



Randomized numerical linear algebra Bootcamp, Part III



Daniel Kressner

M  R A^T

Mathematics of Randomized
linear Algebra Techniques

EPFL

Institute of Mathematics

EPFL *ETH zürich*

Subspace Embeddings and Linear Regression

Overdetermined least-squares problems

Consider

$$\min_{x \in \mathbb{R}^k} \|Ax - b\|_2, \quad A \in \mathbb{R}^{n \times k}, \quad b \in \mathbb{R}^n, \quad n \gg k, \quad \text{rank}(A) = k.$$

Classical approach: Compute thin QR factorization $A = QR$, $Q \in \mathbb{R}^{n \times k}$, $Q^\top Q = I$, and solve $Rx = Q^\top b$.



$$Q^\top [A \mid b] = [R \mid Q^\top b]$$

$$\|Ax - b\|_2^2 = \|Rx - Q^\top b\|_2^2 + \|(I - QQ^\top)b\|_2^2.$$

$x = R^{-1}Q^\top b$ annihilates first component; second component is independent of x .

Overdetermined least-squares problems

Idea of **sketching**: Replace $Q^\top \in \mathbb{R}^{k \times n}$ by (cheaper) sketching matrix $S^\top \in \mathbb{R}^{m \times n}$ with $k \leq m \ll n$ and solve sketched problem

$$\min\{\|S^\top A\tilde{x} - S^\top b\|_2 : \tilde{x} \in \mathbb{R}^k\}. \quad (1)$$

A good sketching matrix S should approximately capture the range of A and* b . A common way to quantify this is the subspace embedding property.

Given (low-dimensional) subspace $\mathcal{U} \subset \mathbb{R}^n$, S is called an **ϵ -subspace embedding** for $0 < \epsilon < 1$ if

$$(1 - \epsilon)\|u\|_2^2 \leq \|S^\top u\|_2^2 \leq (1 + \epsilon)\|u\|_2^2 \quad \forall u \in \mathcal{U}.$$

*Capturing b is often not needed, for example in preconditioning approach discussed later.

Overdetermined least-squares problems

Lemma (sketch-and-solve)

Suppose that $\tilde{x}$ solves the sketched least-squares problem (1) with an ϵ -subspace embedding S of the subspace $\text{range}([A, b])$. Then

$$\|A\tilde{x} - b\|_2^2 \leq \frac{1+\epsilon}{1-\epsilon} \|Ax_\star - b\|_2^2,$$

where $x_\star$ minimizes $\|Ax - b\|_2$.

Proof of Lemma.

$$\begin{aligned} \|A\tilde{x} - b\|_2^2 &\leq \frac{1}{1-\epsilon} \|S^\top A\tilde{x} - S^\top b\|_2^2 \stackrel{\text{opt. of } \tilde{x}}{\leq} \frac{1}{1-\epsilon} \|S^\top Ax_\star - S^\top b\|_2^2 \\ &\leq \frac{1+\epsilon}{1-\epsilon} \|Ax_\star - b\|_2^2, \end{aligned}$$

where we used that $Ax_\star - b, A\tilde{x} - b \in \text{range}([A, b])$.

◇

Subspace embeddings and injections

A trivial subspace embedding (with $\epsilon = 0$) is to take $S = \tilde{Q}$, where $\tilde{Q}$ is orthonormal basis of $\text{range}([A, b])$. This is *not* what the lemma is aiming for. We aim for (much) cheaper constructions that (preferably) use no information about A, b .

The roles of ϵ in the two sides of the ϵ -subspace embedding property are unbalanced. While $\epsilon < 1$ is absolutely essential for lhs, could comfortably afford larger ϵ on rhs.

Motivates the following definition: Given a subspace $\mathcal{U} \subset \mathbb{R}^k$, we call $S \in \mathbb{R}^{n \times m}$ an ϵ -subspace injection for $\mathcal{U}$ if

$$(1 - \epsilon)\|u\|_2^2 \leq \|S^\top u\|_2^2 \quad \forall u \in \mathcal{U}.$$

Unlike subspace embedding, an injection only prevents vectors from being nearly annihilated.

Sketch-and-solve with subspace injections

Lemma (sketch-and-solve with injection)

Let $x_\star$ minimize $\|Ax - b\|_2$ and suppose that $\tilde{x}$ solves the sketched least-squares problem (1), where S is an ϵ -subspace injection of $\text{range}(A)$ and $\|S^\top r_\star\|_2^2 \leq \gamma \|r_\star\|_2^2$ for $r_\star := Ax_\star - b$. Then

$$\|A\tilde{x} - b\|_2^2 \leq \left(1 + \frac{\gamma}{1 - \epsilon}\right) \|Ax_\star - b\|_2^2.$$

Proof. Set $d = \tilde{x} - x_\star$. Using that the sketched residual $S^\top(A\tilde{x} - b)$ is orthogonal to the range of $S^\top A$, Pythagoras implies

$$\|S^\top r_\star\|_2^2 = \|S^\top(A\tilde{x} - b)\|_2^2 + \|S^\top Ad\|_2^2 \geq \|S^\top Ad\|_2^2.$$

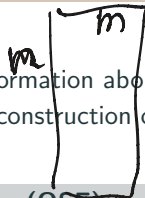
Thus,

$$(1 - \epsilon)\|Ad\|_2^2 \leq \|S^\top Ad\|_2^2 \leq \|S^\top r_\star\|_2^2 \leq \gamma \|r_\star\|_2^2.$$

This implies the result because $\|A\tilde{x} - b\|_2^2 = \|r_\star\|_2^2 + \|Ad\|_2^2$. $\diamond$

Oblivious Subspace Embeddings

Subspace embedding is called *oblivious* if no information about the subspace beyond its dimension k is used in the construction of the embedding.



Definition of Oblivious Subspace Embedding (OSE)

A random matrix $\Omega \in \mathbb{R}^{n \times m}$ is called a (k, ϵ, δ) -OSE if the following holds: For *fixed but arbitrary* k -dimensional subspace $\mathcal{U} \subset \mathbb{R}^n$,

$$(1 - \epsilon) \|u\|_2^2 \leq \|\Omega^\top u\|_2^2 \leq (1 + \epsilon) \|u\|_2^2 \quad \forall u \in \mathcal{U} \quad (2)$$

holds[†] with probability at least $1 - \delta$.

One hopes for $m \ll n$. Question for reflection: Why is Ω necessarily random? What is this emphasis on *fixed but arbitrary* all about?

[†]Some define this inequality without the squares. Only makes a marginal difference.

$$u = Uz$$

$$\| \Sigma^T v \|_2^2 = v^T \Sigma \Sigma^T u$$

$$\frac{z^T U \Sigma \Sigma^T U z}{z^T z}$$

$$u \in \mathbb{R}^n$$

Rayleigh quotient

$$[1 - \varepsilon, 1 + \varepsilon]$$

↑ eigen values of $U^T \Sigma \Sigma^T U$

$A^T A$ has eigen values in
 $[1 - \epsilon, 1 + \epsilon]$

A has singular values
 $[\sqrt{1 - \epsilon}, \sqrt{1 + \epsilon}]$

OSE: Important characterizations

Let $U \in \mathbb{R}^{n \times k}$ be ONB for $\mathcal{U}$. Each of the following is equivalent to the defining property (2) of OSE:

1. $\max\{|\|\Omega^\top u\|_2^2 - 1| : u \in \mathcal{U}, \|u\|_2 = 1\} \leq \epsilon$
2. $\sigma_{\min}(\Omega^\top U)^2 \geq 1 - \epsilon, \sigma_{\max}(\Omega^\top U)^2 \leq 1 + \epsilon$
3. $\lambda_{\min}(U^\top \Omega \Omega^\top U) \geq 1 - \epsilon, \lambda_{\max}(U^\top \Omega \Omega^\top U) \leq 1 + \epsilon$
4. $\|I - U^\top \Omega \Omega^\top U\|_2 \leq \epsilon$

Subspace injection is called *oblivious* if it works for arbitrary subspace. In contrast to OSE, we replace the uniform upper bound in the embedding property by norm preservation in expectation (isotropy).

Definition of Oblivious Subspace Injection (OSI)

A random matrix $\Omega \in \mathbb{R}^{n \times m}$ is called a (k, ϵ, δ) -OSI if the following hold:

1. $\mathbb{E}[\|\Omega^\top z\|_2^2] = \|z\|_2^2$ for every $z \in \mathbb{R}^n$ (isotropy);
2. For *fixed but arbitrary* k -dimensional subspace $\mathcal{U} \subset \mathbb{R}^n$,

$$(1 - \epsilon)\|u\|_2^2 \leq \|\Omega^\top u\|_2^2 \quad \forall u \in \mathcal{U}$$

holds with probability at least $1 - \delta$.

Task for reflection: Adapt the characterizations of OSE to the injection property of OSI.

OSE and a second encounter with JL

For $k = 1$, OSE reduces to JL.

Johnson-Lindenstrauss (JL) property

A random matrix Ω satisfies (ϵ, δ) -JL property if for *fixed but arbitrary* vector u the inequalities

$$(1 - \epsilon)\|u\|_2^2 \leq \|\Omega^\top u\|_2^2 \leq (1 + \epsilon)\|u\|_2^2 \quad (3)$$

hold with probability at least $1 - \delta$.

We have already seen that $\Omega = \Omega_g / \sqrt{m}$ for $n \times m$ Gaussian random matrix Ω_g satisfies (ϵ, δ) -JL property if

$$m \geq 8\epsilon^{-2} \log(2\delta^{-1}).$$

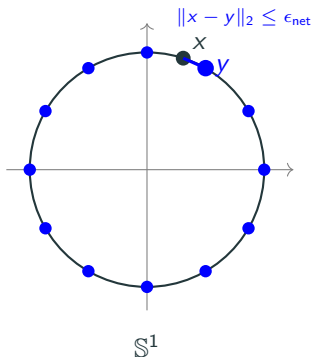
Bad dependence on ϵ ; great dependence on δ and n !

When extending JL to OSE for k -dimensional subspace $\mathcal{U} \subset \mathbb{R}^n$, union bound suffers from the obvious problem that $\mathcal{U}$ contains infinitely many vectors. Popular techniques to overcome such problems: [epsilon nets](#), [chaining](#), and [approximate matrix multiplication](#).

Idea of ϵ -nets: Given ONB $U \in \mathbb{R}^{n \times k}$, every (normalized) vector in $\mathcal{U}$ takes the form Ux , $x \in S^{k-1}$ (unit sphere in $\mathbb{R}^k$). Cover the sphere with vectors up to a distance $\epsilon = \mathcal{O}(1)$ and use union bound. MANY vectors will be needed, but $O(|\log \delta|)$ dependence of JL saves us!

From JL to OSE: subspace dimension $k = 2$

Let $U \in \mathbb{R}^{N \times 2}$ be ONB for two-dimensional subspace $\mathcal{U}$. Every unit vector in $\mathcal{U}$ has the form Ux for $x \in \mathbb{S}^1$.



Place M equally spaced points

$$x_j = \begin{bmatrix} \cos(2\pi j/M) \\ \sin(2\pi j/M) \end{bmatrix}, \quad j = 0, \dots, M-1.$$

Every $x \in \mathbb{S}^1$ has a grid point y with angular distance at most π/M . Hence

$$\|x - y\|_2 \leq \pi/M.$$

Thus, choosing

$$M \geq \frac{\pi}{\epsilon_{\text{net}}}$$

gives an ϵ_{net} -net of $\mathbb{S}^1$.

From JL to OSE: general k ?

For k -dimensional subspace with ONB $U \in \mathbb{R}^{n \times k}$, we need to control Ux for $x \in \mathbb{S}^{k-1}$.

A crude construction: Put a Cartesian grid of spacing

$$h = \frac{\epsilon_{\text{net}}}{\sqrt{k}}$$

on the cube $[-1, 1]^k$.

For every $x \in \mathbb{S}^{k-1}$, rounding every coordinate to the nearest grid point gives z with

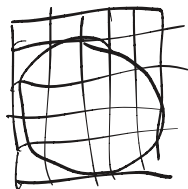
$$\|x - z\|_2 \leq \sqrt{k} \frac{h}{2} = \frac{\epsilon_{\text{net}}}{2}.$$

After normalizing z to $y = z/\|z\|_2 \in \mathbb{S}^{k-1}$,

$$\|x - y\|_2 \leq \epsilon_{\text{net}}.$$

The number of grid points is at most

$$\left(1 + \frac{2\sqrt{k}}{\epsilon_{\text{net}}}\right)^k.$$



$$\left(1 + \frac{2\sqrt{k}}{\epsilon_{\text{net}}}\right)^k$$

Lemma 5.2 in Vershynin'2012

S^{k-1} has an epsilon net $\mathcal{N}_{\epsilon_{\text{net}}} \subset S^{k-1}$ of cardinality at most $(1 + 2/\epsilon_{\text{net}})^k$, that is, for every $x \in S^{k-1}$ there is $y \in \mathcal{N}_{\epsilon_{\text{net}}}$ such that

$$\|x - y\|_2 \leq \epsilon_{\text{net}}.$$

Think of ϵ_{net} not too small, like $\epsilon_{\text{net}} = 1/2$ or $\epsilon_{\text{net}} = 1/4$.

For symmetric $k \times k$ matrix C : $\|C\|_2 = \max\{|x^\top Cx| : x \in S^{k-1}\}$.

Corollary (capturing spectral norm with quadratic forms)

Let $\mathcal{N}_{\epsilon_{\text{net}}}$ be epsilon net from Lemma and assume $\epsilon_{\text{net}} < 1/\sqrt{2}$. Then

$$\|C\|_2 \leq (1 - 2\epsilon_{\text{net}}^2)^{-1} \max\{|y^\top Cy| : y \in \mathcal{N}_{\epsilon_{\text{net}}}\}.$$

Proof. Let $x \in S^{k-1}$ s.t. $Cx = \lambda_1 x$ and $\|C\|_2 = |\lambda_1|$. Let $y \in \mathcal{N}_{\epsilon_{\text{net}}}$ s.t. $\|x - y\|_2 \leq \epsilon_{\text{net}}$. Then

$$|y^\top Cy - x^\top Cx| = |(x-y)^\top (C - \lambda_1 I)(x-y)| \leq 2\|C\|_2 \|x-y\|_2^2 \leq 2\epsilon_{\text{net}}^2 \|C\|_2.$$

This implies $|y^\top Cy| \geq |x^\top Cx| - |x^\top Cx - y^\top Cy| \geq (1 - 2\epsilon_{\text{net}}^2) \|C\|_2. \quad \diamond$

From JL to OSE

Theorem: $\text{JL} \rightarrow \text{OSE}$

Any random matrix Ω satisfying the $(\epsilon/2, \delta/5^k)$ -JL property also satisfies the (k, ϵ, δ) -OSE property.

Proof. By JL we have

$$|y^\top (I - U^\top \Omega \Omega^\top U)y| \leq \epsilon/2 \quad \text{with probability} \geq 1 - \delta/5^k.$$

for arbitrary $y \in S^{k-1}$. By the union bound this shows JL holds for all vectors in $\mathcal{N}_{1/2}$ with prob $\geq 1 - \delta$. By the corollary,

$$\|I - U^\top \Omega \Omega^\top U\|_2 \leq 2 \max\{|y^\top (I - U^\top \Omega \Omega^\top U)y| : y \in \mathcal{N}_{1/2}\} \leq \epsilon,$$

which shows OSE. $\diamond$

By JL for Gaussian random matrices, the theorem establishes OSE if

$$m \geq 32\epsilon^{-2}(k \log 5 + \log 2/\delta) = \mathcal{O}(\epsilon^{-2}(k + \log \delta^{-1})).$$

This general construction gets the asymptotics right, but the constants are somewhat too large.

Tighter bounds on Gaussian embeddings

Let Ω be an $n \times m$ Gaussian random matrix scaled by $1/\sqrt{m}$. Then, by invariance of Gaussian random vectors under rotations, $G = \Omega^\top U$ is $m \times k$ Gaussian random matrix scaled by $1/\sqrt{m}$.

Section 7.3 of [Vershynin'2018] shows

$$\mathbb{E}\|G\|_2 \leq 1 + \sqrt{k/m}, \quad \mathbb{E}\sigma_{\min}(G) \geq 1 - \sqrt{k/m}.$$

[Davidson–Szarek'2001; Vershynin'2018]: For every $t \geq 0$,

$$\Pr \left\{ 1 - \sqrt{\frac{k}{m}} - t \leq \sigma_{\min}(G) \leq \sigma_{\max}(G) \leq 1 + \sqrt{\frac{k}{m}} + t \right\} \geq 1 - 2e^{-mt^2/2}.$$

After some calculations, this shows that

$$m \geq \left(\frac{2 + \epsilon/2}{\epsilon} \right)^2 \left(\sqrt{k} + \sqrt{2 \log(2/\delta)} \right)^2 = \mathcal{O}(\epsilon^{-2}(k + \log \delta^{-1}))$$

ensures (k, ϵ, δ) -OSE.

Marchenko–Pastur and Gaussian subspace embeddings

Let

$$G \in \mathbb{R}^{m \times k}, \quad G_{ij} \stackrel{\text{iid}}{\sim} \mathcal{N}(0, 1/m), \quad \rho := \frac{k}{m} \in (0, 1).$$

The Gram matrix $W = G^\top G \in \mathbb{R}^{k \times k}$ is a (scaled) **Wishart / sample-covariance matrix**.

Marchenko–Pastur law

Suppose $k, m \rightarrow \infty$ with $k/m \rightarrow \rho \in (0, 1)$. Then the empirical eigenvalue distribution of W converges almost surely to the density

$$f_\rho(\lambda) = \frac{\sqrt{(b - \lambda)(\lambda - a)}}{2\pi\rho\lambda}, \quad a \leq \lambda \leq b,$$

where

$$a = (1 - \sqrt{\rho})^2, \quad b = (1 + \sqrt{\rho})^2.$$

Moreover, the extreme eigenvalues converge to the edges:

$$\lambda_{\min}(G^\top G) \rightarrow (1 - \sqrt{\rho})^2, \quad \lambda_{\max}(G^\top G) \rightarrow (1 + \sqrt{\rho})^2.$$

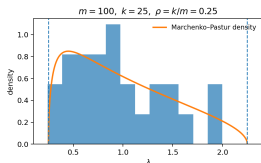
Marchenko–Pastur and Gaussian subspace embeddings

Fix

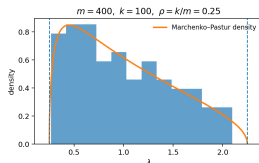
$$\rho = \frac{k}{m} = \frac{1}{4}.$$

Then the limiting spectral interval is

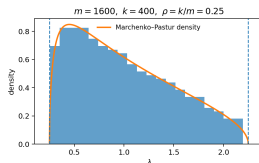
$$[(1 - \sqrt{\rho})^2, (1 + \sqrt{\rho})^2] = [0.25, 2.25].$$



$m = 100, \quad k = 25$



$m = 400, \quad k = 100$



$m = 1600, \quad k = 400$

Marchenko–Pastur eigenvalue distribution gets approached for a single realization as k, m grow.

Marchenko–Pastur and Gaussian subspace embeddings

Revisit $G := \Omega^\top U \in \mathbb{R}^{m \times k}$, Gaussian matrix scaled by $1/\sqrt{m}$. When $\rho = k/m$, Marchenko–Pastur predicts

$$\text{spec}(U^\top \Omega \Omega^\top U) \approx [(1 - \sqrt{\rho})^2, (1 + \sqrt{\rho})^2].$$

For $\rho \ll 1$, $(1 \pm \sqrt{\rho})^2 = 1 \pm 2\sqrt{\rho} + \rho$. Hence, the typical distortion is

$$\epsilon \approx 2\sqrt{\frac{k}{m}},$$

or equivalently

$$m \approx 4\epsilon^{-2}k.$$

The nonasymptotic Gaussian OSE bounds seen before are high-probability, non-asymptotic versions of this result.

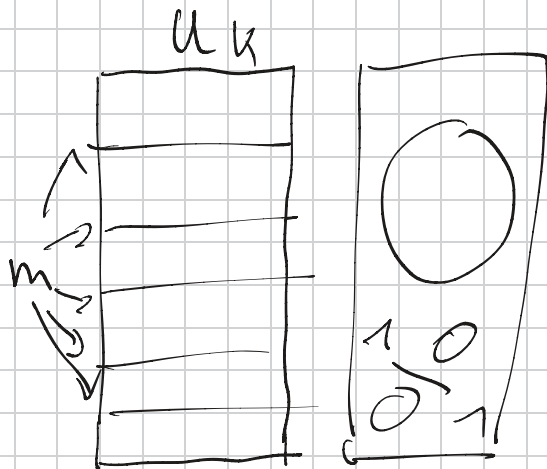
OSE beyond Gaussians

Results on Gaussians essentially extend to matrices with sub-Gaussian iid entries (e.g., Rademacher).

Sketching arbitrary $n \times p$ matrix A with $m \times n$ (sub-)Gaussian matrix $\Omega^\top$ requires $\mathcal{O}(mnp) = \mathcal{O}(knp)$ operations. Can be reduced by imposing structure on Ω .

Popular choices for structured random embeddings:

- Coordinate sampling
- Subsampled unitary transforms
- Sparse sign matrices (including CountSketch)
- Khatri-Rao products, tensor structured sketches, ...



Coordinate sampling

A particularly cheap choice:

$$\Omega = \left[\omega_1 \mid \omega_2 \mid \cdots \mid \omega_m \right], \quad \omega_i \text{ are iid with } \mathbb{P}\{\omega_i = e_j / \sqrt{p_j m}\} = p_j,$$

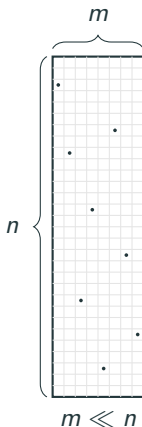
for unit vectors $e_1, \dots, e_n \in \mathbb{R}^n$ and prescribed sampling probabilities $p_1, \dots, p_n$. Scaling ensures $\mathbb{E}[\Omega\Omega^\top] = I_n$.

Choice of $p_1, \dots, p_n$ problematic:

- Uniform sampling: Estimator will perform poorly if there are a few rows $u_j^\top$ of U that dominate the rest. Measured by coherence statistics:

$$\mu(U) := n \cdot \max_j \|u_j\|_2^2, \quad k \leq \mu(U) \leq n.$$

- Leverage score sampling makes embedding non-oblivious!



Subsampled trigonometric transforms

Idea: First apply (random) orthogonal transformation to assure incoherence. Then use uniform sampling.

Popular approach:

$$\Omega^\top = R^\top F D,$$

where

- R is $n \times m$ coordinate sampling matrix with uniform sampling probabilities;
- F is deterministic $n \times n$ orthogonal/unitary matrix (fast transform);
- D diagonal with iid Rademacher diagonal entries.

Results by [Avron et al. 2010], [Tropp'2011]:

Small $\nu = n \max_{ij} |f_{ij}|^2 \rightsquigarrow$ small coherence $\mu \lesssim k\nu$ whp.

Subsampled trigonometric transforms

Most popular choices for F :

- **SRFT** (Subsampled Randomized Fourier Transform): F is the discrete Fourier transform, which achieves minimal value $\eta = 1$. Applying subsampled Fourier transform $R^\top F$ to a vector can be carried out in $\mathcal{O}(n \log m)$ ops. With $m \sim (k + \log n) \log k$, need

$$\mathcal{O}(n(\log(k + \log n) + \log \log k)) \approx \mathcal{O}(n \log k)$$

ops to apply RFD to a vector.

- Subsampled Randomized Hartley Transform / Subsampled Randomized Cosine Transform = real variants of SRFT.
- **SRHT** (Subsampled randomized Hadamard transform)

$$\Omega^\top = R^\top H D, \text{ where } H = \frac{1}{\sqrt{n}} \begin{bmatrix} 1 & 1 \\ 1 & -1 \end{bmatrix} \otimes \begin{bmatrix} 1 & 1 \\ 1 & -1 \end{bmatrix} \otimes \cdots \otimes \begin{bmatrix} 1 & 1 \\ 1 & -1 \end{bmatrix}$$

(zero padding if n is not a power of 2)

For all these constructions:

OSE holds for $m = \mathcal{O}(\epsilon^{-2}(k + \log(n/\delta)) \log(k/\delta))$.

Sparse sign matrices

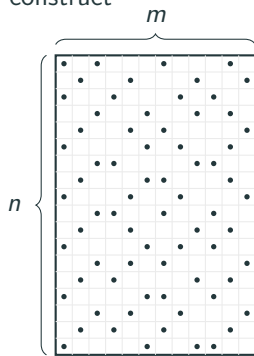
SparseUniform / sparse sign construction: Fix row sparsity $\zeta \ll m$. For every $i = 1, \dots, n$, choose uniformly a subset $J_i \subset \{1, \dots, m\}$ with $|J_i| = \zeta$ and iid Rademacher variables $\{\varepsilon_{ij} : j \in J_i\}$ to construct $\Omega \in \mathbb{R}^{n \times m}$ with [Kane–Nelson'2014]:

$$\Omega_{ij} = \begin{cases} \varepsilon_{ij} / \sqrt{\zeta}, & j \in J_i, \\ 0, & j \notin J_i. \end{cases}$$

Then $\mathbb{E}[\Omega\Omega^\top] = I_n$, $\mathbb{E}\|\Omega^\top x\|_2^2 = \|x\|_2^2$.

$\zeta = 1 \iff \text{CountSketch.}$

Applying sparse sign sketch to $A \in \mathbb{R}^{n \times p}$ requires $\mathcal{O}(\zeta np)$ ops.



$\zeta = 4$ nonzeros/row

Practical rule of thumb: $\boxed{\zeta = 4}$ often sufficient.

Sparse sign matrices: Theory

SparseStack: Write $m = \zeta b$, split m columns of Ω in ζ blocks, and place *one* random signed nonzero in each block. Thus every row of Ω still has exactly ζ nonzeros.

OSI

[Tropp'2026]: For fixed k -dimensional subspace, SparseStack is (k, ϵ) -OSI for fixed failure prob δ provided

$$\zeta \sim \epsilon^{-2} \log(k), \quad m \sim \epsilon^{-2} k$$

Nelson–Nguyen conjecture states for sparse sign matrices that

$$m = \mathcal{O}(\epsilon^{-2}(k + \log(1/\delta))), \quad \zeta = \mathcal{O}(\epsilon^{-1} \log(k/\delta)).$$

implies (k, ϵ, δ) -OSE.

Important caveat: constant ζ with $m = \mathcal{O}(k)$ cannot give OSI. Thus $\zeta = 4$ is only a *practical* recommendation [Huang–Rudelson–Tikhomirov'2026]

Sketch-and-precondition

Overdetermined least-squares revisited

$$\min\{\|Ax - b\|_2 : x \in \mathbb{R}^k\}, \quad A \in \mathbb{R}^{n \times k}, \quad b \in \mathbb{R}^n.$$

Recall result of Lemma (sketch-and-solve):

$$\|A\tilde{x} - b\|_2^2 \leq \frac{1 + \epsilon}{1 - \epsilon} \|Ax_\star - b\|_2^2.$$

Two issues:

- Quality of sketch directly impacts quality of the LSQ solution.
- VERY expensive to get high relative accuracy! ($m \sim \epsilon^{-2}$)

Idea: Use sketching as preconditioner in iterative solver $\rightsquigarrow$
BLENDENPIK [Avron/Maymounkov/Toledo'2010].

Iterative solvers for LSQ problems

Consider $k \times k$ SPD linear system $Cx = d$. The method of conjugate gradients (CG) only requires j matrix-vector products with C and $\mathcal{O}(k)$ extra storage to produce approximation x_j satisfying

$$\|x - x_j\|_C \leq 2\|x - x_0\|_C \left(\frac{\sqrt{\kappa(C)} - 1}{\sqrt{\kappa(C)} + 1} \right)^j,$$

with condition number

$$\kappa(C) = \frac{\sigma_{\max}(C)}{\sigma_{\min}(C)} = \frac{\lambda_{\max}(C)}{\lambda_{\min}(C)},$$

see, e.g., [Golub/Van Loan'2013].

LSQR [Paige/Saunders] for solving LSQ problem is *mathematically* equivalent to applying CG to normal equations $A^\top Ax = A^\top b$.

Convergence:

$$\|b - Ax_j\|_2 \leq 2\|A(x_0 - x_\star)\|_2 \left(\frac{\kappa(A) - 1}{\kappa(A) + 1} \right)^j,$$

Iterative solvers for LSQ problems

Use sketching as preconditioner: Compute QR decomposition

$$S^T A = \hat{Q} \hat{R}$$

and precondition least-squares problem

$$\min \|Ax - b\|_2 = \min \underbrace{\|A\hat{R}^{-1}\|}_{:=\tilde{x}} \|\hat{R}x - b\|_2$$

Let S be ϵ -embedding of $\text{range}(A)$ and consider QR decomposition $A = QR$. Then

$$\kappa(A\hat{R}^{-1}) = \kappa(QR\hat{R}^{-1}) = \kappa(R\hat{R}^{-1}) = \kappa(\hat{Q}\hat{R}R^{-1}) = \kappa(S^T Q) \leq \sqrt{\frac{1+\epsilon}{1-\epsilon}}.$$

Choose $\epsilon = 1/2$. Then LSQR applied to $\min \|A\hat{R}^{-1}\tilde{x} - b\|_2$ converges at rate ≈ 0.27 .

To attain accuracy $\varepsilon \rightsquigarrow \sim |\log \varepsilon|$ iterations needed. Complexity depends on $|\log \varepsilon|$ instead of ε^{-2} ! More detailed discussion of complexity in Sec 8.3 of [Tropp'2020].

Randomized low-rank approximation

HMT'2011 Halko/Martinsson/Tropp. Finding Structure with Randomness: Probabilistic algorithms for constructing approximate matrix decompositions. SIAM Review.

Nakatsukasa'2020 Fast and stable randomized low-rank matrix approximation. arXiv.

Tropp/Webber'2023 Randomized algorithms for low-rank matrix approximation: Design, analysis, and applications. arXiv.

Low-Rank Approximation

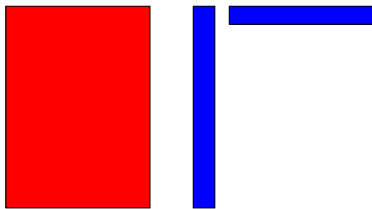
Low-rank factorization

Lemma: Rank- r factorization

A matrix[‡] $A \in \mathbb{R}^{n \times n}$ of rank k admits a factorization of the form

$$A = BC^\top, \quad B \in \mathbb{R}^{n \times k}, \quad C \in \mathbb{R}^{n \times k}.$$

We say that A has **low rank** if $\text{rank}(A) \ll n$.



	A	$BC^\top$
#entries	n^2	$2nk$

Generically, A has **full rank**, so aim instead at **approx.** A by low rank.

[‡]Nothing dramatic happens when moving to rectangular matrices.

Low-rank approximation

Task: Perform low-rank approximation of $n \times n$ matrix A .

Why? How?

The singular value decomposition (SVD)

Theorem: SVD

Let $A \in \mathbb{R}^{n \times n}$. Then there are orthogonal matrices $U, V \in \mathbb{R}^{n \times n}$ such that

$$A = U \Sigma V^T, \quad \text{with} \quad \Sigma = \begin{bmatrix} \sigma_1 & & \\ & \ddots & \\ & & \sigma_n \end{bmatrix}$$

and $\sigma_1 \geq \sigma_2 \geq \dots \geq \sigma_n \geq 0$.

- $\sigma_1, \dots, \sigma_n$ are called singular values
- $u_1, \dots, u_n$ are called *left* singular vectors
- $v_1, \dots, v_n$ are called *right* singular vectors
- Singular values are always uniquely defined by A .
- Singular vectors are *never* unique. If $\sigma_1 > \sigma_2 > \dots > \sigma_n > 0$ then unique up to $u_i \leftarrow \pm u_i, v_i \leftarrow \pm v_i$.

The singular value decomposition (SVD)

Rewriting the SVD:

$$A = U_k \Sigma_k V_k^\top + U_\perp \Sigma_\perp V_\perp^\top$$

$$U_k = \begin{bmatrix} u_1 & \cdots & u_k \end{bmatrix}, \quad \Sigma_k = \begin{bmatrix} \sigma_1 & & \\ & \ddots & \\ & & \sigma_k \end{bmatrix}, \quad V_k = \begin{bmatrix} v_1 & \cdots & v_k \end{bmatrix}$$

Theorem: Eckart-Young-Mirsky

Let $A \in \mathbb{R}^{n \times n}$ and set $\mathcal{T}_k(A) := U_k \Sigma_k V_k^\top$

$$\|A - \mathcal{T}_k(A)\| = \min \{ \|A - B\| : B \in \mathbb{R}^{n \times n} \text{ has rank at most } k \}$$

holds for any unitarily invariant norm $\|\cdot\|$.

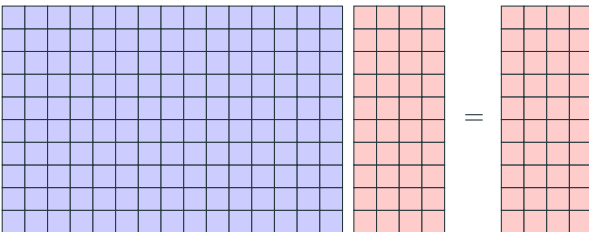
Best rank- k truncation error will be gold standard. In particular:

$$\|A - \mathcal{T}_k(A)\|_F = \sqrt{\sigma_{k+1}^2 + \sigma_{k+2}^2 + \cdots}.$$

Randomized SVD

Two basic ideas

1. If A has exact rank k then $\text{range}(A\Omega) = \text{range}(A)$ holds for generic $\Omega \in \mathbb{R}^{n \times k}$.

$$Y = A\Omega =$$


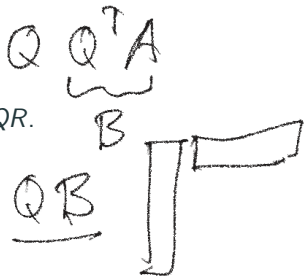
The diagram illustrates the matrix multiplication $Y = A\Omega$. On the left, a large blue grid represents matrix A (10 rows by 10 columns). To its right is a smaller red grid representing matrix Ω (10 rows by 4 columns). An equals sign follows, and to the right is a red grid representing the product matrix Y (10 rows by 4 columns).

2. If Q is an orthonormal basis of range of A then

$$A = QQ^{\top}A = Q(A^{\top}Q)^{\top}$$

Basic randomized algorithm for low-rank approx

1. Draw Gaussian random matrix $\Omega \in \mathbb{R}^{n \times k}$.
2. Perform block mat-vec $Y = A\Omega$.
3. Compute (economic) QR decomposition $Y = QR$.
4. Form $B = Q^T A$.
5. Return $\hat{A} = QB$ (in factorized form)

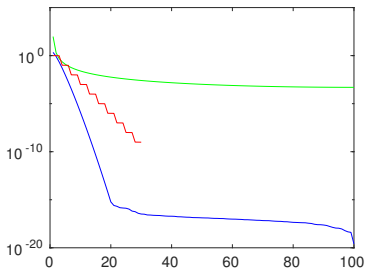


Exact recovery: If A has rank k , we recover $\hat{A} = A$ with probability 1.

Three test matrices

- (a) The 100×100 Hilbert matrix A defined by $A(i, j) = 1/(i + j - 1)$.
- (b) The matrix A defined by $A(i, j) = \exp(-\gamma|i - j|/n)$ with $\gamma = 0.1$.
- (c) 30×30 diagonal matrix with diagonal entries

$$1, 0.99, 0.98, \frac{1}{10}, \frac{0.99}{10}, \frac{0.98}{10}, \frac{1}{100}, \frac{0.99}{100}, \frac{0.98}{100}, \dots$$



Singular values of test matrices

Randomized algorithm applied to test matrices

errors measured in spectral norm:

(a) Hilbert matrix, $k = 5$:

Optimal	mean	std
0.0019	0.0092	0.0099

(b) Matrix with slower decay, $k = 25$:

Optimal	mean	std
0.0034	0.012	0.002

(c) Matrix with staircase sv, $k = 7$:

Optimal	mean	std
0.010	0.038	0.025

Randomized algorithm applied to test matrices

errors measured in Frobenius norm:

(a) Hilbert matrix, $k = 5$:

Optimal	mean	std
0.0019	0.0093	0.0099

(b) Matrix with slower decay, $k = 25$:

Optimal	mean	std
0.011	0.024	0.001

(c) Matrix with staircase sv, $k = 7$:

Optimal	mean	std
0.014	0.041	0.024

Basic randomized algorithms for low-rank approx

Add oversampling. (usually small) integer p

Randomized Algorithm:

1. Draw standard Gaussian random matrix $\Omega \in \mathbb{R}^{n \times (k+p)}$.
2. Perform block mat-vec $Y = A\Omega$.
3. Compute (economic) QR decomposition $Y = QR$.
4. Form $B = Q^\top A$.
5. Return $\hat{A} = QB$ (in factorized form)

Problem: $\hat{A}$ (usually) has rank $k + p > k$.

Solution: Compress $B \approx \mathcal{T}_k(B) \rightsquigarrow Q\mathcal{T}_k(B)$ has rank k .

Error:

$$\begin{aligned}\|Q\mathcal{T}_k(B) - A\| &= \|Q\mathcal{T}_k(B) - QB + QB - A\| \\ &\leq \|\mathcal{T}_k(B) - B\| + \|(I - QQ^\top)A\|\end{aligned}$$

Basic randomized algorithms for low-rank approx

Add oversampling. (usually small) integer p

Randomized Algorithm:

1. Draw standard Gaussian random matrix $\Omega \in \mathbb{R}^{n \times (k+p)}$.
2. Perform block mat-vec $Y = A\Omega$.
3. Compute (economic) QR decomposition $Y = QR$.
4. Form $B = Q^\top A$.
5. Return $\hat{A} = Q\mathcal{T}_k(B)$ (in factorized form)

Gold standard best rank- k approximation error:

- spectral norm: σ_{k+1}
- Frobenius norm: $\sqrt{\sigma_{k+1}^2 + \cdots + \sigma_n^2}$.

Randomized algorithm applied to test matrices

errors measured in spectral norm:

(a) Hilbert matrix, $k = 5$:

Optimal	mean	std	
0.0019	0.0092	0.0099	$p = 0$
0.0019	0.0026	0.0019	$p = 1$
0.0019	0.0019	0.0001	$p = 2$

(b) Matrix with slower decay, $k = 25$:

Optimal	mean	std	
0.0034	0.012	0.002	$p = 0$
0.0034	0.010	0.0015	$p = 2$
0.0034	0.0037	0.0002	$p = 25$

(c) Matrix with staircase sv, $k = 7$:

Optimal	mean	std	
0.010	0.038	0.025	$p = 0$
0.010	0.021	0.012	$p = 1$
0.010	0.012	0.005	$p = 2$

Analysis: general considerations

Goal: Say something sensible about $\|(I - QQ^\top)A\|$. Expected value, tail bounds, ... wrt random matrix Ω .

Significantly more complicated than what we have seen so far! Analysis of randomized low-rank approximation can be separated into two phases:

1. **Structural bound:** Derive bound that holds for (almost) every Ω .

This bound usually depends on Ω and dependence needs to be simple enough to facilitate 2nd phase.

2. **Stochastic analysis:** Derive expected value, tail bounds for structural bound using random matrix theory, concentration results, ...

1. Structural bound: construct a competitor

Write

$$A = \underbrace{U_k \Sigma_k V_k^\top}_{\mathcal{T}_k(A)} + \underbrace{U_\perp \Sigma_\perp V_\perp^\top}_E, \quad Y = A\Omega, \quad Q = \text{orth}(Y),$$

and assume that $\Omega_1 := V_k^\top \Omega$ has full row rank.

Since $QQ^\top A$ is best approximation to A with range contained in $\text{range}(Y)$:

$$\|(I - QQ^\top)A\|_F \leq \|A - YZ^\top\|_F \quad \text{for every } Z.$$

Choose $Z^\top := \Omega_1^\dagger V_k^\top$. Then

$$\begin{aligned} YZ^\top &= A\Omega\Omega_1^\dagger V_k^\top = U_k \Sigma_k \underbrace{V_k^\top \Omega \Omega_1^\dagger}_{I_k} V_k^\top + E\Omega\Omega_1^\dagger V_k^\top \\ &= \mathcal{T}_k(A) + \tilde{E}, \end{aligned}$$

where $\tilde{E} := U_\perp \Sigma_\perp (V_\perp^\top \Omega)(V_k^\top \Omega)^\dagger V_k^\top$.

The competitor recovers the dominant rank- k part exactly.

1. Structural bound

From the previous slide,

$$A - YZ^T = E - \tilde{E}.$$

The two terms are orthogonal in Frobenius inner product because $V_{\perp}^T V_k = 0$. Hence, $\|E - \tilde{E}\|_F^2 = \|E\|_F^2 + \|\tilde{E}\|_F^2$ and

$$\|(I - QQ^T)A\|_F^2 \leq \|\Sigma_{\perp}\|_F^2 + \|\Sigma_{\perp}(V_{\perp}^T \Omega)(V_k^T \Omega)^{\dagger}\|_F^2.$$

$$\underbrace{\|\Sigma_{\perp}\|_F^2}_{\text{best rank-}r \text{ error}} + \underbrace{\|\Sigma_{\perp}(V_{\perp}^T \Omega)(V_k^T \Omega)^{\dagger}\|_F^2}_{\text{penalty from sketching}}.$$

Key point: $(V_k^T \Omega)^{\dagger}$ only requires a lower singular-value bound $\implies$ subspace injection.

2. Stochastic analysis: why subspace injection suffices

Set $\Omega_1 = V_k^\top \Omega$, $\Omega_2 = V_\perp^\top \Omega$. The structural bound is

$$\|(I - QQ^\top)A\|_F^2 \leq \|\Sigma_\perp\|_F^2 + \|\Sigma_\perp \Omega_2 \Omega_1^\dagger\|_F^2.$$

Suppose Ω is isotropic and, with probability at least $3/4$, is an ϵ -subspace injection for $\text{range}(V_k)$. Then

$$\|\Omega_1^\dagger\|_2^2 \leq \frac{1}{1 - \epsilon}.$$

Moreover, isotropy gives

$$\mathbb{E}\|\Sigma_\perp \Omega_2\|_F^2 = \|\Sigma_\perp\|_F^2.$$

By Markov,

$$\|\Sigma_\perp \Omega_2\|_F^2 \leq 4\|\Sigma_\perp\|_F^2$$

with probability at least $3/4$.

Hence, with probability at least $1/2$,

$$\|(I - QQ^\top)A\|_F^2 \leq \left(1 + \frac{4}{1 - \epsilon}\right) \|A - \mathcal{T}_k(A)\|_F^2.$$

2. Stochastic analysis: Gaussian test matrix

Let

$$\Omega \in \mathbb{R}^{n \times (k+p)}, \quad p \geq 2,$$

be standard Gaussian. By rotational invariance,

$$G_1 := V_k^\top \Omega \in \mathbb{R}^{k \times (k+p)}, \quad G_2 := V_\perp^\top \Omega$$

are independent standard Gaussian matrices.

Conditioning on G_1 ,

$$\mathbb{E}_{G_2} \|\Sigma_\perp G_2 G_1^\dagger\|_F^2 = \|\Sigma_\perp\|_F^2 \|G_1^\dagger\|_F^2.$$

Inverse-Wishart identity [Muirhead] gives $\mathbb{E} \|G_1^\dagger\|_F^2 = \frac{k}{p-1}$. Therefore

$$\mathbb{E} \|(I - QQ^\top)A\|_F^2 \leq \left(1 + \frac{k}{p-1}\right) \|A - \mathcal{T}_k(A)\|_F^2.$$

In particular,

$$\mathbb{E} \|(I - QQ^\top)A\|_F \leq \sqrt{1 + \frac{k}{p-1}} \|A - \mathcal{T}_k(A)\|_F.$$

General symmetric matrices

If A is symmetric, the approximation $QQ^\top A$ of the randomized SVD does *not* preserve symmetry.

Idea: Use Q from randomized SVD and approximate range+co-range of A at the same time:

$$QQ^\top AQQ^\top.$$

This produces an error on the level of the randomized SVD error

$$\varepsilon := \|A - QQ^\top A\|_F:$$

$$\begin{aligned}\|A - QQ^\top AQQ^\top\|_F^2 &= \|A - QQ^\top A + QQ^\top A - QQ^\top AQQ^\top\|_F^2 \\ &= \|(I - QQ^\top)A\|_F^2 + \|QQ^\top A(I - QQ^\top)\|_F^2 \\ &\leq 2\|(I - QQ^\top)A\|_F^2,\end{aligned}$$

and hence $\|A - QQ^\top AQQ^\top\|_F \leq \sqrt{2}\varepsilon$.

SPSD matrices: Nyström

For a symmetric positive semi-definite (SPSD) matrix A , one can improve upon $QQ^T A QQ^T$ on two aspects: (1) halving #matvecs, and (2) ensuring streaming property.

Basic idea: If SPSPD A has rank k and Ω is $n \times (k + p)$ Gaussian with $p \geq 0$ then $A\Omega(\Omega^T A\Omega)^\dagger \Omega^T$ is almost surely an oblique projector onto range of A . Hence,

$$\hat{A} := A\Omega(\Omega^T A\Omega)^\dagger (A\Omega)^T$$

is (almost surely) equal to A .

For general SPSPD A , approximation $\hat{A}$ is called **Nyström approximation**. Choose $p \geq 2$.

Highly recommended reading on Nyström: Tropp/Webber arXiv'2023.

SPSD matrices: Error analysis of Nyström

Error of Nyström satisfies

$$A - \hat{A} = A^{1/2}(I - \Pi_{A^{1/2}\Omega})A^{1/2}, \quad \Pi_{A^{1/2}\Omega} := A^{1/2}\Omega(\Omega^\top A\Omega)^\dagger(A^{1/2}\Omega)^\top$$

Note: $\Pi_{A^{1/2}\Omega}$ is orth. projector (onto range of $\Pi_{A^{1/2}\Omega}$). This implies:

$I - \Pi_{A^{1/2}\Omega}$ is orth. projector $\Rightarrow I - \Pi_{A^{1/2}\Omega}$ is SPSP $\Rightarrow A - \hat{A}$ is SPSP.

Measure error in nuclear norm:

$$\|A - \hat{A}\|_* = \sigma_1(A - \hat{A}) + \dots + \sigma_n(A - \hat{A}) = \lambda_1(A - \hat{A}) + \dots + \lambda_n(A - \hat{A}) = \text{trace}(A - \hat{A})$$

Hence,

$$\|A - \hat{A}\|_* = \|A^{1/2}(I - \Pi_{A^{1/2}\Omega})(I - \Pi_{A^{1/2}\Omega})A^{1/2}\|_* = \|(I - \Pi_{A^{1/2}\Omega})A^{1/2}\|_F^2.$$

But this is the error of the randomized SVD applied to $A^{1/2}$! Hence,

$$\begin{aligned} \mathbb{E}\|A - \hat{A}\|_* &= \mathbb{E}\|(I - \Pi_{A^{1/2}\Omega})A^{1/2}\|_F^2 \leq \left(1 + \frac{k}{p-1}\right) \|A^{1/2} - \mathcal{T}_k(A^{1/2})\|_F^2 \\ &= \left(1 + \frac{k}{p-1}\right) (\lambda_{r+1}(A) + \dots + \lambda_n(A)) \end{aligned}$$

Second term is best rank- k approx error in nuclear norm.