

SCALABLE KRONECKER-FACTORED FISHER APPROXIMATION FOR NEURAL NETWORK PARAMETER SENSITIVITY



CODE AND MODELS

Viktoriia Chekalina^{1,2}, Daniil Moskovskiy^{3,4}, Tatyana Matveeva⁵,
Andrey Kuznetsov^{1,6}, Evgeny Frolov^{3,5}

¹Fusion Brain, ²MSU, ³AXXX, ⁴Applied AI Institute, ⁵HSE University, ⁶Innopolis University

Task description

Structural data — such as layer weights and activations — can be modeled using a matrix normal distribution:

$$p(\mathbf{W}) = \frac{1}{(2\pi)^{nm/2} |\Sigma_1 \otimes \Sigma_2|^{1/2}} \exp\left(-\frac{1}{2}(\mathbf{W} - \mu)^\top (\Sigma_1 \otimes \Sigma_2)^{-1} (\mathbf{W} - \mu)\right) \quad (1)$$

with:

- Σ_1 models **row-wise dependencies**, Σ_2 models **column-wise dependencies**.
- $\mathcal{I}(\theta) = \Sigma_{\text{col}}^{-1} \otimes \Sigma_{\text{row}}^{-1}$ — Fisher Information (Hessian for PTQ tasks).
- Correlation information is encoded in the **off-diagonal elements**, so the **full Hessian** needs to be considered.

Matrix-Free Fisher Factorization (MFF): The Idea

Problem. The **full Hessian** of a linear layer has quadratic size relative to the layer matrix — it cannot be loaded into GPU memory.

- **Idea.** Kronecker decomposition reduces to truncated SVD, and SVD reduces to matvec / rmatvec products.
- We exploit the following property:

$$(\mathbf{K} \otimes \mathbf{L}) \text{vec}(\mathbf{C}) = \text{vec}(\mathbf{K}^\top \mathbf{C} \mathbf{L})$$

- Matrix–vector product becomes (full Hessian is never materialized):

$$\tilde{\mathcal{L}}_F \mathbf{z} = \frac{1}{|\mathcal{D}|} \sum_{i=1}^{|\mathcal{D}|} \text{vec}(\mathbf{G}_i^\top \mathbf{z} \mathbf{G}_i)$$

Time Complexity

- Theoretical complexity reduction: $\mathcal{O}(m^2 n^2) \rightarrow \mathcal{O}(mn^2 + m^2 n)$
- Practical speedup:

Model	Params in layer	Params in Hess	Decomp. time (s)
BERT	2.3×10^6	5.5×10^{12}	43
Llama2 7B	45×10^6	2.0×10^{15}	183
Llama3.1 8B	58×10^6	3.4×10^{15}	249
Llama2 13B	70.8×10^6	4.9×10^{15}	313

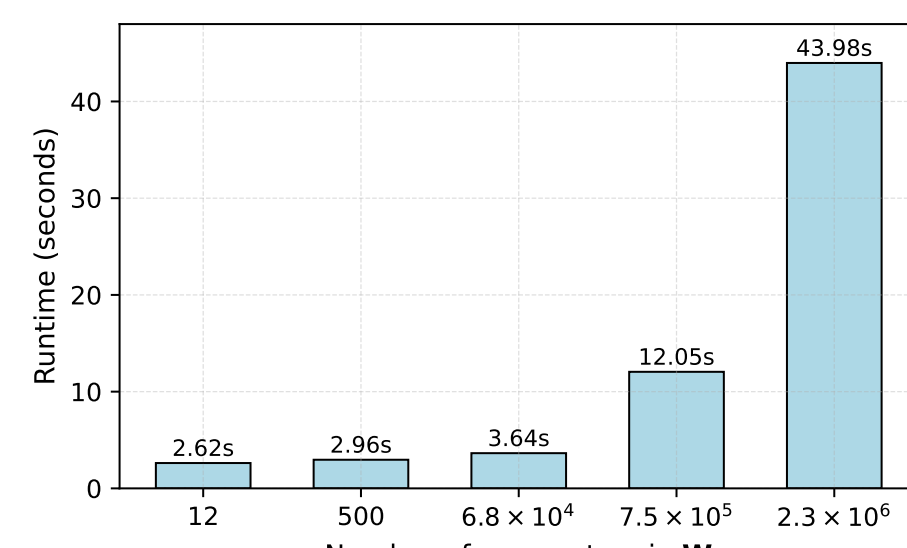


Figure 1: Empirical runtime for Kronecker decomposition.

Generalized Fisher-Weighted SVD (GFWSVD)

- We find factors Σ_1, Σ_2 such that $\mathbf{L}_{\Sigma_1}^\top \mathbf{L}_{\Sigma_1} = \Sigma_1$ and $\mathbf{L}_{\Sigma_2}^\top \mathbf{L}_{\Sigma_2} = \Sigma_2$, which gives the curvature-aware low-rank approximation:

$$\tilde{\mathbf{W}}_r = \text{SVD}_r(\mathbf{L}_{\Sigma_1}^\top \mathbf{W} \mathbf{L}_{\Sigma_2})$$

- **SVD-based pruning:** a linear layer $\mathbf{W} \in \mathbb{R}^{n \times m}$ factorizes into two smaller layers

$$\mathbf{W}_2 = \mathbf{U}_r \sqrt{\Sigma_r} \in \mathbb{R}^{r \times m}, \quad \mathbf{W}_1 = \sqrt{\Sigma_r} \mathbf{V}_r^\top \in \mathbb{R}^{n \times r}$$

- This SVD pruning is a **generalization** of weighted-pruning methods (proof in the paper).

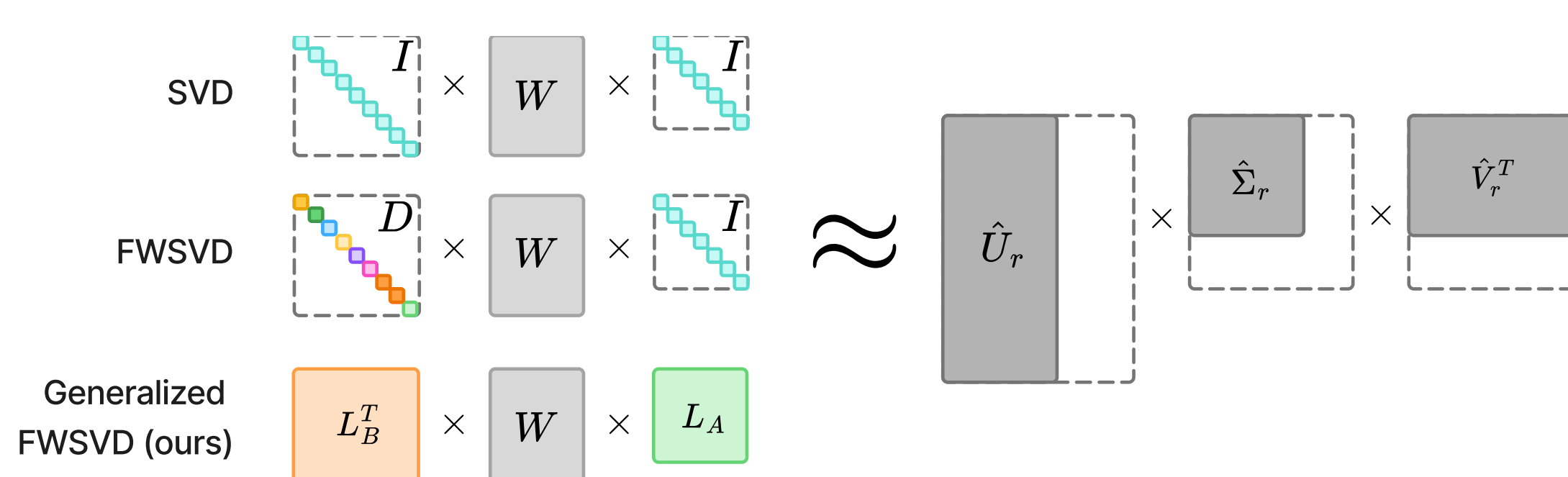


Figure 2: Generalization of the weighted SVD framework. For standard SVD both transformation matrices are identity; for FWSVD the left matrix is diagonal (non-identity) and the right is identity; for GFWSVD both matrices are non-diagonal.

GFWSVD: Off-Diagonal FIM on Synthetic Data

To isolate off-diagonal FIM elements, we train MLPs from scratch on binary classification under three covariance regimes: isotropic, anisotropic, and correlated.

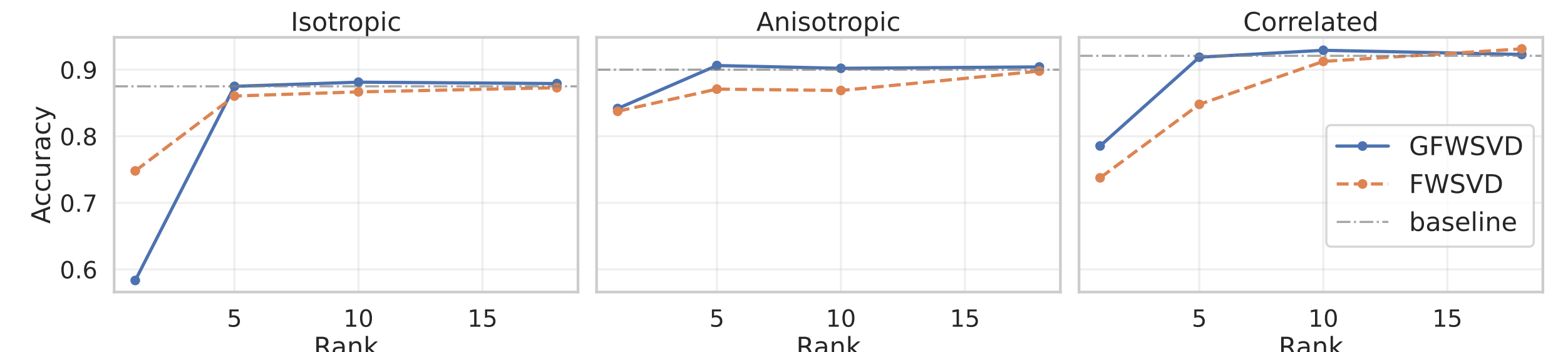


Figure 3: Comparison of diagonal and full Kronecker decompositions of a linear layer in an MLP.

Full Kronecker decomposition captures cross-feature curvature that the diagonal one ignores. The gap **widens at higher compression**, and is **largest in the correlated regime**.

GFWSVD: Results on BERT Encoders (GLUE)

Compressing BERT across ranks and measuring macro-averaged GLUE, our method shows **supremacy at strong compression** — the more aggressive the compression, the larger our advantage.

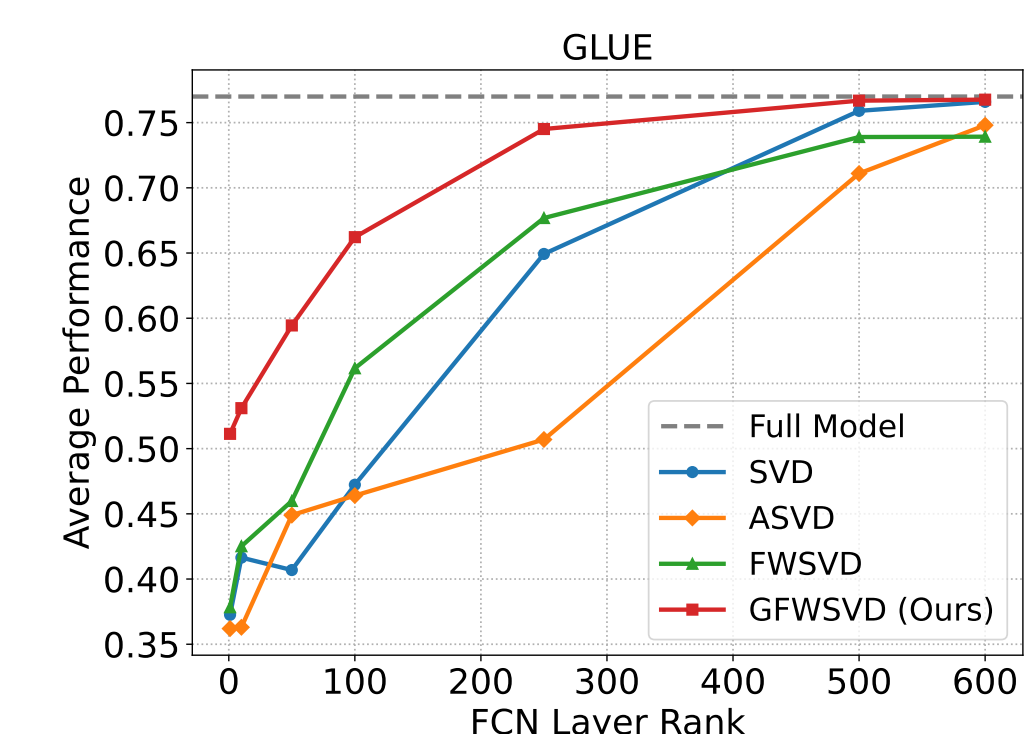


Figure 4: Macro-averaged GLUE performance of BERT for different compression ranks.

GFWSVD: Results on LLaMA-3.1-8B (Common Sense Reasoning)

LLaMA-3.1-8B Instruct is a **dense model with correlated weights** — exactly the regime where modeling off-diagonal curvature pays off. We compress it at ratios from **20% to 50%**.

Method	C. Ratio	ARC-C ↑	ARC-E ↑	HellaSwag ↑	PIQA ↑	WinoG. ↑	AVG ↑
Full model	0%	0.52	0.81	0.59	0.79	0.73	0.63
FWSVD	20%	0.21	0.38	0.20	0.60	0.52	0.35
ASVD		0.21	0.33	0.27	0.61	0.54	0.35
B. Sharing		0.34	0.68	0.42	0.70	0.65	0.52
GFWSVD		0.35	0.68	0.45	0.75	0.63	0.53
FWSVD	30%	0.21	0.30	0.26	0.57	0.51	0.33
ASVD		0.22	0.30	0.25	0.58	0.52	0.34
B. Sharing		0.29	0.52	0.43	0.63	0.60	0.46
GFWSVD		0.33	0.61	0.42	0.71	0.58	0.48
FWSVD	40%	0.21	0.27	0.26	0.54	0.48	0.32
ASVD		0.23	0.27	0.26	0.55	0.49	0.33
B. Sharing		0.24	0.39	0.33	0.56	0.56	0.39
GFWSVD		0.25	0.41	0.32	0.61	0.55	0.39
FWSVD	50%	0.20	0.27	0.26	0.50	0.51	0.31
ASVD		0.22	0.26	0.26	0.50	0.48	0.31
B. Sharing		0.23	0.30	0.29	0.52	0.53	0.35
GFWSVD		0.24	0.31	0.28	0.55	0.54	0.36

Across every compression ratio, GFWSVD attains the best average common-sense reasoning accuracy.