

Solving Differential Equations using Quantum Algorithms

A concise overview presentation



CrossMark

OPEN ACCESS

RECEIVED

22 October 2024

REVISED

28 November 2024

ACCEPTED FOR PUBLICATION

27 December 2024

PUBLISHED

13 January 2025

Original Content from
this work may be used
under the terms of the
Creative Commons
Attribution 4.0 licence.

Any further distribution
of this work must
maintain attribution to
the author(s) and the title
of the work, journal
citation and DOI.



PAPER

Self-adaptive physics-informed quantum machine learning for solving differential equations

Abhishek Setty^{1,2,3,*} , Rasul Abdusalamov and Felix Motzoi^{2,3} ¹ Department of Continuum Mechanics, RWTH Aachen University, Aachen 52062, Germany² Forschungszentrum Jülich, Institute of Quantum Control (PGI-8), Jülich D-52425, Germany³ Institute for Theoretical Physics, University of Cologne, Cologne D-50937, Germany

* Author to whom any correspondence should be addressed.

E-mail: a.setty@fz-juelich.de**Keywords:** quantum machine learning, physics-informed neural networks, variational quantum algorithms, differential equations

Abstract

Chebyshev polynomials have shown significant promise as an efficient tool for both classical and quantum neural networks to solve linear and nonlinear differential equations (DEs). In this work, we adapt and generalize this framework in a quantum machine learning setting for a variety of problems, including the 2D Poisson's equation, second-order linear DE, system of DEs, nonlinear Duffing and Riccati equation. In particular, we propose in the quantum setting a modified Self-Adaptive Physics-Informed Neural Network approach, where self-adaptive weights are applied to problems with multi-objective loss functions. We further explore capturing correlations in our loss function using a quantum-correlated measurement, resulting in improved accuracy for initial value problems. We analyse also the use of entangling layers and their impact on the solution accuracy for second-order DEs. The results indicate a promising approach to the near-term evaluation of DEs on quantum devices.

Overview

Differential Equations



PINNs



Quantum PINNs



VQC Training



Results

Why Differential Equations?

Differential Equations: The language of Physical Systems

Physical System	Example
Classical Mechanics	Planetary motion, pendulums
Heat Transfer	Diffusion of heat in materials
Fluid Dynamics	Weather prediction, aerodynamics
Electromagnetism	Wireless communication, optics
Quantum Mechanics	Atoms, molecules, quantum devices

Understanding and predicting physical phenomena often reduces to solving differential equations.

How Do We Solve Differential Equations?

Given a differential equation, how can we obtain its solution?

Analytical Solutions

Idea: Find a closed-form expression satisfying the equation.

$$\frac{dy}{dx} = ky$$

$$y(x) = Ce^{kx}$$

Advantages

Exact solution

Limitations

Available only for a limited class of problems

Numerical Methods

Idea: Discretize the domain and approximate derivatives.

Examples:

- Finite Difference Method (FDM)
- Finite Element Method (FEM)
- Spectral Methods

Advantages

Applicable to complex systems

Limitations

Computational cost grows rapidly with dimension

Can Machine Learning Help?

Instead of solving the differential equation directly, can we learn the solution?

Differential Equation



Neural Network

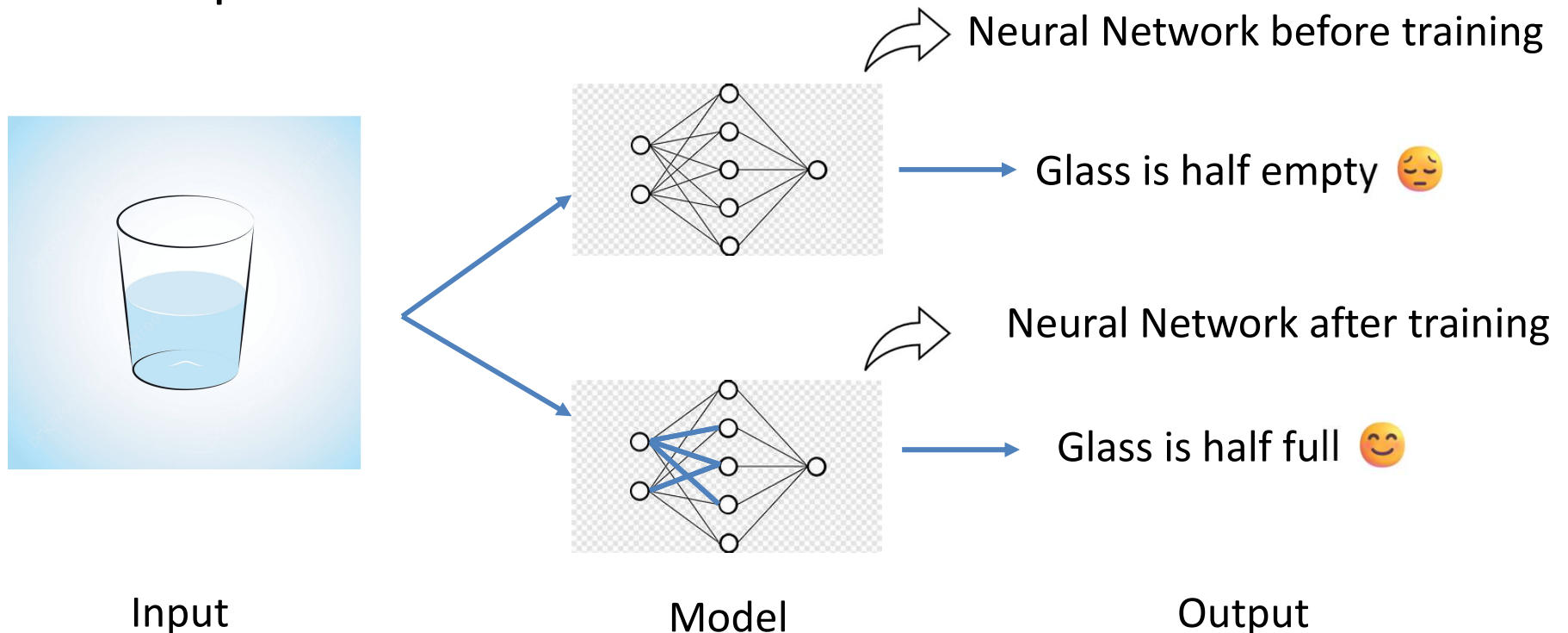


Approximate Solution

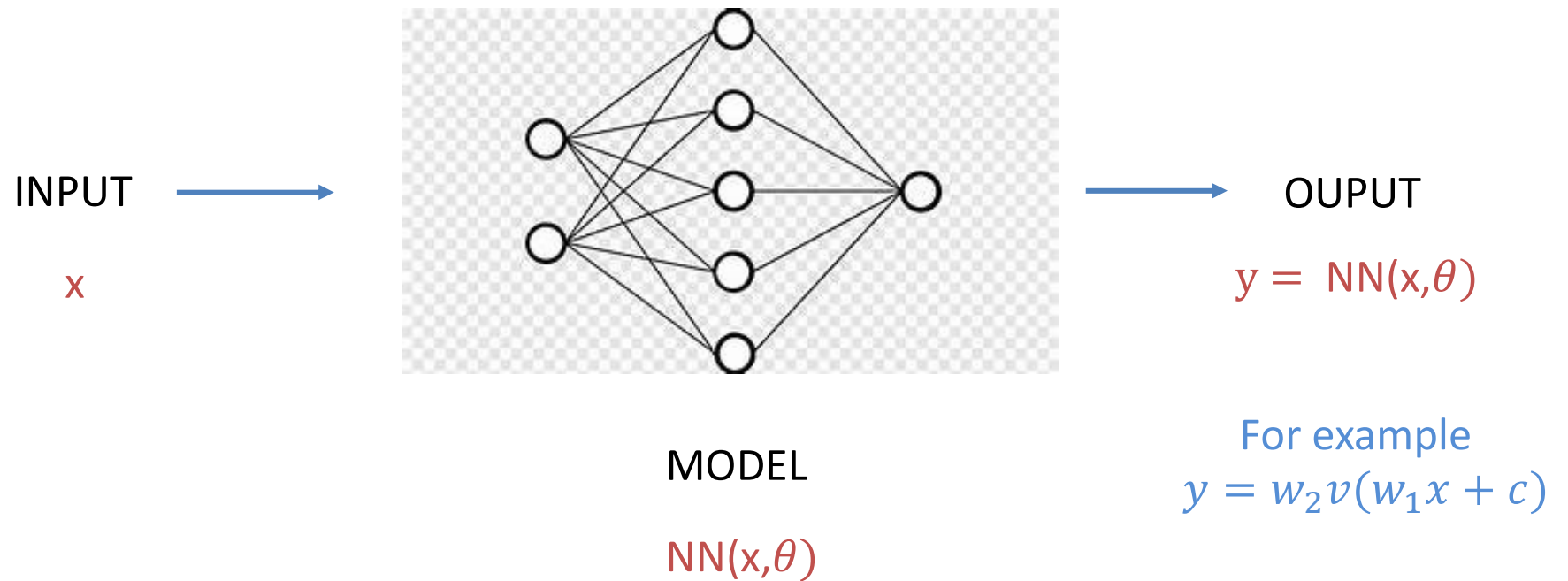
Physics-Informed Neural Network

Neural Network as a Function Approximator

- The Therapist Network



Takeaway: Neural Network is just a trainable function

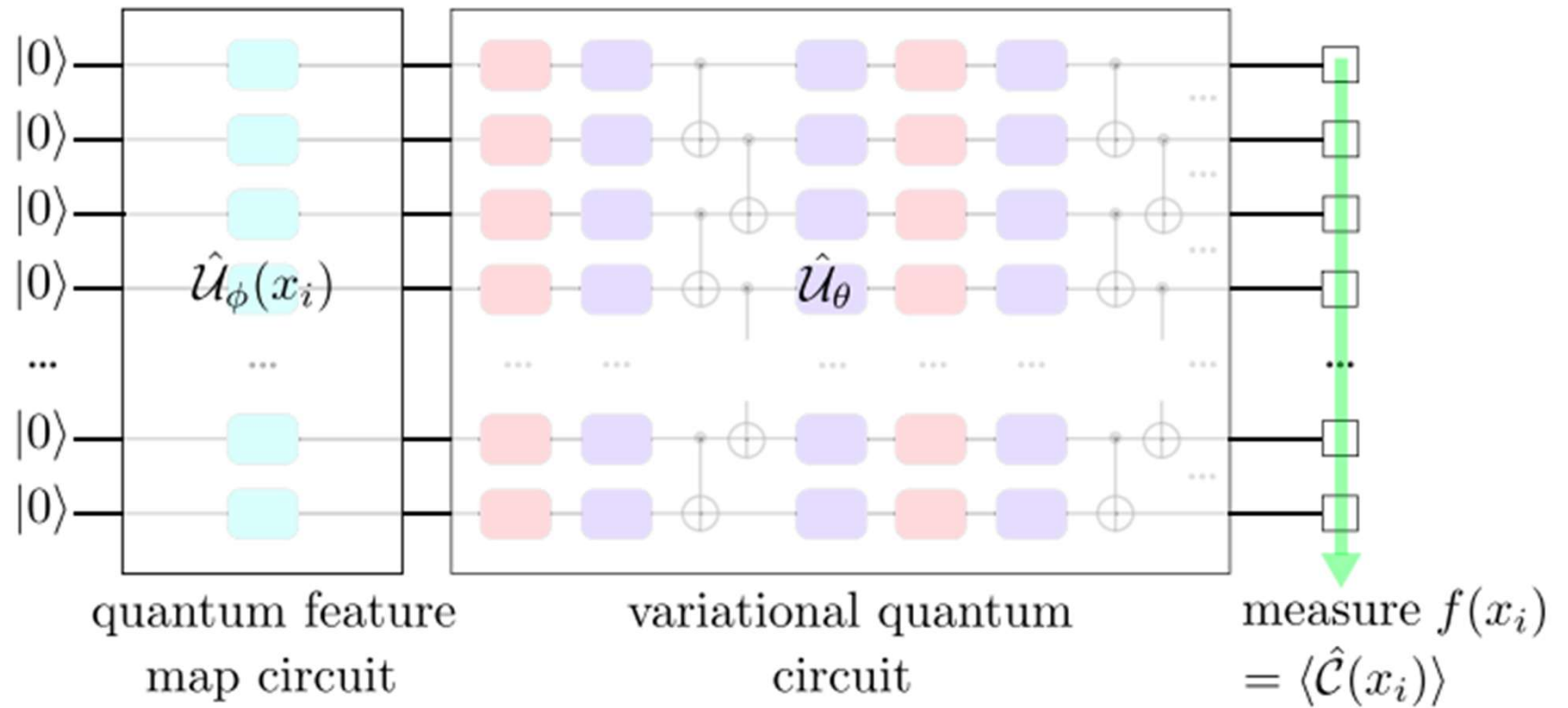


Trained using:

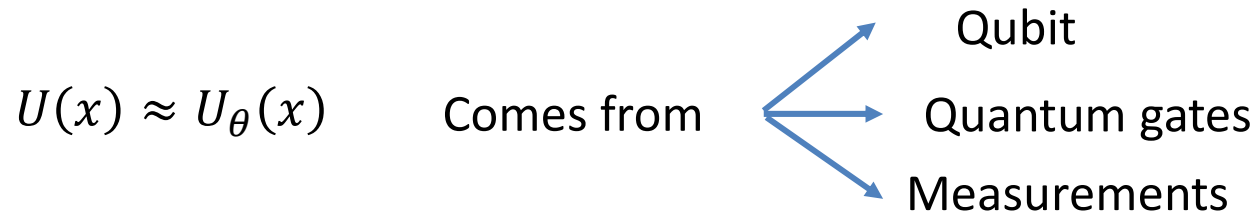
- Gradient descent approach (stochastic) with an appropriate loss function

Can we have a quantum analog of
PINNs?

Quantum Circuit for solving DEs



VQC uses a quantum circuit to represent the function solution, so the Quantum circuit acts as a quantum function approximator.



Step 1 : Encode the Input

DE variable : x , Needs to be encoded in the quantum circuit.

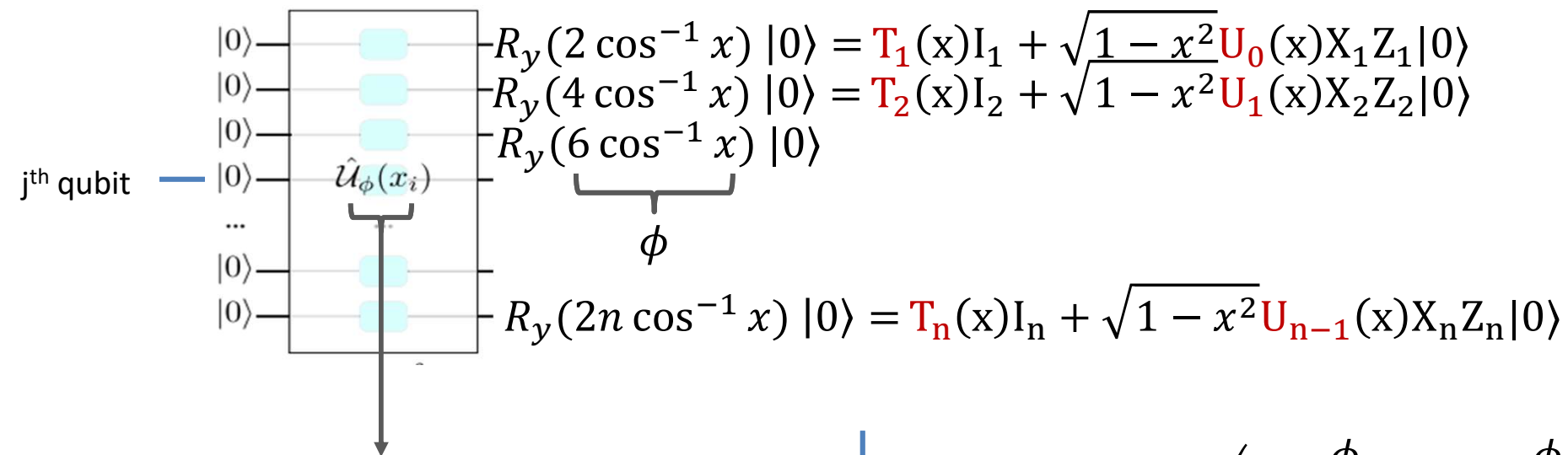
Chebyshev Feature Map encoding

Motivation: Are powerful basis functions for approximating solutions of differential Equations.

Many functions can be approximated efficiently as

$$f(x) \approx \sum_k c_k T_k(x)$$

This is analogous to Fourier expansions.



$$R_y(2 \cos^{-1} x) |0\rangle = \mathbf{T}_1(x)I_1 + \sqrt{1-x^2}\mathbf{U}_0(x)X_1Z_1|0\rangle$$

$$R_y(4 \cos^{-1} x) |0\rangle = \mathbf{T}_2(x)I_2 + \sqrt{1-x^2}\mathbf{U}_1(x)X_2Z_2|0\rangle$$

$$R_y(6 \cos^{-1} x) |0\rangle$$

$$R_y(2n \cos^{-1} x) |0\rangle = \mathbf{T}_n(x)I_n + \sqrt{1-x^2}\mathbf{U}_{n-1}(x)X_nZ_n|0\rangle$$

$$\hat{U}_{\phi(x)} = \bigotimes_{j=1}^N R_{y,j}(2n[j] \cos^{-1} x)$$

Rotation of the j^{th} qubit about the y-axis by angle $2n[j] \cos^{-1} x$.

We now have information stored (via qubits) in basis of Chebyshev's polynomials. These are the building blocks.

$$R_y = \begin{pmatrix} \cos(\frac{\phi}{2}) & -\sin(\frac{\phi}{2}) \\ \sin(\frac{\phi}{2}) & \cos(\frac{\phi}{2}) \end{pmatrix}$$

$$T_n(x) = \cos(ncos^{-1}(x))$$

$$\sin\theta U_n(\cos\theta) = \sin(n+1)\theta$$

$$T_0(x) = 1,$$

$$T_1(x) = x,$$

$$T_2(x) = 2x^2 - 1$$

$$T_3(x) = 4x^3 - 3x$$

$$U_0(x) = 1$$

$$U_1(x) = 2x$$

$$U_2(x) = 4x^2 - 1$$

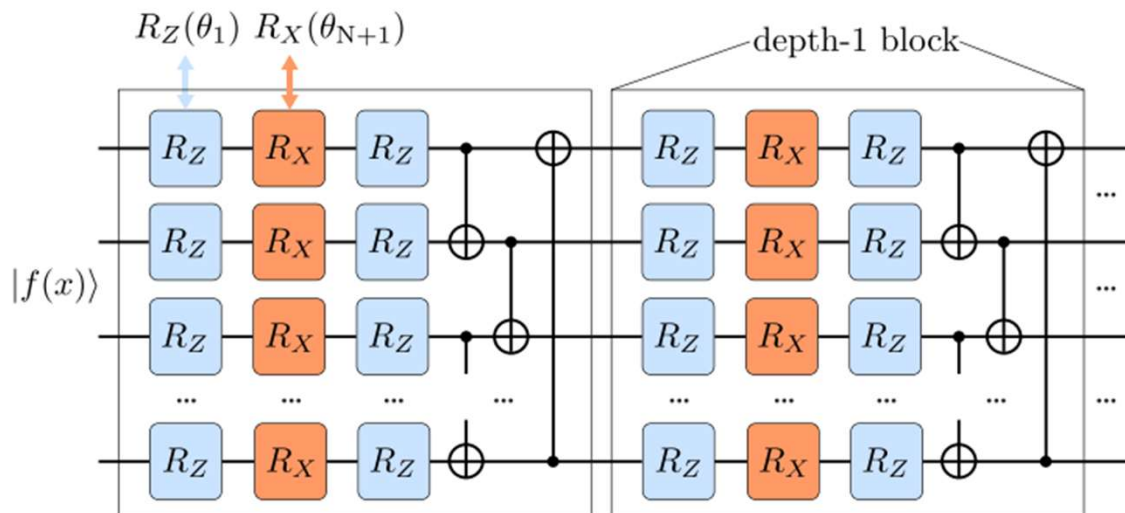
$$U_3(x) = 8x^3 - 4x$$

We Have the Features... Now What?

Now we need a recipe that is :

- Expressive enough
- Trainable
- Hardware-friendly

These requirements motivate the use of a Hardware-Efficient Ansatz.



Now imagine the encoded state is a collection of ingredients.
The job is to repeatedly do two things:

Step 1: Mix each ingredient

This is the rotation layer.

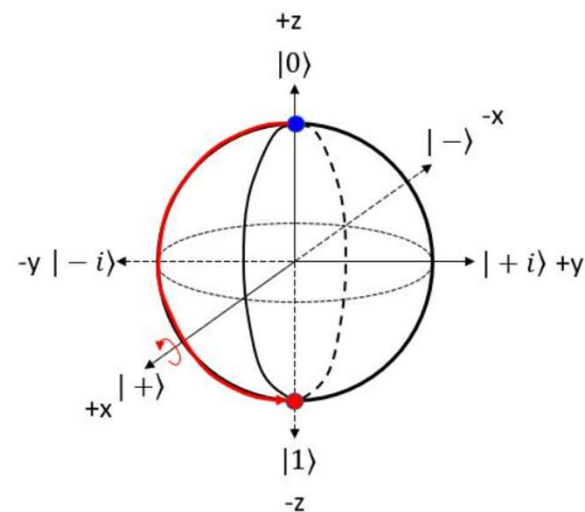
$$R_Z(\theta_1)R_X(\theta_2)R_Z(\theta_3)$$

on every qubit.

Provides Maximum local flexibility for each qubit

Interpretation:

Adjust the contribution of each encoded feature independently.

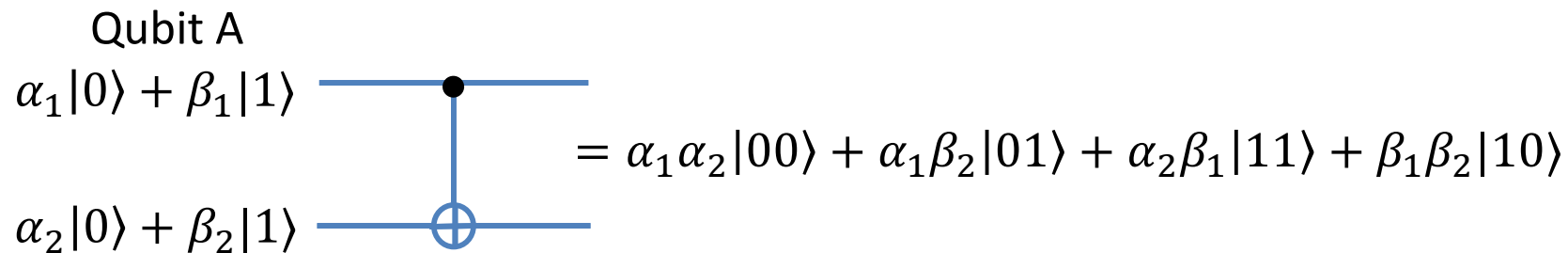


Step 2: Allow ingredients to interact

This is the CNOT layer.

Interpretation:

Create correlations between features.



Mixture of the amplitudes. Now more complicated structures can emerge.

Now repeat the block (rotation+ entanglement) d times, for even richer expressivity.

We have a quantum state prepared ... Now what?

Now is the time for extraction of it's properties

Measurements

Difference between a quantum and classical algorithm:

In Classical algorithm,

input $\rightarrow$ computation $\rightarrow$ **output** = a number, a vector , direct approximation of the answer etc.

In Quantum algorithm,

input $\rightarrow$ parameterized circuit $\rightarrow$ $|\psi(\theta)\rangle \rightarrow$ measurement $\rightarrow$ samples $\rightarrow$ classical estimate.

But what since we know the whole circuit and gates and everything used, so in principle we can just keep following it and end up with the amplitude form of the final state, isn't it?

Umm... not really!

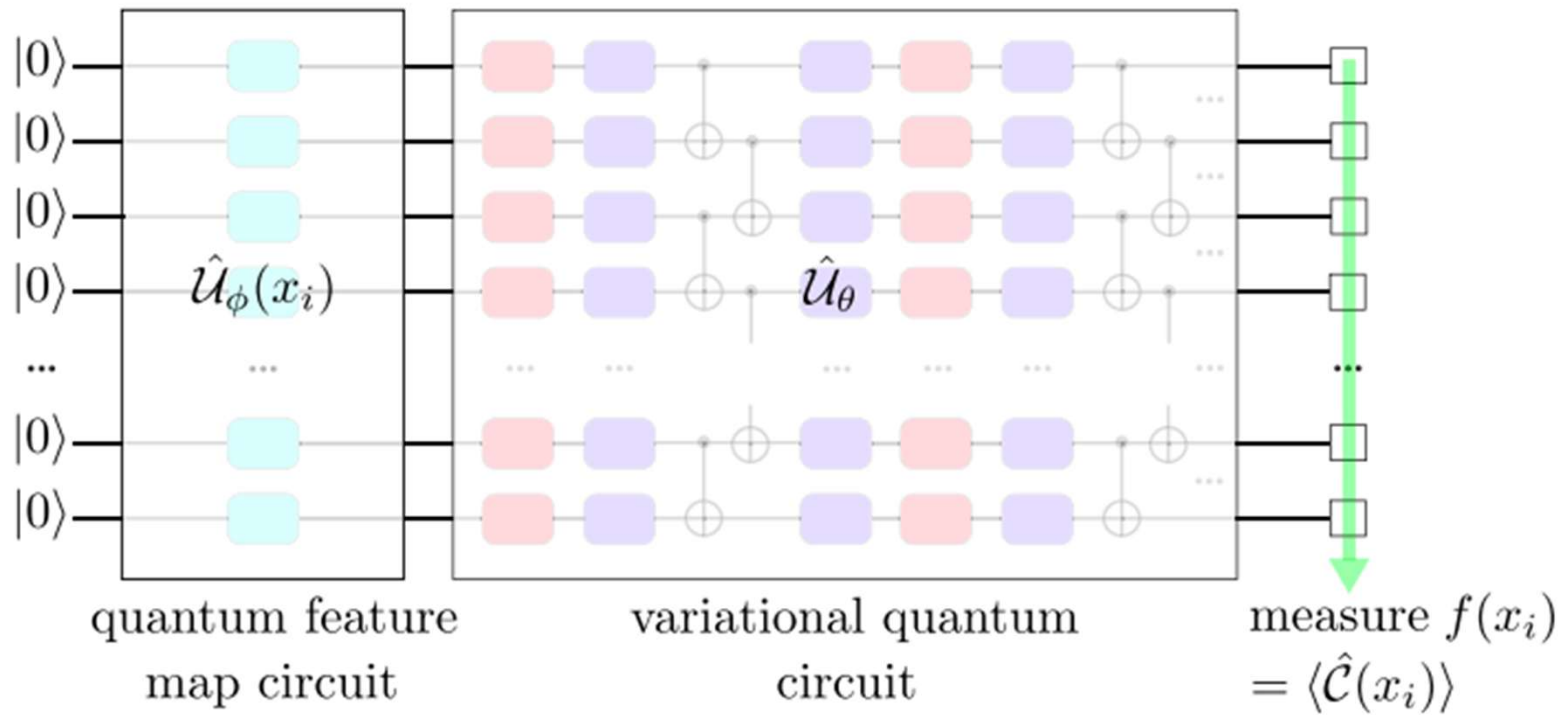
This is where the problem of simulating a quantum system on a classical Machine comes Into the picture

For n qubits, the state vector has 2^n complex amplitudes:

- so, for $n=10$: 1024 amplitudes (easy)
- For $n=20$: 1 million amplitudes (ok-ish)
- For $n=50$: 1 quadrillion amplitudes (impossible)
- For $n=100$: astronomically impossible !

Classical difficulty arises not only from the size of the quantum state, but also from the structure of entanglement and interference, which makes even sampling and expectation estimation computationally challenging

Recap



Encodes Chebyshev's features
with input x

Tunes and mixes the features

Output = $f(x)$

Optimization of f(x)

Loss function

To simulate the system, the quantum circuit is trained using the following loss function:

$$L_{\theta}[d_x f, f, x] = L_{\theta}^{(diff)}[d_x f, f, x] + L_{\theta}^{(boundary)}[f, x]$$

$$L_{\theta}^{(diff)}[d_x f, f, x] = \frac{1}{M} \sum_{i=1 \text{ to } M} L(F[d_x f(x_i), f(x_i), x_i], 0)$$

$$L_{\theta}^{(boundary)}[f, x] = \eta L(f(x_0), u_0)$$

$L(a, b)$ = function to describe the distance between a and b .

$$L(a, b) = (a - b)^2$$

$$L(a, b) = |a - b|$$

For example: Harmonic Oscillator:

Consider:

$$\frac{d^2u}{dx^2} + u(x) = 0$$

With boundary conditions,

$$u(0) = 0, \quad u\left(\frac{\pi}{2}\right) = 1$$

Then loss function would be:

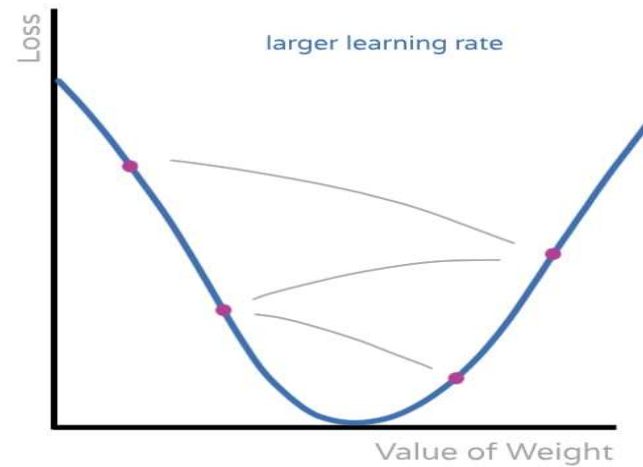
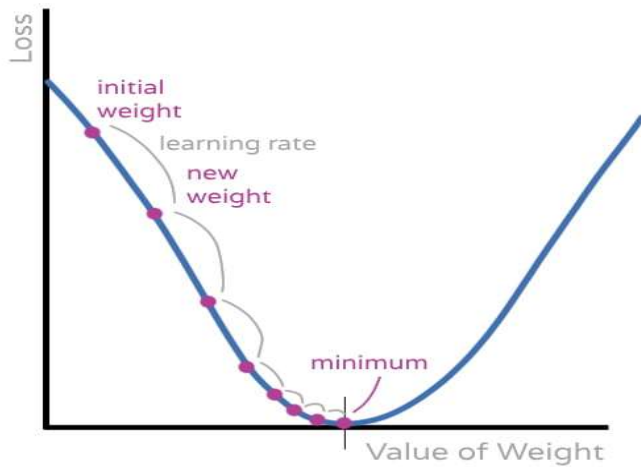
$$L = \frac{1}{N_f} \sum_i \left(\frac{d^2 u_{pred}}{dx^2} + u_{pred}(x) \right)^2 + (u_{pred}(0) - 0)^2 + \left(u_{pred}\left(\frac{\pi}{2}\right) - 1 \right)^2$$

The parameter update rule for gradient descent approach :

$$\theta_i^{(t+1)} = \theta_i^{(t)} - \eta \frac{\partial L}{\partial \theta_i}$$

η : learning rate

$\frac{\partial L}{\partial \theta_i}$: gradient of the loss with respect to parameter θ_i



Now how do we compute the derivatives of the classical function $f(x)$?

$$\frac{df}{dx} = \sum_j \frac{\partial U_\phi}{\partial \phi_j} \frac{d\phi_j}{dx}$$

As the function f is not given analytically. It is produced through quantum measurements, Therefore we need to differentiate the encoder circuit $\hat{U}_\phi$.

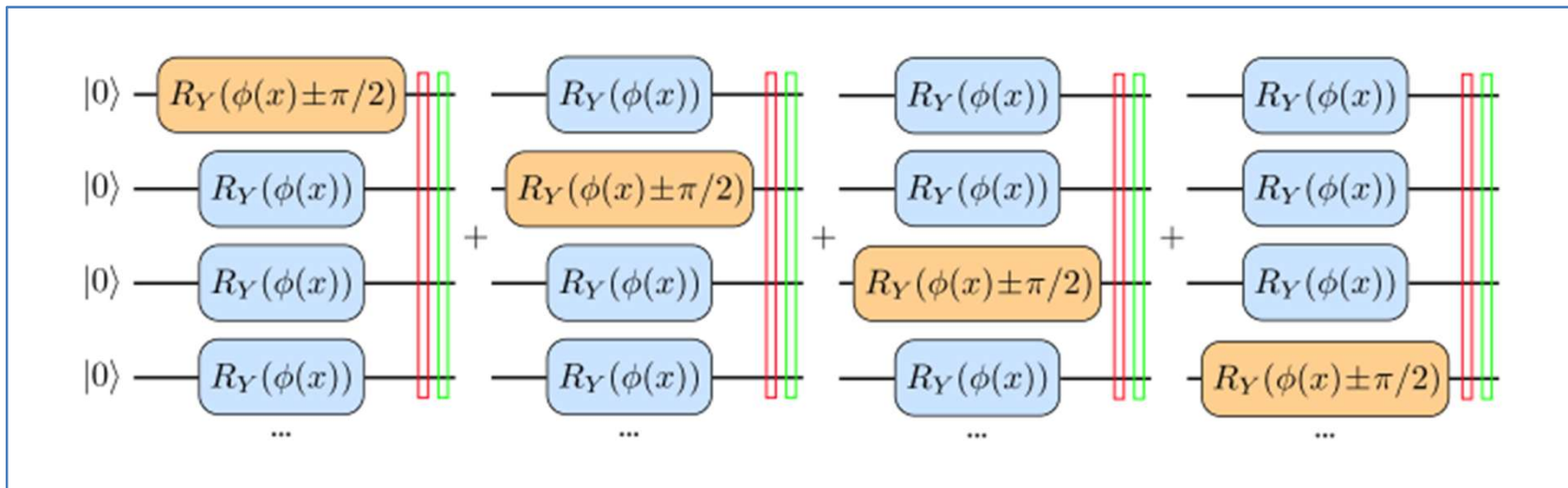
Parameter shift rule for the derivative of the encoder circuit:

For string of Pauli operators or involutory operators in the encoder circuit,
Then the derivative are written as,

$$\frac{\partial U_\phi}{\partial \phi_j} = \langle \psi_{d\phi,j,\theta}^+(x) | \hat{C} | \psi_{d\phi,j,\theta}^+(x) \rangle - \langle \psi_{d\phi,j,\theta}^-(x) | \hat{C} | \psi_{d\phi,j,\theta}^-(x) \rangle$$

State prepared with a phase shift in
 j^{th} qubit angle $\phi(x)$ to $\phi(x) + \frac{\pi}{2}$

State prepared with a phase shift in j^{th} qubit
angle $\phi(x)$ to $\phi(x) - \frac{\pi}{2}$



$$\frac{d^2 f}{dx^2} = \frac{d}{dx} \left(\frac{df}{dx} \right) = \frac{d}{dx} \left(\sum_j \frac{dU_{\phi_j}}{d\phi_j} \frac{d\phi_j}{dx} \right)$$

$$= \sum_{j,k} \frac{d^2 U_{\phi}}{d\phi_j d\phi_k} \left[\frac{4jk}{1-x^2} \right] + \sum_j \frac{dU_{\phi}}{d\phi_j} \left[-\frac{2jx}{(1-x^2)^{\frac{1}{2}}} \right]$$

Four shifted terms per qubit

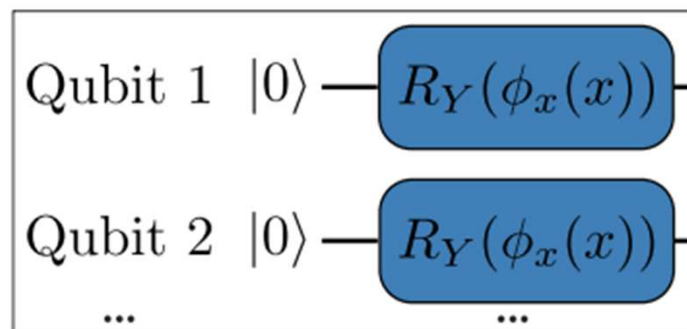
Two shifted terms per qubit

What about the PDE? How can we encode other independent variables?

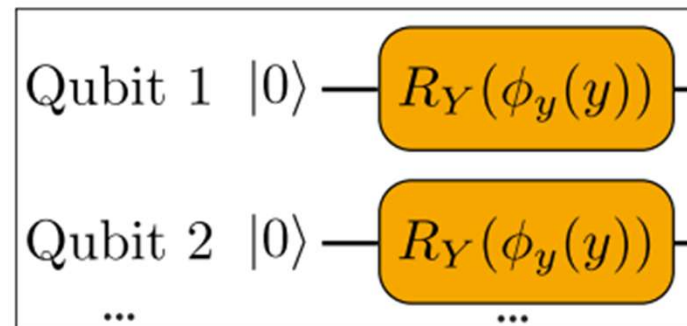
We define two unitary operators $\hat{U}_{\phi_x}(x)$ and $\hat{U}_{\phi_y}(y)$.

$$\hat{U}_{\phi(x,y)}(x,y) = \hat{U}_{\phi_x}(x) \otimes \hat{U}_{\phi_y}(y)$$

Feature map
circuit for
variable x



Feature map
circuit for
variable y



Variate

Measure

The calculation of the first and second-order derivatives of quantum circuits concerning all qubits are described by equations,

$$\begin{bmatrix} \frac{\partial u}{\partial x} \\ \frac{\partial u}{\partial y} \end{bmatrix} = \begin{bmatrix} \frac{\partial u_{\text{pred}}}{\partial \phi_{xj}(x_f^i)} \\ \frac{\partial u_{\text{pred}}}{\partial \phi_{ym}(y_f^i)} \end{bmatrix},$$

$$\begin{bmatrix} \frac{\partial^2 u}{\partial x^2} & \frac{\partial^2 u}{\partial x \partial y} \\ \frac{\partial^2 u}{\partial y \partial x} & \frac{\partial^2 u}{\partial y^2} \end{bmatrix} = \begin{bmatrix} \frac{\partial^2 u_{\text{pred}}}{\partial \phi_{xj}(x_f^i) \partial \phi_{xk}(x_f^i)} & \frac{\partial^2 u_{\text{pred}}}{\partial \phi_{xj}(x_f^i) \partial \phi_{yn}(y_f^i)} \\ \frac{\partial^2 u_{\text{pred}}}{\partial \phi_{ym}(y_f^i) \partial \phi_{xk}(x_f^i)} & \frac{\partial^2 u_{\text{pred}}}{\partial \phi_{ym}(y_f^i) \partial \phi_{yn}(y_f^i)} \end{bmatrix}$$

Self-adaptive weights for multi-objective loss function

$$\mathcal{L}(\theta, \lambda_f, \lambda_b) = \mathcal{L}_f(\theta, \lambda_f) + \mathcal{L}_b(\theta, \lambda_b)$$

$$\mathcal{L}_f(\theta, \lambda_f) = \frac{1}{N_f} \sum_{i=1}^{N_f} m(\lambda_f^i) \left(f(x_f^i; \theta) \right)^2,$$

$$\mathcal{L}_b(\theta, \lambda_b) = \frac{1}{N_b} \sum_{i=1}^{N_b} m(\lambda_b^i) \left(u_{\text{pred}}(x_b^i; \theta) - u_b(x_b^i) \right)^2$$



Mask functions

Strictly increasing
function of λ .

$$\theta^{k+1} = \theta^k - \eta^k \nabla_{\theta} \mathcal{L}(\theta^k, \lambda_f^k, \lambda_b^k),$$

$$\lambda_f^{k+1} = \lambda_f^k + \rho_f^k \nabla_{\lambda_f} \mathcal{L}(\theta^k, \lambda_f^k, \lambda_b^k),$$

$$\lambda_b^{k+1} = \lambda_b^k + \rho_b^k \nabla_{\lambda_b} \mathcal{L}(\theta^k, \lambda_f^k, \lambda_b^k),$$

Parameter λ maximizes whereas
Parameter θ minimizes.

Results

1. Riccati equation:

Why Study the Riccati Equation?

- Appears in optimal control and state estimation.
- Used in robotics, aerospace, autonomous systems, and quantum control.
- Provides a nonlinear model for many physical processes.

$$\frac{du}{dx} - 4u + 6u^2 - \sin(50x) - u \cos(25x) + \frac{1}{2} = 0,$$
$$u(x=0) = u_0$$

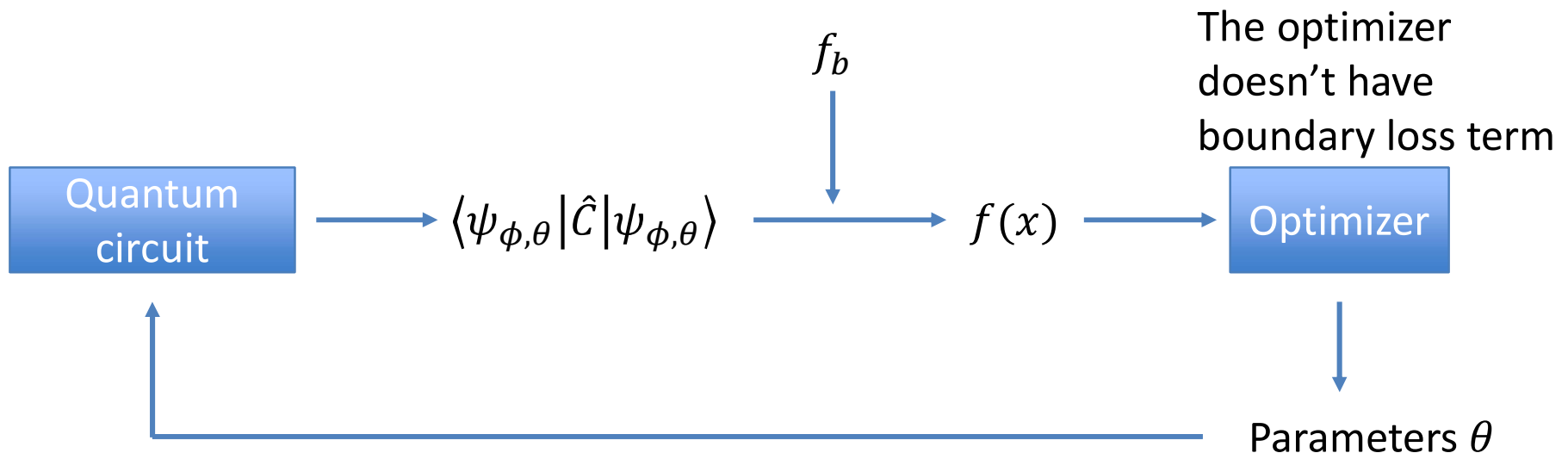
Floating boundary handling

In multiple loss function, trying to optimize one loss function will affect the Performance of the other.

But, if DE has $f(x)$ term in it then there is a possibility to solve the problem In a more efficient way.

$$f(x) = f_b + \langle \psi_{\phi,\theta} | \hat{C} | \psi_{\phi,\theta} \rangle$$

$$f_b = u_0 - \langle \psi_{\theta,\phi}(x_0) | \hat{C} | \psi_{\theta,\phi}(x_0) \rangle$$



Earlier Loss term is given by (here u means f):

$$L = \left(\sum_x \frac{du}{dx} - 4u + 6u^2 - \sin(50x) - u \cos(25x) \right) + \frac{1}{2} + (u(x=0) - u_0)$$

Comes from the differential eqn

Boundary condition

With modified $u(x)$:

$$u_m(x) = (u_0 - u_{pred}(x_0)) + u_{pred}(x)$$

The loss term becomes:

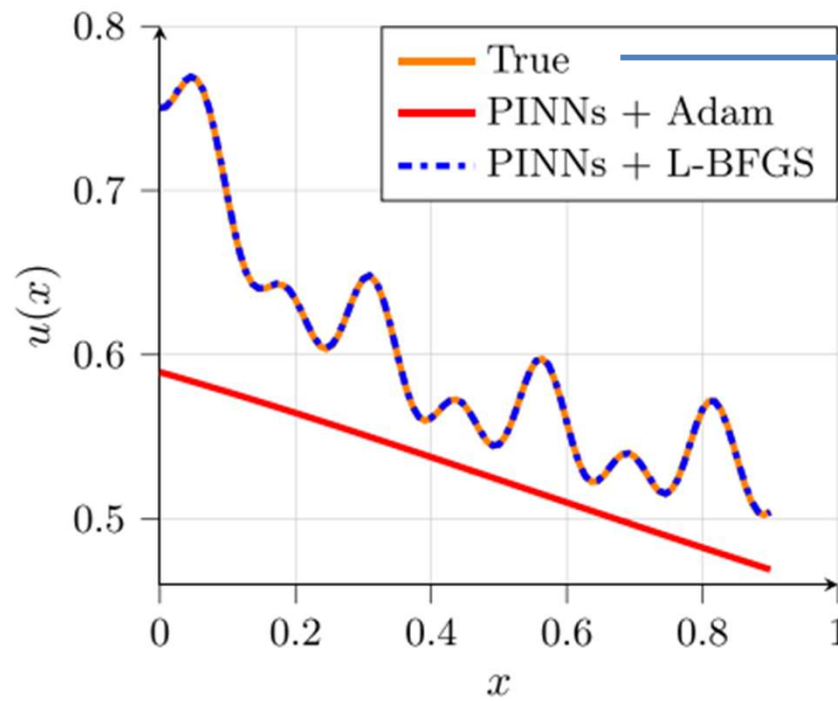
$$L = \sum_x \frac{du_m}{dx} - 4u_m + 6u_m^2 - \sin(50x) - u_m \cos(25x) + \frac{1}{2}$$

Addition:

$$\hat{C} = \sum_j Z_j$$

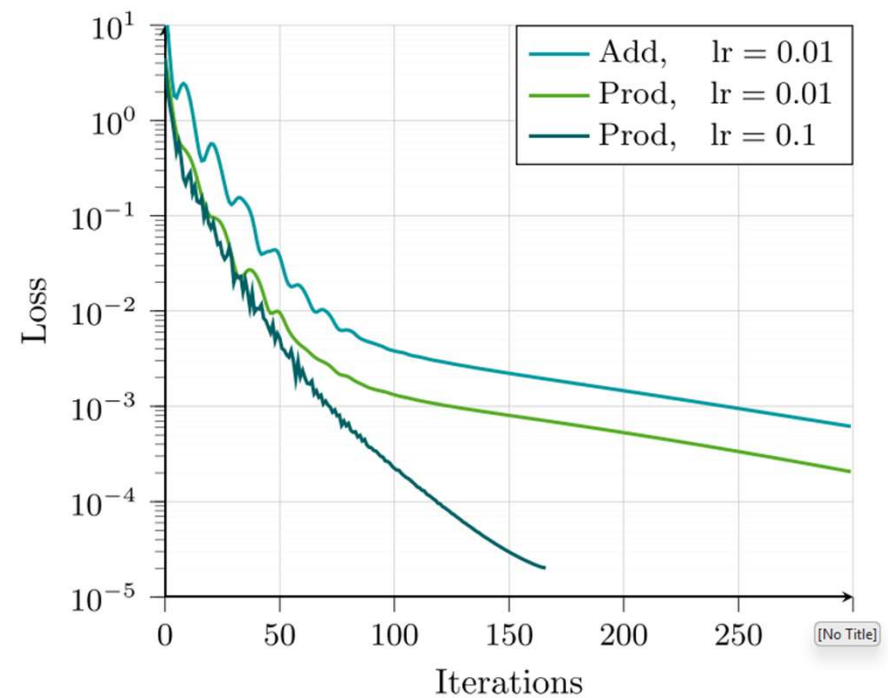
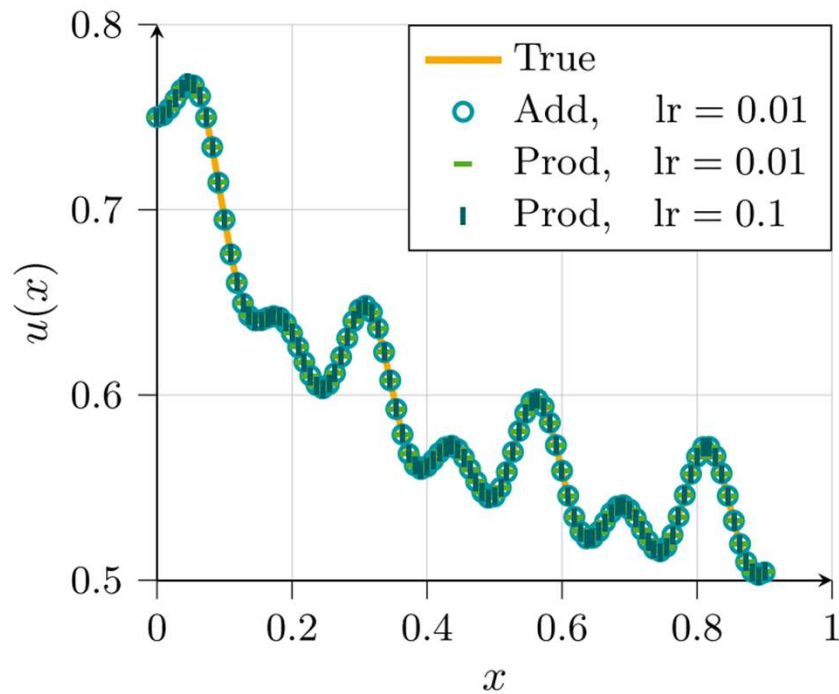
Product:

$$\hat{C} = \bigotimes_{j=1}^N Z_j$$



From Runge-kutta third order

8 Qubits and 8 block depth



Key takeaways:

- Classical ML PINNs requires high precision update direction by L-BFGS optimizer
- Choice of observable as the tensor product outperformed the addition choice.

Why so?

Because it can show correlations between the qubits thus expressing Entanglement.

as an example, say for a bell state: $|\phi\rangle = \frac{|00\rangle}{\sqrt{2}} + \frac{|11\rangle}{\sqrt{2}}$

$$\langle Z_1 \rangle + \langle Z_2 \rangle = 0 + 0 = 0$$

but

$$\langle Z_1 Z_2 \rangle = 1$$

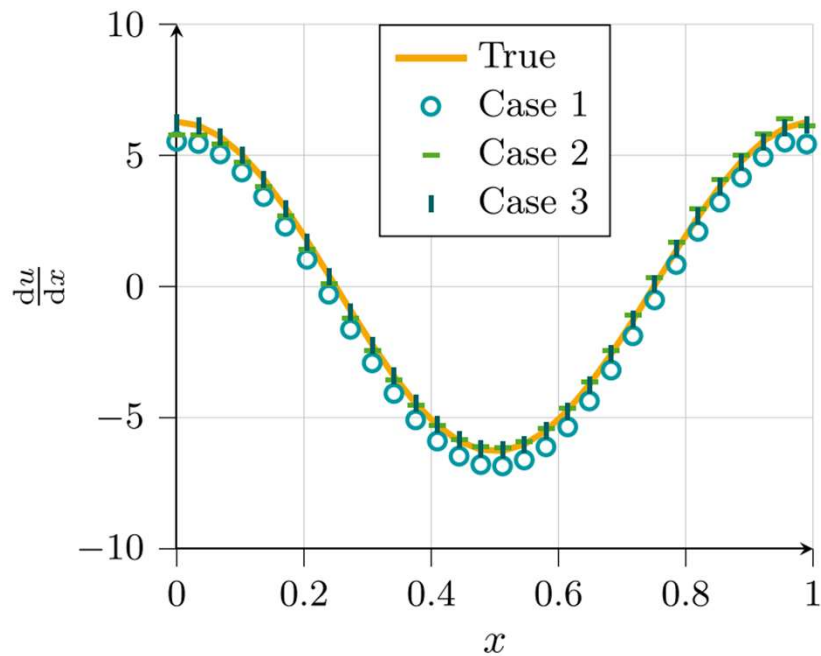
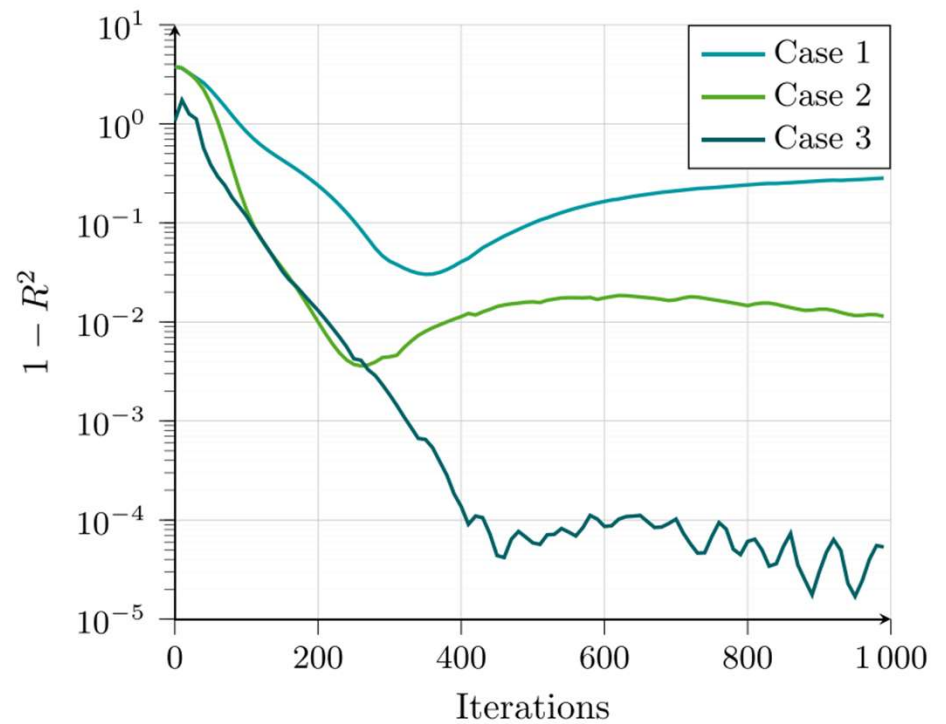
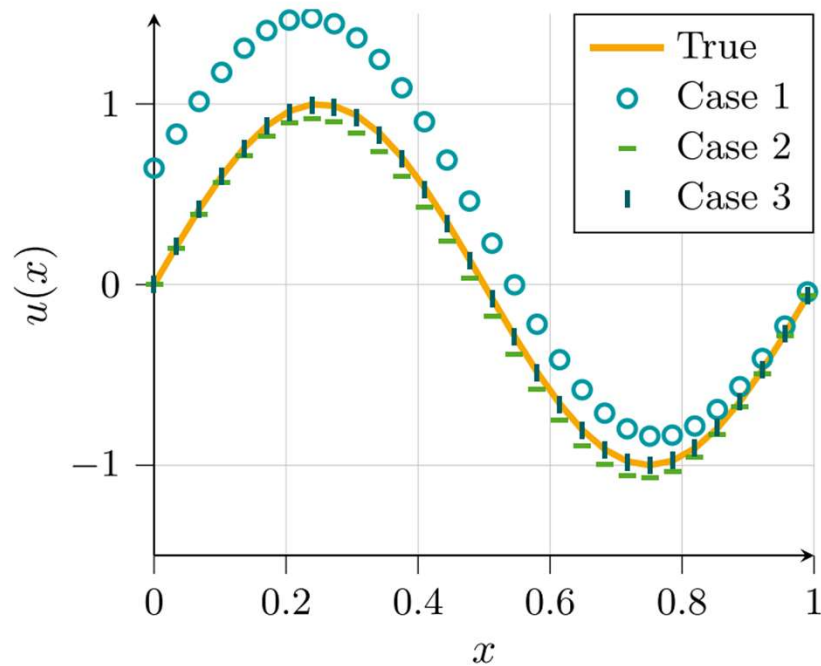
Meaning qubits are perfectly correlated

2. Second-order linear DE:

$$\frac{d^2u}{dx^2} + 4\pi^2 \sin(2\pi x) = 0$$

Loss term (here no floating boundary handling):

$$\begin{aligned} Loss = & \frac{\alpha_f}{N_f} \left[\frac{d^2u}{dx^2} + 4\pi^2 \sin(2\pi x) \right] \\ & + \alpha_b [(u(x=0) - 0)^2 + (u(x=0.99) - \sin 2\pi 0.99)^2] \end{aligned}$$



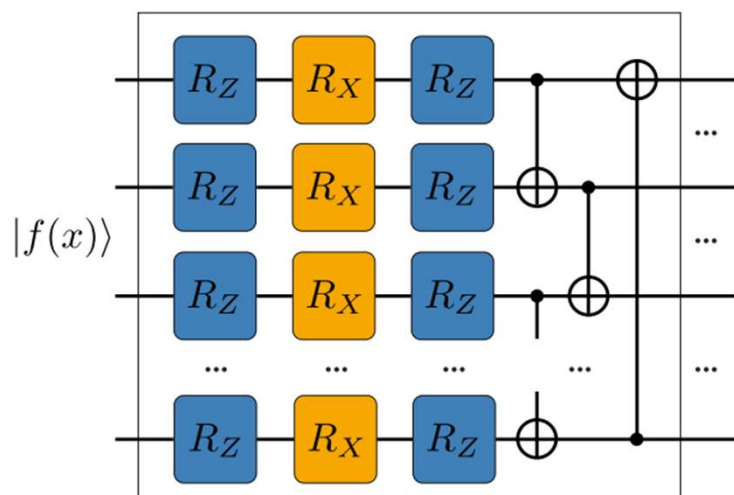
Case 1:

balanced weights of $\alpha_f = \alpha_b = 1$

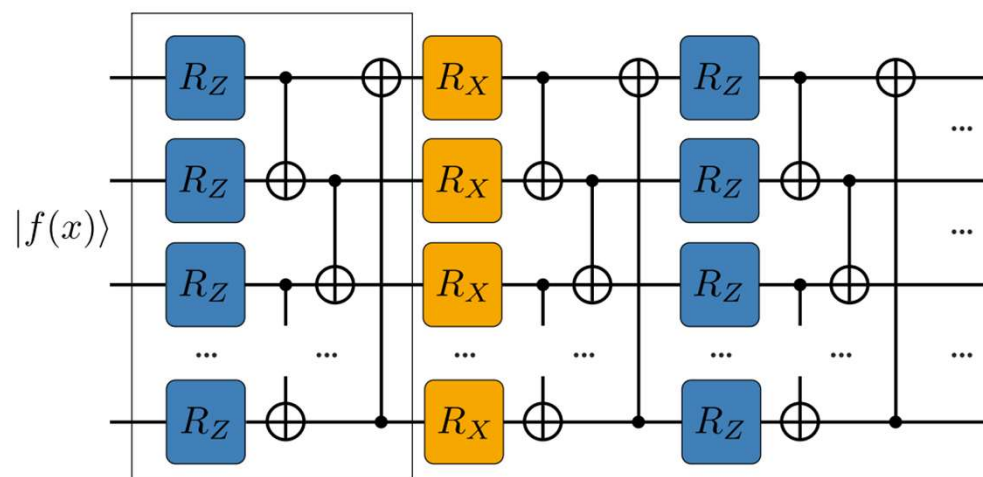
Case 2:

balanced weights of $\alpha_f = 10^{-1}$ and $\alpha_b = 10^3$
(No Self-adaptive yet!)

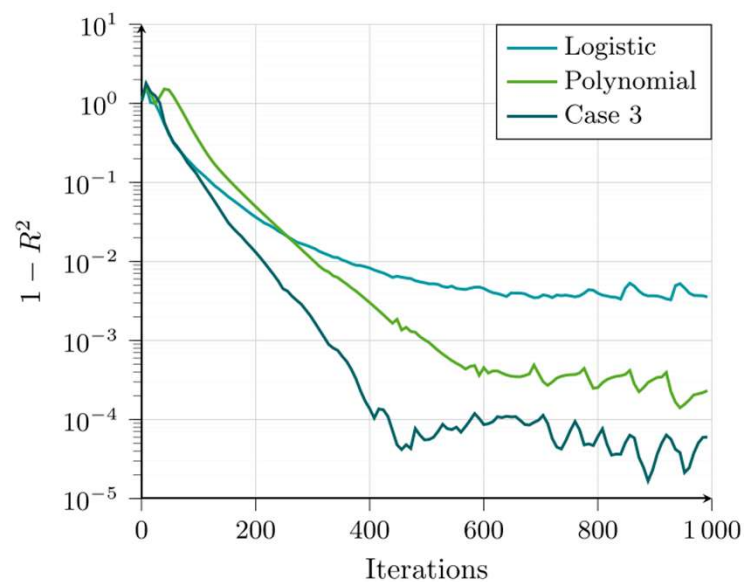
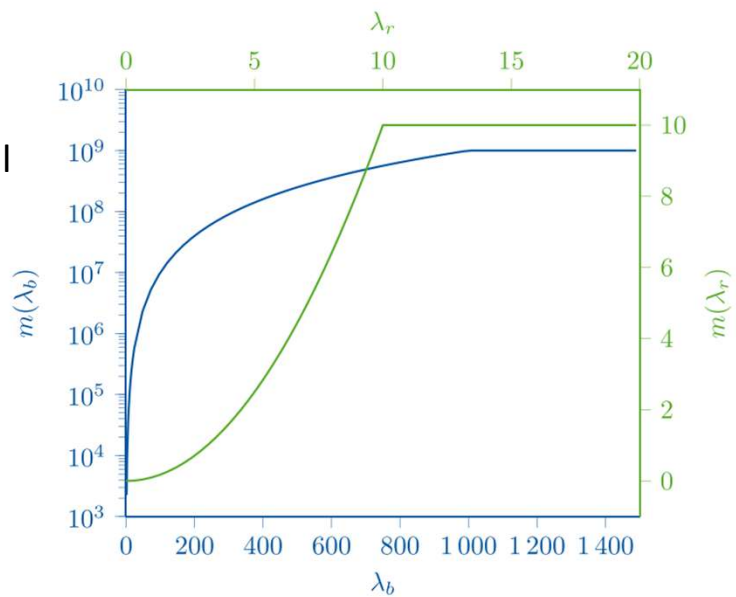
Case 1 and 2



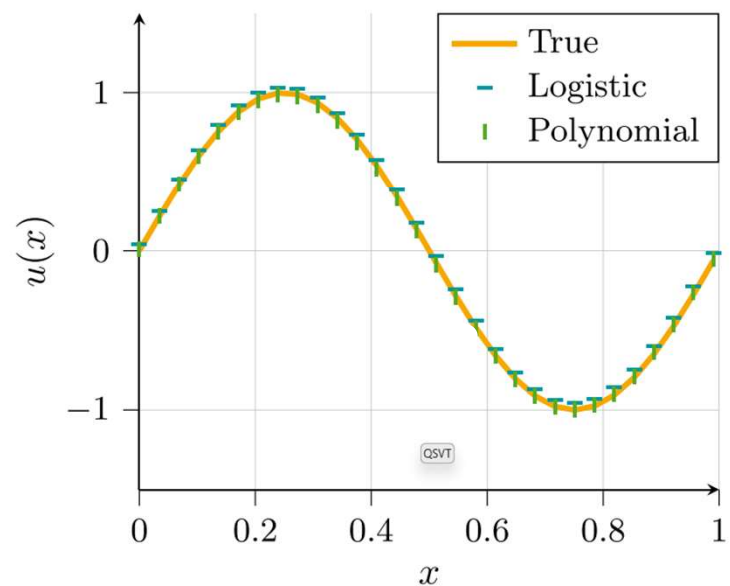
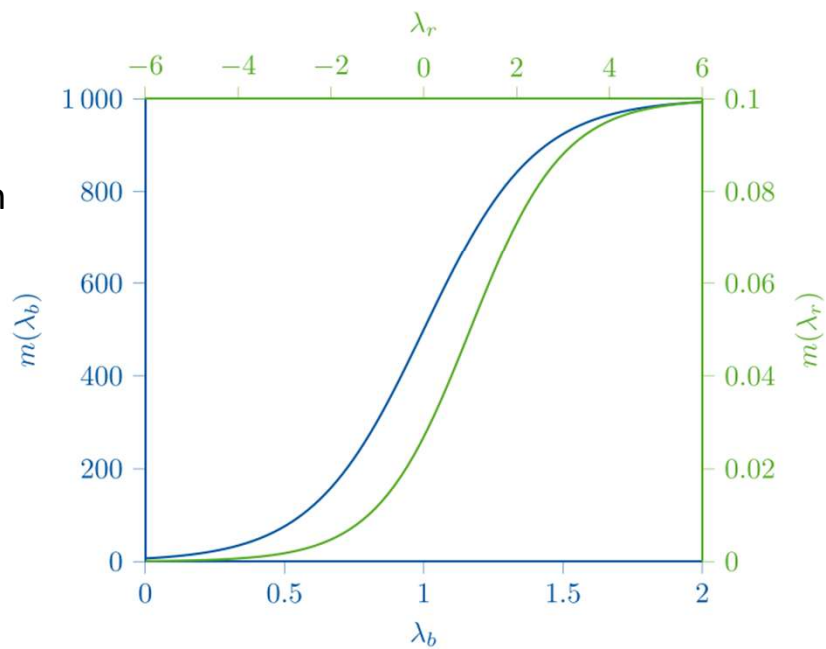
Case 3



Polynomial
Mask fn



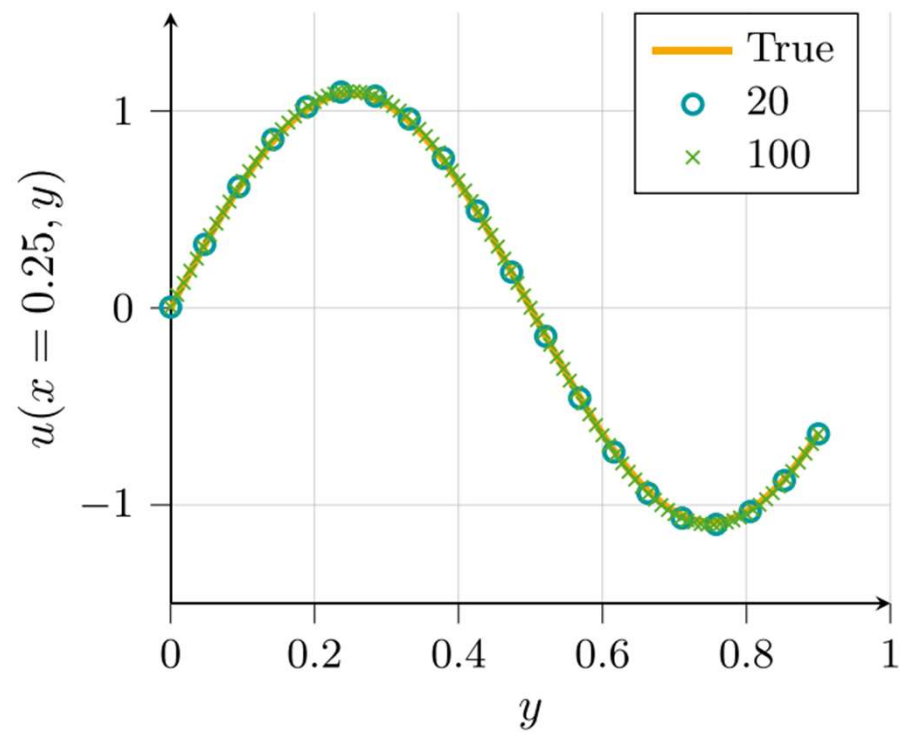
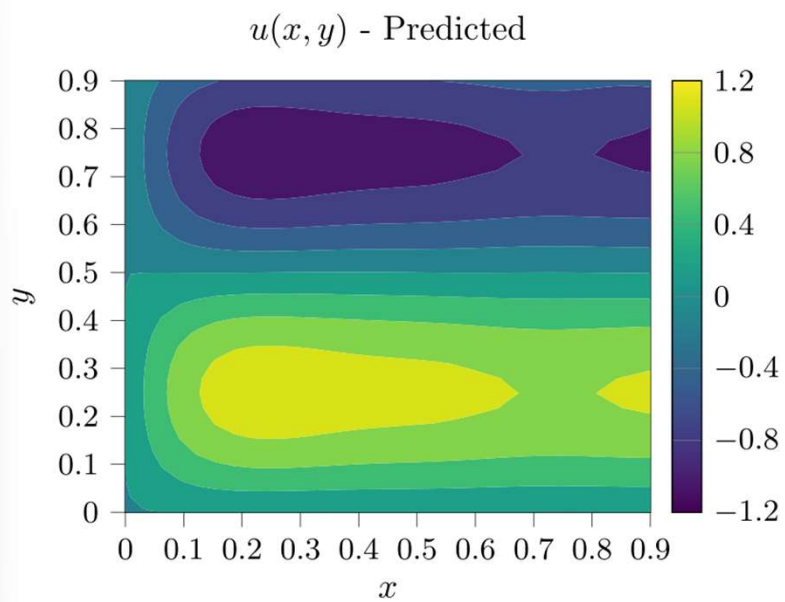
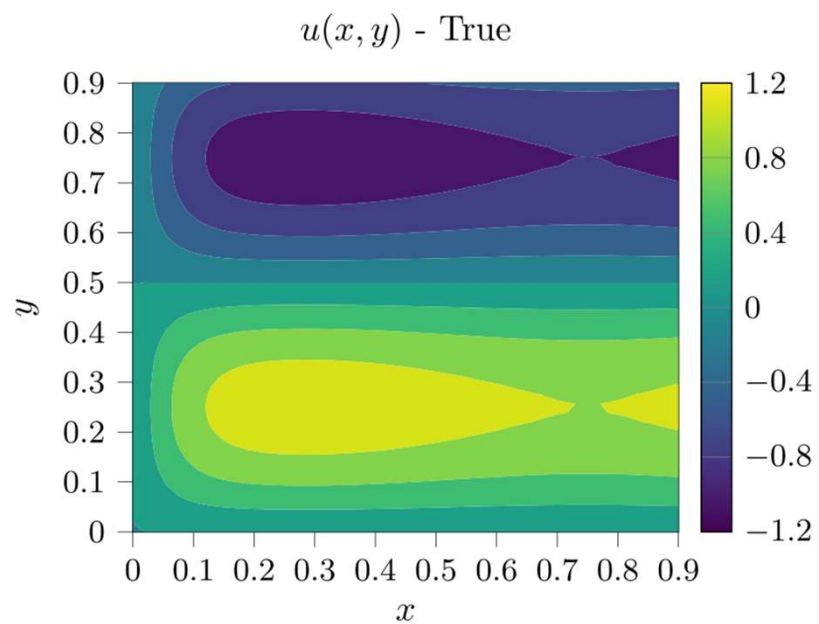
Logistic
Mask fn

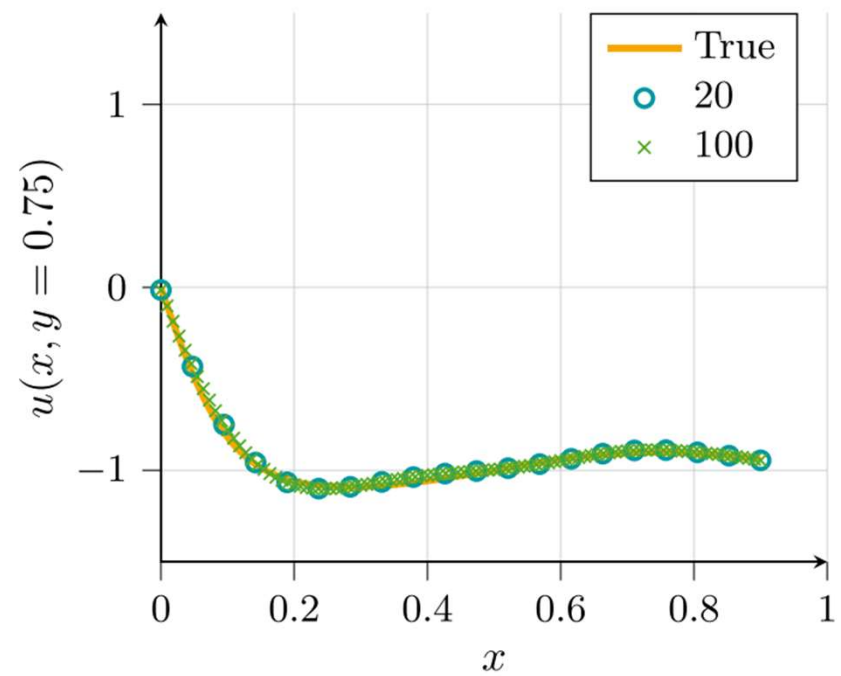
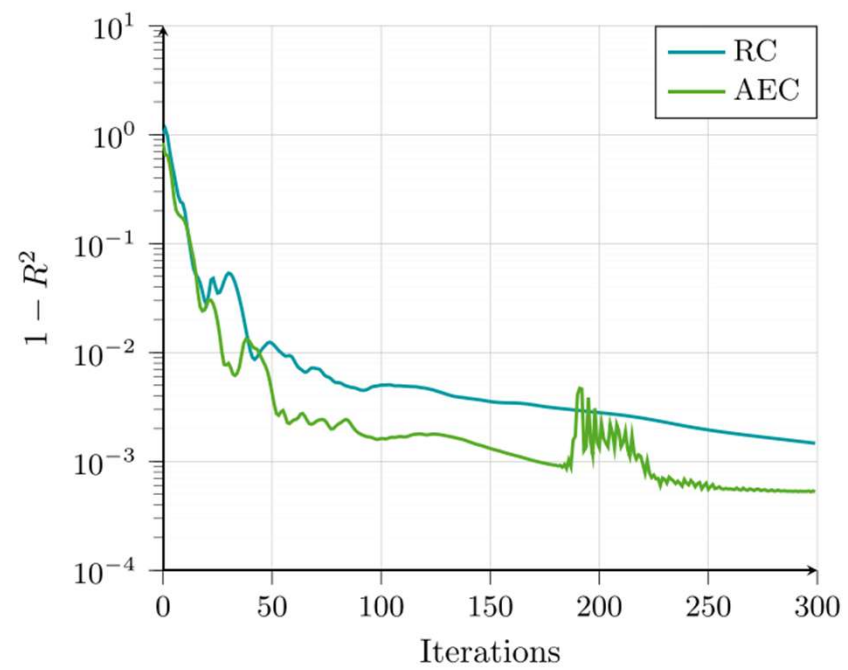


3. Second-order PDE:

$$\Delta u(x, y) = g(x, y) \quad \text{with } x, y \in [0, 0.9] \times [0, 0.9]$$

- 900 collocation points(i.e 30×30 equidistant points)
- Dirichlet boundary conditions given by $u_b(x, 0)$, $u_b(x, 0.9)$, $u_b(0, y)$, and $u_b(0.9, y)$.
- $C = \alpha_f MSE_f + \alpha_b MSE_b$, $\alpha_f = 10^{-1}$; $\alpha_b = 10^3$
- The circuit architecture for this problem is chosen to be 3 qubits for variable x and 3 qubits for variable y with a depth of 7 blocks in the variational circuit.





Conclusions:

Key Takeaways

- Differential equations are solved by minimizing a physics-informed loss function.
- The VQC acts as a quantum ansatz for the solution $u(x)$.
- Function values are extracted through expectation value measurements.
- PDE residuals require derivatives, obtained through the chain rule and parameter-shift rule.
- Training is performed using a hybrid quantum-classical optimization loop.

Outlook

- More expressive ansatz and feature maps (The number of layers d is directly proportional to the circuit's expressive power. However, as the parameters increase, the optimization will suffer with a problem known as barren plateau)
- Improved trainability and noise resilience.
- A further question that still needs to be addressed is whether a potential quantum advantage can be achieved.

Thank you!