

# Data-efficient Kernel Methods for Learning Differential Equations and their Solution Operators

Bamdad Hosseini  
Department of Applied Mathematics  
University of Washington

Data Assimilation and Inverse Problems for Digital Twins Workshop  
Institute for Mathematical and Statistical Innovation

Oct 8, 2025

*New talk ...*  
*new typos ...*

---

*Matthew Thorpe*

# Acknowledgements

Yasamin Jalalian, Juan Felipe Osorio Ramirez, Alexander Hsu, Bamdad Hosseini, and Houman Owhadi, **Data-Efficient Kernel Methods for Learning Differential Equations and Their Solution Operators: Algorithms and Error Analysis**,  
arXiv:2503.01036, 2025

- AH, JFOP, and BH were supported by
  - NSF-DMS-2208535, "Machine Learning for Bayesian Inverse Problems"
  - NSF-DMS-2337678, "CAREER: Gaussian Processes for Scientific Machine Learning: Theoretical Analysis and Computational Algorithms"
- YJ and HO were supported by
  - NSF-GEO-2425908, "CAIG: Discovering the Law of Stress Transfer and Earthquake Dynamics in a Fault Network using a Computational Graph Discovery Approach"
  - FA9550-20-1-0358, "Machine Learning and Physics-Based Modeling and Simulation"
  - FOA-AFRL-AFOSR-2023-0004, "Mathematics of Digital Twins"
  - DE-SC0023163, "SEA-CROGS: Scalable, Efficient, and Accelerated Causal Reasoning Operators, Graphs and Spikes for Earth and Embedded Systems"
  - Jet Propulsion Laboratory PDRDF 24AW0133

# Outline

- ① Overview
- ② Algorithms
- ③ Theory

# Overview

# Solving PDEs / Learning PDEs / Operator Learning

Generic nonlinear PDE

$$\mathfrak{P}(u) = f$$

- (PDE solvers) Given  $\mathfrak{P}, f$  find  $\hat{u}$  such that  $\mathfrak{P}(\hat{u}) \approx f$

$$\widehat{\mathfrak{P}^{-1}} : f \mapsto u$$

- (Learning PDEs) Given training data  $\{u_m, f_m\}_{m=1}^M$  find  $\hat{\mathfrak{P}}$  such that  $\hat{\mathfrak{P}}(u) \approx f$

$$\hat{\mathfrak{P}} : u \mapsto f$$

- (Operator learning) Given training data  $\{u_m, f_m\}_{m=1}^M$  find  $\widehat{\mathfrak{P}^{-1}}$  such that  $\widehat{\mathfrak{P}^{-1}}(f) \approx u$

$$\widehat{\mathfrak{P}^{-1}} : f \mapsto u$$

---

<sup>1</sup>Brunton, Proctor, and Kutz, “Discovering governing equations from data by sparse identification of nonlinear dynamical systems”.

<sup>2</sup>Kovachki, Lanthaler, and Stuart, “Operator learning: Algorithms and analysis”.

- Sparse observation meshes  
 $X_m := \{x_{m,1}, \dots, x_{m,N}\} \subset \Omega$

## Goal

Given sparse and noisy observations

$$\{u_m(X_m) + \epsilon_m, f_m\}_{m=1}^M$$

learn  $\hat{\mathfrak{P}} \approx \mathfrak{P}$  and  $\hat{\mathfrak{P}}^{-1} \approx \mathfrak{P}^{-1}$ .

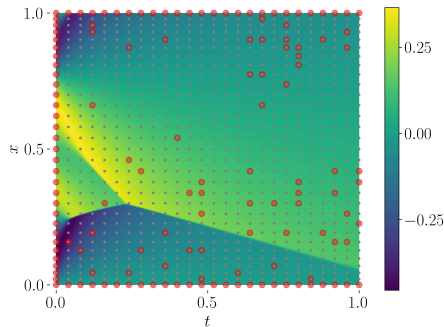
## Idea

First learn  $\hat{\mathfrak{P}}$  assuming it is local. Then  
"invert"  $\hat{\mathfrak{P}}$  to obtain  $\hat{\mathfrak{P}}^{-1}$ .

E.g. Burgers' PDE

$$\mathfrak{P}(u) = u_t + uu_x - \nu u_{xx} = 0$$

**While  $\mathfrak{P}$  is simple, the solution map  $\mathfrak{P}^{-1} : u(x, 0) \mapsto u(x, t)$  is complex.**



## Problem Setup

$$\mathfrak{P}(u)(x) = f(x) \quad x \in \Omega$$

$$\begin{aligned}\mathfrak{P}(u)(x) &= p \circ (x, Du(x), D^2u(x), \dots, ) \\ &= p \circ \Phi(u, x)\end{aligned}$$

- **(known)** Mapping  $\Phi : \mathcal{U} \times \Omega \rightarrow \mathbb{R}^Q$ , linear in  $u$  and nonlinear in  $x$
- **(unknown)** Nonlinear function  $p : \mathbb{R}^Q \rightarrow \mathbb{R}$

E.g. Nonlinear, Variable Coefficient, Elliptic PDE

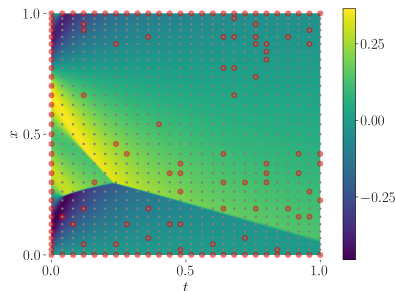
$$\mathfrak{P}(u) = -\partial_x (a(x)\partial_x u) + \alpha u^3 = f$$

$$\Phi(u, x) = (x, u(x), u_x(x), u_{xx}(x))$$

$$p(s_1, s_2, s_3, s_4) = -a_x(s_1)s_3 - a(s_1)s_4 + \alpha s_2^3$$

# Kernel Equation Learning (KEqL): Formulation

- Assumption  $\mathfrak{P} = p \circ \Phi(u, x) = f(x)$
- Training data  $(u_m(X_m), f_m)_{m=1}^M$
- Banach spaces  $(\mathcal{U}, \|\cdot\|_{\mathcal{U}})$  and  $(\mathcal{P}, \|\cdot\|_{\mathcal{P}})$
- Optimal recovery problem



$$(\hat{u}_m, \hat{p}) = \arg \min_{v_m \in \mathcal{U}, q \in \mathcal{P}} \|q\|_{\mathcal{P}} + \lambda \sum_{m=1}^M \|v_m\|_{\mathcal{U}}$$

s.t.

$$v_m(X_m) = u_m(X_m)$$

(Observation/data)

$$q \circ \Phi(v_m, x) = f_m(x), \quad \forall x \in \Omega$$

(PDE constraint)

# Reproducing Kernel Hilbert Spaces (RKHS)

- Set  $\Omega$
- Positive definite and symmetric (PDS) kernel  
 $K : \Omega \times \Omega \rightarrow \mathbb{R}$ :
  - $K(x, x') = K(x', x), \forall x, x' \in \Omega$
  - For any set of points  $X = \{x_1, \dots, x_N\} \subset \Omega$  the matrix  $K(X, X) = (K(x_i, x_j))$  is PDS.
- Pre-Hilbert space

$$\mathcal{K}_0 := \left\{ f : \Omega \rightarrow \mathbb{R} \mid f = \sum_{j=1}^{N_f} \alpha_{f,j} K(\cdot, x_{f,j}) = K(\cdot, X_f) \alpha_f \right\}$$

- Define inner product

$$\langle f, g \rangle_{\mathcal{K}_0} := \alpha_f^T K(X_f, X_g) \alpha_g$$

- Complete  $\mathcal{K}_0$  to get RKHS  $\mathcal{K}$
- $\mathcal{K}$  is uniquely defined by  $K$
- Reproducing property:  
 $\forall f \in \mathcal{K}$

$$f(x) = \langle f, K(\cdot, x) \rangle_{\mathcal{K}}$$

- In fact, for any  $\phi \in \mathcal{K}^*$

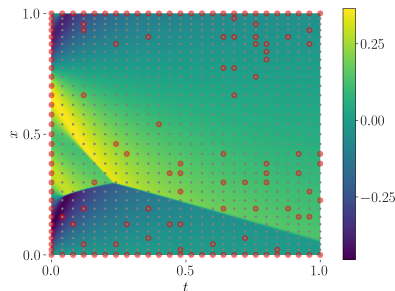
$$\phi(f) = \langle f, K(\cdot, \phi) \rangle_{\mathcal{K}}$$

where

$$K(x, \phi) = \phi(K(x, \cdot))$$

# Kernel Equation Learning (KEqL): Formulation

- Assumption  $\mathfrak{P} = p \circ \Phi(u, x) = f(x)$
- Training data  $(u_m(X_m), f_m)_{m=1}^M$
- RKHS space  $\mathcal{U}$  with kernel  $U$
- RKHS space  $\mathcal{P}$  with kernel  $P$
- Collocation mesh  $X \supseteq \cup_m X_m$



$$(\hat{u}_m, \hat{p}) = \arg \min_{v_m \in \mathcal{U}, q \in \mathcal{P}} \|q\|_{\mathcal{P}}^2 + \lambda \sum_{m=1}^M \|v_m\|_{\mathcal{U}}^2$$

s.t.

$$v_m(X_m) = u_m(X_m)$$

(Observation/data)

$$q \circ \Phi(v_m, X) = f_m(X)$$

(Discrete PDE constraint)

# KEqL: All-at-once Inversion

$$\begin{aligned} (\hat{u}_m, \hat{p}) = \arg \min_{v_m \in \mathcal{U}, q \in \mathcal{P}} & \|q\|_{\mathcal{P}}^2 + \lambda \sum_{m=1}^M \|v_m\|_{\mathcal{U}}^2 \\ \text{s.t.} \quad & v_m(X_m) = u_m(X_m) \quad (\text{Observation/data}) \\ & q \circ \Phi(v_m, X) = f_m(X) \quad (\text{Discrete PDE constraint}) \end{aligned}$$

---

<sup>3</sup>Chen, Liu, and Sun, “Physics-informed learning of governing equations from scarce data”.

<sup>4</sup>Sun, Liu, and Sun, “Physics-informed Spline Learning for Nonlinear Dynamics Discovery”.

<sup>5</sup>Haber and Oldenburg, “Joint inversion: a structural approach”.

<sup>6</sup>Kaltenbacher, “Regularization based on all-at-once formulations for inverse problems”.

# KEqL: Algorithm Summary

- Relax equality constraints

$$(\hat{u}_m, \hat{p}) = \arg \min_{v_m \in \mathcal{U}, q \in \mathcal{P}} \|q\|_{\mathcal{P}}^2 + \sum_{m=1}^M \lambda \|v_m\|_{\mathcal{U}}^2 \\ + \lambda' \|u_m(X^m) - v_m(X^m)\|_2^2 + \lambda'' \|q \circ \Phi(v_m, X) - f_m(X)\|_2^2$$

- Apply kernel trick (representer theorem) or use feature maps
- Solve using a Levenberg–Marquardt-type algorithm

# Operator Learning via KEqL

- Given KEqL solution  $\hat{p}$  approximate PDE solution map  $\mathfrak{P}^{-1}$  with pseudo inverse of

$$\hat{\mathfrak{P}} := \hat{p} \circ \Phi.$$

$$\hat{\mathfrak{P}}^\dagger(f) := \arg \min_{v \in \mathcal{U}} \|v\|_{\mathcal{U}}^2 + \lambda'' \|\hat{p} \circ \Phi(v, X) - f(X)\|_2^2.$$

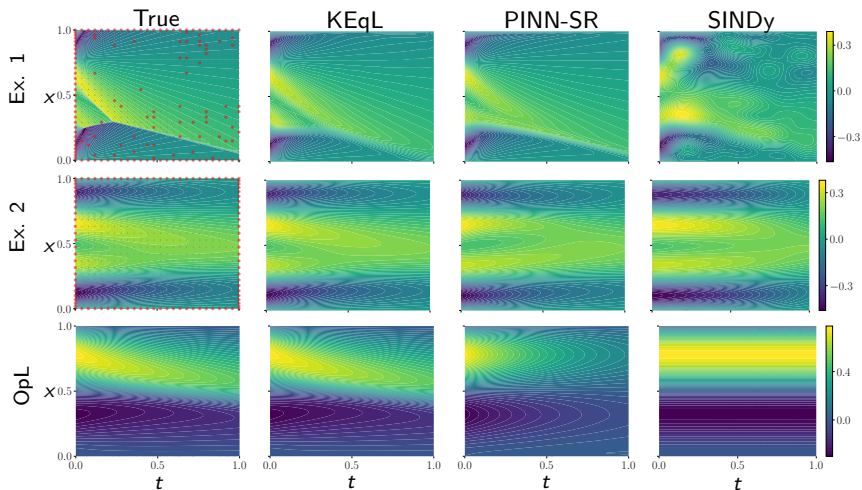
---

<sup>7</sup>Chen et al., “Solving and learning nonlinear PDEs with Gaussian processes”.

<sup>8</sup>Long et al., “A kernel framework for learning differential equations and their solution operators”.

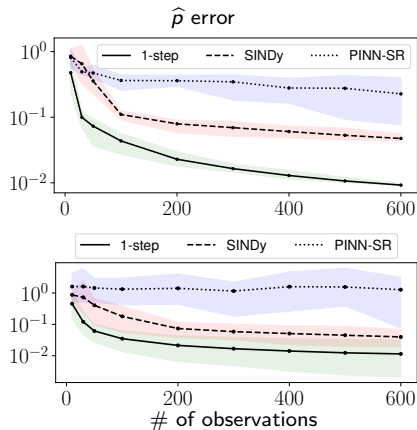
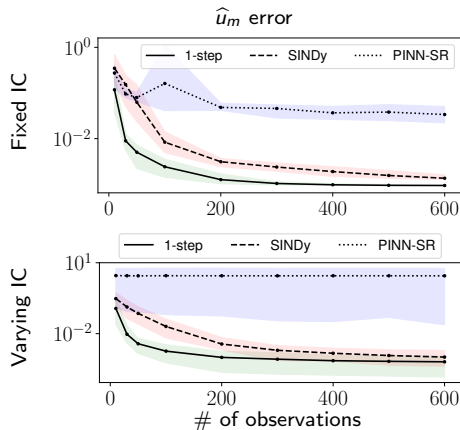
# Burgers' Benchmark: Shocks with Sparse Data

$$\mathfrak{P}(u) = u_t + uu_x - 0.1u_x x = 0 \quad \text{plus B.C and I.C}$$



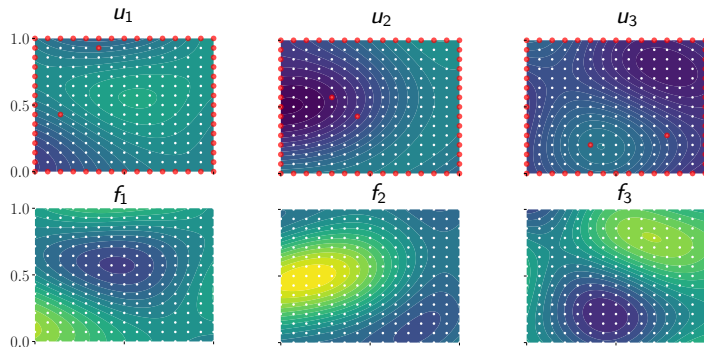
# Burgers' Benchmark: Comparison to PINN-SR and SINDy

- Large performance gap with PINN model



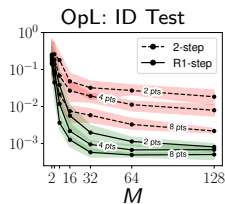
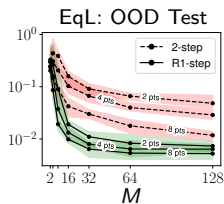
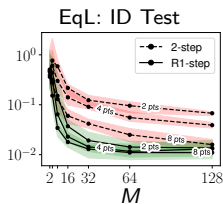
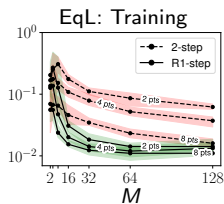
# Darcy Flow PDE: Variable Coefficients with Sparse Data

$$\mathfrak{P}(u) = -\operatorname{div} a \nabla u = f \quad a(x) = \exp(\sin(\cos(x_1) + \cos(x_2))) \quad + \text{B.C.}$$



# Darcy Flow PDE: EqL and OpL Errors

- $M$  – # of solution and source pairs  $(u_m(X), f_m)_{m=1}^M$
- 2-step – Methods like SINDy that infer  $u_m$  first before learning  $p$
- R1-step – KEqL with reduced basis for better performance



# Quantitative Error Analysis

## Theorem [JOHHO] (Part 1)

Consider  $\{u_m(X), f_m\}_{m=1}^M$  defined on a smooth domain  $\Omega \subset \mathbb{R}^d$  and observation mesh  $X_{\text{obs}} \subset \Omega$  with  $|X_{\text{obs}}| = N$ . Suppose that  $u^m \in H^\gamma(\Omega)$  for  $\gamma > d/2 + \text{order of } \mathfrak{P}$  and  $\mathfrak{P} = p \circ \Phi$  such that  $\Phi(u, x) \in \mathbb{R}^Q$ . Define the fill distance

$$\rho := \sup_{x' \in \Omega} \inf_{x \in X_{\text{obs}}} |x - x'|,$$

Then under sufficient smoothness assumptions it holds for  $0 \leq \gamma' \leq \gamma$

$$\sum_{m=1}^M \|u_m - \hat{u}_m\|_{H^{\gamma'}(\Omega)}^2 \lesssim \rho^{2(\gamma-\gamma')} \left( \|p\|_{\mathcal{P}}^2 + \sum_{m=1}^M \|u_m\|_{\mathcal{U}}^2 \right).$$

---

<sup>9</sup>Zhang and Schaeffer, “On the convergence of the SINDy algorithm”.

<sup>10</sup>Scholl et al., “The uniqueness problem of physical law learning”.

<sup>11</sup>He, Zhao, and Zhong, “How much can one learn a partial differential equation from its solution?”

<sup>12</sup>Boullé, Halikias, and Townsend, “Elliptic PDE learning is provably data-efficient”.

# Quantitative Error Analysis

## Theorem [JOHHO] (Part 2)

Consider  $\{u_m(X), f_m\}_{m=1}^M$  defined on a smooth domain  $\Omega \subset \mathbb{R}^d$  and observation mesh  $X_{\text{obs}} \subset \Omega$  with  $|X_{\text{obs}}| = N$ . Suppose that  $u^m \in H^\gamma(\Omega)$  for  $\gamma > d/2 + \text{order of } \mathfrak{P}$  and  $\mathfrak{P} = p \circ \Phi$  such that  $\Phi(u, x) \in \mathbb{R}^Q$  and  $p \in H^\eta(\mathbb{R}^Q)$  for  $\eta > Q/2$ . Define the set  $S := \bigcup_{m=1}^M \bigcup_{x \in X_{\text{obs}}} \Phi(u_m, x)$  along with the fill distances

$$\rho := \sup_{x' \in \Omega} \inf_{x \in X_{\text{obs}}} |x - x'|,$$

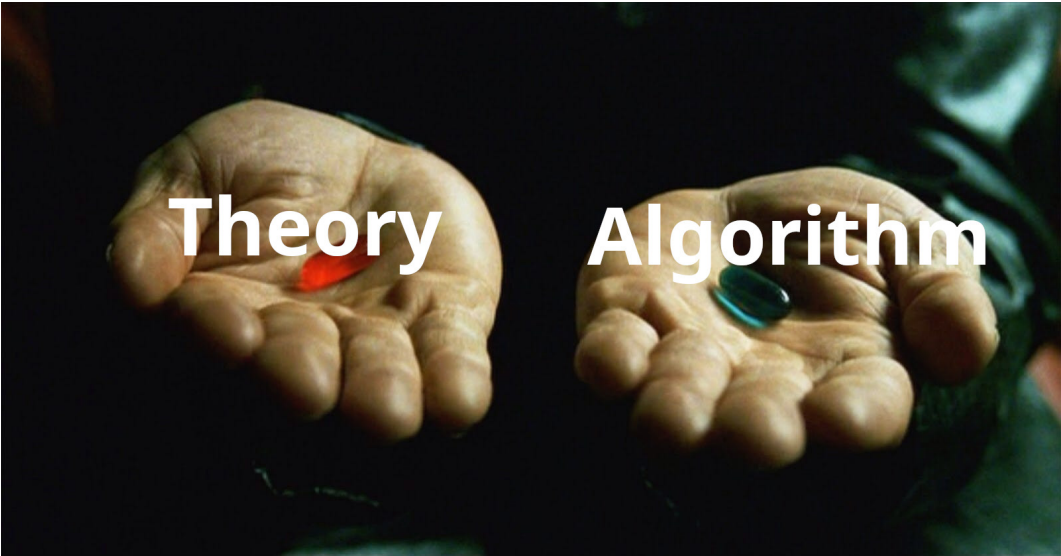
$$\varrho(B) := \sup_{s' \in B} \inf_{s \in S \cap B} |s - s'|,$$

for any smooth and bounded set  $B \subset \mathbb{R}^Q$ . Then under sufficient smoothness assumptions it holds for  $d/2 + \text{order of } \mathfrak{P} \leq \gamma' \leq \gamma$

$$\sum_{m=1}^M \|p - \hat{p}\|_{L^\infty(B)} \lesssim \left[ \rho^{\gamma-\gamma'} + \varrho(B)^{\eta-Q/2} \right] \left( \|p\|_{\mathcal{P}} + \sum_{m=1}^M \|u_m\|_{\mathcal{U}} \right).$$

# Summary

- A kernel method for learning PDEs and filtering solutions
- Joint/simultaneous recovery of solutions and PDE form
- Successful recovery of PDEs and their solution maps with very scarce data
- Optimization problem is amenable to quasi-newton algorithms, easier and more efficient training
- Better performance compared to some neural net models or two step learning
- Operator learning through equation learning, going well beyond typical operator learning setups
- Quantitative convergence analysis, reminiscent of Sobolev sampling inequalities for scattered data approximation

A close-up photograph of two human hands, palms up, holding small, oval-shaped pills. The left hand holds a red pill, and the right hand holds a green pill. The background is dark and out of focus. The word 'Theory' is overlaid in white text on the left hand, and the word 'Algorithm' is overlaid in white text on the right hand.

**Theory**

**Algorithm**

- ① Overview
- ② Algorithms
- ③ Theory

# Algorithms

# RKHS Representer Theorems

$$\hat{f} = \arg \min_{f \in \mathcal{K}} \|f\|_{\mathcal{K}} \quad \text{s.t.} \quad \phi(f) = \mathbf{z}$$

- Shorthand notation  $\phi = (\phi_1, \dots, \phi_N) \in (\mathcal{K}^*)^N$
- Data  $\mathbf{z} \in \mathbb{R}^N$
- Closed form solution

$$\hat{f} = K(\cdot, \phi) K(\phi, \phi)^{-1} \mathbf{z}, \quad \|\hat{f}\|_{\mathcal{K}}^2 = \mathbf{z}^T K(\phi, \phi)^{-1} \mathbf{z}$$

- Vector field  $K(\cdot, \phi) = (K(\cdot, \phi_1), \dots, K(\cdot, \phi_N)) \in \mathcal{K}^N$
- Matrix  $K(\phi, \phi) \in \mathbb{R}^{N \times N}$

$$K(\phi, \phi)_{ij} = \phi_i(K(\cdot, \phi_j))$$

# Feature Map Perspective

$$\hat{f} = \arg \min_{f \in \mathcal{K}} \|f\|_{\mathcal{K}} \quad \text{s.t.} \quad \phi(f) = \mathbf{z}$$

- Solution

$$\hat{f} = K(\cdot, \phi) K(\phi, \phi)^{-1} \mathbf{z}, \quad \|\hat{f}\|_{\mathcal{K}}^2 = \mathbf{z}^T K(\phi, \phi)^{-1} \mathbf{z}$$

- Coefficient vector  $\alpha = K(\phi, \phi)^{-1} \mathbf{z}$

$$\hat{f} = K(\cdot, \phi) \alpha = \sum_{j=1}^N \alpha_j K(\cdot, \phi_j)$$

- Using reproducing property

$$\|\hat{f}\|_{\mathcal{K}}^2 = \alpha^T K(\phi, \phi) \alpha$$

## Reformulate KEqL

$$(\hat{u}_m, \hat{p}) = \arg \min_{v_m \in \mathcal{U}, q \in \mathcal{P}} \|q\|_{\mathcal{P}}^2 + \sum_{m=1}^M \|v_m\|_{\mathcal{U}}^2 \\ + \|u_m(X^m) - v_m(X^m)\|_2^2 + \|q \circ \Phi(v_m, X) - f_m(X)\|_2^2$$

- Take  $\lambda = \lambda' = \lambda'' = 1$  for simplicity
- Reduced model  $v_m = U(\cdot, X)\alpha_m$  and define  $\alpha = (\alpha_1, \dots, \alpha_M)$
- Recall  $X$  is the collocation mesh
- Define  $S(\alpha) = \bigcup_{m=1}^M \bigcup_{x \in X} \Phi(K(\cdot, X)\alpha, x)$  and  $S(\alpha_m) = \bigcup_{x \in X} \Phi(K(\cdot, X)\alpha_m, x)$
- Take  $q = P(\cdot, S(\alpha))\beta$

$$(\hat{\alpha}, \hat{\beta}) = \arg \min_{\alpha, \beta} \beta^T P(S(\alpha), S(\alpha))\beta + \sum_{m=1}^M \alpha_m^T U(X, X)\alpha_m \\ + \|u_m(X^m) - U(X^m, X)\alpha_m\|_2^2 + \|P(S(\alpha_m), S(\alpha))\beta - f_m(X)\|_2^2$$

## Reformulate KEqL

$$(\hat{\alpha}, \hat{\beta}) = \arg \min_{\alpha, \beta} \beta^T P(S(\alpha), S(\alpha)) \beta + \sum_{m=1}^M \alpha_m^T U(X, X) \alpha_m \\ + \|u_m(X^m) - U(X^m, X) \alpha_m\|_2^2 + \|P(S(\alpha_m), S(\alpha)) \beta - f_m(X)\|_2^2$$

- Compositional structure
- Quadratic in  $\alpha$  for fixed  $\beta$
- Quadratic in  $\beta$  for fixed  $\alpha$
- Propose sequential quadratic approximations
- Additional regularization to control deviation in each step, similar to Levenberg–Marquardt (LM)

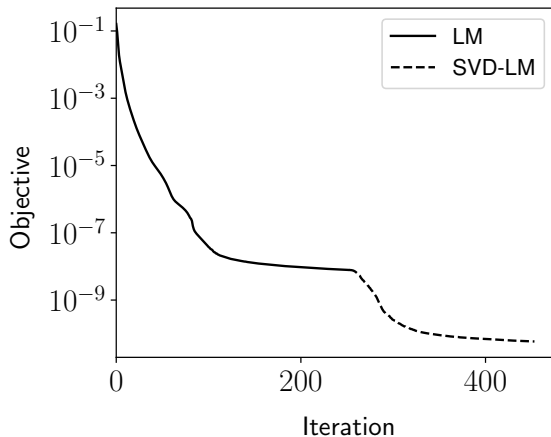
$$(\hat{\alpha}, \hat{\beta}) = \arg \min_{\alpha, \beta} \beta^T P(S(\alpha), S(\alpha)) \beta + \sum_{m=1}^M \alpha_m^T U(X, X) \alpha_m \\ + \|u_m(X^m) - U(X^m, X) \alpha_m\|_2^2 + \|P(S(\alpha_m), S(\alpha)) \beta - f_m(X)\|_2^2$$

Minimizing sequence  $(\alpha^{(k)}, \beta^{(k)})$  given by

$$(\hat{\alpha}^{(k+1)}, \hat{\beta}^{(k+1)}) := \arg \min_{\alpha, \beta} \beta^T P(S(\alpha^{(k)}), S(\alpha^{(k)})) \beta + \sum_{m=1}^M \alpha_m^T U(X, X) \alpha_m \\ + \|u_m(X^m) - U(X^m, X) \alpha_m\|_2^2 \\ + \left\| \left[ P(S(\alpha_m^{(k)}), S(\alpha^{(k)})) + \nabla_{\alpha} P(S(\alpha_m^{(k)}), S(\alpha^{(k)})) (\alpha - \alpha^{(k)}) \right] \beta - f_m(X) \right\|_2^2 \\ + \sigma_k (\alpha - \alpha^{(k)})^T U(X, X) (\alpha - \alpha^{(k)}) + \sigma'_k (\beta - \beta^{(k)})^T P(S(\alpha_m^{(k)}), S(\alpha^{(k)})) (\beta - \beta^{(k)})$$

Choose  $\sigma_k, \sigma'_k$  based on decrease of objective, established heuristics for LM

# LM for KEqL



# Theory

# Sobolev Sampling Inequalities

## Sobolev Sampling Inequality

Suppose  $\Omega \subset \mathbb{R}^d$  is a bounded set with Lipschitz boundary and consider a set of points  $X = \{x_1, \dots, x_N\} \subset \overline{\Omega}$  with fill distance  $h_X := \sup_{x \in \Omega} \inf_{x' \in X} \|x - x'\|_2$ . Let  $u|_X$  denote the restriction of  $u$  to the set  $X$ . Further consider  $\gamma > d/2$  and  $0 \leq \gamma' \leq \gamma$  and let  $u \in H^\gamma(\Omega)$ .

- ① (Noiseless) Suppose  $u|_X = 0$ . Then there exists  $h_0 > 0$  so that whenever  $h_X \leq h_0$  we have  $\|u\|_{H^{\gamma'}(\Omega)} \leq C_\Omega h_X^{\gamma-\gamma'} \|u\|_{H^\gamma(\Omega)}$ , where  $C_\Omega > 0$  is a constant that depends only on  $\Omega$ .
- ② (Noisy) Suppose  $u|_X \neq 0$ . Then there exists  $h_0 > 0$  so that whenever  $h_X \leq h_0$  we have  $\|u\|_{L^\infty(\Omega)} \leq C_\Omega h_X^{\gamma-d/2} \|u\|_{H^\gamma(\Omega)} + 2\|u|_X\|_\infty$ , where  $C_\Omega > 0$  is a constant that depends only on  $\Omega$ .

# Controlling the Filtering Error

$$(\hat{u}_m, \hat{p}) = \arg \min_{v_m \in \mathcal{U}, q \in \mathcal{P}} \|q\|_{\mathcal{P}}^2 + \sum_{m=1}^M \|v_m\|_{\mathcal{U}}^2$$

$$\text{s.t.} \quad v_m(X_{\text{obs}}) = u_m(X_{\text{obs}}) \quad (\text{Observation/data})$$

$$q \circ \Phi(v_m, X_{\text{obs}}) = f_m(X_{\text{obs}}) \quad (\text{Discrete PDE constraint})$$

- Apply Sobolev sampling inequality, recalling  $\rho = \sup_{x' \in \Omega} \inf_{x \in X_{\text{obs}}} |x - x'|$

$$\|\hat{u}_m - u_m\|_{H^{\gamma'}(\Omega)} \lesssim \rho^{(\gamma-\gamma')} \|\hat{u}_m - u_m\|_{H^{\gamma}(\Omega)}$$

- Sum up, use assumed embedding  $\mathcal{U} \subset H^{\gamma}(\Omega)$ , and optimality

$$\begin{aligned} \sum_m \|\hat{u}_m - u_m\|_{H^{\gamma'}(\Omega)}^2 &\lesssim \rho^{2(\gamma-\gamma')} \left( \sum_m \|\hat{u}_m\|_{H^{\gamma}(\Omega)}^2 + \|\hat{u}_m\|_{H^{\gamma}(\Omega)}^2 \right) \\ &\lesssim \rho^{2(\gamma-\gamma')} \left( \sum_m \|\hat{u}_m\|_{\mathcal{U}}^2 + \|u_m\|_{\mathcal{U}}^2 \right) \lesssim \rho^{2(\gamma-\gamma')} \left( \|p\|_{\mathcal{P}}^2 + \sum_m \|u_m\|_{\mathcal{U}}^2 \right) \end{aligned}$$

# Controlling the Equation Error

$$(\hat{u}_m, \hat{p}) = \arg \min_{v_m \in \mathcal{U}, q \in \mathcal{P}} \|q\|_{\mathcal{P}}^2 + \sum_{m=1}^M \|v_m\|_{\mathcal{U}}^2$$

$$\text{s.t.} \quad v_m(X_{\text{obs}}) = u_m(X_{\text{obs}})$$

$$q \circ \Phi(v_m, X_{\text{obs}}) = f_m(X_{\text{obs}}) = p \circ \Phi(u_m, X_{\text{obs}})$$

- Observe this is "almost" an interpolation problem for  $p$
- $S = \bigcup_m \bigcup_{x \in X_{\text{obs}}} \Phi(u_m, x)$  then constraint is

$$q(S + \text{noise}) = p(S)$$

- Akin to total least squares!
- We need to do more work

# Controlling the Equation Error

- Basic idea

$$\begin{aligned}\widehat{p}(S + \delta S) &= p(S) \\ \widehat{p}(S) + \nabla \widehat{p}(S) \delta S &\approx p(S) \\ \Rightarrow |\widehat{p}(S) - p(S)| &\approx |\nabla \widehat{p}(S) \delta S|\end{aligned}$$

- Use Sobolev sampling inequality again but with **noisy RHS**
- Local Lipschitz assumption for  $q \in \mathcal{P}$ :

$$|q(s) - q(s')| \leq C(B) \|q\|_{\mathcal{P}} |s - s'|, \quad \forall s, s' \in B,$$

for any bounded set  $B \subset \mathbb{R}^Q$

# Controlling the Equation Error

$$\hat{p} \circ \Phi(\hat{u}_m, x_k) = p \circ \Phi(u_m, x_k) = f(x_k)$$

$$\Rightarrow |(\hat{p} - p) \circ \Phi(u_m, x_k)| = |\hat{p} \circ \Phi(\hat{u}_m, x_k) - \hat{p} \circ \Phi(u_m, x_k)|$$

$$\text{(Lipschitz assumption)} \quad \leq C(B) \|\hat{p}\|_{\mathcal{P}} |\Phi(\hat{u}_m, x_k) - \Phi(u_m, x_k)|$$

Suppose  $\Phi(u, x) = (x, u(x), \partial^{\mathbf{a}} u(x))$  for some multi-index  $\mathbf{a}$  s.t.  $|\mathbf{a}| \leq k \in \mathbb{N}$

$$|(\hat{p} - p) \circ \Phi(u_m, x_k)| \lesssim C(B) \|\hat{p}\|_{\mathcal{P}} \|\hat{u}_m - u_m\|_{C^k(\Omega)}$$

$$\text{(Sobolev embedding)} \quad \lesssim C(B) \|\hat{p}\|_{\mathcal{P}} \|\hat{u}_m - u_m\|_{H^{\gamma'}(\Omega)}$$

$$\text{(From filtering bound)} \quad \lesssim C(B) \|\hat{p}\|_{\mathcal{P}} \rho^{\gamma - \gamma'} (\|p\|_{\mathcal{P}} + \sum_m \|u_m\|_{\mathcal{U}})$$

# Controlling the Equation Error

$$|(\hat{p} - p) \circ \Phi(u_m, x_k)| \lesssim C(B) \|\hat{p}\|_{\mathcal{P}} \rho^{\gamma-\gamma'} \left( \|p\|_{\mathcal{P}} + \sum_m \|u_m\|_{\mathcal{U}} \right)$$

- Extra lemma:  $\|\hat{p}\|_{\mathcal{P}} \leq \|p\|_{\mathcal{P}} + \rho^{\gamma} \sum_m \|u_m\|_{\mathcal{U}^2}^2$  where  $\|\cdot\|_{\mathcal{U}^2}$  is stronger norm on  $\mathcal{U}$ .
- Up to leading order

$$|(\hat{p} - p) \circ \Phi(u_m, x_k)| \lesssim C(B) \|p\|_{\mathcal{P}} \rho^{\gamma-\gamma'} \left( \|p\|_{\mathcal{P}} + \sum_m \|u_m\|_{\mathcal{U}} \right)$$

- Apply sampling inequality for  $p$  and  $\hat{p}$  with fill distance

$$\varrho(B) := \sup_{s' \in B} \inf_{m,k} |s' - \Phi(u_m, x_k)|$$

$$\|\hat{p} - p\|_{L^\infty(\Omega)} \leq C(B) \left( \varrho(B)^{\eta-Q/2} + \rho^{(\gamma-\gamma')} \right) \left( \|p\|_{\mathcal{P}} + \sum_m \|u_m\|_{\mathcal{U}} \right)$$

# Thank you

Yasamin Jalalian, Juan Felipe Osorio Ramirez, Alex Hsu,, Bamdad Hosseini, and Houman Owhadi, **Data-Efficient Kernel Methods for Learning Differential Equations and Their Solution Operators: Algorithms and Error Analysis**, arXiv:2503.01036, 2025

# References I

- [1] S. L. Brunton, J. L. Proctor, and J. N. Kutz. “Discovering governing equations from data by sparse identification of nonlinear dynamical systems”. In: *Proceedings of the National Academy of Sciences* 113.15 (2016), pp. 3932–3937. DOI: 10.1073/pnas.1517384113.
- [2] N. B. Kovachki, S. Lanthaler, and A. M. Stuart. “Operator learning: Algorithms and analysis”. In: *arXiv preprint arXiv:2402.15715* (2024).
- [3] Z. Chen, Y. Liu, and H. Sun. “Physics-informed learning of governing equations from scarce data”. In: *Nature communications* 12.1 (2021), p. 6136.
- [4] F. Sun, Y. Liu, and H. Sun. “Physics-informed Spline Learning for Nonlinear Dynamics Discovery”. In: *Proceedings of the Thirtieth International Joint Conference on Artificial Intelligence, IJCAI-21*. Ed. by Z.-H. Zhou. Main Track. International Joint Conferences on Artificial Intelligence Organization, Aug. 2021, pp. 2054–2061.

## References II

- [5] E Haber and D Oldenburg. “Joint inversion: a structural approach”. In: *Inverse problems* 13.1 (1997), p. 63.
- [6] B. Kaltenbacher. “Regularization based on all-at-once formulations for inverse problems”. In: *arXiv [math.NA]* (Mar. 16, 2016).
- [7] Y. Chen et al. “Solving and learning nonlinear PDEs with Gaussian processes”. In: *Journal of Computational Physics* 447 (2021), p. 110668.
- [8] D. Long et al. “A kernel framework for learning differential equations and their solution operators”. In: *Physica D: Nonlinear Phenomena* 460 (2024), p. 134095.
- [9] L. Zhang and H. Schaeffer. “On the convergence of the SINDy algorithm”. In: *Multiscale Modeling & Simulation* 17.3 (2019), pp. 948–972.
- [10] P. Scholl et al. “The uniqueness problem of physical law learning”. In: *ICASSP 2023-2023 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*. IEEE. 2023, pp. 1–5.

## References III

- [11] Y. He, H. Zhao, and Y. Zhong. “How much can one learn a partial differential equation from its solution?” In: *Foundations of Computational Mathematics* 24.5 (2024), pp. 1595–1641.
- [12] N. Boullé, D. Halikias, and A. Townsend. “Elliptic PDE learning is provably data-efficient”. In: *Proceedings of the National Academy of Sciences* 120.39 (2023), e2303904120.