

Speeding Up Learning and Inference

Danna Gurari

University of Colorado Boulder
Spring 2025



<https://dannagurari.colorado.edu/course/neural-networks-and-deep-learning-spring-2025/>

Review

- Last lecture on model compression:
 - Motivation
 - Pruning
 - Knowledge distillation
 - Final project report: Overleaf tutorial
- Assignments (Canvas):
 - Final project outline due in one week
- Expect a Piazza announcement from Supriya Naidu about course feedback
- Questions?

Today's Topics

- Motivation
- Hardware tricks
- Architectural tricks
- Training tricks
- Programming tutorial

Today's Topics

- Motivation
- Hardware tricks
- Architectural tricks
- Training tricks
- Programming tutorial

Trend: More Compute-Intensive Models

Training Time
(more FLOPs)

Inference Time
(more generated tokens)



e.g., A Common Training Situation



Boss: What did you do last month?

You: Trained the model for one epoch.



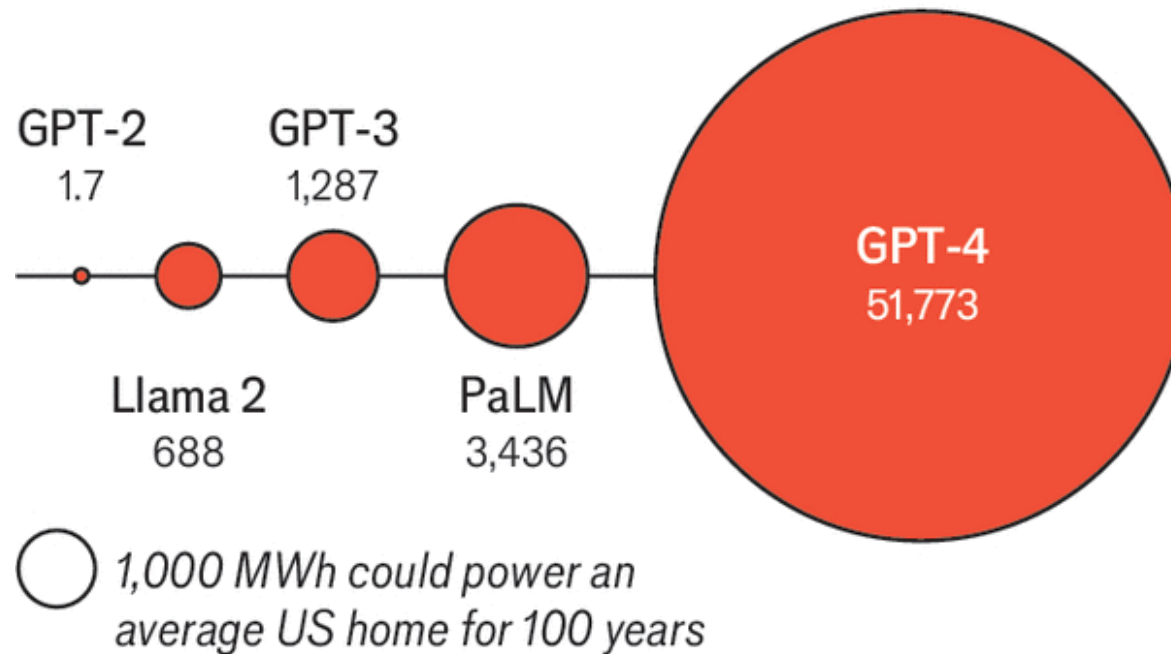
Why Is Extensive Compute Undesirable?

- Time-consuming
- Expensive
- Increased environmental impact from carbon emissions

e.g., A Common Training Situation

Burning up

Energy used to train models, MWh



How to develop and use AI
models with less compute?

Today's Topics

- Motivation
- Hardware tricks
- Architectural tricks
- Training tricks
- Programming tutorial

Ideas

- **Quantization**: reduce precision
- **Operator fusion**: reduce memory read/writes
- **Caching**: reduce computation
- **Distributed optimization**: parallelize computation

