

limitations of the sigmoid function in deep learning and solutions

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

Compound of Interest

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

[Get Quote](#)

Technical Support Center: Deep Learning Activation Functions

This guide provides troubleshooting for common issues encountered during deep learning experiments related to the sigmoid activation function. It is intended for researchers, scientists, and drug development professionals.

Frequently Asked Questions (FAQs) & Troubleshooting

Q1: My deep neural network is training very slowly, and the performance of earlier layers isn't improving. What could be the cause?

A: This is a classic symptom of the vanishing gradient problem.^{[1][2][3][4]} When using sigmoid activation functions in deep networks, the gradients of the loss function can become extremely small as they are propagated backward from the output layer to the earlier layers.^{[1][2][5]}

The derivative of the sigmoid function has a maximum value of 0.25, and for inputs far from zero, the derivative approaches zero.^{[1][6]} During backpropagation, these small derivatives are multiplied together across many layers, causing the gradient to "vanish" exponentially.^{[5][7][8]} As a result, the weights of the initial layers do not get updated effectively, and the network fails to learn.^{[5][7][9]}

Q2: My model's convergence is slow and exhibits a zig-zagging behavior during training. What's happening?

A: This issue is likely due to the non-zero-centered output of the sigmoid function.^{[3][6][10][11]} The sigmoid function outputs values in the range (0, 1), which are always positive.^{[12][13]}

When the inputs to a neuron are always positive, the gradients for the weights during backpropagation will all have the same sign (either all positive or all negative).^{[10][11][14]} This restricts the weight updates to a specific direction, leading to inefficient, zig-zagging dynamics in the gradient descent process and slower convergence.^{[6][10][14]}

Q3: Are there more computationally efficient alternatives to the sigmoid function for hidden layers?

A: Yes. The sigmoid function involves calculating an exponential, which can be computationally expensive compared to simpler functions.^{[1][3][15][16]} For deep networks that require numerous calculations of activation functions, this can significantly impact training time.^[15] A highly efficient and popular alternative is the Rectified Linear Unit (ReLU) and its variants.^{[15][17][18]}

Solutions & Experimental Protocols

Solution 1: Replace Sigmoid with ReLU Family Activation Functions

The most common solution to the vanishing gradient and computational cost issues is to replace the sigmoid function in the hidden layers with a member of the Rectified Linear Unit (ReLU) family.^{[2][17][19]}

Below is a sample protocol for changing the activation function in a neural network layer using popular deep learning frameworks.

Objective: Replace sigmoid with ReLU in the hidden layers of a sequential model.

Framework: TensorFlow (with Keras API)

Framework: PyTorch

Troubleshooting Guide for ReLU-based Activations

While ReLU is a powerful default, you may encounter the "Dying ReLU" problem, where neurons become inactive and always output zero.[\[6\]](#)[\[16\]](#)[\[20\]](#) This happens if a large gradient update causes the weights to be adjusted such that the neuron's input is always negative.[\[6\]](#)

Q: My network's performance has stalled after switching to ReLU. Some neurons seem to be inactive. What should I do?

A: This is likely the "Dying ReLU" problem.[\[20\]](#)[\[21\]](#) Consider using variants of ReLU designed to mitigate this issue.

- Leaky ReLU: Introduces a small, non-zero slope for negative inputs, ensuring that neurons always have a small gradient and can continue to learn.[\[17\]](#)[\[22\]](#)[\[23\]](#)[\[24\]](#)
- Exponential Linear Unit (ELU): Uses an exponential curve for negative inputs, which can lead to faster convergence and better generalization.[\[12\]](#)[\[20\]](#)[\[25\]](#)[\[26\]](#) It allows for negative outputs, which can help push the mean activation closer to zero.[\[27\]](#)
- Gaussian Error Linear Unit (GELU): A smooth activation function that weights inputs by their magnitude rather than gating them by their sign.[\[28\]](#) It is commonly used in transformer models like BERT and GPT.[\[21\]](#)[\[28\]](#)

Data Presentation: Comparison of Activation Functions

Activation Function	Output Range	Zero-Centered	Vanishing Gradient Problem	Dying Neuron Problem	Computational Cost	Primary Use Case (Hidden Layers)
Sigmoid	(0, 1)	No	Yes[2][3][16]	No	High (Exponential)[3][16]	Not recommended for hidden layers in deep networks. [29]
Tanh	(-1, 1)	Yes[30]	Yes (less severe than Sigmoid)[7]	No	High (Exponential)	Recurrent Neural Networks (RNNs). [30][31]
ReLU	[0, ∞)	No	No (for positive inputs)[2][16]	Yes[6][16]	Low[15][18]	Default choice for CNNs and MLPs.[31][32]
Leaky ReLU	(-∞, ∞)	Almost	No	No[22][23][33]	Low	When "Dying ReLU" is an issue. [23][24]
ELU	(-α, ∞)	Closer to zero	No	No[26][27]	Medium (Exponential for x<0)	A strong alternative to ReLU. [12][27]
GELU	(-0.17, ∞)	Yes	No	No	High	Transformer-based

models.

[\[21\]](#)[\[28\]](#)[\[34\]](#)

Solution 2: Implement Batch Normalization

Batch Normalization (BN) is a technique that can mitigate the problems of sigmoid and other activation functions.[\[35\]](#)[\[36\]](#) It normalizes the inputs to each layer to have a mean of zero and a variance of one.[\[37\]](#)[\[38\]](#)

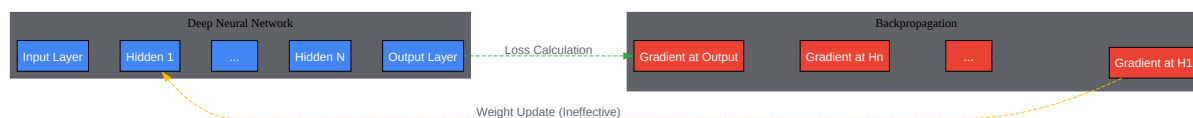
How it helps:

- Reduces Internal Covariate Shift: BN stabilizes the distribution of inputs to each layer, which accelerates training.[\[35\]](#)
- Mitigates Vanishing Gradients: By keeping the inputs to the sigmoid function centered around zero, where the gradient is steepest, BN ensures a healthier gradient flow.[\[37\]](#)[\[38\]](#)
- Allows Higher Learning Rates: BN makes the model more robust to the scale of parameters, enabling faster training with higher learning rates.[\[36\]](#)

Objective: Add a Batch Normalization layer before the activation function.

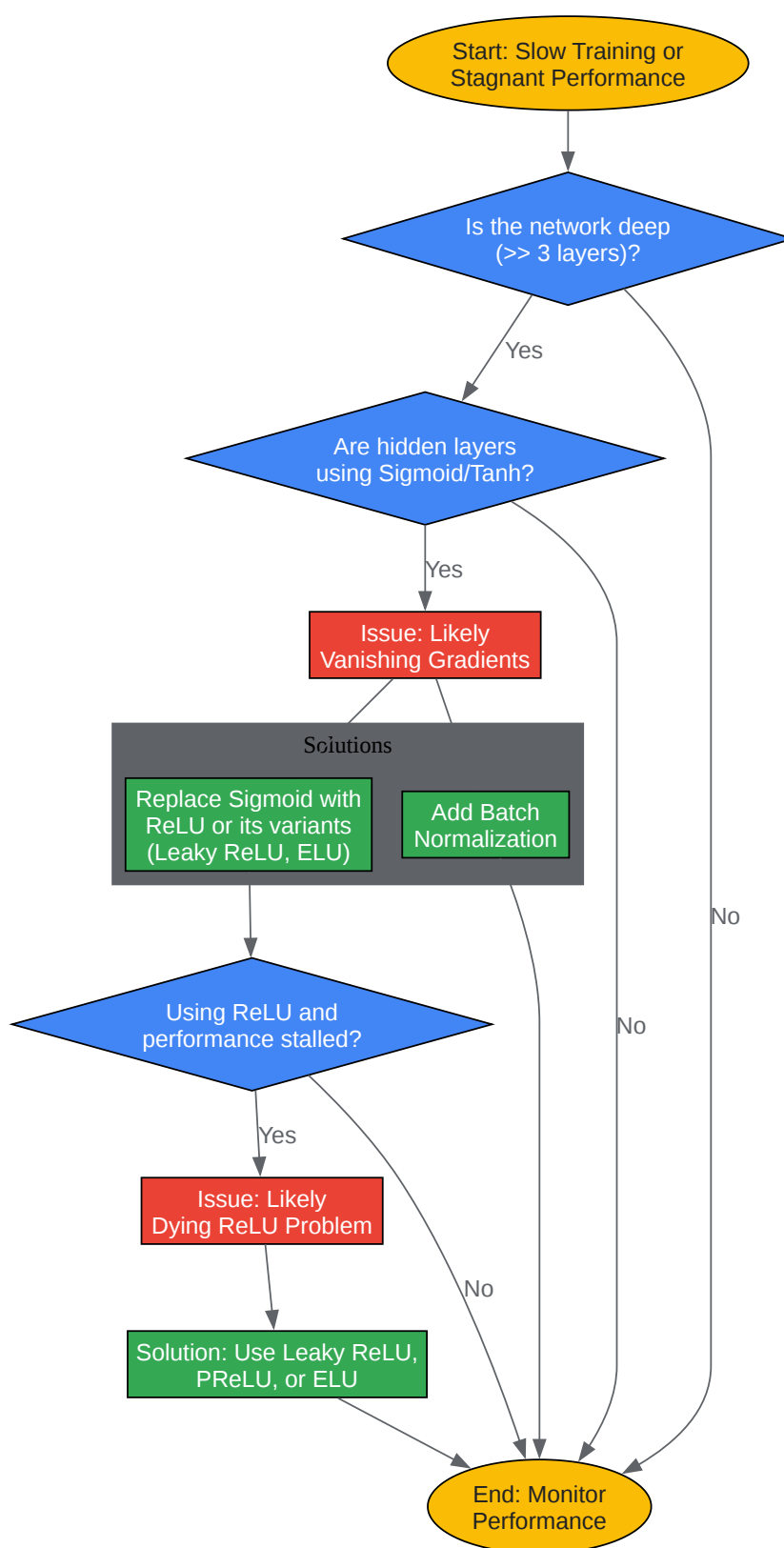
Framework: TensorFlow (with Keras API)

Visualizations (Graphviz DOT Language)



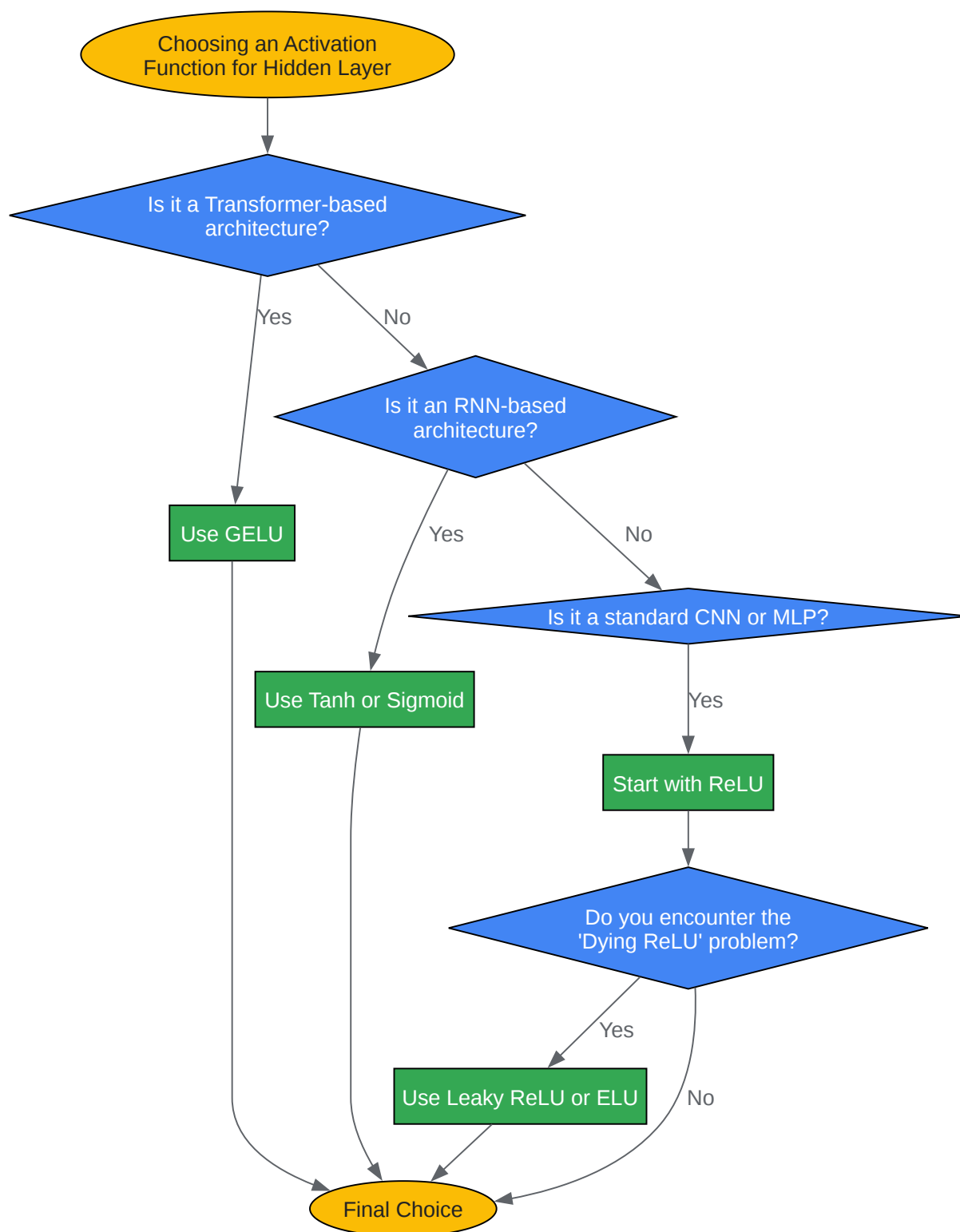
[Click to download full resolution via product page](#)

Caption: The vanishing gradient problem in a deep network using sigmoid (σ).



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

Caption: Troubleshooting workflow for slow training in deep neural networks.



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

Caption: Decision tree for selecting an appropriate activation function. Caption: Decision tree for selecting an appropriate activation function.

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. What are the potential drawbacks of using Sigmoid activation functions in deep neural networks? - Infermatic \[infermatic.ai\]](#)
- [2. medium.com \[medium.com\]](#)
- [3. nachi-keta.medium.com \[nachi-keta.medium.com\]](#)
- [4. quora.com \[quora.com\]](#)
- [5. What is Vanishing Gradient Problem? \[mygreatlearning.com\]](#)
- [6. Don't use sigmoid: Neural Nets \[kharshit.github.io\]](#)
- [7. Vanishing and Exploding Gradients Problems in Deep Learning - GeeksforGeeks \[geeksforgeeks.org\]](#)
- [8. Vanishing gradient problem - Wikipedia \[en.wikipedia.org\]](#)
- [9. medium.com \[medium.com\]](#)
- [10. neural networks - Why are non zero-centered activation functions a problem in backpropagation? - Cross Validated \[stats.stackexchange.com\]](#)
- [11. neural networks - Non zero centered activation functions - Cross Validated \[stats.stackexchange.com\]](#)
- [12. Activation Functions — ML Glossary documentation \[ml-cheatsheet.readthedocs.io\]](#)
- [13. Tanh vs. Sigmoid vs. ReLU - GeeksforGeeks \[geeksforgeeks.org\]](#)
- [14. medium.com \[medium.com\]](#)
- [15. stats.stackexchange.com \[stats.stackexchange.com\]](#)
- [16. medium.com \[medium.com\]](#)
- [17. What are the alternatives to the Sigmoid activation function in machine learning? - Massed Compute \[massedcompute.com\]](#)

- [18. medium.com \[medium.com\]](#)
- [19. kdnuggets.com \[kdnuggets.com\]](#)
- [20. Elu Activation Function in Neural Network - GeeksforGeeks \[geeksforgeeks.org\]](#)
- [21. nachi-keta.medium.com \[nachi-keta.medium.com\]](#)
- [22. Leaky Relu Activation Function in Deep Learning - GeeksforGeeks \[geeksforgeeks.org\]](#)
- [23. Leaky ReLU Activation Function Explained | Ultralytics \[ultralytics.com\]](#)
- [24. quora.com \[quora.com\]](#)
- [25. koshurst.medium.com \[koshurst.medium.com\]](#)
- [26. closeheat.com \[closeheat.com\]](#)
- [27. tungmphung.com \[tungmphung.com\]](#)
- [28. GELU: Gaussian Error Linear Unit Explained | Ultralytics \[ultralytics.com\]](#)
- [29. medium.com \[medium.com\]](#)
- [30. medium.com \[medium.com\]](#)
- [31. machinelearningmastery.com \[machinelearningmastery.com\]](#)
- [32. machinemindscape.com \[machinemindscape.com\]](#)
- [33. medium.com \[medium.com\]](#)
- [34. pub.towardsai.net \[pub.towardsai.net\]](#)
- [35. knowledge.kevinjosephthomas.com \[knowledge.kevinjosephthomas.com\]](#)
- [36. quora.com \[quora.com\]](#)
- [37. What are the differences in batch normalization effects on models with sigmoid and tanh activation functions? - Massed Compute \[massedcompute.com\]](#)
- [38. towardsdatascience.com \[towardsdatascience.com\]](#)
- To cite this document: BenchChem. [limitations of the sigmoid function in deep learning and solutions]. BenchChem, [2026]. [Online PDF]. Available at: [\[https://www.benchchem.com/product/b074413/docs#limitations-of-the-sigmoid-function-in-deep-learning-and-solutions\]](https://www.benchchem.com/product/b074413/docs#limitations-of-the-sigmoid-function-in-deep-learning-and-solutions)

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](#)