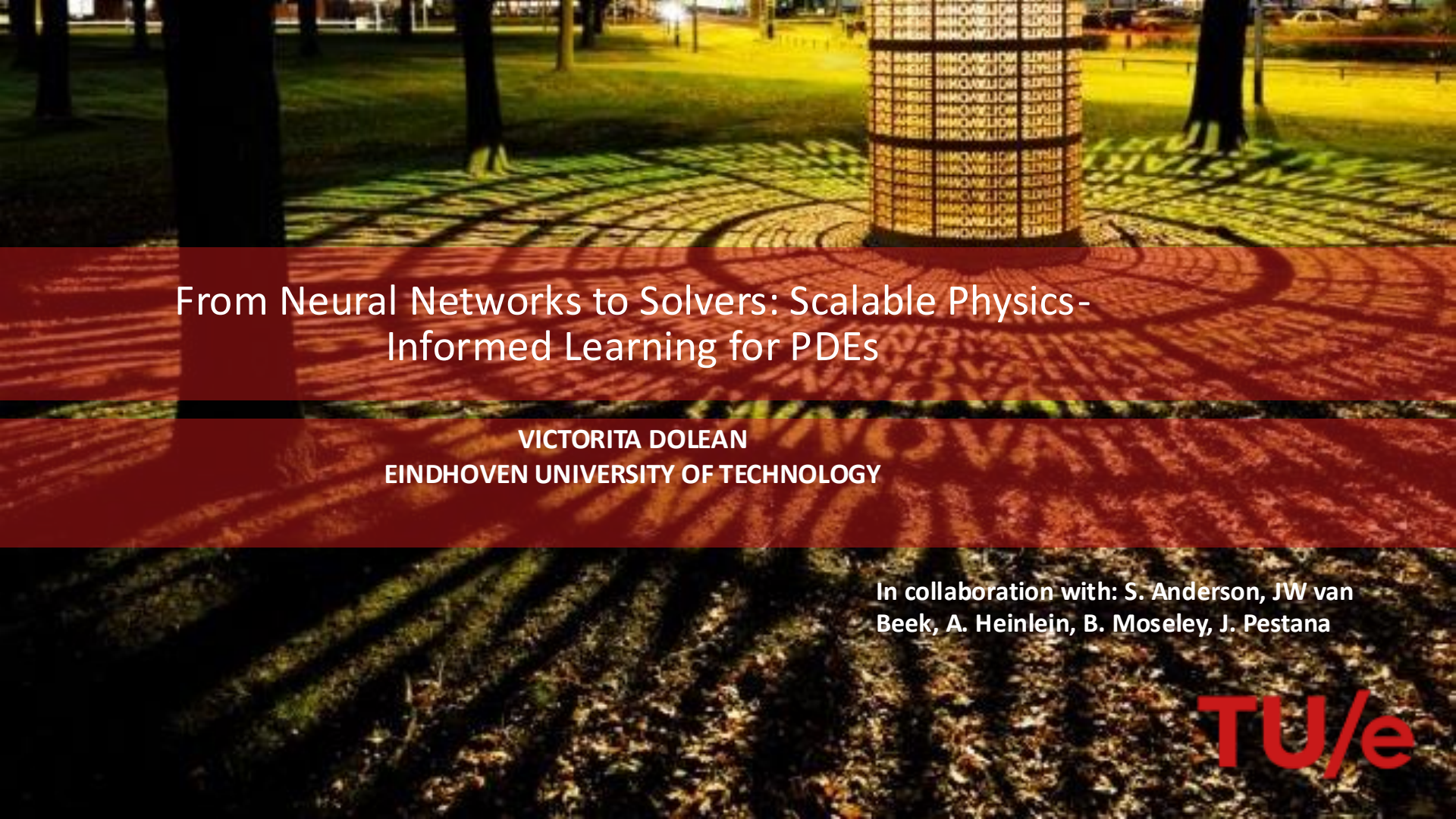


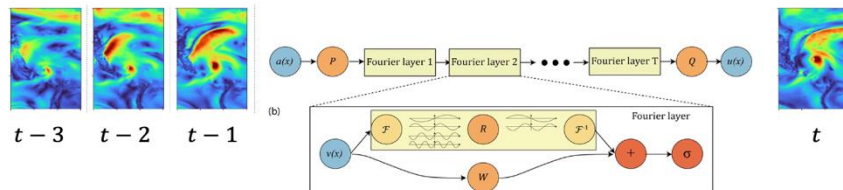
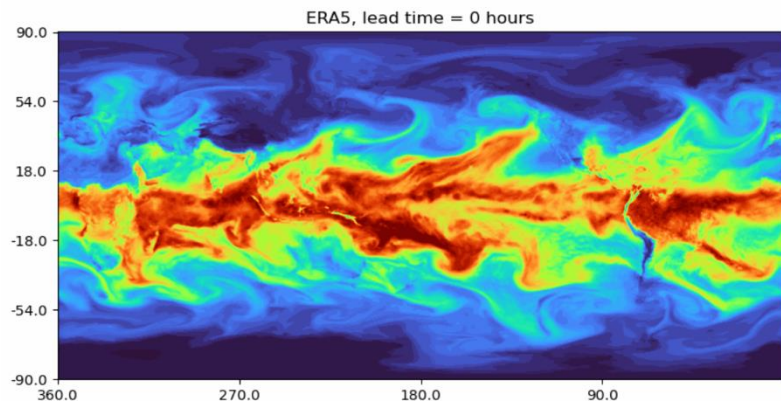
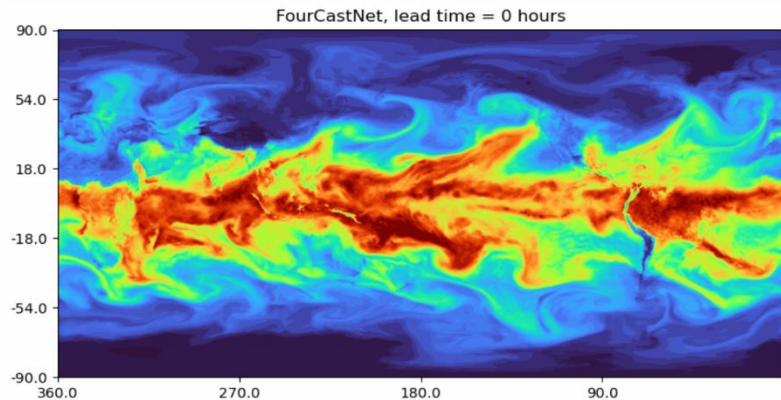


# From Neural Networks to Solvers: Scalable Physics-Informed Learning for PDEs

VICTORITA DOLEAN  
EINDHOVEN UNIVERSITY OF TECHNOLOGY

In collaboration with: S. Anderson, JW van Beek, A. Heinlein, B. Moseley, J. Pestana

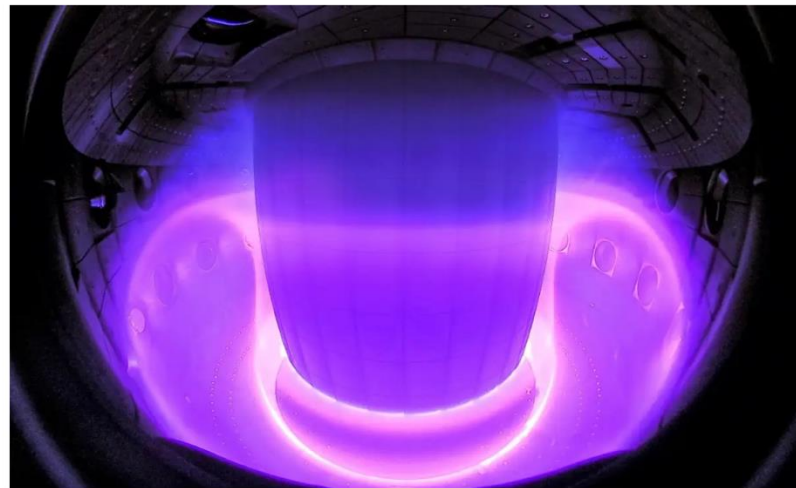
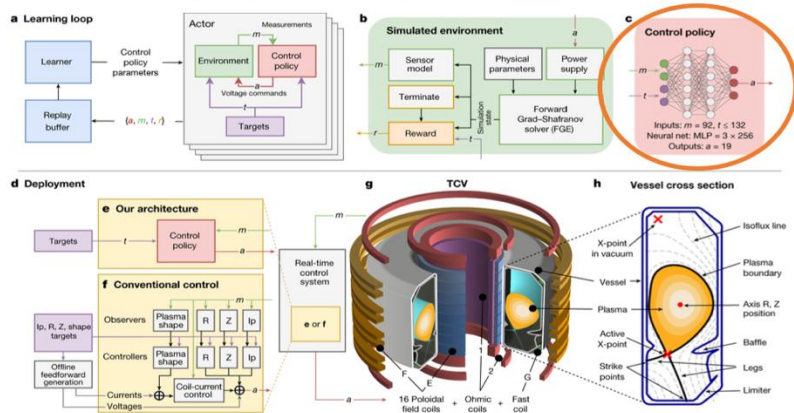
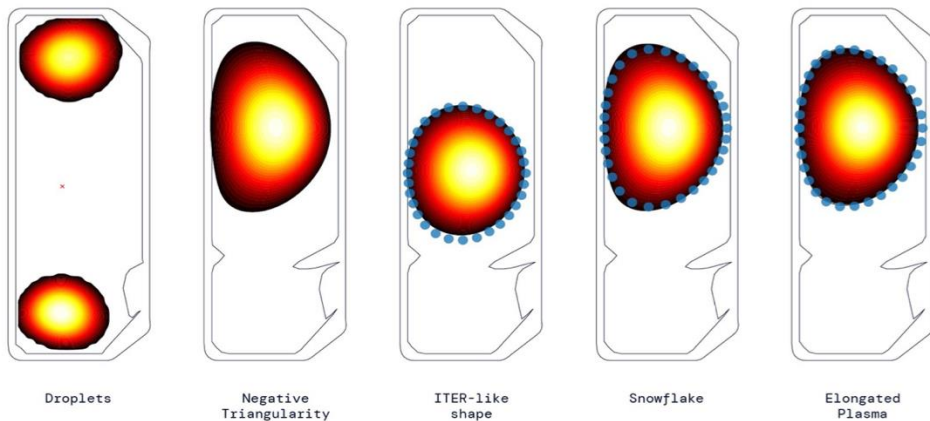
# AI for science: a revolution?



Pathak et al, FourCastNet: A Global Data-driven High-resolution Weather Model using Adaptive Fourier Neural Operators, ArXiv (2022)



# AI for science: a revolution?



Variable Configuration Tokamak (TCV) in Lausanne, Switzerland  
Source: DeepMind & SPC/EPFL

nature

Explore content ▾ About the journal ▾ Publish with us ▾

nature > articles > article

Article | [Open access](#) | [Published: 16 February 2022](#)

## Magnetic control of tokamak plasmas through deep reinforcement learning

Jonas Degraeve, Federico Felici , Jonas Buchli , Michael Neunert, Brendan Tracey , Francesco Carpanese, Timo Ewalds, Roland Hafner, Abbas Abdolmaleki, Diego de las Casas, Craig Donner, Leslie Fritz, Cristian Galperti, Andrea Huber, James Keeling, Maria Tsimpoukelli, Jackie Kay, Antoine Merle, Jean-Marc Moret, Seb Noury, Federico Pesamosca, David Pfau, Olivier Sauter, Cristian Sommariva, ... Martin Riedmiller  + Show authors

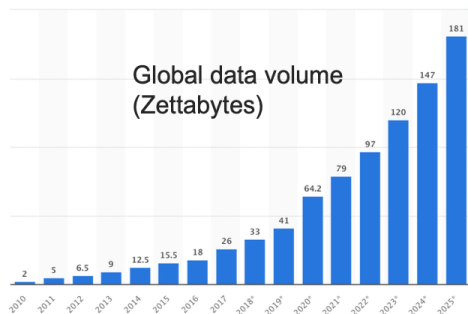
[Nature](#) 602, 414–419 (2022) | [Cite this article](#)

206k Accesses | 223 Citations | 2430 Altmetric | [Metrics](#)

# Why now?

Neural networks date back to the 1950's – so why is deep learning so popular today?

Rapidly increasing  
amounts of data



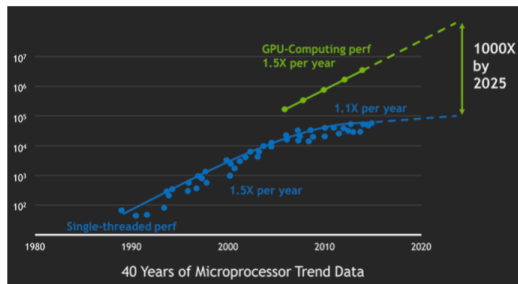
Source: Statista

IMAGENET



WIKIPEDIA  
The Free Encyclopedia

Hardware  
improvements



Source: NVIDIA

- Graphical processing units (GPUs)
- Highly optimised for deep learning (massively parallel)

Software  
improvements

PyTorch



TensorFlow

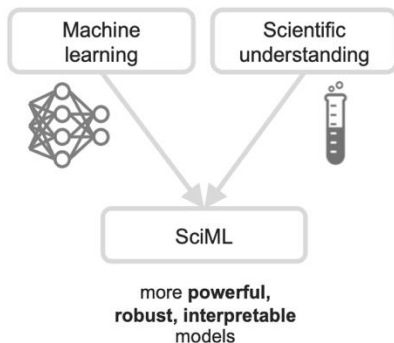
K Keras



- Mature deep learning frameworks
- Better training algorithms
- Deeper and more sophisticated architectures

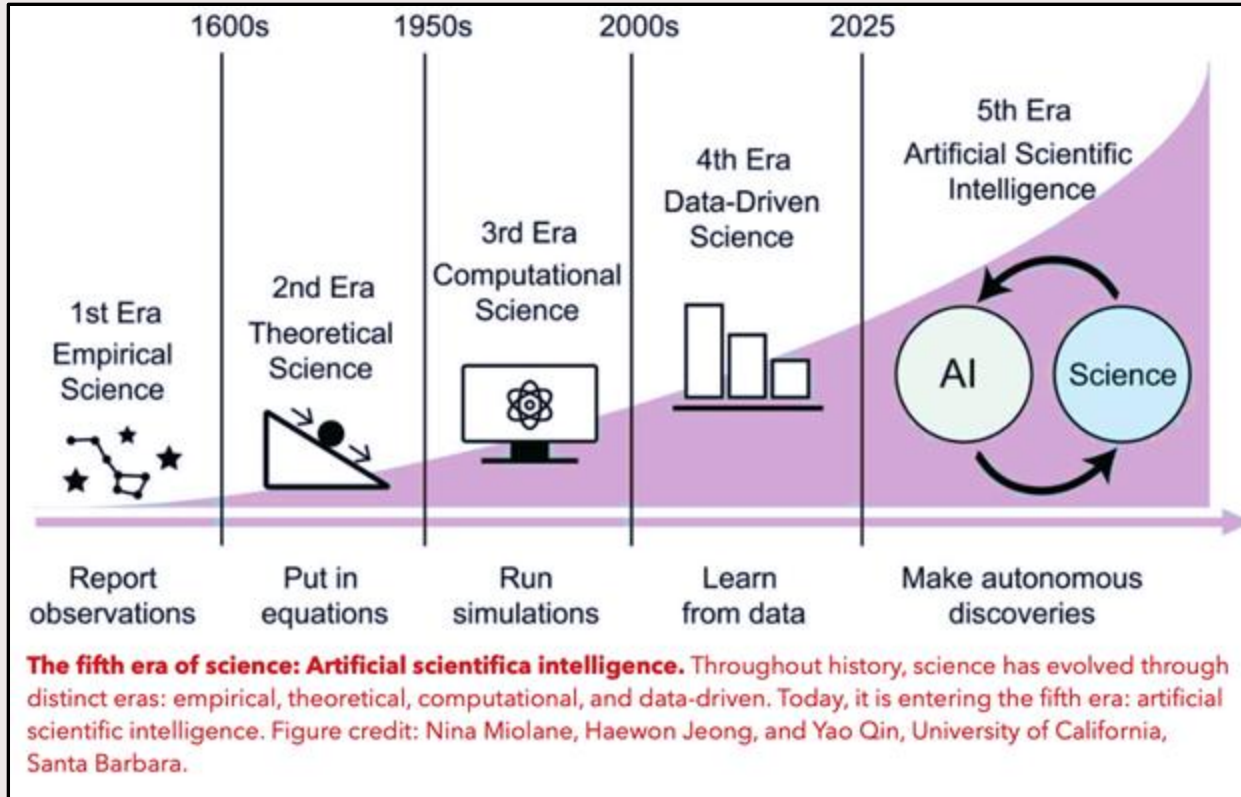
# Scientific machine learning

Hamiltonian neural networks  
Learned sub-grid processes  
Hidden physics models  
DeepONets  
PDE-Net  
Algorithm unrolling  
Differentiable simulation  
Fourier neural operators  
Encoding conservation laws  
Physics-informed neural networks  
Physics-constrained Gaussian processes  
AI Feynman  
AlphaFold  
Learned regularisation  
Physics-informed neural operators  
Neural ODEs  
Solver-in-the-loop



# Scientific Machine Learning

(part of Artificial Scientific Intelligence: AI4Science – Science4AI)



# Classical Solvers vs Scientific ML

## Model-first (Classical solvers)

- ✓ Transparent: you can see why it works
- ☀ Physics built-in from the start
- 📄 Slower per run, but predictable
- 📄 No training data needed
- 🛡 High trust: proofs & error estimates

## Data+Model (Scientific ML)

- 🔍 Often a black box – needs explainers
- ☀ Physics can be added (PINNs, operators)
- ⚡ Very fast after training
- 📖 Learns from data (needs examples)
- 🧭 Trust is active research (UQ, robustness)

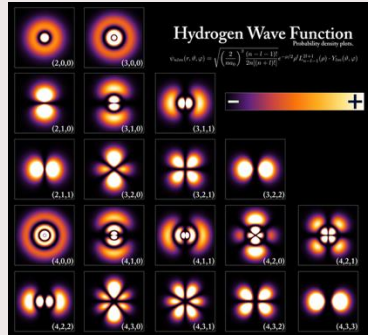
More physics



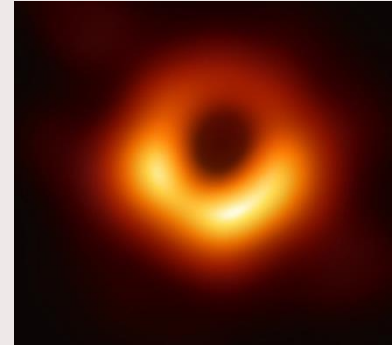
More data



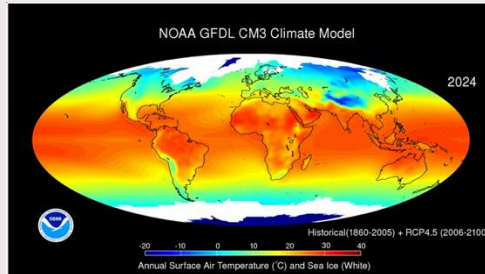
# Importance of PDEs



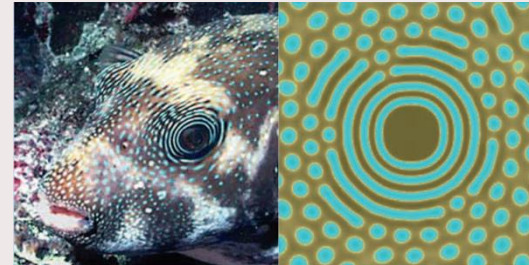
Source: Wikipedia  
Schrödinger equation



Source: The Event Horizon Telescope (2019)  
Einstein field equations



Source: NOAA  
Navier-Stokes equations



Source: Kondo and Miura, Science (2010)  
Reaction-diffusion equation



# A rapidly growing field

ICLR 2024 Workshop on  
AI4DifferentialEquations in Science

Machine Learning and the Physical Sciences  
Workshop at the 37th conference on Neural Information Processing Systems (NeurIPS)  
December 15, 2023

Synergy of Scientific and Machine  
Learning Modeling

ICML 2023 Workshop, July 28 2023, Room 320 of the Hawaii Convention Center



AI for Science

NeurIPS 2021

ICML 2022

NeurIPS 2022

NeurIPS 2023

ICLR 2023 Workshop on Physics for Machine Learning

Physics4ML

The Symbiosis of Deep Learning and Differential Equations (DLDE)

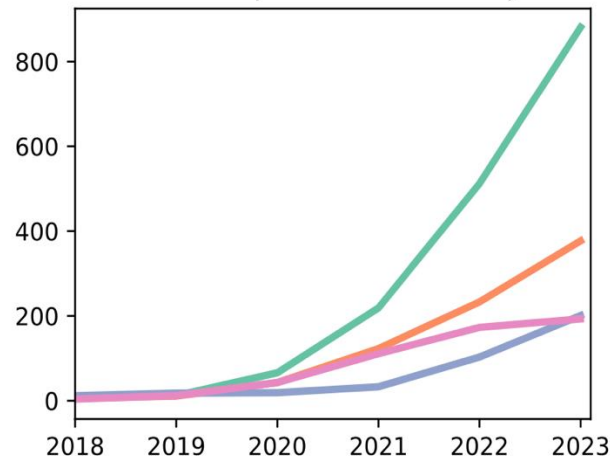
NeurIPS 2022 Workshop

AAAI-MLPS 2021

AAAI 2021 Spring Symposium on

Combining Artificial Intelligence and Machine Learning with Physics Sciences

Number of publications each year

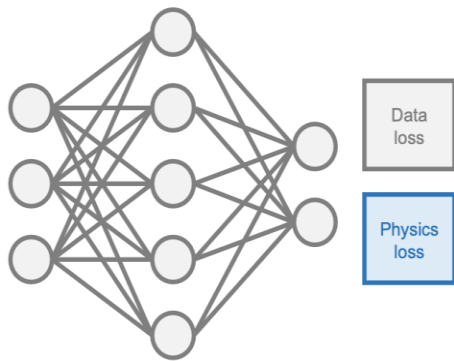


- physics-informed neural networks
- scientific machine learning / physics-informed ML / AI for science
- operator learning / neural operators
- differentiable physics / neural differential equations

Source: Scopus keyword search (Feb 2024)

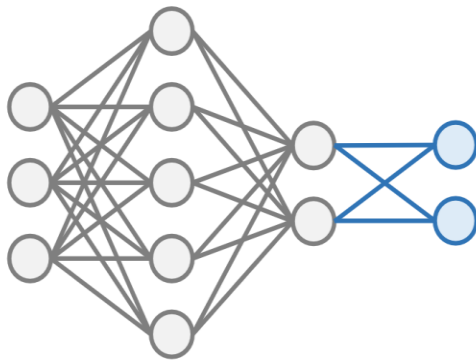
# Ways to incorporate scientific principles into machine learning

## Loss function



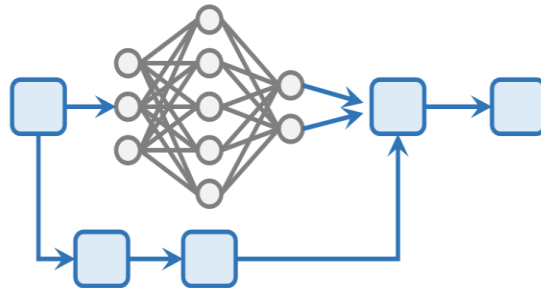
Example:  
Physics-informed neural networks  
(add governing equations to loss  
function)

## Architecture



Example:  
Encoding symmetries / conservation laws  
(e.g. energy conservation, rotational  
invariance)

## Hybrid approaches



Example:  
Neural differential equations  
(incorporating neural networks into PDE  
models)

# Can physics-informed neural networks (PINNs) beat finite difference / finite element methods?

## Can physics-informed neural networks beat the finite element method?

Tamara G Grossmann , Urszula Julia Komorowska, Jonas Latz, Carola-Bibiane Schönlieb

*IMA Journal of Applied Mathematics*, Volume 89, Issue 1, January 2024, Pages 143–174,  
<https://doi.org/10.1093/imamat/hxae011>

**Published:** 23 May 2024 **Article history** ▼

## Solving inverse problems in physics by optimizing a discrete loss: Fast and accurate learning without neural networks

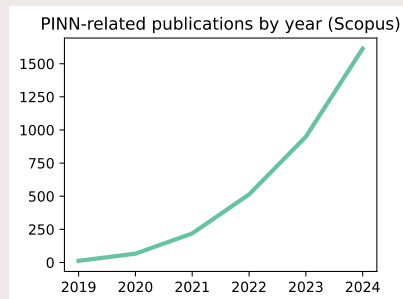
Petr Karnakov, Sergey Litvinov, Petros Koumoutsakos  **Author Notes**

*PNAS Nexus*, Volume 3, Issue 1, January 2024, pgae005,  
<https://doi.org/10.1093/pnasnexus/pgae005>

**Published:** 11 January 2024 **Article history** ▼

## 7. Discussion and conclusions

After having investigated each of the PDEs on its own, let us now discuss and draw conclusions from the results as a whole. Considering the solution time and accuracy, PINNs are not able to beat FEM in our study. In all the examples that we have studied, the FEM solutions were faster at the same or at a higher accuracy.



## Conclusion

We introduce the ODIL framework for solving inverse problems for PDEs by casting their discretization as an optimization problem and applying optimization techniques that are widely available in machine-learning software. The concept of casting the PDE as is closely related to the neural network formulations proposed by (15–17) and recently revived as PINNs. However, the fact that we use the discrete approximation of the equations allows for ODIL to be orders of magnitude more efficient in terms of computational cost and accuracy compared to the PINN for which complex flow problems “remain elusive” (71).

TITLE-ABS-KEY ( "physics-informed neural network"  
OR "physics informed neural network" )

# What is a physics-informed neural network?

Problem: damped harmonic oscillator



$$m \frac{d^2 u}{dt^2} + \mu \frac{du}{dt} + ku = 0$$

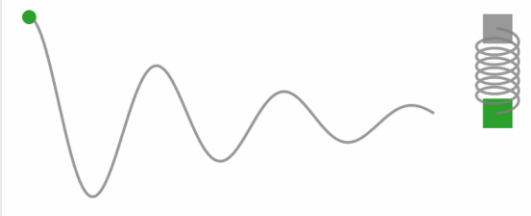
Raissi et al, Physics-informed neural networks: A deep learning framework for solving forward and inverse problems involving nonlinear partial differential equations, JCP (2018)

Lagaris et al, Artificial neural networks for solving ordinary and partial differential equations, IEEE (1998)

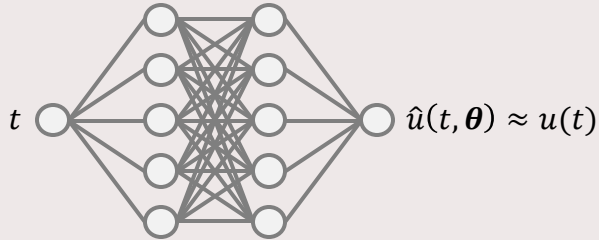


# What is a physics-informed neural network?

Problem: damped harmonic oscillator



$$m \frac{d^2 u}{dt^2} + \mu \frac{du}{dt} + ku = 0$$



Raissi et al, Physics-informed neural networks: A deep learning framework for solving forward and inverse problems involving nonlinear partial differential equations, JCP (2018)

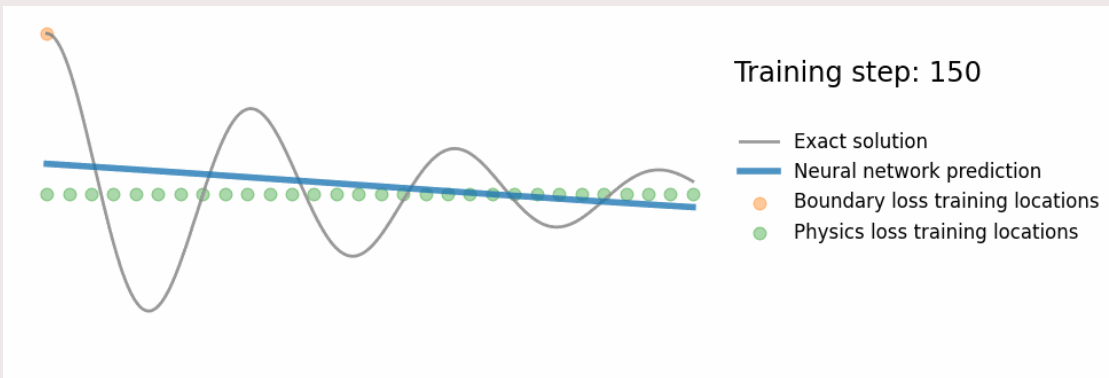
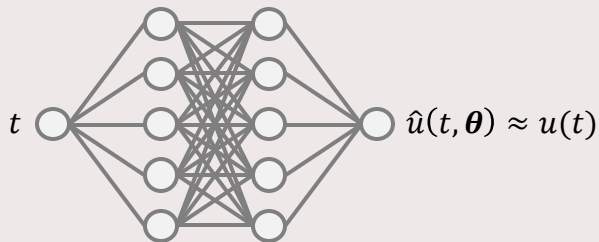
Lagaris et al, Artificial neural networks for solving ordinary and partial differential equations, IEEE (1998)

# What is a physics-informed neural network?

Problem: damped harmonic oscillator



$$m \frac{d^2 u}{dt^2} + \mu \frac{du}{dt} + ku = 0$$



$$\text{Boundary loss} \left\{ \begin{aligned} L(\theta) &= \lambda_1 (\hat{u}(t=0, \theta) - 1)^2 \\ &+ \lambda_2 \left( \frac{d\hat{u}}{dt}(t=0, \theta) - 0 \right)^2 \end{aligned} \right.$$

$$\text{Physics loss (aka PDE residual)} \left\{ \begin{aligned} &+ \frac{1}{N_p} \sum_i \left( \left[ m \frac{d^2}{dt^2} + \mu \frac{d}{dt} + k \right] \hat{u}(t_i, \theta) \right)^2 \end{aligned} \right.$$

Raissi et al, Physics-informed neural networks: A deep learning framework for solving forward and inverse problems involving nonlinear partial differential equations, JCP (2018)

Lagaris et al, Artificial neural networks for solving ordinary and partial differential equations, IEEE (1998)

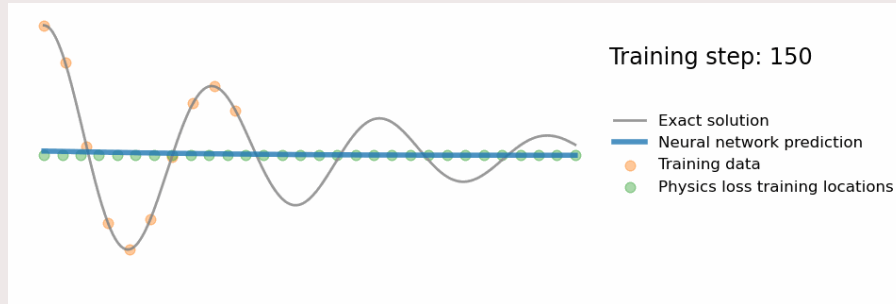
# Scalability challenges of PINNs

## Advantages of PINNs

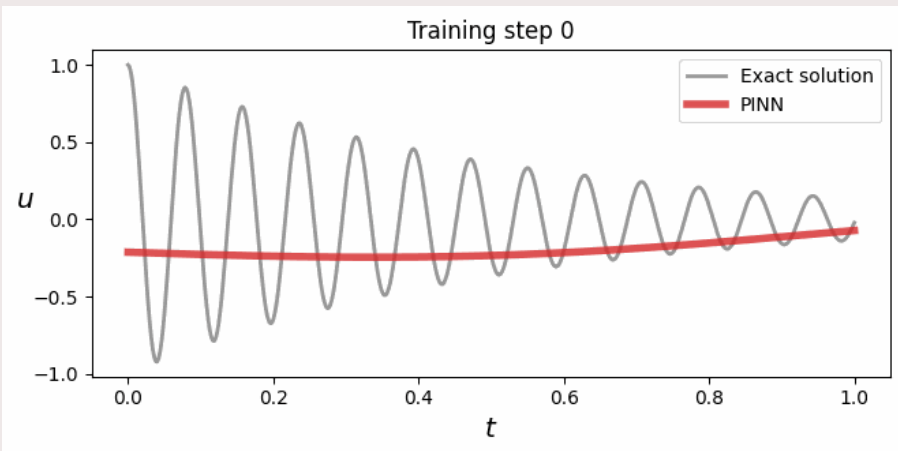
- **Mesh-free**
- Can solve forward and inverse problems, and seamlessly incorporate **observational data**
- Mostly **unsupervised**
- Can perform well for high-dimensional PDEs

## Limitations of PINNs

- **Computational cost** often high (especially for forward-only problems)
- Can be hard to **optimise**
- Challenging to **scale** to high-frequency, multi-scale problems



# Scaling PINNs to high frequency / multiscale problems

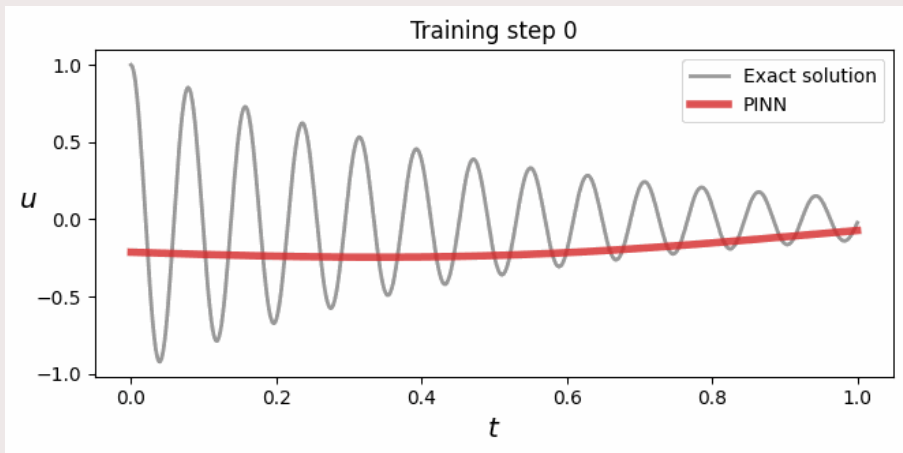


Problem: PINNs **struggle** to solve high-frequency / multiscale problems



Damped harmonic oscillator

# Scaling PINNs to high frequency / multiscale problems



Spectral bias:

NNs tend to converge much **slower** on high frequencies than on low frequencies

Rahaman, N., et al, On the spectral bias of neural networks. 36th International Conference on Machine Learning, ICML (2019)

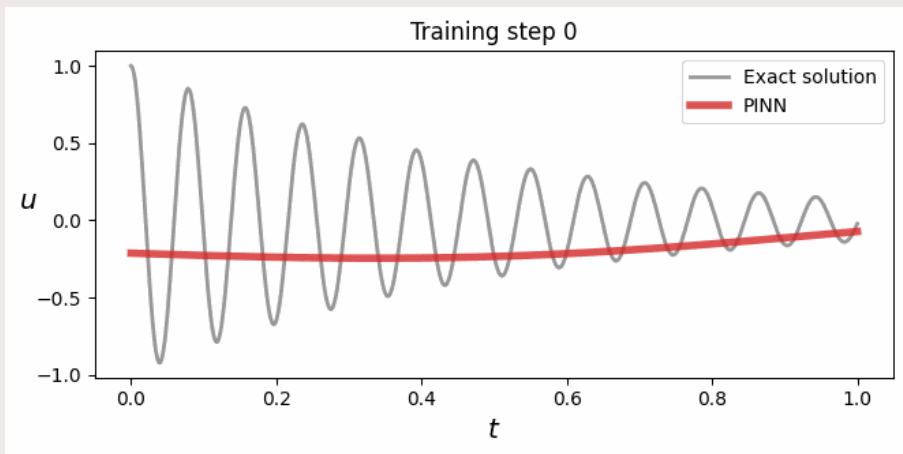
Problem: PINNs **struggle** to solve high-frequency / multiscale problems



Damped harmonic oscillator



# Scaling PINNs to high frequency / multiscale problems



Problem: PINNs **struggle** to solve high-frequency / multiscale problems

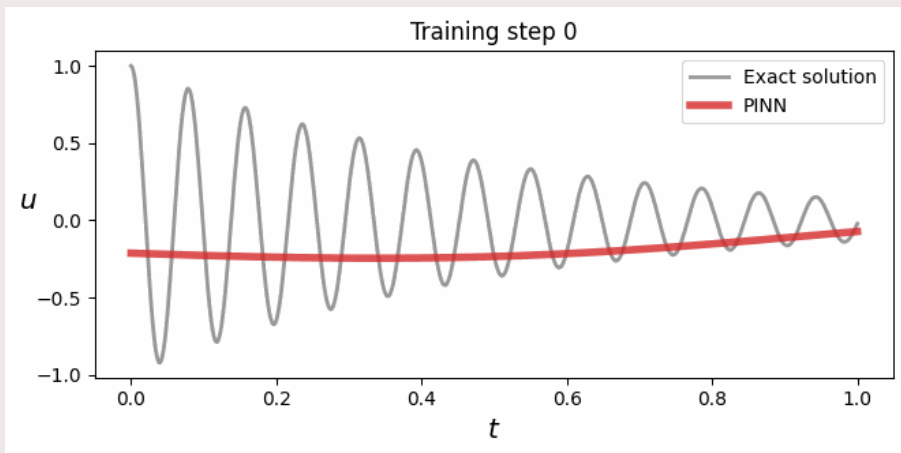
As higher frequencies are added:

- More collocation points required
- Larger neural network required
- Spectral bias slows convergence

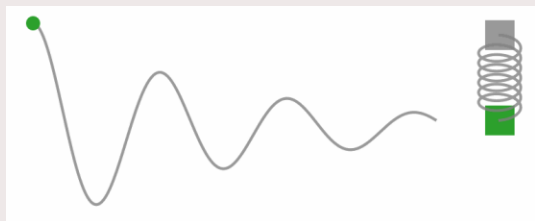


Damped harmonic oscillator

# Scaling PINNs to high frequency / multiscale problems



Problem: PINNs **struggle** to solve high-frequency / multiscale problems



Damped harmonic oscillator

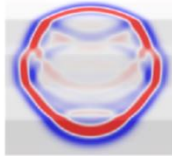
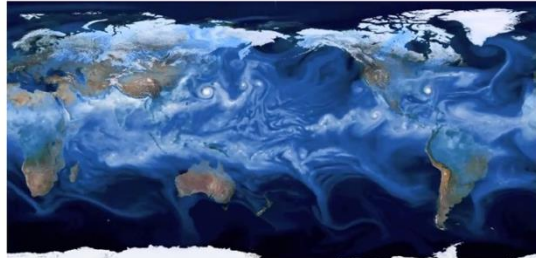
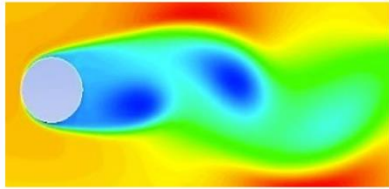
As higher frequencies are added:

- More collocation points required ( $\propto \omega^d$ )
- Larger neural network required ( $\propto f(\omega)$ )
- Spectral bias slows convergence ( $\propto s(\omega)$ )

$\Rightarrow$  Empirically, cost of training often  
 $\sim \mathcal{O}(\omega^d f(\omega) s(\omega))$

c.f. FD simulation, where cost of simulation  
can scale like  $\sim \mathcal{O}(\omega^d)$

# Scaling to more complex problems



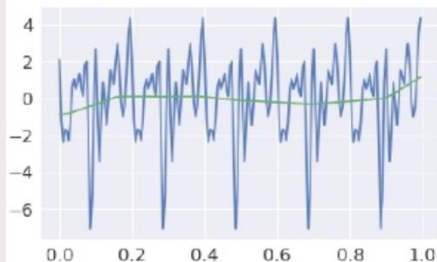
Majority of PINN research focuses on **toy/simplified** problems,  
as proof-of-principle studies

It is often challenging to **scale**  
traditional scientific algorithms to:

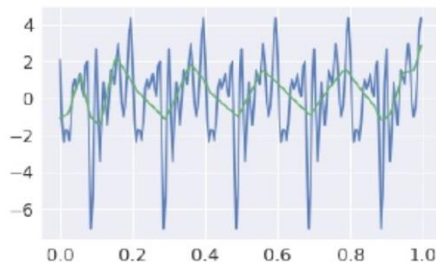
- More complex phenomena  
(**multi-scale, multi-physics**)
- Large domains / higher  
frequency solutions
- Incorporate **real, noisy** and  
**sparse** data

How do PINNs cope in this setting?

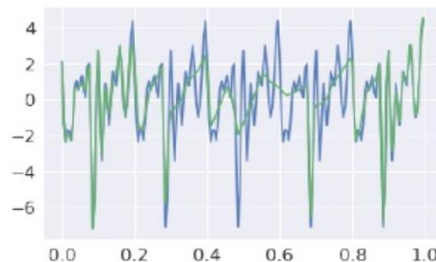
# Spectral bias issue



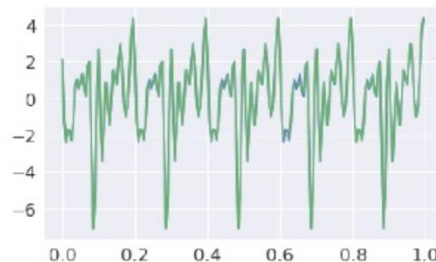
(a) Iteration 100



(b) Iteration 1000



(c) Iteration 10000

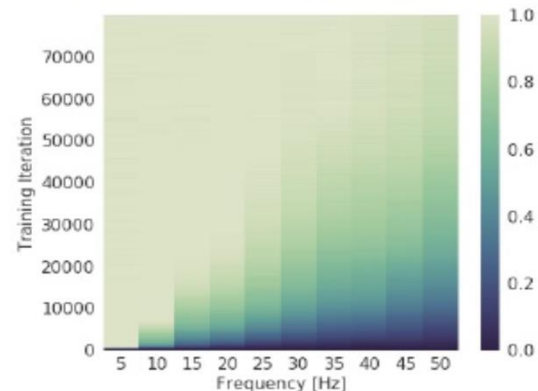


(d) Iteration 80000

NNs prioritise learning **lower** frequency functions first

Under certain assumptions can be proved via neural tangent kernel theory

Rahaman, N., et al, On the spectral bias of neural networks. 36th International Conference on Machine Learning, ICML (2019)



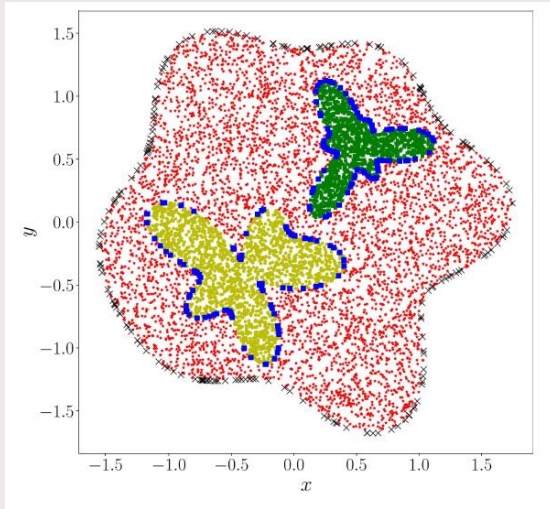
This behaviour can limit the model's ability to scale effectively, especially when the goal is to approximate complex scientific data with both large-scale and fine-scale dynamics.

# Divide and conquer approaches tackle spectral bias

**Partition of Unity Networks (POUNets)** offer several mechanisms to mitigate spectral bias, especially in the context of scaling SciML models.

- **Partitioning the domain** into smaller regions,
- **Allowing localized learning** of high-frequency features.
- **Leveraging multiscale representations**, and
- **Ensuring smooth transitions across regions**

# PINNs + domain decomposition



Jagtap, A., et al., Extended physics-informed neural networks (XPINNs): A generalized space-time domain decomposition based deep learning framework for nonlinear partial differential equations. Communications in Computational Physics (2020)

Idea:

Take a “**divide-and-conquer**” strategy to model more complex problems:

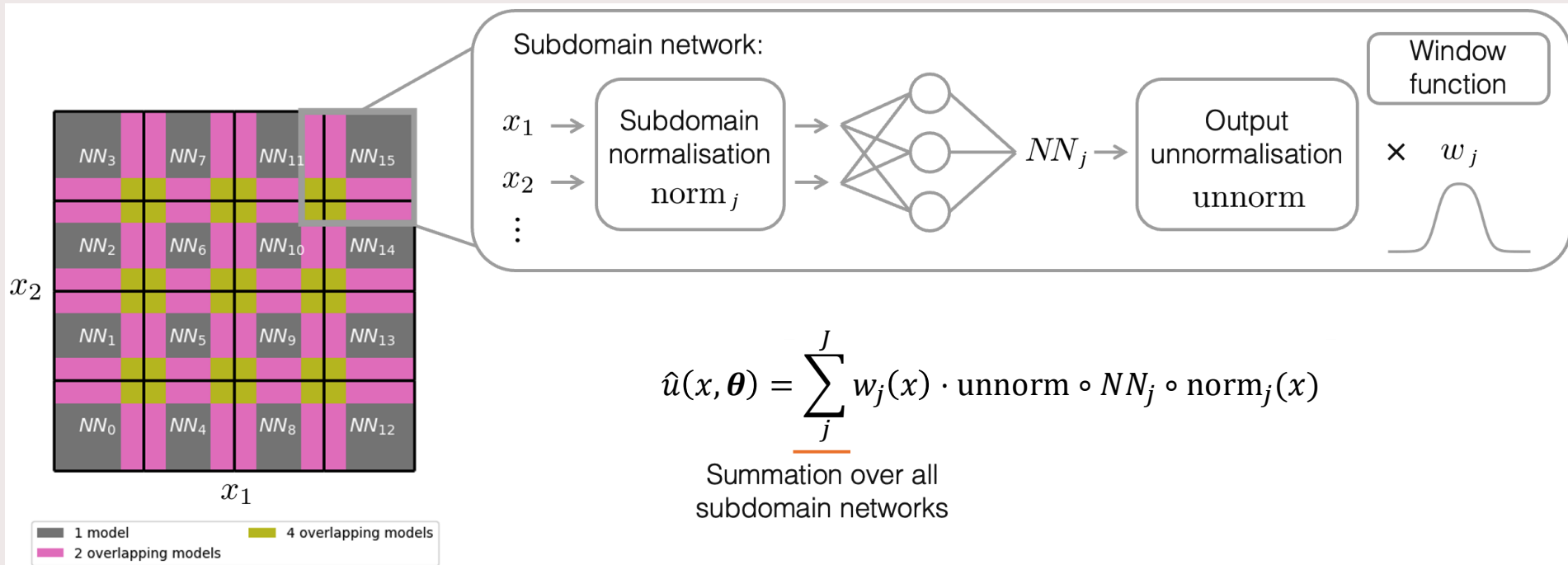
1. Divide modelling domain into many smaller **subdomains**
2. Use a separate neural network in each subdomain to model the solution

Hypothesis:

The resulting (coupled) local optimization problems are easier to solve than a single global problem



# Finite basis PINNs (FBPINNs)

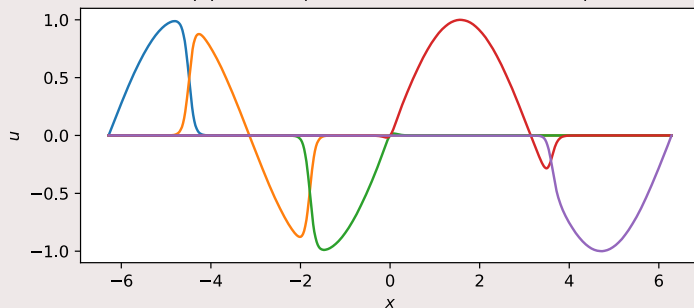


Moseley et al, Finite Basis Physics-Informed Neural Networks (FBPINNs): a scalable domain decomposition approach for solving differential equations, ACM (2023)

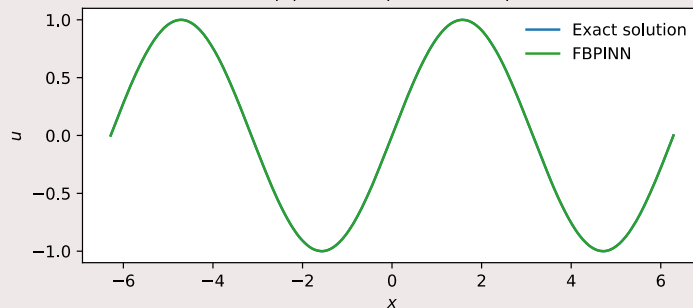
Idea: use **overlapping** subdomains and a **globally** defined solution ansatz

# FBPINNs in 1D

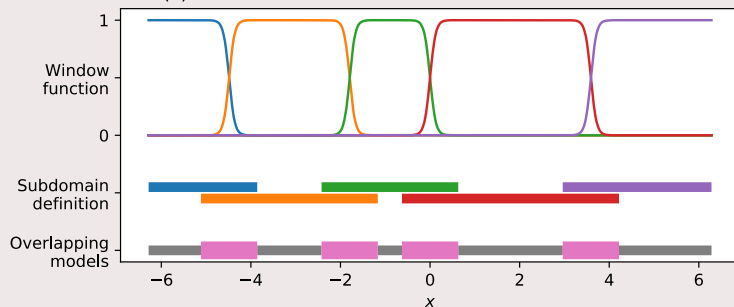
(a) FBPINN (individual network solutions)



(b) FBPINN (full solution)



(c) FBPINN subdomain definition and window functions



$$\hat{u}(x, \theta) = \sum_j^J w_j(x) \cdot \text{unnorm} \circ NN_j \circ \text{norm}_j(x)$$

Window  
function

Subdomain  
network

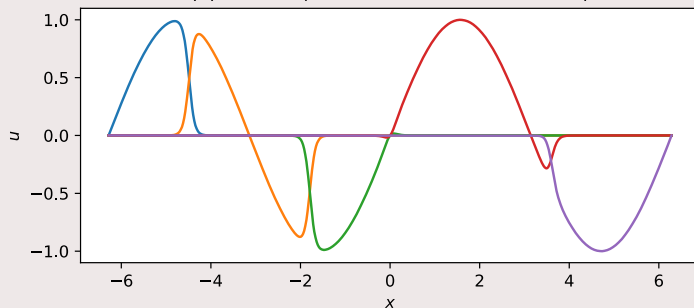
Individual subdomain  
normalisation

Idea: use **overlapping** subdomains and a **globally** defined solution ansatz

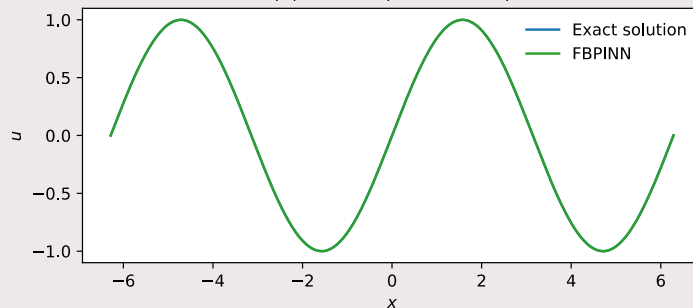
Moseley et al, Finite Basis Physics-Informed Neural Networks (FBPINNs): a scalable domain decomposition approach for solving differential equations, ACM (2023)

# FBPINNs in 1D

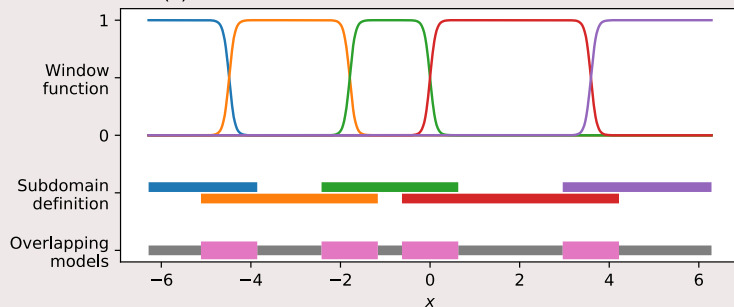
(a) FBPINN (individual network solutions)



(b) FBPINN (full solution)



(c) FBPINN subdomain definition and window functions



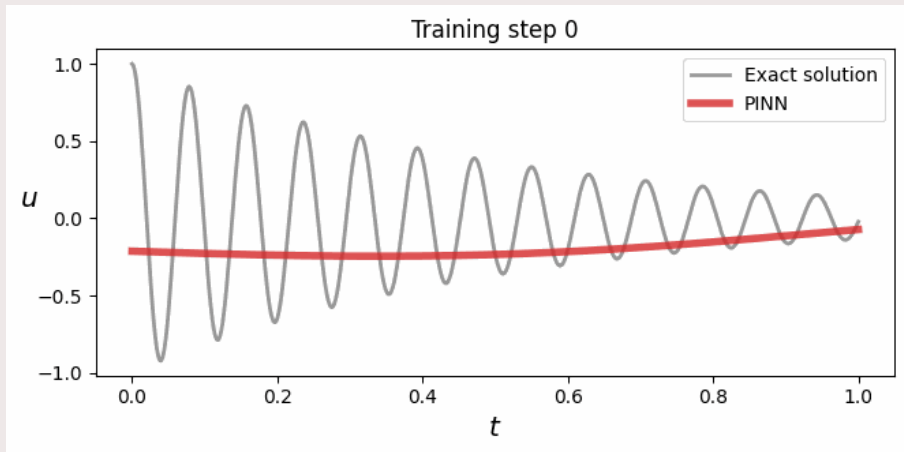
$$\hat{u}(x, \theta) = \sum_j^J w_j(x) \cdot \text{unnorm} \circ NN_j \circ \text{norm}_j(x)$$

Window function    Subdomain network    Individual subdomain normalisation

Notes:

- FBPINNs can be trained with same loss function as PINNs
- And can simply be thought of as a “custom architecture”

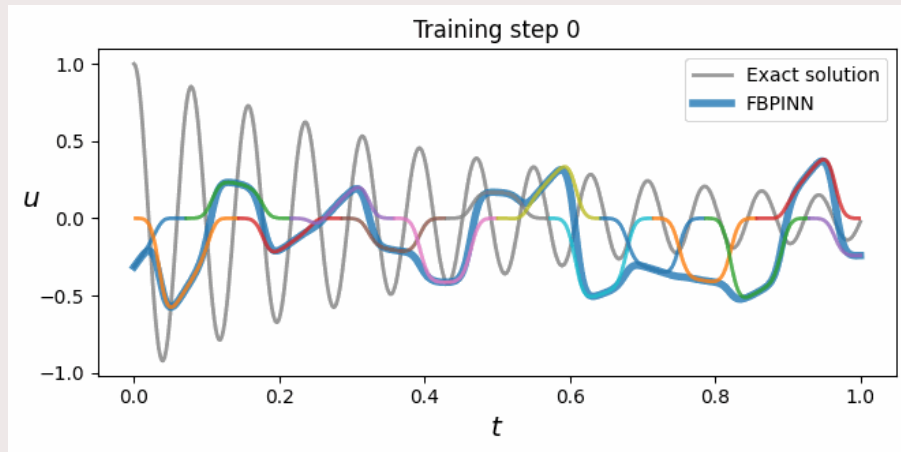
# FBPINNs vs PINNs



Problem: PINNs **struggle** to solve high-frequency / multiscale problems



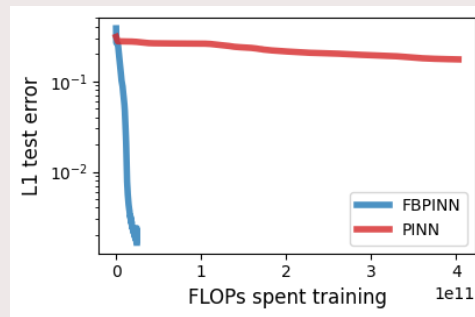
Damped harmonic oscillator



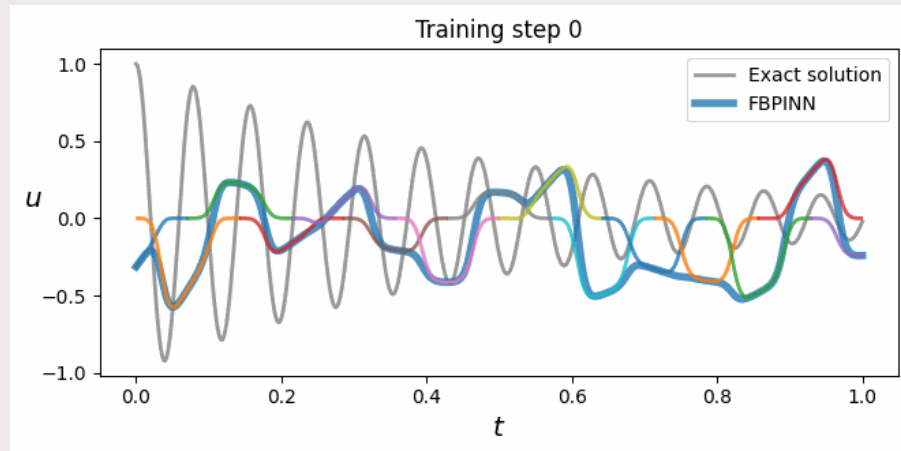
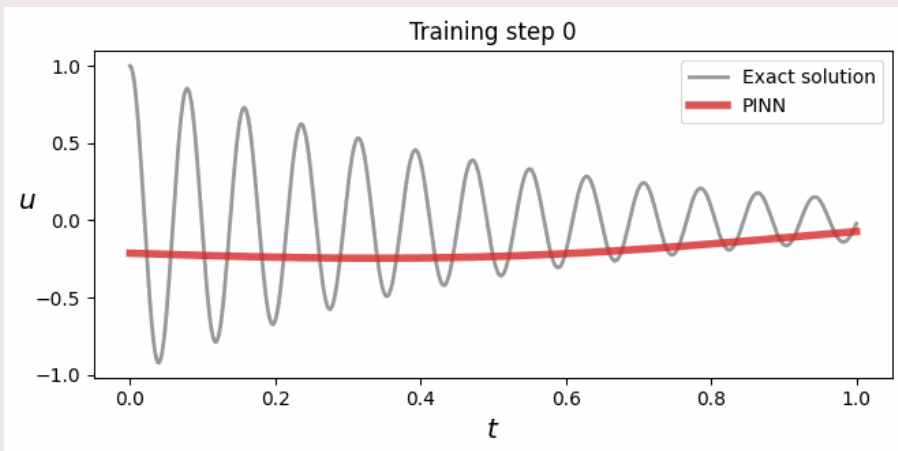
FBPINN solution

Number of subdomains: 15

Subdomain networks: 1 hidden layer, 32 hidden units



# FBPINNs vs PINNs

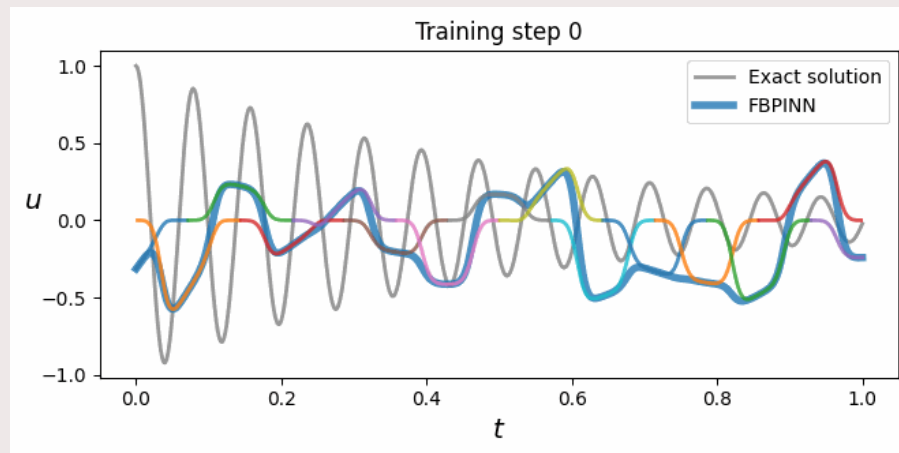
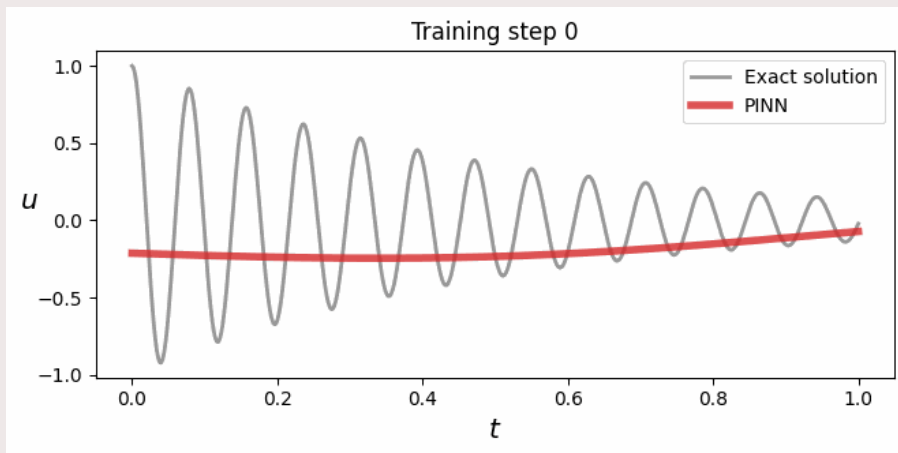


As higher frequencies are added:

- More collocation points required ( $\propto \omega^d$ )
- Larger neural network required ( $\propto f(\omega)$ )
- Spectral bias slows convergence ( $\propto s(\omega)$ )

⇒ Empirically, cost of training often  $\sim \mathcal{O}(\omega^d f(\omega) s(\omega))$

# FBPINNs vs PINNs



As higher frequencies are added:

- More collocation points required ( $\propto \omega^d$ )
- Larger neural network required ( $\propto f(\omega)$ )
- Spectral bias slows convergence ( $\propto s(\omega)$ )

$\Rightarrow$  Empirically, cost of training often  $\sim \mathcal{O}(\omega^d f(\omega) s(\omega))$

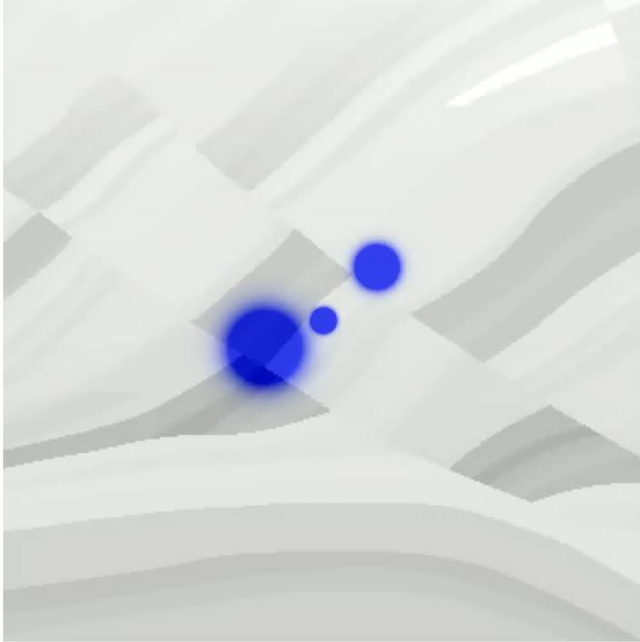
As higher frequencies are added:

- More collocation points required ( $\propto \omega^d$ )
- Same size network can be used in each subdomain
- Domain decomposition alleviates spectral bias

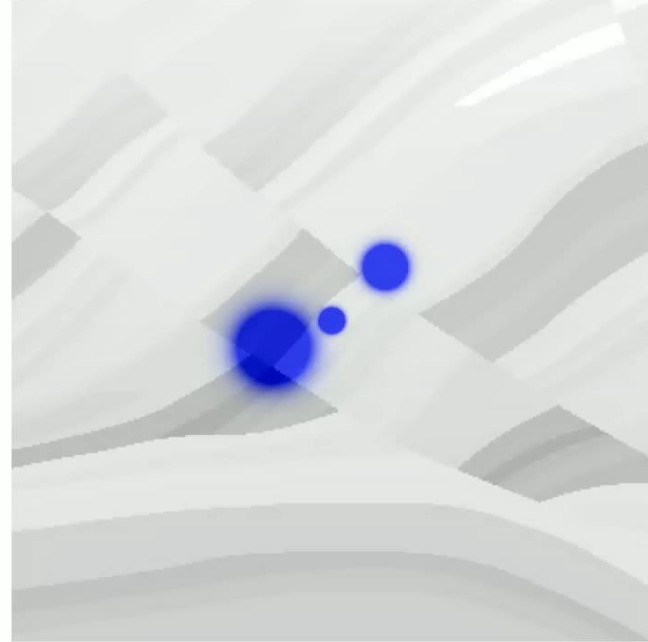
$\Rightarrow$  Empirically, cost of training can be closer to  $\sim \mathcal{O}(\omega^d)$

# Multi-scale simulation with FBPINNs

FBPINN solution



FD simulation



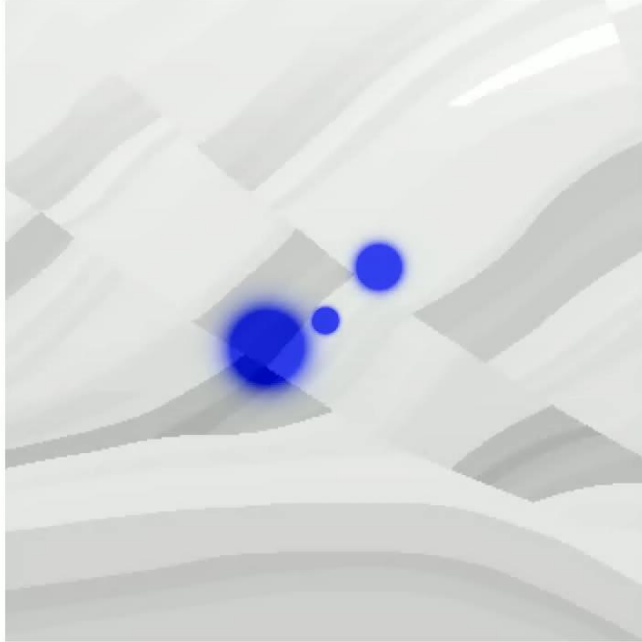
Number of subdomains:  $60 \times 60 \times 60 = 216,000$

Total number of trainable parameters: 9 M

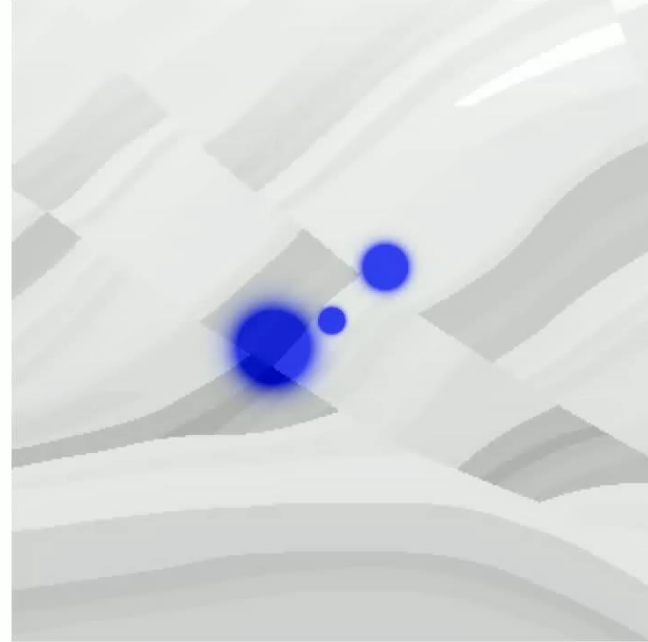


# Multi-scale simulation with FBPINNs

FBPINN solution



FD simulation

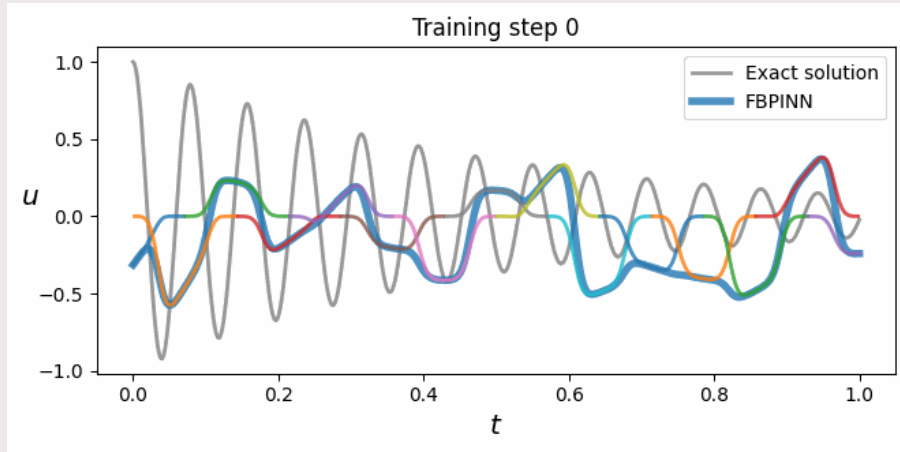


Training time: **~2 hrs** on GPU (with optimised code)

Number of subdomains:  $60 \times 60 \times 60 = 216,000$   
Total number of trainable parameters: 9 M

FD simulation time: ~5 mins on CPU!

# Why are FBPINNs still slow?



As higher frequencies are added:

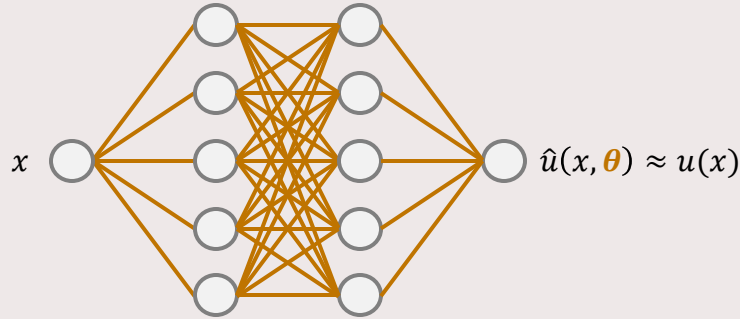
- More collocation points required ( $\propto \omega^d$ )
- Same size network can be used in each subdomain
- Domain decomposition alleviates spectral bias

⇒ Empirically, cost of training can be closer to  $\sim \mathcal{O}(\omega^d)$

**BUT** gradient descent is a slow optimiser (non-convex loss requires lots of iterations + backprop introduces lots of overhead)

..can we **avoid** gradient descent altogether?

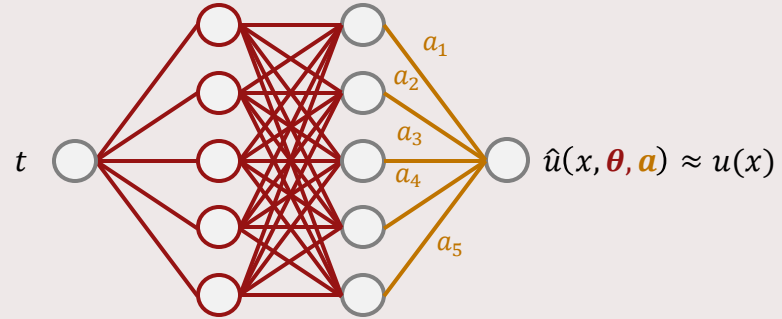
# Idea – Extreme learning machines



Neural network

All weights **trainable**

$$\hat{u} = NN(x, \theta)$$



Extreme learning machine

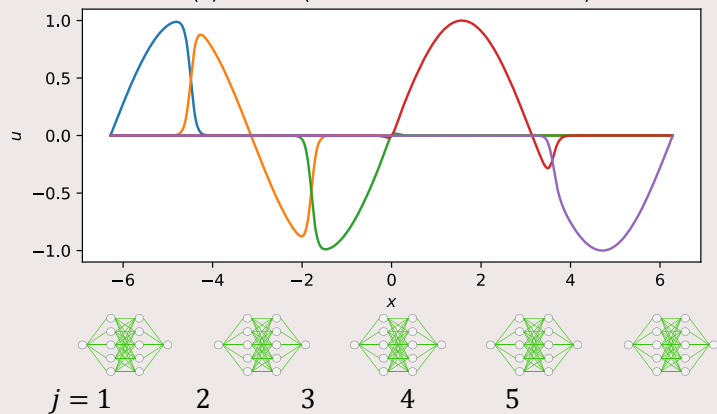
Hidden weights are **randomly** initialised and **fixed**

Only last layer **trainable**

$$\hat{u} = \sum_k^K a_k \phi(x, \theta_k)$$

# FBPINN

(a) FBPINN (individual network solutions)



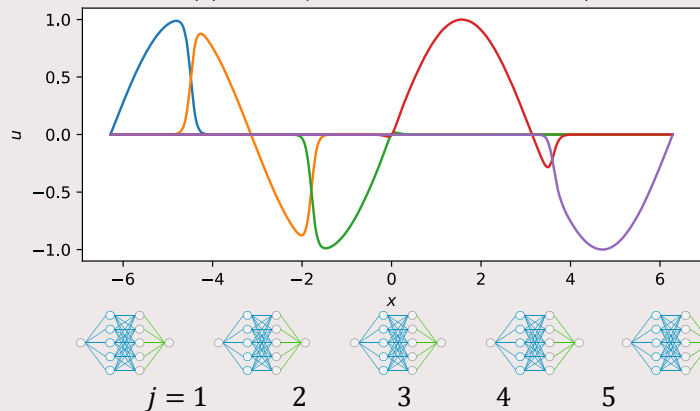
$$\hat{u}(x, \theta) = \sum_j^J w_j(x) NN_j(x, \theta_j) \quad \text{FBPINN}$$

Window function      Subdomain network

(ignoring normalization functions for simplicity)

# Extreme learning machine (ELM) FBPINNs

(a) FBPINN (individual network solutions)



$$\hat{u}(x, \mathbf{a}) = \sum_j^J w_j(x) \sum_k^K a_{jk} \phi(x, \theta_{jk}) \quad \text{ELM-FBPINN}$$

**ELM in each  
subdomain**

$J$  = total number of subdomains

$K$  = number of basis functions per subdomain

Dwivedi, V., and Srinivasan, B. Physics Informed Extreme Learning Machine (PIELM)—A rapid method for the numerical solution of partial differential equations. *Neurocomputing*. (2020).

Dong, S., and Li, Z. Local extreme learning machines and domain decomposition for solving linear and nonlinear partial differential equations. *Computer Methods in Applied Mechanics and Engineering*. (2021).

# Extreme learning machine (ELM) FBPINNs



$$\hat{u}(x, \mathbf{a}) = \sum_j^J w_j(x) \sum_k^K \mathbf{a}_{jk} \phi(x, \boldsymbol{\theta}_{jk})$$

ELM-FBPINN

ELM in **each** subdomain

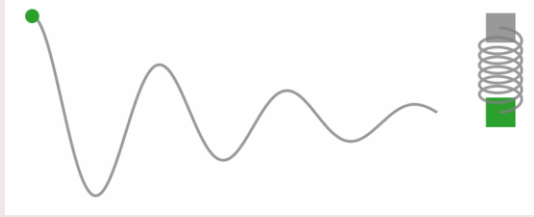
$J$  = total number of subdomains

$K$  = number of basis functions per subdomain

$N$  = number of collocation points

$$\begin{aligned} L &= \sum_i^N \left( \left[ m \frac{d^2}{dt^2} + \mu \frac{d}{dt} + k \right] \hat{u}(t_i, \boldsymbol{\theta}) \right)^2 \quad (\text{ignoring boundary loss for simplicity}) \\ &= \sum_i^N \left( \left[ m \frac{d^2}{dt^2} + \mu \frac{d}{dt} + k \right] \sum_j^J w_j(t_i) \sum_k^K a_{jk} \phi(x, \boldsymbol{\theta}_{jk}) \right)^2 \\ &= \sum_i^N \left( \sum_j^J \sum_k^K a_{jk} \mathcal{N} w_j(t_i) \phi(t_i, \boldsymbol{\theta}_{jk}) \right)^2 = \left\| \begin{pmatrix} \mathcal{N} w_j(t_0) \phi(t_0, \boldsymbol{\theta}_{00}) & \dots & \mathcal{N} w_j(t_0) \phi(t_0, \boldsymbol{\theta}_{JK}) \\ \vdots & \ddots & \vdots \\ \mathcal{N} w_j(t_N) \phi(t_N, \boldsymbol{\theta}_{00}) & \dots & \mathcal{N} w_j(t_N) \phi(t_N, \boldsymbol{\theta}_{JK}) \end{pmatrix} \begin{pmatrix} a_{00} \\ \vdots \\ a_{JK} \end{pmatrix} - \begin{pmatrix} 0 \\ \vdots \\ 0 \end{pmatrix} \right\|^2 \end{aligned}$$

# Extreme learning machine (ELM) FBPINNs



$$\hat{u}(x, \mathbf{a}) = \sum_j^J w_j(x) \sum_k^K \mathbf{a}_{jk} \phi(x, \boldsymbol{\theta}_{jk}) \quad \text{ELM-FBPINN}$$

ELM in **each** subdomain

$J$  = total number of subdomains

$K$  = number of basis functions per subdomain

$N$  = number of collocation points

$$\begin{aligned} L &= \sum_i^N \left( \left[ m \frac{d^2}{dt^2} + \mu \frac{d}{dt} + k \right] \hat{u}(t_i, \boldsymbol{\theta}) \right)^2 \\ &= \sum_i^N \left( \left[ m \frac{d^2}{dt^2} + \mu \frac{d}{dt} + k \right] \underbrace{\sum_j^J w_j(t_i) \sum_k^K \mathbf{a}_{jk} \phi(x, \boldsymbol{\theta}_{jk})}_{\text{ELM-FBPINN}} \right)^2 \\ &= \sum_i^N \left( \sum_j^J \sum_k^K \mathbf{a}_{jk} \mathcal{N} w_j(t_i) \phi(t_i, \boldsymbol{\theta}_{jk}) \right)^2 = \left\| \begin{pmatrix} \mathcal{N} w_j(t_0) \phi(t_0, \boldsymbol{\theta}_{00}) & \dots & \mathcal{N} w_j(t_0) \phi(t_0, \boldsymbol{\theta}_{JK}) \\ \vdots & \ddots & \vdots \\ \mathcal{N} w_j(t_N) \phi(t_N, \boldsymbol{\theta}_{00}) & \dots & \mathcal{N} w_j(t_N) \phi(t_N, \boldsymbol{\theta}_{JK}) \end{pmatrix} \begin{pmatrix} a_{00} \\ \vdots \\ a_{JK} \end{pmatrix} - \begin{pmatrix} 0 \\ \vdots \\ 0 \end{pmatrix} \right\|^2 \end{aligned}$$

# Extreme learning machine (ELM) FBPINNs



$$\hat{u}(x, \mathbf{a}) = \sum_j^J w_j(x) \sum_k^K a_{jk} \phi(x, \boldsymbol{\theta}_{jk})$$

ELM-FBPINN

ELM in **each** subdomain

$J$  = total number of subdomains

$K$  = number of basis functions per subdomain

$N$  = number of collocation points

$$\begin{aligned} L &= \sum_i^N \left( \left[ m \frac{d^2}{dt^2} + \mu \frac{d}{dt} + k \right] \hat{u}(t_i, \boldsymbol{\theta}) \right)^2 \\ &= \sum_i^N \left( \left[ m \frac{d^2}{dt^2} + \mu \frac{d}{dt} + k \right] \sum_j^J w_j(t_i) \sum_k^K a_{jk} \phi(x, \boldsymbol{\theta}_{jk}) \right)^2 \\ &= \sum_i^N \left( \sum_j^J \sum_k^K a_{jk} \mathcal{N} w_j(t_i) \phi(t_i, \boldsymbol{\theta}_{jk}) \right)^2 = \left\| \begin{pmatrix} \mathcal{N} w_j(t_0) \phi(t_0, \boldsymbol{\theta}_{00}) & \dots & \mathcal{N} w_j(t_0) \phi(t_0, \boldsymbol{\theta}_{JK}) \\ \vdots & \ddots & \vdots \\ \mathcal{N} w_j(t_N) \phi(t_N, \boldsymbol{\theta}_{00}) & \dots & \mathcal{N} w_j(t_N) \phi(t_N, \boldsymbol{\theta}_{JK}) \end{pmatrix} \begin{pmatrix} a_{00} \\ \vdots \\ a_{JK} \end{pmatrix} - \begin{pmatrix} 0 \\ \vdots \\ 0 \end{pmatrix} \right\|^2 \end{aligned}$$

Assuming  $\mathcal{N}$  is a linear operator,  $\mathcal{N} = \left[ m \frac{d^2}{dt^2} + \mu \frac{d}{dt} + k \right]$



# Extreme learning machine (ELM) FBPINNs



$$\hat{u}(x, \mathbf{a}) = \sum_j^J w_j(x) \sum_k^K \mathbf{a}_{jk} \phi(x, \boldsymbol{\theta}_{jk}) \quad \text{ELM-FBPINN}$$

ELM in **each** subdomain

$J$  = total number of subdomains

$K$  = number of basis functions per subdomain

$N$  = number of collocation points

$$\begin{aligned} L &= \sum_i^N \left( \left[ m \frac{d^2}{dt^2} + \mu \frac{d}{dt} + k \right] \hat{u}(t_i, \boldsymbol{\theta}) \right)^2 \\ &= \sum_i^N \left( \left[ m \frac{d^2}{dt^2} + \mu \frac{d}{dt} + k \right] \sum_j^J w_j(t_i) \sum_k^K \mathbf{a}_{jk} \phi(x, \boldsymbol{\theta}_{jk}) \right)^2 \\ &= \sum_i^N \left( \sum_j^J \sum_k^K \mathbf{a}_{jk} \mathcal{N} w_j(t_i) \phi(t_i, \boldsymbol{\theta}_{jk}) \right)^2 = \left\| \begin{pmatrix} \mathcal{N} w_j(t_0) \phi(t_0, \boldsymbol{\theta}_{00}) & \dots & \mathcal{N} w_j(t_0) \phi(t_0, \boldsymbol{\theta}_{JK}) \\ \vdots & \ddots & \vdots \\ \mathcal{N} w_j(t_N) \phi(t_N, \boldsymbol{\theta}_{00}) & \dots & \mathcal{N} w_j(t_N) \phi(t_N, \boldsymbol{\theta}_{JK}) \end{pmatrix} \begin{pmatrix} \mathbf{a}_{00} \\ \vdots \\ \mathbf{a}_{JK} \end{pmatrix} - \begin{pmatrix} 0 \\ \vdots \\ 0 \end{pmatrix} \right\|^2 \end{aligned}$$

$$\equiv \|M\mathbf{a} - \mathbf{f}\|^2$$

$$M: N \times JK \quad \mathbf{a}: JK \quad \mathbf{f}: N$$

# Quadratic optimisation

This is a linear least squares problem!

$$L(\mathbf{a}) = \|\mathbf{M}\mathbf{a} - \mathbf{f}\|^2$$

The global minima is given by solving

$$\mathbf{A} \mathbf{a}^* = \mathbf{b} \quad (\text{normal equation})$$

Where,

$$\begin{aligned} \mathbf{A} &= \mathbf{M}^T \mathbf{M} & JK \times JK \\ \mathbf{b} &= \mathbf{M}^T \mathbf{f} \end{aligned}$$

# Quadratic optimisation

This is a linear least squares problem!

$$L(\mathbf{a}) = \|M\mathbf{a} - \mathbf{f}\|^2$$

The global minima is given by solving

$$A \mathbf{a}^* = \mathbf{b} \quad (\text{normal equation})$$

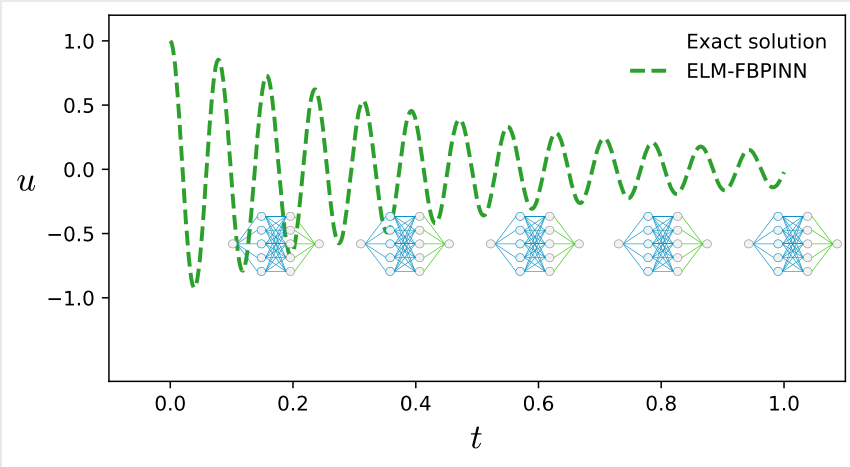
Where,

$$A = M^T M \quad JK \times JK$$

$$\mathbf{b} = M^T \mathbf{f} \quad JK$$

- By using a linear combination of fixed basis functions, we have turned the loss function from non-convex to convex (quadratic)
- I.e., we can now use **linear solvers** to train ELM-FBPINNs, instead of gradient descent! 

# Example – 1D harmonic oscillator



FBPINN / ELM-FBPINN:

20 subdomains

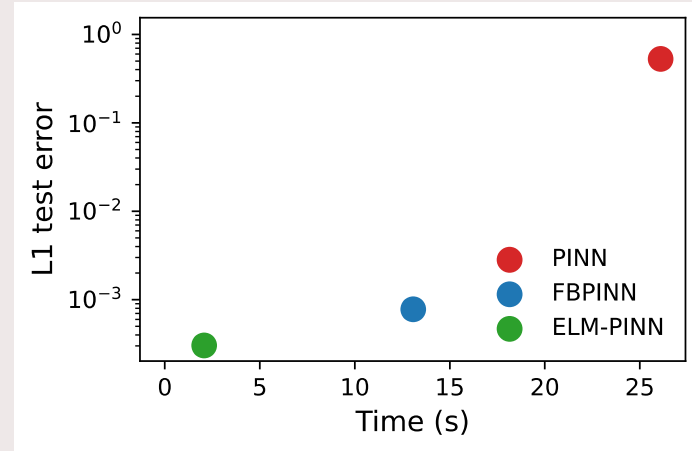
1 hidden layer, 8 hidden units (=basis functions)

Tanh activation

PINN:

2 hidden layers, 64 hidden units

Tanh activation



PINN / FBPINN: Adam optimiser, 0.001 learning rate

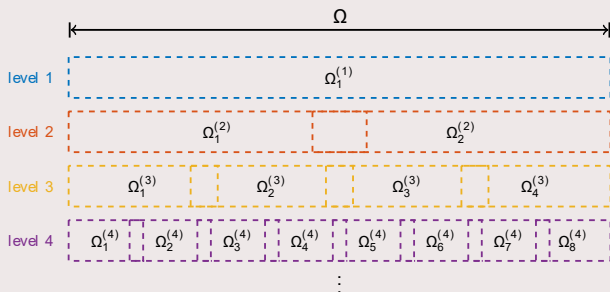
ELM-FBPINN: Conjugate gradient linear solver

# Example – 2D multi-scale Laplace

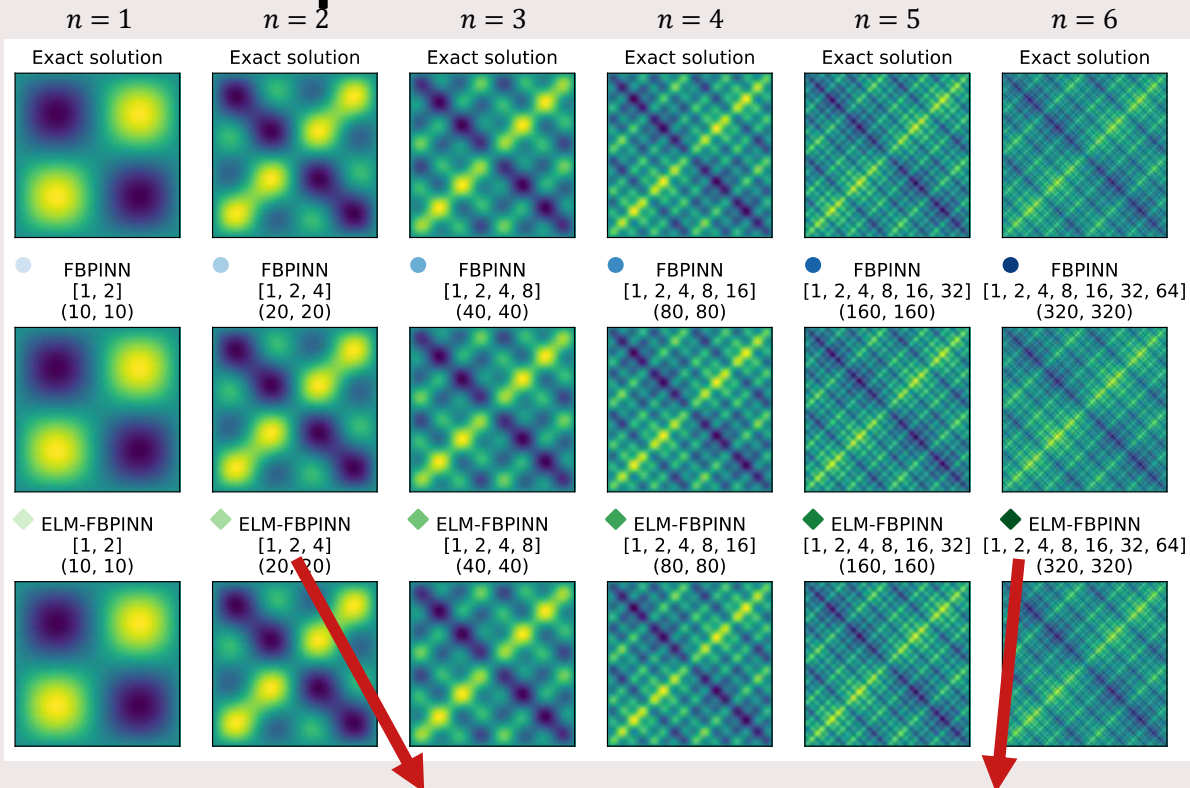
**Multi-scale problem:**

$$\nabla^2 u(x_1, x_2) = -\frac{2}{n} \sum_{i=1}^n (2^i \pi)^2 \sin(2^i \pi x_1) \sin(2^i \pi x_2)$$

**Multilevel domain decomposition:**



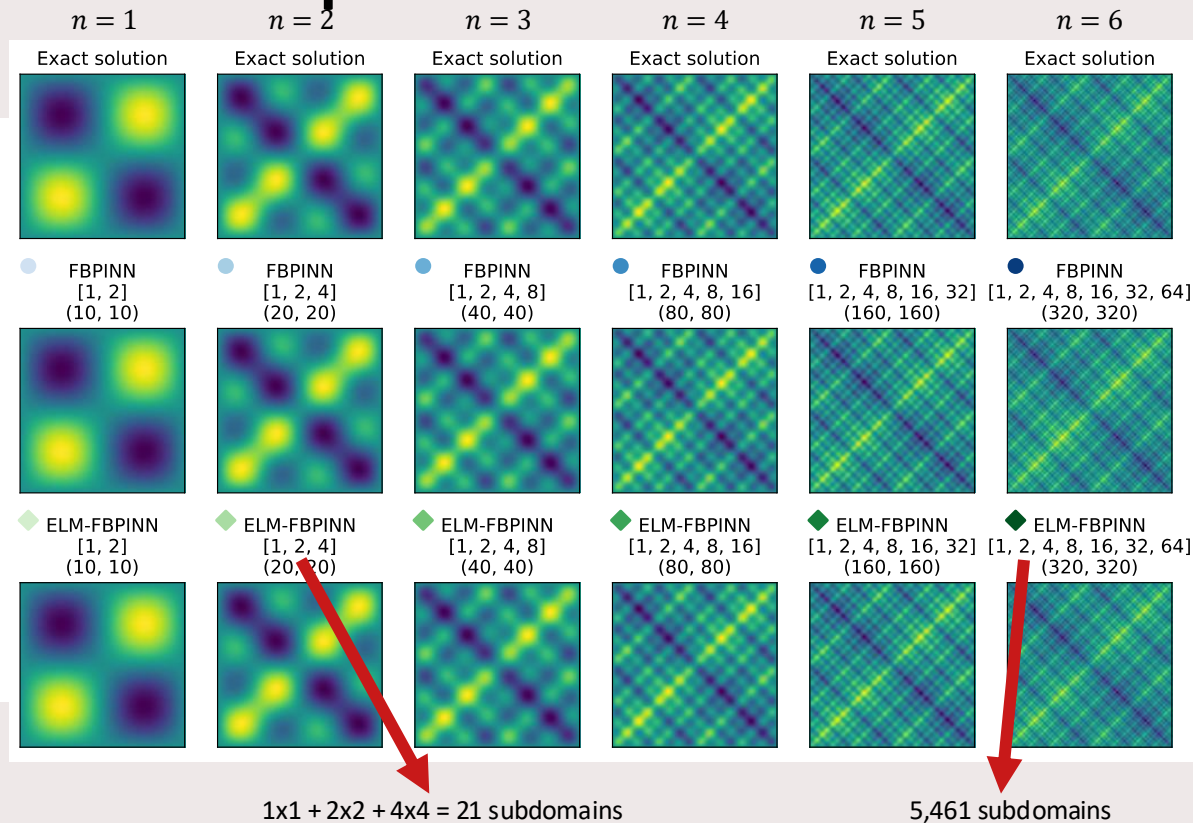
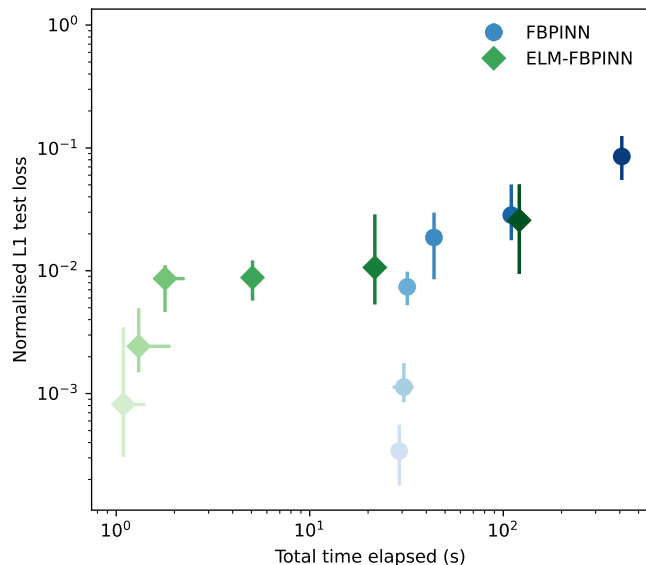
Dolean, V., et al, Multilevel domain decomposition-based architectures for physics-informed neural networks, CMAME (2024)



$1 \times 1 + 2 \times 2 + 4 \times 4 = 21$  subdomains

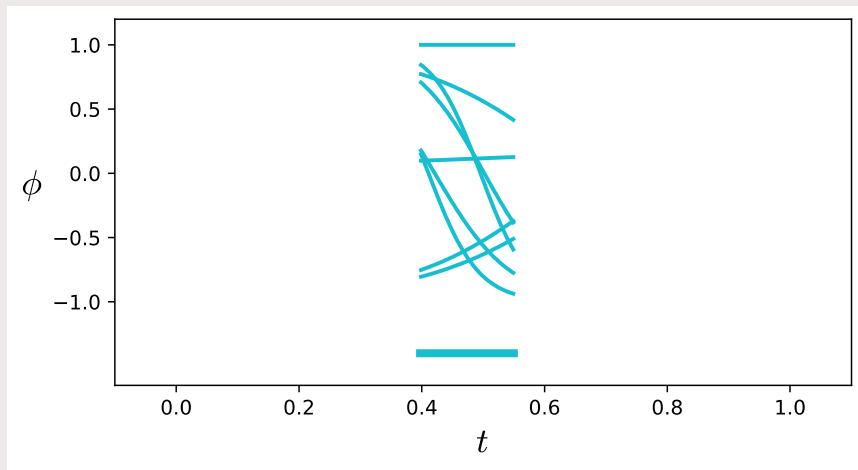
5,461 subdomains

# Example – 2D multi-scale Laplace

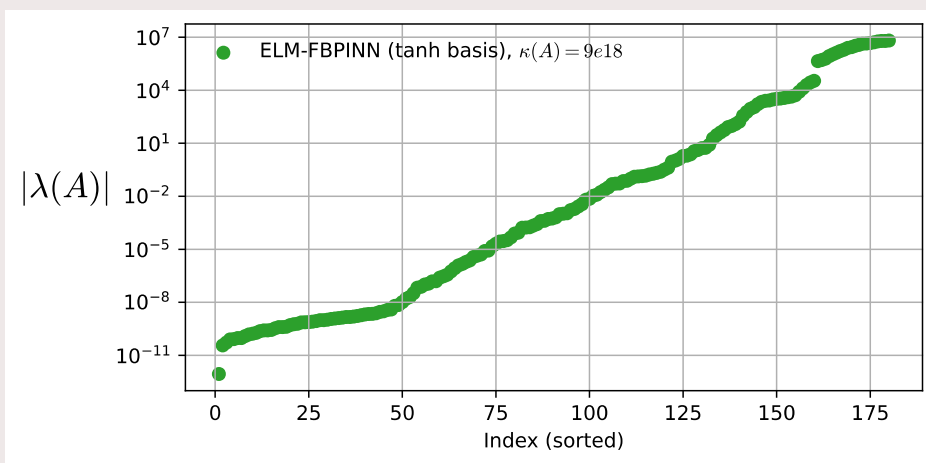


# Challenges with ELM-FBPINNs

Challenge 1: Linear dependence between basis functions  $\Rightarrow$  **poorly conditioned matrix**  $A \mathbf{a}^* = \mathbf{b}$



Conjugate gradient solver requires **~5000** iterations

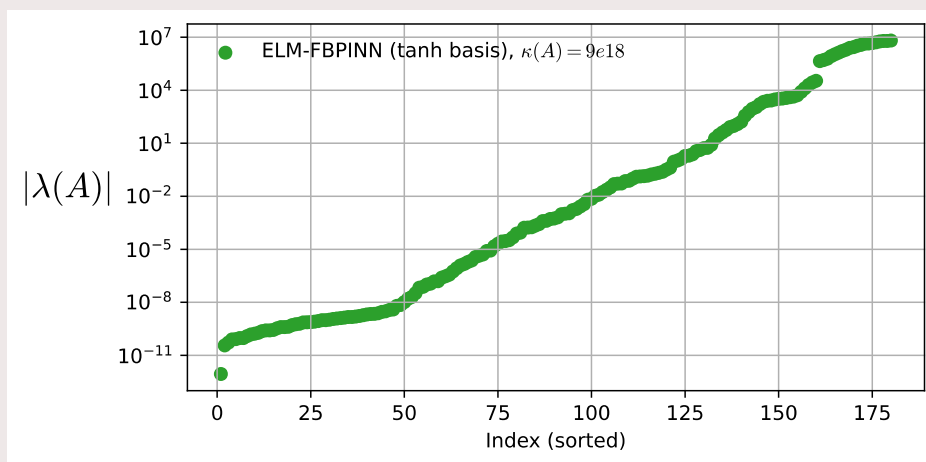
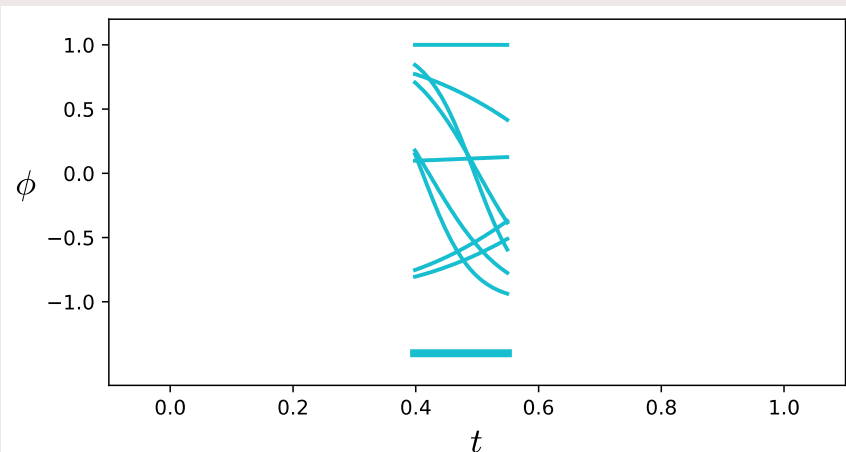


$$M = \begin{pmatrix} \mathcal{N}w_j(t_0)\phi(t_0, \boldsymbol{\theta}_{00}) & \dots & \mathcal{N}w_j(t_0)\phi(t_0, \boldsymbol{\theta}_{JK}) \\ \vdots & \ddots & \vdots \\ \mathcal{N}w_j(t_N)\phi(t_N, \boldsymbol{\theta}_{00}) & \dots & \mathcal{N}w_j(t_N)\phi(t_N, \boldsymbol{\theta}_{JK}) \end{pmatrix},$$

$$A = M^T M$$

# Challenges with ELM-FBPINNs

Challenge 1: Linear dependence between basis functions  $\Rightarrow$  **poorly conditioned matrix**  $A \mathbf{a}^* = \mathbf{b}$



Possible solution: use **principal component analysis / preconditioning**, see:

Shang, Y., Heinlein, A., Mishra, S., & Wang, F. Overlapping Schwarz Preconditioners for Randomized Neural Networks with Domain Decomposition. ArXiv. (2024).

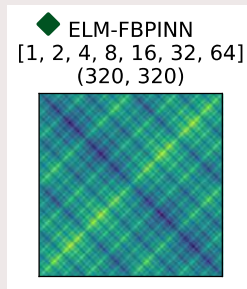
$$M = \begin{pmatrix} \mathcal{N}w_j(t_0)\phi(t_0, \boldsymbol{\theta}_{00}) & \dots & \mathcal{N}w_j(t_0)\phi(t_0, \boldsymbol{\theta}_{JK}) \\ \vdots & \ddots & \vdots \\ \mathcal{N}w_j(t_N)\phi(t_N, \boldsymbol{\theta}_{00}) & \dots & \mathcal{N}w_j(t_N)\phi(t_N, \boldsymbol{\theta}_{JK}) \end{pmatrix},$$
$$A = M^T M$$



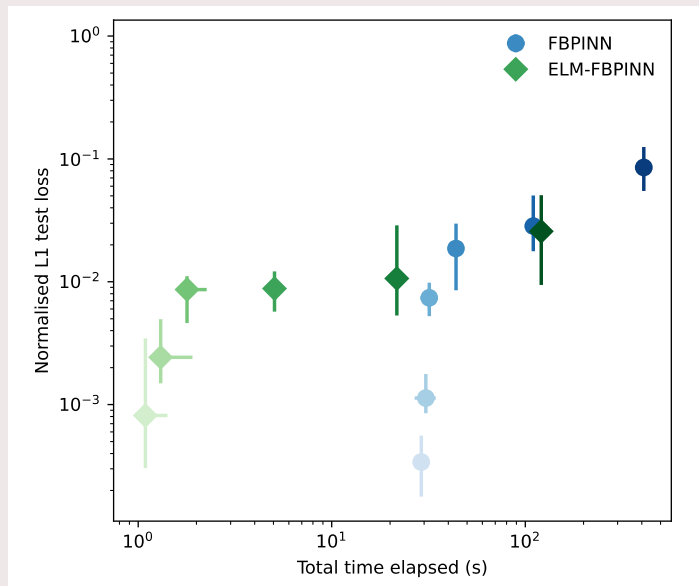
# Challenges with ELM-FBPINNs

Challenge 2: **Big matrix!**  $A \mathbf{a}^* = \mathbf{b}$

$A: JK \times JK$   $\mathbf{b}: JK$



$J = 5,461$ ,  $K = 6$   
 $A: 32,766 \times 32,766 = 1 \text{ billion elements!}$

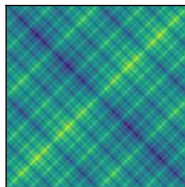


# Challenges with ELM-FBPINNs

Challenge 2: **Big matrix!**  $A \mathbf{a}^* = \mathbf{b}$

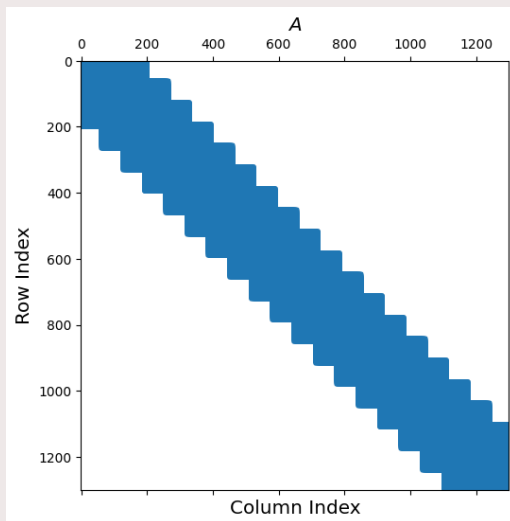
$$A: JK \times JK \quad \mathbf{b}: JK$$

◆ ELM-FBPINN  
[1, 2, 4, 8, 16, 32, 64]  
(320, 320)



$$J = 5,461, \quad K = 6$$
$$A: 32,766 \times 32,766 = \mathbf{1 \text{ billion elements!}}$$

Solution: exploit **sparsity**



Use sparse solver which only uses matvec products

```
scipy.sparse.linalg.  
cg  
  
cg(A, b, x0=None, *, rtol=1e-05, atol=0.0, maxiter=None, M=None,  
callback=None) \[source\]  
  
Use Conjugate Gradient iteration to solve  $A\mathbf{x} = \mathbf{b}$ .  
  
Parameters:  
A : {sparse array, ndarray, LinearOperator}  
The real or complex N-by-N matrix of the linear system. A must represent a  
hermitian, positive definite matrix. Alternatively, A can be a linear operator which can  
produce  $A\mathbf{x}$  using, e.g., scipy.sparse.linalg.LinearOperator.
```

# Can physics-informed neural networks (PINNs) beat finite difference / finite element methods?

## Can physics-informed neural networks beat the finite element method?

Tamara G Grossmann , Urszula Julia Komorowska, Jonas Latz, Carola-Bibiane Schönlieb

*IMA Journal of Applied Mathematics*, Volume 89, Issue 1, January 2024, Pages 143–174,

<https://doi.org/10.1093/imamat/hxae011>

**Published:** 23 May 2024    **Article history** ▼

## Solving inverse problems in physics by optimizing a discrete loss: Fast and accurate learning without neural networks

Petr Karnakov, Sergey Litvinov, Petros Koumoutsakos     **Author Notes**

*PNAS Nexus*, Volume 3, Issue 1, January 2024, pgae005,

<https://doi.org/10.1093/pnasnexus/pgae005>

**Published:** 11 January 2024    **Article history** ▼

## 7. Discussion and conclusions

After having investigated each of the PDEs on its own, let us now discuss and draw conclusions from the results as a whole. Considering the solution time and accuracy, PINNs are not able to beat FEM in our study. In all the examples that we have studied, the FEM solutions were faster at the same or at a higher accuracy.

## Conclusion

We introduce the ODIL framework for solving inverse problems for PDEs by casting their discretization as an optimization problem and applying optimization techniques that are widely available in machine-learning software. The concept of casting the PDE as is closely related to the neural network formulations proposed by (15–17) and recently revived as PINNs. However, the fact that we use the discrete approximation of the equations allows for ODIL to be orders of magnitude more efficient in terms of computational cost and accuracy compared to the PINN for which complex flow problems “remain elusive” (71).

# Are PINNs becoming FEM?

ELM-FBPINN

$$m \frac{d^2 u}{dt^2} + \mu \frac{du}{dt} + ku = f$$

$$\hat{u}(t, \mathbf{a}) = \sum_j^J w_j(t) \sum_k^K \mathbf{a}_{jk} \phi(t, \boldsymbol{\theta}_{jk})$$

$$\begin{aligned} L(\mathbf{a}) &= \sum_i^N \left( \left[ m \frac{d^2}{dt^2} + \mu \frac{d}{dt} + k \right] \hat{u}(t_i, \mathbf{a}) - f(t_i) \right)^2 \\ &= \|M\mathbf{a} - \mathbf{f}\|^2 \\ &\Rightarrow A\mathbf{a} - \mathbf{b} = \mathbf{0} \end{aligned}$$

$A: JK \times JK$        $\mathbf{b}: JK$   
(sparse & symmetric)

Finite element method

$$m \frac{d^2 u}{dt^2} + \mu \frac{du}{dt} + ku = f$$

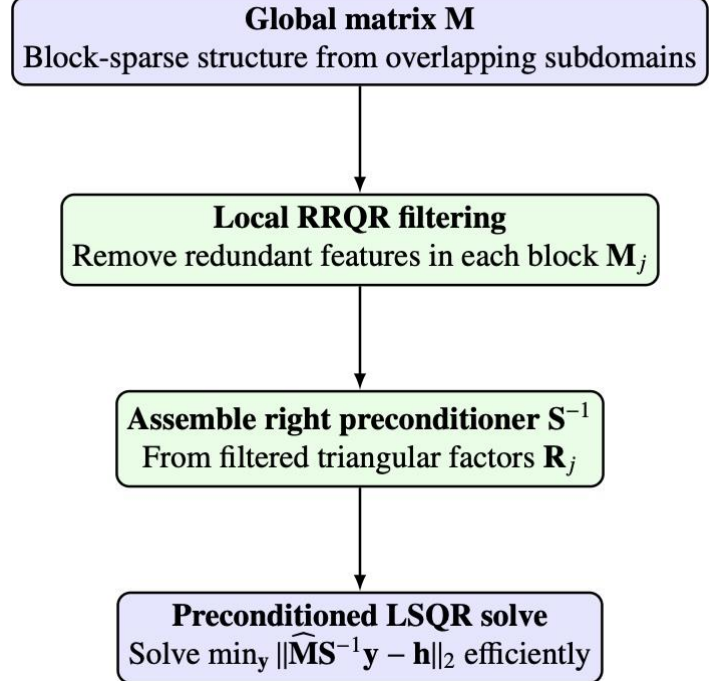
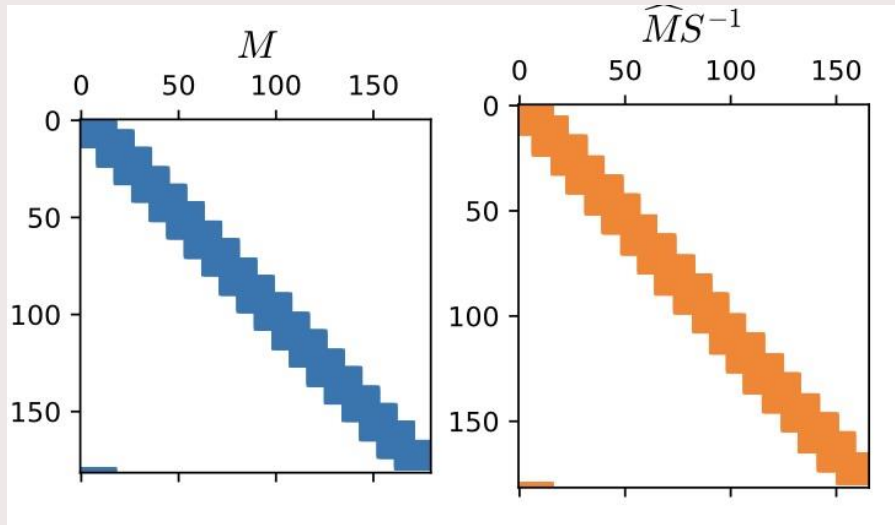
$$\hat{u}(t, \mathbf{a}) = \sum_j^J \mathbf{a}_j \phi_j(t)$$

$$\begin{aligned} \int_0^T \phi_i(t) \left[ m \frac{d^2}{dt^2} + \mu \frac{d}{dt} + k \right] \hat{u}(t, \mathbf{a}) dt &= \int_0^T \phi_i(t) f dt \quad \forall i = 0, \dots, J \\ \Rightarrow (-mL + \mu D + kM)\mathbf{a} &= \mathbf{b} \end{aligned}$$

$L, D, M: J \times J$        $\mathbf{b}: J$   
(sparse & symmetric)

# Can we do better i.e. how about preconditioning?

**Idea:** filter the redundant features in each block then precondition directly the least squares problem  
Sparsity is preserved.

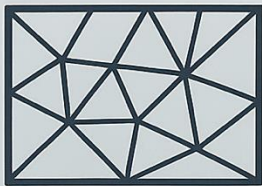


# Takeaways

- 💡 **From optimization to linear algebra:** reframing learning as a *numerical linear algebra* problem enables new algorithmic insights and requires new methodologies.
- ⚙️ **Scalability & robustness:** extending to *multilevel* and *overlapping* decompositions could improve performance in large-scale or high-dimensional settings.
- ⚡ **Acceleration & theory:** GPU implementations and convergence-rate analysis are key next steps.
- 🔄 **Beyond linear problems:** the framework naturally extends to nonlinear PDEs via Newton-type iterations.

# A change in perspective

Classical solver



$$\Delta u = f$$

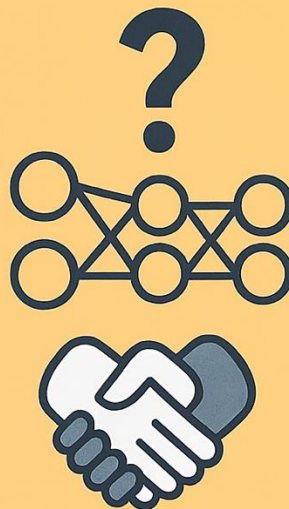


Proof



Error estimate

Neural network



Change of  
perspective

# Hybrid future

*The integration of machine learning (Keplerian paradigm) and more general artificial intelligence technologies with physical modelling based on first principles (Newtonian paradigm) will impact scientific computing in engineering in fundamental ways. (Stefan Kurz, ETH Zurich)*

