
FP64 is All You Need: Rethinking Failure Modes in Physics-Informed Neural Networks

Chenhui Xu Dancheng Liu Amir Nassereldine Jinjun Xiong*
University at Buffalo, SUNY
Buffalo, NY, USA, 14226
{cxu26, jinjun}@buffalo.edu

Abstract

Physics-Informed Neural Networks (PINNs) often exhibit “failure modes” in which the PDE residual loss converges while the solution error stays large, a phenomenon traditionally blamed on local optima separated from the true solution by steep loss barriers. We challenge this understanding by demonstrate that the real culprit is insufficient arithmetic precision: with standard FP32, the L-BFGS optimizer prematurely satisfies its convergence test, freezing the network in a spurious failure phase. Simply upgrading to FP64 rescues optimization, enabling vanilla PINNs to solve PDEs without any failure modes. These results reframe PINN failure modes as precision-induced stalls rather than inescapable local minima and expose a three-stage training dynamic—un-converged, failure, success—whose boundaries shift with numerical precision. Our findings emphasize that rigorous arithmetic precision is the key to dependable PDE solving with neural networks. Our code is available at https://github.com/miniHuiHui/PINN_FP64.

1 Introduction

Physics-Informed Neural Networks (PINNs) [27] have gained wide attention and applications in recent years as a novel numerical solver for partial differential equations (PDEs). PINNs can find a numerical solution by optimizing the residual loss defined by the PDE, leveraging the nature of the universal approximation [16] of neural networks and the automatic differentiation [4] provided by the deep learning framework such PyTorch [26]. Although theoretically it is capable of providing exact numerical solutions for PDEs, the researchers have identified several *failure modes* for PINNs [20]. In these cases presenting a failure mode, the PDE residual loss is optimized to a very small value, but the numerical solution provided by PINN is over trivial and has a huge error with the true solution.

To mitigate these failure modes, various methods based on optimization [29, 28], regional gradient [33, 34], sampling [12, 32, 9], and model architectures [35, 36, 7, 23] have been proposed. These methods are based on a consensual understanding of the PINN’s loss landscape: such failure modes occur when the model becomes trapped in a local optimum located within an extremely sharply descending loss basin [2, 20]. The optimizer typically has difficulty climbing out of this loss basin, resulting in a model that eventually converges to a totally wrong solution. Based on this hypothesis, there should be a significant loss barrier [11] between the failure modes and the true solution of the PDE.

However, our empirical results show that there is no such a loss barrier between the failure mode and the true solution that can block the optimizer from surmounting. As shown in Fig. 1 (a)&(b), we find that the model that eventually converges to an ideal numerical solution also has experienced a failure mode pattern during the training process. There is a similar non-synchronous decrease in residual loss and error during the training process in both success and failure mode cases. As in Fig. 1 (c)&(d), the expected sudden rise in loss in line with the traditional loss landscape understanding does not occur during the transition of the model from a failure mode pattern to a success case.

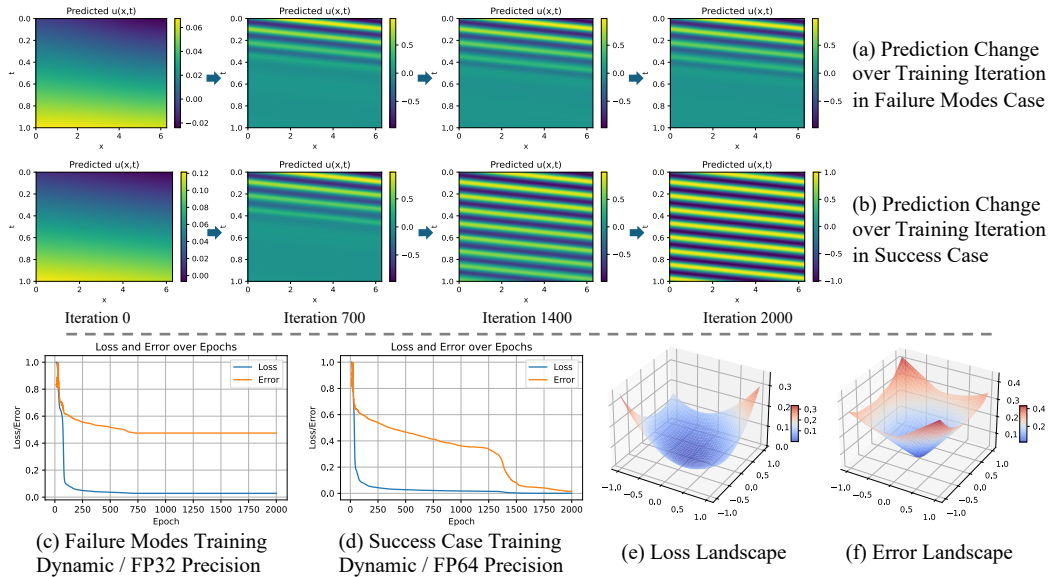


Figure 1: Training Dynamic of PINN's failure mode case and success case.

This phenomenon directly challenges the current understanding of the loss landscape of PINN failure cases in the machine learning community. The cause of the PINN failure case is not a local optimal solution. Instead, as shown in Fig. 1 (e)&(f), there is a significant discrepancy between the loss landscape and error landscape of a PINN. It suggests that the so-called failure mode may just be an intermediate stage of PINN optimization. The model ends up in failure mode due to an unanticipated stop in the optimization process. Therefore, this leads to a fundamental question:

What makes PINN's optimization process unexpectedly stop at a failure mode?

We found the answer to this question surprisingly simple: The **arithmetic precision** of the model is not enough to maintain the optimization process any longer. Arithmetic precision is generally considered not to play a decisive role in neural networks in vision and language, where using FP32, FP16 and even lower precision is common practice. But the problem PINN is dealing with is scientific computing with high precision requirements. By using the default precision setting provided in the deep learning framework, the convergence condition of L-BFGS [21], which is the mostly used robust optimizer for PINN, is triggered prematurely due to the lack of enough arithmetic precision. To illustrate that, we show in Fig. 1, with exactly the same problem setting, models with FP32 precision consistently present failure modes, while the models with FP64 precision always succeed.

Regarding the training dynamic, we identify that the training process of PINN will undergo three phases: un-converged phase, failure phase, and success phase. With the same initialization and optimization setup, models with different precisions will finally stop in different stages of optimization. Failure modes problems that used to be considered more difficult (e.g., with a higher-frequency convection parameter) have longer failure phases and more gentle loss plain. As a result, the required arithmetic precision is positively correlated with the difficulty of solving a failure mode problem.

Further, we find that all these known failure modes can be solved by vanilla PINN with sufficient arithmetic precision using the L-BFGS optimizer. Performance of vanilla PINNs with FP64 precision on PDEs like convection can surpass most state-of-the-art model architectures that claim an enhancement on PINN failure modes. With reduced arithmetic precision, models like PINNsFormer [36] show a typical failure mode feature again, indicating that it may not solve the failure modes fundamentally.

Contributions. In this paper, we mainly made the following contributions:

- We revolutionize the understanding of PINN failure modes. We reveal that, on the loss landscape of PINN, instead of being in a separate loss basin, the failure mode has a pathway to the optimal solution that can be found by the optimizer.

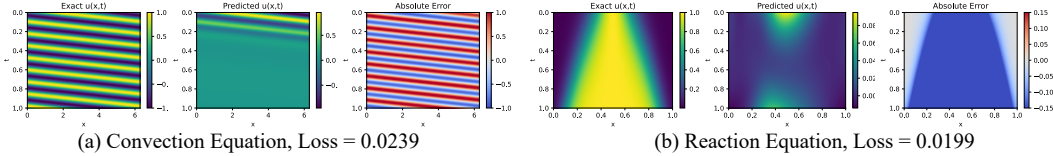


Figure 2: Some cases of PINN failure modes.

- We reveal that unexpected stopping of the optimizer due to insufficient arithmetic precision is the core cause of PINN failure modes. With suitable precision, vanilla PINN is strong enough to avoid failure modes with the L-BFGS optimizer.
- We show that the optimization process of PINN is divided into three stages related to difficulty, and different precisions will make the optimization stop at different stages.

2 Related Works

Physics-Informed Neural Networks. Physics-Informed Neural Networks[27] embed the governing differential equations of a physical system directly into a neural-network training objective. Instead of fitting a network solely to observational data, PINNs minimize a residual loss that penalizes violations of the PDE (or ODE/integro-differential law) at a set of collocation points. This simple idea turns a neural network $u_\theta(x, t)$ into a mesh-free solver that learns a function satisfying

$$\mathcal{F}(u_\theta(x, t)) = 0, \forall (x, t) \in \Omega; \mathcal{B}(u_\theta(x, t)) = 0, \forall (x, t) \in \partial\Omega, \quad (1)$$

where $(x, t) \in \Omega$ is the spatial-temporal coordinate, $\partial\Omega$ is the boundary of Ω . The $\mathcal{F}$, and $\mathcal{B}$ denote the operators defined by PDE equations, initial conditions, and boundary conditions, respectively.

The training of PINNs leverages automatic differentiation techniques [25] available within contemporary deep learning frameworks, such as PyTorch [26], to perform numerical differentiation and thereby construct the residual loss at selected collection points. Recent studies have demonstrated that second-order optimization methods, notably quasi-Newton methods such as SSBFGS [24], Broyden [5], and L-BFGS [21], can significantly enhance the stability and fitting accuracy of PINN training [18, 15, 33, 20], compared to relying exclusively on first-order optimizers such as Adam [17].

Failure Modes in Physics-Informed Neural Networks. Despite the promising performance and broad applicability, PINNs are susceptible to several failure modes that hinder their effectiveness in practice [20]. The failure modes of PINN refer to a phenomenon in which the PDE residual loss of a model is optimized to a very low level (close to 0), but its resulting solution is still far from the expected true solution. As in Fig. 2, 1-d convection and reaction equations are optimized to close-zero level, but their approximations to the equation's true solutions present the over-simplified patterns.

In response to this phenomenon, researchers have come up with a very wide range of conjectures. Early attempts involve curriculum regularization and sequence-to-sequence learning [20], which breaks down complex PINN problems into simple ones and conquers them one at a time. Subsequently, R3 [9] tries to address this issue with resampling on the failure area. Then, several model architectures like PINNsFormer [36], ProPINN [34], PirateNet [30], and PINNMamba [35] are proposed to help the propagation of the correct pattern among the collection points. From an optimization perspective, researchers also propose regional optimization [33] and Adam+L-BFGS [28] to help the convergence.

Yet, these approaches rely on a hypothesis that the failure mode of PINN is a local minimal that completely isolated from the true solution in loss landscape. During the optimization, model either goes towards a failure mode loss basin or a true solution basin. These approaches intend to induce the model towards the true solutions with an inductive bias. While these methods have achieved some results on some part of failure mode cases, their understanding is not intrinsic. We argue that the failure modes are in the same loss basin as the true solution, and are caused by early stopping.

Arithmetic Precision in Neural Networks. Arithmetic precision is commonly thought to have a subsidiary impact on the mainstream applications of neural networks [14]. Although early deep-learning systems relied almost exclusively on 32-bit floating-point arithmetic, a large body of empirical work now shows that reducing numerical precision rarely degrades the accuracy of modern models.

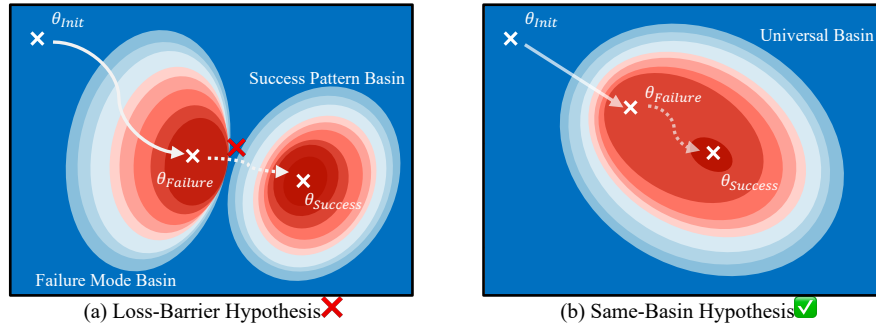


Figure 3: A schematic that illustrates two kinds of loss landscape hypothesis.

Production-scale vision, speech, and language pipelines routinely deploy INT8/FP8/FP16/BF16 kernels or mixed-precision execution and deliver FP32-level quality with multi-fold improvements in speed and energy efficiency [37, 13, 10]. On the GPU-design side, hardware is now co-engineered for these low-precision formats: NVIDIA’s Hopper architecture [8] adds FP8/INT8 Tensor Cores and a Transformer Engine for efficiency-accuracy trade-off. This co-exploration of lower-precision from both software and hardware hints that, in general application of neural networks, arithmetic precision might not be an essential determinant of the model performance. However, we identify that this trend may not fit the scientific machine learning. Numerical methods based on neural networks need to take the same high accuracy calculations as traditional methods (e.g., finite element methods [3]).

3 Problem and Experiment Settings

3.1 Problem Settings

The PINN employs a neural network parameterized by θ to approximate the solution of a PDE defined as Eq. 1 by optimizing the following residual loss on 101×101 collection points (x_i, t_i) :

$$\begin{aligned} \text{minimize}_{\theta} \quad & \mathcal{L}(u_{\theta}) = \lambda_{\mathcal{F}} \mathcal{L}_{\mathcal{F}}(u_{\theta}) + \lambda_{\mathcal{B}} \mathcal{L}_{\mathcal{B}}(u_{\theta}), \\ \text{where } \mathcal{L}_{\mathcal{F}}(u_{\theta}) = & \frac{1}{|\chi|} \sum_{(x_i, t_i) \in \chi} \|\mathcal{F}(u_{\theta}(x_i, t_i))\|^2; \mathcal{L}_{\mathcal{B}}(u_{\theta}) = \frac{1}{|\partial\chi|} \sum_{(x_i, t_i) \in \partial\chi} \|\mathcal{B}(u_{\theta}(x_i, t_i))\|^2; \end{aligned} \quad (2)$$

To evaluate the performance of the models, we take relative Mean Absolute Error (rMAE, a.k.a ℓ_1 relative error) and relative Root Mean Square Error (rRMSE, a.k.a ℓ_2 relative error) formulated as:

$$\text{rMAE}(\hat{u}) = \frac{\sum_{n=1}^N |u_{\theta}(x_n, t_n) - u(x_n, t_n)|}{\sum_{n=1}^N |u(x_n, t_n)|}, \quad \text{rRMSE}(\hat{u}) = \sqrt{\frac{\sum_{n=1}^N |u_{\theta}(x_n, t_n) - u(x_n, t_n)|^2}{\sum_{n=1}^N |u(x_n, t_n)|^2}}, \quad (3)$$

where N is the number of test points, $u(x, t)$ is the ground truth, and $\hat{u}(x, t)$ is the model’s prediction.

In this paper, we investigate four famous failure mode PDEs that are widely studied in the PINN community: convection equations, reaction equations, wave equations, and Allen-Cahn [20, 36, 34]. The descriptions and details of these problems can be found in the Appendix A.

3.2 Experiments Settings

We initialize the vanilla PINN with an MLP with 3 hidden layers and 512 neurons in each layer. We include advanced PINN neural architectures like PINNsFormer [36], KAN [22], PINNMamba [35] to test the generalization ability of our findings, following their original settings. We train the models with L-BFGS optimizer following common practice [20, 33, 36, 34, 35]. All the experiments are implemented on an NVIDIA H100 GPU, with CUDA version 12.8 and Pytorch version 2.1.1.

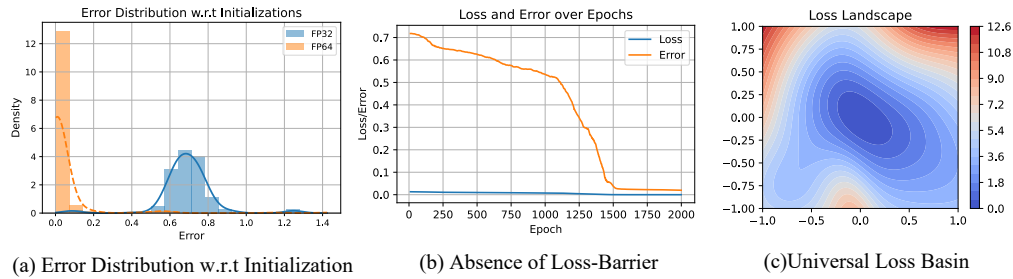


Figure 4: Several empirical results on convection support Same-Basin Hypothesis.

4 Understanding PINN Failure Modes from Loss Landscape Perspective

In this section, we demonstrate that the failure modes of PINNs are not attributed to the local minima with extremely steep loss landscape, which is a view widely accepted by scholars who have studied this area. Instead, we point out that, the failure modes are in the same loss basin as the true solution. The model finally present a failure mode is ascribed to the early stopping of the optimization process.

4.1 Loss-Barrier Hypothesis vs Same-Basin Hypothesis: Which is the True Story?

We first state two hypotheses for failure modes of PINNs: (1) the *Loss-Barrier Hypothesis*, which is widely accepted by the PINN researchers, and (2) the *Same-Basin Hypothesis*, which we propose.

Loss-Barrier Hypothesis. PINN failures are initially attributed to that “the optimizer has gotten stuck in a local minima with a very high loss function” [20]. This assertion is then been overthrown because the observation of near-zero empirical residual loss in failure modes experiment cases [33, 34, 35, 36]. This hints that the cause of PINN failure modes is not that PDE residual loss cannot be optimized. Instead, as shown in Fig. 3 (a), it seems that the failure modes are because the optimizer is trapped in a local minimum of the loss landscape, which is nearly as low as the true solution’s empirical residual loss on collection points. Under this hypothesis, the optimizer is not able to distinguish the failure modes and the true solution minimum. Therefore, models cannot find the true solution, because there is a **loss barrier** between the failure mode local minimal and the true solution minimal. The model’s loss faces a huge penalty if its parameters are optimized to cross such a loss barrier.

Same-Basin Hypothesis. Based on the intuition that the loss basins are generally connected with a simple perturbation [1], we propose an alternative same-basin hypothesis. It says that the failure modes parameters are in the same loss basin as the true solution, as shown in Fig. 3 (b). There is a large flat plain at the bottom of this loss basin, and the loss of the true solution is only slightly lower than that of the failure modes. Due to some reasons, the optimizer early stops when it reaches the flat plain in the iterated optimization process. In this hypothesis, there should not be a significant loss barrier between the failure mode and the true solution. As a result, the optimizer with a global convergence guarantee should be able to find the global minimum empirically with ideal computational conditions.

4.2 Physics-Informed Neural Networks Follow Same-Basin Hypothesis

Although the Loss-Barrier Hypothesis is generally believed to hold, we identify three important clues pointing to the Same-Basin Hypothesis: (1) initialization insensitivity, (2) absence of loss barrier, and (3) loss connectivity. This hints that the failure modes are caused by insufficient optimization.

PINNs Fail with any Random Initializations. If the failure mode patterns and the true solution patterns are in separate loss basins, it should be the case that, with some initialization, the model can be optimized to the true solution. But in practice, the situation is that the model is always trapped in the failure mode pattern, no matter how we conduct the random initializations. As shown in Fig.4 (a), with 2000 different random initializations, the trained PINNs always fail on convection equations. This implies the extreme sparsity of the global optimum over the loss landscape. Under this implication, Loss-Barrier Hypothesis completely contradicts the intrinsic lottery ticket phenomena of neural networks [11]. This suggests choosing the alternative hypothesis: Same-Basin Hypothesis.

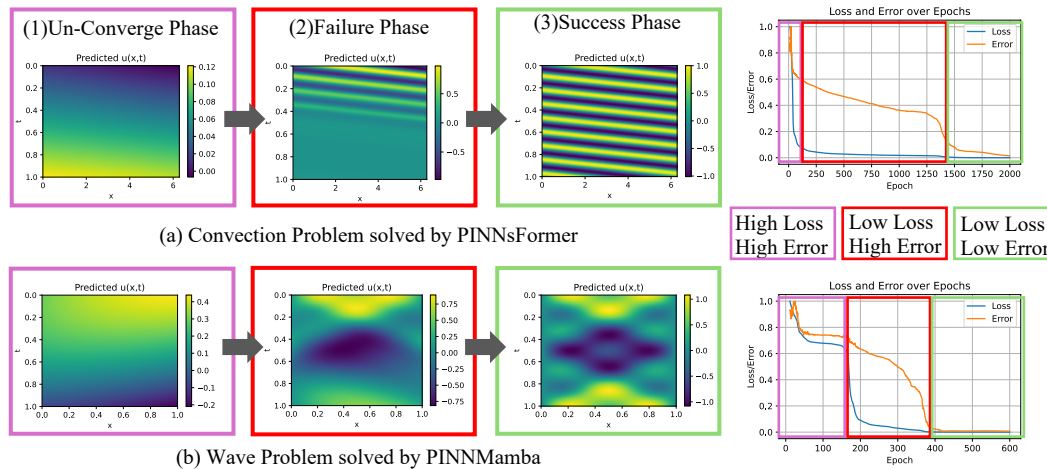


Figure 5: Loss-Error Dynamics of PINN models can be in three phases.

Non-existent Loss Barrier from Failure Mode to Optimized Solution. A gravely important finding is that the expected loss barrier in the Loss-Barrier Hypothesis doesn't empirically exist. To illustrate this finding, we initialize the PINN model with failure mode parameters that have been fully optimized for the convection problem, where there is a convergence signal of L-BFGS optimizer. Then we turn the arithmetic precision from FP32 to FP64 (an important finding we will discuss in Section 5), and continue the training process for 2000 iterations. We find that, unlike the loss function that is no longer updated at all as shown in Fig. 1 (c), the loss shows a further trace decline. As shown in Fig. 4 (b), during this further decline process, we observe a steep decline of the error (rMAE). This is direct evidence against the Loss-Barrier Hypothesis. The absence of a loss barrier in practice hints that the failure mode minimal and the true solution minimal are connected in the loss landscape.

Universal Trough in Loss Landscape. Further, if the Loss-Barrier Hypothesis holds, there should be a very large number of loss basins that are close to zero at their lowest point. However, as shown in Fig. 4 (c), we can only observe one trough in the loss landscape. The minimum basins are not separated, instead, there is a connected trough in the loss landscape. This directly disproves the assertion that PINN's loss is too rugged, making model optimization difficult. This phenomenon suggests that the failure modes from PINN and the true solution should be in the same basin, and there should be an optimization pathway between them. We further visualized the loss and error landscape of a well-trained model, as shown in Fig. 1 (e)&(f). We found that the loss is relatively level and smooth, and does not have the expected extremely rugged loss landscape. In contrast, the error landscape is much steeper at the position near the true solution optimal. This means that a small perturbation to a parameter can have a much larger effect on the error of the model than the loss. At the same time, the true solution exists in only a very small region on this large loss plain.

The three clues above show, directly or indirectly, that the real reason why PINNs can have failure modes is not an insurmountable barrier in the loss landscape, in other words, it is not due to being stuck in a local optimum. Instead, the PINN true solution of the equation and the failure modes are in the same loss basin. This means that when considering the optimization problem of PINN, the researcher should accept the Same-Basin Hypothesis for model design and tuning.

4.3 The Loss-Error Dynamics of PINN's Training

Based on Same-Basin Hypothesis, the whole training dynamics of the PINN model can be divided into three phases: (1) Un-Converged phase, where both the loss and error are still in high level, suggesting the model parameters are still not optimized; (2) Failure phase, where the loss is optimized to near-zero level but the error is still in the high level, and the model presents failure modes such as naive solution; and (3) Success phase, where the model parameter can both achieve near-zero level loss and error. This implies that every model that is optimizable to arrive at a near-true solution also experiences a pattern of failure modes at some point during the training process, yet the model is able

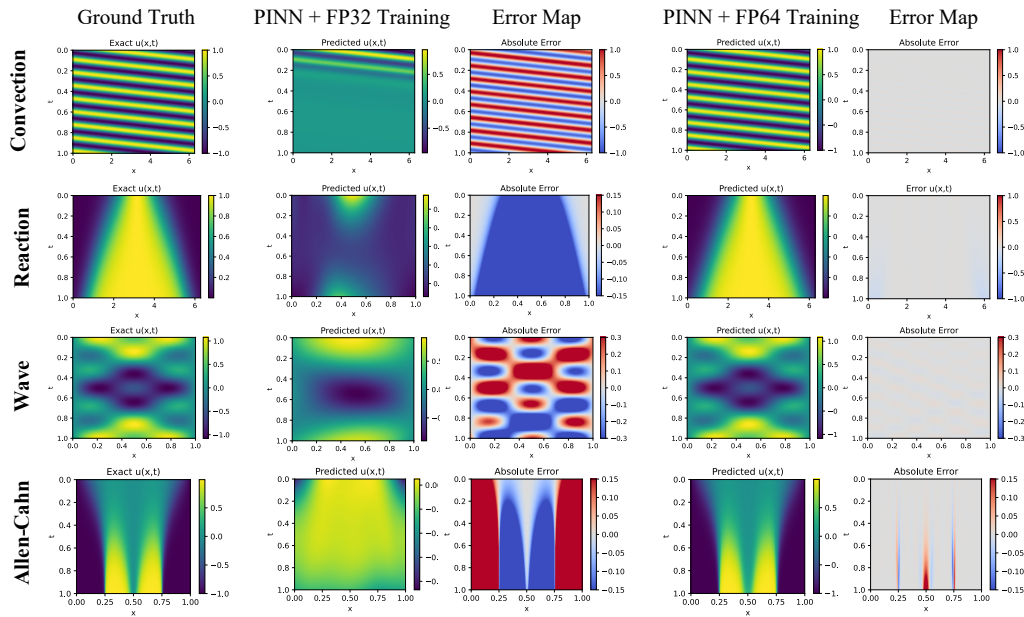


Figure 6: Results for solving 4 failure mode PDE problems with different arithmetic precision.

to be further optimized and eventually arrives at the optimal solution of the equations that the model is capable of giving. Thus, the failure modes are intermediate states of the PINNs training dynamics.

We empirically confirm these 3 phases with different PINN models and PDE problem settings. As shown in Fig. 5, we train a PINNsFormer [36] and a PINNMamba [35], which have been empirically proven to solve the failure modes problem, for the convection and wave problem, respectively. We note that the division of training dynamics of PINN into three phases is common in various problem settings. Different problem settings have different phase duration characteristics. Those PDEs with longer failure phase are more sensitive and likely to be affected by failure modes. When the PINN failure modes were first identified, it was concluded that the difficulty of modeling the PDE increases with the PDE parameter [20]. We attribute this to that the larger PDE parameter will extend the duration of the failure phase. Therefore, the optimization processes are more likely to stop at the failure phase.

5 Precision Serves a Key Factor in PINN's Training

We now know that the cause of failure modes is the premature stopping of the model optimization process. This leads to another question: what contributes to the premature stopping of PINN model optimization? In this section, we reveal that the arithmetic precision is a key issue of the PINN optimization. Unlike neural networks in CV and NLP, single-precision (FP32) training does not guarantee that all neural networks corresponding to PDE problems are efficiently optimized. Instead, double-precision (FP64) training is a stable strategy for various model and problem settings.

5.1 Failure Modes Disappear when Using Double-Precision Training

We first show that the most basic vanilla PINN can solve all currently known failure modes problems when using double precision training. As shown in Fig. 6, with all these PDE problems where the vanilla PINN exhibits failure modes, the use of double-precision training yields solutions that are highly consistent with ground truth. This suggests that using double-precision training allows the model optimization to go through all three phases, rather than stopping at the failure phase. Moreover, as shown in Table 1, the vanilla PINN trained with double precision can outperform all state-of-the-art methods in terms of prediction error. It hints that the inductive-biased approach that was once based on the understanding of the Loss-Barrier Hypothesis is not fundamental to the addressing of PINN

Table 1: Comparison with baseline methods on 4 failure mode problems.

Model	Convection			Reaction		
	Loss	rMAE	rRMSE	Loss	rMAE	rRMSE
PINN [27]	0.0133 $\pm$ 0.0055	0.6904 $\pm$ 0.0826	0.7640 $\pm$ 0.0694	0.1991 $\pm$ 0.0001	0.9788 $\pm$ 0.0019	0.9778 $\pm$ 0.0018
QRes [6]	0.0153 $\pm$ 0.0027	0.7498 $\pm$ 0.0464	0.8184 $\pm$ 0.0382	0.1991 $\pm$ 0.0001	0.9826 $\pm$ 0.0023	0.9830 $\pm$ 0.0026
PINNsFormer [36]	0.0009 $\pm$ 0.0001	0.0327 $\pm$ 0.0068	0.0435 $\pm$ 0.0073	3.0e-6 $\pm$ 1.0e-6	0.0147 $\pm$ 0.0013	0.0296 $\pm$ 0.0027
KAN [22]	0.0250 $\pm$ 0.0042	0.6213 $\pm$ 0.0675	0.6985 $\pm$ 0.0701	7.0e-6 $\pm$ 1.0e-6	0.0167 $\pm$ 0.0014	0.0312 $\pm$ 0.0034
PirateNet [30]	0.0347 $\pm$ 0.0061	0.9704 $\pm$ 0.1826	0.9740 $\pm$ 0.1894	4.0e-6 $\pm$ 1.0e-6	0.0178 $\pm$ 0.0023	0.0443 $\pm$ 0.0064
RoPINN [33]	0.0189 $\pm$ 0.0062	0.6251 $\pm$ 0.0940	0.7204 $\pm$ 0.0941	4.8e-5 $\pm$ 9.0e-6	0.0589 $\pm$ 0.0161	0.0965 $\pm$ 0.0310
PINN Mamba [35]	0.0001 $\pm$ 2.0e-5	0.0184 $\pm$ 0.0037	0.0197 $\pm$ 0.0038	1.0e-6 $\pm$ 1.0e-6	0.0092 $\pm$ 0.0017	0.0213 $\pm$ 0.0036
PINN_FP64	5.0e-6 $\pm$ 1.0e-6	0.0059 $\pm$ 0.0013	0.0072 $\pm$ 0.0017	1.0e-5 $\pm$ 5.0e-6	0.0271 $\pm$ 0.0063	0.0502 $\pm$ 0.0111

Model	Wave			Allen-Cahn		
	Loss	rMAE	rRMSE	Loss	rMAE	rRMSE
PINN [27]	0.0174 $\pm$ 0.0061	0.2746 $\pm$ 0.0574	0.2837 $\pm$ 0.0571	0.4703 $\pm$ 0.2986	0.9720 $\pm$ 0.0370	0.9662 $\pm$ 0.0300
QRes [6]	0.0965 $\pm$ 0.0192	0.5335 $\pm$ 0.1230	0.5273 $\pm$ 0.1172	0.9887 $\pm$ 0.0021	0.9821 $\pm$ 0.0089	0.9846 $\pm$ 0.0092
PINNsFormer [36]	0.0231 $\pm$ 0.0017	0.3492 $\pm$ 0.0871	0.3571 $\pm$ 0.0872	0.4625 $\pm$ 0.2875	0.9908 $\pm$ 0.0401	0.9913 $\pm$ 0.0420
KAN [22]	0.0067 $\pm$ 0.0012	0.1475 $\pm$ 0.0354	0.1489 $\pm$ 0.0357	0.0234 $\pm$ 0.0031	0.3166 $\pm$ 0.0233	0.5661 $\pm$ 0.0440
PirateNet [30]	0.0153 $\pm$ 0.0051	0.2544 $\pm$ 0.0471	0.2637 $\pm$ 0.0480	0.0017 $\pm$ 0.0001	0.1088 $\pm$ 0.0109	0.1889 $\pm$ 0.0180
RoPINN [33]	0.0015 $\pm$ 0.0005	0.0631 $\pm$ 0.0226	0.0642 $\pm$ 0.0238	-	-	-
PINN Mamba [35]	0.0002 $\pm$ 3e-5	0.0193 $\pm$ 0.0033	0.0195 $\pm$ 0.0033	0.0027 $\pm$ 0.0002	0.1432 $\pm$ 0.0123	0.2645 $\pm$ 0.0201
PINN_FP64	4.2e-5 $\pm$ 1.6e-5	0.0080 $\pm$ 0.0032	0.0081 $\pm$ 0.0031	1.3e-5 $\pm$ 4.0e-6	0.0157 $\pm$ 0.0036	0.0545 $\pm$ 0.0112

failure modes. Instead, insufficient arithmetic precision is a straightforward contributor to the failure modes of the PINN. As with traditional numerical methods such as Finite Element Analysis, PINN also requires the use of double precision to achieve accurate and trustworthy results.

5.2 Optimization Stops Early without Sufficient Arithmetic Precision

Based on the analysis in Section 4, the failure modes of PINN are due to premature stopping of the optimization process. In order to find the reason why arithmetic precision affects the optimization stopping of the PINN, we target the optimizer and its implementation in a deep learning framework.

PINNs Commonly Use L-BFGS Optimizers. The PDE residual contains high-order differential operators whose spectra can span many orders of magnitude, giving a Hessian with extreme eigenvalue ratios. First-order methods (Adam/SGD) therefore move in directions dominated by the largest eigenvalues and make little progress on the stiff directions. Rightly, second-order optimizers such as Newton, BFGS, and L-BFGS implicitly (or explicitly) invert an estimate of the Hessian, scaling each parameter update by the local curvature and thereby shrinking the stiff directions while enlarging the flat ones. Among all the second-order optimizers, L-BFGS is the most commonly utilized for its excellent performance while maintaining relatively low-level time and space complexity. L-BFGS [21] approximates the inverse Hessian using a history of the past m updates of the position and gradient. At each iteration, it constructs the Hessian approximation implicitly using the stored history. We include the details of the L-BFGS optimizer in Appendix B.

Existing Implementation of L-BFGS Optimizer. In deep learning framework’s implementation of the L-BFGS algorithm, such as ‘torch.optim.LBFGS’, algorithmic control is divided between user-defined outer loops and library-defined inner loops. The user governs the outer loop by repeatedly invoking ‘optimizer.step(closure)’, where each call corresponds to a single quasi-Newton update that assimilates one curvature pair for the inverse Hessian. During a single call to ‘optimizer.step(closure)’, it may perform up to ‘max_iter’ inner loops that repeatedly reevaluate the objective and gradient, compute a quasi-Newton search direction via the two-loop recursion with the latest ‘history_size’ curvature pairs, optionally conduct a Strong-Wolfe line search [31], and update both parameters and curvature history. The inner loop halts when convergence criteria on gradient norm or successive parameter changes are met, after which control returns to the user-level outer loop.

Low-Precision Causes the Inner Loop Early Stopping. We found that insufficient arithmetic precision prematurely triggers the convergence criteria of the inner loop. As shown in Fig. 7 (a), FP32 triggers the inner loop early stopping (the count of inner loops ≤ 5) much earlier than FP64 does. This is due to the fact that the trigger condition for convergence in PINN (‘tolerance_change’) has a value of $1e-7$ that is smaller than the machine unit ε for single precision floating point numbers. Machine unit ε is the smallest positive floating-point number that, when added to 1, yields a representable

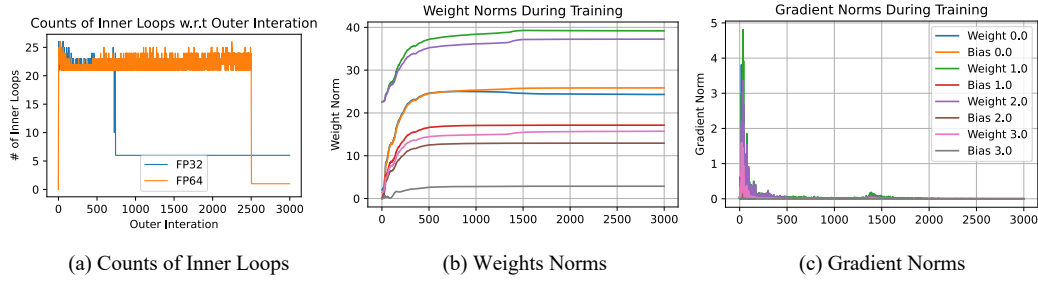


Figure 7: The values of the PINN weights are orders of magnitude different than the values of the gradients, and the ratio is also incremental, causing single precision to underflow earlier.

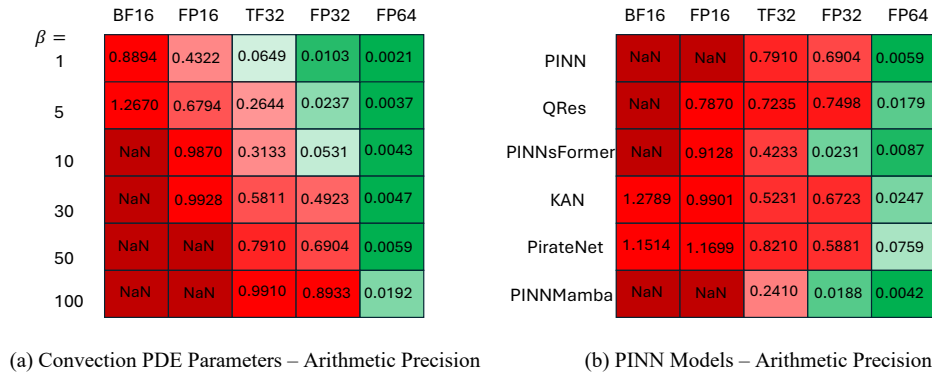


Figure 8: Results on convection equations with different arithmetic precision: (a) PINN’s rMAE on various convection parameters β . (b) Different models’ rMAE on convection with $\beta = 50$.

value strictly greater than 1 on a given hardware/format:

$$\varepsilon = \min\{\varepsilon > 0 : 1 + \varepsilon \text{ is representable and } 1 + \varepsilon \neq 1\}. \quad (4)$$

For a radix $\beta (= 2)$ and p binary digits in the significand, $\varepsilon = \frac{1}{2} \beta^{1-p}$, which means the FP32 has the machine $\varepsilon = 1.19 \text{ e-}7 > \text{‘tolerance_change’} = 1\text{e-}7$, while FP64’s machine $\varepsilon = 2.22 \text{ e-}16$.

As shown in Fig. 7 (b), during the optimization of PINN, the Norms of the weights are incrementally increasing in the order of magnitude of $1\text{e}+1$ level, which is greater than 1. This means that the computation effective numerical precision is even more coarse-grained than machine ε , and thus computing the changes of weights will be easier to underflow. In Fig. 7 (c), we show that the norms of the gradient are decreasing to near-zero level gradually, which will cause the PINN optimization problem to be extremely ill-conditioned. Therefore, the amount of weight change can quickly drop below machine ε with FP32 training, since the weight change is positively correlated with gradient.

5.3 Harder Equations Require Higher Arithmetic Precision

It has been shown that, for a parameterized PDE, the PINN optimization of the equations becomes progressively more difficult as the parameters of the equations become larger, and even failure modes occur [20, 7]. The reason for this phenomenon is that the frequency of model changes will become higher, and the pattern of the model will become more versatile. We attribute this to the fact that as the parameter of the equation gets larger, the separation of its error and loss landscapes also gets progressively larger. This means that for more difficult problems with a larger parameter, an early stopping point in the optimization process will result in a larger error. As shown in Fig. 8 (a), we confirm this problem’s complex trend with varying degrees of arithmetic precision from BF16 to FP64. We find that for more complex equation problems (like $\beta = 100$), increasing the arithmetic precision of the model to double precision can effectively solve the failure modes of PINN.

Table 2: Comparison of FP32 and FP64 configurations in terms of relative RMSE, training time per iteration, and memory usage across PDE benchmarks.

FP32			
Equation	rRMSE	Training Time/Iter (s)	Memory (MB)
convection	0.7640	0.29	1609
reaction	0.9778	0.26	1629
wave	0.2837	0.47	2295
Allen-Cahn	0.9662	0.40	1975
FP64			
Equation	rRMSE	Training Time/Iter (s)	Memory (MB)
convection	0.0072	0.32	2441
reaction	0.0502	0.29	2481
wave	0.0081	0.61	3845
Allen-Cahn	0.0545	0.48	3167

5.4 Arithmetic Precision’s Impact on Various Model Architectures.

In order to verify whether the multiple understandings of model architectures [36, 35] claiming to address PINN failure modes are intrinsic, and whether our proposed method for improving numerical precision is orthogonal to other methods, as shown in Fig. 8, we evaluate a lot of models on convection equations with parameter β being set to 50. We find that just a slight decrease in the precision of PINNsFormer and PINNMamba (from FP32 to TF32) can make the failure modes reproduce, which suggests that the method of constructing inductive bias based on the Loss-Barrier Hypothesis is not an intrinsic solution to the failure modes problem. Using double-precision training, all models converge to the correct solution, which illustrates that our understanding of loss landscape and model precision holds the correctness, generalizability, and orthogonality to other methods.

5.5 Discussion of Computation

As shown in Table 2, we evaluate the FP32 and FP64 training overhead on a Nvidia H100 PCIe GPU. The FP64 latency per iteration is roughly 1.1-1.3 times that of FP32, and the GPU memory footprint is about 1.5-1.7 times that of FP32. This is due to the fact that, with Tensor Core, the FP64 peak arithmetic on H100 PCIe is 60 TFLOPS, the same as the FP32 peak arithmetic. As a result, FP64 is not very much slower than FP32. We observe about 10-30% more latency of FP64 than FP32 that we believe is due to larger memory accesses. This indicates that training with FP64 does incur an increase in computational overhead, but under modern GPU architectures the increase is marginal.

6 Conclusion

In this paper, through hypothesis testing on the loss landscape of PINN, we demystify that the real cause of PINN’s failure modes is the early stopping of the optimization process. We reveal that this early stopping phenomenon can be easily solved by improving the training arithmetic precision of the model. Training with double precision will enable the vanilla PINN to beat state-of-the-art PINN model architectures on various PDE problem settings. This prompts the broader AI for Science community that PINN is still essentially a numerical solution method for PDEs, and that computational precision plays just as crucial a role as it does in traditional methods.

Acknowledgements

This work is supported, in part, by the National Science Foundation and the Institute of Education Sciences under Grant 2229873 (AI4ExceptionalEd), National Science Foundation under Grant 2235364 (FuSe-TG), and SUNY-IBM AI Collaborative Research Award. Any opinions, findings and conclusions or recommendations expressed in this material are those of the author(s) and do not necessarily reflect the views of the sponsors.

References

- [1] S. Ainsworth, J. Hayase, and S. Srinivasa. Git re-basin: Merging models modulo permutation symmetries. In *The Eleventh International Conference on Learning Representations*, 2023.
- [2] S. Basir and I. Senocak. Critical investigation of failure modes in physics-informed neural networks. In *AiAA SCITECH 2022 Forum*, page 2353, 2022.
- [3] F. Bassi and S. Rebay. A high-order accurate discontinuous finite element method for the numerical solution of the compressible navier–stokes equations. *Journal of computational physics*, 131(2):267–279, 1997.
- [4] A. G. Baydin, B. A. Pearlmutter, A. A. Radul, and J. M. Siskind. Automatic differentiation in machine learning: a survey. *Journal of machine learning research*, 18(153):1–43, 2018.
- [5] C. G. Broyden. A class of methods for solving nonlinear simultaneous equations. *Mathematics of computation*, 19(92):577–593, 1965.
- [6] J. Bu and A. Karpatne. Quadratic residual networks: A new class of neural networks for solving forward and inverse problems in physics involving pdes. In *Proceedings of the 2021 SIAM International Conference on Data Mining (SDM)*, pages 675–683. SIAM, 2021.
- [7] W. Cho, M. Jo, H. Lim, K. Lee, D. Lee, S. Hong, and N. Park. Parameterized physics-informed neural networks for parameterized pdes. In *International Conference on Machine Learning*, pages 8510–8533. PMLR, 2024.
- [8] J. Choquette. Nvidia hopper h100 gpu: Scaling performance. *IEEE Micro*, 43(3):9–17, 2023.
- [9] A. Daw, J. Bu, S. Wang, P. Perdikaris, and A. Karpatne. Mitigating propagation failures in physics-informed neural networks using retain-resample-release (r3) sampling. In *International Conference on Machine Learning*, pages 7264–7302. PMLR, 2023.
- [10] M. Fishman, B. Chmiel, R. Banner, and D. Soudry. Scaling fp8 training to trillion-token llms. In *Proceedings of the International Conference on Learning Representations (ICLR)*, 2025.
- [11] J. Frankle, G. K. Dziugaite, D. Roy, and M. Carbin. Linear mode connectivity and the lottery ticket hypothesis. In *International Conference on Machine Learning*, pages 3259–3269. PMLR, 2020.
- [12] Z. Gao, L. Yan, and T. Zhou. Failure-informed adaptive sampling for pinns. *SIAM Journal on Scientific Computing*, 45(4):A1971–A1994, 2023.
- [13] A. Gholami, S. Kim, Z. Dong, Z. Yao, M. W. Mahoney, and K. Keutzer. A survey of quantization methods for efficient neural network inference. *arXiv preprint arXiv:2103.13630*, 2021.
- [14] S. Gupta, A. Agrawal, K. Gopalakrishnan, and P. Narayanan. Deep learning with limited numerical precision. In *International conference on machine learning*, pages 1737–1746. PMLR, 2015.
- [15] Z. Hao, C. Ying, H. Su, J. Zhu, J. Song, and Z. Cheng. Bi-level physics-informed neural networks for pde constrained optimization using broyden’s hypergradients. In *International Conference on Learning Representations*, 2023.
- [16] K. Hornik, M. Stinchcombe, and H. White. Multilayer feedforward networks are universal approximators. *Neural networks*, 2(5):359–366, 1989.
- [17] D. P. Kingma and J. Ba. Adam: A method for stochastic optimization. *arXiv preprint arXiv:1412.6980*, 2014.
- [18] E. Kiyani, K. Shukla, J. F. Urbán, J. Darbon, and G. E. Karniadakis. Which optimizer works best for physics-informed neural networks and kolmogorov-arnold networks? *arXiv preprint arXiv:2501.16371*, 2025.
- [19] D. A. Kopriva. *Implementing spectral methods for partial differential equations: Algorithms for scientists and engineers*. Springer Science & Business Media, 2009.
- [20] A. Krishnapriyan, A. Gholami, S. Zhe, R. Kirby, and M. W. Mahoney. Characterizing possible failure modes in physics-informed neural networks. *Advances in neural information processing systems*, 34:26548–26560, 2021.
- [21] D. C. Liu and J. Nocedal. On the limited memory bfgs method for large scale optimization. *Mathematical programming*, 45(1):503–528, 1989.

- [22] Z. Liu, Y. Wang, S. Vaidya, F. Ruehle, J. Halverson, M. Soljacic, T. Y. Hou, and M. Tegmark. KAN: Kolmogorov–arnold networks. In *The Thirteenth International Conference on Learning Representations*, 2025.
- [23] P. C. Nguyen, X. Cheng, S. Azarfar, P. Seshadri, Y. T. Nguyen, M. Kim, S. Choi, H. Udaykumar, and S. Baek. Parcv2: Physics-aware recurrent convolutional neural networks for spatiotemporal dynamics modeling. In *International Conference on Machine Learning*, pages 37649–37666. PMLR, 2024.
- [24] J. Nocedal and Y.-x. Yuan. Analysis of a self-scaling quasi-newton method. *Mathematical Programming*, 61(1):19–37, 1993.
- [25] A. Paszke, S. Gross, S. Chintala, G. Chanan, E. Yang, Z. DeVito, Z. Lin, A. Desmaison, L. Antiga, and A. Lerer. Automatic differentiation in pytorch. 2017.
- [26] A. Paszke, S. Gross, F. Massa, A. Lerer, J. Bradbury, G. Chanan, T. Killeen, Z. Lin, N. Gimelshein, L. Antiga, A. Desmaison, A. Kopf, E. Yang, Z. DeVito, M. Raison, A. Tejani, S. Chilamkurthy, B. Steiner, L. Fang, J. Bai, and S. Chintala. Pytorch: An imperative style, high-performance deep learning library. In H. Wallach, H. Larochelle, A. Beygelzimer, F. d'Alché-Buc, E. Fox, and R. Garnett, editors, *Advances in Neural Information Processing Systems*, volume 32. Curran Associates, Inc., 2019.
- [27] M. Raissi, P. Perdikaris, and G. E. Karniadakis. Physics-informed neural networks: A deep learning framework for solving forward and inverse problems involving nonlinear partial differential equations. *Journal of Computational physics*, 378:686–707, 2019.
- [28] P. Rathore, W. Lei, Z. Frangella, L. Lu, and M. Udell. Challenges in training pinns: A loss landscape perspective. *arXiv preprint arXiv:2402.01868*, 2024.
- [29] C. Wang, S. Li, D. He, and L. Wang. Is l2 physics informed loss always suitable for training physics informed neural network? *Advances in Neural Information Processing Systems*, 35:8278–8290, 2022.
- [30] S. Wang, B. Li, Y. Chen, and P. Perdikaris. Piratenets: Physics-informed deep learning with residual adaptive networks. *Journal of Machine Learning Research*, 25(402):1–51, 2024.
- [31] P. Wolfe. Convergence conditions for ascent methods. *SIAM review*, 11(2):226–235, 1969.
- [32] C. Wu, M. Zhu, Q. Tan, Y. Kartha, and L. Lu. A comprehensive study of non-adaptive and residual-based adaptive sampling for physics-informed neural networks. *Computer Methods in Applied Mechanics and Engineering*, 403:115671, 2023.
- [33] H. Wu, H. Luo, Y. Ma, J. Wang, and M. Long. Ropinn: Region optimized physics-informed neural networks. *arXiv preprint arXiv:2405.14369*, 2024.
- [34] H. Wu, Y. Ma, H. Zhou, H. Weng, J. Wang, and M. Long. Propinn: Demystifying propagation failures in physics-informed neural networks. *arXiv preprint arXiv:2502.00803*, 2025.
- [35] C. Xu, D. Liu, Y. Hu, J. Li, R. Qin, Q. Zheng, and J. Xiong. Sub-sequential physics-informed learning with state space model. In *International Conference on Machine Learning*. PMLR, 2025.
- [36] Z. Zhao, X. Ding, and B. A. Prakash. Pinnsformer: A transformer-based framework for physics-informed neural networks. In *The Twelfth International Conference on Learning Representations*, 2024.
- [37] N. Zmora, H. Wu, and J. Rodge. Achieving fp32 accuracy for int8 inference using quantization aware training with NVIDIA TensorRT, 2021.

A PDE Setups

A.1 1-D Convection

The one-dimensional convection (or advection) equation characterises the transport of a scalar field $u(x, t)$ —such as temperature, concentration, or momentum—under a uniform velocity β . Widely studied in fluid dynamics and transport theory, it is expressed as

$$\begin{aligned}\frac{\partial u}{\partial t} + \beta \frac{\partial u}{\partial x} &= 0, \quad \forall x \in [0, 2\pi], t \in [0, 1], \\ u(x, 0) &= \sin x, \quad \forall x \in [0, 2\pi], \\ u(0, t) &= u(2\pi, t), \quad \forall t \in [0, 1],\end{aligned}\tag{5}$$

where β denotes the constant advection speed. Larger values of β increase the difficulty for PINNs, making this equation a common benchmark with known failure modes. Following prevailing practice [36, 33], we set $\beta = 50$.

Its analytical solution is

$$u_{\text{ana}}(x, t) = \sin(x - \beta t).\tag{6}$$

A.2 1-D Reaction

The one-dimensional reaction equation models the temporal evolution of a reacting species along a single spatial dimension:

$$\begin{aligned}\frac{\partial u}{\partial t} - \rho u(1 - u) &= 0, \quad \forall x \in [0, 2\pi], t \in [0, 1], \\ u(x, 0) &= \exp\left(-\frac{(x - \pi)^2}{2(\pi/4)^2}\right), \quad \forall x \in [0, 2\pi], \\ u(0, t) &= u(2\pi, t), \quad \forall t \in [0, 1],\end{aligned}\tag{7}$$

with ρ representing the growth-rate coefficient. Increasing ρ likewise poses greater challenges for PINNs; we adopt the standard choice $\rho = 5$ in accordance with [36, 33]. The closed-form solution is

$$u_{\text{ana}} = \frac{\exp\left(-\frac{(x - \pi)^2}{2(\pi/4)^2}\right) \exp(\rho t)}{\exp\left(-\frac{(x - \pi)^2}{2(\pi/4)^2}\right) (\exp(\rho t) - 1) + 1}.\tag{8}$$

A.3 1-D Wave

The one-dimensional wave equation describes wave propagation—for example, vibrations on a string—and in our study takes the form

$$\begin{aligned}\frac{\partial^2 u}{\partial t^2} - 4 \frac{\partial^2 u}{\partial x^2} &= 0, \quad \forall x \in [0, 1], t \in [0, 1], \\ u(x, 0) &= \sin(\pi x) + \frac{1}{2} \sin(\beta \pi x), \quad \forall x \in [0, 1], \\ \frac{\partial u(x, 0)}{\partial t} &= 0, \quad \forall x \in [0, 1], \\ u(0, t) &= u(1, t) = 0, \quad \forall t \in [0, 1],\end{aligned}\tag{9}$$

where β controls the second harmonic; we set $\beta = 3$ following [36, 33]. Because the PDE involves second-order derivatives and the initial condition contains first-order derivatives, optimisation is notoriously difficult [33]. Nevertheless, this case demonstrates that PINNMamba's matrix-defined time differentiation—with uniform scaling across derivative orders—can better capture temporal dynamics.

Its analytical solution is

$$u_{\text{ana}}(x, t) = \sin(\pi x) \cos(2\pi t) + \sin(\beta \pi x) \cos(2\beta \pi t).\tag{10}$$

A.4 Allen–Cahn

The Allen–Cahn equation is a canonical reaction–diffusion benchmark:

$$\begin{aligned}\frac{\partial u}{\partial t} - 0.0001 \frac{\partial^2 u}{\partial x^2} + 5u^3 - 5u &= 0, \quad x \in (-1, 1), \quad t \in (0, 1), \\ u(x, 0) &= x^2 \cos(\pi x), \quad x \in [-1, 1], \\ u(-1, t) &= u(1, t), \quad t \in [0, 1], \\ \frac{\partial u(-1, t)}{\partial x} &= \frac{\partial u(1, t)}{\partial x}, \quad t \in [0, 1].\end{aligned}\tag{11}$$

Because no closed-form analytic solution exists, we adopt a high-resolution spectral approximation [19] as reference data, following [27, 34]. The sharp interface and double entrance makes this PDE a widely used stress test for PINNs and a representative failure mode.

B L-BFGS Optimizer

For an unconstrained, twice-differentiable objective $f : \mathbb{R}^n \rightarrow \mathbb{R}$, Newton’s method uses the exact Hessian $H_k = \nabla^2 f(x_k)$ to obtain the step

$$x_{k+1} = x_k - \alpha_k H_k^{-1} \nabla f(x_k),\tag{12}$$

yielding a quadratic local model and (under standard assumptions) quadratic convergence. Quasi-Newton methods avoid forming or inverting the expensive H_k by maintaining an approximation $B_k \approx H_k$ (or its inverse $H_k \approx B_k^{-1}$) that is updated from first-order information only. The most successful member is the BFGS update

$$H_{k+1} = (I - \rho_k s_k y_k^\top) H_k (I - \rho_k y_k s_k^\top) + \rho_k s_k s_k^\top, \quad \rho_k = \frac{1}{y_k^\top s_k},\tag{13}$$

with $s_k = x_{k+1} - x_k$ and $y_k = \nabla f(x_{k+1}) - \nabla f(x_k)$. BFGS enjoys global convergence and super-linear local convergence, but it stores a dense $n \times n$ matrix—requiring $O(n^2)$ memory and time per iteration—making it impractical when n is large.

L-BFGS eliminates the quadratic storage by discarding all but the most recent m curvature pairs $\{(s_i, y_i)\}_{i=k-m}^{k-1}$. Instead of storing H_k explicitly, it represents the inverse Hessian implicitly through these vectors and a diagonal scaling H_k^0 (often $\gamma_k I$ with $\gamma_k = \frac{s_{k-1}^\top y_{k-1}}{y_{k-1}^\top y_{k-1}}$).

In compact format, let:

$$S_k = [s_{k-m}, \dots, s_{k-1}] \in \mathbb{R}^{n \times m}, \quad Y_k = [y_{k-m}, \dots, y_{k-1}], \quad R_k = S_k^\top Y_k \text{ (upper-triangular)},\tag{14}$$

$$D_k = \text{diag}(s_{k-m}^\top y_{k-m}, \dots, s_{k-1}^\top y_{k-1}).\tag{15}$$

With an initial diagonal scaling $H_k^0 = \gamma_k I$ (usually $\gamma_k = \frac{s_{k-1}^\top y_{k-1}}{y_{k-1}^\top y_{k-1}}$), the inverse Hessian can be written without loss of information as

$$H_k = H_k^0 + [S_k \ H_k^0 Y_k] \begin{bmatrix} R_k^{-\top} (D_k + Y_k^\top H_k^0 Y_k) R_k^{-1} & -R_k^{-\top} \\ -R_k^{-1} & 0 \end{bmatrix} \begin{bmatrix} S_k^\top \\ Y_k^\top H_k^0 \end{bmatrix}.\tag{16}$$

The memory requirement reduces to $O(nm)$ floats, while the computational cost per iteration becomes $O(nm)$ flops. This scaling is linear in the number of parameters when m is fixed, with typical values $m \in [5, 20]$.

C Broader Social Impact

Making FP64 the default for PINNs turns a previously brittle solver into a dependable scientific tool: the same vanilla architecture that failed at FP32 converges robustly on convection, reaction, wave and Allen-Cahn equations when upgraded to FP64. Reliable neural PDE solvers promise faster, mesh-free simulations for climate modelling, renewable-energy design, and personalised biomedical devices, potentially shortening innovation cycles and widening access to high-fidelity analysis in domains where traditional finite-element codes are prohibitively slow.

The work also reframes precision as a first-class hyper-parameter in numerical machine learning, aligning PINNs with classical finite-element practice where double precision is standard for safety-critical engineering. This alignment reduces the risk of deploying under-resolved models in high-stakes settings. Yet the same fidelity could be misapplied—for instance, to accelerate the design of destructive fluid-dynamic systems. We therefore release code under a license prohibiting military applications and encourage transparent reporting of energy usage and precision settings to foster reproducibility and responsible adoption.

D Limitations

However, double precision is more energy-intensive and currently restricted to high-end hardware such as NVIDIA's H100 GPUs, on which our experiments were run. Increased demand for FP64 compute could raise the carbon footprint of scientific machine learning and exacerbate resource inequity between well-funded labs and those with limited budgets. Conversely, our findings may motivate hardware–software co-design efforts that deliver energy-efficient FP64 pipelines, ultimately mitigating this environmental cost.

NeurIPS Paper Checklist

The checklist is designed to encourage best practices for responsible machine learning research, addressing issues of reproducibility, transparency, research ethics, and societal impact. Do not remove the checklist: **The papers not including the checklist will be desk rejected.** The checklist should follow the references and follow the (optional) supplemental material. The checklist does NOT count towards the page limit.

Please read the checklist guidelines carefully for information on how to answer these questions. For each question in the checklist:

- You should answer [Yes], [No], or [NA].
- [NA] means either that the question is Not Applicable for that particular paper or the relevant information is Not Available.
- Please provide a short (1–2 sentence) justification right after your answer (even for NA).

The checklist answers are an integral part of your paper submission. They are visible to the reviewers, area chairs, senior area chairs, and ethics reviewers. You will be asked to also include it (after eventual revisions) with the final version of your paper, and its final version will be published with the paper.

The reviewers of your paper will be asked to use the checklist as one of the factors in their evaluation. While "[Yes]" is generally preferable to "[No]", it is perfectly acceptable to answer "[No]" provided a proper justification is given (e.g., "error bars are not reported because it would be too computationally expensive" or "we were unable to find the license for the dataset we used"). In general, answering "[No]" or "[NA]" is not grounds for rejection. While the questions are phrased in a binary way, we acknowledge that the true answer is often more nuanced, so please just use your best judgment and write a justification to elaborate. All supporting evidence can appear either in the main paper or the supplemental material, provided in appendix. If you answer [Yes] to a question, in the justification please point to the section(s) where related material for the question can be found.

IMPORTANT, please:

- **Delete this instruction block, but keep the section heading "NeurIPS Paper Checklist",**
- **Keep the checklist subsection headings, questions/answers and guidelines below.**
- **Do not modify the questions and only use the provided macros for your answers.**

1. Claims

Question: Do the main claims made in the abstract and introduction accurately reflect the paper's contributions and scope?

Answer: [Yes]

Justification: We attribute the failure modes of PINN to insufficient arithmetic precision, and analyzed the loss landscape of PINN. We include that in the abstract and introduction.

Guidelines:

- The answer NA means that the abstract and introduction do not include the claims made in the paper.
- The abstract and/or introduction should clearly state the claims made, including the contributions made in the paper and important assumptions and limitations. A No or NA answer to this question will not be perceived well by the reviewers.
- The claims made should match theoretical and experimental results, and reflect how much the results can be expected to generalize to other settings.
- It is fine to include aspirational goals as motivation as long as it is clear that these goals are not attained by the paper.

2. Limitations

Question: Does the paper discuss the limitations of the work performed by the authors?

Answer: [Yes]

Justification: We create a section in Appendix to discuss the limitation of the paper.

Guidelines:

- The answer NA means that the paper has no limitation while the answer No means that the paper has limitations, but those are not discussed in the paper.
- The authors are encouraged to create a separate "Limitations" section in their paper.
- The paper should point out any strong assumptions and how robust the results are to violations of these assumptions (e.g., independence assumptions, noiseless settings, model well-specification, asymptotic approximations only holding locally). The authors should reflect on how these assumptions might be violated in practice and what the implications would be.
- The authors should reflect on the scope of the claims made, e.g., if the approach was only tested on a few datasets or with a few runs. In general, empirical results often depend on implicit assumptions, which should be articulated.
- The authors should reflect on the factors that influence the performance of the approach. For example, a facial recognition algorithm may perform poorly when image resolution is low or images are taken in low lighting. Or a speech-to-text system might not be used reliably to provide closed captions for online lectures because it fails to handle technical jargon.
- The authors should discuss the computational efficiency of the proposed algorithms and how they scale with dataset size.
- If applicable, the authors should discuss possible limitations of their approach to address problems of privacy and fairness.
- While the authors might fear that complete honesty about limitations might be used by reviewers as grounds for rejection, a worse outcome might be that reviewers discover limitations that aren't acknowledged in the paper. The authors should use their best judgment and recognize that individual actions in favor of transparency play an important role in developing norms that preserve the integrity of the community. Reviewers will be specifically instructed to not penalize honesty concerning limitations.

3. Theory assumptions and proofs

Question: For each theoretical result, does the paper provide the full set of assumptions and a complete (and correct) proof?

Answer: [NA]

Justification: No theory in mathematical format is proposed in this paper.

Guidelines:

- The answer NA means that the paper does not include theoretical results.
- All the theorems, formulas, and proofs in the paper should be numbered and cross-referenced.
- All assumptions should be clearly stated or referenced in the statement of any theorems.
- The proofs can either appear in the main paper or the supplemental material, but if they appear in the supplemental material, the authors are encouraged to provide a short proof sketch to provide intuition.
- Inversely, any informal proof provided in the core of the paper should be complemented by formal proofs provided in appendix or supplemental material.
- Theorems and Lemmas that the proof relies upon should be properly referenced.

4. Experimental result reproducibility

Question: Does the paper fully disclose all the information needed to reproduce the main experimental results of the paper to the extent that it affects the main claims and/or conclusions of the paper (regardless of whether the code and data are provided or not)?

Answer: [Yes]

Justification: We provide an anonymous link to the code.

Guidelines:

- The answer NA means that the paper does not include experiments.

- If the paper includes experiments, a No answer to this question will not be perceived well by the reviewers: Making the paper reproducible is important, regardless of whether the code and data are provided or not.
- If the contribution is a dataset and/or model, the authors should describe the steps taken to make their results reproducible or verifiable.
- Depending on the contribution, reproducibility can be accomplished in various ways. For example, if the contribution is a novel architecture, describing the architecture fully might suffice, or if the contribution is a specific model and empirical evaluation, it may be necessary to either make it possible for others to replicate the model with the same dataset, or provide access to the model. In general, releasing code and data is often one good way to accomplish this, but reproducibility can also be provided via detailed instructions for how to replicate the results, access to a hosted model (e.g., in the case of a large language model), releasing of a model checkpoint, or other means that are appropriate to the research performed.
- While NeurIPS does not require releasing code, the conference does require all submissions to provide some reasonable avenue for reproducibility, which may depend on the nature of the contribution. For example
 - (a) If the contribution is primarily a new algorithm, the paper should make it clear how to reproduce that algorithm.
 - (b) If the contribution is primarily a new model architecture, the paper should describe the architecture clearly and fully.
 - (c) If the contribution is a new model (e.g., a large language model), then there should either be a way to access this model for reproducing the results or a way to reproduce the model (e.g., with an open-source dataset or instructions for how to construct the dataset).
 - (d) We recognize that reproducibility may be tricky in some cases, in which case authors are welcome to describe the particular way they provide for reproducibility. In the case of closed-source models, it may be that access to the model is limited in some way (e.g., to registered users), but it should be possible for other researchers to have some path to reproducing or verifying the results.

5. Open access to data and code

Question: Does the paper provide open access to the data and code, with sufficient instructions to faithfully reproduce the main experimental results, as described in supplemental material?

Answer: [\[Yes\]](#)

Justification: We provide an anonymous link to the code.

Guidelines:

- The answer NA means that paper does not include experiments requiring code.
- Please see the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- While we encourage the release of code and data, we understand that this might not be possible, so “No” is an acceptable answer. Papers cannot be rejected simply for not including code, unless this is central to the contribution (e.g., for a new open-source benchmark).
- The instructions should contain the exact command and environment needed to run to reproduce the results. See the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- The authors should provide instructions on data access and preparation, including how to access the raw data, preprocessed data, intermediate data, and generated data, etc.
- The authors should provide scripts to reproduce all experimental results for the new proposed method and baselines. If only a subset of experiments are reproducible, they should state which ones are omitted from the script and why.
- At submission time, to preserve anonymity, the authors should release anonymized versions (if applicable).

- Providing as much information as possible in supplemental material (appended to the paper) is recommended, but including URLs to data and code is permitted.

6. Experimental setting/details

Question: Does the paper specify all the training and test details (e.g., data splits, hyper-parameters, how they were chosen, type of optimizer, etc.) necessary to understand the results?

Answer: [\[Yes\]](#)

Justification: In section 3.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The experimental setting should be presented in the core of the paper to a level of detail that is necessary to appreciate the results and make sense of them.
- The full details can be provided either with the code, in appendix, or as supplemental material.

7. Experiment statistical significance

Question: Does the paper report error bars suitably and correctly defined or other appropriate information about the statistical significance of the experiments?

Answer: [\[Yes\]](#)

Justification: We include the CI in the main results in Table 1.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The authors should answer "Yes" if the results are accompanied by error bars, confidence intervals, or statistical significance tests, at least for the experiments that support the main claims of the paper.
- The factors of variability that the error bars are capturing should be clearly stated (for example, train/test split, initialization, random drawing of some parameter, or overall run with given experimental conditions).
- The method for calculating the error bars should be explained (closed form formula, call to a library function, bootstrap, etc.)
- The assumptions made should be given (e.g., Normally distributed errors).
- It should be clear whether the error bar is the standard deviation or the standard error of the mean.
- It is OK to report 1-sigma error bars, but one should state it. The authors should preferably report a 2-sigma error bar than state that they have a 96% CI, if the hypothesis of Normality of errors is not verified.
- For asymmetric distributions, the authors should be careful not to show in tables or figures symmetric error bars that would yield results that are out of range (e.g. negative error rates).
- If error bars are reported in tables or plots, The authors should explain in the text how they were calculated and reference the corresponding figures or tables in the text.

8. Experiments compute resources

Question: For each experiment, does the paper provide sufficient information on the computer resources (type of compute workers, memory, time of execution) needed to reproduce the experiments?

Answer: [\[Yes\]](#)

Justification: In Section 3, our experiments can be conduct on any GPU with double precision support, but Nvidia A100 or H100 is recommended.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The paper should indicate the type of compute workers CPU or GPU, internal cluster, or cloud provider, including relevant memory and storage.

- The paper should provide the amount of compute required for each of the individual experimental runs as well as estimate the total compute.
- The paper should disclose whether the full research project required more compute than the experiments reported in the paper (e.g., preliminary or failed experiments that didn't make it into the paper).

9. Code of ethics

Question: Does the research conducted in the paper conform, in every respect, with the NeurIPS Code of Ethics <https://neurips.cc/public/EthicsGuidelines?>

Answer: [Yes]

Justification: Yes.

Guidelines:

- The answer NA means that the authors have not reviewed the NeurIPS Code of Ethics.
- If the authors answer No, they should explain the special circumstances that require a deviation from the Code of Ethics.
- The authors should make sure to preserve anonymity (e.g., if there is a special consideration due to laws or regulations in their jurisdiction).

10. Broader impacts

Question: Does the paper discuss both potential positive societal impacts and negative societal impacts of the work performed?

Answer: [Yes]

Justification: In Appendix.

Guidelines:

- The answer NA means that there is no societal impact of the work performed.
- If the authors answer NA or No, they should explain why their work has no societal impact or why the paper does not address societal impact.
- Examples of negative societal impacts include potential malicious or unintended uses (e.g., disinformation, generating fake profiles, surveillance), fairness considerations (e.g., deployment of technologies that could make decisions that unfairly impact specific groups), privacy considerations, and security considerations.
- The conference expects that many papers will be foundational research and not tied to particular applications, let alone deployments. However, if there is a direct path to any negative applications, the authors should point it out. For example, it is legitimate to point out that an improvement in the quality of generative models could be used to generate deepfakes for disinformation. On the other hand, it is not needed to point out that a generic algorithm for optimizing neural networks could enable people to train models that generate Deepfakes faster.
- The authors should consider possible harms that could arise when the technology is being used as intended and functioning correctly, harms that could arise when the technology is being used as intended but gives incorrect results, and harms following from (intentional or unintentional) misuse of the technology.
- If there are negative societal impacts, the authors could also discuss possible mitigation strategies (e.g., gated release of models, providing defenses in addition to attacks, mechanisms for monitoring misuse, mechanisms to monitor how a system learns from feedback over time, improving the efficiency and accessibility of ML).

11. Safeguards

Question: Does the paper describe safeguards that have been put in place for responsible release of data or models that have a high risk for misuse (e.g., pretrained language models, image generators, or scraped datasets)?

Answer: [NA]

Justification: The paper poses no such risks.

Guidelines:

- The answer NA means that the paper poses no such risks.

- Released models that have a high risk for misuse or dual-use should be released with necessary safeguards to allow for controlled use of the model, for example by requiring that users adhere to usage guidelines or restrictions to access the model or implementing safety filters.
- Datasets that have been scraped from the Internet could pose safety risks. The authors should describe how they avoided releasing unsafe images.
- We recognize that providing effective safeguards is challenging, and many papers do not require this, but we encourage authors to take this into account and make a best faith effort.

12. Licenses for existing assets

Question: Are the creators or original owners of assets (e.g., code, data, models), used in the paper, properly credited and are the license and terms of use explicitly mentioned and properly respected?

Answer: [\[Yes\]](#)

Justification: Use the code of PINNMamba, which is on CC-BY 4.0 and MIT.

Guidelines:

- The answer NA means that the paper does not use existing assets.
- The authors should cite the original paper that produced the code package or dataset.
- The authors should state which version of the asset is used and, if possible, include a URL.
- The name of the license (e.g., CC-BY 4.0) should be included for each asset.
- For scraped data from a particular source (e.g., website), the copyright and terms of service of that source should be provided.
- If assets are released, the license, copyright information, and terms of use in the package should be provided. For popular datasets, paperswithcode.com/datasets has curated licenses for some datasets. Their licensing guide can help determine the license of a dataset.
- For existing datasets that are re-packaged, both the original license and the license of the derived asset (if it has changed) should be provided.
- If this information is not available online, the authors are encouraged to reach out to the asset's creators.

13. New assets

Question: Are new assets introduced in the paper well documented and is the documentation provided alongside the assets?

Answer: [\[NA\]](#)

Justification: The paper does not release new assets.

Guidelines:

- The answer NA means that the paper does not release new assets.
- Researchers should communicate the details of the dataset/code/model as part of their submissions via structured templates. This includes details about training, license, limitations, etc.
- The paper should discuss whether and how consent was obtained from people whose asset is used.
- At submission time, remember to anonymize your assets (if applicable). You can either create an anonymized URL or include an anonymized zip file.

14. Crowdsourcing and research with human subjects

Question: For crowdsourcing experiments and research with human subjects, does the paper include the full text of instructions given to participants and screenshots, if applicable, as well as details about compensation (if any)?

Answer: [\[NA\]](#)

Justification: The paper does not involve crowdsourcing nor research with human subjects

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Including this information in the supplemental material is fine, but if the main contribution of the paper involves human subjects, then as much detail as possible should be included in the main paper.
- According to the NeurIPS Code of Ethics, workers involved in data collection, curation, or other labor should be paid at least the minimum wage in the country of the data collector.

15. Institutional review board (IRB) approvals or equivalent for research with human subjects

Question: Does the paper describe potential risks incurred by study participants, whether such risks were disclosed to the subjects, and whether Institutional Review Board (IRB) approvals (or an equivalent approval/review based on the requirements of your country or institution) were obtained?

Answer: [NA]

Justification: The paper does not involve crowdsourcing nor research with human subjects

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Depending on the country in which research is conducted, IRB approval (or equivalent) may be required for any human subjects research. If you obtained IRB approval, you should clearly state this in the paper.
- We recognize that the procedures for this may vary significantly between institutions and locations, and we expect authors to adhere to the NeurIPS Code of Ethics and the guidelines for their institution.
- For initial submissions, do not include any information that would break anonymity (if applicable), such as the institution conducting the review.

16. Declaration of LLM usage

Question: Does the paper describe the usage of LLMs if it is an important, original, or non-standard component of the core methods in this research? Note that if the LLM is used only for writing, editing, or formatting purposes and does not impact the core methodology, scientific rigorousness, or originality of the research, declaration is not required.

Answer: [NA]

Justification: The core method development in this research does not involve LLMs as any important, original, or non-standard components.

Guidelines:

- The answer NA means that the core method development in this research does not involve LLMs as any important, original, or non-standard components.
- Please refer to our LLM policy (<https://neurips.cc/Conferences/2025/LLM>) for what should or should not be described.