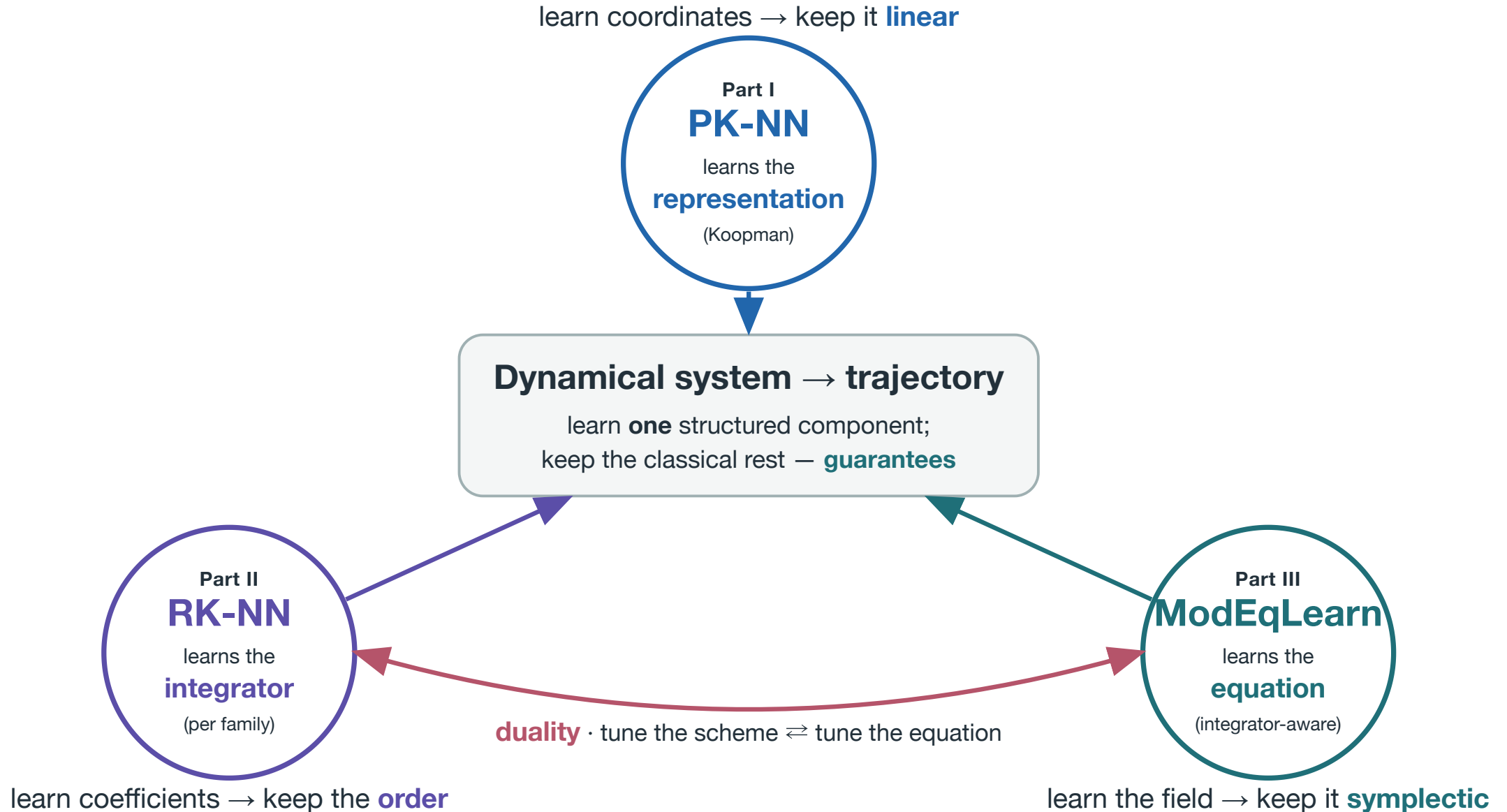
Staggered and eclipsed conformations define distinct interatomic-distance fingerprints.

h (fs)	Method	50 fs	100 fs	200 fs	500 fs
0.5	MACE pretrained	<10 ⁻⁴	<10 ⁻⁴	<10 ⁻⁴	0.107
	Classical	<10 ⁻⁴	<10 ⁻⁴	0.071	0.071
	Our method	0.107	0.214	0.143	0.143
1.0	MACE pretrained	<10 ⁻⁴	<10 ⁻⁴	0.036	0.071
	Classical	<10 ⁻⁴	<10 ⁻⁴	0.107	<10 ⁻⁴
	Our method	0.107	0.214	0.107	0.143
2.5	MACE pretrained	0.071	---	---	---
	Classical	---	---	---	---
	Our method	0.036	0.143	0.071	0.179
5.0	MACE pretrained	---	---	---	---
	Classical	---	---	---	---
	Our method	---	---	---	---

Table 4.2: $\Delta\psi$ for eclipsed conformation; lower is better, --- indicates blowup.

$$\psi(r) = \frac{1}{N(N-1)} \sum_{i=1}^N \sum_{j \neq i}^N \delta(r - \|\mathbf{x}_i - \mathbf{x}_j\|)$$
$$\Delta\psi = \int_{r_{\min}}^{r_{\max}} |\psi_{\text{pred}}(r) - \psi_{\text{true}}(r)| dr$$

Summary



Acknowledgements

