

Learning incomplete factorization preconditioners for GMRES

Aleix Nieto Juscafresa^{*}, Uppsala University, Sweden
Joint work with Paul Häusner^{*} and Jens Sjölund

^{*}Joint first authors

Numerical Analysis and Scientific Computation with Applications (NASCA26)

June 11, 2026



UPPSALA
UNIVERSITET

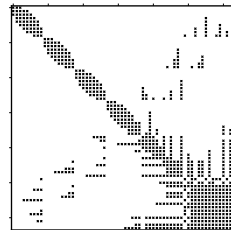
The problem

Solving $Ax = b$

PROBLEM

Solve $Ax = b$ *efficiently*, where $A \in \mathbb{R}^{n \times n}$ is **large** and **sparse**.

- ▶ A is **sparse**: $\text{nnz}(A) \ll n^2$
- ▶ Arises in PDE discretizations, optimization problems (Newton, IPMs), Gaussian processes, ...
- ▶ Often **many related systems** share sparsity pattern, e.g. parameter studies, optimization iterations, ...



Sparsity pattern

Direct vs. iterative solvers

✓ Direct solvers

Compute the solution via matrix inversion or factorization (LU, QR, Cholesky).

- ▶ Accurate and reliable
- ▶ **Do not fully exploit sparsity**: factorization introduces **fill-in**, increasing memory and computational cost

Direct vs. iterative solvers

✔ Direct solvers

Compute the solution via matrix inversion or factorization (LU, QR, Cholesky).

- ▶ Accurate and reliable
- ▶ **Do not fully exploit sparsity:** factorization introduces **fill-in**, increasing memory and computational cost

↻ Iterative solvers (e.g. GMRES)

Refine an initial guess to approximate the solution, efficient when approximations are sufficient.

- ▶ Cheap per iteration, low memory
- ▶ **Exploit sparsity:** dominated by sparse matrix-vector products; scales with $\text{nnz}(\mathbf{A})$

Direct vs. iterative solvers

✔ Direct solvers

Compute the solution via matrix inversion or factorization (LU, QR, Cholesky).

- ▶ Accurate and reliable
- ▶ **Do not fully exploit sparsity**: factorization introduces **fill-in**, increasing memory and computational cost

↻ Iterative solvers (e.g. GMRES)

Refine an initial guess to approximate the solution, efficient when approximations are sufficient.

- ▶ Cheap per iteration, low memory
- ▶ **Exploit sparsity**: dominated by sparse matrix-vector products; scales with $\text{nnz}(\mathbf{A})$

BOTTLENECK

Iterative solvers scale, but their convergence depends on the **spectrum** of $\mathbf{A}$.

GMRES

DEFINITION

For a matrix A and a vector v , the k -th **Krylov subspace** $\mathcal{K}_k(A, v)$ is defined as

$$\text{span}\{v, Av, \dots, A^{k-1}v\}.$$

¹Yousef Saad: GMRES: A Generalized Minimal Residual Algorithm for Solving Non-symmetric Linear Systems, 1986.

GMRES

DEFINITION

For a matrix A and a vector v , the k -th **Krylov subspace** $\mathcal{K}_k(A, v)$ is defined as

$$\text{span}\{v, Av, \dots, A^{k-1}v\}.$$

- ▶ GMRES¹ searches for increasingly better approximations within expanding Krylov subspaces
- ▶ At step k , it minimizes $\|b - Ax\|_2$ over $x_0 + \mathcal{K}_k(A, r_0)$, where $r_0 = b - Ax_0$

¹Yousef Saad: GMRES: A Generalized Minimal Residual Algorithm for Solving Non-symmetric Linear Systems, 1986.

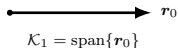
GMRES

DEFINITION

For a matrix A and a vector v , the k -th **Krylov subspace** $\mathcal{K}_k(A, v)$ is defined as

$$\text{span}\{v, Av, \dots, A^{k-1}v\}.$$

- ▶ GMRES¹ searches for increasingly better approximations within expanding Krylov subspaces
- ▶ At step k , it minimizes $\|b - Ax\|_2$ over $x_0 + \mathcal{K}_k(A, r_0)$, where $r_0 = b - Ax_0$



A horizontal line with an arrow pointing to the right, labeled r_0 at the tip. Below the line, the text $\mathcal{K}_1 = \text{span}\{r_0\}$ is written.

¹Yousef Saad: GMRES: A Generalized Minimal Residual Algorithm for Solving Non-symmetric Linear Systems, 1986.

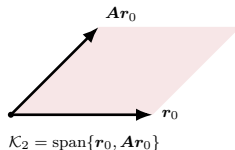
GMRES

DEFINITION

For a matrix A and a vector v , the k -th **Krylov subspace** $\mathcal{K}_k(A, v)$ is defined as

$$\text{span}\{v, Av, \dots, A^{k-1}v\}.$$

- ▶ GMRES¹ searches for increasingly better approximations within expanding Krylov subspaces
- ▶ At step k , it minimizes $\|b - Ax\|_2$ over $x_0 + \mathcal{K}_k(A, r_0)$, where $r_0 = b - Ax_0$



¹Yousef Saad: GMRES: A Generalized Minimal Residual Algorithm for Solving Non-symmetric Linear Systems, 1986.

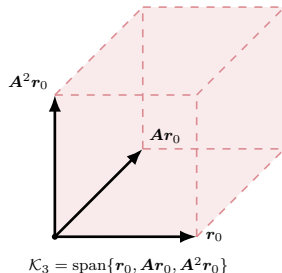
GMRES

DEFINITION

For a matrix A and a vector v , the k -th **Krylov subspace** $\mathcal{K}_k(A, v)$ is defined as

$$\text{span}\{v, Av, \dots, A^{k-1}v\}.$$

- ▶ GMRES¹ searches for increasingly better approximations within expanding Krylov subspaces
- ▶ At step k , it minimizes $\|b - Ax\|_2$ over $x_0 + \mathcal{K}_k(A, r_0)$, where $r_0 = b - Ax_0$



¹Yousef Saad: GMRES: A Generalized Minimal Residual Algorithm for Solving Non-symmetric Linear Systems, 1986.

Why GMRES?

- ▶ One of the most widely used Krylov subspace methods for **general** linear systems $Ax = b$
- ▶ No structure assumed on $A \rightarrow$ works for **arbitrary** (non-symmetric) matrices

Why GMRES?

- ▶ One of the most widely used Krylov subspace methods for **general** linear systems $Ax = b$
- ▶ No structure assumed on $A \rightarrow$ works for **arbitrary** (non-symmetric) matrices

ASSUMPTIONS ON A

- ▶ **Invertible** ($A \in \text{GL}(n, \mathbb{R})$)
- ▶ **Sparse**

Preconditioning

GMRES is broadly applicable, but its performance depends heavily on preconditioning.

Preconditioning

IDEA

Transform the original system $Ax = b$ into an **equivalent system**, with improved spectral properties

$$AP^{-1}z = b, \quad x = P^{-1}z.$$

Preconditioning

IDEA

Transform the original system $Ax = b$ into an **equivalent system**, with improved spectral properties

$$AP^{-1}z = b, \quad x = P^{-1}z.$$

- ▶ $P = I$: no preconditioning
- ▶ $P = A$: direct method
- ▶ **Classical choices:** Jacobi, incomplete factorizations (ILU, IC)

Preconditioning

IDEA

Transform the original system $Ax = b$ into an **equivalent system**, with improved spectral properties

$$AP^{-1}z = b, \quad x = P^{-1}z.$$

- ▶ $P = I$: no preconditioning
- ▶ $P = A$: direct method
- ▶ **Classical choices:** Jacobi, incomplete factorizations (ILU, IC)

But: time to compute the preconditioner P
vs.
 acceleration obtained in the iterative solver

Preconditioning

IDEA

Transform the original system $Ax = b$ into an **equivalent system**, with improved spectral properties

$$AP^{-1}z = b, \quad x = P^{-1}z.$$

- ▶ $P = I$: no preconditioning
- ▶ $P = A$: direct method
- ▶ **Classical choices:** Jacobi, incomplete factorizations (ILU, IC)

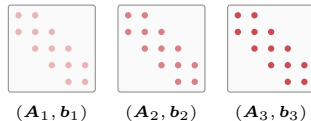
But: time to compute the preconditioner P
vs.
acceleration obtained in the iterative solver

Classical preconditioners are *hand-designed heuristics*!

Research question

Research question

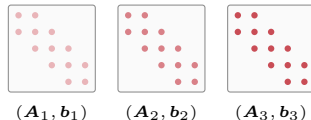
- ▶ Many real-world problems give rise to **related** linear systems
- ▶ Classical preconditioners (ILU, Jacobi) ignore this structure
- ▶ Prior **learned** preconditioners assume **SPD** matrices (CG), *general non-symmetric systems remain open*
- ▶ A learned model can **amortize** training cost across solves



Same sparsity pattern,
different coefficients

Research question

- ▶ Many real-world problems give rise to **related** linear systems
- ▶ Classical preconditioners (ILU, Jacobi) ignore this structure
- ▶ Prior **learned** preconditioners assume **SPD** matrices (CG), *general non-symmetric systems remain open*
- ▶ A learned model can **amortize** training cost across solves



Same sparsity pattern,
different coefficients

Can we learn a preconditioner from data that accelerates GMRES and competes with classical heuristics on related problems?

Our approach

Learning-to-optimize^{2,3}

1. Use machine learning to accelerate numerical solvers by exploiting cross-problem similarity.
2. We view preconditioner construction as an amortized optimization problem.

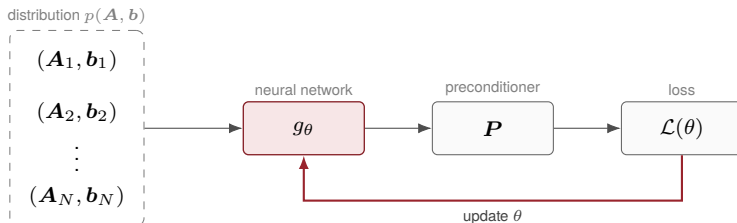
²Brandon Amos, Tutorial on amortized optimization, Foundations and Trends in Machine Learning, 2023.

³Tianlong Chen et al., Learning to optimize: a primer and a benchmark, JMLR, 2022.

Learning-to-optimize^{2,3}

1. Use machine learning to accelerate numerical solvers by exploiting cross-problem similarity.
2. We view preconditioner construction as an amortized optimization problem.

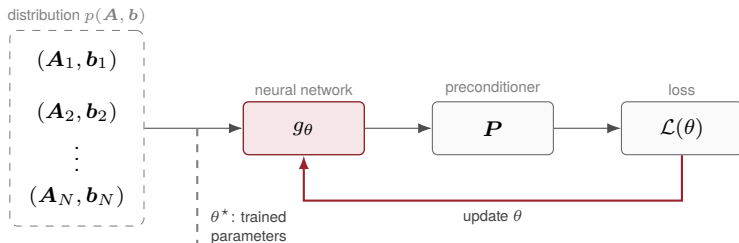
Training (offline)



Learning-to-optimize^{2,3}

1. Use machine learning to accelerate numerical solvers by exploiting cross-problem similarity.
2. We view preconditioner construction as an amortized optimization problem.

Training (offline)



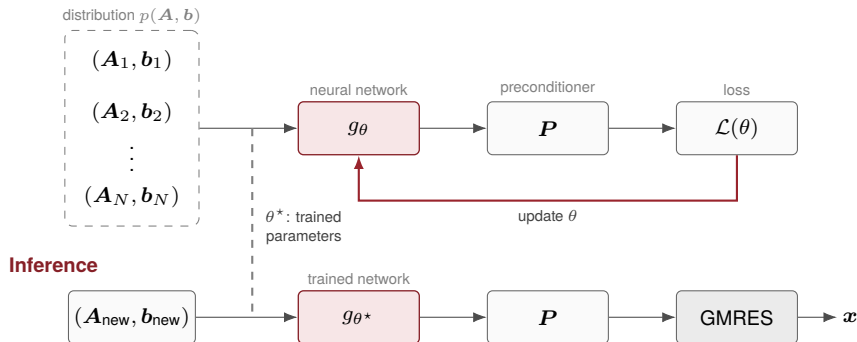
Inference



Learning-to-optimize^{2,3}

1. Use machine learning to accelerate numerical solvers by exploiting cross-problem similarity.
2. We view preconditioner construction as an amortized optimization problem.

Training (offline)

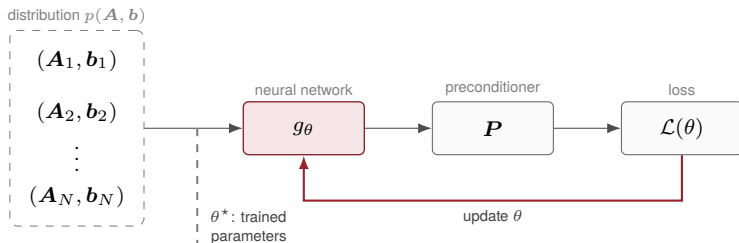


Train **once, offline**, then generate preconditioners for new problems with **low inference cost**.

Learning-to-optimize^{2,3}

1. Use machine learning to accelerate numerical solvers by exploiting cross-problem similarity.
2. We view preconditioner construction as an amortized optimization problem.

Training (offline)



Inference

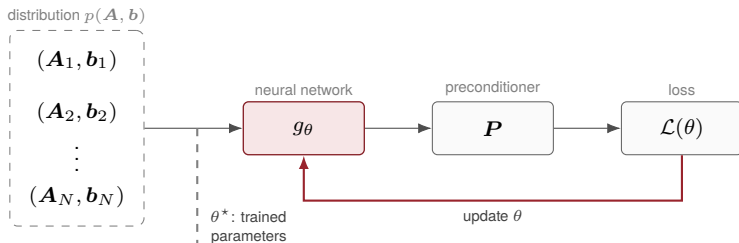


But what is g_θ ?

Learning-to-optimize^{2,3}

1. Use machine learning to accelerate numerical solvers by exploiting cross-problem similarity.
2. We view preconditioner construction as an amortized optimization problem.

Training (offline)



Inference



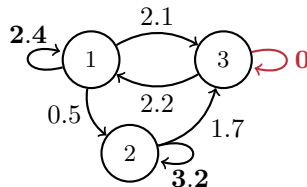
But what is g_θ ? → **a graph neural network.**

Graph neural networks and linear algebra

Graph neural networks and linear algebra

- Every matrix A defines a weighted directed graph (**Coates** representation)

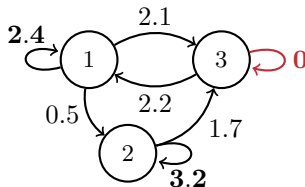
$$\begin{bmatrix} 2.4 & 0 & 2.2 \\ 0.5 & 3.2 & 0 \\ 2.1 & 1.7 & 0 \end{bmatrix}$$



Graph neural networks and linear algebra

- Every matrix A defines a weighted directed graph (**Coates** representation)

$$\begin{bmatrix} 2.4 & 0 & 2.2 \\ 0.5 & 3.2 & 0 \\ 2.1 & 1.7 & 0 \end{bmatrix}$$

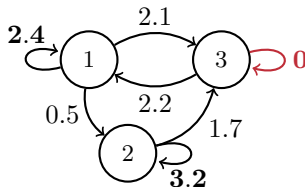


Rows / columns of $A \leftrightarrow$ **nodes**

Graph neural networks and linear algebra

- Every matrix A defines a weighted directed graph (**Coates** representation)

$$\begin{bmatrix} 2.4 & 0 & 2.2 \\ 0.5 & 3.2 & 0 \\ 2.1 & 1.7 & 0 \end{bmatrix}$$

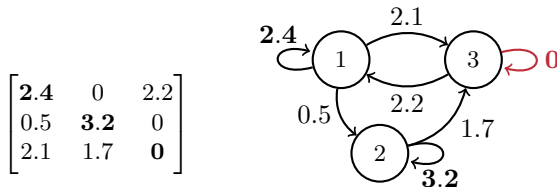


Rows / columns of $A \leftrightarrow$ **nodes**

Non-zero entries $\leftrightarrow$ **edges**

Graph neural networks and linear algebra

- Every matrix A defines a weighted directed graph (**Coates** representation)

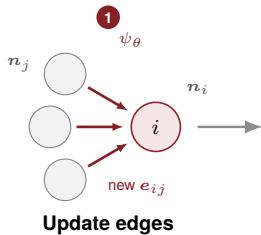


Rows / columns of $A \leftrightarrow$ **nodes**

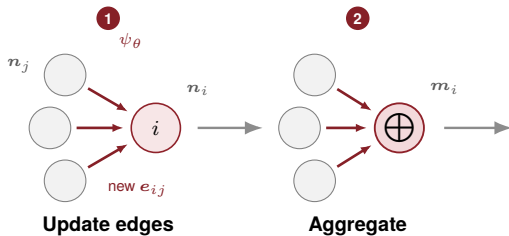
Non-zero entries $\leftrightarrow$ **edges**

Value of each entry $\leftrightarrow$ **edge feature**

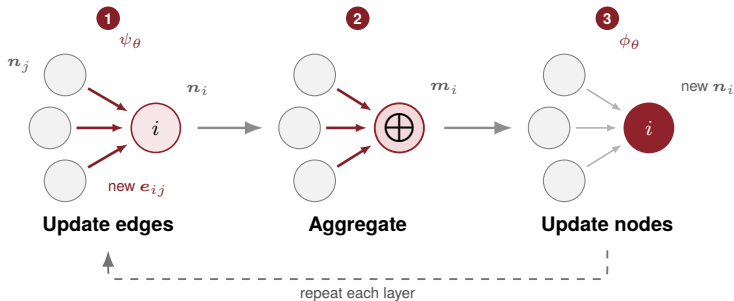
Graph neural networks (GNNs)



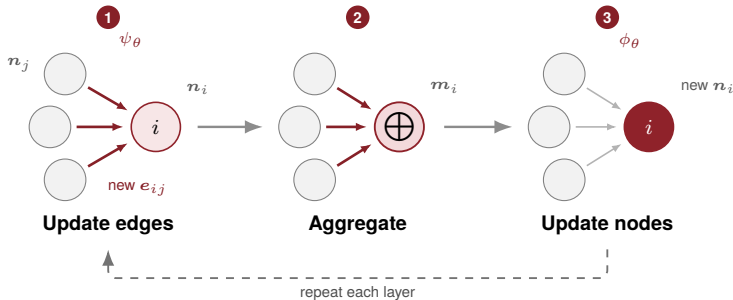
Graph neural networks (GNNs)



Graph neural networks (GNNs)



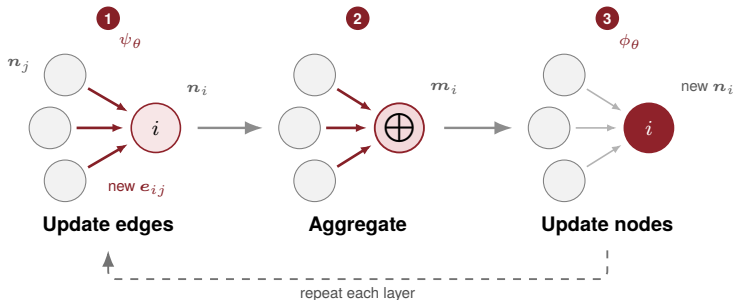
Graph neural networks (GNNs)



A GNN repeats **message passing** each layer, updating *both* edge and node features.

Node features n_i on nodes, edge features e_{ij} on edges (= matrix entries A_{ij}).

Graph neural networks (GNNs)



A GNN repeats **message passing** each layer, updating *both* edge and node features.

Node features n_i on nodes, edge features e_{ij} on edges (= matrix entries A_{ij}).

The $\oplus$ is **permutation-invariant** — therefore independent of node ordering.

Learning factorization preconditioners

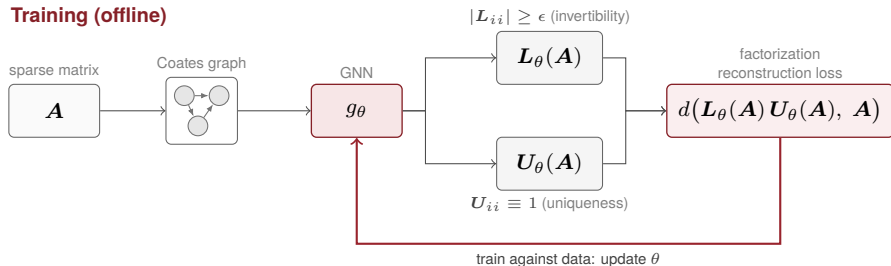
Learned preconditioners

Goal: learn a parameterized mapping g_θ that takes a matrix A and predicts a suitable preconditioner.

Learned preconditioners

Goal: learn a parameterized mapping g_θ that takes a matrix A and predicts a suitable preconditioner.

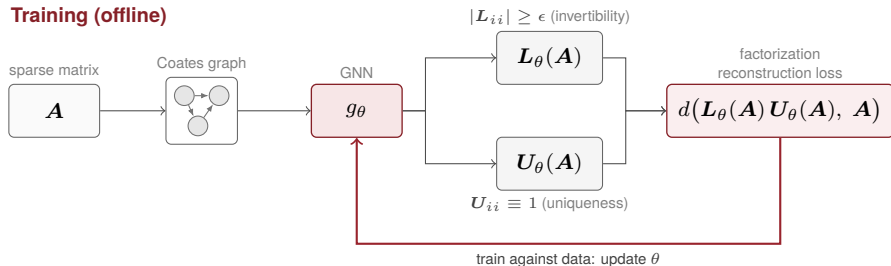
Training (offline)



Learned preconditioners

Goal: learn a parameterized mapping g_θ that takes a matrix A and predicts a suitable preconditioner.

Training (offline)

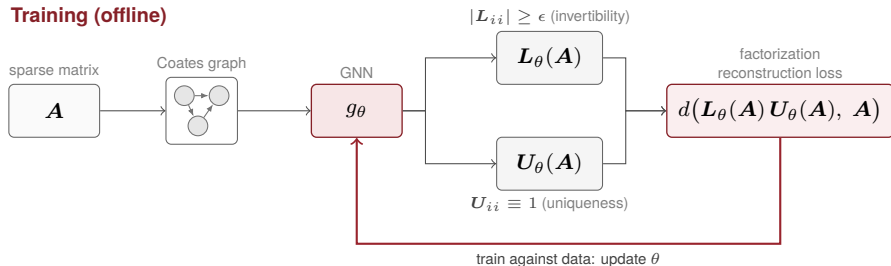


- ▶ **One GNN, two factors:** edges with $i > j$ form L , edges with $i < j$ form U
- ▶ **Invertibility:** non-zero diagonals make L, U invertible $\Rightarrow P = LU$ is a valid preconditioner
- ▶ **Why factorize?** Triangular $L, U \rightarrow$ **efficient sparse triangular solves**; forward/back-substitution

Learned preconditioners

Goal: learn a parameterized mapping g_θ that takes a matrix A and predicts a suitable preconditioner.

Training (offline)



The sparsity of L and U is inherited from A via the GNN architecture⁴
$$\begin{cases} L_{ij} = 0 & \text{if } A_{ij} = 0, \\ U_{ij} = 0 & \text{if } A_{ij} = 0. \end{cases}$$

⁴P. Häusner et al., Neural incomplete factorization: learning preconditioners for CG, TMLR, 2024.

Loss functions

We design loss functions such that $\mathbf{P} = \mathbf{LU} \approx \mathbf{A}$ and analyze their effect on the singular values of the preconditioned system $\mathbf{AP}^{-1}$:

$$\mathcal{L}_{\max}(\mathbf{P}, \mathbf{A}) = \|\mathbf{P} - \mathbf{A}\|_F$$

Reduces large singular values.

- + Easy to compute
- Classical preconditioners are good at it

$$\mathcal{L}_{\min}(\mathbf{P}, \mathbf{A}) = \|\mathbf{PA}^{-1} - \mathbf{I}\|_F$$

Encourages larger small singular values.

- + Small singular values affect convergence
- Requires computations with the inverse

Loss functions

We design loss functions such that $\mathbf{P} = \mathbf{LU} \approx \mathbf{A}$ and analyze their effect on the singular values of the preconditioned system $\mathbf{AP}^{-1}$:

$$\mathcal{L}_{\max}(\mathbf{P}, \mathbf{A}) = \|\mathbf{P} - \mathbf{A}\|_F$$

Reduces large singular values.

- + Easy to compute
- Classical preconditioners are good at it

$$\mathcal{L}_{\min}(\mathbf{P}, \mathbf{A}) = \|\mathbf{PA}^{-1} - \mathbf{I}\|_F$$

Encourages larger small singular values.

- + Small singular values affect convergence
- Requires computations with the inverse

Our proposal: combine both losses $\longrightarrow \mathcal{L}_{\text{com}} \approx \mathcal{L}_{\max} + \alpha \mathcal{L}_{\min}$.

We set $\alpha = 1/\tau$, chosen on the validation set to balance the two terms; the combined loss $\mathcal{L}_{\text{com}}$ can be unstable to train (sensitive to α).

Numerical efficiency

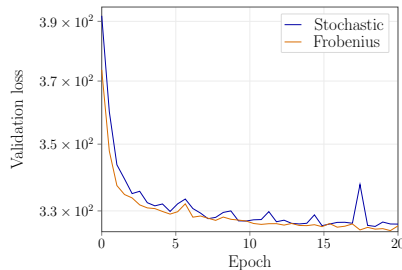
PROBLEM

The Frobenius norm scales computationally poorly.

Unbiased approximation using Hutchinson's trace estimator:

$$\|M\|_F^2 \approx \|Mw\|_2^2, \quad w \sim \mathcal{N}(\mathbf{0}, \mathbf{I}).$$

Intuition: estimate a Frobenius norm from a single matrix-vector product, without ever forming M .



Efficient model training and inference:

- ▶ Training needs only matrix-vector products (or precomputed solves)
- ▶ Linear inference-time complexity in the number of non-zero elements

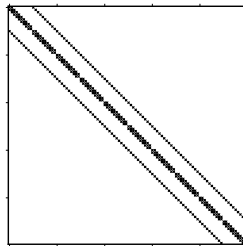
Data

Synthetic dataset inspired by the **Poisson equation**:

- ▶ FEM discretization, $n = 2\,500$ ($\approx 50\,000$ – $150\,000$ non-zeros)
- ▶ Perturb *only* the existing non-zeros with Gaussian noise X_{ij} — the sparsity pattern is preserved:

$$a'_{ij} = \begin{cases} a_{ij} + X_{ij} & \text{if } a_{ij} \neq 0, \\ 0 & \text{if } a_{ij} = 0. \end{cases}$$

Non-symmetry is introduced through coefficient perturbations.



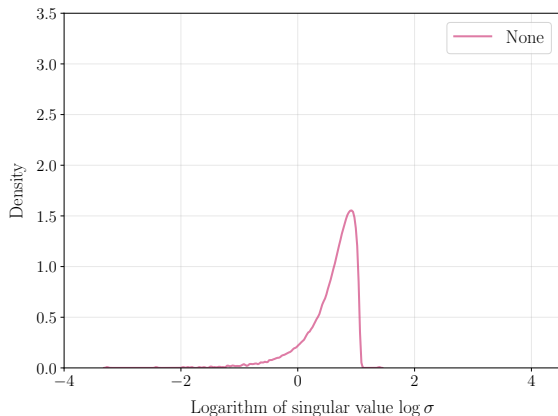
Sparsity pattern across the dataset

Results

Singular value distributions

Small synthetic dataset inspired by Poisson PDE discretization.

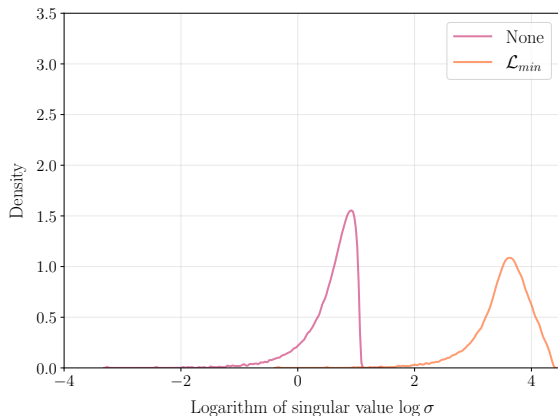
Goal: cluster the singular values and keep them away from zero — better conditioning → fewer iterations



Singular value distributions

Small synthetic dataset inspired by Poisson PDE discretization.

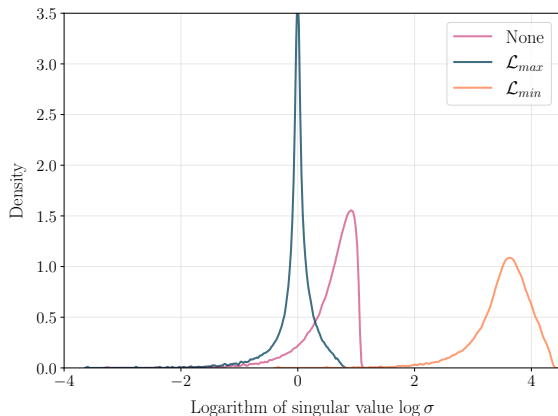
Goal: cluster the singular values and keep them away from zero — better conditioning $\rightarrow$ fewer iterations



Singular value distributions

Small synthetic dataset inspired by Poisson PDE discretization.

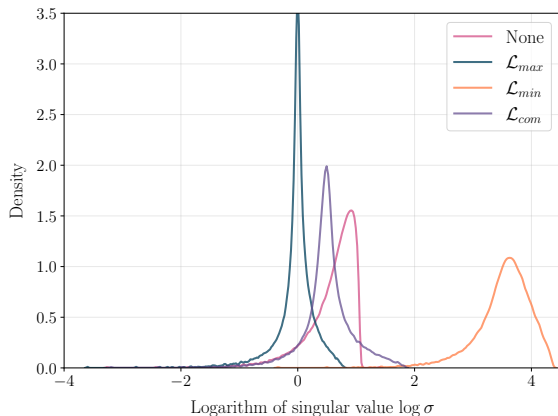
Goal: cluster the singular values and keep them away from zero — better conditioning $\rightarrow$ fewer iterations



Singular value distributions

Small synthetic dataset inspired by Poisson PDE discretization.

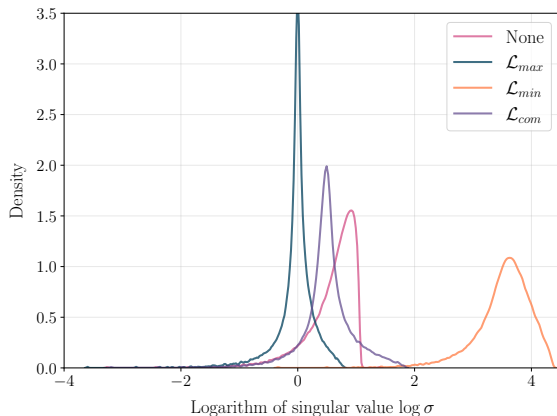
Goal: cluster the singular values and keep them away from zero — better conditioning $\rightarrow$ fewer iterations



Singular value distributions

Small synthetic dataset inspired by Poisson PDE discretization.

Goal: cluster the singular values and keep them away from zero — better conditioning $\rightarrow$ fewer iterations



Learned losses reshape the distribution of singular values in different ways.

GMRES convergence performance

Performance in terms of GMRES iterations and time:

	Method	$\sigma_{\min} \uparrow$	$\sigma_{\max} \downarrow$	$\kappa \downarrow$	$\ P - A\ _F \downarrow$	$\ PA^{-1} - I\ _F \downarrow$	Time $\downarrow$	Iterations $\downarrow$
Precond.	No preconditioner	0.0014	20.70	31 152.94	255.26	1 577.65	30.85	1 153
	Jacobi	0.0003	5.16	31 166.13	205.83	6 319.45	30.12	1 152
	ILU(0)	0.0006	30.32	120 740.90	138.31	3 688.45	3.33	413
	Learned IC	0.0006	7.58	27 405.57	143.23	3 719.47	12.84	692
Loss	$\mathcal{L}_{\max}$	0.0007	5.00	16 139.46	88.76	3 261.23	3.67	437
	$\mathcal{L}_{\min}$	1.1375	20 318.88	37 030.72	287.62	50.05	24.41	1 054
	$\mathcal{L}_{\text{com}}$	0.0017	52.92	66 691.71	197.82	1 240.60	3.42	418

GMRES convergence performance

Performance in terms of GMRES iterations and time:

	Method	$\sigma_{\min} \uparrow$	$\sigma_{\max} \downarrow$	$\kappa \downarrow$	$\ P - A\ _F \downarrow$	$\ PA^{-1} - I\ _F \downarrow$	Time $\downarrow$	Iterations $\downarrow$
Precond.	No preconditioner	0.0014	20.70	31 152.94	255.26	1 577.65	30.85	1 153
	Jacobi	0.0003	5.16	31 166.13	205.83	6 319.45	30.12	1 152
	ILU(0)	0.0006	30.32	120 740.90	138.31	3 688.45	3.33	413
	Learned IC	0.0006	7.58	27 405.57	143.23	3 719.47	12.84	692
Loss	$\mathcal{L}_{\max}$	0.0007	5.00	16 139.46	88.76	3 261.23	3.67	437
	$\mathcal{L}_{\min}$	1.1375	20 318.88	37 030.72	287.62	50.05	24.41	1 054
	$\mathcal{L}_{\text{com}}$	0.0017	52.92	66 691.71	197.82	1 240.60	3.42	418

Our loss functions are working as expected, pushing the singular values in the intended direction.

GMRES convergence performance

Performance in terms of GMRES iterations and time:

	Method	$\sigma_{\min} \uparrow$	$\sigma_{\max} \downarrow$	$\kappa \downarrow$	$\ P - A\ _F \downarrow$	$\ PA^{-1} - I\ _F \downarrow$	Time $\downarrow$	Iterations $\downarrow$
Precond.	No preconditioner	0.0014	20.70	31 152.94	255.26	1 577.65	30.85	1 153
	Jacobi	0.0003	5.16	31 166.13	205.83	6 319.45	30.12	1 152
	ILU(0)	0.0006	30.32	120 740.90	138.31	3 688.45	3.33	413
	Learned IC	0.0006	7.58	27 405.57	143.23	3 719.47	12.84	692
Loss	$\mathcal{L}_{\max}$	0.0007	5.00	16 139.46	88.76	3 261.23	3.67	437
	$\mathcal{L}_{\min}$	1.1375	20 318.88	37 030.72	287.62	50.05	24.41	1 054
	$\mathcal{L}_{\text{com}}$	0.0017	52.92	66 691.71	197.82	1 240.60	3.42	418

The learned preconditioners achieve the best singular values clustering and condition number, while already matching classical methods (ILU(0)) on GMRES iterations and runtime.

Conclusion

Conclusion

- ▶ A **GNN architecture** that outputs a **sparse, non-singular** incomplete LU factorization
- ▶ Studied how different **loss functions** shape the “**spectrum**” of the preconditioned system AP^{-1}
- ▶ Learned preconditioners nearly match GMRES convergence with ILU(0) on related systems

Future work

Future work

- ▶ Take the right-hand side $\mathbf{b}$ into account during training
- ▶ Apply to new problem domains such as Newton's method, IPMs, GPs
- ▶ Learn additional heuristics, such as the sparsity pattern of the preconditioner directly
- ▶ Combine data-driven and classical algorithms⁵

The development of data-driven preconditioners has just begun!

⁵Vladislav Trifonov et al., Learning from linear algebra: a graph neural network approach to preconditioner design for conjugate gradient solvers, arXiv, 2024.

Thank you for your attention!

Aleix Nieto Juscafresa, Uppsala University

E-mail: `aleix.nieto.juscafresa@it.uu.se`



Supported by the Wallenberg AI, Autonomous Systems and Software Program (WASP) funded by Knut and Alice Wallenberg Foundation