

Technical Support Center: Accelerating PINN Training with HPC

Author: BenchChem Technical Support Team. **Date:** April 2026

Compound of Interest

Compound Name: *Pdic-NN*

Cat. No.: *B15614678*

[Get Quote](#)

This technical support center provides troubleshooting guidance and answers to frequently asked questions for researchers, scientists, and drug development professionals who are using High-Performance Computing (HPC) to accelerate the training of Physics-Informed Neural Networks (PINNs).

Troubleshooting Guide

Q: My PINN training is extremely slow, even on an HPC cluster. What are the common bottlenecks?

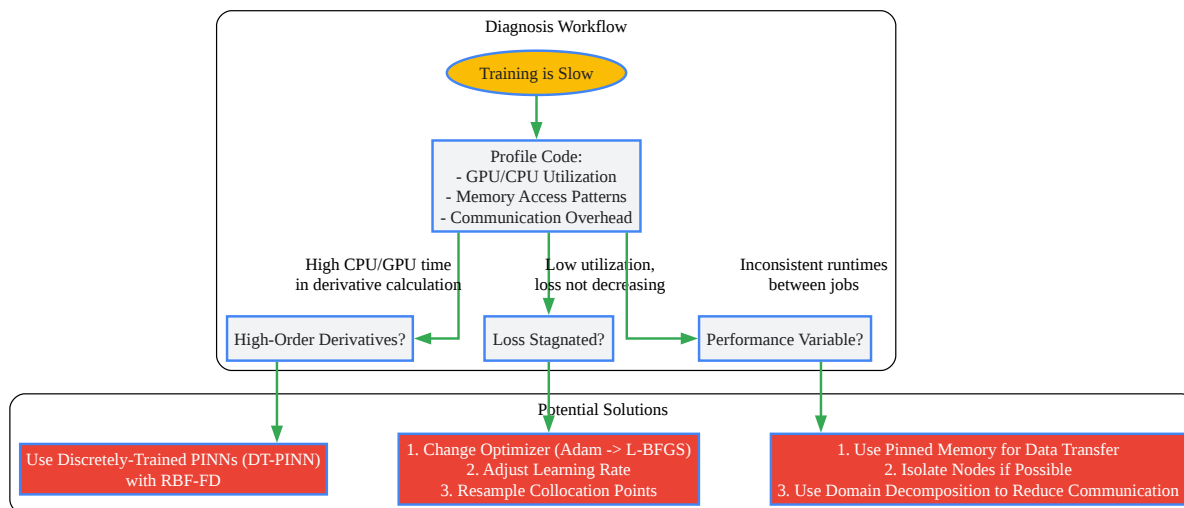
A: Slow training on HPC systems can stem from various bottlenecks that are not always immediately obvious. Common issues include:

- **Computational Cost of Automatic Differentiation (AD):** The repeated calculation of partial derivatives in the loss function via automatic differentiation is a primary cause of slowdowns, especially for PDEs with higher-order derivatives.^{[1][2][3]} This is a well-known computational expense in PINN training.^[1]
- **System-Level Bottlenecks:** When scaling deep learning workloads, you can encounter different bottlenecks related to memory capacity, communication overhead between nodes,

I/O limitations for large datasets, or even the compute capabilities for specific operations.[4]

- HPC Environment Variability: Performance on HPC clusters can be inconsistent. Other jobs running on the cluster can interfere with yours, especially if they are network-intensive, causing contention for shared resources.[5]
- Ill-Conditioned Loss Landscape: The loss function in PINNs can be difficult to minimize, a problem known as ill-conditioning, which is often caused by the differential operators in the PDE residual.[6][7][8] This can lead to slow convergence for gradient-based optimizers.[8]

Here is a workflow to diagnose and address training bottlenecks:



[Click to download full resolution via product page](#)

Caption: Workflow for diagnosing and resolving slow PINN training.

Q: My training process is failing with "out of memory" errors on the GPU. What should I do?

A: Out-of-memory errors are common when dealing with large models or complex domains. Here are some strategies:

- **Reduce Batch Size:** This is the simplest approach, but it may affect convergence speed and stability.
- **Use Domain Decomposition:** Techniques like Conservative PINNs (cPINNs) and eXtended PINNs (XPINNs) break the computational domain into smaller subdomains.^[9] Each subdomain is assigned its own smaller neural network, reducing the memory footprint on a single GPU and allowing for parallel training.^{[9][10]}
- **Optimize Data Transfers:** For HPC environments with GPUs, using pinned (or page-locked) memory can significantly accelerate data transfers between the CPU and GPU.^[11] In CUDA, this can be done with `cudaMallocHost` or `cudaHostAlloc`.^[11] However, be cautious not to overallocate pinned memory, as this can reduce the amount of memory available to the operating system and lead to instability.^[11]
- **Model Parallelism:** For very large models, consider model parallelism, where different parts of the neural network are placed on different GPUs. This is more complex to implement than data parallelism but can be effective for memory-intensive models.

Q: The accuracy of my PINN is poor, or the training fails to converge. How can I fix this?

A: Poor accuracy or convergence failure often points to issues with the loss landscape, network architecture, or sampling strategy.

- **Optimizer Choice:** Standard optimizers like Adam can struggle with the ill-conditioned loss landscapes of PINNs.^{[7][8]} A common and effective strategy is to begin training with Adam

and then switch to a quasi-Newton method like L-BFGS for fine-tuning.[12] The combination of Adam followed by L-BFGS has been shown to be superior to using either one alone.[6][7]

- **Network Architecture:** Deeper networks are more prone to vanishing or exploding gradients, a problem that is amplified in PINNs due to the multiple orders of differentiation required.[13] It is often better to use shallower, wider networks (e.g., 3-4 layers with 256 nodes each) as the patterns PINNs need to learn are often simpler than those in fields like computer vision. [13]
- **Activation Functions:** The choice of activation function is critical because it is differentiated multiple times.[13] Ensure the function has at least $n+1$ non-zero derivatives, where n is the order of the PDE.[13]
- **Adaptive Sampling:** Instead of a fixed set of collocation points, use an adaptive sampling method. These techniques progressively add more points to areas where the model has a higher error (i.e., high PDE residual).[14][15] This focuses the network's attention on the most difficult parts of the domain.[14]
- **Input Normalization:** Normalizing all inputs (spatial and temporal) to a range like $[-1, 1]$ at the beginning of the network can improve accuracy and training stability.[12][16]

Frequently Asked Questions (FAQs)

Q: What is Domain Decomposition for PINNs and how does it accelerate training?

A: Domain decomposition is a "divide and conquer" strategy that breaks a large, complex computational domain into multiple smaller, simpler subdomains.[10][17] For PINNs, this involves training a separate, smaller neural network for each subdomain.[9][17] These networks are trained in parallel, and consistency is enforced by adding interface conditions to the loss function that ensure the solutions match at the boundaries between subdomains.[10]

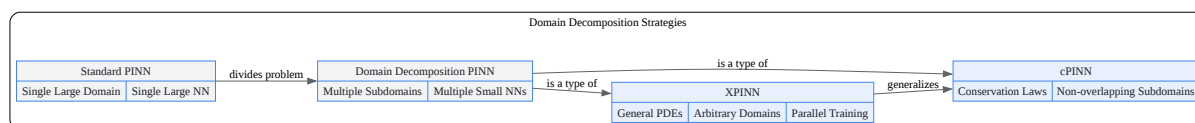
This approach accelerates training in several ways:

- **Parallelization:** Each subdomain's network can be trained independently and in parallel on different compute nodes or GPUs, which is a natural fit for HPC architectures.[9]

- Improved Accuracy: By using separate networks for different regions, the model can better capture complex or localized solution features, which can improve overall accuracy.[10][17]
- Reduced Complexity: Solving multiple smaller problems can be more computationally tractable and stable than solving one large, complex one.[10]

Popular domain decomposition frameworks include:

- cPINNs (Conservative PINNs): Uses a non-overlapping Schwarz-based decomposition approach.[9]
- XPINNs (eXtended PINNs): Extends the cPINN methodology to more general PDEs and arbitrary space-time domains, associating each subdomain with its own sub-PINN.[9]



[Click to download full resolution via product page](#)

Caption: Relationship between PINN and Domain Decomposition methods.

Q: Are there faster alternatives to Automatic Differentiation for calculating PDE residuals?

A: Yes. While automatic differentiation (AD) is a core component of vanilla PINNs, it can be computationally expensive.[1] A powerful alternative is to use Discretely-Trained PINNs (DT-PINNs).[1][2]

DT-PINNs replace the exact spatial derivatives computed by AD with high-order numerical discretizations.[1][3] A common method is to use meshless radial basis function-finite

differences (RBF-FD), which can be applied via sparse matrix-vector multiplication.[\[1\]\[2\]](#) This approach is effective even for irregular domain geometries.[\[1\]\[3\]](#)

Technique	Derivative Calculation	Precision	Relative Speed
Vanilla PINN	Automatic Differentiation (AD)	32-bit (fp32)	1x (Baseline)
DT-PINN	Numerical Discretization (RBF-FD)	64-bit (fp64)	2-4x Faster [1][3]

Table 1: Comparison of Vanilla PINN and DT-PINN performance. DT-PINNs can achieve similar or better accuracy with significantly faster training times on a GPU.[\[1\]\[3\]](#)

Q: How should I choose and distribute my collocation points for optimal performance?

A: The distribution of training points (collocation points) is critical for both accuracy and convergence speed.[\[18\]\[19\]](#)

- Fixed Sampling: Simple methods like grid sampling or random sampling using techniques like Latin Hypercube Sampling are easy to implement but may not be efficient, as they can ignore important features of the PDE's solution.[\[15\]\[18\]\[19\]](#)
- Adaptive Sampling: More advanced strategies involve adapting the point distribution during training. A highly effective approach is to add more collocation points in regions where the PDE residual is highest, forcing the model to improve in areas where it performs poorly.[\[14\]](#) Adversarial training can also be used to find these "failure regions" and generate new samples there.[\[15\]](#)
- Re-sampling: For random sampling methods, re-sampling the points at each iteration is a cheap operation that helps ensure the entire domain is covered over the course of training and can better capture localized features.[\[13\]](#)

Experimental Protocols

Protocol: Evaluating Training Acceleration with eXtended PINNs (XPINN)

This protocol outlines the methodology for comparing the performance of a standard "vanilla" PINN against an XPINN that uses domain decomposition.

1. Objective: To quantify the reduction in training time and improvement in accuracy when using an XPINN compared to a vanilla PINN for solving a complex PDE on an HPC cluster.

2. Methodology:

- Problem Definition: Select a challenging 2D or 3D PDE, such as the Navier-Stokes equations for fluid dynamics, defined over a large computational domain.[\[17\]](#)
- Vanilla PINN Setup:
 - Construct a single, large feed-forward neural network to approximate the solution over the entire domain.
 - Define the loss function based on the PDE residual, boundary conditions, and initial conditions across the whole domain.
- XPINN Setup:
 - Decompose the computational domain into N smaller, non-overlapping subdomains.[\[9\]](#)[\[17\]](#)
 - For each subdomain i, instantiate a separate, smaller neural network (sub-PINN).[\[9\]](#)
 - Define a loss function for each sub-PINN that includes the PDE residual and boundary/initial conditions relevant to that subdomain.
 - Add interface loss terms that enforce continuity of the solution and its derivatives between adjacent subdomains.[\[10\]](#)
- Training and Parallelization:

- For the XPINN, assign each of the N sub-PINNs to a separate GPU or compute node in the HPC cluster for parallel training.[9]
- Train the vanilla PINN on a single, comparable compute node.
- Use an identical optimization strategy for both models (e.g., Adam for 10k iterations followed by L-BFGS) to ensure a fair comparison.[12]
- Data Collection:
 - Record the total wall-clock time required for each model to reach a target loss value.
 - After training, calculate the final prediction error (e.g., L2 error) for both models against a known analytical solution or a high-fidelity numerical simulation.
 - Monitor GPU utilization, memory usage, and inter-node communication overhead during the training process.

3. Expected Outcome: The XPINN approach is expected to demonstrate significantly reduced training time due to parallelization and potentially higher accuracy, as the ensemble of smaller networks can better approximate a complex solution.[10][17]

Need Custom Synthesis?

BenchChem offers custom synthesis for rare earth carbides and specific isotopic labeling.

Email: info@benchchem.com or [Request Quote Online](#).

References

- [1. proceedings.neurips.cc](https://proceedings.neurips.cc) [proceedings.neurips.cc]
- [2. researchgate.net](https://researchgate.net) [researchgate.net]
- [3. \[2205.09332\] Accelerated Training of Physics-Informed Neural Networks \(PINNs\) using Meshless Discretizations](#) [arxiv.org]
- [4. m.youtube.com](https://m.youtube.com) [m.youtube.com]
- [5. repositum.tuwien.at](https://repositum.tuwien.at) [repositum.tuwien.at]

- 6. [2402.01868] Challenges in Training PINNs: A Loss Landscape Perspective [arxiv.org]
- 7. arxiv.org [arxiv.org]
- 8. Challenges in Training PINNs: A Loss Landscape Perspective [arxiv.org]
- 9. epubs.siam.org [epubs.siam.org]
- 10. pubs.aip.org [pubs.aip.org]
- 11. What are the best practices for using pinned memory in a high-performance computing environment? - Massed Compute [massedcompute.com]
- 12. medium.com [medium.com]
- 13. towardsdatascience.com [towardsdatascience.com]
- 14. hpcwire.com [hpcwire.com]
- 15. scispace.com [scispace.com]
- 16. medium.com [medium.com]
- 17. pubs.aip.org [pubs.aip.org]
- 18. Strategies for training point distributions in physics-informed neural networks [arxiv.org]
- 19. arxiv.org [arxiv.org]
- To cite this document: BenchChem. [Technical Support Center: Accelerating PINN Training with HPC]. BenchChem, [2026]. [Online PDF]. Available at: [https://www.benchchem.com/product/b15614678/docs#technical-support-center-accelerating-pinn-training-with-hpc]

Disclaimer & Data Validity:

The information provided in this document is for Research Use Only (RUO) and is strictly not intended for diagnostic or therapeutic procedures. While BenchChem strives to provide accurate protocols, we make no warranties, express or implied, regarding the fitness of this product for every specific experimental setup.

Technical Support: The protocols provided are for reference purposes. Unsure if this reagent suits your experiment?

Need Industrial/Bulk Grade? [Request Custom Synthesis Quote](#)

BenchChem

Our mission is to be the trusted global source of essential and advanced chemicals, empowering scientists and researchers to drive progress in science and industry.

Contact

Address: 3281 E Guasti Rd

Ontario, CA 91761, United States

Phone: (601) 213-4426

Email: info@benchchem.com

[Contact our Ph.D. Support Team for a compatibility check](#)