



DEEP
LEARNING
INSTITUTE



BROWN



MathWorks®

Deep Learning for Science and Engineering Teaching Kit

Deep Learning for Scientists and Engineers

DAY 1: PINNs and PIKANs

Instructors: George Em Karniadakis, Khemraj (Raj) Shukla





The Deep Learning for Science and Engineering Teaching Kit is licensed by NVIDIA and Brown University under the [Creative Commons Attribution-NonCommercial 4.0 International License](https://creativecommons.org/licenses/by-nc/4.0/).

Course Roadmap

Lecture 1: PINNs and PIKANs: 90 Minutes

Lecture 1: Hands On: 90 Mins

Lecture 2: Neural Operators: 105 Minutes

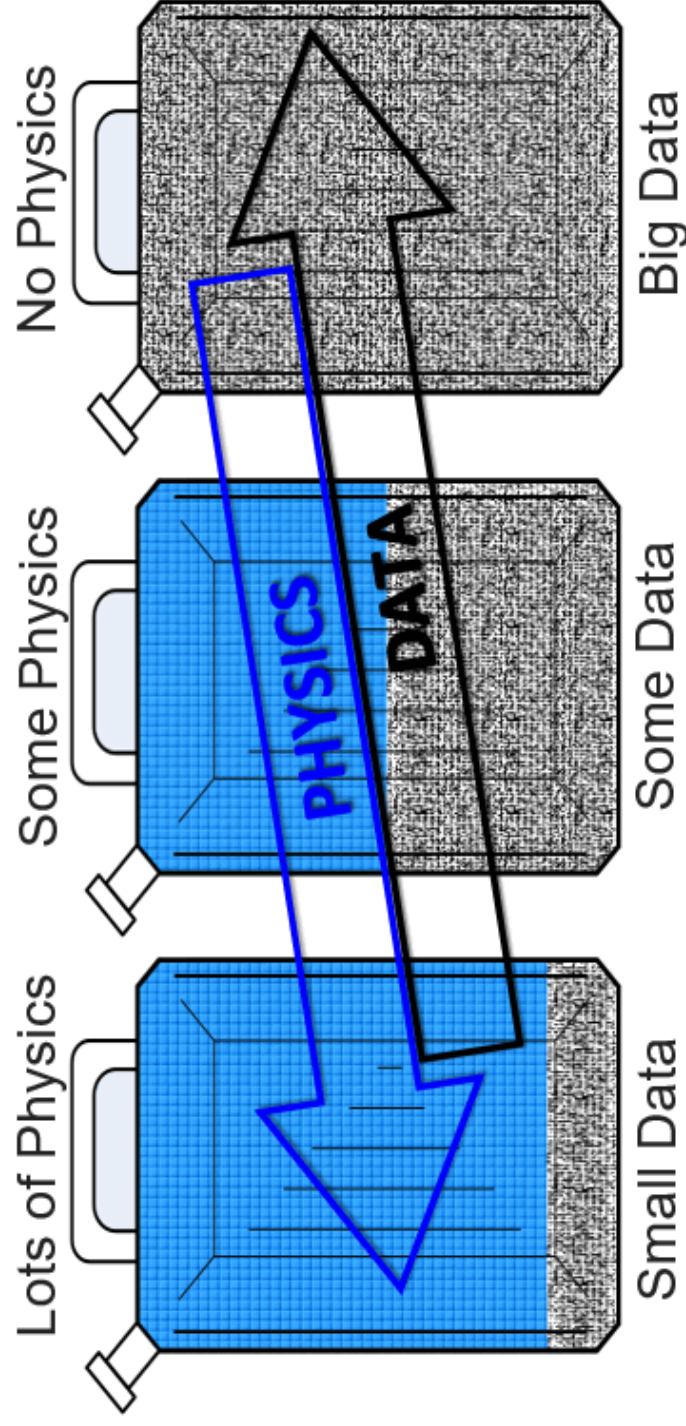
Lecture 2: Hands On: 75 Mins

Contents

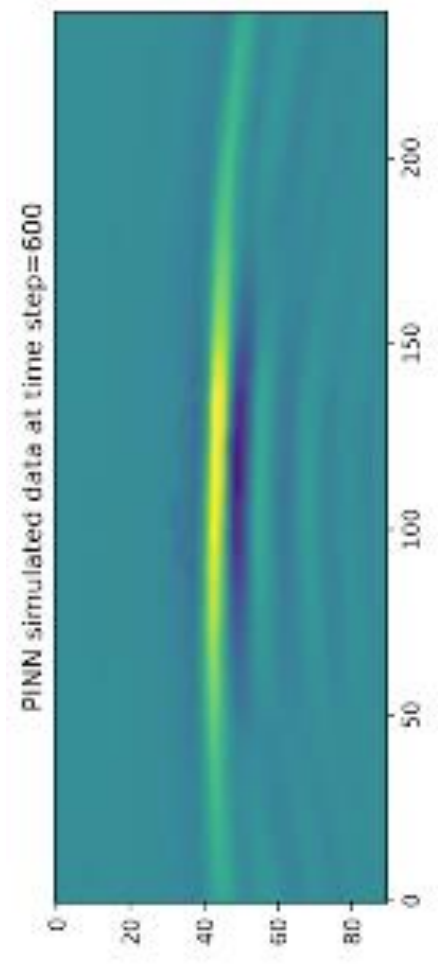
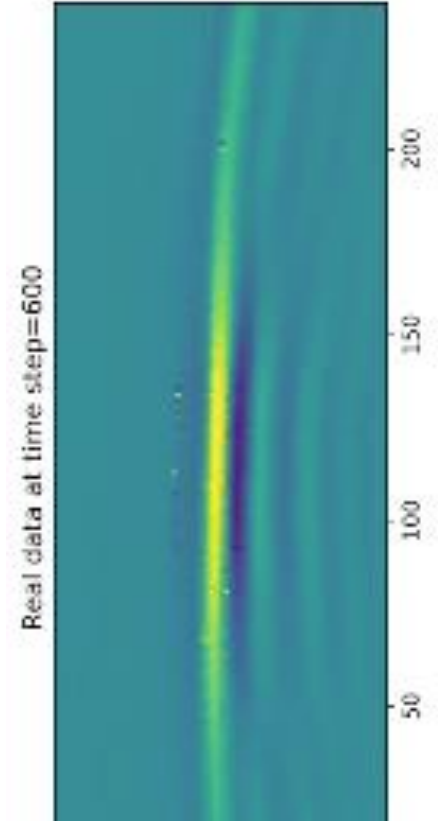
- Data + Physical Laws
- Data + Physical Laws + Neural Networks
- What is a PINN and Why PINNs
- Real Data – PINN Application
- PINNs for Burgers Equation
- Optimizers
- PINNs for Boundary Value Problems
- Soft Constraints and Weights – IB Theory
- Residual Based Weights & Learning Theory
- Weighted Residual Methods
- hp-VPINNs: Domain Decomposition
- Variational Neural Networks
- Convergence Theory of PINNs
- Convergence Theory of hp-VPINNs
- Error Decomposition
- PINNs with Domain Decomposition
- Separable PINNs
- Kolmogorov-Arnold-Networks (KANs)
- Physics-Informed KANs (PIKANs)
- Summary
- References

Data + Physical Laws

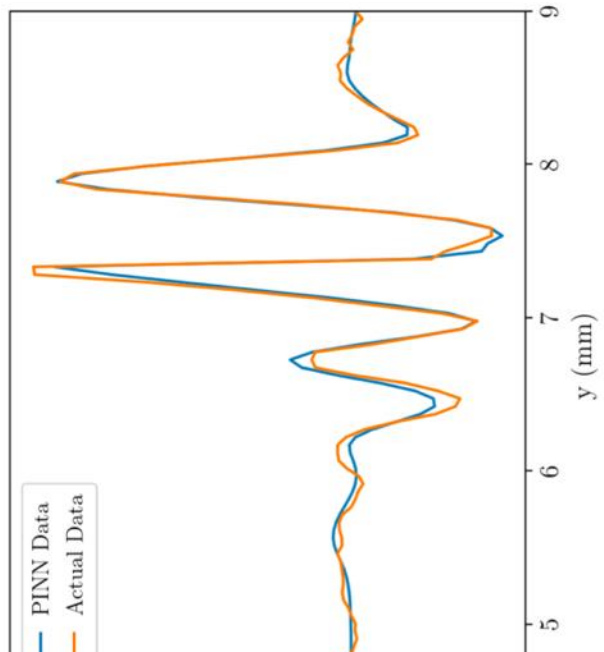
Three scenarios of Physics-Informed Learning Machines



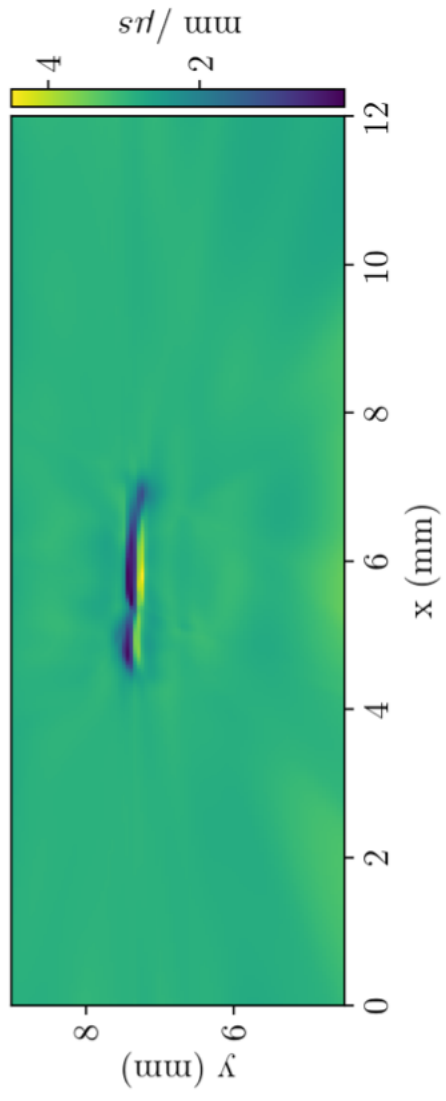
Crack characterization: 0° Pre-crack + Post-crack + $v(x, y)$



$$f := \frac{\partial^2 \Psi_{\text{NN}}}{\partial t^2} - c^2$$



ces of (a) and (b) at $x = 6$ mm at $t = 12.38 \mu\text{s}$.



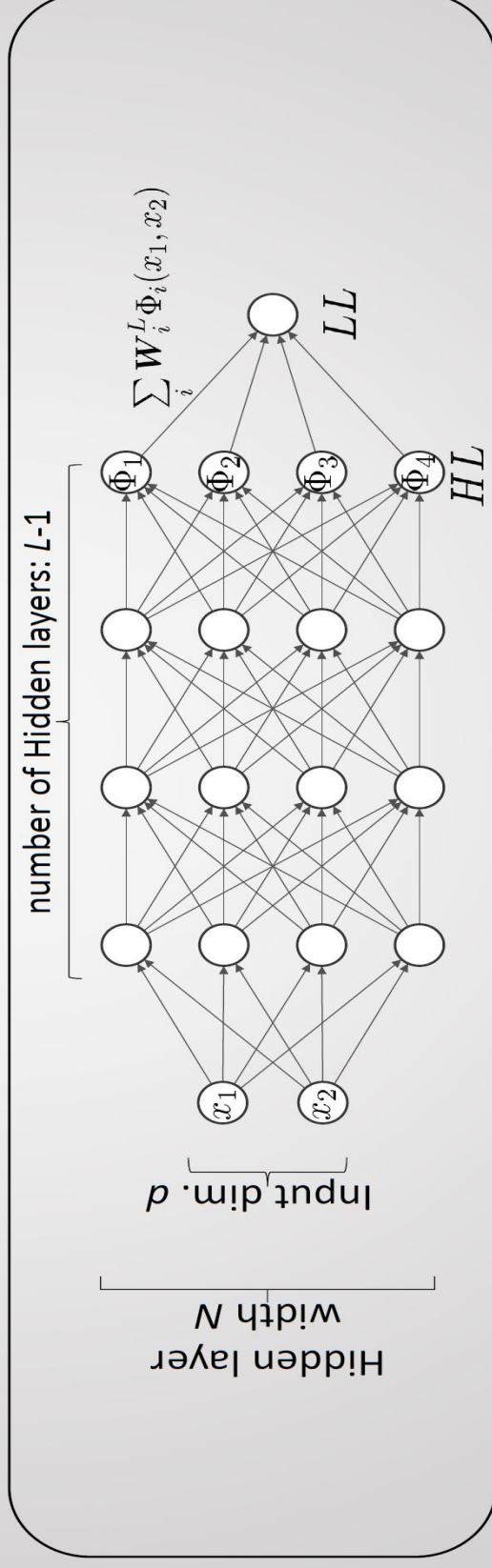
(d) Speed $v(x, y)$ recovered from PINN simulation.

Adaptive Basis Viewpoint: What is a Neural Network?

- We consider a family of neural networks $\mathcal{N}_\theta : \mathbb{R}^d \rightarrow \mathbb{R}$ consisting of $L - 1$ hidden layers of width N composed with a final linear layer, admitting the representation

$$\mathcal{N}_\theta(x) = \sum_{i=1}^N \mathbf{W}_i^L \Phi_i(x; \theta_H)$$

where $\mathbf{W}^L$ and θ_H are the parameter corresponding to the final linear layer and the hidden layers respectively. We interpret θ as a concatenation of $\mathbf{W}^L$ and θ_H .



- This view point makes it clear that θ_H parameterizes the basis (like FEM mesh & Shape functions), while $\mathbf{W}^L$ are just coefficients for these basis functions.

Connecting Neural Networks with Finite Elements (FEM) – 1D

Recall the hat function

$$\varphi(x) = \begin{cases} 2x & x \in [0, \frac{1}{2}] , \\ 2(1-x) & x \in [\frac{1}{2}, 1] , \\ 0, & \text{otherwise.} \end{cases}$$

$$\varphi(x) = 2\operatorname{ReLU}(x) - 4\operatorname{ReLU}(x - \frac{1}{2}) + 2\operatorname{ReLU}(x - 1).$$

- Linear finite element basis function φ_i

$$\varphi_i = \varphi\left(\frac{x - x_{i-1}}{2h}\right) = \varphi(w_h x + b_i), \quad w_h = \frac{1}{2h}, \quad b_i = \frac{-(i-1)}{2}.$$

$$\varphi_i \in \operatorname{Span}\{\operatorname{ReLU}(wx + b), w, b \in \mathbb{R}\}$$

- The Linear FEM space in 1D is a subspace of shallow neural networks with the ReLU activation function.

Connecting Neural Networks with Finite Elements (FEM) - 2D/3D

◇ For $d \geq 2$, the linear FEM space **is not** a subspace of **shallow** ReLU networks (He et al. 2020)

$$\left\{ \sum_{i=1}^n a_i \operatorname{ReLU}(w_i^T x + b_i), w_i \in \mathbb{R}^{1 \times d}, a_i, b_i \in \mathbb{R} \right\}$$

◇ For $d \geq 2$, the linear FEM space **is** a subspace of **deep** ReLU networks (Arora et al. 2016).

- Arora R, Basu A, Mianjy P, Mukherjee A. Understanding deep neural networks with rectified linear units. arXiv preprint arXiv:1611.01491. 2016.
- He J, Li L, Xu J, Zheng C. ReLU deep neural networks and linear finite elements. arXiv preprint arXiv:1807.03973. 2018 Jul 11.



Data + Neural Networks + Physical Laws

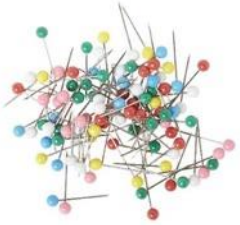


PINNs

Physics-informed neural networks: A deep learning framework for solving forward and inverse problems involving nonlinear partial differential equations

M Raissi, P Perdikaris, GE Karniadakis
Journal of Computational Physics 378, 686-707, **2019**

arXiv:1711.10561; arXiv:1711.10566 (2017)



PINNs: Physics-Informed Neural Networks

AI FOR SCIENCE

“ ... All of the knowledge are led to **physics equations** and now **embedded into the neural networks**. And it gives a gigantic head start, **physics informed neural networks...** ”

EXPERIMENTAL DATA

Real-time Steering

NEURAL ESTIMATION

Fast Approximation

SIMULATION

DATA

INVERSE PROBLEMS
LIGO — Gravitational Waves

DESIGN SPACE EXPLORATION
ICF — MERLIN — Fusion

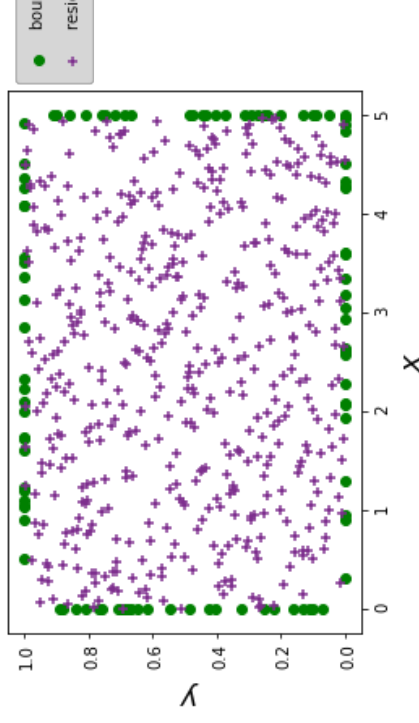
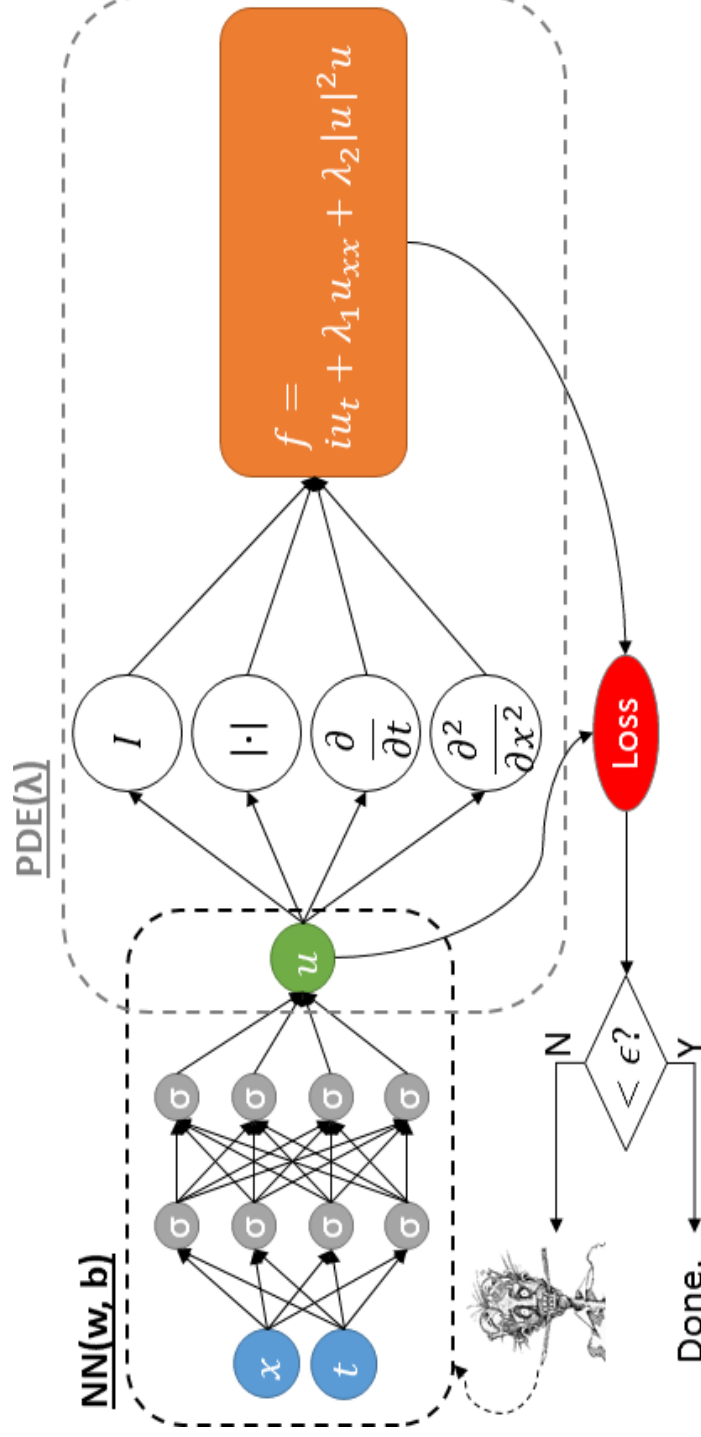
FASTER PREDICTION
ANI — MD — Chemistry

REAL-TIME STEERING
ITER — Fusion Energy

NVIDIA CEO J. Huang at SC19,
the annual supercomputing conference, Nov 19

What is a PINN? Physics-Informed Neural Network

We employ two (or more) NNs that share the same parameters

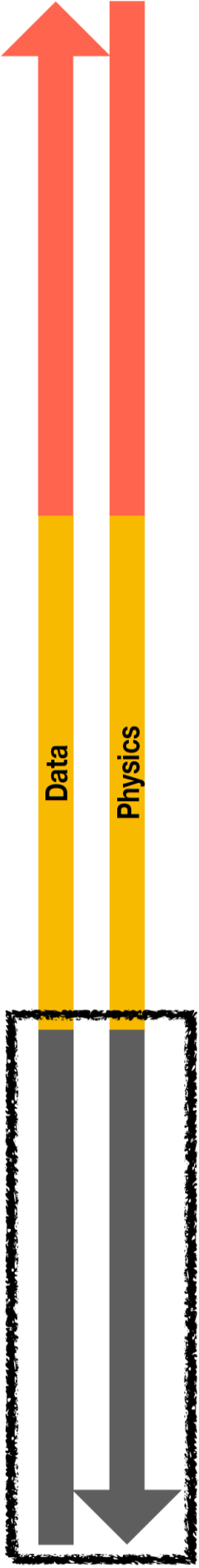


➤ Minimize:

$$MSE_u = \frac{1}{N_u} \sum_{i=1}^{N_u} |u(t_u^i, x_u^i) - u^i|^2,$$

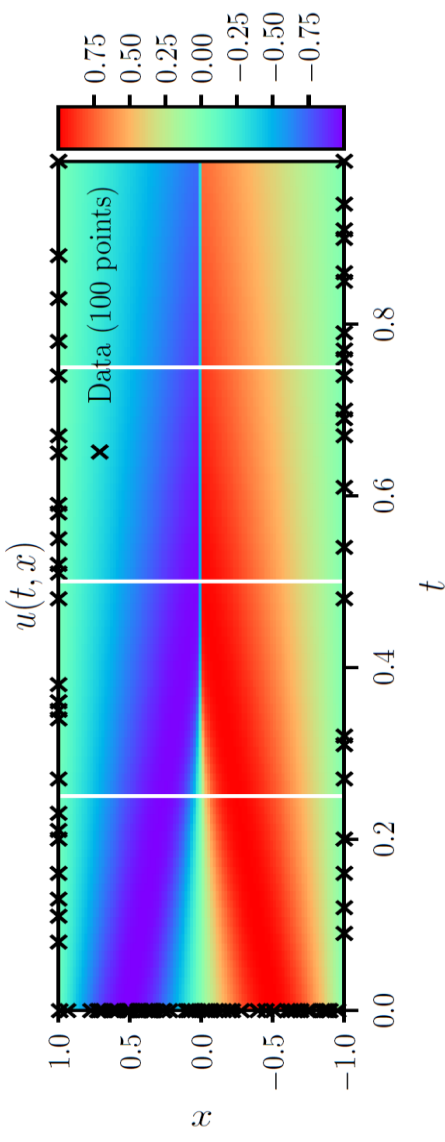
$$MSE_f = \frac{1}{N_f} \sum_{i=1}^{N_f} |f(t_f^i, x_f^i)|^2.$$

$$MSE = MSE_u + MSE_f,$$



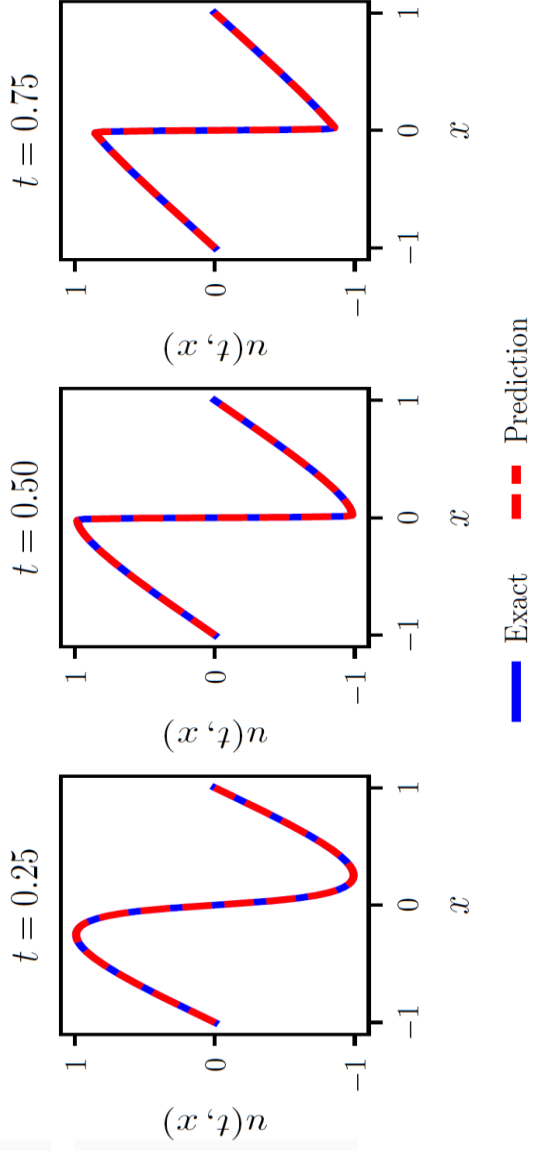
Physics Informed Neural Networks

$$u_t + uu_x - (0.01/\pi)u_{xx} = 0$$



```
def u(t, x):
    u = neural_net(tf.concat([t, x], 1), weights, biases)
    return u
```

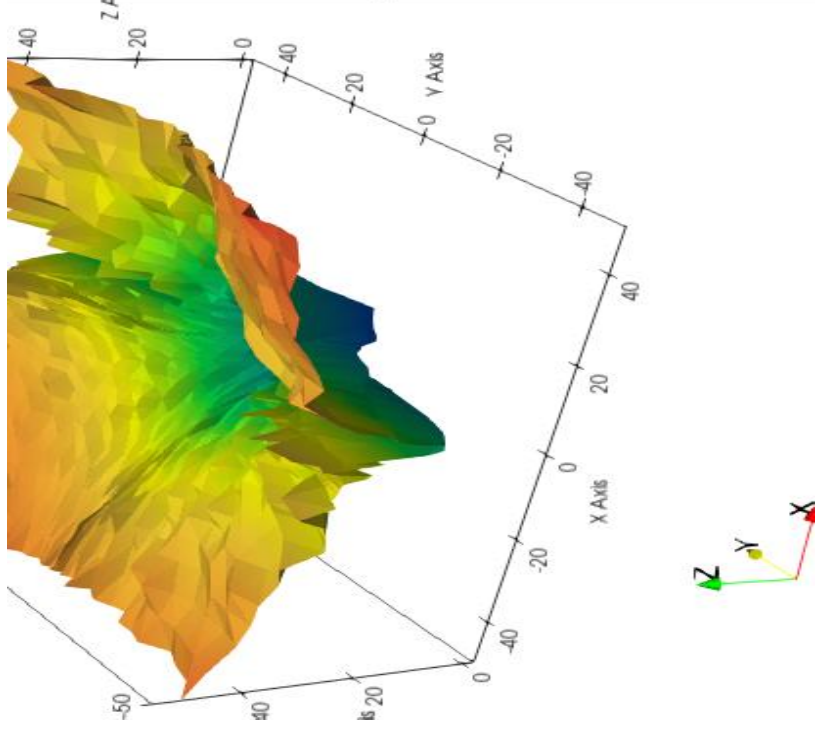
```
def f(t, x):
    u = u(t, x)
    u_t = tf.gradients(u, t)[0]
    u_x = tf.gradients(u, x)[0]
    u_xx = tf.gradients(u_x, x)[0]
    f = u_t + u*u_x - (0.01/tf.pi)*u_xx
    return f
```



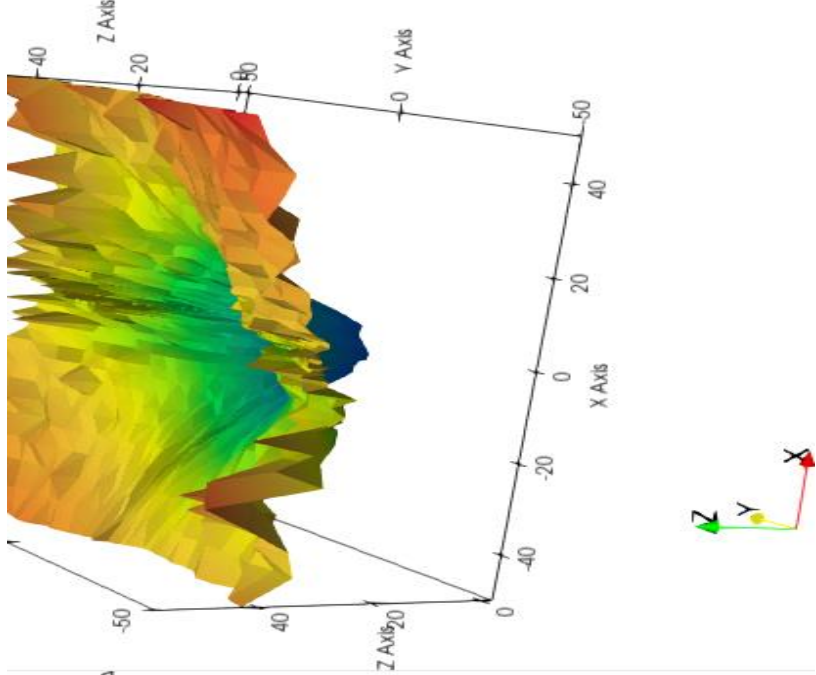
$$\sum_{i=1}^{N_u} |u(t_u^i, x_u^i) - u^i|^2 + \sum_{i=1}^{N_f} |f(t_f^i, x_f^i)|^2$$

Visualizing PINN Landscapes: 1-D viscous Burgers Equation

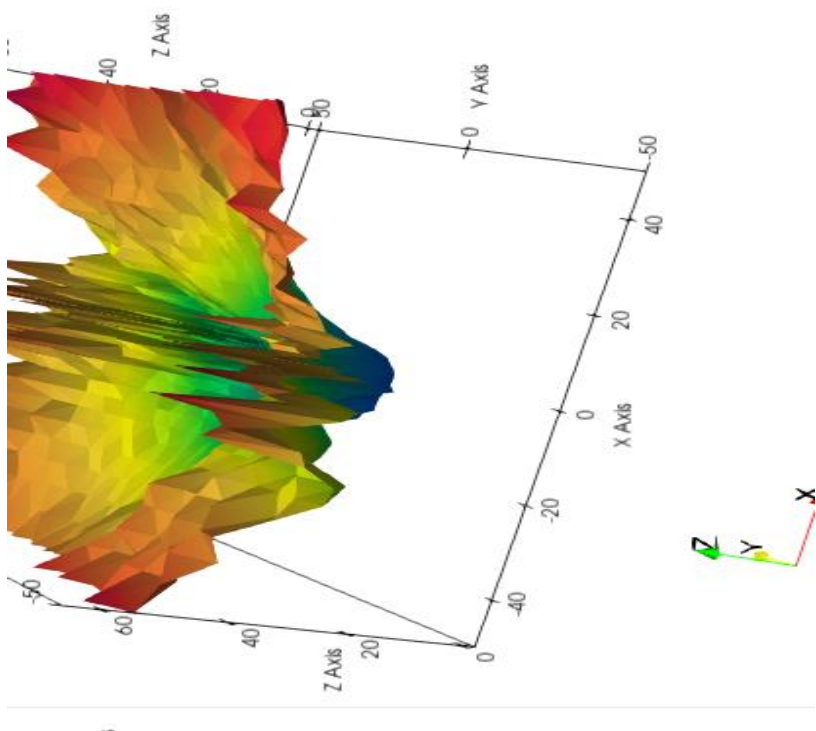
$$\nu = 10^{-2}$$



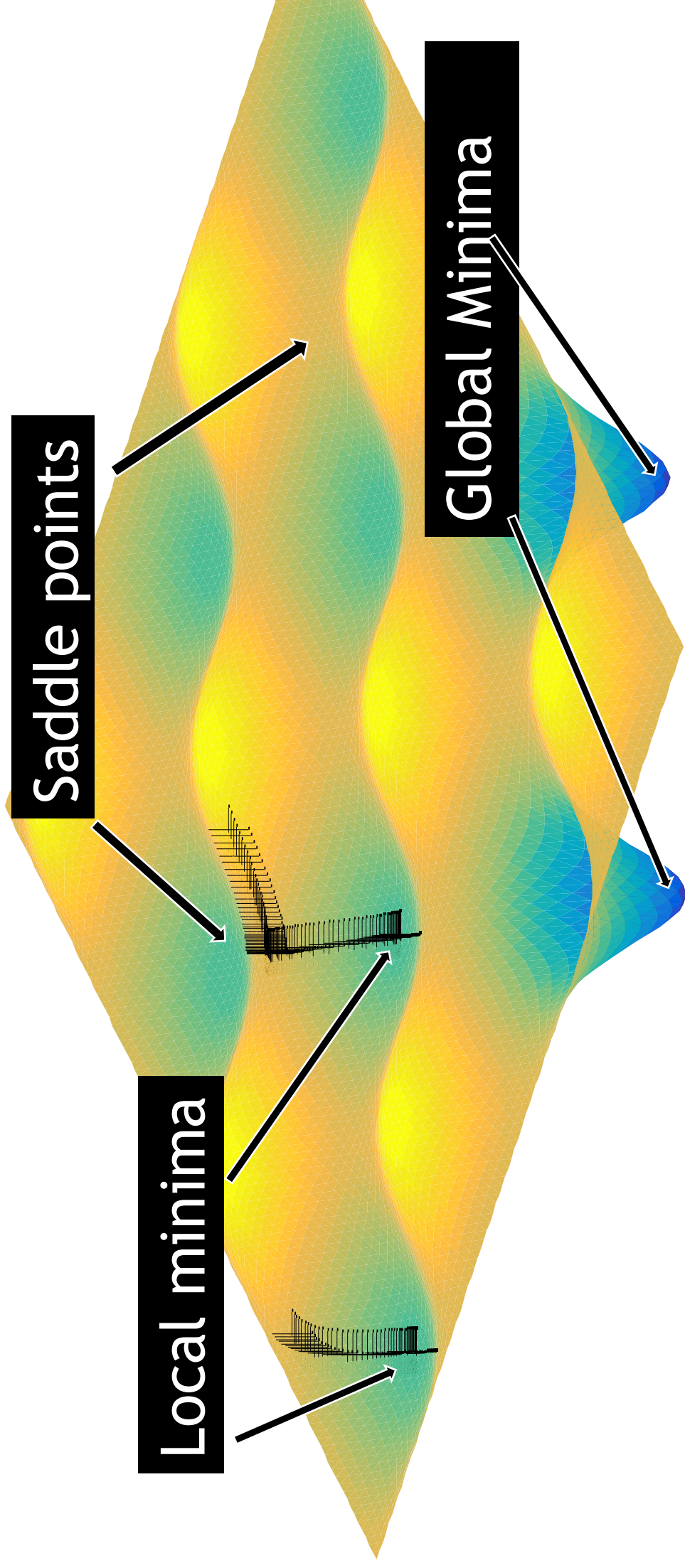
$$\nu = 10^{-3}$$



$$\nu = 10^{-8}$$



Types of Stationary Points

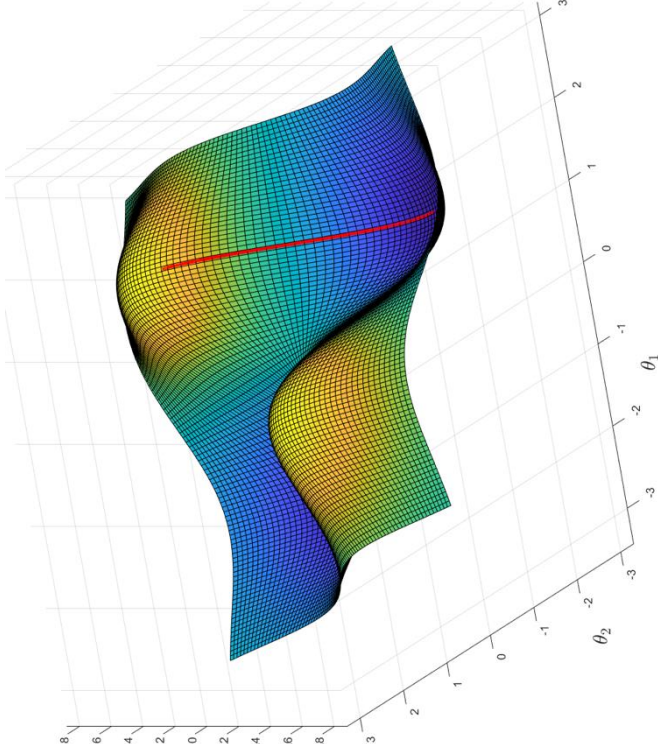


- x-y plane = search space for optimization
- z axis = value of objective function to be minimized

Gradient Descent (GD)

$$\theta^* = \arg \min_{\theta} \mathcal{L}(\theta)$$

$$\theta_{n+1} = \theta_n - \eta \nabla_{\theta} \mathcal{L}(\theta)$$



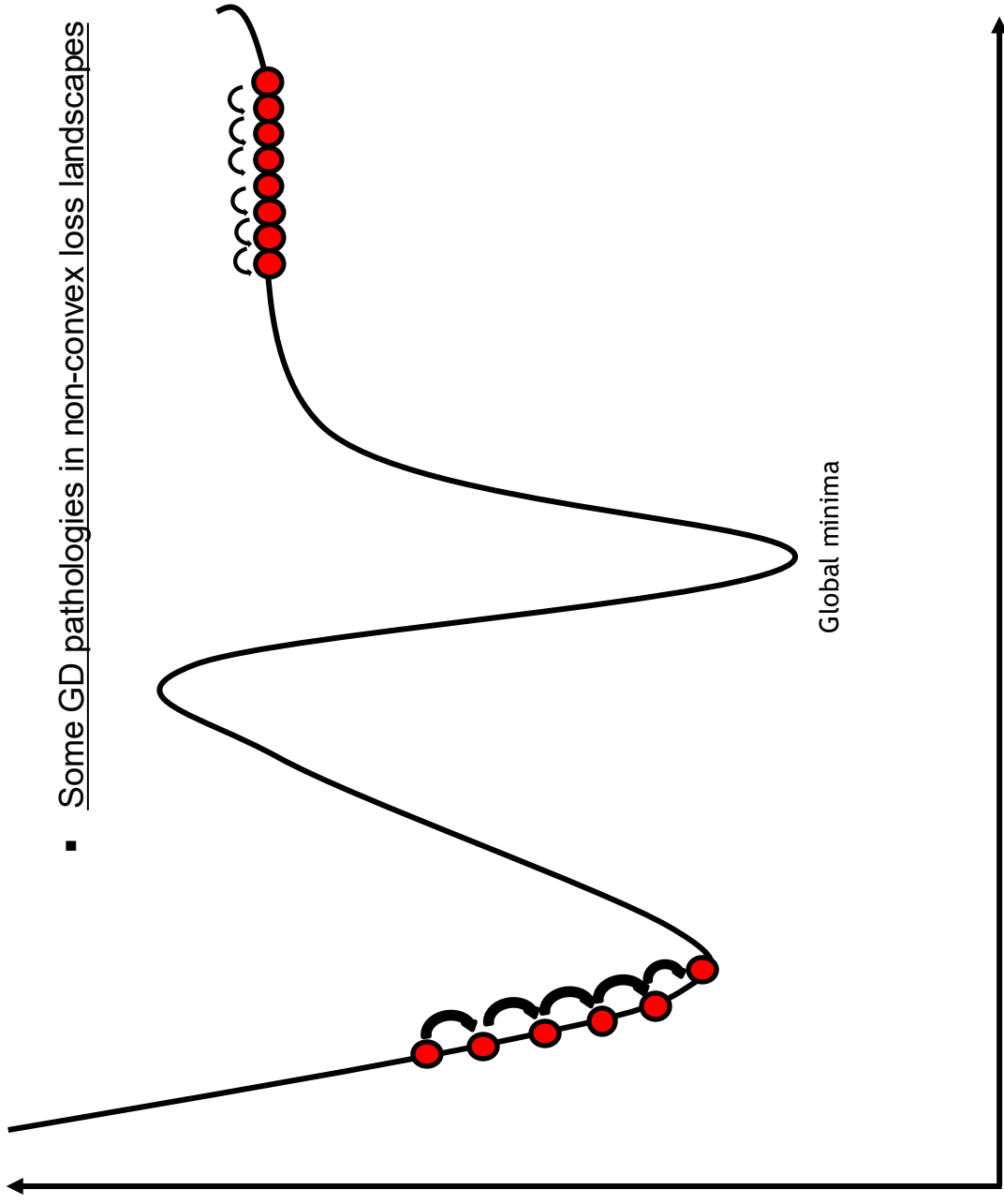
```
tf.keras.optimizers.SGD(learning_rate=0.01)
```



DEEP
LEARNING
INSTITUTE



```
torch.optim.SGD(params, lr=0.01)
```

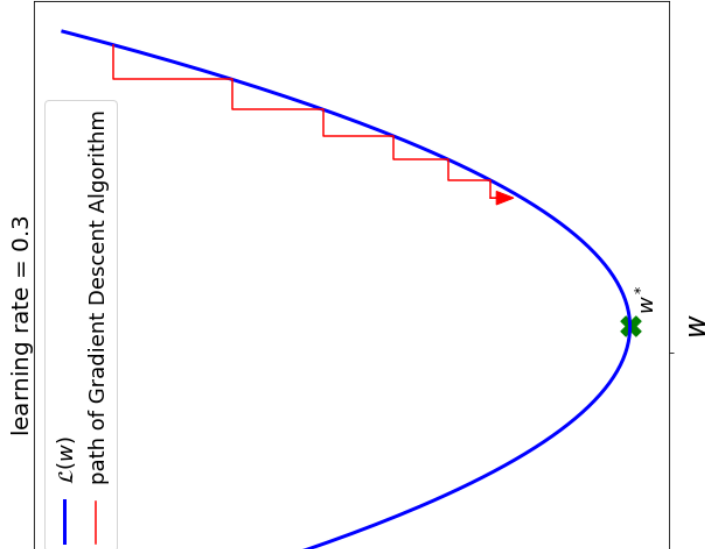


- Some GD pathologies in non-convex loss landscapes

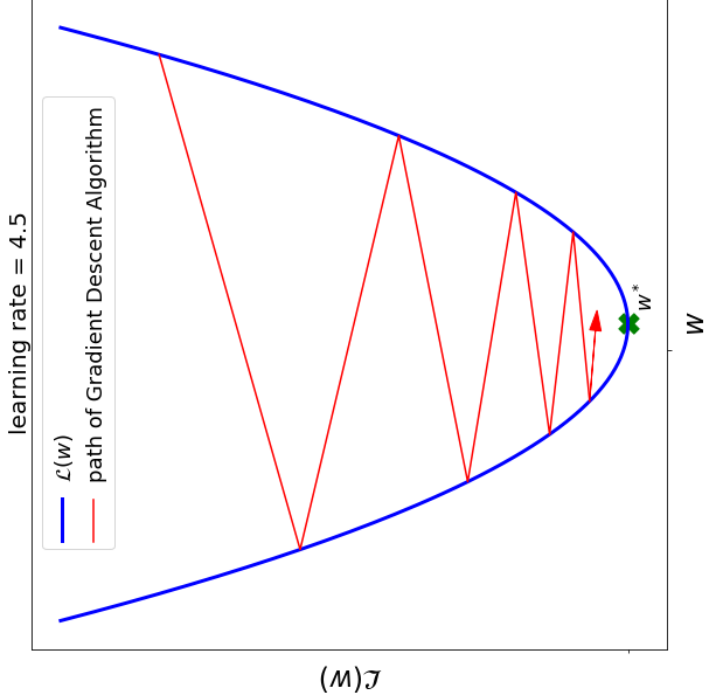
Effect of Learning Rate

In linear regression we have [convexity](#) (hence global minimum) but still we should scale all features for faster convergence

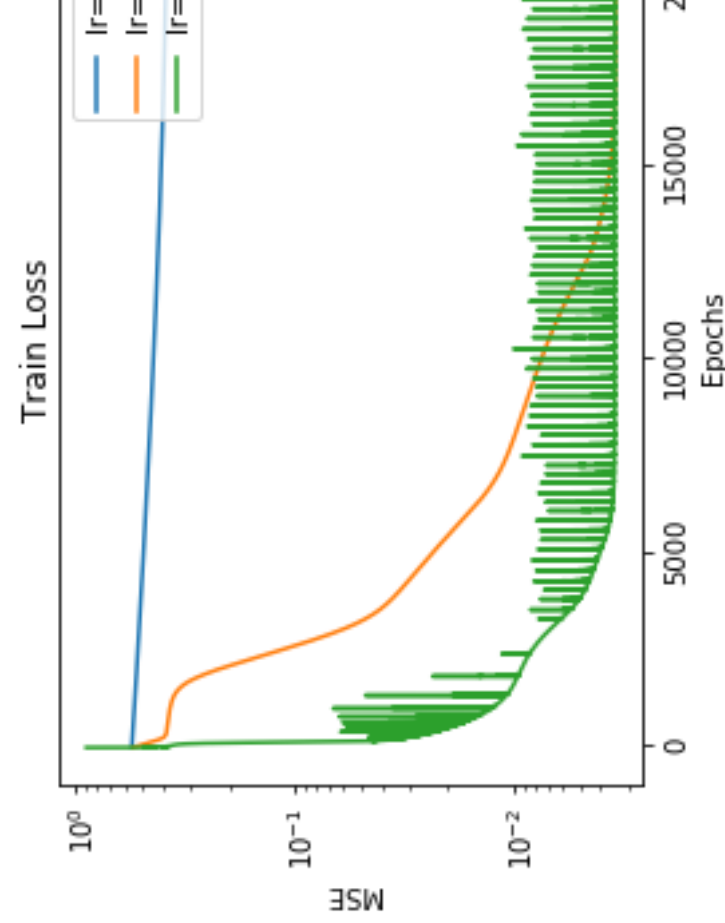
Loss plot associated with learning the sine function with noise using a fully connected neural network



Learning rate is too small



Learning rate is too large

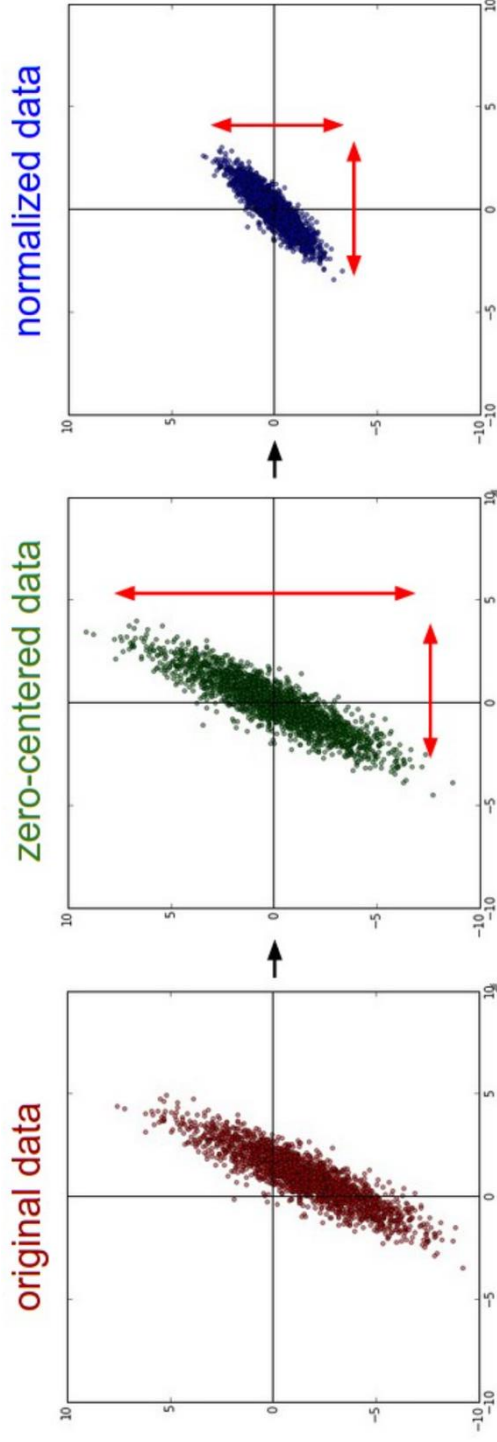


Convergence depends strongly on the learning rate

❑ An effective strategy is to use a variable/decaying learning rate

a Normalization

Ref: <https://zaffnet.github.io/batch-normalization>



Common data preprocessing pipeline. **Left:** Original toy, 2-dimensional input data. **Middle:** The data is zero-centered by subtracting the mean in each dimension. The data cloud is now centered around the origin. **Right:** Each dimension is additionally scaled by its standard deviation. The red lines indicate the extent of the data - they are of unequal length in the middle, but of equal length on the right.

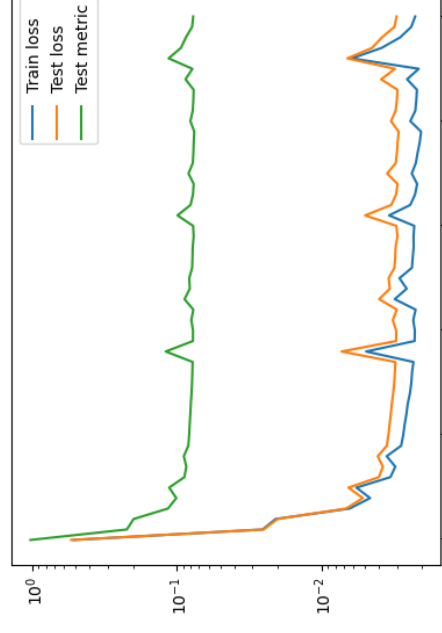
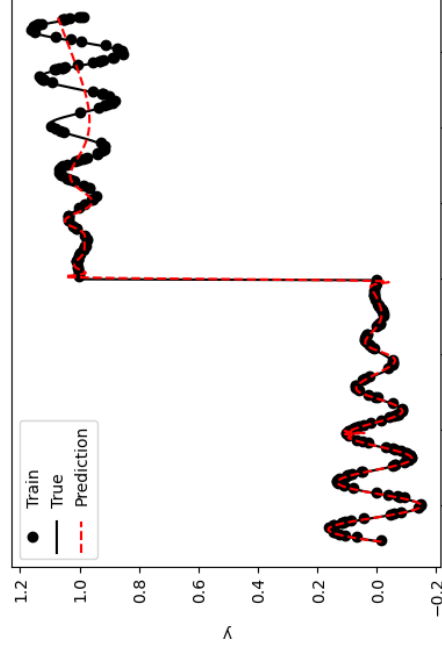
In 1998, Yan LeCun in his paper, Efficient BackProp, highlighted the importance of normalizing the inputs. Preprocessing of the inputs using normalization is a standard machine learning procedure and is known to help in faster convergence. Normalization is done to achieve the following objectives:

- ❑ The average of each input variable (or feature) over the training set is close to zero (Mean subtraction).
- ❑ Covariances of the features are same (Scaling).
- ❑ Decorrelate the features (Whitening – not required for CNNs).

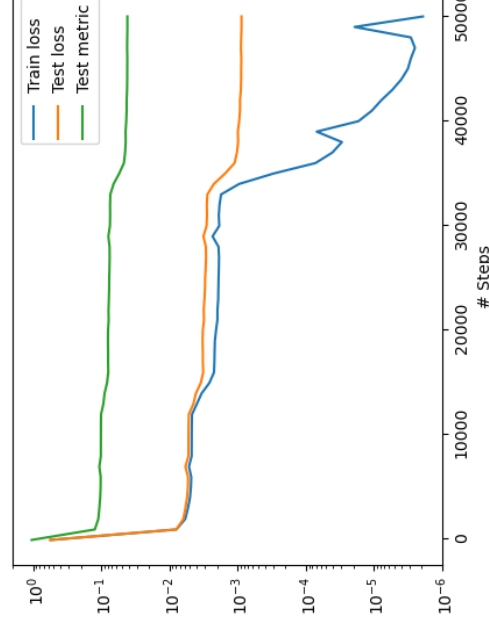
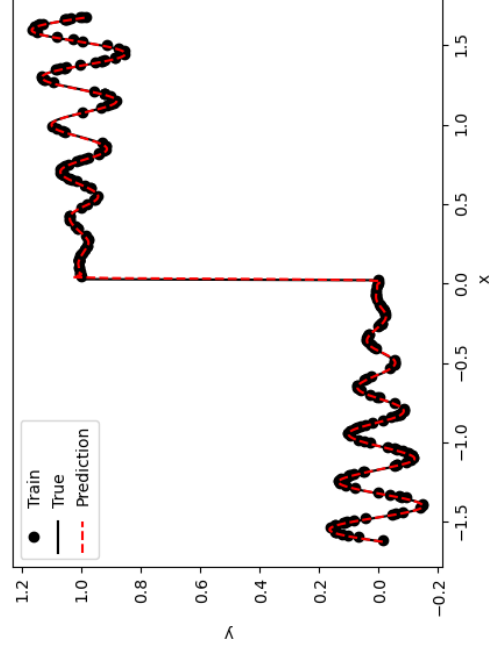
a Normalization - Example

Without data-normalization

$$y = \begin{cases} \frac{1}{10}x\sin(20x) & \text{if } x < 0 \\ 0.5 & \text{if } x = 0.5 \\ 1 + \frac{1}{10}x\sin(20x) & \text{if } x > 0 \end{cases}$$



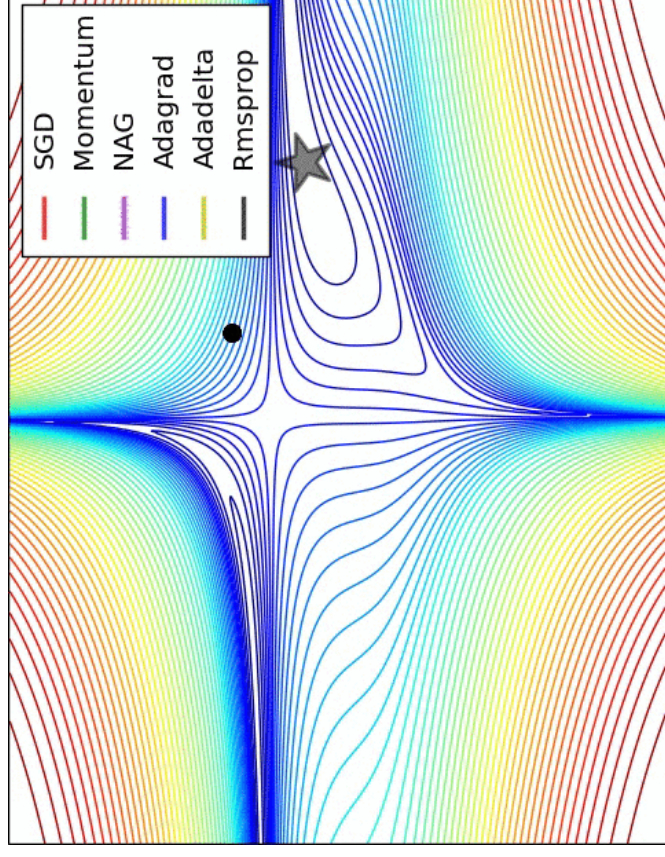
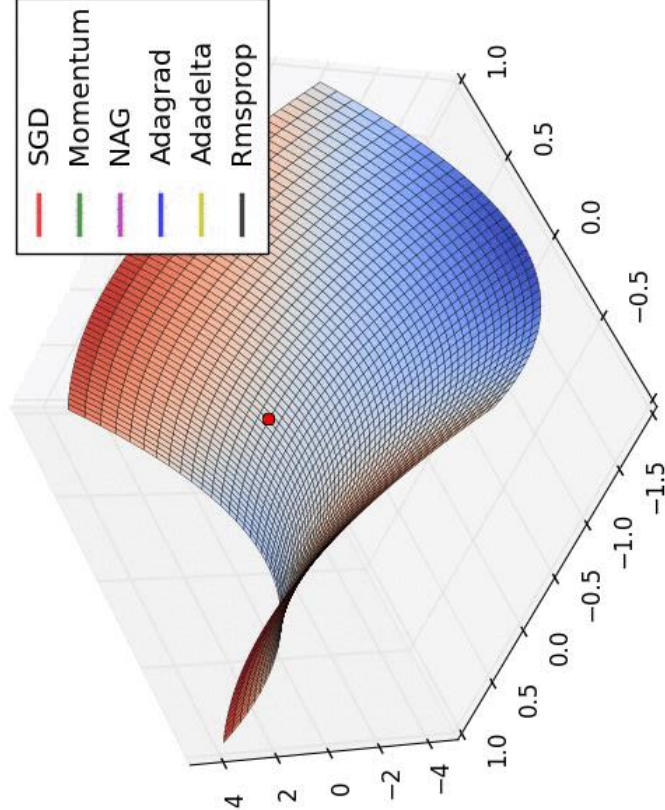
With data-normalization



Overview of Gradient Descent Optimization Algorithms

[org/abs/1609.04747](https://arxiv.org/abs/1609.04747)

post explores how many of the most popular gradient-based optimization algorithms actually work.



movie shows the **behaviour of the algorithms at a saddle point** Notice that momentum, and NAG find it difficult to break symmetry, although the two latter manage to escape the saddle point, while Adagrad, RMSprop, and Adadelata head down the negative slope.

B. In this movie, we see their **behavior on the contours of a loss surface (the Beale function) over time**. Note that Adagrad, Adadelata, and RMSprop almost immediately head off in the right direction and converge similarly while Momentum and NAG are led off-track, evoking the image of a ball rolling down the hill. NAG, however, is quickly able to correct its course due to increased responsiveness by looking ahead and heads to the minimum.

What Optimizer to Use?

On faster optimizers : [second-order](#)

Hessian based optimizers

Optimization technique discussed so far is based on first-order partial derivatives (Jacobians).

There are optimizers, which are based on the second-order partial derivatives (the Hessian).

These algorithms are very expensive as there are n^2 Hessians to be computed for each output (n is number of parameters) in comparison to n Jacobians per output.

Next, we will discuss a couple of Hessian-based optimizer and their implementation in TensorFlow and PyTorch.

$$\nabla_{\theta}^2 f(\theta) = \begin{bmatrix} \frac{\partial^2 f(\theta)}{\partial \theta_1^2} & \frac{\partial^2 f(\theta)}{\partial \theta_1 \partial \theta_2} & \dots & \frac{\partial^2 f(\theta)}{\partial \theta_1 \partial \theta_n} \\ \frac{\partial^2 f(\theta)}{\partial \theta_2 \partial \theta_1} & \frac{\partial^2 f(\theta)}{\partial \theta_2^2} & \dots & \frac{\partial^2 f(\theta)}{\partial \theta_2 \partial \theta_d} \\ \vdots & \vdots & \ddots & \vdots \\ \frac{\partial^2 f(\theta)}{\partial \theta_d \partial \theta_1} & \frac{\partial^2 f(\theta)}{\partial \theta_d \partial \theta_2} & \dots & \frac{\partial^2 f(\theta)}{\partial \theta_d^2} \end{bmatrix}$$

What Optimizer to Use?

Newton's algorithm

Newton's algorithm uses iterative algorithms, which consider a sequence of $\mathbf{x}_n$ converging to $\mathbf{x}^*$. Therefore, $\forall f''(\mathbf{x}) \neq 0$, the current estimate being $\mathbf{x}_n$, the next estimate $\mathbf{x}_{n+1}$ should satisfy

$$f(\mathbf{x}_{n+1}) < f(\mathbf{x}_n)$$

Using Taylor Approximation

$$f(\mathbf{x} + \Delta \mathbf{x}) \approx f(\mathbf{x}) + \Delta \mathbf{x}^\top \nabla f(\mathbf{x}) + \frac{1}{2} \Delta \mathbf{x}^\top (\nabla^2 f(\mathbf{x})) \Delta \mathbf{x}$$

Newton's algorithm produces a sequence of such quadratic approximations $\mathbf{h}_n$ and considers $\mathbf{x}_{n+1} = \mathbf{x}_n + \Delta \mathbf{x}$;

$$\mathbf{h}_n(\Delta \mathbf{x}) = f(\mathbf{x}_n) + \Delta \mathbf{x}^\top \mathbf{g}_n + \frac{1}{2} \Delta \mathbf{x}^\top \mathbf{H}_n \Delta \mathbf{x}$$

Where $\mathbf{g}_n$ and $\mathbf{H}_n$ are gradient and Hessian of f at $\mathbf{x}_n$.

$$\frac{\partial \mathbf{h}_n(\Delta \mathbf{x})}{\partial \Delta \mathbf{x}} = \mathbf{g}_n + \mathbf{H}_n \Delta \mathbf{x} = 0$$

optimality condition

$$\Delta \mathbf{x} = -\mathbf{H}_n^{-1} \mathbf{g}_n$$

$$\mathbf{x}_{n+1} = \mathbf{x}_n - \alpha (\mathbf{H}_n^{-1} \mathbf{g}_n)(\mathbf{x}_n)$$

What Optimizer to Use?

Pros and Cons of Iterative algorithms:

The main issue with iterative algorithm is that it needs to compute the **inverse** Hessian matrix.

In DL applications, the dimensionality of the input corresponds to model parameters. It is quite common to have hundreds of millions of parameters.

For these reasons, computing the Hessian or its inverse is often impractical.

Therefore, Iterative algorithm in this form is rarely used in practice to optimize functions corresponding to large problems.

Unluckily, the above algorithm can still work even if $\mathbf{H}_n^{-1}$ does not correspond to the exact inverse Hessian at $\mathbf{x}_n$, but is instead a good approximation. These are handled by **quasi-Newton methods** and are implemented through **BFGS** and **L-BFGS** methods.

What Optimizer to Use?

Newton Methods

General form of quasi-Newton methods is

$$\mathbf{x}_{n+1} = \mathbf{x}_n - \alpha \mathbf{H}_n^{-1} \nabla f(\mathbf{x}_n)$$

where α is line search parameter and $\mathbf{H}_n^{-1}$ is the inverse Hessian or some approximation.

Quasi-Newton methods lose their quadratic convergence of Newton's method, but super linear convergence is often achieved.

we will discuss two forms of quasi-Newton methods:

Broyden-Fletcher-Goldfarb-Shannon (BFGS) method

Limited-memory Broyden-Fletcher-Goldfarb-Shannon (L-BFGS) method

What Optimizer to Use?

GS Method

General form of quasi-Newton methods is

$$\mathbf{x}_{n+1} = \mathbf{x}_n - \alpha \mathbf{H}_n^{-1} \nabla f(\mathbf{x}_n)$$

where α is line search parameter and is $\mathbf{H}_n^{-1}$ Hessian or some approximation to the Hessian.

assume

$$\begin{aligned} \mathbf{H}_n \mathbf{s}_n &= -\nabla f(\mathbf{x}_n), & \mathbf{x}_{n+1} &= \mathbf{x}_n + \alpha \mathbf{s}_n \\ \mathbf{y}_n &= \nabla f(\mathbf{x}_{n+1}) - \nabla f(\mathbf{x}_n) \end{aligned}$$

The updated approximated Hessian should satisfy the following under-determined system

$$\mathbf{H}_{n+1} \mathbf{s}_n = \mathbf{y}_n$$

The objective is to adjust $\mathbf{H}_n$ minimally to satisfy the above condition by ensuring $\mathbf{H}_n$ remains symmetric-positive-definite.

What Optimizer to Use?

GS Method

GS proposed

$$H_{n+1} = H_n + \beta \mathbf{u} \mathbf{u}^\top + \gamma \mathbf{v} \mathbf{v}^\top$$

Substitute $\mathbf{u} = \mathbf{y}_n$ and $\mathbf{v} = H_n \mathbf{s}_n$ and $H_n^\top = H_n$ to obtain

$$H_{n+1} = H_n + \beta \mathbf{y}_n \mathbf{y}_n^\top + \gamma H_n \mathbf{s}_n \mathbf{s}_n^\top H_n$$

ce

$$\begin{aligned} \underline{\mathbf{y}_n} &= \underline{H_{n+1} \mathbf{s}_n} = \underline{H_n \mathbf{s}_n} + \underline{\beta \mathbf{y}_n \mathbf{y}_n^\top \mathbf{s}_n} + \underline{\gamma H_n \mathbf{s}_n \mathbf{s}_n^\top H_n \mathbf{s}_n} \\ \beta &= \frac{1}{\mathbf{y}_n^\top \mathbf{s}_n} \quad \gamma = -\frac{1}{\mathbf{s}_n^\top H_n \mathbf{s}_n} \end{aligned}$$

Since the update approximation to Hessian is

$$H_{n+1} = H_n + \frac{\mathbf{y}_n \mathbf{y}_n^\top}{\mathbf{y}_n^\top \mathbf{s}_n} - \frac{H_n \mathbf{s}_n \mathbf{s}_n^\top H_n}{\mathbf{s}_n^\top H_n \mathbf{s}_n}$$

What Optimizer to Use?

BFGS optimizer

BFGS is the most popular of all Quasi-Newton methods but has storage complexity.

BFGS (Limited memory BFGS), which does not require to explicitly store H^{-1} but instead stores the previous data $\nabla f(x_i))\}_{i=1}^k$ and manages to compute $d = H^{-1} \nabla f(x)$ directly from this data. L-BFGS has storage complexity of $\mathcal{O}(n)$

BFGS implementation is not straightforward in PyTorch and TF2. A detailed implementation is discussed in Lecture 4, where we provide a simple API for both.

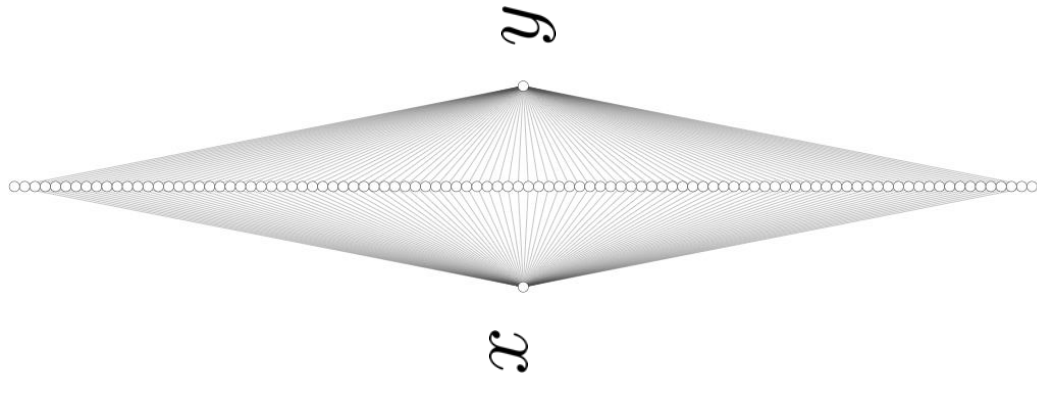
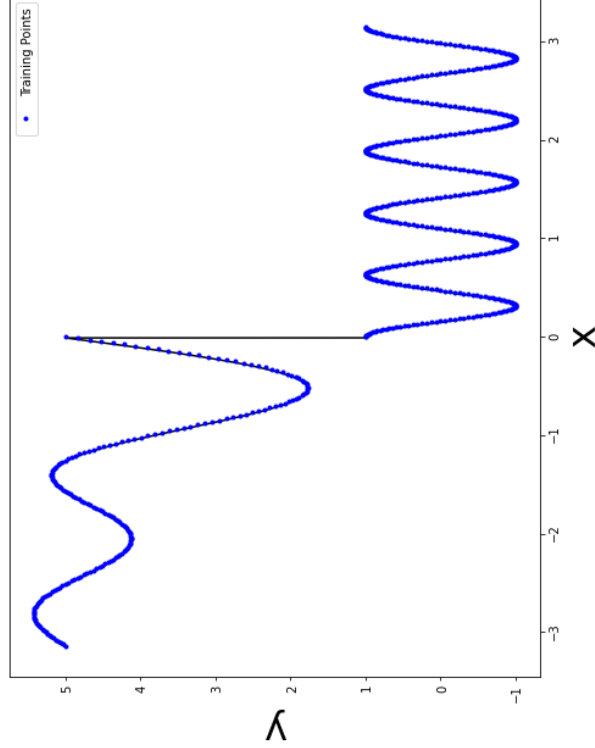


<pre> def bfgs(f: Callable[[np.ndarray], float], x0: np.ndarray, tolerance: float, max_iter: int, line_search: Callable[[np.ndarray, np.ndarray], float] = None, history_size: int = 100, line_search_fn: Callable[[np.ndarray], float] = None,) -> np.ndarray: """BFGS algorithm""" x = x0 for i in range(max_iter): d = -bfgs_step(f, x, tolerance, history_size, line_search_fn) x = x + d if np.linalg.norm(d) < tolerance: return x return x </pre>	<pre> torch.optim.LBFGS(params, lr=1, max_iter= 20, max_eval=None, tolerance_grad=1e- 07, tolerance_change=1e- 09, history_size=100, line_search_fn=None) </pre>	L-BFGS implementation in TensorFlow is provided through TF Probability package <pre> tfp.optimizer.lbfgs_minimize(f, initial_position=self.get_weights(), num_correction_pairs=50, max_iterations=2000) </pre>
---	--	--

+ Function Approximation + Adam + L-BFGS

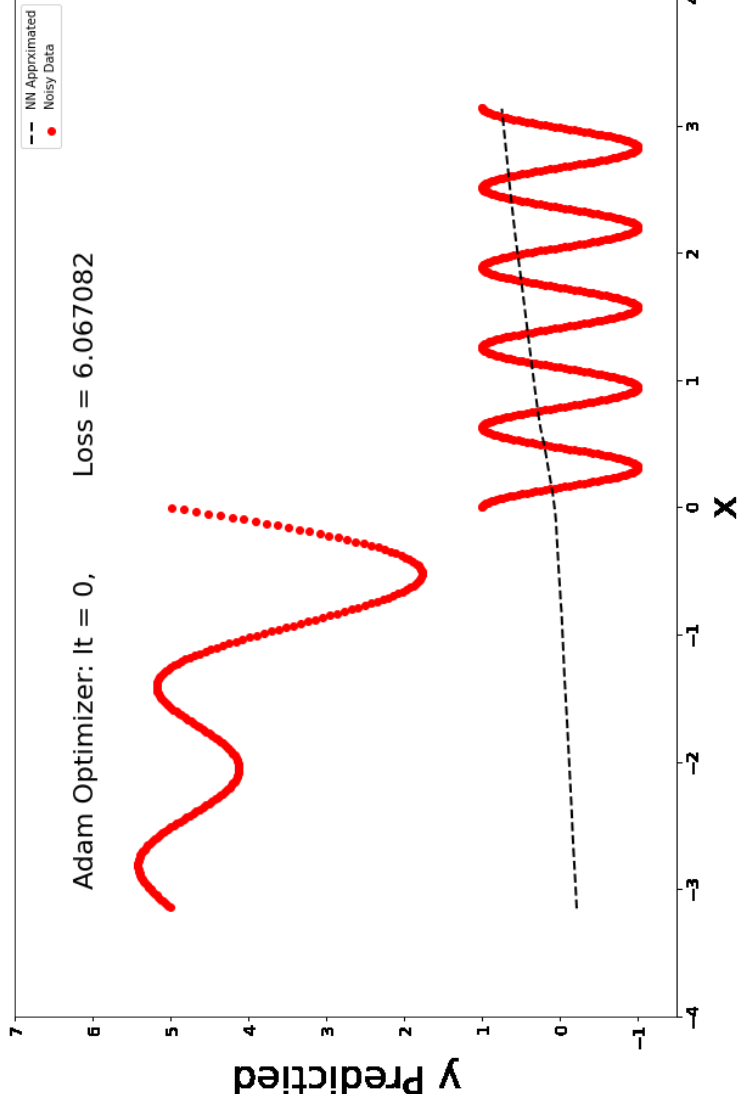
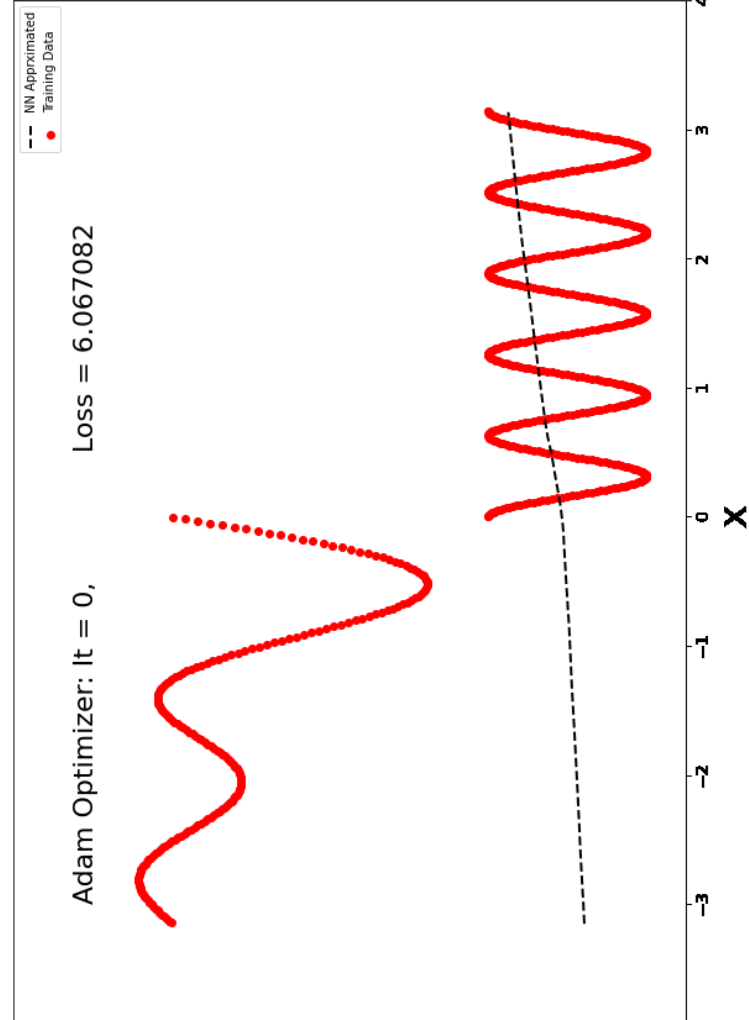
$$y = 5 + \sum_{k=1}^4 \sin(kx), \quad x < 0$$

$$y = \cos(10x), \quad x \geq 0$$



+ Function Approximations + Adam + L-BFGS

Adam (30 Iterations) + 50 L-BFGS (50 Iterations)



FASTER OPTIMIZATION: Self-Scaled Broyden's method

- Instead of computing the inverse Hessian matrix directly, the algorithm reformulates the scaling parameter calculation, avoiding the $O(N^3)$ complexity.

$$H_{k+1} = H_k + \underbrace{\sigma_k}_{\text{red circle}} \frac{(y_k - H_k s_k)(y_k - H_k s_k)^T}{(y_k - H_k s_k)^T s_k} - \frac{H_k s_k s_k^T H_k}{s_k^T H_k s_k}$$

$$\sigma_k = 1 + \underbrace{a_k \theta_k}_{\text{red wavy line}}$$

$$b_k = \frac{s_k \cdot H_k^{-1} s_k}{y_k \cdot s_k} = -\alpha_k \frac{s_k \cdot \nabla \mathcal{J}(\Theta_k)}{y_k \cdot s_k},$$

$$h_k = \frac{y_k \cdot H_k y_k}{y_k \cdot s_k},$$

$$\rightarrow a_k = h_k b_k - 1$$

$$c_k = \sqrt{\frac{a_k}{a_k + 1}},$$

$$\rho_k^- = \min(1, h_k(1 - c_k)),$$

$$\theta_k^- = \frac{\rho_k^- - 1}{a_k},$$

$$\theta_k^+ = \frac{1}{\rho_k^-},$$

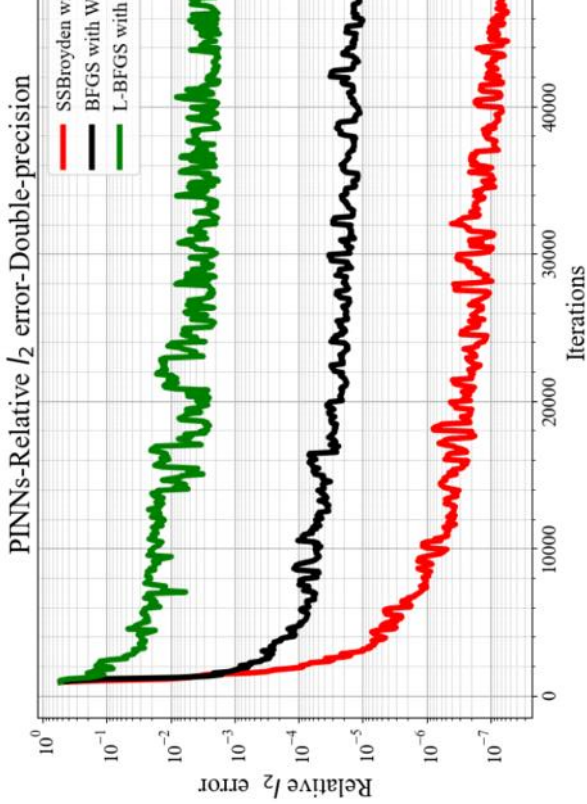
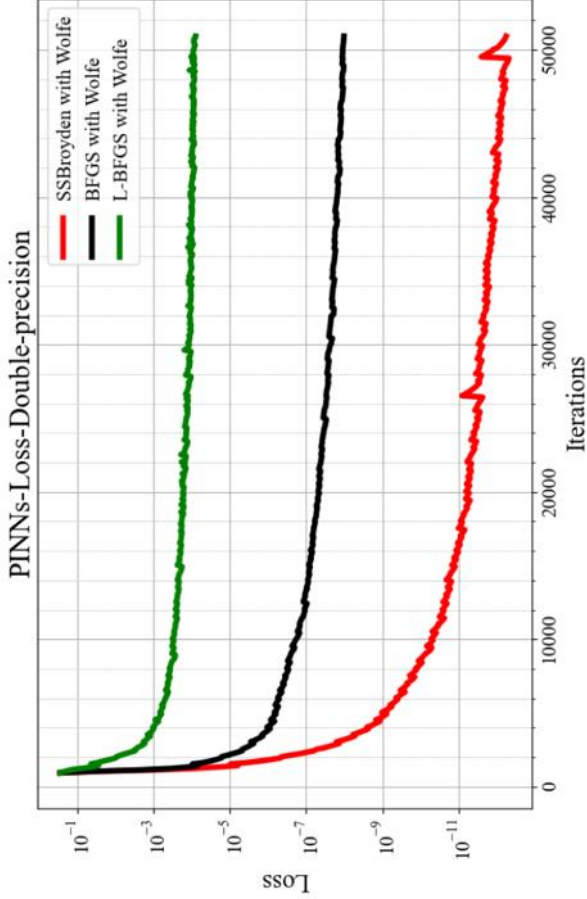
$$\rightarrow \theta_k = \max\left(\theta_k^-, \min\left(\theta_k^+, \frac{b_k}{1 - b_k}\right)\right)$$

Urbán, J. F., Stefanou, P., and Pons, J. A. (2024).

M. Al-Baali, Variational quasi-newton methods for unconstrained optimization, Journal of Optimization Theory and Applications 77 (1) (1993) 127–143. doi:10.1007/BF00940782.

izing PINNs for the Burgers Equation using SSBroyden

$$+ u \frac{\partial u}{\partial x} = \nu \frac{\partial^2 u}{\partial x^2}$$
$$(x, 0) = -\sin(\pi x)$$
$$= \frac{0.01}{\pi} \approx 0.003$$



Case	Optimizer[# Iters.], Line-search [#Iters.]	Relative l ₂ error	Training time (s)	Total par
2	Adam [1000] + BFGS with Wolfe [50000]	1.50 × 10 ⁻⁵	1292	1,341
2	Adam [1000] + SSBroyden with Wolfe [50000]	7.57 × 10 ⁻⁸	1354	1,341
2	Adam [1000] + L-BFGS with Wolfe [50000]	2.05 × 10 ⁻³	713	1,341
3	Adam [1000] + BFGS with Wolfe [30000]	2.21 × 10 ⁻⁵	774	1,341
3	Adam [1000] + SSBroyden with Strong Wolfe [30000]	2.12 × 10 ⁻⁷	819	1,341
4	BFGS with Wolfe [30000]	1.73 × 10 ⁻⁵	863	1,341
4	SSBroyden with Wolfe [30000]	3.59 × 10 ⁻⁷	862	1,341
5	Adam [1000] + BFGS with Wolfe [30000]	8.52 × 10 ⁻⁶	3049	3,021
5	Adam [1000] + SSBroyden with Wolfe [30000]	1.62 × 10 ⁻⁸	2812	3,021

Optimizing PINNs for the Kuramoto-Sivashinsky Equation using SSBroyden

$$-\alpha u \frac{\partial u}{\partial x} + \beta \frac{\partial^2 u}{\partial x^2} + \gamma \frac{\partial^4 u}{\partial x^4} = 0$$

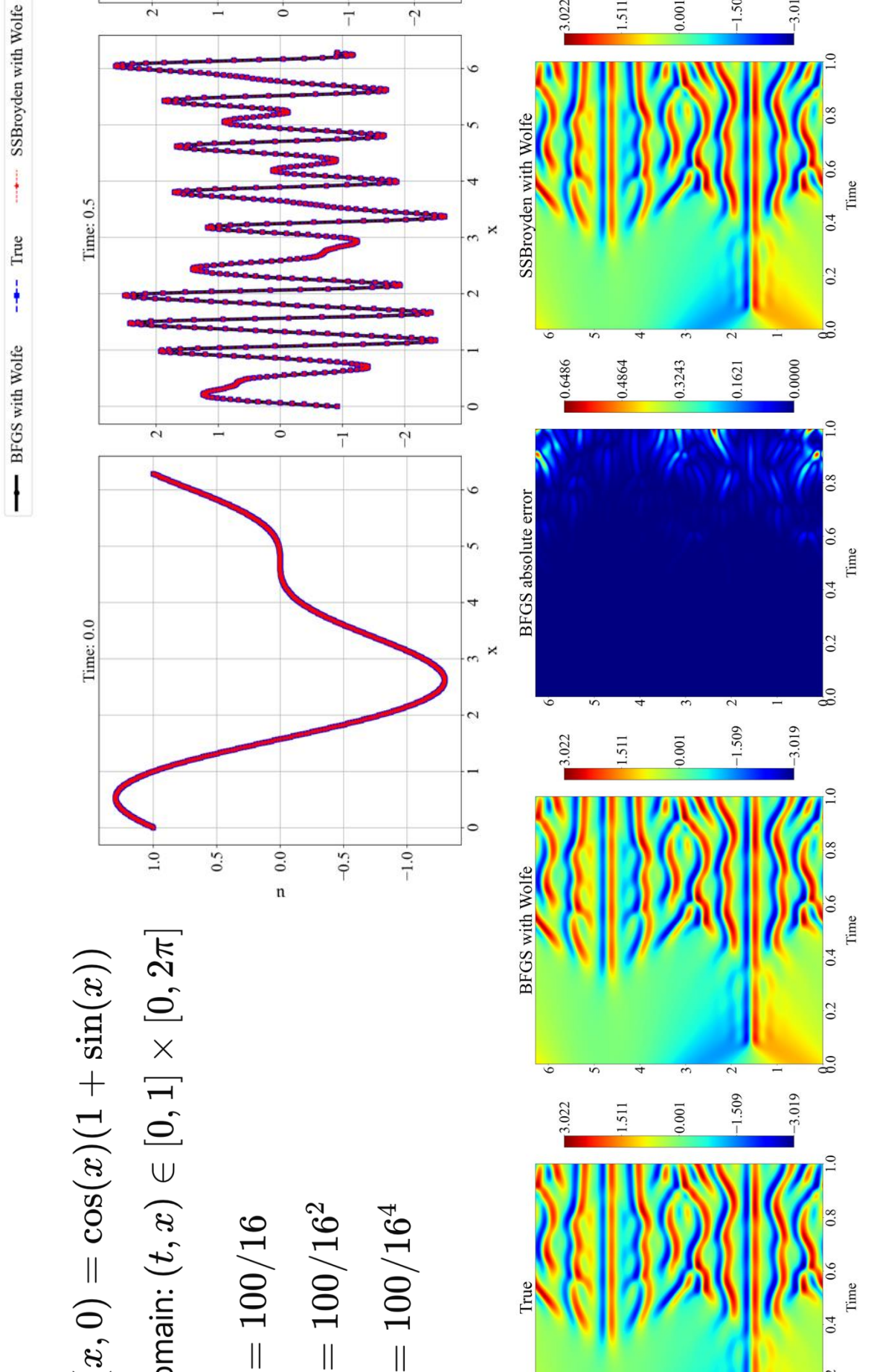
$$u(x, 0) = \cos(x)(1 + \sin(x))$$

$$\text{Domain: } (t, x) \in [0, 1] \times [0, 2\pi]$$

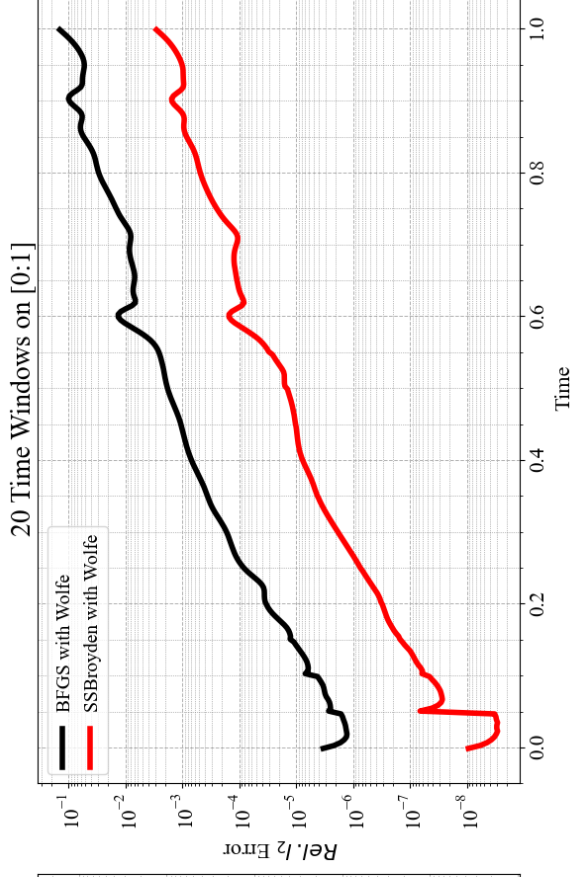
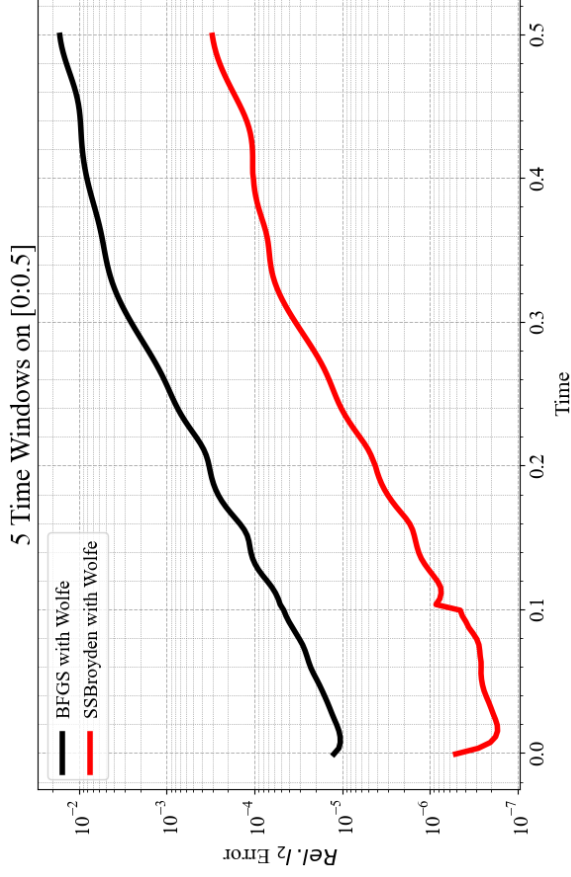
$$= 100/16$$

$$= 100/16^2$$

$$= 100/16^4$$

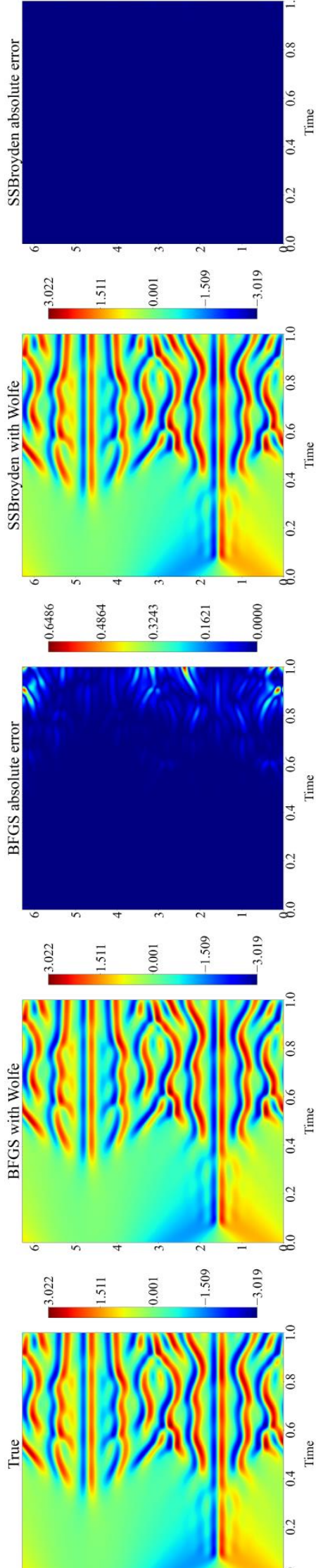


izing PINNs for the Kuramoto-Sivashinsky Equation using SSBroyden



Case	Optimizer[# Iters.], Line-search [#Iters.]	Relative l_2 error	Training time (s)	Total parameters
1	BFGS with Wolfe [20000] (20 windows)	$8.75e \times 10^{-2}$	102177	4,411
1	SSBroyden with Wolfe [20000] (20 windows)	2.13×10^{-3}	113214	4,411
2	BFGS with Wolfe [30000] (20 windows)	3.65×10^{-2}	119286	4,411
2	SSBroyden with Wolfe [30000] (20 windows)	6.51×10^{-4}	130222	4,411
3	Adam [1000] + BFGS with Wolfe [30000] (20 windows)	6.15×10^{-2}	119353	4,411
3	Adam [1000] + SSBroyden with Wolfe [30000] (20 windows)	7.53×10^{-4}	127372	4,411
4	BFGS with Wolfe [20000] (5 windows)	7.54×10^{-3}	17793	4,411
4	SSBroyden with Wolfe [20000] (5 windows)	1.24×10^{-4}	19048	4,411
5	BFGS with Wolfe [30000] (5 windows)	2.39×10^{-3}	26697	4,411
5	SSBroyden with Wolfe [30000] (5 windows)	2.65×10^{-5}	30883	4,411

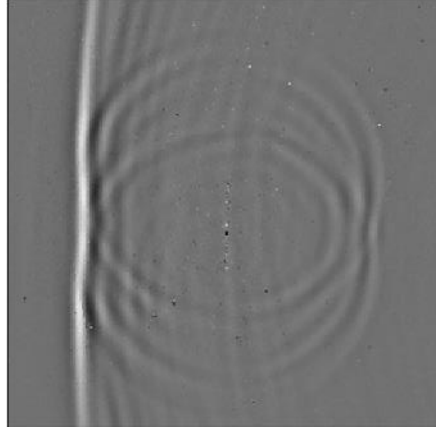
PINNs 2.0 – Better Optimizers



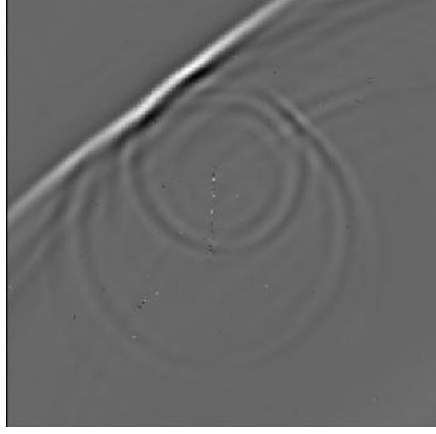
Benchmark	$\mathcal{L}_2$ Error (SOAP [56])	$\mathcal{L}_2$ Error (SSbroyden with Wolfe)
Burgers Equation	8.06×10^{-6}	4.19×10^{-8}
Allen-Cahn Equation	3.48×10^{-6}	9.43×10^{-7}
Kuramoto-Sivashinsky Equation	3.86×10^{-2} on $t \in [0, 0.8]$	2.65×10^{-5} on $t \in [0, 1]$
Ginzburg-Landau Equation	4.78×10^{-3}	2.33×10^{-3}

Crack characterization in an Aluminum Alloy using PINNs

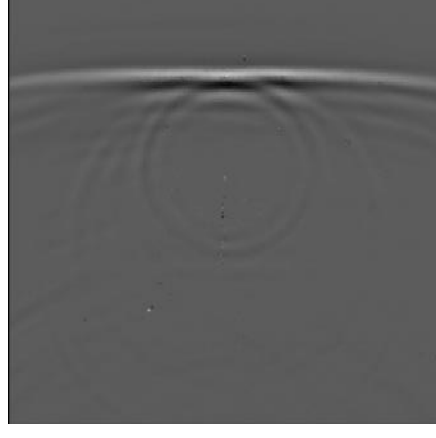
- **Objective:** Identifying and characterizing a surface breaking crack in a 7075-T651 aluminum alloy metal plate.
- **Dataset:** Particle displacement induced by surface acoustic waves of 5 MHz at three incidence angle 0° , 45° , and 90° .
- **Approach:** Reconstruct the wave field and compute the speed of sound using Inverse PINNs



0° incidence



45° incidence

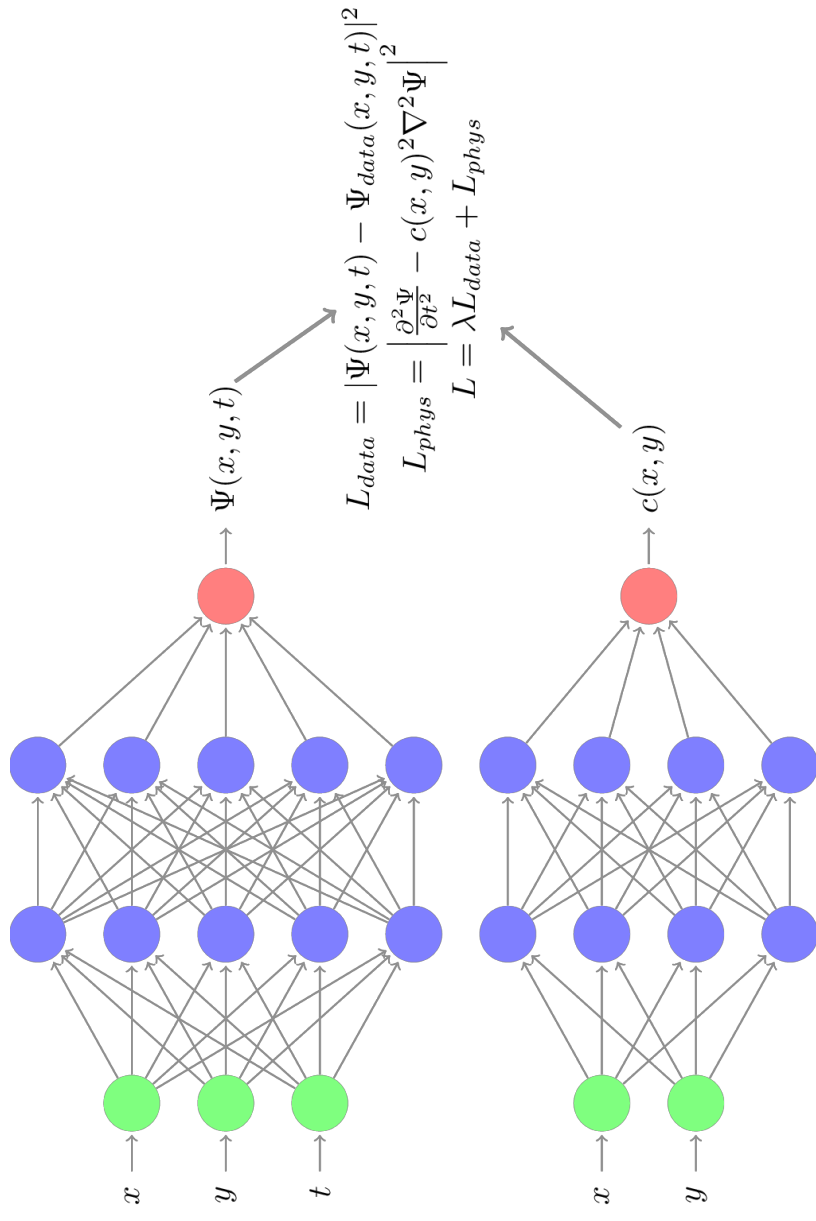


90° incidence

- K. Shukla, P. Clark Di Leoni, J. Blackshire, D. Sparkman, G. E. Karniadakis, “Physics-informed Neural Network for Ultrasound Nondestructive Quantification of Surface Breaking Cracks. *J Nondestruct Eval* 39, 61 (2020). <https://doi.org/10.1007/s10921-020-00705-1>

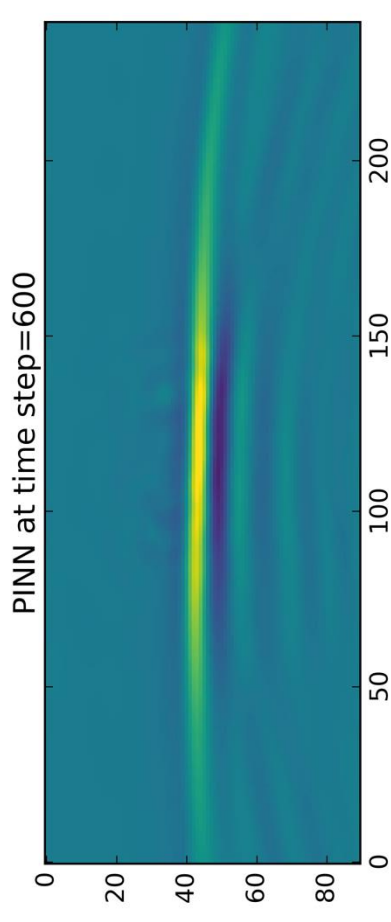
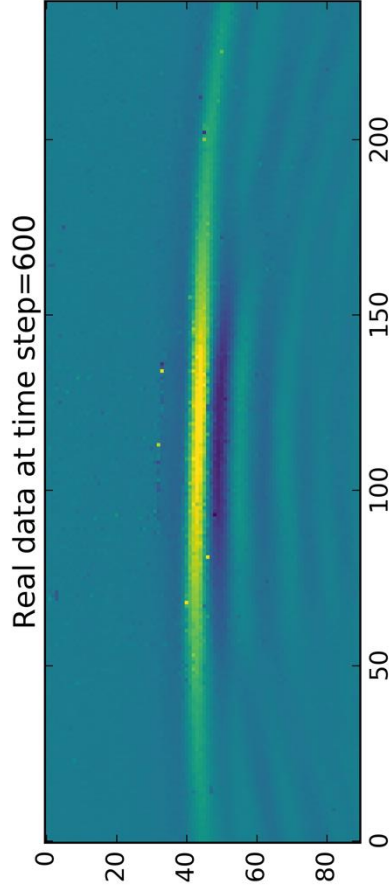
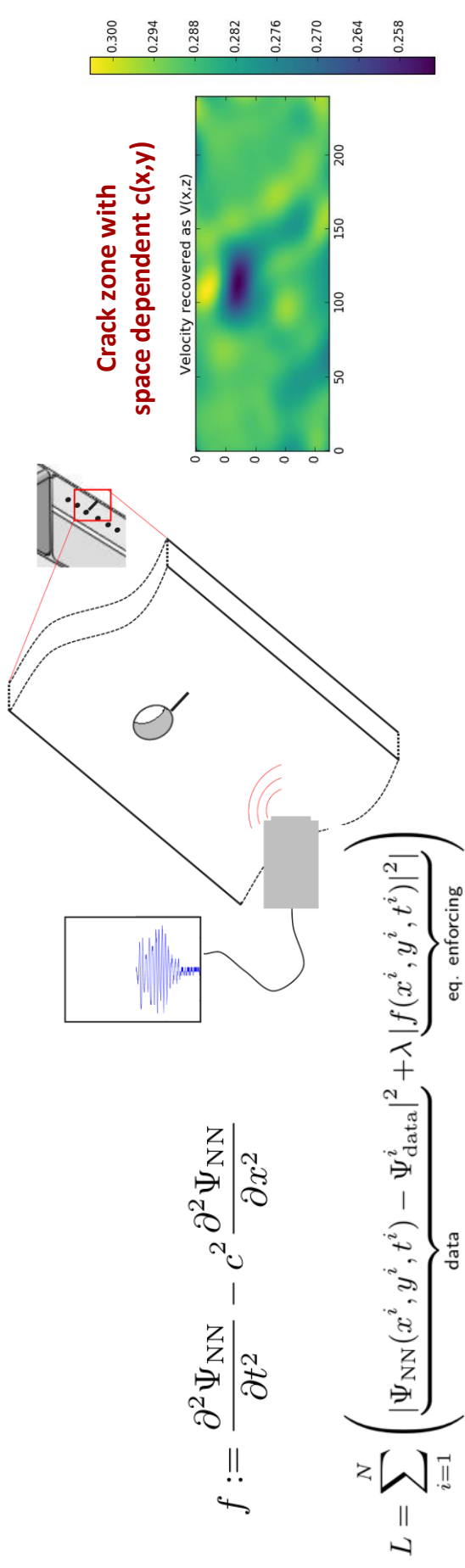
Problem Setup and Method

Physics-informed Neural Network



- Train a Neural Network informed with the acoustic wave equation using sparse data. The speed of sound in the material is approximated by a second neural network

Ultra-Sound Testing of Materials – WPAFB Real Data



PINNs: Training For Boundary Value Problems

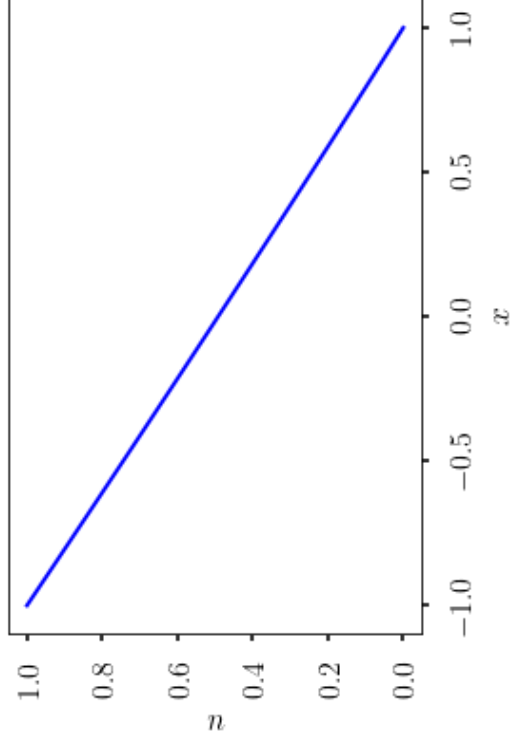
- Define an ODE for given boundary conditions

$$\nu \frac{d^2 u}{dx^2} - u = e^x,$$

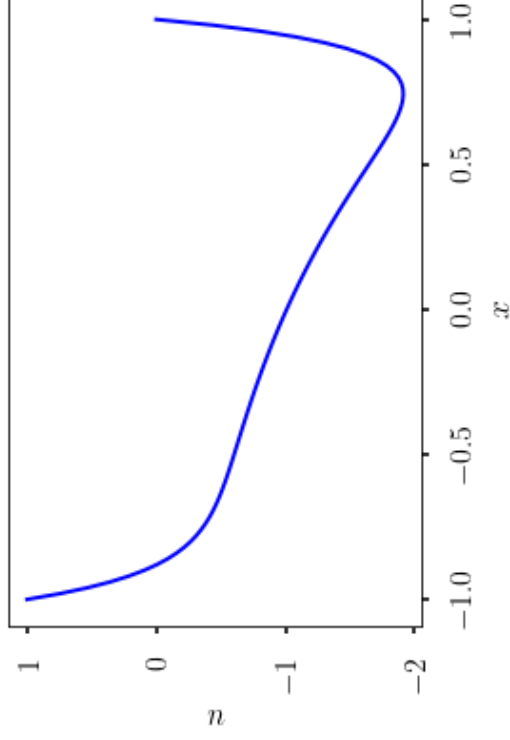
where $x \in [-1, 1]$, $u(-1) = 1, u(1) = 0$.

- Boundary Layer vs Viscosity

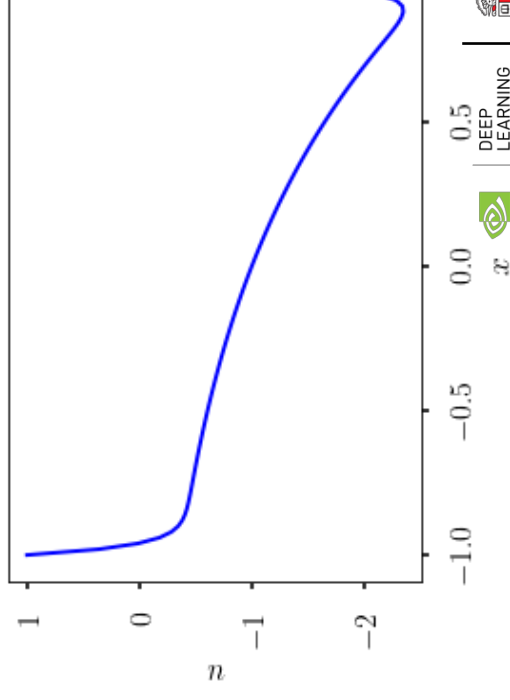
$\nu = 10^2$



$\nu = 10^{-2}$

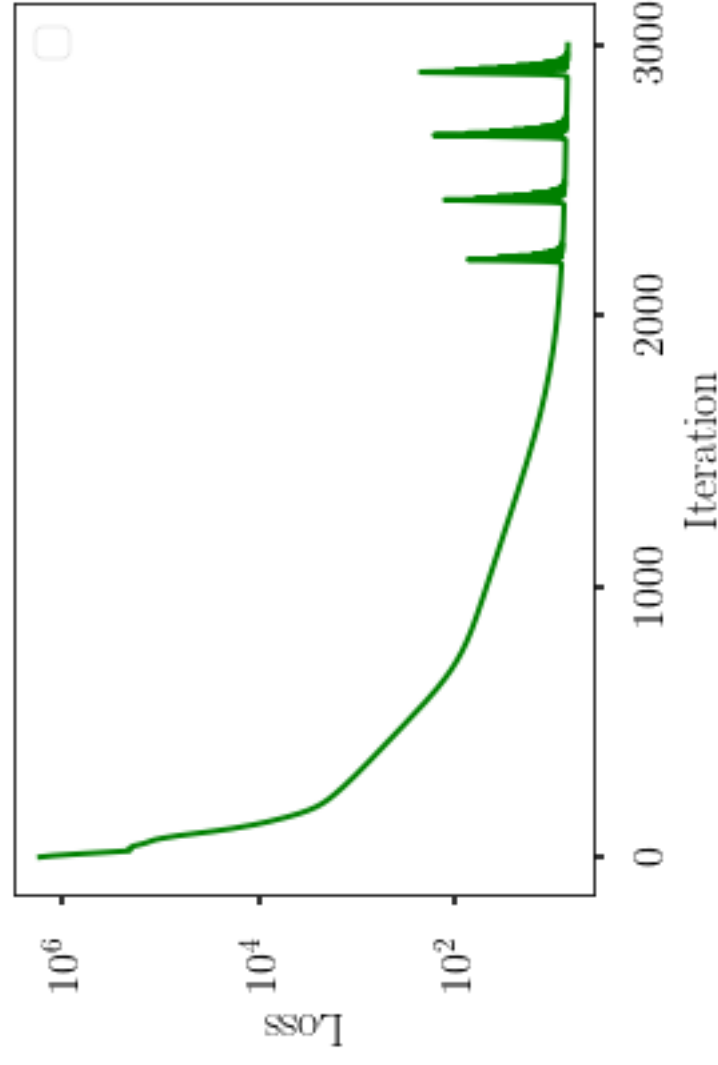
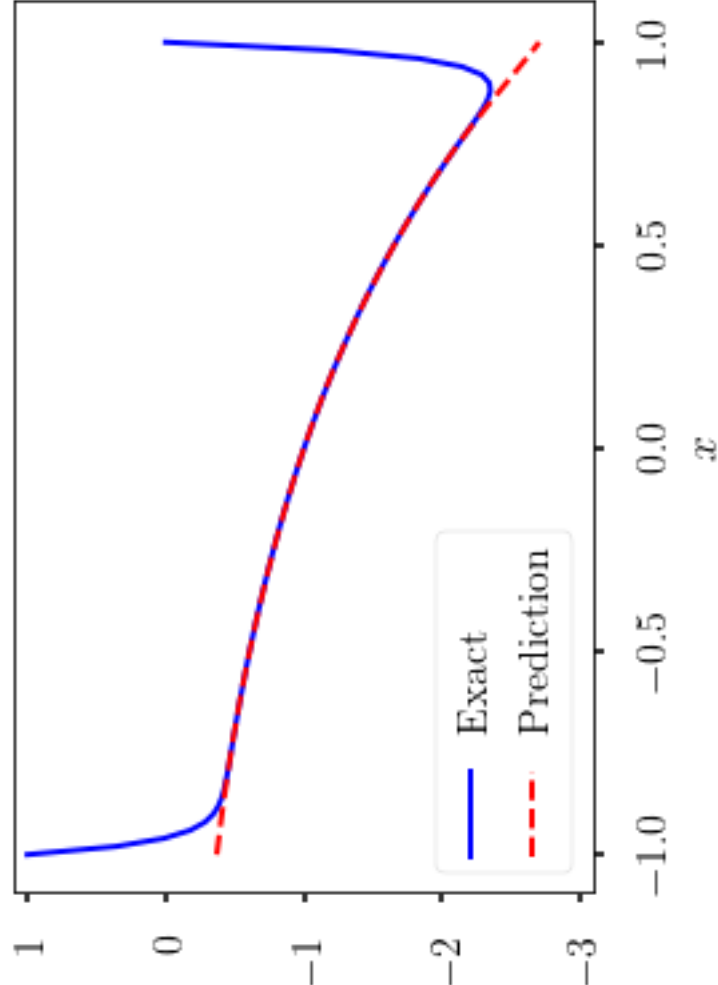


$\nu = 10^{-3}$



Approach 1: PINNs with Soft Constraints

- Results



proach 2: Self-Adaptive PINNs

is Defined as

$$\begin{aligned}\mathcal{N}_{\mathbf{x},t}[u(\mathbf{x},t)] &= 0, & \mathbf{x} \in \Omega, t \in [0,T] \\ u(\mathbf{x},t) &= g(\mathbf{x},t), & \mathbf{x} \in \partial\Omega, t \in [0,T] \\ u(\mathbf{x},0) &= h(\mathbf{x}), & \mathbf{x} \in \Omega\end{aligned}$$

function

$$\mathcal{L}(\mathbf{w}) = \mathcal{L}_s(\mathbf{w}) + \mathcal{L}_r(\mathbf{w}) + \mathcal{L}_b(\mathbf{w}) + \mathcal{L}_0(\mathbf{w})$$

$$\mathcal{L}_s(\mathbf{w}) = \frac{1}{N_s} \sum_{i=1}^{N_s} \left| \underbrace{u(\mathbf{x}_s^i, t_s^i; \mathbf{w})}_{\text{physics}} - \underbrace{y_s^i}_{\text{data}} \right|^2$$

$$\mathcal{L}_r(\mathbf{w}) = \frac{1}{N_r} \sum_{i=1}^{N_r} \left| \underbrace{r(\mathbf{x}_r^i, t_r^i; \mathbf{w})}_{\text{physics}} \right|^2$$

$$\mathcal{L}_b(\mathbf{w}) = \frac{1}{N_b} \sum_{i=1}^{N_b} \left| \underbrace{u(\mathbf{x}_b^i, t_b^i; \mathbf{w})}_{\text{physics}} - \underbrace{g_b^i}_{\text{data}} \right|^2$$

$$\mathcal{L}_0(\mathbf{w}) = \frac{1}{N_0} \sum_{i=1}^{N_0} \left| \underbrace{u(\mathbf{x}_0^i, 0; \mathbf{w})}_{\text{physics}} - \underbrace{h_0^i}_{\text{data}} \right|^2$$

- McClenny L, Braga-Neto U. Self-adaptive physics-informed neural networks using a soft attention mechanism. arXiv preprint arXiv:2009.04544. 2020.

Self-adaptive PINN utilizes the following loss function

$$\mathcal{L}(w, \lambda_r, \lambda_b, \lambda_0) = \mathcal{L}_s(w) + \mathcal{L}_r(w, \lambda_r) + \mathcal{L}_b(w, \lambda_b) + \mathcal{L}_0(w, \lambda_0)$$

where $\lambda_r = (\lambda_r^1, \dots, \lambda_r^{N_r})$, $\lambda_b = (\lambda_b^1, \dots, \lambda_b^{N_b})$, and $\lambda_0 = (\lambda_0^1, \dots, \lambda_0^{N_0})$ are trainable self-adaptation weights for the initial, boundary, and collocation points, respectively, and

$$\begin{aligned}\mathcal{L}_r(w, \lambda_r) &= \frac{1}{N_r} \sum_{i=1}^{N_r} [\lambda_r^i r(\mathbf{x}_r^i, t_r^i; w)]^2 \\ \mathcal{L}_b(w, \lambda_b) &= \frac{1}{N_b} \sum_{i=1}^{N_b} [\lambda_b^i (u(\mathbf{x}_b^i, t_b^i; w) - g_b^i)]^2 \\ \mathcal{L}_0(w, \lambda_0) &= \frac{1}{N_0} \sum_{i=1}^{N_0} [\lambda_0^i (u(\mathbf{x}_0^i, 0; w) - h_0^i)]^2\end{aligned}$$

The key feature of self-adaptive PINNs is that the loss $\mathcal{L}(w, \lambda, \lambda_b, \lambda_0)$ is minimized with respect to the network weights w , as usual, but is maximized with respect to the self-adaptation weights λ_r , in other words, training seeks a saddle point

$$\min_w \max_{\lambda_r, \lambda_b, \lambda_0} \mathcal{L}(w, \lambda_r, \lambda_b, \lambda_0).$$

This can be accomplished by a gradient descent/ascent procedure, with updates given by:

gradient descent

$$w^{k+1} = w^k - \eta_k \nabla_w \mathcal{L}(w^k, \lambda_r^k, \lambda_b^k, \lambda_0^k)$$

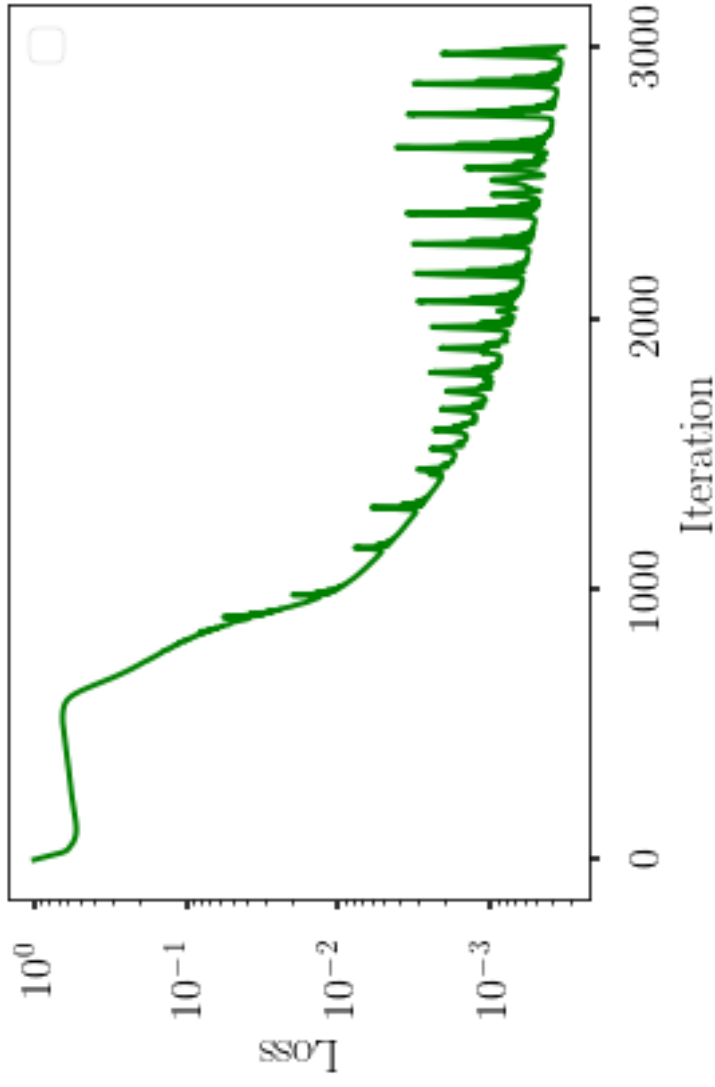
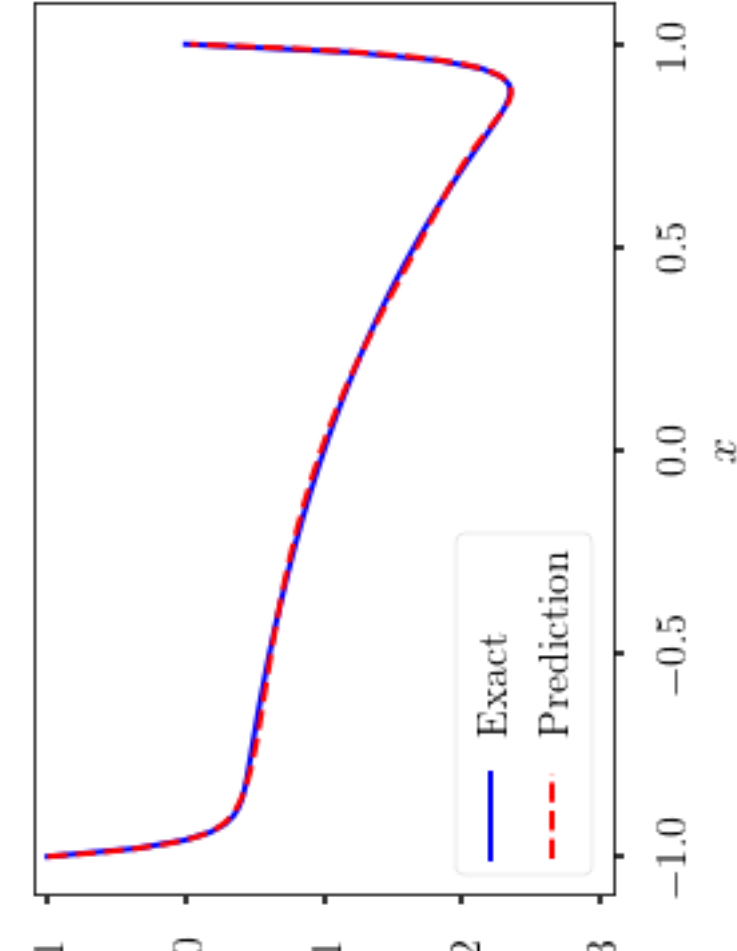
$$\lambda_r^{k+1} = \lambda_r^k + \eta_k \nabla_{\lambda_r} \mathcal{L}(w^k, \lambda_r^k, \lambda_b^k, \lambda_0^k)$$

$$\lambda_b^{k+1} = \lambda_b^k + \eta_k \nabla_{\lambda_b} \mathcal{L}(w^k, \lambda_r^k, \lambda_b^k, \lambda_0^k)$$

$$\lambda_0^{k+1} = \lambda_0^k + \eta_k \nabla_{\lambda_0} \mathcal{L}(w^k, \lambda_r^k, \lambda_b^k, \lambda_0^k)$$

gradient ascent

proach 2: Results



Self-Adaptive PINN Architecture	
Hyper-parameters	Value
No. of Layers	8
Neurons per layer	8
No. of Iterations	1500
Learning rate	1×10^{-3}
No. of Residuals Points	500

Residual Based Attention (RBA) Weights

(Anagnostopoulos et al., CMAME 2024)

Weighting strategies:

Modify loss terms contribution during training

Fixed global multipliers (per loss term):

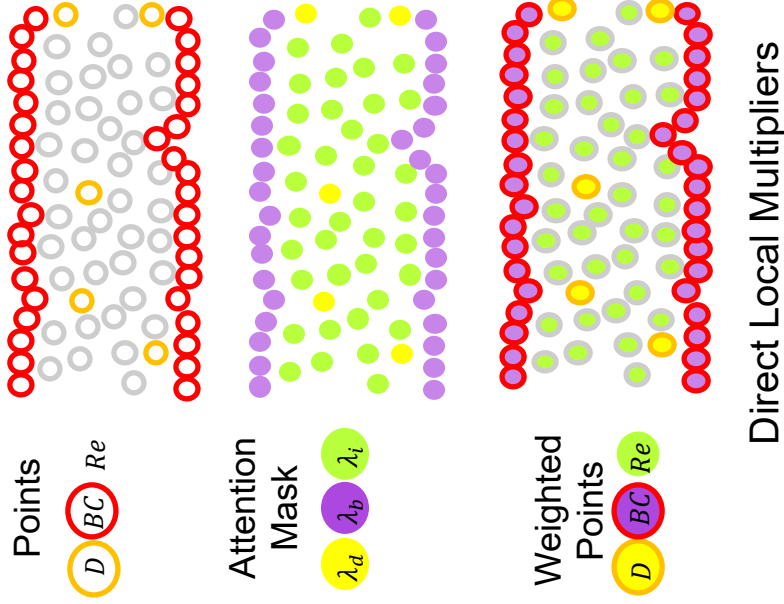
$$\mathcal{L} = \lambda_D \mathcal{L}_{Data} + \lambda_B \mathcal{L}_{BCS} + \lambda_{PDE} \mathcal{L}_{PDE}$$

Adaptive local multipliers (per training point):

$$\mathcal{L} = \langle (\lambda_d \cdot e_d)^2 \rangle + \langle (\lambda_b \cdot e_b)^2 \rangle + \langle (\lambda_i \cdot r_i)^2 \rangle$$

Requirements for λ_i^k :

- Identify high error regions $\Rightarrow f(r_i)$
- Efficient and easy to compute $\Rightarrow |r_i|$
- Stable $\Rightarrow \eta |r_i| / |r_{max}|$
- Remember the previous iterations $f(\lambda_i^{k-1})$
- Relate to the current training state $\gamma f(\lambda_i^{k-1})$



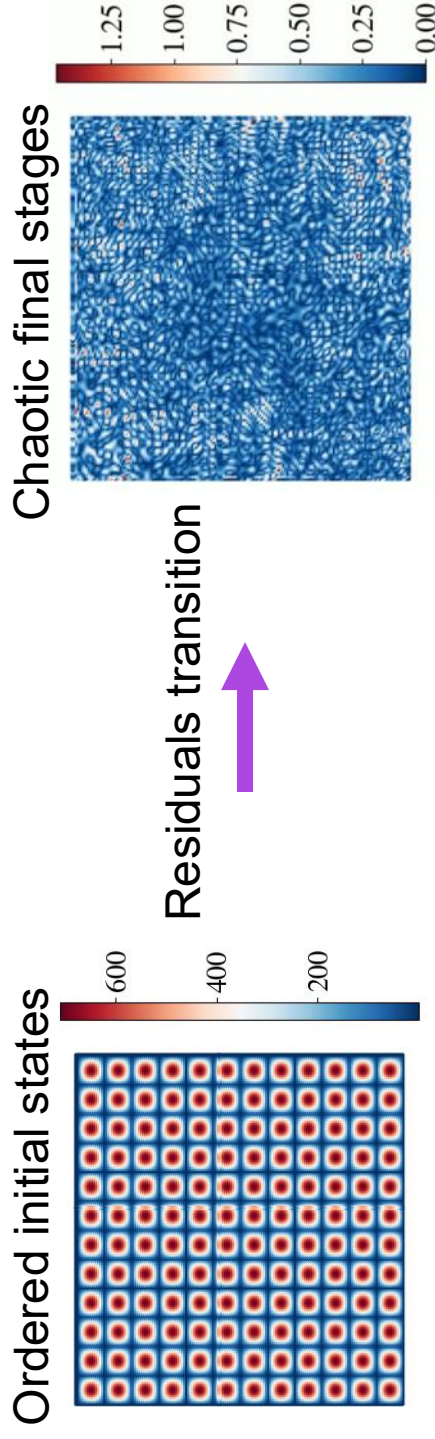
Residual Based Attention (RBA) Weights

(Anagnostopoulos et al., CNAME 2024)

Residual-based attention weights (Cumulative residuals)

$$\lambda_i^k = \gamma \lambda_i^{k-1} + \frac{|r_i|}{\max_{0 \leq j \leq N} |r_j|}$$

Helmholtz
Equation

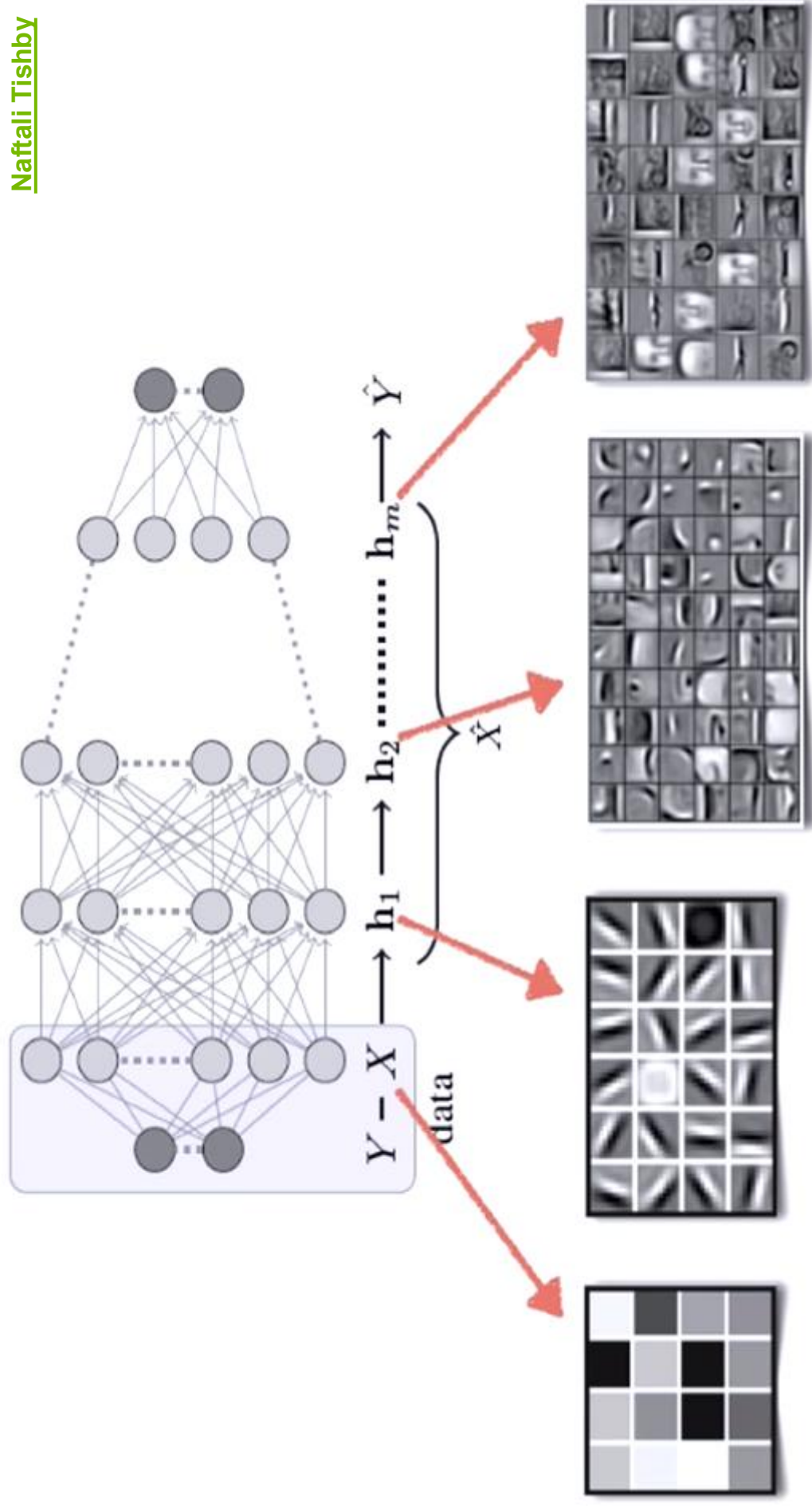


Weight and residual evolution reveal two distinct phases akin to “fitting” and “diffusion”.

Information Bottleneck Theory for Deep Learning



[Naftali Tishby](#)



Information Bottleneck Theory for Deep Learning

□ Entropy: $H(p) = - \sum_i p_i \log_2(p_i)$

□ Cross-Entropy: $H(p, q) = - \sum_i p_i \log_2(q_i)$

□ KL Divergence: $D(p||q) = H(p, q) - H(p)$

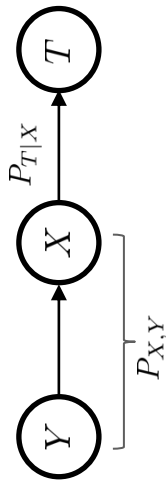
□ KL-distribution divergence: for any two distributions $p(x), q(x)$ over X :

$$D[p(x)||q(x)] = \sum_x p(x) \log \frac{p(x)}{q(x)} \geq 0$$

Information Bottleneck Theory for Deep Learning

for any two random variables X, Y

The mutual Information: $I(X; Y) = D[p(x, y) || p(x)p(y)] = D[p(x|y) || p(x)] = D[p(y|x) || p(y)]$



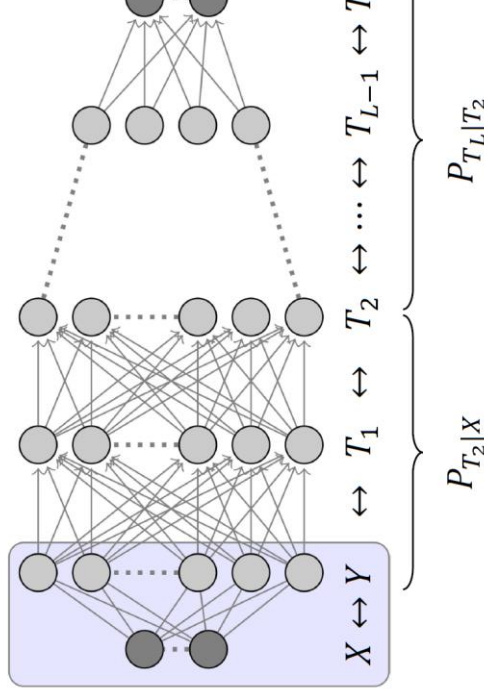
Graphical representation of probabilistic relations in framework. The triple (X, Y, T) is jointly distributed according to $P_{X,Y}P_{T|X}$, thus forming a Markov chain $Y \leftrightarrow X \leftrightarrow T$. The goal is to find a compressed representation T of X (via the transition kernel $P_{T|X}$), i.e., minimize $I(X; T)$, while preserving at least α bits of the information about Y , i.e., $I(Y; T) \geq \alpha$

Data Processing Inequality (DPI) & Invariance:

for any Markov chain $X \leftrightarrow Y \leftrightarrow Z$: $I(X; Y) \geq I(X; Z)$

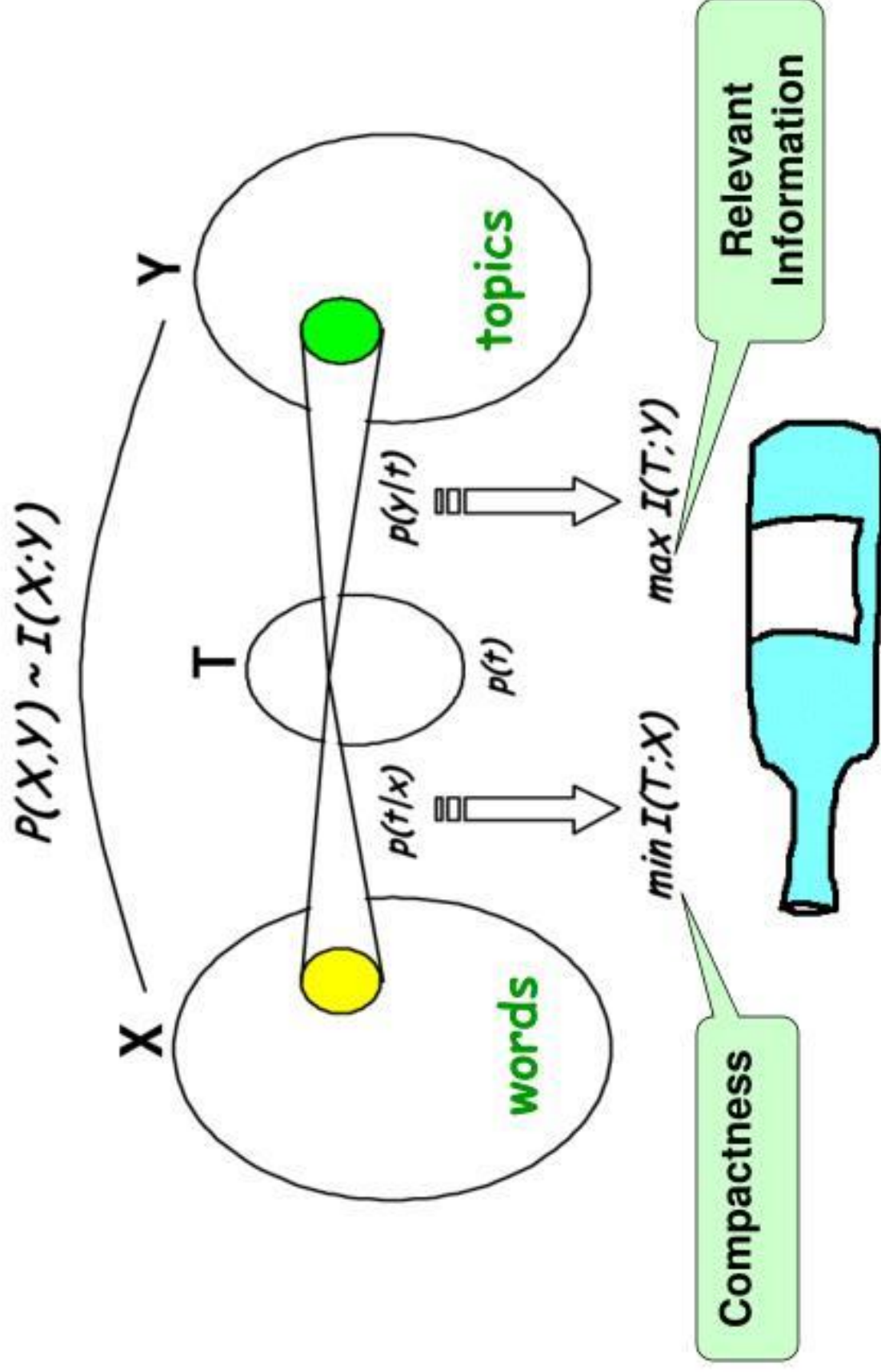
Parameterization Invariance for invertible ϕ, ψ : $I(X; Y) = I(\phi(x); \psi(x))$

Information Plane: For large typical X , the sample complexity of a DNN is completely determined by the encoder mutual information $I(X; T)$, of the last hidden layer. The accuracy (generalization error) is determined by the decoder information $I(T; Y)$ of the last hidden layer



1. Tishby, Naftali, and Noga Zaslavsky. "Deep learning and the information bottleneck principle." *Information Theory Workshop (ITW), 2015 IEEE*. IEEE, 2015.
2. Shwartz-Ziv, Ravid, and Naftali Tishby. "Opening the black box of deep neural networks via information." *arXiv preprint arXiv:1703.00810* (2017).

Information Bottleneck



Grange Dual Form and Phase Transition

$$\mathcal{L}_\beta(P_{T|X}) := I(X; T) - \beta I(T; Y)$$

With this definition, the IB problem can be recast as minimizing $\mathcal{L}_\beta(P_{T|X})$ over all possible $P_{T|X}$

β controls the amount of compression in the representation T of X :

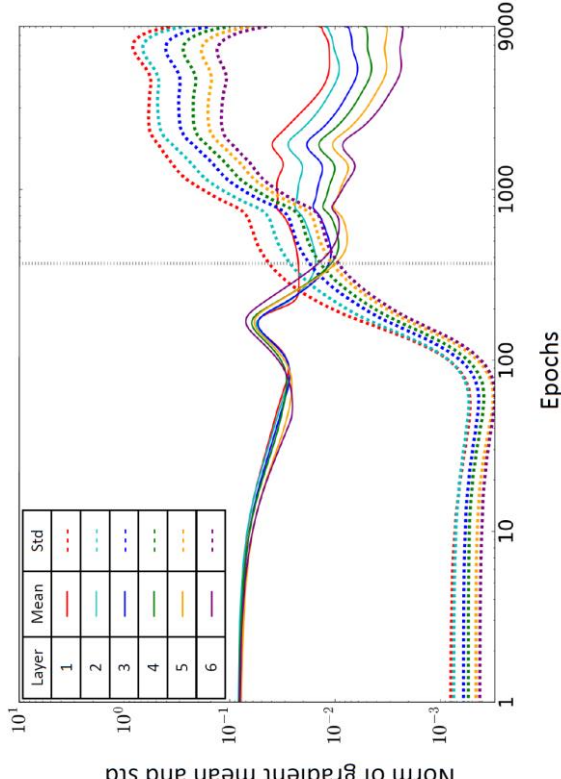
- ❑ Small β implies more compression (sacrificing informativeness)
- ❑ Larger β pushes towards finer representation granularity that favors informativeness.
- ❑ Varying β regulates the tradeoff between informativeness and compression

The working assumption of IB theory is that by optimizing the DNN's layers $T_l \{l = 1, \dots, L\}$ for that task via standard SGD, the encoder-decoder pair for each hidden layer converges to optimal IB solution.

The compression phase was further claimed to be responsible for the outstanding generalization performance of deep networks although no rigorous connection between $I(X; T_l)$ and the generalization error is currently known

A. M. Saxe, Y. Bansal, J. Dapello, M. Advani, A. Kolchinsky, B. D. Tracey, and D. D. Cox, "On the information bottleneck theory of deep learning," in Proceedings of the International Conference on Learning Representations (ICLR), 2018.
B. Z. Goldfeld, E. van den Berg, K. Greenewald, I. Melnyk, N. Nguyen, B. Kingsbury, and Y. Polyanskiy, "Estimating information flow in neural networks," in Proceedings of the International Conference on Machine Learning (ICML-2019), vol. Long Beach, California, USA, Jun. 2019.

Magrange Dual Form and Phase Transition: Summary



Norms of gradient mean (solid lines) and standard deviation (dashed lines) for the different layers. For each layer, the values are normalized by the L2 norm of the corresponding weight matrix. The grey vertical line marks the epoch when a phase transition between a high-SNR to a low-SNR occurs.

- ❑ First, they observed that the training process comprises two consecutive phases, termed ‘fitting’ and ‘compression’: after a quick initial fitting phase, almost the entire training process is dedicated to compressing X , while staying informative about Y .
- ❑ Second, it was observed that the compression phase starts when the mean and variance of the stochastic gradient undergoes a phase transition from high to low signal-to-noise (SNR).
- ❑ Third, it was observed that the two training phases accelerate when more layers are added to the network, which led to an argument that the benefit of deeper architectures is computational.

NEW: More recent results identified that for deterministic DNNs, IB does not apply. They also disputed the compression phase and argued that the geometric clustering of representations is the fundamental phenomenon of interest, while elucidated some of the machinery DNNs employ for learning. Leveraging the clustering perspective, evidence shows that compression and generalization may not be causally related.

A. M. Saxe, Y. Bansal, J. Dapello, M. Advani, A. Kolchinsky, B. D. Tracey, and D. D. Cox, “On the information bottleneck theory of deep learning,” in Proceedings of the International Conference on Learning Representations (ICLR), 2018.

B. Z. Goldfeld, E. van den Berg, K. Greenewald, I. Melnik, N. Nguyen, B. Kingsbury, and Y. Polyanskiy, “Estimating information flow in neural networks,” in Proceedings of the International Conference on Machine Learning (ICML-2019), vol. Long Beach, California, USA, Jun. 2019.

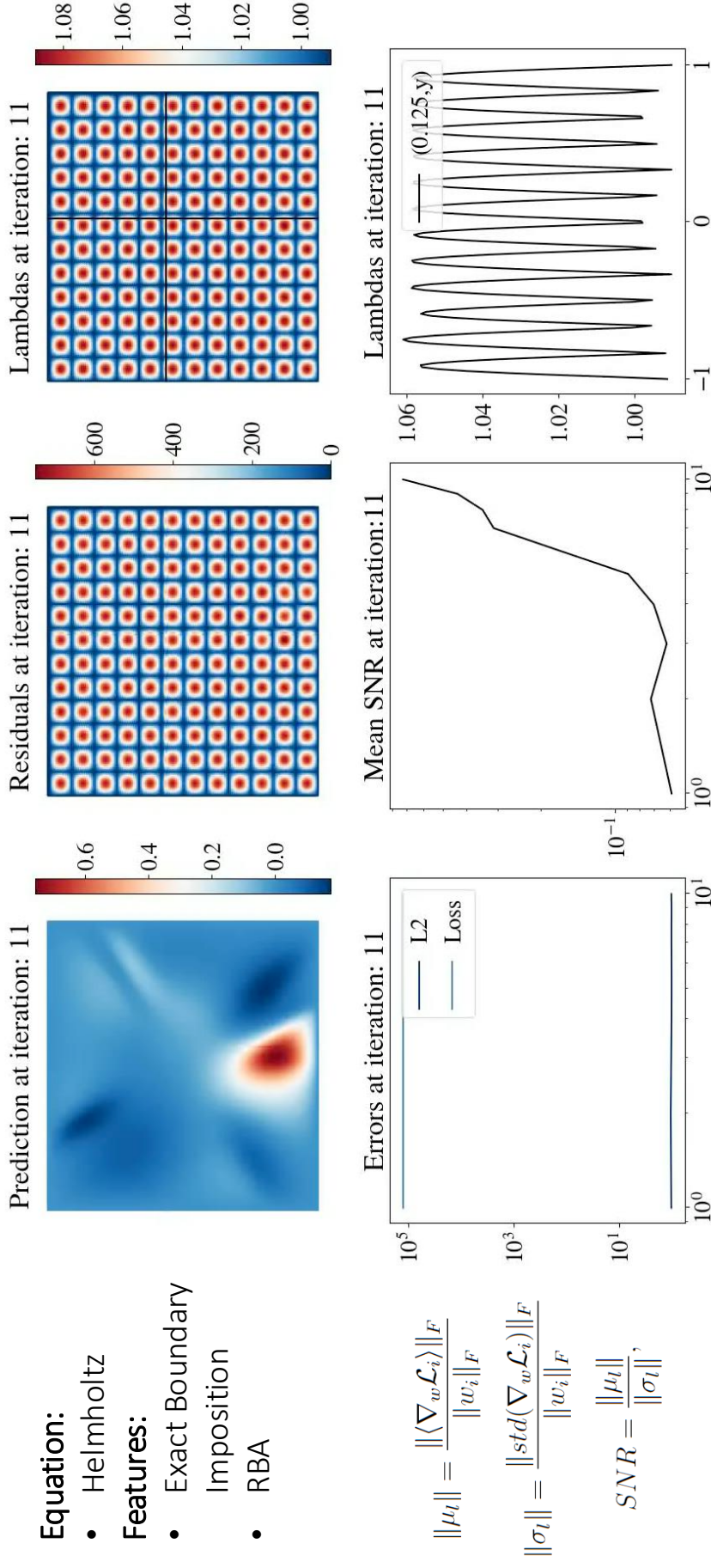
Information Bottleneck Theory

Equation:

- Helmholtz

Features:

- Exact Boundary Imposition
- RBA



$$\|\mu_l\| = \frac{\|\langle \nabla_w \mathcal{L}_i \rangle\|_F}{\|w_i\|_F}$$

$$\|\sigma_l\| = \frac{\|std(\nabla_w \mathcal{L}_i)\|_F}{\|w_i\|_F}$$

$$SNR = \frac{\|\mu_l\|}{\|\sigma_l\|},$$

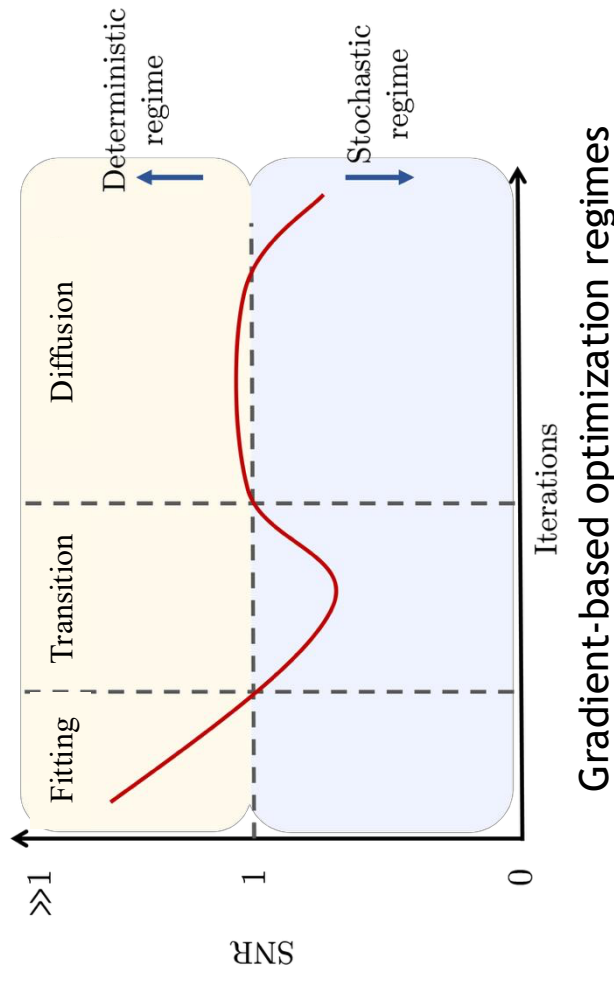
IB Theory: Learning Regimes

- **Fitting:** Neural gradients undergo a phase transition from high to low SNR
- **Transition:** The gradient fluctuations dominate the signal.
- **Diffusion:** Sudden jump on SNR due to gradient homogeneity (i.e., inner batches agree with each other)

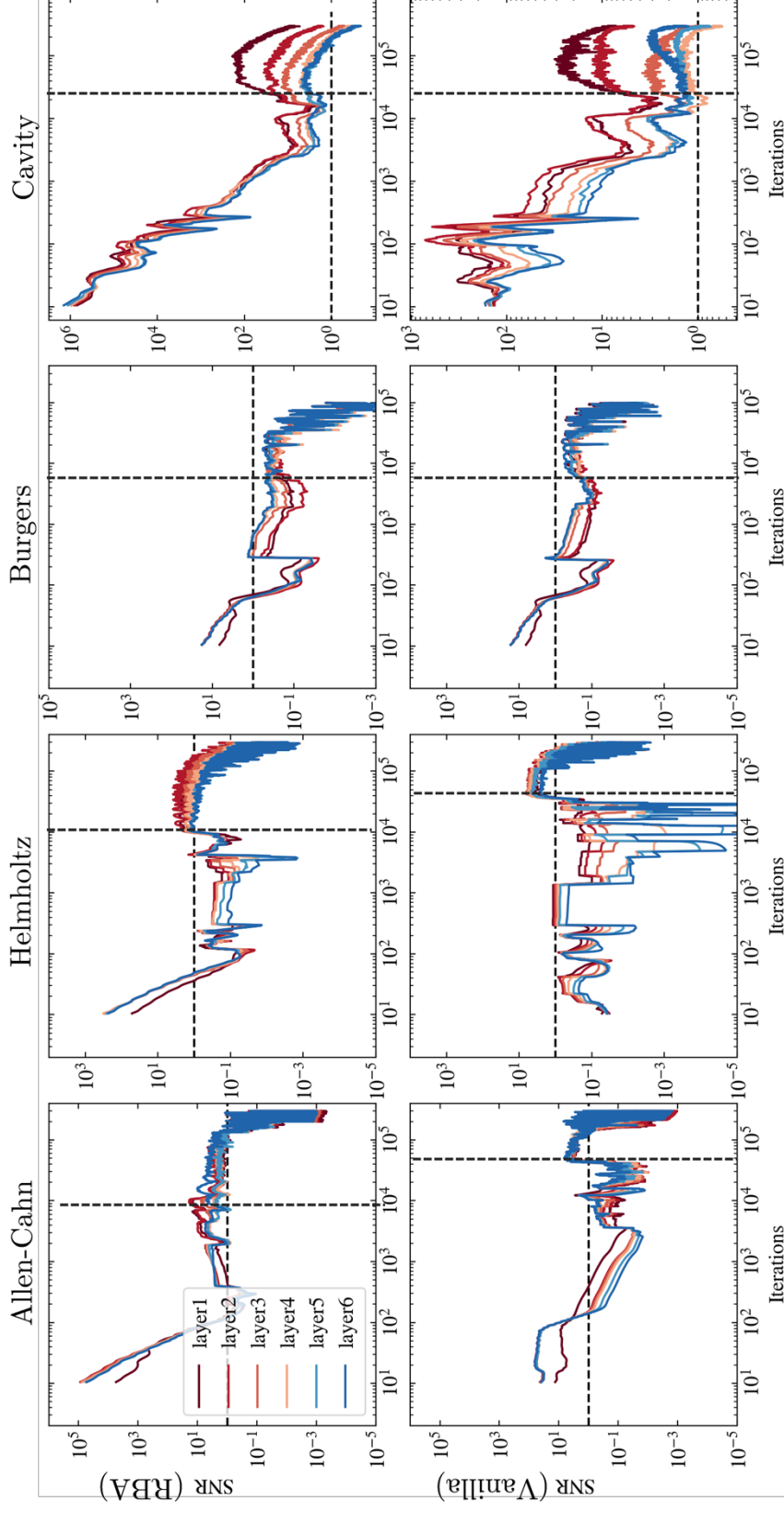
$$\|\mu_l\| = \frac{\|\langle \nabla_w \mathcal{L}_i \rangle\|_F}{\|w_i\|_F}$$

$$\|\sigma_l\| = \frac{\|std(\nabla_w \mathcal{L}_i)\|_F}{\|w_i\|_F}$$

$$SNR = \frac{\|\mu_l\|}{\|\sigma_l\|},$$

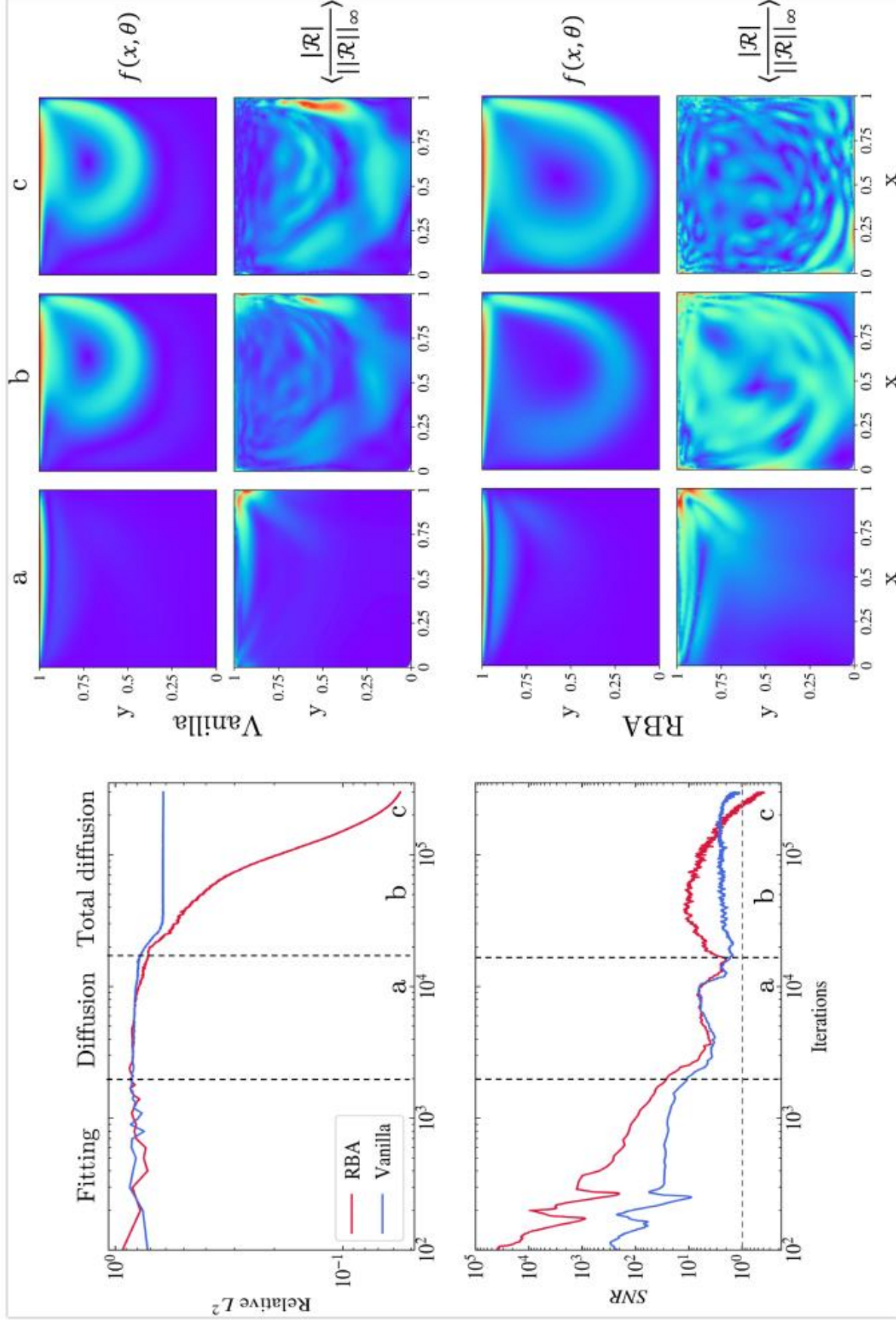


IB Theory: Learning Regimes



Layer-wise SNR -tanh activation function. (Anagnostopoulos et al., 2024)

IB Theory: Cavity Flow, Re=1000



Weighted Residual Methods

Problem Definition:

$$\begin{array}{lll} \mathcal{L}^q u(\mathbf{x}, t) = f(\mathbf{x}, t), & (\mathbf{x}, t) \in \Omega \times (0, T] & \text{:PDE} \\ u(\mathbf{x}, t) = h(\mathbf{x}, t), & (\mathbf{x}, t) \in \partial\Omega \times [0, T] & \text{:B.C.} \\ u(\mathbf{x}, 0) = g(\mathbf{x}), & \mathbf{x} \in \Omega & \text{:I.C.} \end{array}$$

$\Omega \subset \mathbb{R}^d$ – Bounded domain with boundaries $\partial\Omega$

$\mathcal{L}^q$ – Differential/Integral/Identity operator with parameters q

$u(\mathbf{x}, t)$ – Solution to be obtained

f, h, g – Known source, boundary and initial conditions

Given the PDE and boundary/initial conditions, what is the solution $u(\mathbf{x}, t)$?

Ns Family: Weighted Residual Methods

pproximation: $u \approx \tilde{u} = u_{NN}$

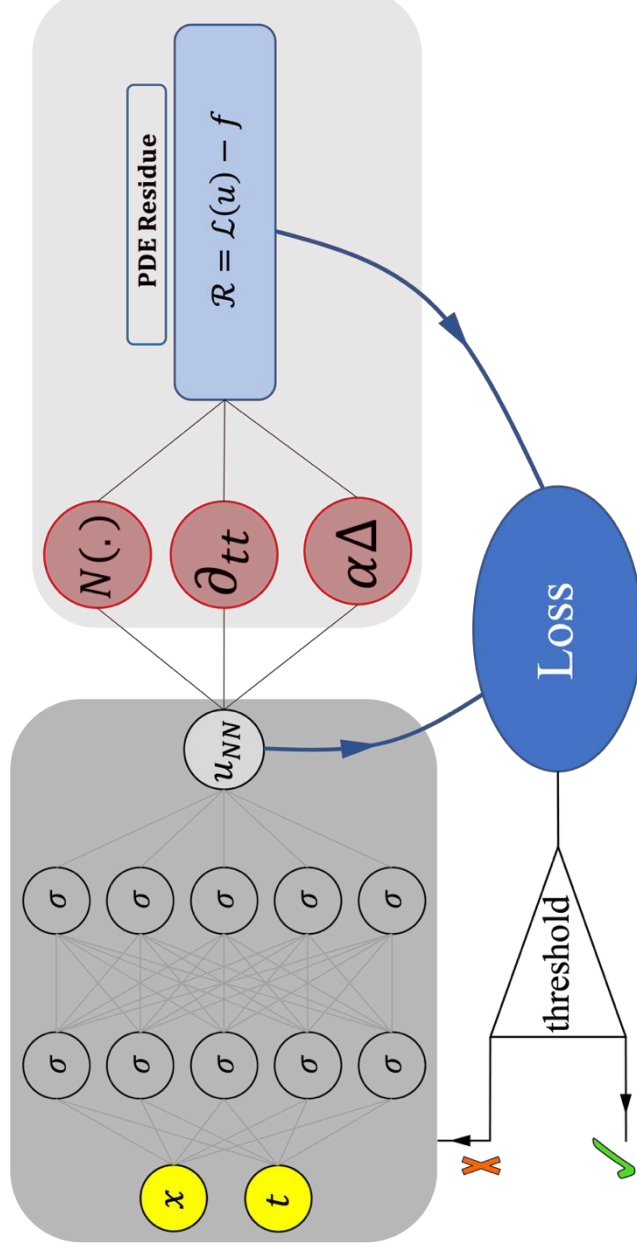
Strong-Form Residuals:

- $\mathcal{L}^q(\tilde{u}) = \mathcal{L}^q \tilde{u} - f$: PDE residual
- $\tilde{u} - h$: boundary residual
- $\tilde{u} - g$: initial residual

Loss Function:

$$L^s(\tilde{u}) = \frac{1}{N_r} \sum_{i=1}^{N_r} \underbrace{|r(\mathbf{x}_r^i, t_r^i)|^2}_{\text{physics}} + \tau_b \frac{1}{N_b} \sum_{i=1}^{N_b} \underbrace{|r_b(\mathbf{x}_b^i, t_b^i)|^2}_{\text{physics}} + \tau_0 \frac{1}{N_0} \sum_{i=1}^{N_0} \underbrace{|r_0(\mathbf{x}_0^i)|^2}_{\text{physics}}$$

minimization problem: $\min L^s(\tilde{u})$



Ns Family: Weighted Residual Methods

Weighted/Variational Residuals:
projections onto space of test function V

$$\mathcal{R}_j(\tilde{u}) = \int_{\Omega \times (0, T]} r(\tilde{u}) v_j dx dt$$

$$\mathcal{R}_{b,j}(\tilde{u}) = \int_{\partial\Omega \times (0, T]} r_b(\tilde{u}) v_j dx dt$$

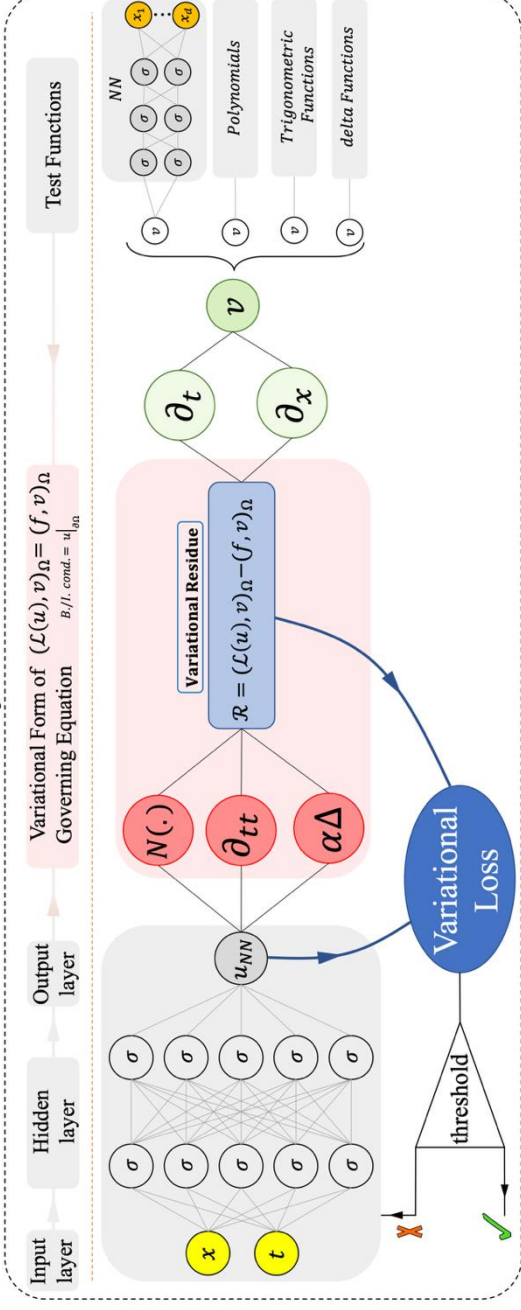
$$\mathcal{R}_{0,j}(\tilde{u}) = \int_{\Omega} r_0(\tilde{u}) v_j dx$$

Loss Function:

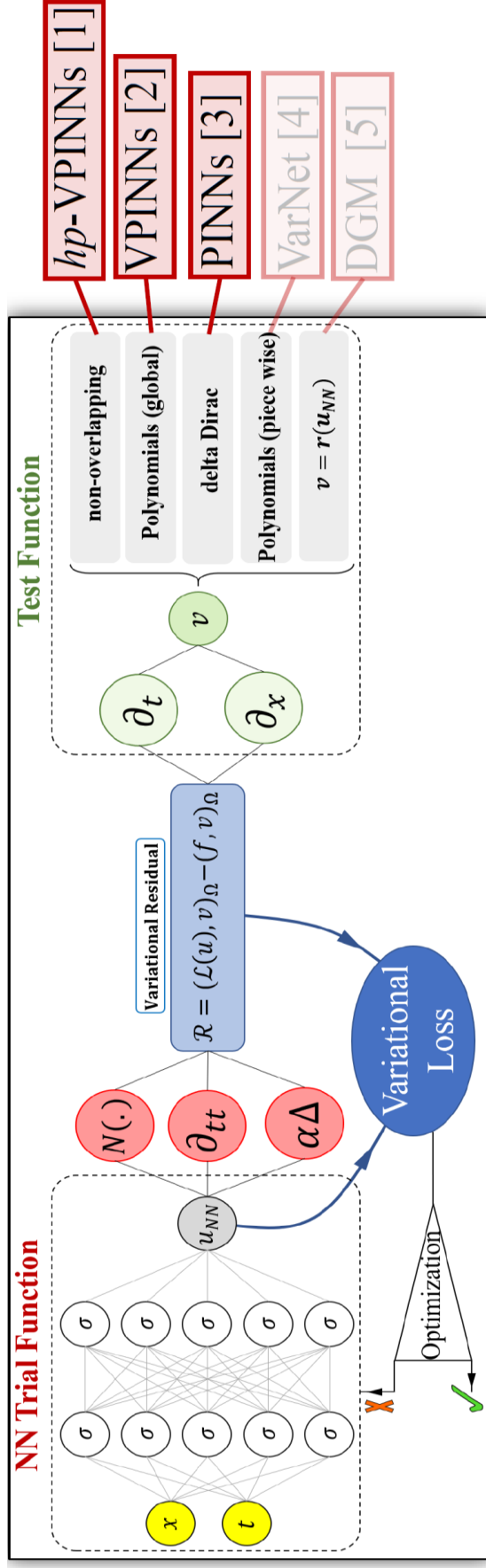
$$L^v(\tilde{u}, v) = w \sum_{j=1}^{N_r} \mathcal{R}_j^2(\tilde{u}) + w_b \sum_{j=1}^{N_b} \mathcal{R}_{b,j}^2(\tilde{u}) + w_0 \sum_{j=1}^{N_0} \mathcal{R}_{0,j}^2(\tilde{u})$$

minimization problem: $\min L^v(\tilde{u}, v)$

VPINN: Variational Physics – Informed Neural Network



VPINNs: A General Framework for Solving PDEs



	Local NN	Global NN
Local v	Conservative VPINNs [6]	hp -VPINNs [5]
Global v	-	VPINNs[3], VarNet[4], D3M [7]

Global:= single domain

Local:= domain decomposition

[1] Kharazmi et al., CMAME (2020).

[2] Kharazmi et al., arXiv (2019).

[3] Raissi et al., JCP (2019)

[4] Khodayi-Mehr et al., arXiv (2019).

[5] Sirignano et al., JCP (2018)

[6] Jagtap et al., CMAME (2020)

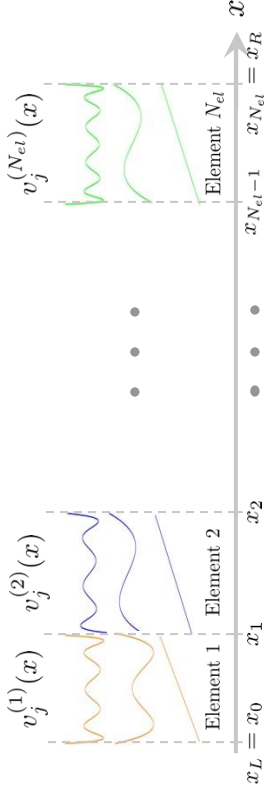
[7] Li et al., arXiv (2019)

VPINNs: A General Framework for Solving PDEs

$$\Omega = \bigcup_e \Omega^{(e)}$$

$$v_j^{(e)} = \begin{cases} \bar{v}_j \neq 0, & x \in \Omega^{(e)} \\ 0, & \text{otherwise} \end{cases}$$

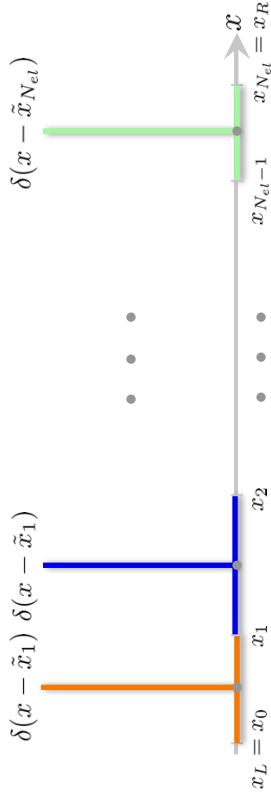
h-refinement & *p*-refinement



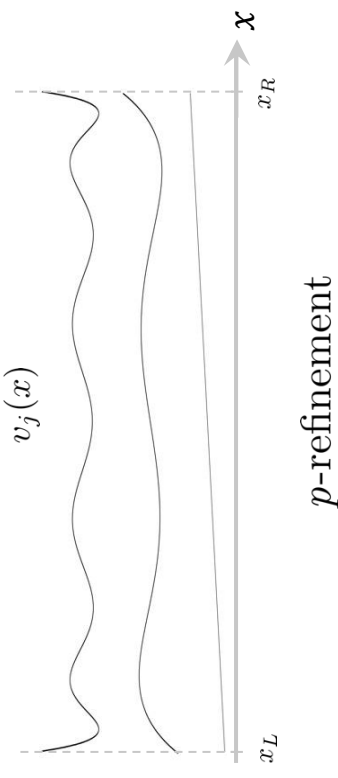
PINNs

hp-VPINNs

VPINNs



Delta Dirac test functions



p-refinement

- nonlinear approximation via neural networks
- *h*-refinement via domain decomposition
- *p*-refinement via projection onto high-order polynomial space

VPINNs: A General Framework for Solving PDEs

Loss Function

elemental/local variational loss

$$\text{VPINNs: } L^v = \boxed{\sum_{e=1}^{N_{el}} \frac{1}{N_{el} K^{(e)}} \sum_{k=1}^{K^{(e)}} |\mathcal{R}_k^{(e)}|^2} + \tau_b \frac{1}{N_b} \sum_{i=1}^{N_b} \underbrace{|r_b(\mathbf{x}_b^i, t_b^i)|^2}_{\text{physics data}} + \tau_0 \frac{1}{N_0} \sum_{i=1}^{N_0} \underbrace{|r_0(\mathbf{x}_0^i)|^2}_{\text{data}}$$

global variational loss

$$\text{VPINNs: } L^v = \boxed{\frac{1}{K} \sum_{j=1}^K |\mathcal{R}_j|^2} + \tau_b \frac{1}{N_b} \sum_{i=1}^{N_b} \underbrace{|r_b(\mathbf{x}_b^i, t_b^i)|^2}_{\text{physics data}} + \tau_0 \frac{1}{N_0} \sum_{i=1}^{N_0} \underbrace{|r_0(\mathbf{x}_0^i)|^2}_{\text{data}}$$

strong-form loss

$$\text{PINNs: } L^s = \boxed{\frac{1}{N_r} \sum_{i=1}^{N_r} \underbrace{|r(\mathbf{x}_r^i, t_r^i)|^2}_{\text{physics data}}} + \tau_b \frac{1}{N_b} \sum_{i=1}^{N_b} \underbrace{|r_b(\mathbf{x}_b^i, t_b^i)|^2}_{\text{physics data}} + \tau_0 \frac{1}{N_0} \sum_{i=1}^{N_0} \underbrace{|r_0(\mathbf{x}_0^i)|^2}_{\text{data}}$$

ational Neural Networks (VNNs): Function Approximation

Variational residual:

$$\mathcal{R}_k^{(e)} = \int_{\Omega_e} (u_{NN}(\mathbf{x}) - u(\mathbf{x})) v_k^{(e)}(\mathbf{x}) d\Omega_e$$

Variational loss:

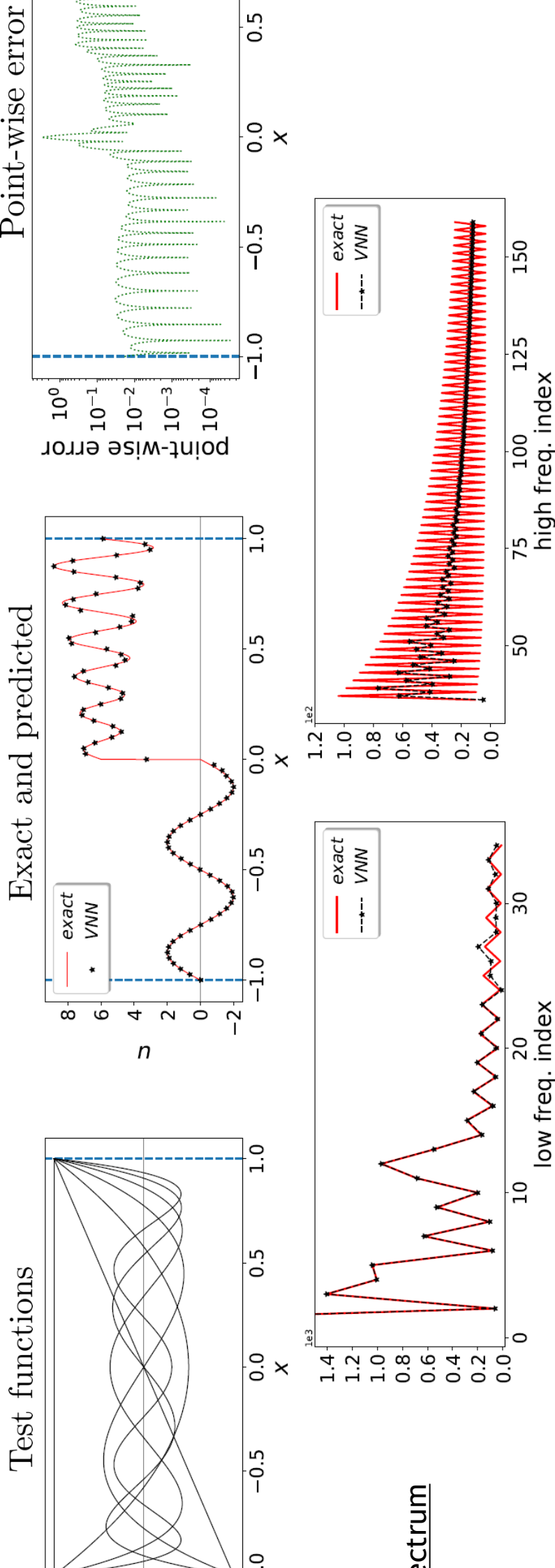
$$L^{\mathbf{v}} = \sum_{e=1}^{N_{el}} \frac{1}{K^{(e)}} \sum_{k=1}^{K^{(e)}} \left| \mathcal{R}_k^{(e)} \right|^2$$

Example (discontinuous function approximation): the target function is

$$u^{exact} = \begin{cases} 2 \sin(4\pi x) & x \in [-1, 0], \\ 6 + \exp^{1.2x} \sin(12\pi x) & x \in (0, 1] \end{cases}$$

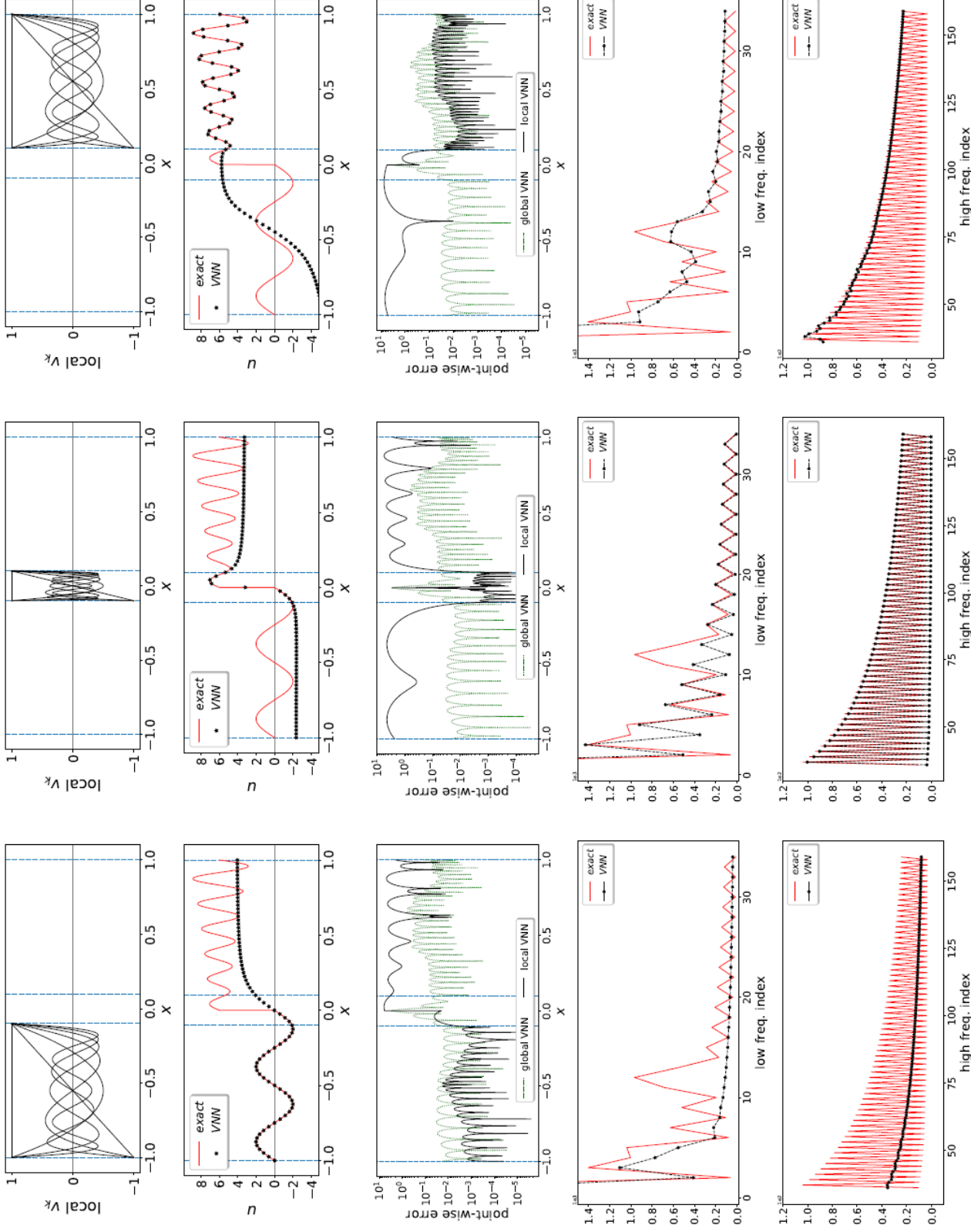
Global test functions (single element): Function Approximation

Global test functions (single element):



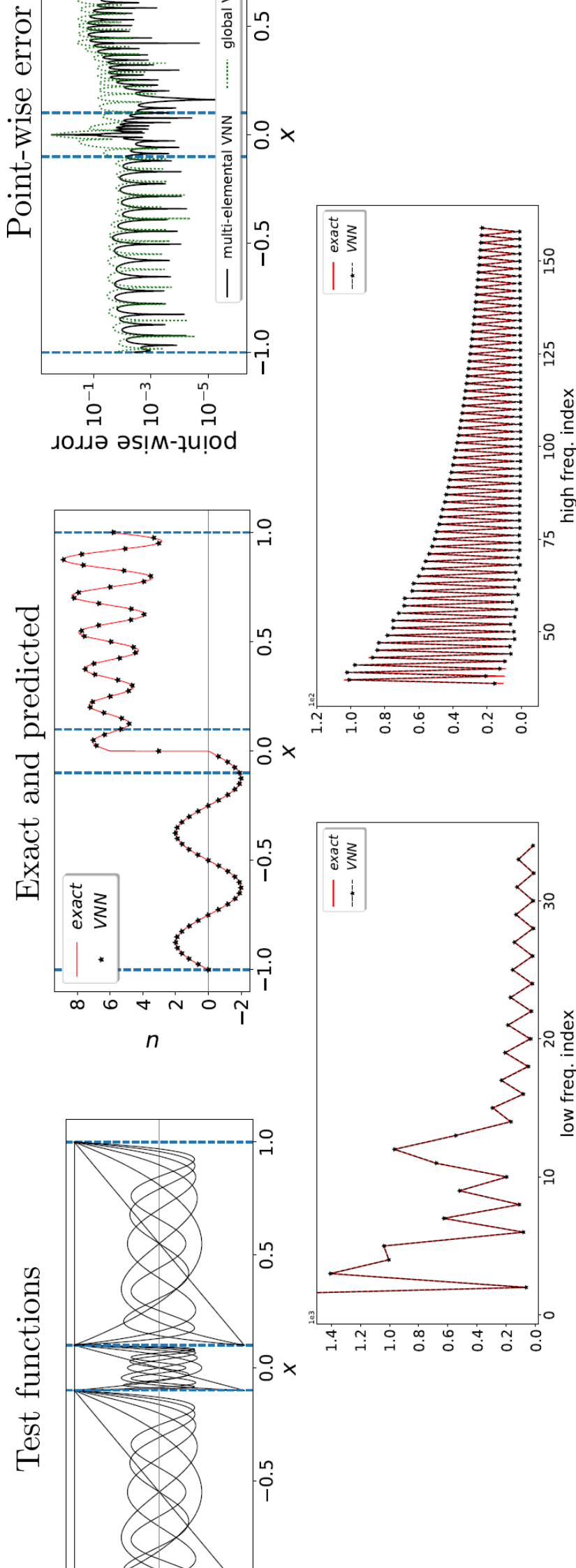
Functional Neural Networks (VNNs): Function Approximation

local (elemental)
st functions:



Multi-elemental Neural Networks (VNNs): Function Approximation

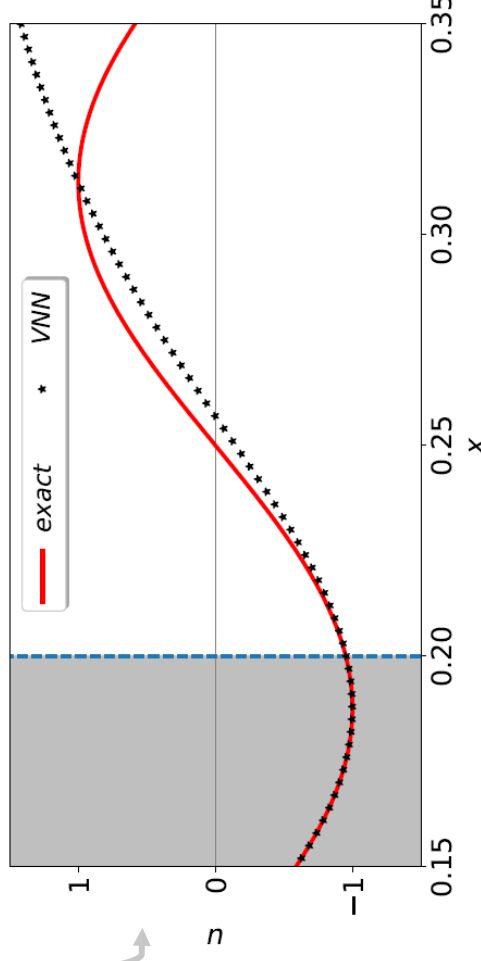
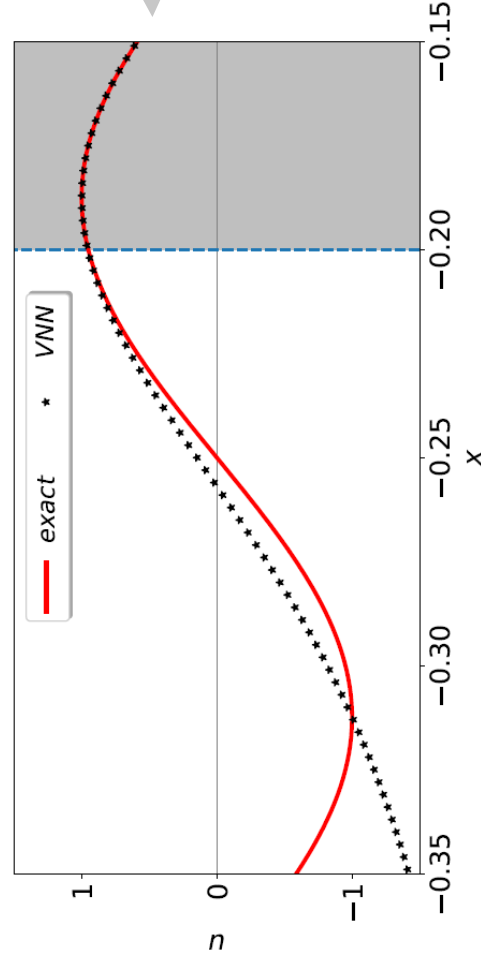
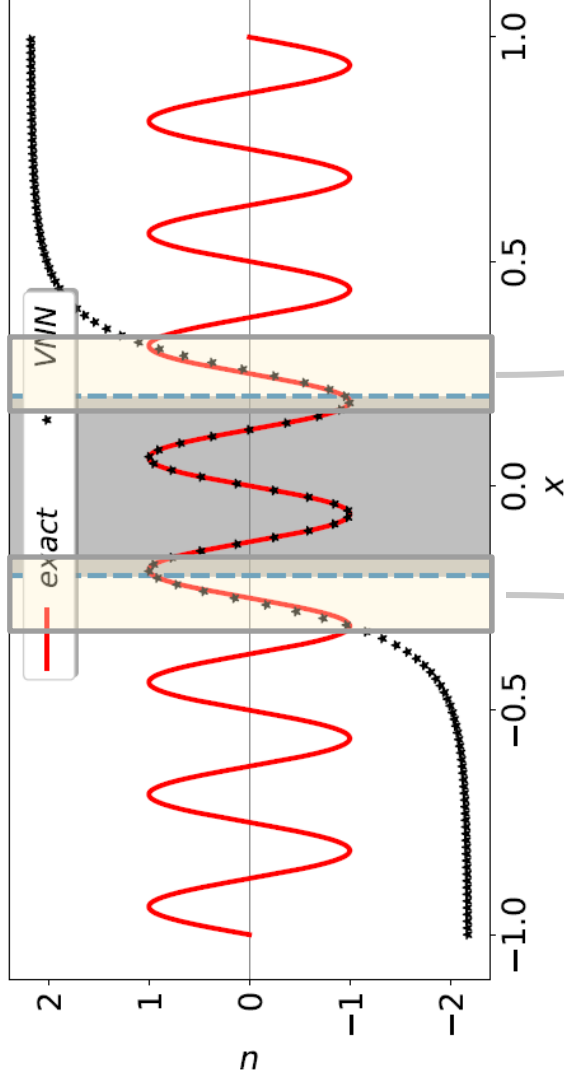
Multi-elemental test functions:



Functional Neural Networks (VNNs): Function Approximation

Localized learning & learning out-of-the-box

- The test functions have **non-zero** values only in the gray box ($-0.2 < x < 0.2$).
- The test functions have **zero** values in the white region.
- It is expected that VNN only learn the function in the gray box.
- But VNN learns the function slightly outside of the box, a notion of generalization to unseen domain is indicated as the yellow region



allow VPINNs: Analytical Approach

Variational formulation requires computing numerical integration.

Q: Can we avoid error due to numerical integration?

Starting with a shallow network to compute the integrals analytically.

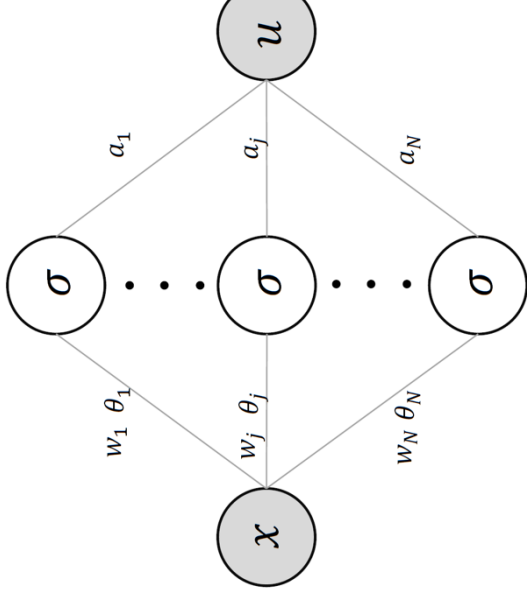
$$u_{NN}(x) = \sum_{j=1}^N u_{NN_j}(x) = \sum_{j=1}^N a_j \sigma(w_j x + \theta_j)$$

Steady-state Burger's equation:

$$\begin{cases} u \frac{du}{dx} - \frac{d^2 u}{dx^2} = f(x) \\ u(-1) = g, \quad u(1) = h \end{cases}$$

$$Residual_k^v = \mathcal{R}_k + \mathcal{R}_k^{NL} - F_k - r^b$$

$$\mathcal{R}_k = - \left(\frac{d^2 u_{NN}(x)}{dx^2}, v_k(x) \right)_{\Omega} \quad \mathcal{R}_k^{NL} = \left(u_{NN} \frac{du_{NN}(x)}{dx}, v_k(x) \right)_{\Omega} \quad F_k = (f(x), v_k(x))_{\Omega}$$



allow VPINNs: Analytical Approach

- Consider **Sine** activation functions with **Sine** test function $v_k(x) = \sin(k\pi x)$, $k = 1, 2, \dots, K$
- Analytical expression of **variational** residuals

$$\mathcal{R}_k^{(1)} = \mathcal{R}_k^{(2)} = 2(-1)^k k\pi \sum_{j=1}^N \frac{a_j w_j^2 \cos(\theta_j) \sin(w_j)}{w_j^2 - k^2 \pi^2},$$

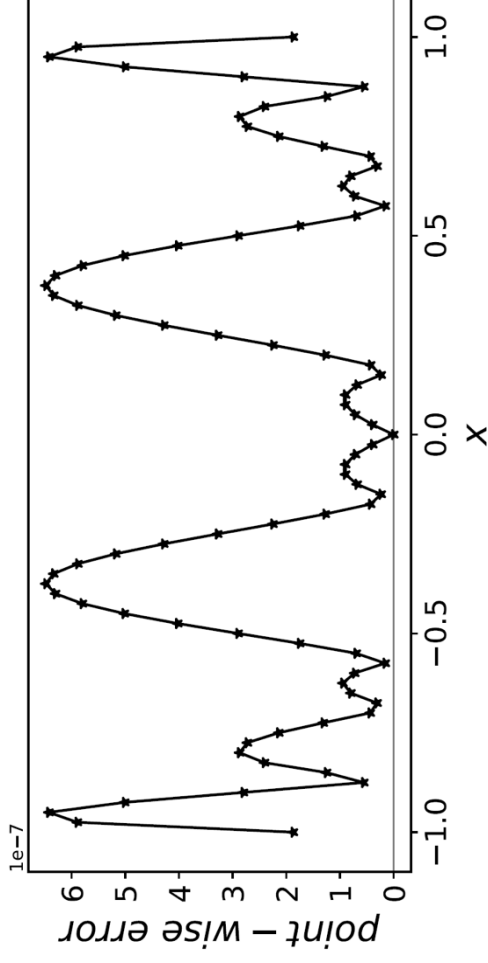
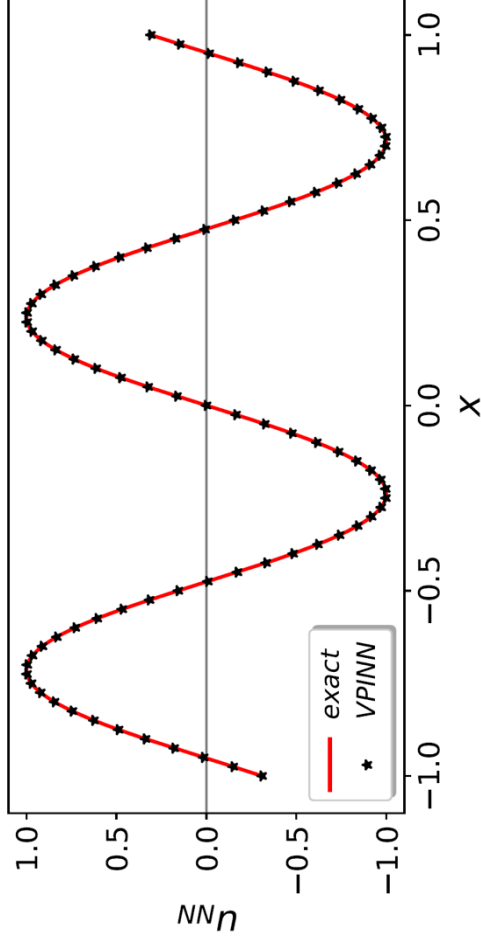
$$\mathcal{R}_k^{(3)} = 2(-1)^k k\pi \sum_{j=1}^N \frac{a_j k^2 \pi^2 \cos(\theta_j) \sin(w_j)}{w_j^2 - k^2 \pi^2} + (-1)^k k\pi (h - g),$$

$$\mathcal{R}_k^{NL} = (-1)^k k\pi \sum_{i=1}^N \sum_{j=1}^N a_i a_j w_i \left[\frac{\sin(w_j + w_i) \cos(\theta_j + \theta_i)}{(w_j + w_i)^2 - k^2 \pi^2} + \frac{\sin(w_j - w_i) \cos(\theta_j - \theta_i)}{(w_j - w_i)^2 - k^2 \pi^2} \right]$$

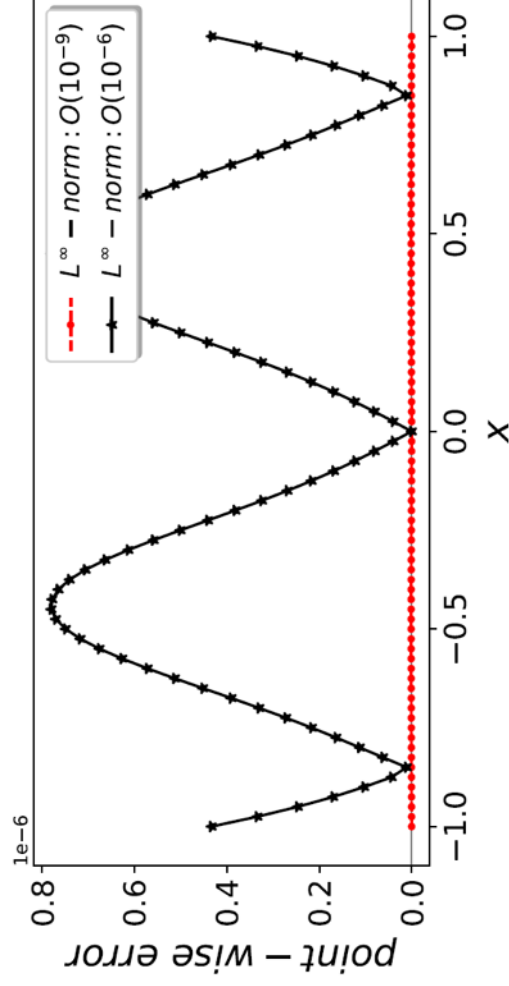
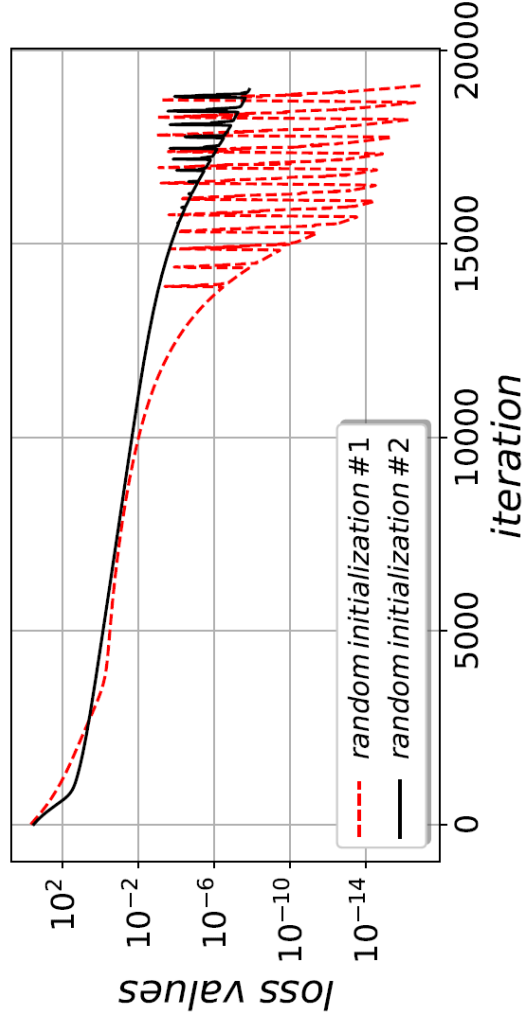
- We can also use **sine** activation functions with **Legendre polynomial** as test functions $v_k(x) = P_{k+1}(x) - P_{k-1}(x)$, $k = 1, 2, \dots, K$.
- We then use the **recursion** formula to obtain the **variational** residuals.

allow VPINNs: Analytical Approach

Example: Steady-state Burger's equation with smooth solution. $u^{exact} = A \sin(\omega x)$



work initializations matters:



VPINNs: Poisson's Equation (1D)

- Governing Equation:

$$-\frac{d^2 u(x)}{dx^2} = f(x)$$

- Boundary conditions:

$$u(-1) = g$$

$$u(1) = h$$

- Variational residual:

$$\mathcal{R}_k = \sum_{e=1}^{N_{el}} \mathcal{R}_k^{(e)} = \sum_{e=1}^{N_{el}} \int_{x_{e-1}}^{x_e} \left(-\frac{d^2 u_{NN}(x)}{dx^2} - f(x) \right) v_k^{(e)}(x) dx,$$

VPINNs: Poisson's Equation (1D)

- Reducing the regularity and simplifying the network by integration-by-parts:

$$^{(1)}\mathcal{R}_k^{(e)} = - \int_{x_{e-1}}^{x_e} \frac{d^2 u_{NN}(x)}{dx^2} v_k^{(e)}(x) dx - F_k^{(e)},$$

$$^{(2)}\mathcal{R}_k^{(e)} = \int_{x_{e-1}}^{x_e} \frac{du_{NN}(x)}{dx} \frac{dv_k^{(e)}(x)}{dx} dx - \frac{du_{NN}(x)}{dx} v_k^{(e)}(x) \Big|_{x_{e-1}}^{x_e} - F_k^{(e)},$$

$$^{(3)}\mathcal{R}_k^{(e)} = - \int_{x_{e-1}}^{x_e} u_{NN}(x) \frac{d^2 v_k^{(e)}(x)}{dx^2} dx - \frac{du_{NN}(x)}{dx} v_k^{(e)}(x) \Big|_{x_{e-1}}^{x_e} + u_{NN}(x) \frac{dv_k^{(e)}(x)}{dx} \Big|_{x_{e-1}}^{x_e} - F_k^{(e)}$$

- Variational loss:

$$L^{\mathbf{v}(i)} = \sum_{e=1}^{N_{el}} \frac{1}{K^{(e)}} \sum_{k=1}^{K^{(e)}} \left| {}^{(i)}\mathcal{R}_k^{(e)} \right|^2 + \frac{\tau_b}{2} \left(|u_{NN}(-1) - g|^2 + |u_{NN}(1) - h|^2 \right), \quad i = 1, 2, 3$$