

Query Efficient Structured Matrix Learning

Conference on Learning Theory 2026

Pratyush Avi

New York University

Joint work with

Noah Amsel¹, Feyza Duman Keles¹, Tyler Chen¹, Chinmay Hegde¹, Cameron Musco², Christopher Musco¹, David Persson³

¹ New York University ² UMass Amherst ³ Flatiron Institute

Structured Matrix Learning

Problem: Let $\mathbf{A}$ be some target matrix and let $\mathcal{F} \subset \mathbb{R}^{n \times n}$ be a family of $n \times n$ matrices. For approximation factor $\gamma > 1$, find $\tilde{\mathbf{B}} \in \mathcal{F}$ satisfying:

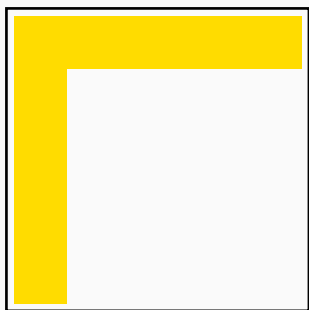
$$\|\mathbf{A} - \tilde{\mathbf{B}}\| \leq \gamma \cdot \min_{\mathbf{B} \in \mathcal{F}} \|\mathbf{A} - \mathbf{B}\|$$

Structured Matrix Learning

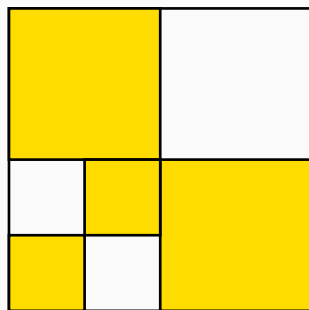
Problem: Let $\mathbf{A}$ be some target matrix and let $\mathcal{F} \subset \mathbb{R}^{n \times n}$ be a family of $n \times n$ matrices. For approximation factor $\gamma > 1$, find $\tilde{\mathbf{B}} \in \mathcal{F}$ satisfying:

$$\|\mathbf{A} - \tilde{\mathbf{B}}\| \leq \gamma \cdot \min_{\mathbf{B} \in \mathcal{F}} \|\mathbf{A} - \mathbf{B}\|$$

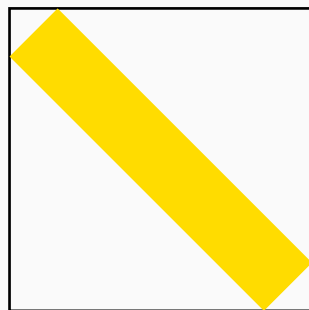
Example families:



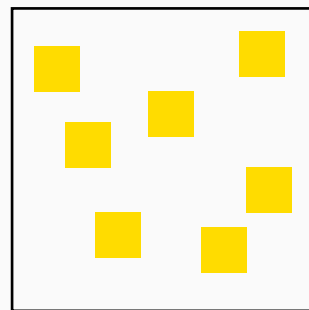
low-rank



hierarchically
low-rank



diagonal



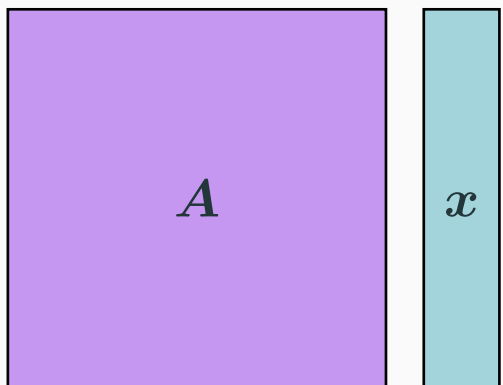
sparse

More e.g.: Toeplitz, butterfly, block-diagonal, tridiagonal, ...

Matrix Learning Access Model

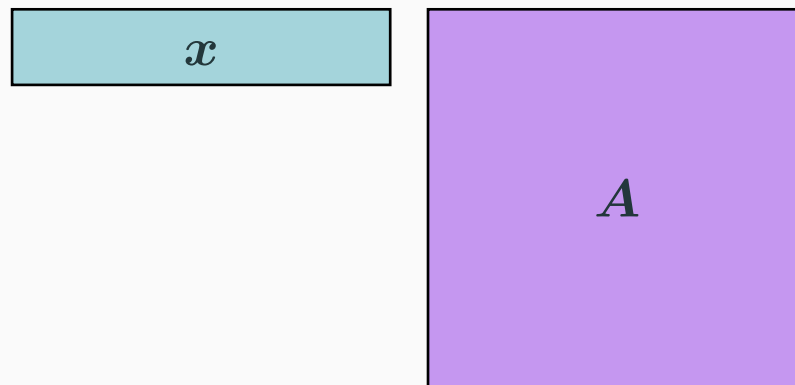
Say we are given access to A via **matrix-vector product** queries.

$$x \mapsto Ax$$



right query

$$x \mapsto x^\top A$$

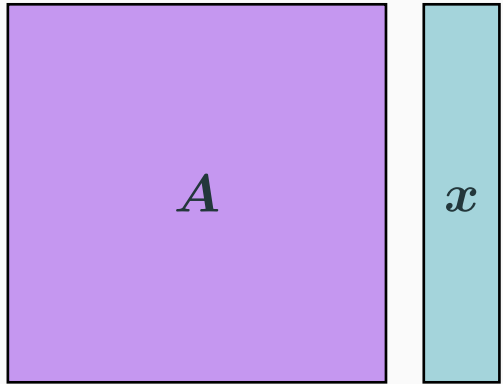


left query

Matrix Learning Access Model

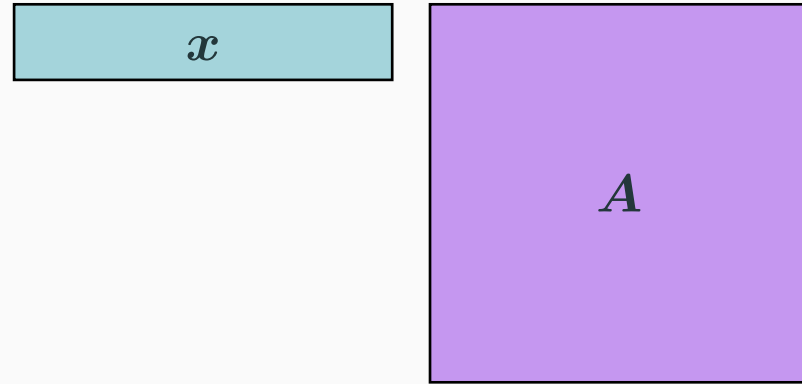
Say we are given access to A via **matrix-vector product** queries.

$$x \mapsto Ax$$



right query

$$x \mapsto x^\top A$$



left query

Question: How many **queries** do we need to find a near-optimal approximation from some hypothesis class $\mathcal{F}$?

Motivation

In many cases it is natural to have matrix-vector product access to $\mathcal{A}$.

Matrix	Oracle for Matrix-Vector Product
Hessian matrix $\mathbf{H}$	Backpropagation to compute $\mathbf{H}\mathbf{x}$
Inverse matrix $\mathbf{B}^{-1}$	Iterative methods to compute $\mathbf{B}^{-1}\mathbf{x}$
Discretization of Linear PDE operator $\mathcal{L}$	PDE solver to compute $\mathcal{L}\mathbf{x}$ $(\mathbf{x}, \mathcal{A}\mathbf{x})$

Baseline: n queries

n matrix-vector product queries are always sufficient. Just multiply with the n standard basis vectors $e_1, \dots, e_n$!

A				
1	0	0	0	0
0	1	0	0	0
0	0	1	0	0
0	0	0	1	0
0	0	0	0	1
e_1	e_2	e_3	e_4	e_5

But we can do better...

Example Hypothesis Class - Diagonal Matrices

Let $\mathcal{F}$ be the class of diagonal matrices.

5	.1	.1	.1	.1
.1	8	.1	.1	.1
.1	.1	-2	.1	.1
.1	.1	.1	3	.1
.1	.1	.1	.1	-4

target matrix $\mathbf{A}$

We want to find $\tilde{\mathbf{B}} \in \mathcal{F}$ satisfying:

$$\|\mathbf{A} - \tilde{\mathbf{B}}\| \leq \gamma \cdot \min_{\mathbf{B} \in \mathcal{F}} \|\mathbf{A} - \mathbf{B}\|$$

Example Hypothesis Class - Diagonal Matrices

Possible procedure [Bekas, Kokiopoulou, Saad 2007]:

- Draw query vector with random ± 1 entries, i.e., $\mathbf{x} \in \{-1, 1\}^n$
- Compute $\mathbf{x} \circ \mathbf{A}\mathbf{x}$

Example Hypothesis Class - Diagonal Matrices

Possible procedure [Bekas, Kokiopoulou, Saad 2007]:

- Draw query vector with random ± 1 entries, i.e., $\mathbf{x} \in \{-1, 1\}^n$
- Compute $\mathbf{x} \circ \mathbf{A}\mathbf{x}$

5	.1	.1	.1	.1
.1	8	.1	.1	.1
.1	.1	-2	.1	.1
.1	.1	.1	3	.1
.1	.1	.1	.1	-4

target matrix $\mathbf{A}$

Example Hypothesis Class - Diagonal Matrices

Possible procedure [Bekas, Kokiopoulou, Saad 2007]:

- Draw query vector with random ± 1 entries, i.e., $\mathbf{x} \in \{-1, 1\}^n$
- Compute $\mathbf{x} \circ \mathbf{A}\mathbf{x}$

5	.1	.1	.1	.1	+1
.1	8	.1	.1	.1	-1
.1	.1	-2	.1	.1	-1
.1	.1	.1	3	.1	+1
.1	.1	.1	.1	-4	+1

target matrix $\mathbf{A}$ $\mathbf{x}$

Example Hypothesis Class - Diagonal Matrices

Possible procedure [Bekas, Kokiopoulou, Saad 2007]:

- Draw query vector with random ± 1 entries, i.e., $\mathbf{x} \in \{-1, 1\}^n$
- Compute $\mathbf{x} \circ \mathbf{A}\mathbf{x}$

5	.1	.1	.1	.1	+1	=	5
.1	8	.1	.1	.1	-1		-7.8
.1	.1	-2	.1	.1	-1		2.2
.1	.1	.1	3	.1	+1		3
.1	.1	.1	.1	-4	+1		-4

target matrix $\mathbf{A}$ $\mathbf{x}$

Example Hypothesis Class - Diagonal Matrices

Possible procedure [Bekas, Kokiopoulou, Saad 2007]:

- Draw query vector with random ± 1 entries, i.e., $\mathbf{x} \in \{-1, 1\}^n$
- Compute $\mathbf{x} \circ \mathbf{A}\mathbf{x}$

5	.1	.1	.1	.1	+1	=	5
.1	8	.1	.1	.1	-1		-7.8
.1	.1	-2	.1	.1	-1		2.2
.1	.1	.1	3	.1	+1		3
.1	.1	.1	.1	-4	+1		-4

target matrix $\mathbf{A}$ $\mathbf{x}$

$$\mathbf{x} \circ \mathbf{A}\mathbf{x} = \sum_{j=1}^n x_j (\mathbf{A}\mathbf{x})_j =$$

Example Hypothesis Class - Diagonal Matrices

Possible procedure [Bekas, Kokiopoulou, Saad 2007]:

- Draw query vector with random ± 1 entries, i.e., $\mathbf{x} \in \{-1, 1\}^n$
- Compute $\mathbf{x} \circ \mathbf{A}\mathbf{x}$

5	.1	.1	.1	.1	+1	=	5
.1	8	.1	.1	.1	-1		-7.8
.1	.1	-2	.1	.1	-1		2.2
.1	.1	.1	3	.1	+1		3
.1	.1	.1	.1	-4	+1		-4

target matrix $\mathbf{A}$ $\mathbf{x}$

$$\mathbf{x} \circ \mathbf{A}\mathbf{x} = \sum_{j=1}^n x_j (\mathbf{A}\mathbf{x})_j =$$

5
7.8
-2.2
3
-4

Example Hypothesis Class - Diagonal Matrices

Possible procedure [Bekas, Kokiopoulou, Saad 2007]:

- Draw query vector with random ± 1 entries, i.e., $\mathbf{x} \in \{-1, 1\}^n$
- Compute $\mathbf{x} \circ \mathbf{Ax}$

5	.1	.1	.1	.1
.1	8	.1	.1	.1
.1	.1	-2	.1	.1
.1	.1	.1	3	.1
.1	.1	.1	.1	-4

$\mathbf{x}$

+1
-1
-1
+1
+1

=

5
-7.8
2.2
3
-4

$\mathbf{x} \circ \mathbf{Ax} = \sum_{j=1}^n x_j (\mathbf{Ax})_j =$

5
7.8
-2.2
3
-4

target matrix $\mathbf{A}$ $\mathbf{x}$

Unbiased estimator: $\mathbb{E}[x_j (\mathbf{Ax})_j] = \sum_k A_{jk} \mathbb{E}[x_j x_k] = A_{jj}$

Example Hypothesis Class - Diagonal Matrices

Possible procedure [Bekas, Kokiopoulou, Saad 2007]:

- Draw query vector with random ± 1 entries, i.e., $\mathbf{x} \in \{-1, 1\}^n$
- Compute $\mathbf{x} \circ \mathbf{Ax}$

5	.1	.1	.1	.1
.1	8	.1	.1	.1
.1	.1	-2	.1	.1
.1	.1	.1	3	.1
.1	.1	.1	.1	-4

+1
-1
-1
+1
+1

 $=$

5
-7.8
2.2
3
-4

$$\mathbf{x} \circ \mathbf{Ax} = \sum_{j=1}^n x_j (\mathbf{Ax})_j =$$

5
7.8
-2.2
3
-4

target matrix $\mathbf{A}$ $\mathbf{x}$

Unbiased estimator: $\mathbb{E}[x_j (\mathbf{Ax})_j] = \sum_k A_{jk} \mathbb{E}[x_j x_k] = A_{jj}$

Repeat $1/\varepsilon$ and average to improve accuracy to get $\gamma = 1 + \varepsilon$ approximation

Existing Results for Popular Classes

Some known results for $\gamma = 1 + \varepsilon$

Structure	Query Complexity	Reference
Diagonal	$O(1/\varepsilon)$	[Bekas et al., 2007] [Dharangutte, Musco 2023]
Rank k	$O(k/\varepsilon^{1/3})$	Randomized SVD, [Bakshi et al., 2022]
s -banded	$O(s/\varepsilon)$	[Amsel et al., 2026]
s -sparse rows	$O(s/\varepsilon)$	[Amsel et al., 2026]
rank- k HODLR	$O(k \log^4 n / \varepsilon^3)$	[Lin, Lu, Ying, 2011] [Chen et al., 2025]
$\vdots$	$\vdots$	$\vdots$

Towards a General Theory of Matrix Learning

Does there exist a general theory to characterize the query complexity of structured matrix approximation?

Something in the spirit of VC dimension, Rademacher complexity, and covering number.

Our Results



Baseline & Setup

We consider:

- Finite hypothesis classes, $\mathcal{F}$
- Perfect matrix-vector product queries

Baseline & Setup

We consider:

- Finite hypothesis classes, $\mathcal{F}$
- Perfect matrix-vector product queries

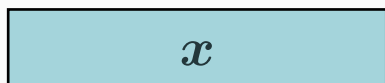
Naive result: $O(\log|\mathcal{F}| / \varepsilon^2)$ queries are sufficient

Imitates standard learning theory results about finite hypothesis classes.

Matrix Learning vs Supervised Learning

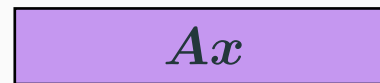
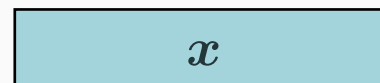
We get significantly more information per query as compared to standard supervised learning with scalar valued outputs.

Supervised learning



scalar output

Matrix learning



n -dim output

e.g. of supervised learning setting: rank-1 matrix sensing

Main Result

The multi-output nature of the problem allows us a **quadratic** improvement in query complexity!

Theorem (Amsel, PA, Duman Keles, Chen, Hegde, Musco, Musco, Persson):

Let $\mathcal{F}$ be a finite matrix family. A near-optimal approximation to $\mathbf{A}$ can be learned with accuracy ϵ with

$$\tilde{O}\left(\sqrt{\log|\mathcal{F}|}\right) \text{ matrix-vector product queries}$$

with high probability.

Main Result

The multi-output nature of the problem allows us a **quadratic** improvement in query complexity!

Theorem (Amsel, PA, Duman Keles, Chen, Hegde, Musco, Musco, Persson):

Let $\mathcal{F}$ be a finite matrix family. A near-optimal approximation to $\mathbf{A}$ can be learned with accuracy 4 with

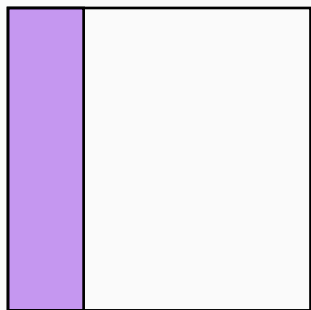
$\tilde{O}\left(\sqrt{\log|\mathcal{F}|}\right)$ matrix-vector product queries

with high probability.

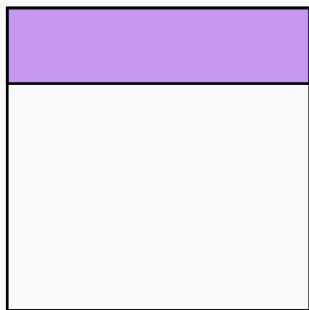
We show that this result is tight up to $\log \log$ factors.

Key Idea

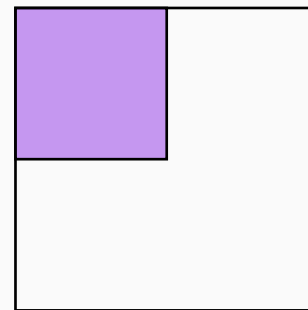
There is some power to using both $\mathbf{Ax}$ and $\mathbf{x}^\top \mathbf{A}$ queries.



Single $\mathbf{Ax}$ query
But n $\mathbf{x}^\top \mathbf{A}$ queries!



n $\mathbf{Ax}$ queries
But single $\mathbf{x}^\top \mathbf{A}$ queries!



Hard case:
Requires $\sqrt{n}$ from both sides

Query Complexity via Main Result

Family	Family Size	Query Complexity
s -sparse	$2^{O(s)}$	$\tilde{O}(\sqrt{s})$
Constant Rank Butterfly	$2^{O(n)}$	$\tilde{O}(\sqrt{n})$
Diagonal	$2^{O(n)}$	$\tilde{O}(\sqrt{n})$
Rank- k	$2^{O(nk)}$	$\tilde{O}(\sqrt{nk})$

Query Complexity via Main Result

Family	Family Size	Query Complexity
s -sparse	$2^{O(s)}$	$\tilde{O}(\sqrt{s})$
Constant Rank Butterfly	$2^{O(n)}$	$\tilde{O}(\sqrt{n})$
Diagonal	$2^{O(n)}$	$\tilde{O}(\sqrt{n})$
Rank- k	$2^{O(nk)}$	$\tilde{O}(\sqrt{nk})$

This bound leads to **tight** results for some families, but **loose** results for others.

Thank you!



Paper link

<https://arxiv.org/abs/2507.19290>

Some open questions:

1. Class size may not be the best complexity measure. Is there something better that can achieve instance optimality?
2. Obtain $(1 + \varepsilon)$ guarantee instead of constant factor
3. Extend this to other norms, e.g.: spectral norm