

Technical Support Center: Troubleshooting Slow Convergence in PINN Training

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

Compound of Interest

Compound Name: *Pdic-NN*

Cat. No.: *B15614678*

[Get Quote](#)

This guide provides troubleshooting steps and frequently asked questions to address slow convergence during the training of Physics-Informed Neural Networks (PINNs). It is intended for researchers, scientists, and professionals in drug development who utilize PINNs in their experiments.

Frequently Asked Questions (FAQs)

Q1: My PINN training is extremely slow and the loss is not decreasing. What are the common causes?

Slow or stalled convergence in PINN training can often be attributed to several factors:

- **Gradient Pathologies:** A significant issue arises from unbalanced gradients between the different terms in your composite loss function (e.g., the PDE residual loss and the boundary/initial condition loss).^{[1][2][3]} This "numerical stiffness" can cause the training to focus on one loss term while neglecting others, leading to poor overall convergence.^{[1][2]}
- **Inappropriate Loss Weighting:** Statically assigning equal or arbitrary weights to different loss components can lead to an imbalance in their contributions to the total loss, hindering effective training.^{[4][5][6]}

- **Suboptimal Neural Network Architecture:** The depth, width, and connectivity of your neural network can impact its ability to learn the solution to the PDE. Very deep networks can suffer from vanishing or exploding gradients.[\[7\]](#)[\[8\]](#)
- **Poor Choice of Activation Function:** The activation function plays a critical role, especially in PINNs where its derivatives are used to compute the PDE residual.[\[9\]](#)[\[10\]](#)[\[11\]](#) An activation function that is not sufficiently differentiable or is ill-suited to the problem can impede learning.[\[9\]](#)
- **Inefficient Optimizer:** The choice of optimization algorithm can significantly affect convergence speed and the final accuracy of the model.[\[12\]](#)[\[13\]](#)[\[14\]](#)[\[15\]](#)
- **Training Point Distribution:** The way collocation points are sampled across the domain can impact the accuracy and convergence of the training process.[\[16\]](#)[\[17\]](#)

Q2: How can I diagnose if I have a problem with unbalanced gradients?

A common symptom of unbalanced gradients is observing one component of your loss function (e.g., boundary condition loss) decreasing while another (e.g., PDE residual loss) remains stagnant or even increases. You can diagnose this by:

- **Monitoring Individual Loss Components:** Plot the values of each loss term separately during training. If they are on vastly different scales or one dominates the others, it's a sign of imbalance.
- **Visualizing Gradient Statistics:** More advanced techniques involve analyzing the back-propagated gradients for each loss term.[\[1\]](#) Histograms of these gradients can reveal if one set of gradients is consistently larger or smaller than the others.[\[18\]](#)

Q3: What strategies can I use to address slow convergence?

Here are several strategies you can employ, often in combination, to improve the convergence of your PINN training:

- **Adaptive Loss Weighting:** Instead of using fixed weights for your loss terms, employ an adaptive weighting scheme. These methods dynamically adjust the weights during training to balance the contribution of each loss component.[4][5][19]
- **Choosing the Right Activation Function:** Consider using adaptive activation functions, which introduce a scalable hyperparameter that can be optimized during training to improve convergence.[20][21][22] Ensure your activation function is sufficiently differentiable for the order of your PDE.[9]
- **Selecting an Appropriate Optimizer:** While Adam is a common starting point, quasi-Newton methods like L-BFGS can be very effective, especially in later stages of training, as they utilize second-order information.[12][13][14] A hybrid approach, starting with Adam and switching to L-BFGS, is often beneficial.
- **Refining the Neural Network Architecture:** Experiment with different network architectures. Shallow and wide networks have been shown to outperform deep and narrow ones in some cases.[8] Incorporating residual connections can also improve gradient flow.[9]
- **Adaptive Learning Rate:** Use a learning rate scheduler, such as ReduceLROnPlateau, which reduces the learning rate when the loss plateaus, allowing for more fine-grained adjustments as training progresses.[9]
- **Adaptive Collocation Point Sampling:** Instead of a fixed grid of collocation points, consider adaptive sampling strategies that place more points in regions where the PDE residual is high.[16][17]

Troubleshooting Guides

Guide 1: Implementing Adaptive Loss Weighting

Problem: The loss for the boundary conditions is decreasing, but the PDE residual loss is stuck.

Methodology:

- **Conceptual Framework:** The core idea is to dynamically update the weights of each loss term based on their training behavior. One approach is to use the magnitude of the gradients

for each loss term to balance their influence.

- Experimental Protocol:
 - At each training step, calculate the gradients of the total loss with respect to the neural network parameters.
 - Also, calculate the gradients for each individual loss component (PDE residual, boundary conditions, etc.).
 - Use these individual gradient statistics to compute a scaling factor for each loss weight. A common technique involves using the inverse of the mean or max of the gradients for each loss term to normalize their magnitudes.
 - Update the loss weights at a certain frequency (e.g., every 100 or 1000 iterations).
 - Monitor the individual loss terms to ensure they are all decreasing over time.

Guide 2: Leveraging Adaptive Activation Functions

Problem: Training is slow from the very beginning, and the network struggles to learn even simple functions.

Methodology:

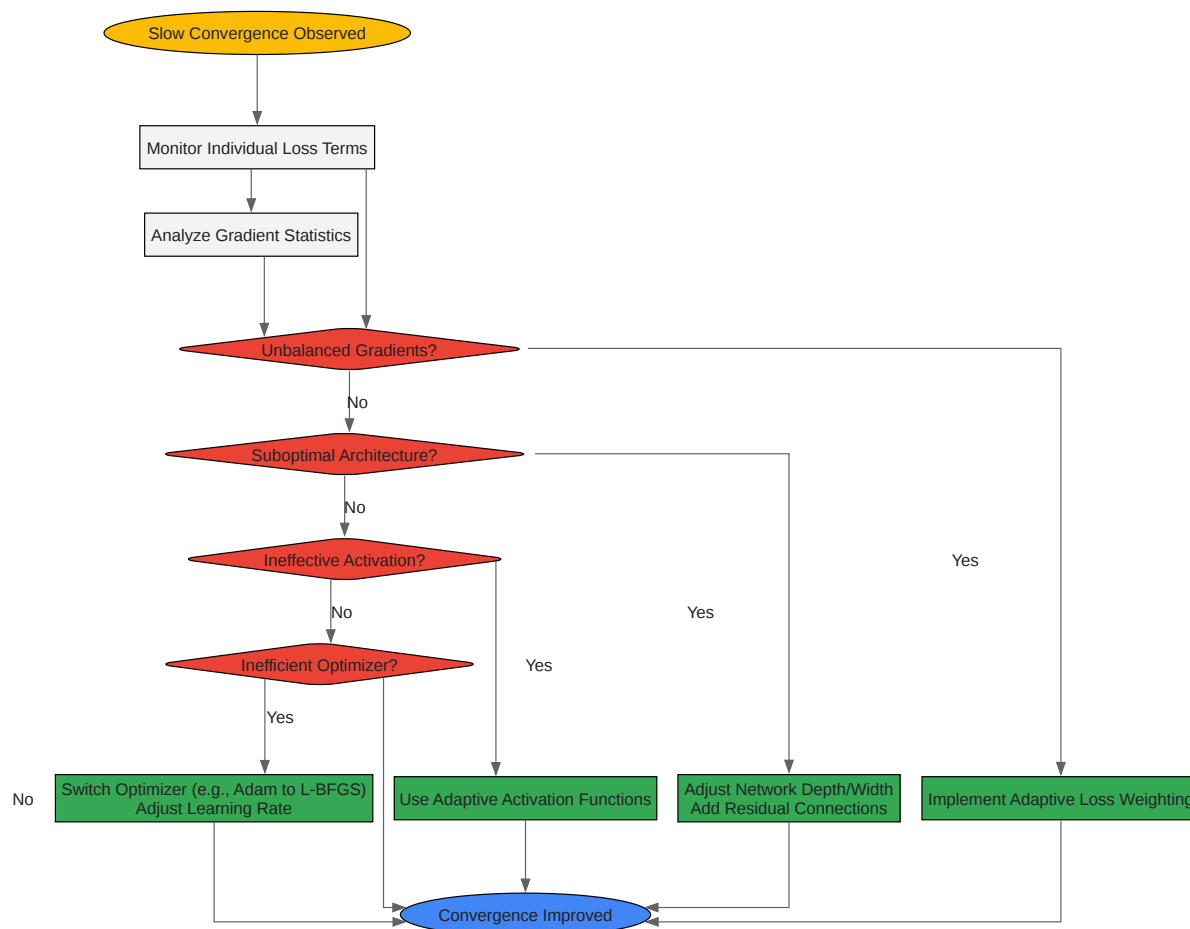
- Conceptual Framework: Introduce a learnable parameter into the activation function. This allows the network to adjust the shape of the activation function during training, which can lead to a more favorable loss landscape and faster convergence.[\[20\]](#)[\[21\]](#)
- Experimental Protocol:
 - Modify your chosen activation function (e.g., tanh or swish) to include a scalable hyperparameter. For example, $\tanh(a * x)$, where 'a' is a trainable parameter.[\[23\]](#)
 - This parameter can be global (one 'a' for the entire network), layer-wise (one 'a' per layer), or neuron-wise (one 'a' for each neuron).[\[22\]](#) A layer-wise approach often provides a good balance between expressiveness and complexity.

- Initialize this parameter (e.g., to 1.0) and allow it to be updated by the optimizer along with the other network weights and biases.
- Compare the convergence rate and final accuracy against a network with a fixed activation function. Studies have shown this can significantly improve the convergence rate, especially in the early stages of training.[\[21\]](#)[\[24\]](#)

Quantitative Data Summary

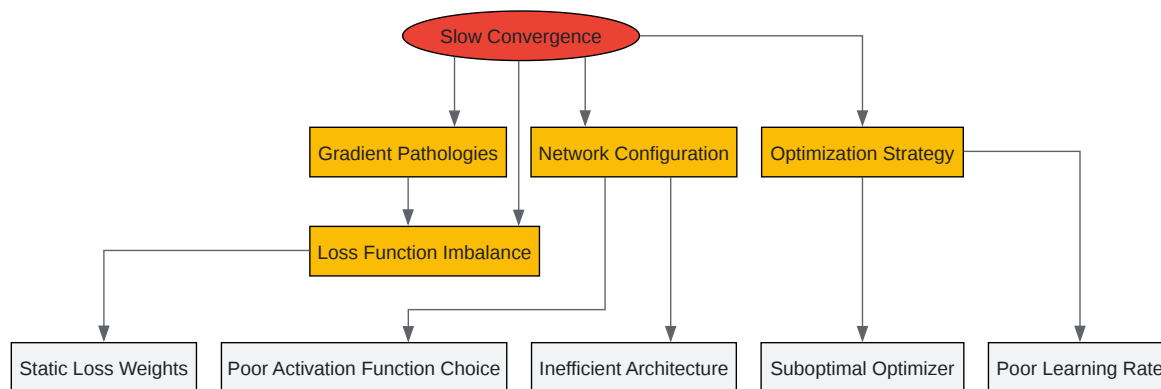
Strategy	Description	Potential Impact on Convergence	Key Hyperparameters
Optimizer Selection	Choice of algorithm to update network weights.	Adam is often good for initial exploration, while L-BFGS can achieve faster convergence near a minimum. [12] [14] A hybrid approach is often effective.	Learning rate, beta1, beta2 (for Adam).
Adaptive Activation Functions	Introducing a learnable parameter in the activation function.	Can significantly accelerate convergence, especially in early training stages. [20] [21] [24]	Initial value of the adaptive parameter, scope (global, layer-wise, or neuron-wise).
Loss Weighting	Method for balancing different loss components.	Adaptive weighting can prevent training from getting stuck by balancing the interplay between different loss terms. [4] [5]	Update frequency of weights, scaling factors.
Network Architecture	Depth and width of the neural network.	Shallow-wide networks may outperform deep-narrow ones for some problems. [8] Residual connections can improve gradient flow. [9]	Number of layers, number of neurons per layer.

Visualizations



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

Caption: A workflow diagram for troubleshooting slow convergence in PINNs.



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

Caption: Logical relationships between causes of slow PINN convergence.

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. [PDF] Understanding and mitigating gradient pathologies in physics-informed neural networks | Semantic Scholar [semanticscholar.org]
- 2. epubs.siam.org [epubs.siam.org]
- 3. researchgate.net [researchgate.net]

- 4. [2104.06217] Self-adaptive loss balanced Physics-informed neural networks for the incompressible Navier-Stokes equations [arxiv.org]
- 5. Dynamic Weight Strategy of Physics-Informed Neural Networks for the 2D Navier–Stokes Equations - PMC [pmc.ncbi.nlm.nih.gov]
- 6. Impact of Loss Weight and Model Complexity on Physics-Informed Neural Networks for Computational Fluid Dynamics [arxiv.org]
- 7. Brain-Inspired Physics-Informed Neural Networks: Bare-Minimum Neural Architectures for PDE Solvers [arxiv.org]
- 8. mdpi.com [mdpi.com]
- 9. medium.com [medium.com]
- 10. [2308.04073] Learning Specialized Activation Functions for Physics-informed Neural Networks [arxiv.org]
- 11. global-sci.com [global-sci.com]
- 12. wias-berlin.de [wias-berlin.de]
- 13. researchgate.net [researchgate.net]
- 14. Which Optimizer Works Best for Physics-Informed Neural Networks and Kolmogorov-Arnold Networks? [arxiv.org]
- 15. datascience.stackexchange.com [datascience.stackexchange.com]
- 16. Strategies for training point distributions in physics-informed neural networks [arxiv.org]
- 17. researchgate.net [researchgate.net]
- 18. app.icerm.brown.edu [app.icerm.brown.edu]
- 19. Improved physics-informed neural network in mitigating gradient-related failures [arxiv.org]
- 20. Adaptive Activation Functions Accelerate Convergence in Deep and Physics-informed Neural Networks (Journal Article) | OSTI.GOV [osti.gov]
- 21. [1906.01170] Adaptive activation functions accelerate convergence in deep and physics-informed neural networks [arxiv.org]
- 22. pubs.aip.org [pubs.aip.org]
- 23. iccs-meeting.org [iccs-meeting.org]
- 24. researchgate.net [researchgate.net]
- To cite this document: BenchChem. [Technical Support Center: Troubleshooting Slow Convergence in PINN Training]. BenchChem, [2026]. [Online PDF]. Available at:

[\[https://www.benchchem.com/product/b15614678/docs#technical-support-center-troubleshooting-slow-convergence-in-pinn-training\]](https://www.benchchem.com/product/b15614678/docs#technical-support-center-troubleshooting-slow-convergence-in-pinn-training)

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](#)