

Technical Support Center: Troubleshooting Vanishing Gradients with Sigmoid Functions

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

Compound of Interest

Compound Name: *Sigmodal*
CAS No.: *1216-40-6*
Cat. No.: *B074413*

[Get Quote](#)

This guide provides troubleshooting steps and frequently asked questions for researchers, scientists, and drug development professionals encountering the vanishing gradient problem, particularly when using sigmoid activation functions in deep neural networks.

Frequently Asked Questions (FAQs) & Troubleshooting Guides

Q1: My deep neural network's accuracy is not improving, and the loss has plateaued very early in training. Why is this happening?

A: This is a classic symptom of the vanishing gradient problem, especially common in deep networks using sigmoid or tanh activation functions.^[1] During backpropagation, the algorithm used to train neural networks, gradients of the loss function are calculated to update the network's weights.^{[1][2]} The sigmoid function squashes its input into a range between 0 and 1.^[3] A critical issue is that its derivative is always small (at most 0.25) and approaches zero for inputs that are large in magnitude (either positive or negative).^{[4][5]}

In a deep network, these small derivatives are multiplied together through many layers.^[2]^[6] This repeated multiplication causes the gradient to shrink exponentially as it propagates backward to the initial layers.^[7]^[8] Consequently, the weights of the early layers are updated by minuscule amounts, causing them to learn extremely slowly or not at all.^[6]^[9] This effectively "stalls" the learning process.^[9]

Q2: How can I concretely diagnose if my network is suffering from vanishing gradients?

A: You can diagnose this issue by monitoring the magnitude of the gradients for weights in different layers during training.

- **Log Gradient Norms:** A common technique is to compute and log the L2 norm of the gradients for each layer's weight matrix after each training batch.
- **Observe the Trend:** If you plot these norms over training epochs, a network with vanishing gradients will show significantly smaller gradient norms for the earlier layers (closer to the input) compared to the later layers (closer to the output).^[3] The norms for the early layers will be close to zero and will not change much, indicating that learning has stopped.^[9]

Here is a logical workflow for diagnosing and addressing the problem:



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

A logical workflow for troubleshooting vanishing gradients.

Q3: What are the most effective solutions to the vanishing gradient problem?

A: Several effective techniques can be used, often in combination. The simplest and most common solution is to replace sigmoid activations in hidden layers with non-saturating alternatives.^{[5][10]}

- **Use Non-Saturating Activation Functions:** The Rectified Linear Unit (ReLU) and its variants (Leaky ReLU, ELU) are the standard choices in modern deep learning.^{[6][11]} For positive inputs, the ReLU function has a derivative of 1, which prevents the gradient from shrinking as it propagates backward.^[12] This allows for faster and more effective training of deep networks.^[13]
- **Proper Weight Initialization:** Poor weight initialization can exacerbate the vanishing gradient problem.^[2] For sigmoid and tanh activations, Xavier (or Glorot) initialization is recommended as it helps keep the signal variance consistent across layers.^{[6][14]} For ReLU-based networks, He initialization is the preferred method.^{[15][16][17]}
- **Batch Normalization:** This technique normalizes the inputs to each layer for each mini-batch.^{[18][19]} It helps by stabilizing the distribution of activation values, which reduces the likelihood of them falling into the saturated (near-zero gradient) regions of the sigmoid function.^{[20][21]} This makes the network more robust to weight initialization and allows for higher learning rates.^{[20][22]}
- **Architectural Changes (Residual Networks):** For very deep networks, even ReLU can have issues. Residual Networks (ResNets) introduce "skip connections" that allow the gradient to flow directly across layers.^{[7][23][24]} This creates an unimpeded path for the gradient to propagate back to the early layers, effectively mitigating the vanishing gradient problem.^{[25][26]}
- **Gating Mechanisms (for RNNs):** In recurrent neural networks (RNNs), Long Short-Term Memory (LSTM) and Gated Recurrent Unit (GRU) architectures use gating mechanisms to control the flow of information and gradients through time, which explicitly addresses the vanishing gradient issue in sequential data.^{[27][28][29][30][31]}

Data Presentation: Comparison of Solutions

The table below summarizes the primary methods to combat vanishing gradients.

Technique	Mechanism of Action	Primary Use Case	Advantages	Considerations
ReLU Activation	Derivative is 1 for positive inputs, preventing gradient shrinkage.[12]	General replacement for sigmoid/tanh in hidden layers.	Computationally efficient, mitigates vanishing gradients.[11][13]	Can suffer from "dying ReLU" problem where neurons become inactive.[32]
Xavier/Glorot Init.	Sets initial weights to maintain variance of activations and back-propagated gradients.[17]	Networks using sigmoid or tanh activation functions.[14]	Reduces the chance of gradients vanishing or exploding early in training.[10]	Less effective for ReLU-based networks.
He Initialization	Accounts for the non-zero centered nature of ReLU activations.[17]	Networks using ReLU or its variants.[16]	Helps maintain signal propagation in deep ReLU networks.	Not designed for saturating activation functions.
Batch Normalization	Normalizes layer inputs to have zero mean and unit variance.[6][19]	Can be used with any activation function in most network types.	Stabilizes and accelerates training, reduces dependence on initialization.[19][20]	Adds computational overhead; behavior can differ between training and inference.
Residual Connections	Creates "shortcuts" for the gradient to bypass layers.[23][24]	Very deep convolutional neural networks (e.g., ResNet).	Allows for the training of extremely deep models (100+ layers).[7][25]	Adds architectural complexity.
Gradient Clipping	Caps the maximum value	Primarily used to prevent exploding	Prevents unstable weight updates from	Does not directly solve the cause

of the gradient
norm.[33]

gradients, but
can help
maintain a stable
gradient flow.[34]

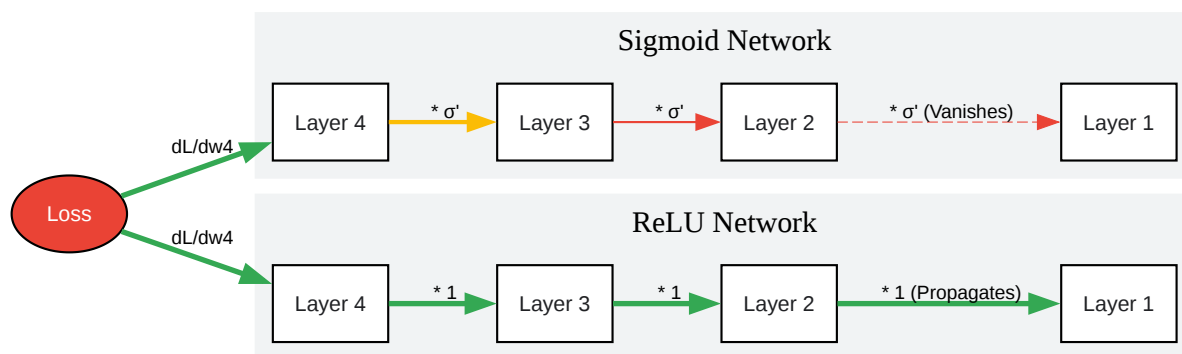
large gradients.
[8]

of vanishing
gradients.[35]

Visualizing the Problem and Solutions

Gradient Flow in Deep Networks

During backpropagation, the gradient is successively multiplied by the derivative of each layer's activation function. With sigmoid, this leads to an exponential decay of the signal.

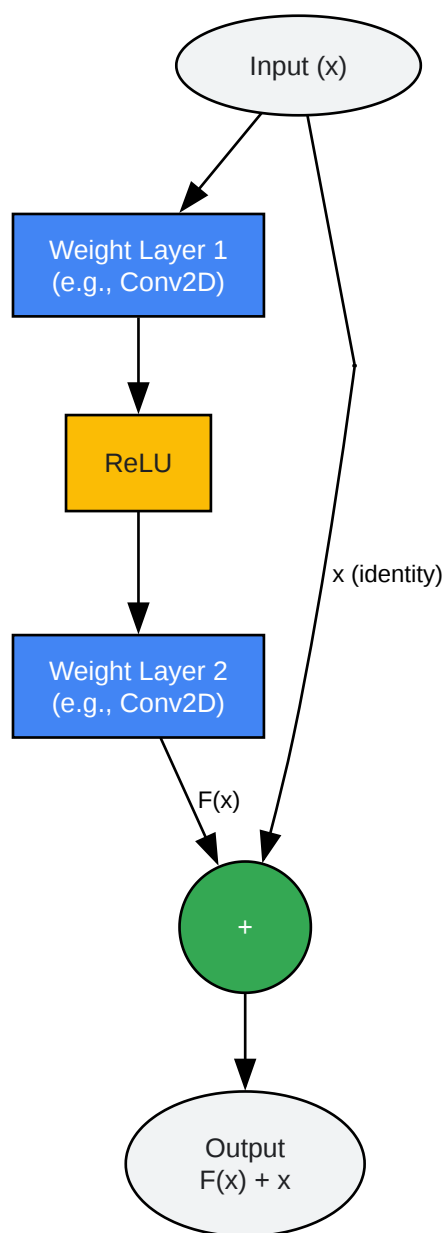


[Click to download full resolution via product page](#)

Comparison of gradient propagation in Sigmoid vs. ReLU networks.

ResNet Skip Connection

A residual block allows the gradient to take a "shortcut," preserving its magnitude across layers.



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

A diagram of a standard ResNet "skip connection" block.

Experimental Protocol: Demonstrating the Vanishing Gradient Problem

This protocol outlines an experiment to visualize the effect of vanishing gradients in a sigmoid-based network and confirm the improvement when using ReLU.

Objective: To compare the magnitude of back-propagated gradients in the hidden layers of a deep neural network using Sigmoid vs. ReLU activation functions.

Methodology:

- Dataset: Use a standard, simple dataset suitable for a multi-layer perceptron (MLP), such as the "make_circles" or "make_moons" dataset from scikit-learn, which requires a non-linear classifier.
- Model Architecture:
 - Construct a deep feed-forward neural network with at least 6-8 hidden layers. Each hidden layer should have a small number of neurons (e.g., 10-20).
 - The output layer should have a single neuron with a sigmoid activation for binary classification.
- Experiment 1: Sigmoid Network
 - Set the activation function for all hidden layers to sigmoid.
 - Use a standard weight initialization, such as the default in Keras/TensorFlow or PyTorch.
 - Train the model for a set number of epochs (e.g., 100).
 - Data Collection: During training, after each weight update (or every N updates), compute and record the L2 norm of the gradients for the weights of each hidden layer.
- Experiment 2: ReLU Network
 - Modify the network from Experiment 1. Change the activation function for all hidden layers to ReLU.[\[1\]](#)
 - Change the weight initialization to He initialization, which is optimized for ReLU.[\[1\]](#)
 - Keep all other parameters (optimizer, learning rate, number of epochs, dataset) identical to Experiment 1.

- Data Collection: Repeat the gradient norm collection process as described in Experiment 1.
- Analysis and Visualization:
 - For each experiment, create a plot with training epochs on the x-axis and the average gradient norm on the y-axis.
 - Plot a separate line for each hidden layer (e.g., Layer 1, Layer 2, ..., Layer 8).
 - Expected Outcome (Sigmoid Plot): The plot for the sigmoid network will show that the gradient norms for the later layers (e.g., 7, 8) are significantly higher than for the early layers (e.g., 1, 2). The lines for the early layers will be flat and close to zero.[6]
 - Expected Outcome (ReLU Plot): The plot for the ReLU network will show that the gradient norms are more evenly distributed across all layers. The early layers will have non-zero gradients, indicating that they are actively learning throughout the training process.

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. machinelearningmastery.com [machinelearningmastery.com]
- 2. [What is Vanishing Gradient Problem?](https://mygreatlearning.com) [mygreatlearning.com]
- 3. [vanishing_grad_example](https://cs224d.stanford.edu) [cs224d.stanford.edu]
- 4. youtube.com [youtube.com]
- 5. kdnuggets.com [kdnuggets.com]
- 6. [Vanishing and Exploding Gradients Problems in Deep Learning - GeeksforGeeks](https://www.geeksforgeeks.org) [geeksforgeeks.org]
- 7. ai.stackexchange.com [ai.stackexchange.com]
- 8. medium.com [medium.com]

- [9. medium.com \[medium.com\]](#)
- [10. researchgate.net \[researchgate.net\]](#)
- [11. Rectified linear unit - Wikipedia \[en.wikipedia.org\]](#)
- [12. Vanishing Gradient Problem in Deep Learning: Explained | DigitalOcean \[digitalocean.com\]](#)
- [13. stats.stackexchange.com \[stats.stackexchange.com\]](#)
- [14. machinelearningmastery.com \[machinelearningmastery.com\]](#)
- [15. Weight Initialization Techniques for Deep Neural Networks - GeeksforGeeks \[geeksforgeeks.org\]](#)
- [16. towardsai.net \[towardsai.net\]](#)
- [17. medium.com \[medium.com\]](#)
- [18. How do batch normalization and gradient clipping help mitigate vanishing gradients? - Massed Compute \[massedcompute.com\]](#)
- [19. How does batch normalization impact the vanishing gradient problem in deep neural networks? - Infermatic \[infermatic.ai\]](#)
- [20. Mastering the Vanishing Gradient Problem: Advanced Strategies for Optimizing Neural Networks - LUNARTECH \[lunartech.ai\]](#)
- [21. datascience.stackexchange.com \[datascience.stackexchange.com\]](#)
- [22. machine learning - How to prevent vanishing gradient or exploding gradient? - Data Science Stack Exchange \[datascience.stackexchange.com\]](#)
- [23. medium.com \[medium.com\]](#)
- [24. medium.com \[medium.com\]](#)
- [25. youtube.com \[youtube.com\]](#)
- [26. aiknow.io \[aiknow.io\]](#)
- [27. spotintelligence.com \[spotintelligence.com\]](#)
- [28. naukri.com \[naukri.com\]](#)
- [29. baeldung.com \[baeldung.com\]](#)
- [30. What is LSTM - Long Short Term Memory? - GeeksforGeeks \[geeksforgeeks.org\]](#)
- [31. ravjot03.medium.com \[ravjot03.medium.com\]](#)
- [32. medium.com \[medium.com\]](#)

- [33. Understanding Gradient Clipping - GeeksforGeeks \[geeksforgeeks.org\]](#)
- [34. Can you explain how gradient clipping can mitigate vanishing gradients? - Infermatic \[infermatic.ai\]](#)
- [35. m.youtube.com \[m.youtube.com\]](#)
- To cite this document: BenchChem. [Technical Support Center: Troubleshooting Vanishing Gradients with Sigmoid Functions]. BenchChem, [2026]. [Online PDF]. Available at: [\[https://www.benchchem.com/product/b074413/docs#technical-support-center-troubleshooting-vanishing-gradients-with-sigmoid-functions\]](https://www.benchchem.com/product/b074413/docs#technical-support-center-troubleshooting-vanishing-gradients-with-sigmoid-functions)

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](#)

