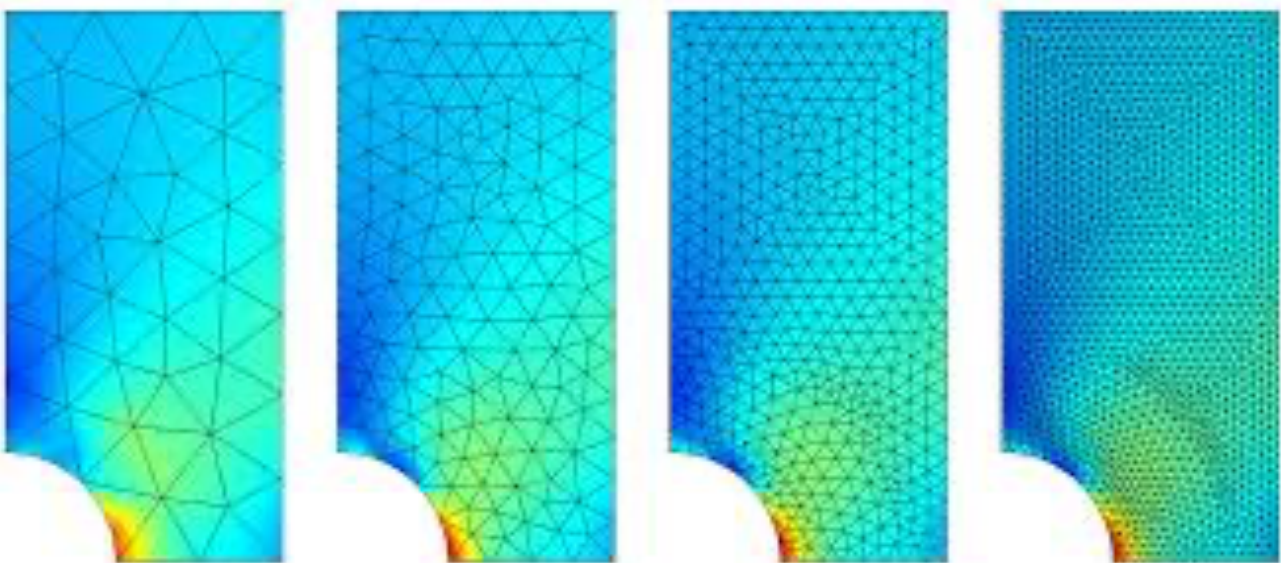


Deep Learning is More Than Just Dense Matrix Multiplication

How to Model Complex (Non-linear) Systems?

Mesh

Piece-wise linear functions



<https://www.comsol.jp/multiphysics/mesh-refinement>

$$\frac{\partial \mathbf{u}}{\partial t} + (\mathbf{u} \cdot \nabla) \mathbf{u} = \nu \nabla^2 \mathbf{u} - \frac{1}{\rho} \nabla p + \mathbf{f}$$

Millennium Prize Problems

Birch and Swinnerton-Dyer conjecture
Hodge conjecture

Navier-Stokes existence and smoothness

P versus NP problem

Poincaré conjecture (solved)

Riemann hypothesis

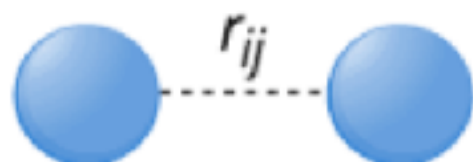
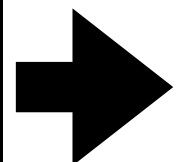
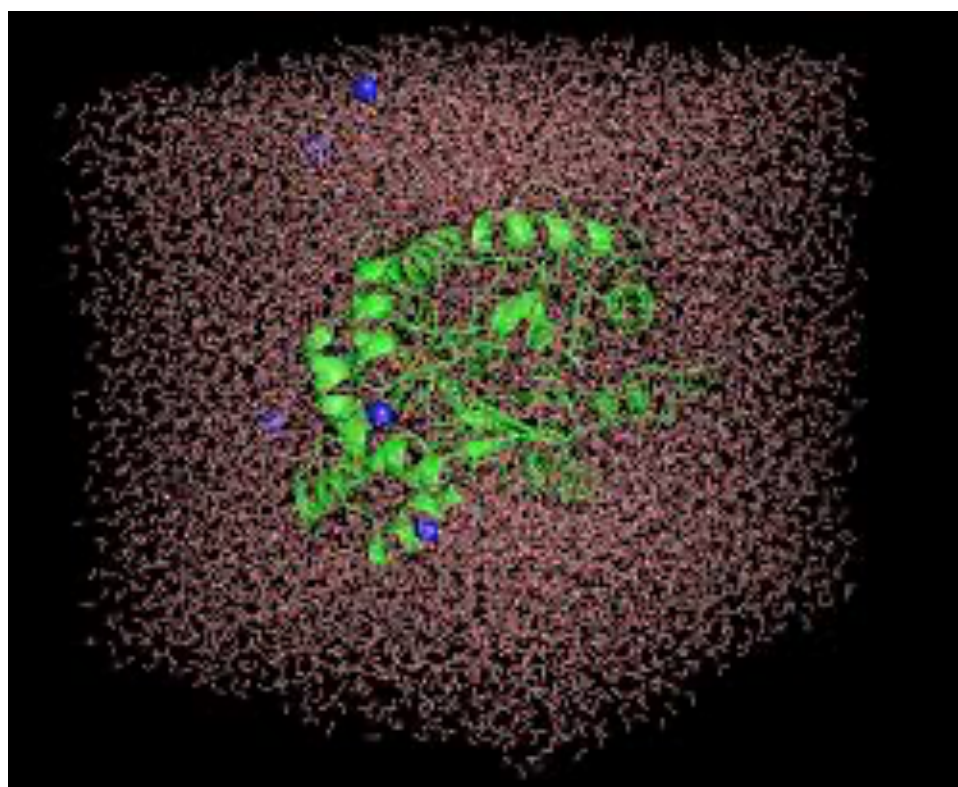
Yang-Mills existence and mass gap

$$Ax = b$$

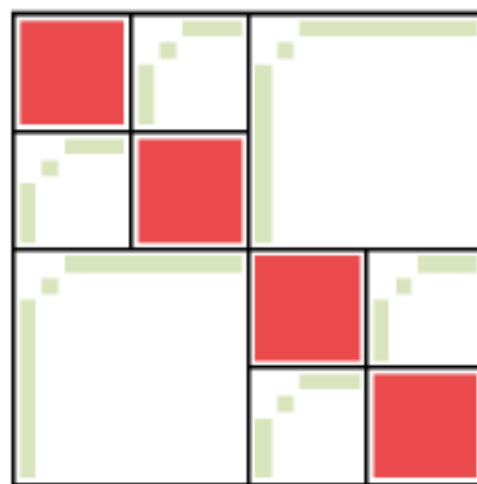
Basic linear algebra
Clear cut API

Particle

Pair-wise linear functions

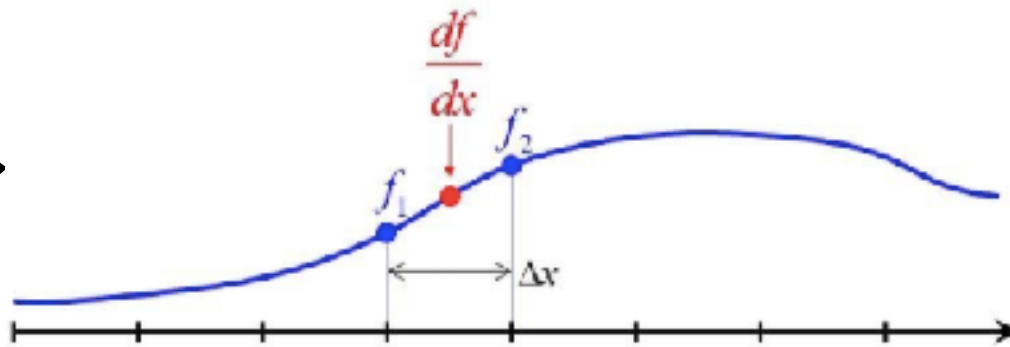


$$F = \frac{1}{4\pi\epsilon_0} \frac{q_i q_j}{r_{ij}^2}$$



<https://qstatix.co.in/molecular-dynamics-simulations/>

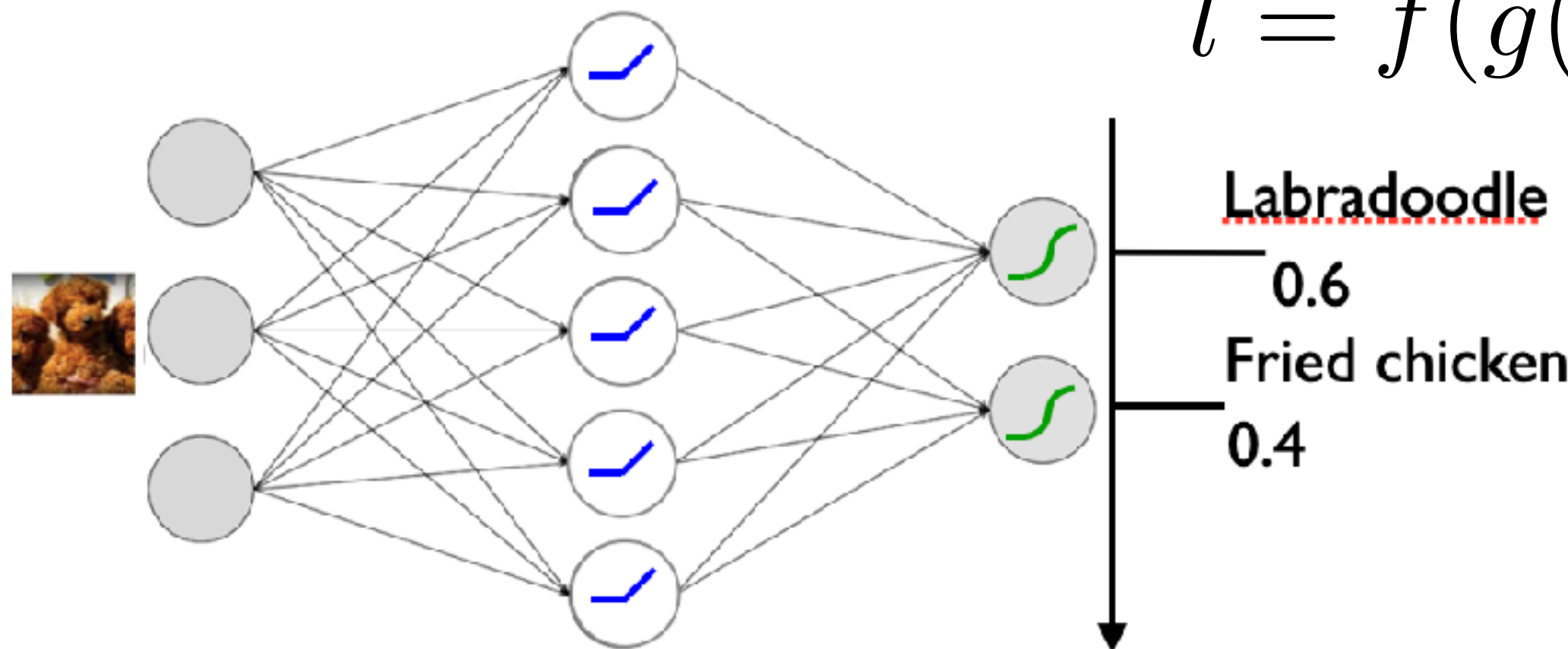
$$\frac{\partial f}{\partial x} = \frac{f_2 - f_1}{\Delta x}$$



Neural Network

Composite (mostly) linear functions

$$l = f(g(h(x)))$$

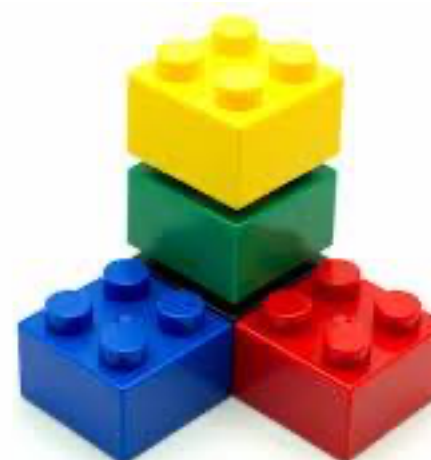


Forward propagation

$$x_0 \rightarrow z_0 = W_0 x_0 \rightarrow x_1 = \text{ReLU}(z_0) \rightarrow z_1 = W_1 x_1 \rightarrow p = \text{softmax}(z_1) \rightarrow l = -\log p$$

Backward propagation

$$\frac{\partial x_1}{\partial z_0} * \frac{\partial z_1}{\partial x_1} * \frac{\partial l}{\partial z_1} = \frac{\partial l}{\partial \theta}$$
$$\otimes \frac{\partial z_0}{\partial W_0} \otimes \frac{\partial z_1}{\partial W_1}$$

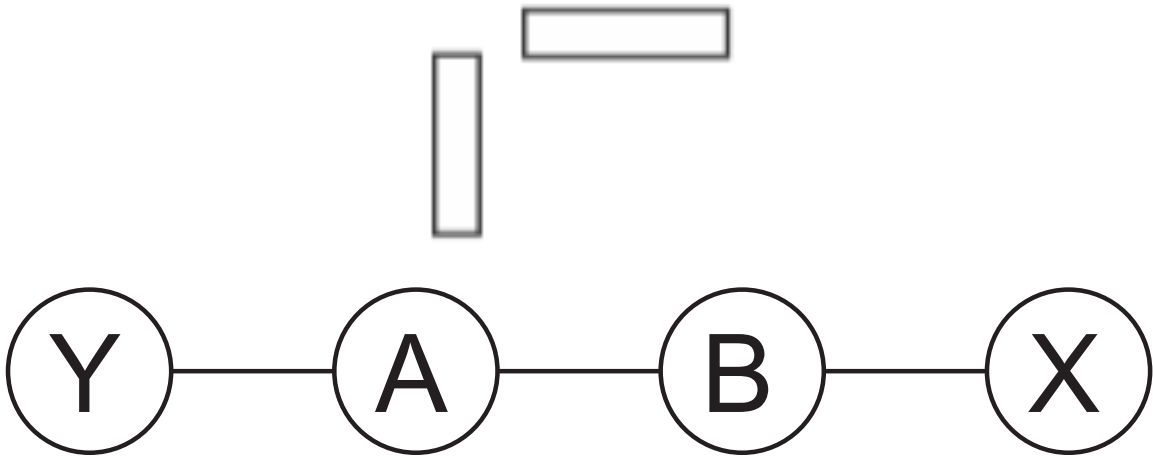


PyTorch

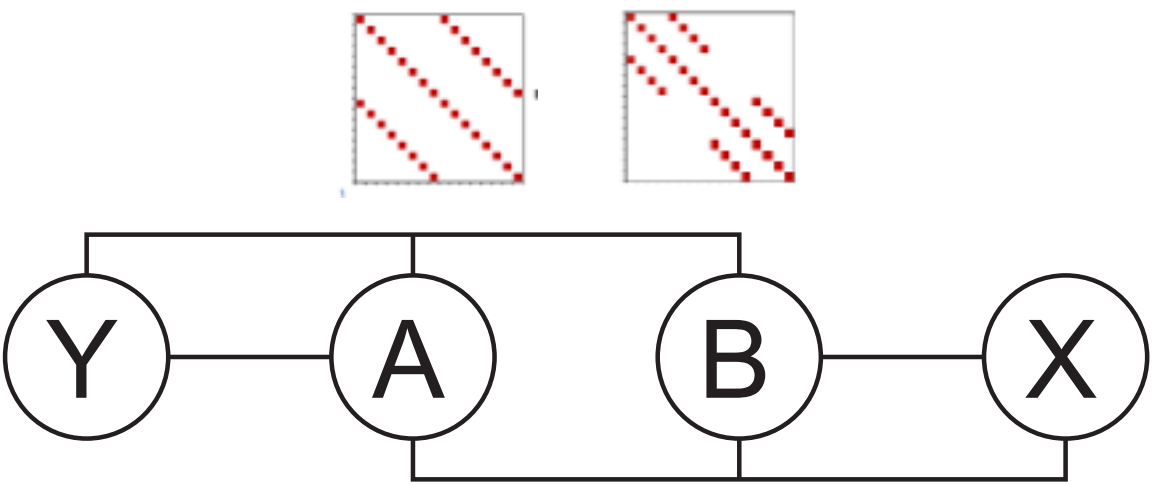
Various Matrix Factorizations

Linear Layer $Y = W X$ $\xrightarrow{\text{Factorize } W}$ $Y = A B X$

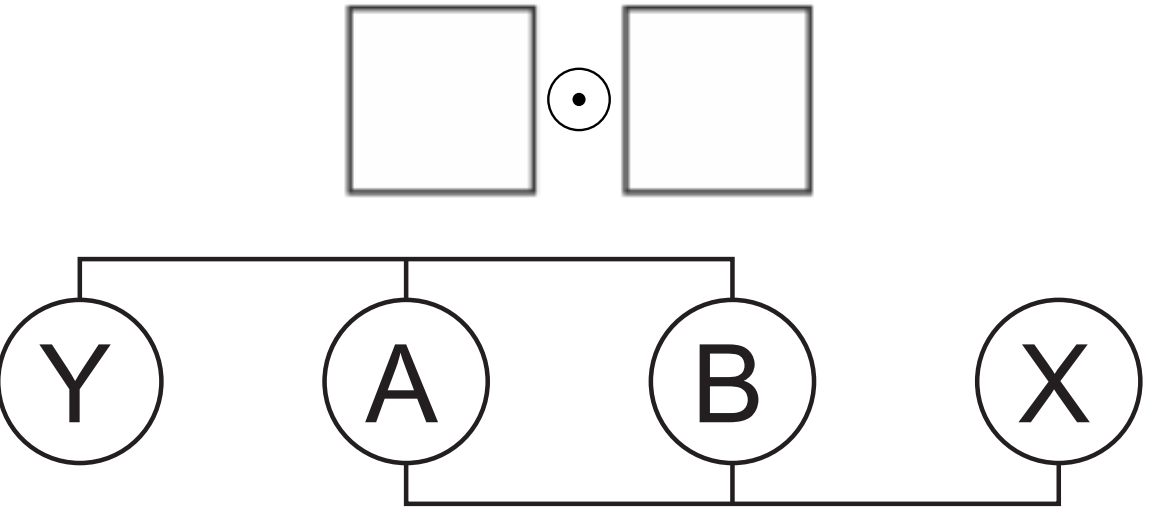
Low-Rank



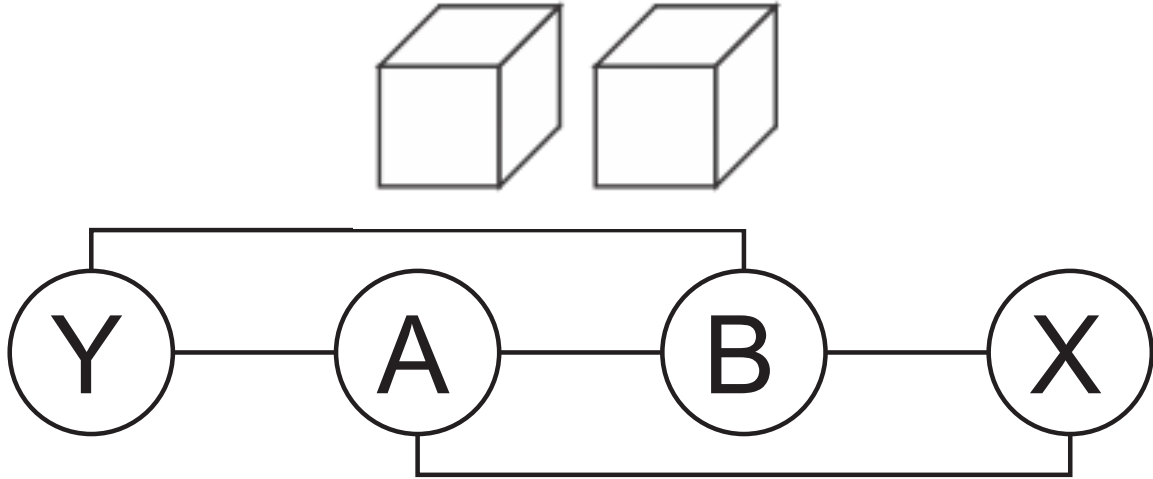
Monarch



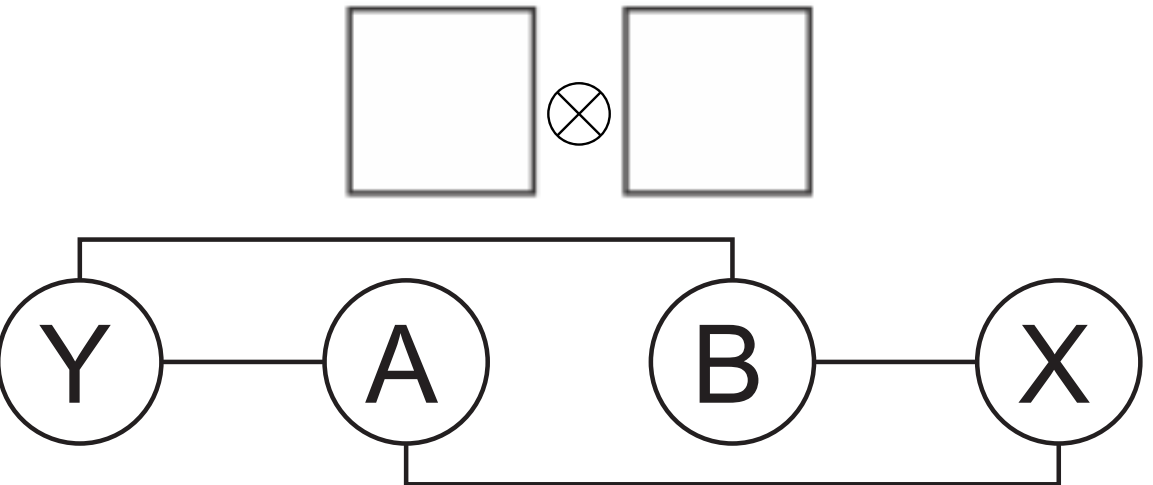
Hadamard



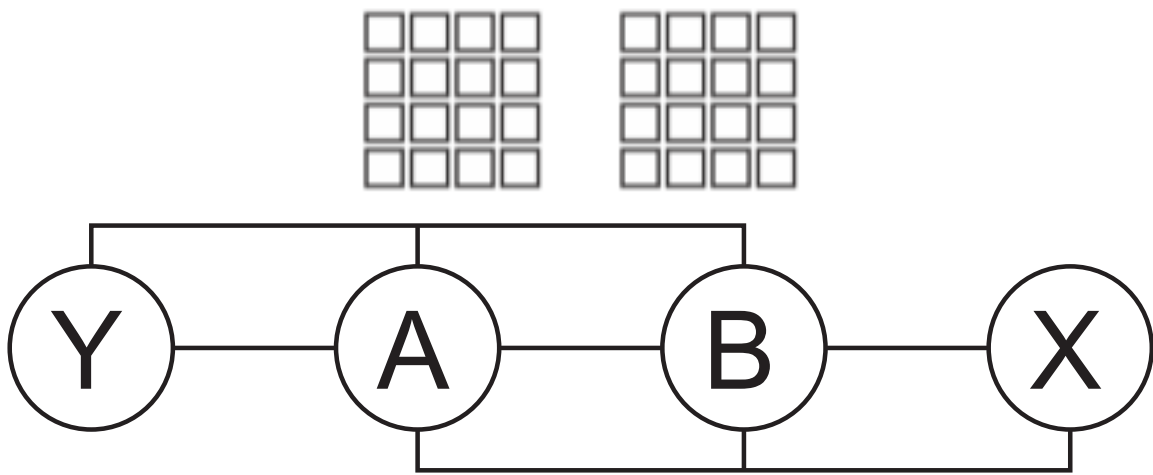
Tensor Train



Kronecker

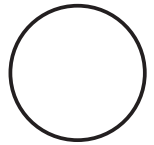


Block Tensor Train



Various Products

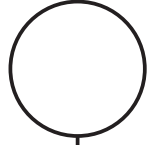
Scalar





a

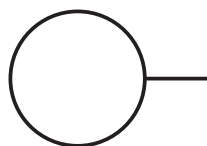
Vector





a_i

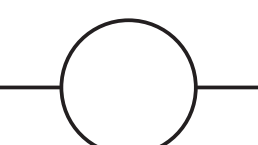
Matrix

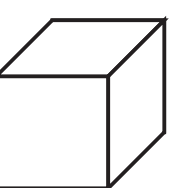




A_{ij}

Tensor

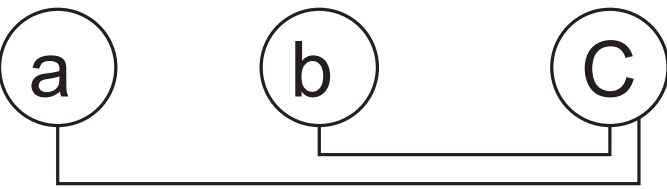


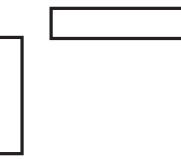


A_{ijk}

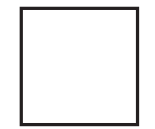
Outer (Kronecker) Product $\otimes$

Vector



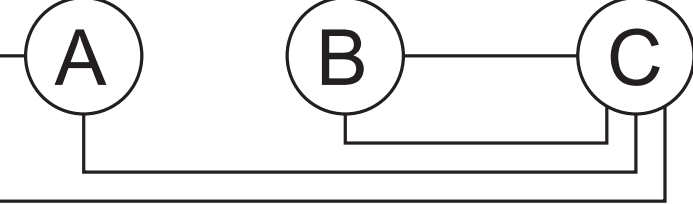


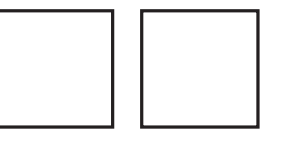
$=$



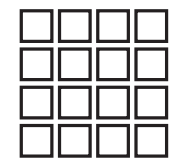
$a_i b_j = C_{ij}$

Matrix



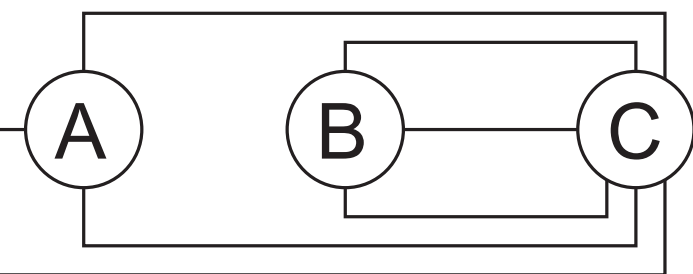


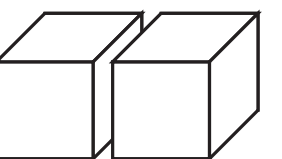
$=$



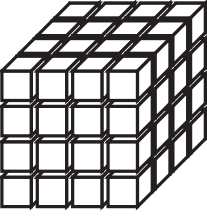
$A_{ij} B_{kl} = C_{ijkl}$

Tensor





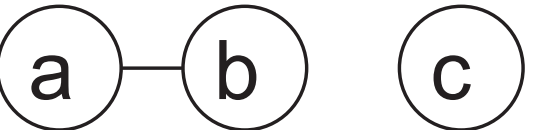
$=$

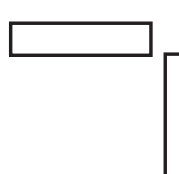


$A_{ijk} B_{lmn} = C_{ijklmn}$

Inner (Dot) Product

Vector



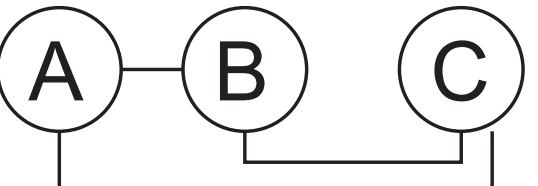


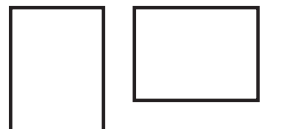
$=$



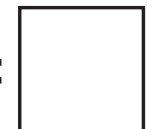
$a_i b_i = c$

Matrix



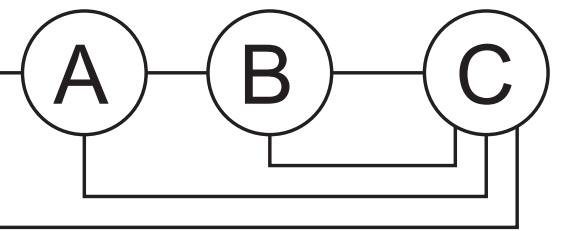


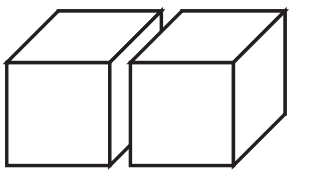
$=$



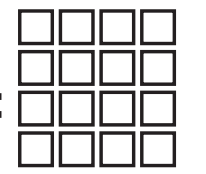
$A_{ij} B_{jk} = C_{ik}$

Tensor





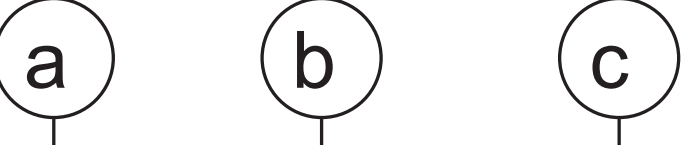
$=$



$A_{ijk} B_{klm} = C_{ijlm}$

Elementwise (Hadamard) Product $\odot$

Vector



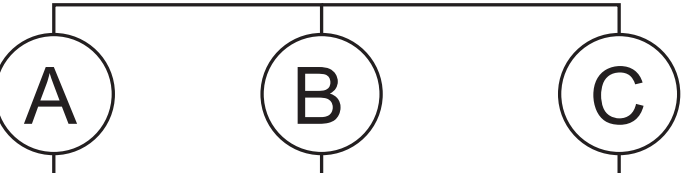


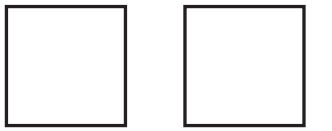
$=$



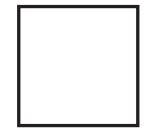
$a_i b_i = c_i$

Matrix



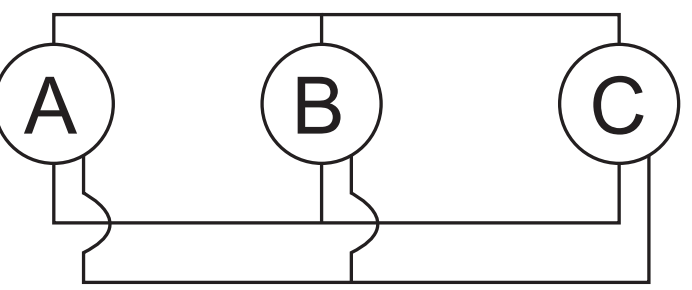


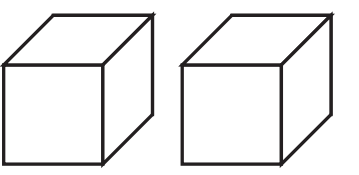
$=$



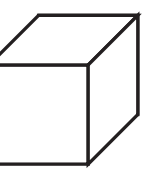
$A_{ij} B_{ij} = C_{ij}$

Tensor





$=$



$A_{ijk} B_{ijk} = C_{ijk}$

Various Matrix Factorizations

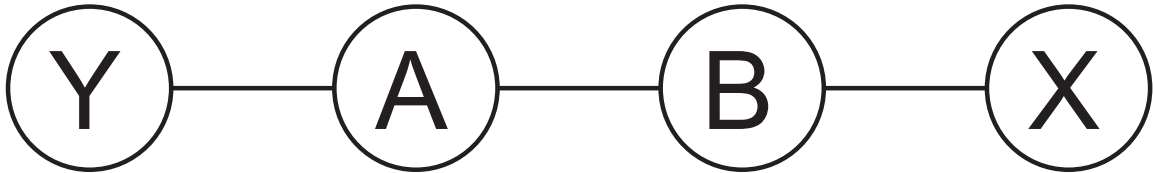
Linear Layer

$Y = WX$

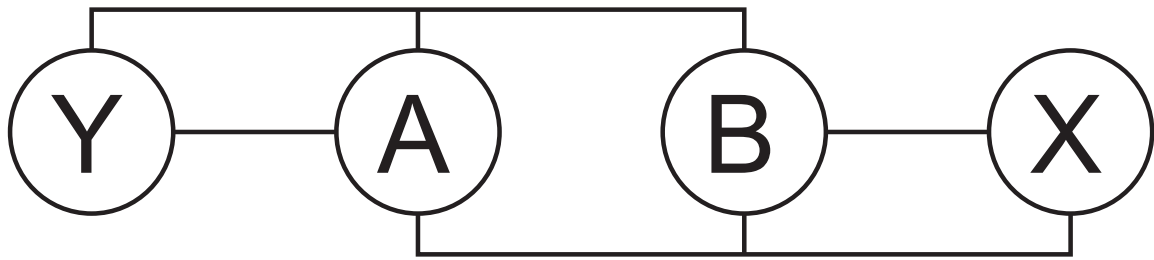
$\xrightarrow{\text{Factorize}}$

$Y = ABX$

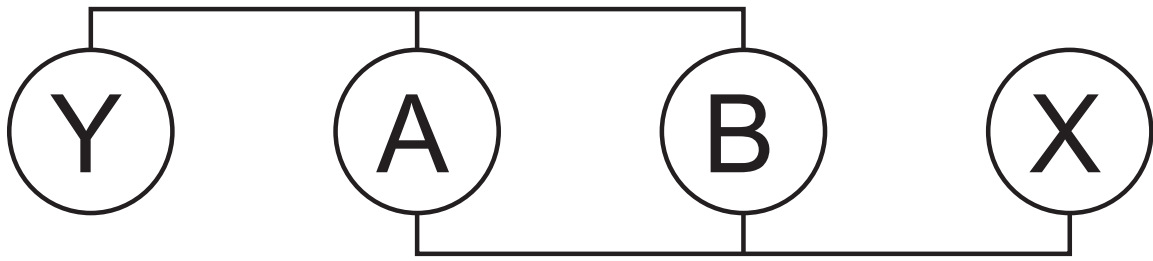
Low-Rank



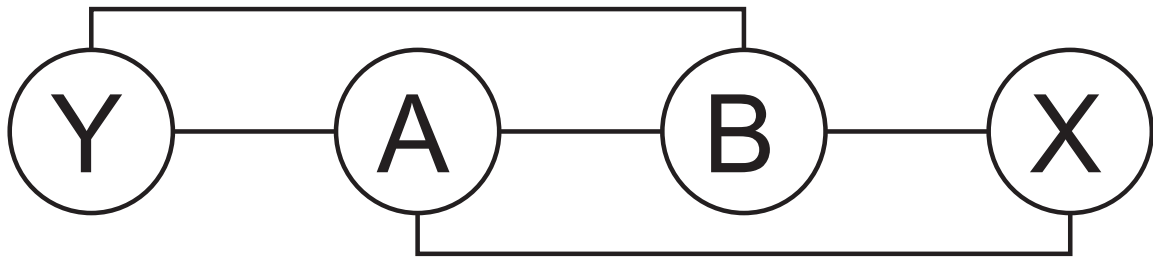
Monarch



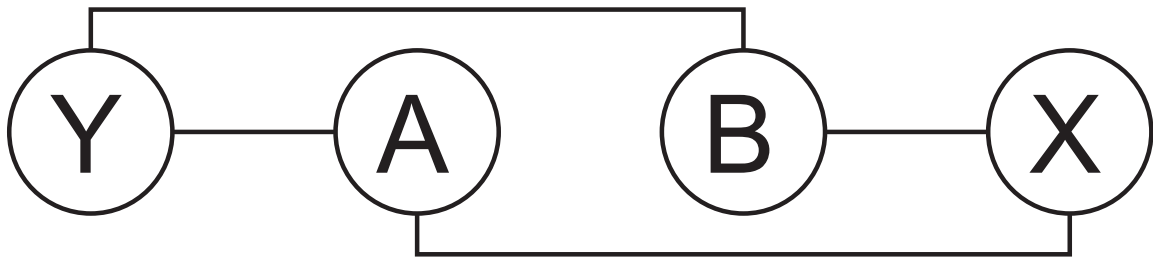
Hadamard



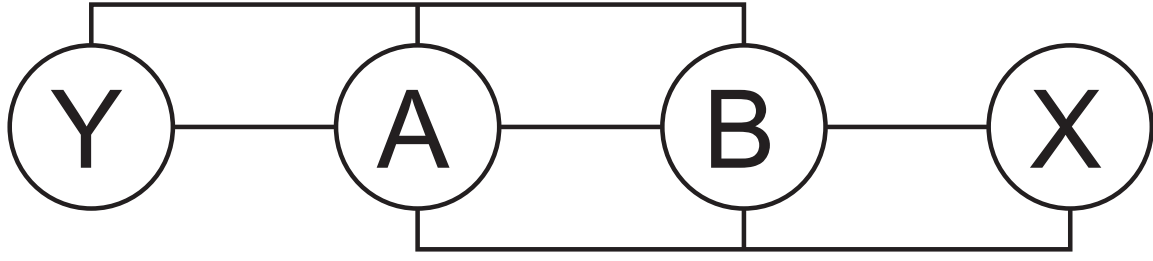
Tensor Train



Kronecker



Block Tensor Train



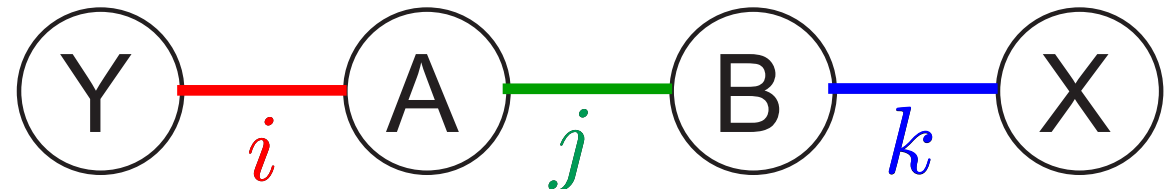
Various Matrix Factorizations

Linear Layer

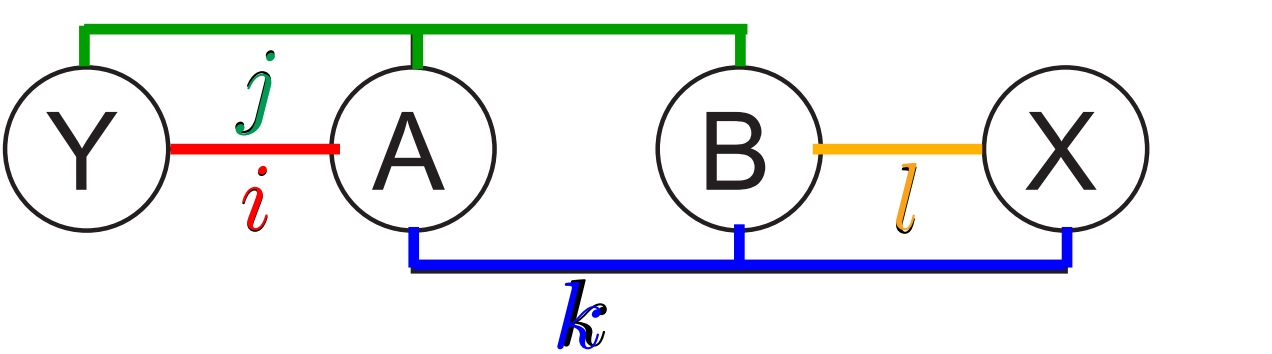
$Y = WX \longrightarrow Y = ABX$

Factorize

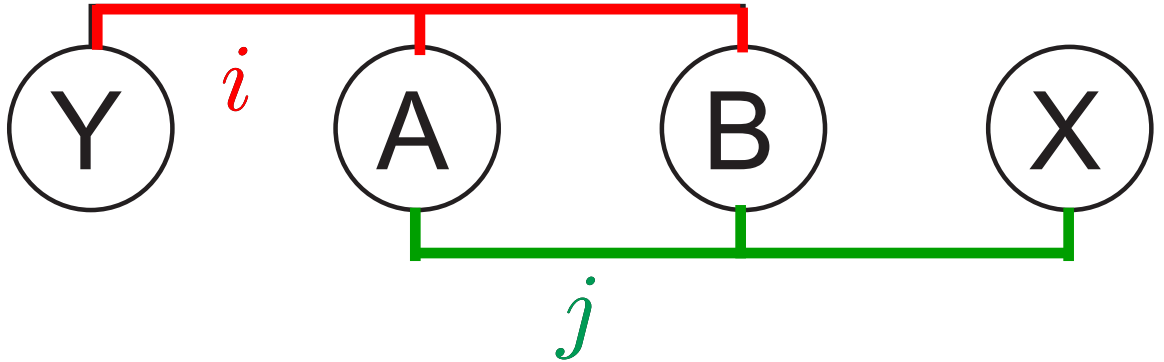
Low-Rank

$$Y_i = A_{ij} B_{jk} X_k$$


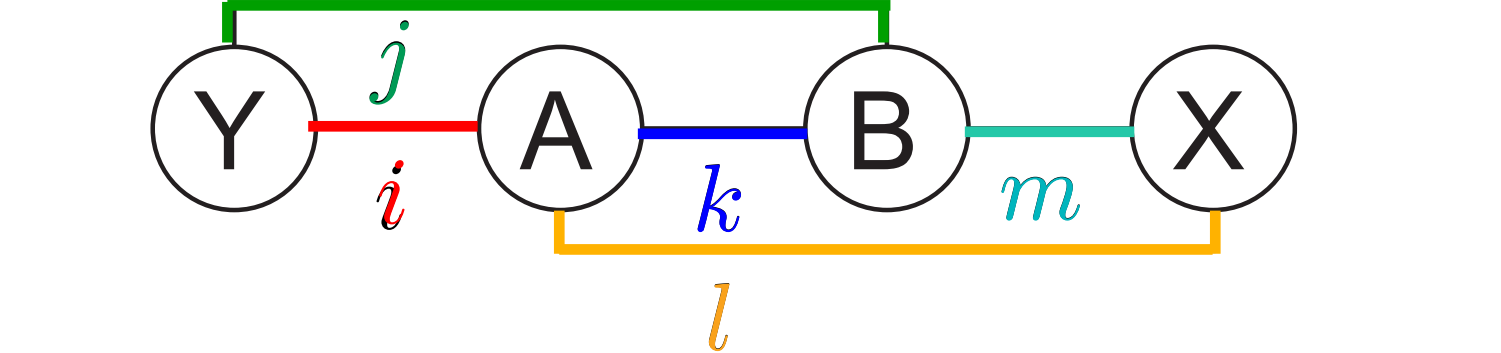
Monarch

$$Y_{ij} = A_{ijk} B_{jkl} X_{kl}$$


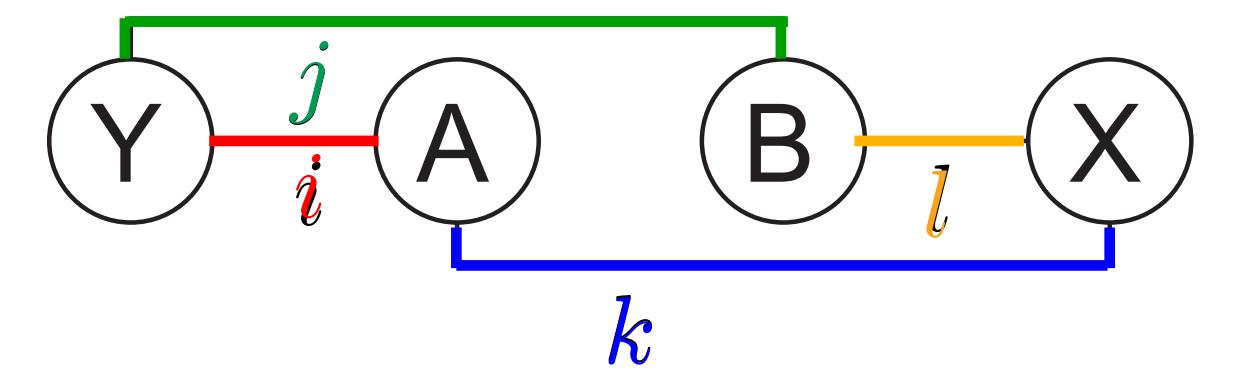
Hadamard

$$Y_i = A_{ij} B_{ij} X_j$$


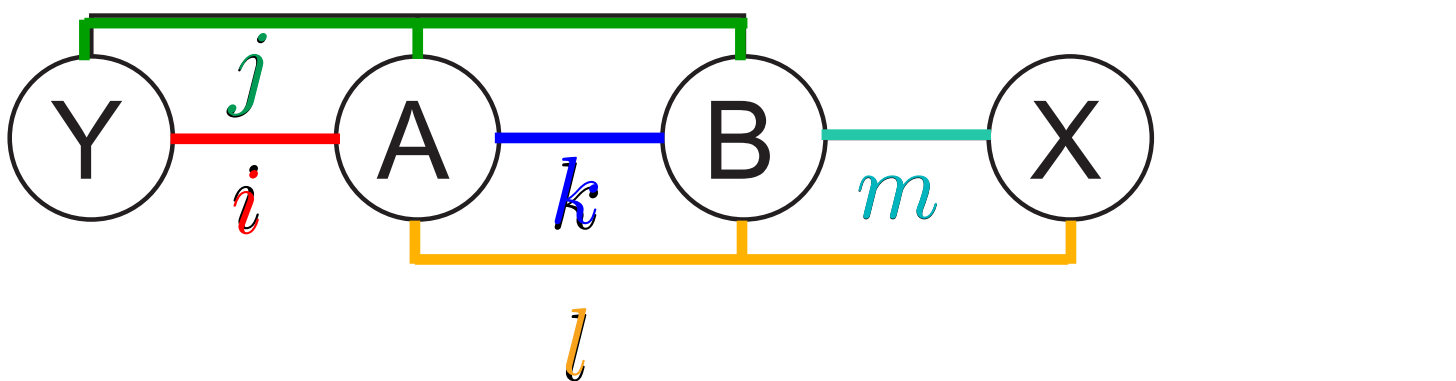
Tensor Train

$$Y_{ij} = A_{ikl} B_{jkm} X_{lm}$$


Kronecker

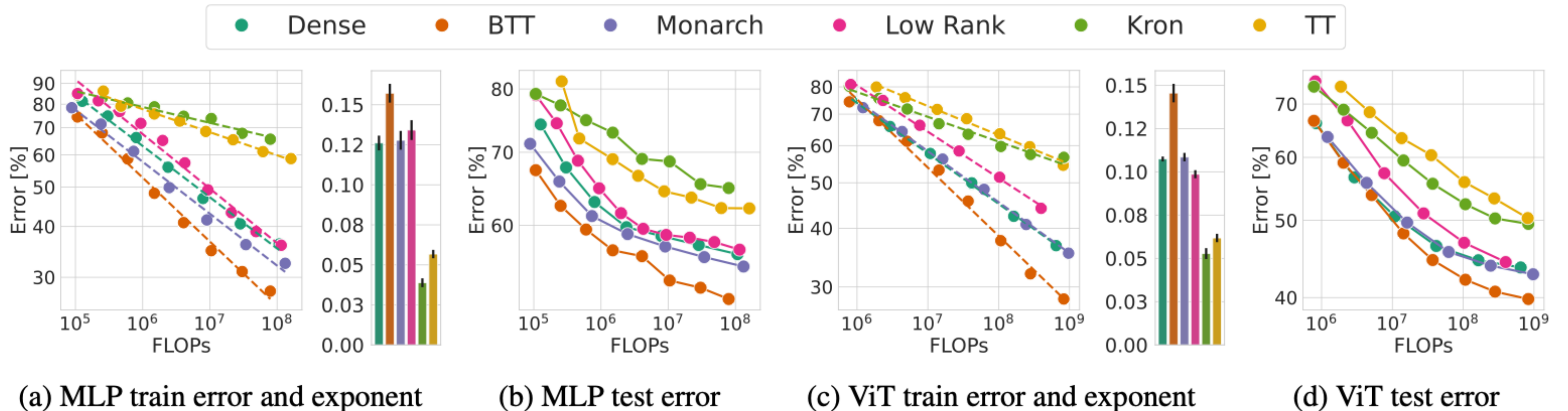
$$Y_{ij} = A_{ik} B_{jl} X_{kl}$$


Block Tensor Train

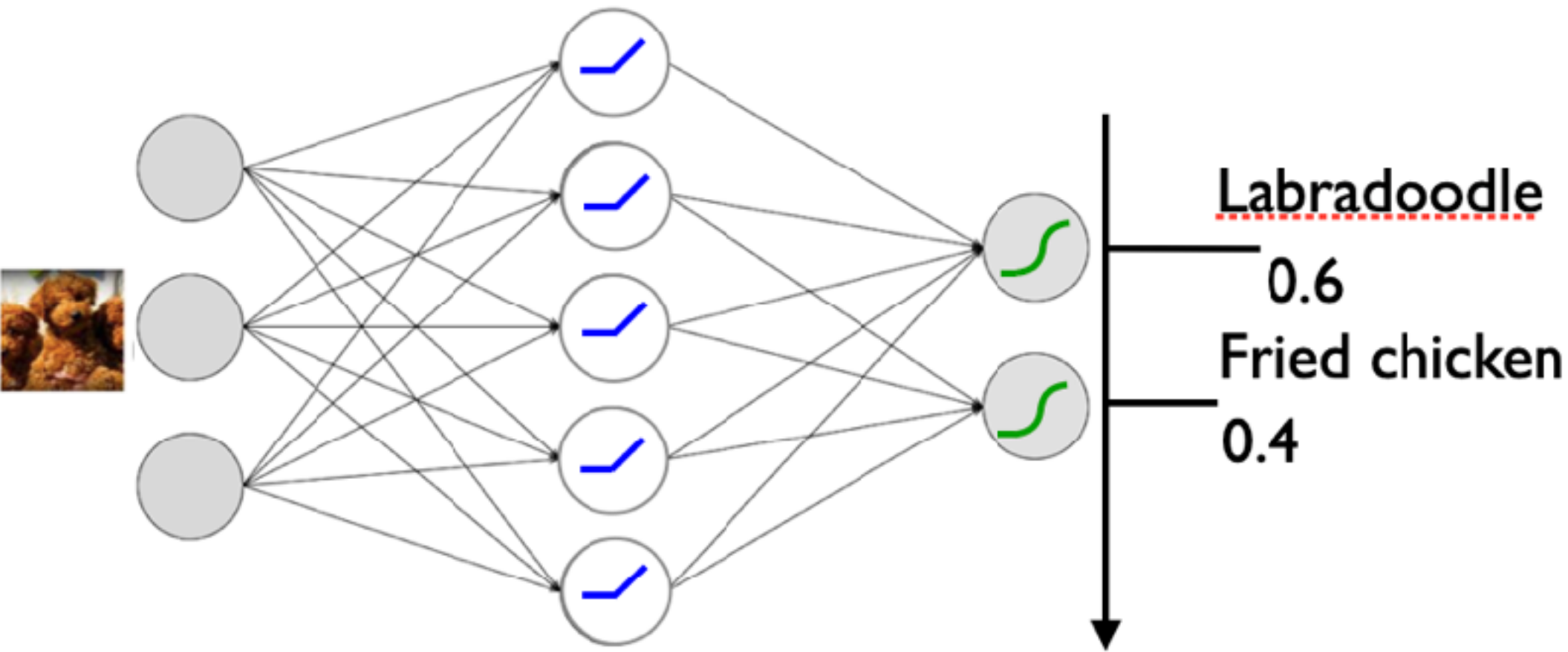
$$Y_{ij} = A_{ijkl} B_{jklm} X_{lm}$$


Scaling Laws on CIFAR-100

- Block Tensor Train has the best scaling
 - $\text{BTT} > \text{Monarch} > \text{Dense} > \text{Low-Rank} > \text{TT} = \text{Kron}$
- Works not only for MLPs, but also ViTs
 - Ranking of methods is also similar
- Test / train gap
 - Similar to existing studies



Backprop and Kronecker products



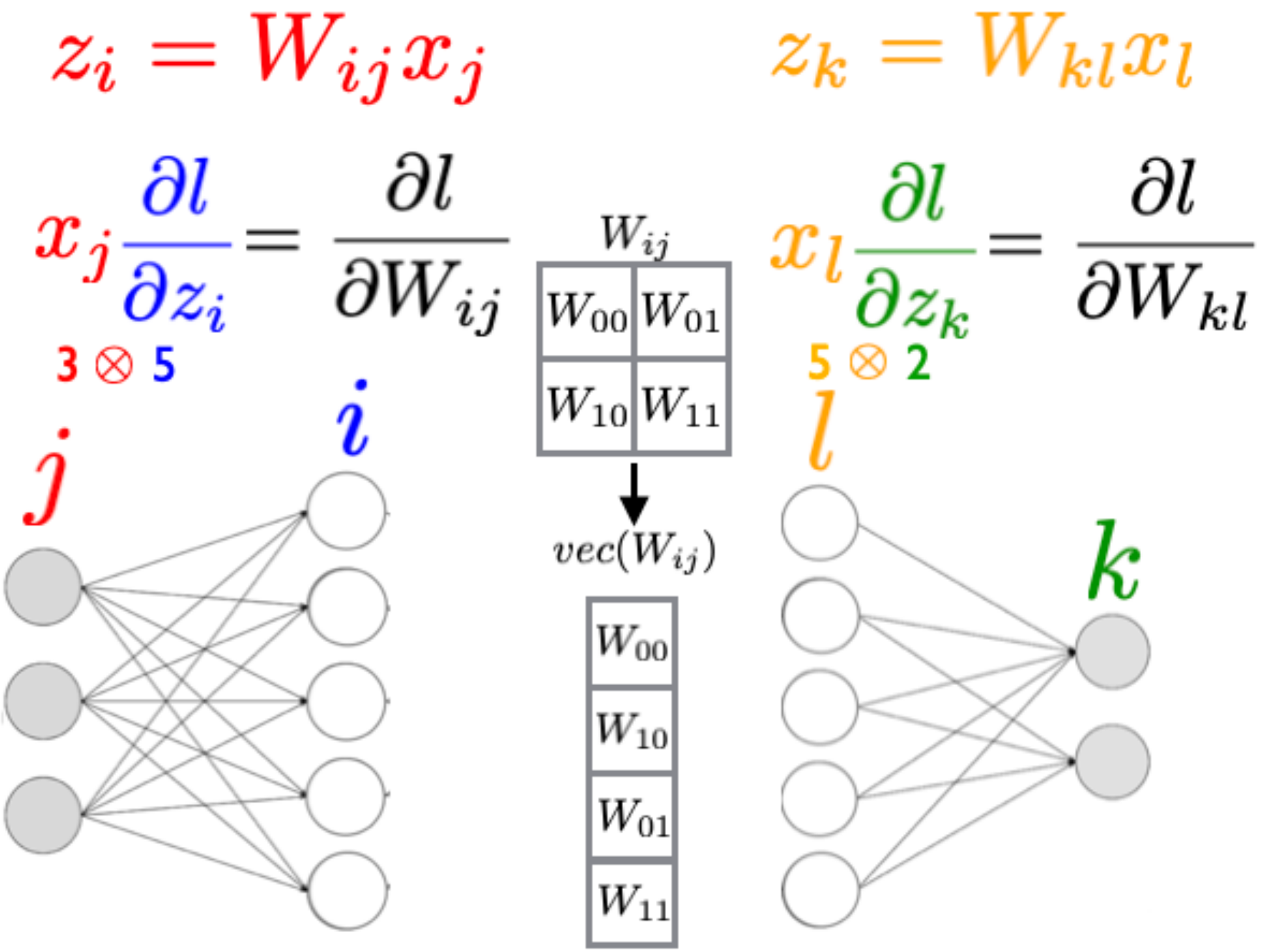
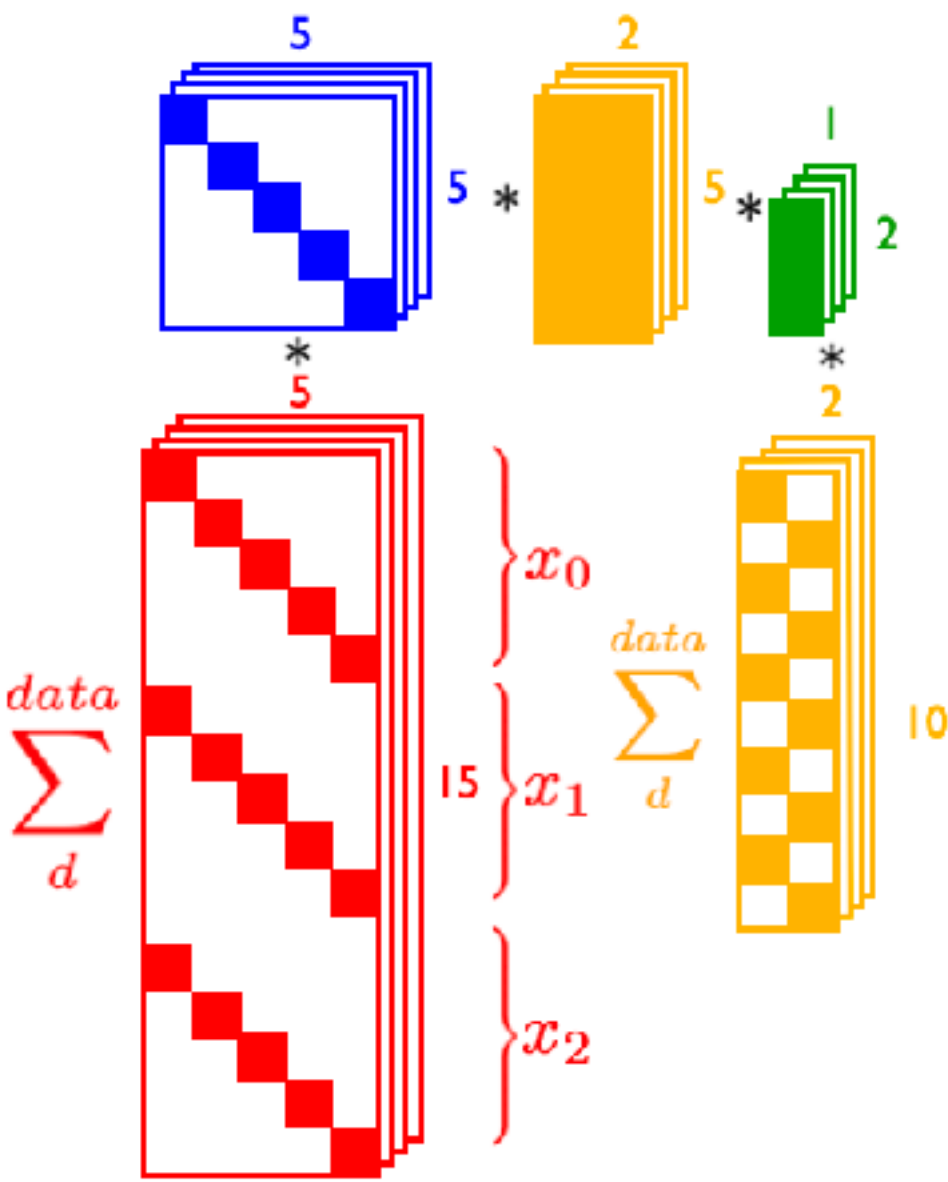
Forward propagation

$x_0 \rightarrow z_0 = W_0 x_0$ $x_1 = \text{ReLU}(z_0) \rightarrow z_1 = W_1 x_1$ $p = \text{softmax}(z_1)$

Backward propagation

$$\frac{\partial x_1}{\partial z_0} \otimes \frac{\partial z_1}{\partial x_1} \otimes \frac{\partial l}{\partial z_1} = \frac{\partial l}{\partial \theta}$$

$\frac{\partial z_0}{\partial W_0}$ $\frac{\partial z_1}{\partial W_1}$



$z_0 = W_{00}x_0 + W_{01}x_1 + W_{02}x_2$
 $z_1 = W_{10}x_0 + W_{11}x_1 + W_{12}x_2$
 $z_2 = W_{20}x_0 + W_{21}x_1 + W_{22}x_2$
 $z_3 = W_{30}x_0 + W_{31}x_1 + W_{32}x_2$
 $z_4 = W_{40}x_0 + W_{41}x_1 + W_{42}x_2$

Approximation of the Hessian/Fisher Matrix

$$l = -\log p$$
$$\frac{\partial l}{\partial p} = -\frac{1}{p}$$
$$\frac{\partial^2 l}{\partial p^2} = \frac{1}{p^2} = \left(\frac{\partial l}{\partial p}\right)^2$$

$$F = \sum^{samples} \left(\frac{\partial l}{\partial W}\right) \left(\frac{\partial l}{\partial W}\right)^T$$
$$\approx \left(\sum^{samples} \frac{\partial l}{\partial W}\right) \left(\sum^{samples} \frac{\partial l}{\partial W}\right)^T$$

$\frac{\partial l}{\partial W}$

Newton-Schultz
 $= USV$

Muon (2025)

$\frac{\partial l}{\partial W}$

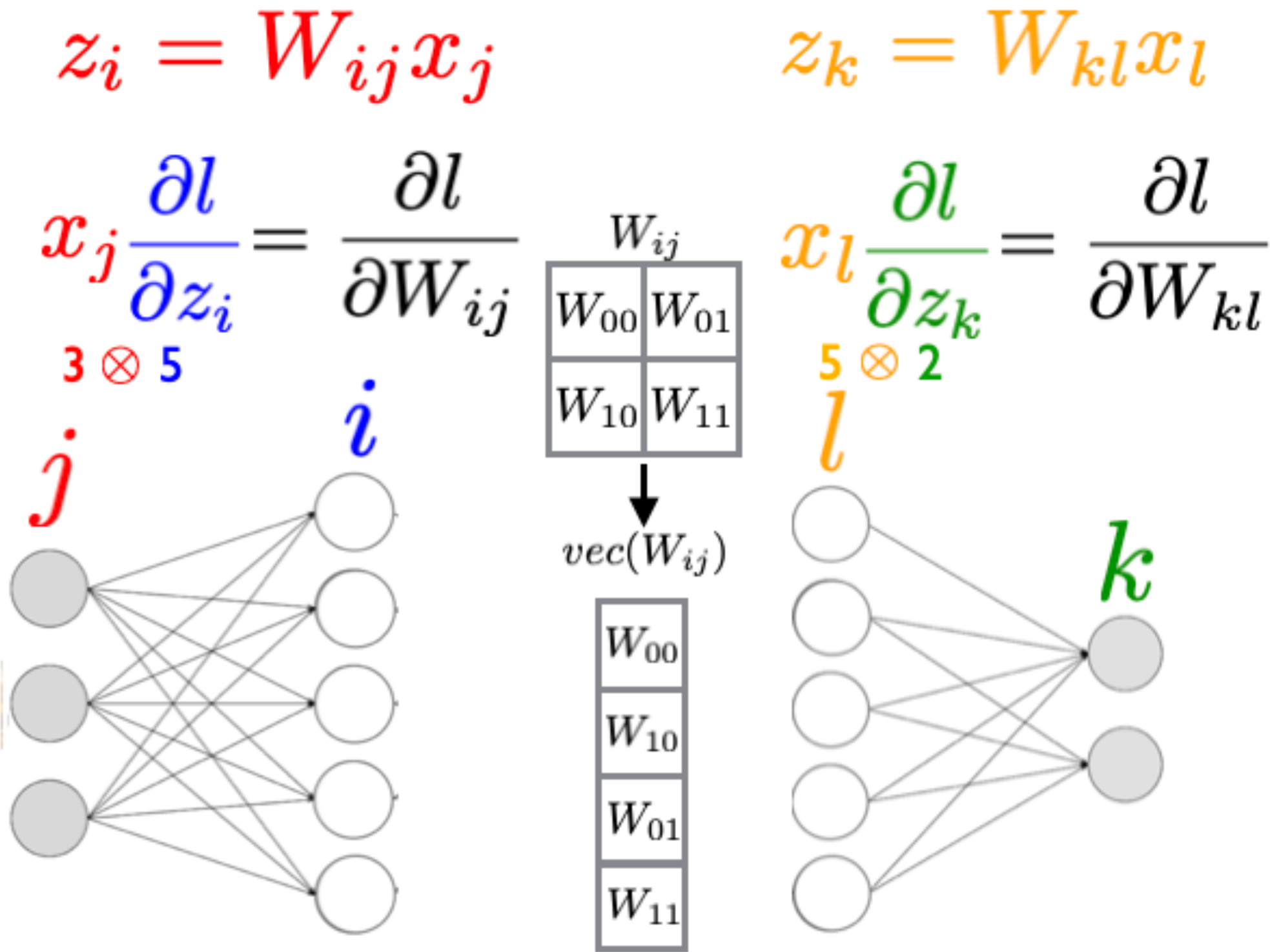
$\frac{\partial l^T}{\partial W}$

Shampoo (2020)
Soap (2023)

$\frac{\partial l}{\partial W}$

IM

IK



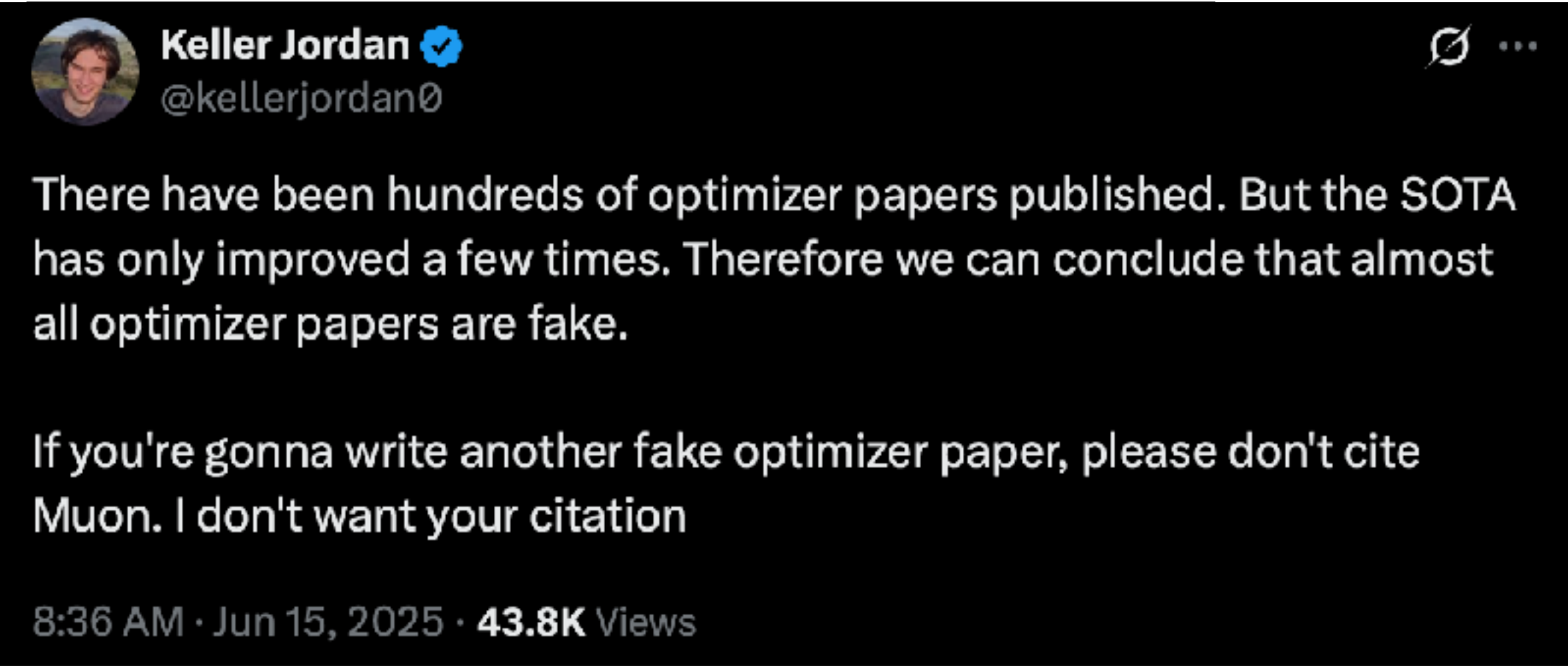
$$F = \sum^{samples} \left(\frac{\partial l}{\partial W}\right) \left(\frac{\partial l}{\partial W}\right)^T$$
$$= \sum^{samples} \left(x \otimes \frac{\partial l}{\partial z}\right) \left(x \otimes \frac{\partial l}{\partial z}\right)^T$$
$$= \sum^{samples} (xx^T) \otimes \left(\frac{\partial l}{\partial z} \frac{\partial l^T}{\partial z}\right)$$
$$\approx \left(\sum^{samples} xx^T\right) \otimes \left(\sum^{samples} \frac{\partial l}{\partial z} \frac{\partial l^T}{\partial z}\right)$$

xx^T

$\frac{\partial l}{\partial z} \frac{\partial l^T}{\partial z}$

IK

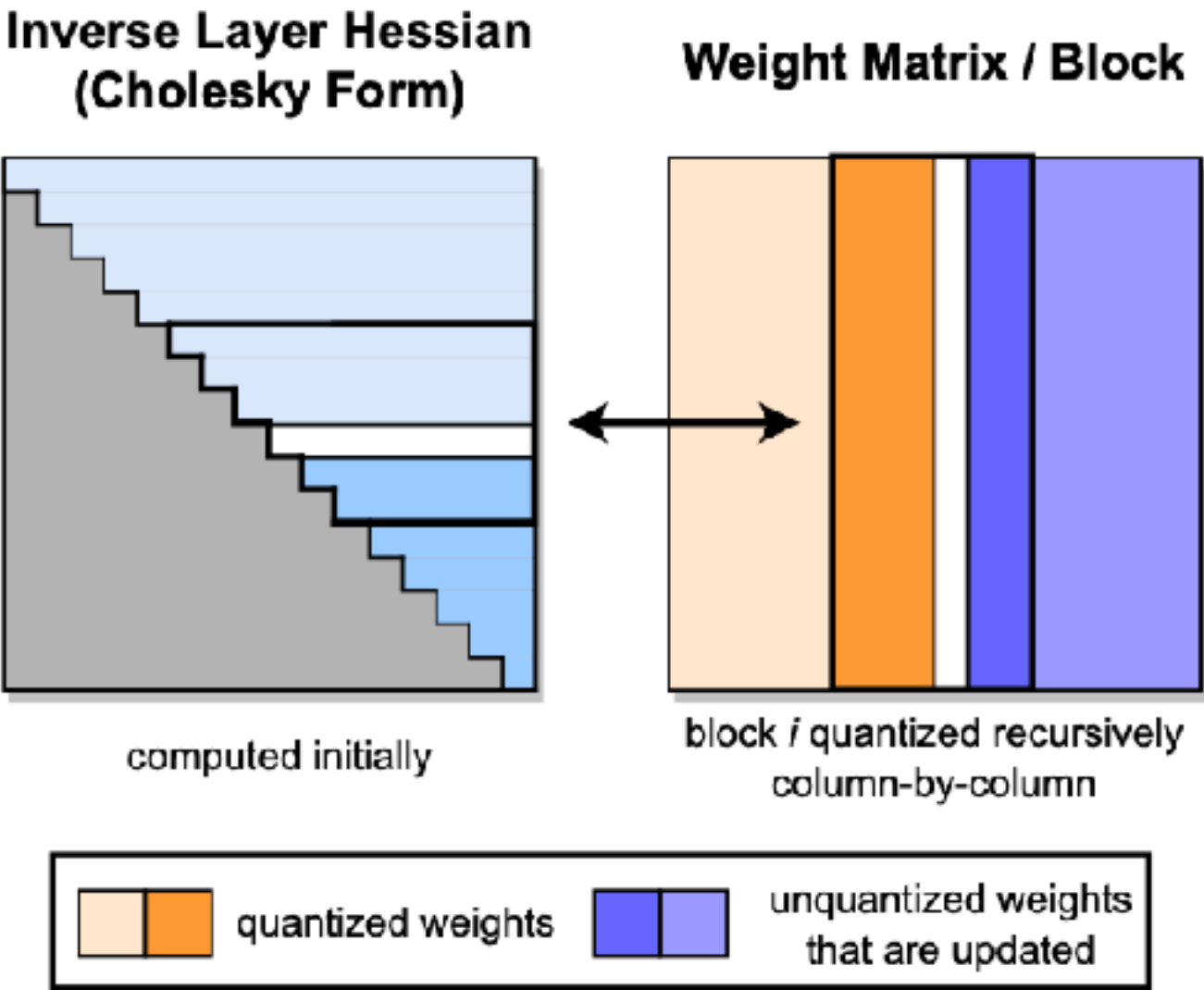
K-FAC (2015)



Using the Matrix for things other than optimization

Quantization

$$\mathbf{H}^{-1} = (2\mathbf{X}\mathbf{X}^\top + \lambda\mathbf{I})^{-1}$$



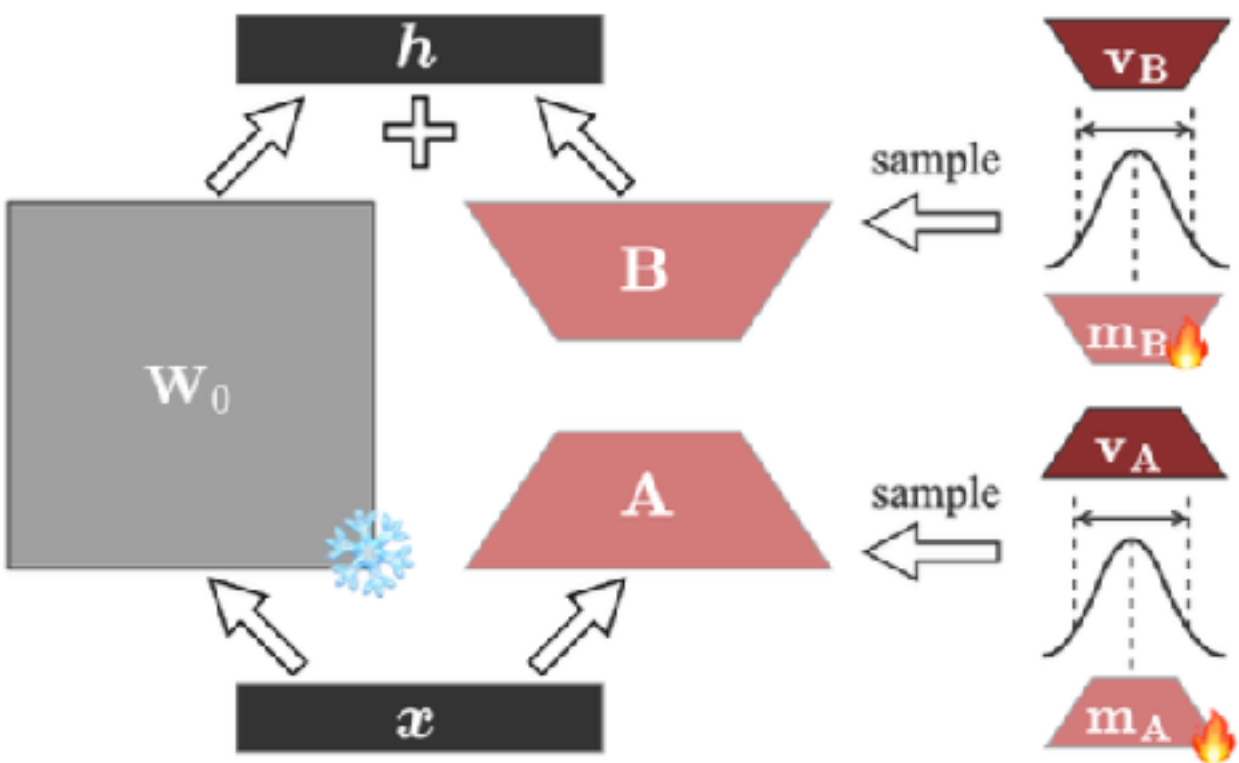
Published as a conference paper at ICLR 2023

GPTQ: ACCURATE POST-TRAINING QUANTIZATION FOR GENERATIVE PRE-TRAINED TRANSFORMERS

Elias Frantar* IST Austria Saleh Ashkboos ETH Zurich Torsten Hoefer ETH Zurich Dan Alistarh IST Austria & NeuralMagic

Pruning

$$\hat{\mathbf{h}} \leftarrow \hat{\nabla} \bar{\ell}(\boldsymbol{\theta}) \cdot \frac{\boldsymbol{\theta} - \mathbf{m}}{\sigma^2}$$



Improving LoRA with Variational Learning

Bai Cong^{1,2} Nico Daheim³ Yuesong Shen⁴
Ryo Yokota¹ Mohammad Emtiyaz Khan² Thomas Möllenhoff²

¹Institute of Science Tokyo ²RIKEN Center for AI Project
³Technical University of Darmstadt ⁴Huawei Hong Kong Research Center

Summary

- Various matrix structures are used in AI
- Matrix structures in HPC do not translate directly to AI
- HPC Matrices
 - Exploit structure in the underlying 3-D geometry
 - Spatial decomposition
- AI Matrices
 - Exploit structure between the billions of dimensions
 - Backprop naturally contains a Kronecker structure
- Changing the model architecture is often the simplest way to improve AI