

ON THE STABILITY OF SINGULAR VALUES AND SUBSPACES UNDER PERTURBATION

ARJUN VELURI

ABSTRACT. In statistics, signal processing, and machine learning, the information we want from a matrix is carried by its singular values and singular vectors, yet the matrix we observe is almost never the one we want: it has been corrupted by measurement error, finite-sample fluctuation, or noise. Matrix perturbation theory makes precise what survives this corruption, and this paper is a self-contained account of that theory for the *singular value decomposition*, organized around one question: when A is replaced by $A + E$, how far can its singular values and singular subspaces move? We construct the decomposition from first principles and prove that it exists for every matrix; we then prove the Eckart–Young theorem, which identifies the best low-rank approximation of a matrix, together with Mirsky’s extension of it to the spectral norm. The perturbation theory itself divides in two: Weyl’s inequality shows that each singular value moves by at most the spectral-norm size of E , while the Davis–Kahan theorem shows that a singular subspace rotates by an angle controlled by that same size divided by the spectral gap isolating it. Gershgorin’s circle theorem enters as a localization tool, confining the eigenvalues of a matrix to discs read directly from its entries. Read together, the theorems answer a question at the center of working with real data: when the input is noisy, which features of a matrix are signal that persists, and which are noise that does not.

1. INTRODUCTION

Perturbing a matrix—replacing it by a nearby matrix and asking how its invariants respond—is one of the recurring problems of applied linear algebra, and it is forced on us by circumstance rather than chosen for interest. In nearly every setting where mathematics meets data, the matrix we can write down is not the one we care about: a measured matrix carries rounding error, a sample covariance differs from the population covariance, an observed signal arrives mixed with noise. We compute not with the exact object but with a corrupted copy of it, and the conclusions we draw are only as trustworthy as their stability under that corruption. The governing question is always the same: does a small change to the input produce a small change in the output, or does the structure we extracted dissolve under it? Matrix perturbation theory exists to answer this question quantitatively, replacing the hope that a method is robust with a bound that says how robust, and under what conditions.

The object at the center of this paper is the *singular value decomposition*. Every real $m \times n$ matrix A factors as

$$A = U\Sigma V^T, \tag{1}$$

where $U \in \mathbb{R}^{m \times m}$ and $V \in \mathbb{R}^{n \times n}$ are orthogonal and $\Sigma \in \mathbb{R}^{m \times n}$ is diagonal with nonnegative entries $\sigma_1 \geq \sigma_2 \geq \dots \geq 0$. The factorization resolves the action of A into three stages—a rotation, a scaling along orthogonal axes, and a second rotation—so that the columns v_i of V are the input directions, the columns u_i of U are the output directions, and the *singular values* σ_i record how far A stretches each one. The singular values are therefore a ranked inventory of the map’s strength: σ_1 marks the direction along which A acts most forcefully, σ_2 the next, and the leading values carry the dominant behaviour of the matrix. What earns the decomposition its central place, rather than a seat among other factorizations, is its universality: it exists for every matrix, square or rectangular, of full rank or not, with no hypothesis required [5, 4].

This ranking turns the decomposition into a tool for approximation. Suppose we cannot afford to keep all of A and must replace it by a matrix of rank at most k ; the natural candidate retains the k largest singular values and discards the rest,

$$A_k = \sum_{i=1}^k \sigma_i u_i v_i^T. \quad (2)$$

The *Eckart–Young theorem* confirms that this candidate is optimal in the Frobenius norm, and Mirsky’s extension shows that the same A_k stays optimal in the spectral norm—indeed in every unitarily invariant norm—so that no matrix of rank at most k approximates A better than its own truncated decomposition [2, 7, 4]. The conclusion is not obvious. The matrices of rank at most k form a vast, curved set with no a priori reason to favour one built from the singular vectors of A ; that the optimum can be read off the decomposition with no search is the first sign that the singular values encode exactly the information a low-rank approximation needs.

Eckart–Young and Mirsky describe a matrix we can see, and in practice we cannot see it. What we measure is a perturbed matrix

$$\tilde{A} = A + E, \quad (3)$$

in which A is the structure of interest and E is an error we neither chose nor know. The approximation machinery still runs, since it returns the best low-rank fit to whatever matrix it receives, but it now acts on $\tilde{A}$, and the truncation we compute is assembled from the singular values and singular vectors of the corrupted matrix rather than the true one. Everything therefore turns on how far those corrupted invariants can stray, and the question separates cleanly in two: by how much can the singular values shift, and by how much can the singular subspaces rotate? Weyl’s inequality answers the first; the Davis–Kahan theorem answers the second.

Both answers are clean enough to state in a line. Weyl’s inequality bounds each singular value on its own,

$$|\tilde{\sigma}_i - \sigma_i| \leq \|E\|_2, \quad (4)$$

so no singular value moves by more than the spectral-norm size of the noise, and the large, informative values are the most stable in relative terms [5]. The subspaces demand more care, because a singular vector is well defined only when its singular value stands apart from the others: the Davis–Kahan theorem bounds the rotation of a subspace by $\|E\|_2$ divided by the gap between its singular values and the rest of the spectrum, making precise the principle that a direction is stable exactly to the degree that it is isolated [1]. Gershgorin’s circle theorem contributes a complementary, entrywise form of control, trapping the eigenvalues of a matrix inside discs whose centers and radii are computed from the entries alone [5]. Taken together, these results form a complete first-order stability theory for the singular value decomposition: a precise account of which parts of the answer a perturbation can move, and by how much.

These are not internal curiosities about matrices. Principal component analysis, dimensionality reduction, covariance estimation, and the spectral methods used for clustering and signal recovery all proceed by computing a singular value decomposition and retaining its leading part, and each is applied to data that is noisy by construction [3]. The perturbation theorems are what license the practice: they explain why a decomposition computed from corrupted measurements still recovers the structure of the clean matrix beneath, and they mark the boundary past which that recovery can no longer be trusted.

The paper builds these results from the ground up.

2. PRELIMINARIES

This section assembles those objects in the order they will be needed: the norms that quantify the size of a matrix and of its perturbation, the orthogonal matrices through which the singular value decomposition is expressed, the spectral theorem that guarantees the decomposition exists, and the

principal angles that record how far one subspace has rotated from another. We work throughout over the real numbers, so that “transpose” may be read everywhere for “conjugate transpose”; the complex case requires no new ideas.

2.1. Vectors, inner products, and norms. We begin with the two quantities every later bound is written in: the length of a vector and the size of a matrix.

Definition 2.1. The *standard inner product* of vectors $x, y \in \mathbb{R}^n$ is

$$\langle x, y \rangle = \sum_{i=1}^n x_i y_i = x^\top y, \quad (5)$$

and the *Euclidean norm* (or ℓ_2 norm) of x is $\|x\| = \sqrt{\langle x, x \rangle}$.

Remark 2.2. The inner product encodes geometry: $\|x\|$ is the length of x , the quantity $\langle x, y \rangle = \|x\| \|y\| \cos \theta$ measures the angle θ between x and y , and x and y are *orthogonal* when $\langle x, y \rangle = 0$. Everything in this paper is ultimately a statement about lengths and angles, and (5) is the instrument that reads them off.

A matrix acts on vectors, and there are two natural ways to measure how large that action is. The first records the largest stretching the matrix can produce.

Definition 2.3. The *spectral norm* of $A \in \mathbb{R}^{m \times n}$ is

$$\|A\|_2 = \max_{\|x\|=1} \|Ax\|. \quad (6)$$

Remark 2.4. The spectral norm is the operator norm induced by the Euclidean norm: it is the worst-case amplification $\|Ax\| / \|x\|$ over all nonzero x , so $\|Ax\| \leq \|A\|_2 \|x\|$ for every x , with equality attained on the direction A stretches most. When A is a perturbation E , the quantity $\|E\|_2$ is exactly the largest distortion the noise can inflict on any single vector, which is why it is the natural yardstick for the bounds to come.

The second measure treats the matrix as a long list of numbers and takes their Euclidean length.

Definition 2.5. The *Frobenius norm* of $A = (a_{ij}) \in \mathbb{R}^{m \times n}$ is

$$\|A\|_F = \left(\sum_{i=1}^m \sum_{j=1}^n a_{ij}^2 \right)^{1/2}. \quad (7)$$

The Frobenius norm has a trace expression that makes its orthogonal invariance, proved below, almost immediate.

Lemma 2.6. For every $A \in \mathbb{R}^{m \times n}$, $\|A\|_F^2 = \text{tr}(A^\top A)$.

Proof. The j th diagonal entry of $A^\top A$ is the inner product of the j th column of A with itself, namely $\sum_{i=1}^m a_{ij}^2$. Summing these diagonal entries over j recovers the double sum in (7). $\square$

The two norms are not independent; each controls the other, with the rank of the matrix governing the gap between them.

Remark 2.7. For every $A \in \mathbb{R}^{m \times n}$ of rank r ,

$$\|A\|_2 \leq \|A\|_F \leq \sqrt{r} \|A\|_2. \quad (8)$$

Both inequalities become transparent once the singular value decomposition is in hand (Section 3), where $\|A\|_2 = \sigma_1$ and $\|A\|_F = (\sigma_1^2 + \dots + \sigma_r^2)^{1/2}$: the left inequality says that a single term is at most the whole sum, and the right that a sum of r terms, each at most σ_1^2 , is at most $r\sigma_1^2$. The spectral norm sees only the leading singular value, while the Frobenius norm pools them all, and the two coincide precisely when A has rank one.

2.2. Orthogonality and orthogonal matrices. The singular value decomposition expresses an arbitrary matrix through orthogonal matrices, the linear maps that move space rigidly. We isolate their defining property and the invariance that we will invoke more than any other fact in the paper.

Definition 2.8. A set $\{q_1, \dots, q_k\} \subseteq \mathbb{R}^n$ is *orthonormal* if $\langle q_i, q_j \rangle = 1$ when $i = j$ and $\langle q_i, q_j \rangle = 0$ otherwise. An orthonormal set of n vectors in $\mathbb{R}^n$ is an *orthonormal basis*. A square matrix $Q \in \mathbb{R}^{n \times n}$ is *orthogonal* if $Q^\top Q = I$, equivalently if its columns form an orthonormal basis of $\mathbb{R}^n$.

Remark 2.9. The condition $Q^\top Q = I$ forces $Q^{-1} = Q^\top$, hence also $QQ^\top = I$, so the rows of an orthogonal matrix are orthonormal as well. Geometrically, the orthogonal matrices are exactly the rotations and reflections of $\mathbb{R}^n$: the maps that fix the origin and preserve all distances.

Lemma 2.10. *If $Q \in \mathbb{R}^{n \times n}$ is orthogonal, then $\langle Qx, Qy \rangle = \langle x, y \rangle$ for all $x, y \in \mathbb{R}^n$, and in particular $\|Qx\| = \|x\|$.*

Proof. A rigid motion can change neither lengths nor angles, and the algebra confirms it: $\langle Qx, Qy \rangle = (Qx)^\top (Qy) = x^\top Q^\top Q y = x^\top y = \langle x, y \rangle$, using $Q^\top Q = I$. Taking $y = x$ gives $\|Qx\|^2 = \|x\|^2$. $\square$

Lemma 2.10 concerns vectors, but its consequence for matrices is what we use repeatedly: multiplying a matrix on either side by an orthogonal matrix changes neither its spectral nor its Frobenius norm.

Lemma 2.11. *Let $A \in \mathbb{R}^{m \times n}$, and let $P \in \mathbb{R}^{m \times m}$ and $Q \in \mathbb{R}^{n \times n}$ be orthogonal. Then $\|PAQ\|_2 = \|A\|_2$ and $\|PAQ\|_F = \|A\|_F$.*

Proof. For the spectral norm, write $\|PAQ\|_2 = \max_{\|x\|=1} \|PAQx\|$. Since P is orthogonal, $\|PAQx\| = \|AQx\|$ by Lemma 2.10; since Q is orthogonal, Qx ranges over the entire unit sphere as x does, so the maximum equals $\max_{\|y\|=1} \|Ay\| = \|A\|_2$. For the Frobenius norm, Lemma 2.6 and the cyclic property of the trace give $\|PAQ\|_F^2 = \text{tr}(Q^\top A^\top P^\top PAQ) = \text{tr}(A^\top A Q Q^\top) = \text{tr}(A^\top A) = \|A\|_F^2$. $\square$

This invariance is the reason the singular value decomposition is so well suited to measuring perturbations: the orthogonal factors U and V in $A = U\Sigma V^\top$ may be stripped away without disturbing either norm, leaving the singular values to carry all the size information by themselves.

2.3. Symmetric matrices and the spectral theorem. The construction of the singular value decomposition in Section 3 rests on a single structural result about symmetric matrices, which we state here.

Definition 2.12. A matrix $A \in \mathbb{R}^{n \times n}$ is *symmetric* if $A^\top = A$. A symmetric matrix A is *positive semidefinite* if $x^\top Ax \geq 0$ for every $x \in \mathbb{R}^n$.

Remark 2.13. Positive semidefiniteness is the matrix analogue of nonnegativity: the quadratic form $x \mapsto x^\top Ax$ never turns a vector against itself. We will meet such matrices in exactly one guise, as products $A^\top A$, and the remark following the spectral theorem explains why.

Theorem 2.14 (Spectral theorem). *Every real symmetric matrix $A \in \mathbb{R}^{n \times n}$ admits a factorization*

$$A = Q\Lambda Q^\top, \tag{9}$$

where $Q \in \mathbb{R}^{n \times n}$ is orthogonal and $\Lambda = \text{diag}(\lambda_1, \dots, \lambda_n)$ is real and diagonal. The columns of Q are orthonormal eigenvectors of A , and the diagonal entries of Λ are the corresponding eigenvalues.

A symmetric matrix thus acts by stretching along n mutually perpendicular axes, the eigenvectors, with the eigenvalues as the stretch factors. We take Theorem 2.14 as the one substantial import from linear algebra and refer the reader to [5] for a proof; everything else in the paper is built from it. The bridge to the next section is the following observation, which turns the spectral theorem, a statement about square symmetric matrices, into a tool for arbitrary rectangular ones.

Remark 2.15. For any $A \in \mathbb{R}^{m \times n}$, the matrix $A^\top A \in \mathbb{R}^{n \times n}$ is symmetric, since $(A^\top A)^\top = A^\top A$, and positive semidefinite, since $x^\top A^\top A x = \|Ax\|^2 \geq 0$. By Theorem 2.14 its eigenvalues are therefore real and nonnegative. This single fact is the seed of the singular value decomposition: Section 3 defines the singular values of A as the square roots of these eigenvalues and the right singular vectors as the eigenvectors of $A^\top A$, so that the decomposition of an arbitrary matrix is manufactured from the spectral decomposition of a symmetric one.

2.4. Rank, subspaces, and the orthogonal complement. The approximation theory of Section 4 compares a matrix against all matrices of a given rank, so we fix the language of rank and of the subspaces attached to a matrix.

Definition 2.16. The *column space* $\text{Col}(A)$ of $A \in \mathbb{R}^{m \times n}$ is the span of its columns, and its *null space* is $\text{Null}(A) = \{x \in \mathbb{R}^n : Ax = 0\}$. The *rank* of A is $\text{rank}(A) = \dim \text{Col}(A)$.

Theorem 2.17 (Rank–nullity). *For every $A \in \mathbb{R}^{m \times n}$, $\text{rank}(A) + \dim \text{Null}(A) = n$.*

We use Theorem 2.17 only as a bookkeeping device and refer to [4] for its proof. The geometry of subspaces enters through orthogonality.

Definition 2.18. The *orthogonal complement* of a subspace $S \subseteq \mathbb{R}^n$ is

$$S^\perp = \{x \in \mathbb{R}^n : \langle x, s \rangle = 0 \text{ for all } s \in S\}. \quad (10)$$

Remark 2.19. Every subspace and its complement together span the whole space: $\mathbb{R}^n = S \oplus S^\perp$, so each vector splits uniquely into a part lying in S and a part orthogonal to it. The map sending a vector to its S -part is the orthogonal projection onto S , and the length of the discarded part measures how far the vector lies from S . This projection picture is the one the subspace distance of the next subsection makes quantitative.

Remark 2.20. The matrices of rank at most k are the natural arena for low-rank approximation: they are precisely the matrices expressible as a sum of k rank-one terms, hence the most that k “directions” can store. Crucially they do not form a subspace, since the sum of two rank- k matrices typically has rank $2k$, and it is this curvature that makes the optimality asserted by the Eckart–Young theorem (Section 4) a genuine theorem rather than an exercise in projection.

2.5. Principal angles and the distance between subspaces. The perturbation theorem for singular subspaces, proved in Section 6, must measure how far a perturbed subspace has turned from the true one. Ordinary subtraction will not serve. A subspace has no canonical basis, and if U and V are matrices whose columns happen to span two subspaces, then $\|U - V\|$ depends on the arbitrary choice of those columns rather than on the subspaces themselves. We need a quantity that depends on the subspaces alone, and the principal angles supply it.

Definition 2.21. Let $U, V \in \mathbb{R}^{n \times d}$ have orthonormal columns, spanning the d -dimensional subspaces $\mathcal{U} = \text{Col}(U)$ and $\mathcal{V} = \text{Col}(V)$. Let the singular values of $U^\top V$ be $1 \geq s_1 \geq \dots \geq s_d \geq 0$. The *principal angles* $\theta_1 \leq \dots \leq \theta_d$ between $\mathcal{U}$ and $\mathcal{V}$ are

$$\theta_i = \arccos(s_i) \in [0, \frac{\pi}{2}], \quad i = 1, \dots, d. \quad (11)$$

Remark 2.22. The singular values of $U^\top V$ lie in $[0, 1]$ because U and V have orthonormal columns, so the angles in (11) are well defined. They interpolate between best and worst alignment: $\theta_1 = 0$ signals a direction common to both subspaces, while a large θ_d exposes a direction in $\mathcal{U}$ nearly orthogonal to all of $\mathcal{V}$. When $d = 1$ the subspaces are lines and θ_1 is the ordinary angle between them, so the principal angles are the honest generalization of that angle to higher dimension. Replacing U by UR for an orthogonal R leaves the singular values of $U^\top V$ unchanged, which is precisely the basis-independence ordinary subtraction lacked: the angles depend on $\mathcal{U}$ and $\mathcal{V}$, not on the matrices that represent them.

The angles are collected into a single matrix whose norm is the distance we will bound.

Definition 2.23. With the principal angles of Definition 2.21, set $\Theta(\mathcal{U}, \mathcal{V}) = \text{diag}(\theta_1, \dots, \theta_d)$ and

$$\sin \Theta(\mathcal{U}, \mathcal{V}) = \text{diag}(\sin \theta_1, \dots, \sin \theta_d). \quad (12)$$

The *distance* between $\mathcal{U}$ and $\mathcal{V}$ is $\|\sin \Theta(\mathcal{U}, \mathcal{V})\|$, measured in either the spectral or the Frobenius norm.

Remark 2.24. This is the right notion of distance for three reasons. It is basis-independent, inheriting that property from the principal angles. It vanishes exactly when $\mathcal{U} = \mathcal{V}$, since every principal angle is then zero, and grows toward its maximum as the subspaces become orthogonal. And it carries a concrete projection meaning: the numbers $\sin \theta_i$ are the singular values of $(I - \mathcal{V}\mathcal{V}^\top)\mathcal{U}$, the result of projecting an orthonormal basis of $\mathcal{U}$ onto $\mathcal{V}^\perp$. The distance therefore measures the part of $\mathcal{U}$ that escapes $\mathcal{V}$, and it is exactly this quantity that the Davis–Kahan theorem controls [1, 8]. Figure 1 illustrates the simplest case.

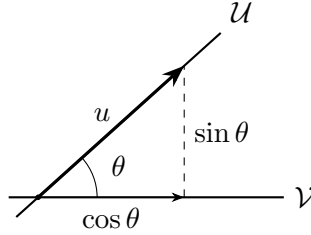


FIGURE 1. For two lines meeting at a single principal angle θ , the unit vector u spanning $\mathcal{U}$ splits into its projection of length $\cos \theta$ onto $\mathcal{V}$ and a perpendicular component of length $\sin \theta$ lying in $\mathcal{V}^\perp$. The distance $\|\sin \Theta(\mathcal{U}, \mathcal{V})\| = \sin \theta$ vanishes exactly when the lines coincide.

2.6. Notation for the rest of the paper. We close by fixing the symbols that recur from Section 3 onward. Throughout, $A \in \mathbb{R}^{m \times n}$ denotes the true matrix whose structure we wish to recover, and $\tilde{A} = A + E$ the matrix we actually observe, with E the perturbation, or noise. Their singular values, always listed in nonincreasing order, are $\sigma_1 \geq \sigma_2 \geq \dots$ for A and $\tilde{\sigma}_1 \geq \tilde{\sigma}_2 \geq \dots$ for $\tilde{A}$, with corresponding left and right singular vectors u_i, v_i and $\tilde{u}_i, \tilde{v}_i$. The leading k left and right singular vectors form the matrices $U_k = [u_1 \ \dots \ u_k]$ and $V_k = [v_1 \ \dots \ v_k]$, whose column spaces are the leading singular subspaces; the perturbed matrices $\tilde{U}_k$ and $\tilde{V}_k$ are defined the same way from $\tilde{A}$. When a perturbation bound for these subspaces requires the singular values they involve to be separated from the rest of the spectrum, we write δ for the size of that separation, the *eigengap*, fixing its precise form where each theorem needs it. Unsubscripted, $\|\cdot\|_2$ is the spectral norm and $\|\cdot\|_F$ the Frobenius norm.

3. THE SINGULAR VALUE DECOMPOSITION

The spectral theorem (Theorem 2.14) gives a complete account of a symmetric matrix: it acts by scaling along orthogonal axes, and nothing more. Most matrices are not symmetric, and many are not even square, so the theorem says nothing about them directly. A rectangular matrix nonetheless maps one space into another, and one may still ask along which directions it acts and by how much. The answer is the singular value decomposition, which factors an arbitrary matrix into a rotation, a nonnegative scaling, and a second rotation, carrying the geometric clarity of the spectral theorem over from symmetric matrices to all of them. What makes the decomposition indispensable rather than merely available is that it asks for nothing in return: no symmetry, no

squareness, no invertibility. The mechanism is already present in Remark 2.15, in the fact that $A^\top A$ is symmetric and positive semidefinite for every A ; the task of this section is to turn that fact into the decomposition itself.

3.1. Existence. The construction follows the route Remark 2.15 marks out. Since $A^\top A$ is symmetric and positive semidefinite, the spectral theorem diagonalizes it with real, nonnegative eigenvalues; their square roots become the singular values, their eigenvectors become the right singular vectors, and the left singular vectors are the normalized images of the right ones under A . The single point requiring care is that A may send some directions to zero, so the construction must still produce a full orthonormal frame in the output space.

Theorem 3.1 (Existence of the singular value decomposition). *Every matrix $A \in \mathbb{R}^{m \times n}$ admits a factorization*

$$A = U \Sigma V^\top, \quad (13)$$

in which $U \in \mathbb{R}^{m \times m}$ and $V \in \mathbb{R}^{n \times n}$ are orthogonal and $\Sigma \in \mathbb{R}^{m \times n}$ has entries $\Sigma_{ii} = \sigma_i$ for $1 \leq i \leq \min(m, n)$ and zeros elsewhere, with $\sigma_1 \geq \sigma_2 \geq \dots \geq \sigma_{\min(m, n)} \geq 0$. The scalars σ_i are the singular values of A , the columns of U its left singular vectors, and the columns of V its right singular vectors.

Proof. By Remark 2.15 the matrix $A^\top A \in \mathbb{R}^{n \times n}$ is symmetric and positive semidefinite, so Theorem 2.14 supplies an orthogonal matrix $V = [v_1 \ \dots \ v_n]$ and real eigenvalues $\lambda_1 \geq \dots \geq \lambda_n \geq 0$ with

$$A^\top A v_i = \lambda_i v_i, \quad i = 1, \dots, n. \quad (14)$$

The eigenvalues are nonnegative because positive semidefiniteness forces $\lambda_i = v_i^\top A^\top A v_i = \|Av_i\|^2 \geq 0$. Define the singular values $\sigma_i = \sqrt{\lambda_i}$, so that $\sigma_1 \geq \dots \geq \sigma_n \geq 0$, and let r be the number of them that are strictly positive, so $\sigma_1 \geq \dots \geq \sigma_r > 0$ and $\sigma_{r+1} = \dots = \sigma_n = 0$.

For each $i \leq r$ the image Av_i is nonzero, and we set

$$u_i = \frac{1}{\sigma_i} Av_i \in \mathbb{R}^m. \quad (15)$$

These vectors are orthonormal. The eigenvector relation (14) lets the inner product of two images be evaluated through $A^\top A$ alone,

$$\langle u_i, u_j \rangle = \frac{1}{\sigma_i \sigma_j} (Av_i)^\top (Av_j) = \frac{1}{\sigma_i \sigma_j} v_i^\top A^\top A v_j = \frac{\lambda_j}{\sigma_i \sigma_j} \langle v_i, v_j \rangle. \quad (16)$$

Because the v_i are orthonormal, the right-hand side is 0 when $i \neq j$, and when $i = j$ it equals $\lambda_i / \sigma_i^2 = 1$. Hence $\{u_1, \dots, u_r\}$ is an orthonormal set in $\mathbb{R}^m$, and in particular $r \leq m$.

Extend this set to an orthonormal basis $u_1, \dots, u_m$ of $\mathbb{R}^m$, which Gram–Schmidt accomplishes after completing the set to any basis [4], and assemble the orthogonal matrix $U = [u_1 \ \dots \ u_m]$. Let $\Sigma \in \mathbb{R}^{m \times n}$ carry $\sigma_1, \dots, \sigma_{\min(m, n)}$ on its diagonal as in the statement.

It remains to confirm (13). Two matrices coincide precisely when they act identically on a basis, so it suffices to check that A and $U \Sigma V^\top$ agree on each v_i . Since the v_i are the columns of the orthogonal matrix V , we have $V^\top v_i = e_i$, and therefore

$$U \Sigma V^\top v_i = U \Sigma e_i = \begin{cases} \sigma_i u_i, & i \leq \min(m, n), \\ 0, & i > \min(m, n), \end{cases} \quad (17)$$

using that the i th column of Σ is $\sigma_i e_i$ when $i \leq \min(m, n)$ and is zero otherwise. Compare this with Av_i . For $i \leq r$, the definition (15) gives $Av_i = \sigma_i u_i$, which matches (17). For $i > r$ the image vanishes, since $\|Av_i\|^2 = \lambda_i = 0$, and the right side of (17) is likewise zero, because $\sigma_i = 0$ in the range $r < i \leq \min(m, n)$ and the column is zero beyond it. Thus $Av_i = U \Sigma V^\top v_i$ for every i , and (13) follows. $\square$

Remark 3.2. The three factors of (13) are the precise form of the rotation, scaling, and rotation described in the introduction. Reading $A = U\Sigma V^\top$ from right to left, $V^\top$ turns the orthonormal right singular vectors $v_1, \dots, v_n$ onto the coordinate axes, Σ scales the i th axis by σ_i , and U turns the scaled axes onto the orthonormal left singular vectors $u_1, \dots, u_m$. Equivalently, the identity $Av_i = \sigma_i u_i$ says that A sends each right singular vector to a multiple of the corresponding left one, so the unit sphere of $\mathbb{R}^n$ maps to an ellipsoid whose semi-axes point along the u_i and have lengths σ_i (Figure 2). The singular values are the lengths of those axes, and they rank the directions of the domain by how forcefully A acts along them.

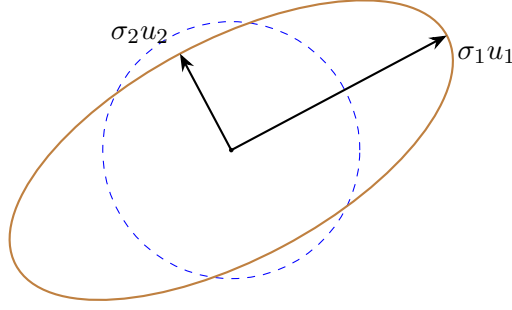


FIGURE 2. The image under A of the unit sphere (dashed) is an ellipsoid (solid) whose semi-axes lie along the left singular vectors u_i and have lengths equal to the singular values σ_i . The factorization $A = U\Sigma V^\top$ reads this picture from right to left: $V^\top$ aligns the right singular vectors with the axes, Σ stretches axis i by σ_i , and U aligns the result with the u_i .

3.2. Uniqueness. The construction made several choices: an ordering among equal eigenvalues, a sign for each eigenvector, a basis for each eigenspace, and a completion of the left vectors. The decomposition remembers only some of them, and the following statement records exactly which.

Proposition 3.3 (Uniqueness). *Let $A = U\Sigma V^\top$ be any singular value decomposition of $A \in \mathbb{R}^{m \times n}$. The singular values $\sigma_1 \geq \sigma_2 \geq \dots \geq 0$ are uniquely determined by A . If a positive singular value σ_i is simple, its left and right singular vectors u_i, v_i are determined up to a common sign. If a positive singular value has multiplicity p , the associated left and right singular vectors are determined only up to a simultaneous orthogonal change of basis within the two p -dimensional subspaces they span.*

Proof. From $A = U\Sigma V^\top$ and the orthogonality of U ,

$$A^\top A = V\Sigma^\top U^\top U\Sigma V^\top = V(\Sigma^\top \Sigma)V^\top, \quad \Sigma^\top \Sigma = \text{diag}(\sigma_1^2, \dots, \sigma_n^2). \quad (18)$$

Equation (18) is therefore a spectral decomposition of $A^\top A$, so the numbers σ_i^2 are its eigenvalues. Eigenvalues are intrinsic to a matrix, so the ordered singular values are determined by A . The columns v_i are orthonormal eigenvectors of $A^\top A$: when σ_i^2 is a simple eigenvalue its unit eigenvector is determined up to sign, and $u_i = \sigma_i^{-1} Av_i$ then inherits that sign; when σ_i^2 has multiplicity p , only its eigenspace is determined, any orthonormal basis of which yields valid right singular vectors, with the matching left vectors spanning the corresponding p -dimensional subspace of $\mathbb{R}^m$. $\square$

Remark 3.4. The dependence on multiplicity is not a technicality but a warning that shapes the perturbation theory to come. When two singular values coincide, the individual singular vectors cease to be defined, and when two singular values are merely close, a small change in A can swing the vectors widely while barely disturbing the subspace they span. A stable measure of how a perturbation moves singular vectors must therefore track subspaces rather than vectors. This is precisely the purpose of the $\sin \Theta$ distance of Definition 2.23, and it is why the Davis–Kahan bound of Section 6 is stated for singular subspaces rather than for the vectors themselves.

3.3. A worked example.

Example 3.5. Consider the matrix

$$A = \begin{pmatrix} 2 & 2 \\ 1 & 0 \\ 0 & 1 \end{pmatrix} \in \mathbb{R}^{3 \times 2}. \quad (19)$$

Following the proof of Theorem 3.1, we begin with

$$A^\top A = \begin{pmatrix} 5 & 4 \\ 4 & 5 \end{pmatrix}. \quad (20)$$

Its characteristic polynomial is $(5 - \lambda)^2 - 16 = \lambda^2 - 10\lambda + 9 = (\lambda - 9)(\lambda - 1)$, so the eigenvalues are $\lambda_1 = 9$ and $\lambda_2 = 1$, with unit eigenvectors

$$v_1 = \frac{1}{\sqrt{2}} \begin{pmatrix} 1 \\ 1 \end{pmatrix}, \quad v_2 = \frac{1}{\sqrt{2}} \begin{pmatrix} 1 \\ -1 \end{pmatrix}. \quad (21)$$

The singular values are their square roots, $\sigma_1 = 3$ and $\sigma_2 = 1$, and $V = [v_1 \ v_2]$. The formula (15) produces the left singular vectors from these,

$$u_1 = \frac{1}{\sigma_1} A v_1 = \frac{1}{3\sqrt{2}} \begin{pmatrix} 4 \\ 1 \\ 1 \end{pmatrix}, \quad (22)$$

$$u_2 = \frac{1}{\sigma_2} A v_2 = \frac{1}{\sqrt{2}} \begin{pmatrix} 0 \\ 1 \\ -1 \end{pmatrix}, \quad (23)$$

and a direct check gives $\|u_1\| = \|u_2\| = 1$ and $\langle u_1, u_2 \rangle = 0$. Because A has rank two, these two vectors fall one short of a basis for $\mathbb{R}^3$, and the proof calls for completing them. The missing direction is any unit vector orthogonal to both, and the normalized cross product supplies

$$u_3 = \frac{1}{3} \begin{pmatrix} -1 \\ 2 \\ 2 \end{pmatrix}, \quad (24)$$

which spans the left null space of A , the orthogonal complement of its column space. With $U = [u_1 \ u_2 \ u_3]$ and

$$\Sigma = \begin{pmatrix} 3 & 0 \\ 0 & 1 \\ 0 & 0 \end{pmatrix}, \quad (25)$$

the product $U\Sigma V^\top$ returns (19), as multiplying out confirms.

The example makes visible something the general proof states only in symbols. The singular values 3 and 1 are unequal, and their ratio is the exact sense in which A favors one direction over another: it stretches the input $v_1 = \frac{1}{\sqrt{2}}(1, 1)^\top$ three times as far as $v_2 = \frac{1}{\sqrt{2}}(1, -1)^\top$. The third left singular vector u_3 , forced in by the completion step, is not idle bookkeeping; it names the single direction of the output space that A never reaches.

3.4. From decomposition to approximation. Every matrix now carries the canonical data of (13): a ranked list of singular values together with the orthonormal frames they connect. The question is what this data is good for, and the first answer concerns approximation. When the singular values decay, the trailing terms contribute little to A , which invites discarding them and keeping only the few largest. Section 4 makes the idea precise and proves the optimality result at the center of the subject: among all matrices of a given rank, the truncated decomposition is the closest to A in both of the norms fixed in Section 2.

4. THE ECKART–YOUNG THEOREM

Section 3 closed with the suggestion that a matrix whose singular values decay is well approximated by discarding the small ones. We now make that suggestion exact. Fix $A \in \mathbb{R}^{m \times n}$ and an integer k with $1 \leq k \leq \text{rank}(A)$, and ask for the matrix of rank at most k that lies closest to A . The matrices of rank at most k are the natural competitors: by Remark 2.20 they are exactly the matrices storable in k directions, so they form the class against which any rank- k compression of A must be judged. The question is which member of the class minimizes the approximation error, measured in one of the norms of Section 2.

Nothing about the question suggests that its answer should be simple. The competitors form an enormous family, and a minimization over all of them might be expected to demand an exhaustive search, with an optimum depending intricately on every entry of A . The content of the Eckart–Young theorem is that no search is needed. The best rank- k approximation is obtained by keeping the k largest singular values of A together with the singular vectors that accompany them and discarding the rest, so the optimum is read directly off the decomposition constructed in Section 3.

4.1. The truncated decomposition.

Definition 4.1. Let $A = U\Sigma V^\top$ be a singular value decomposition (Theorem 3.1) with singular values $\sigma_1 \geq \dots \geq \sigma_r > 0$, where $r = \text{rank}(A)$, and left and right singular vectors u_i, v_i . For $1 \leq k \leq r$, the *truncated singular value decomposition* of rank k is

$$A_k = \sum_{i=1}^k \sigma_i u_i v_i^\top. \quad (26)$$

Remark 4.2. The matrix A_k costs nothing beyond the decomposition itself: it retains the leading k triples (σ_i, u_i, v_i) already produced by Theorem 3.1 and drops the remainder. Its column space is $\text{span}\{u_1, \dots, u_k\}$, so its rank is exactly k . Equivalently $A_k = U\Sigma_k V^\top$, where Σ_k keeps the top k diagonal entries of Σ and zeros the rest.

4.2. Optimality. We first record the error that A_k actually achieves; the rest of the work is to prove that nothing does better.

Lemma 4.3. $\|A - A_k\|_F^2 = \sum_{i=k+1}^r \sigma_i^2$ and $\|A - A_k\|_2 = \sigma_{k+1}$.

Proof. By (26), $A - A_k = \sum_{i=k+1}^r \sigma_i u_i v_i^\top$. The rank-one matrices $u_i v_i^\top$ are orthonormal in the Frobenius inner product $\langle X, Y \rangle_F = \text{tr}(X^\top Y)$, because

$$\langle u_i v_i^\top, u_j v_j^\top \rangle_F = \text{tr}(v_i u_i^\top u_j v_j^\top) = (u_i^\top u_j)(v_j^\top v_i), \quad (27)$$

which is 1 when $i = j$ and 0 otherwise. Expanding the squared norm of the sum therefore leaves only the diagonal terms, $\|A - A_k\|_F^2 = \sum_{i=k+1}^r \sigma_i^2$. For the spectral norm, $A - A_k = U\Sigma' V^\top$ where Σ' carries $\sigma_{k+1}, \dots, \sigma_r$ in its diagonal positions $k+1, \dots, r$, so orthogonal invariance (Lemma 2.11) gives $\|A - A_k\|_2 = \|\Sigma'\|_2 = \sigma_{k+1}$, the largest surviving entry. $\square$

The optimality proof rests on two facts. The first says that once the column space of an approximation is fixed, the best approximation with that column space is the orthogonal projection of A onto it.

Lemma 4.4 (Best approximation with a fixed column space). *Let $C \subseteq \mathbb{R}^m$ be a subspace with orthogonal projector P_C . Among all matrices B with $\text{Col}(B) \subseteq C$, the distance $\|A - B\|_F$ is minimized by $B = P_C A$, and*

$$\|A - P_C A\|_F^2 = \|A\|_F^2 - \text{tr}(P_C A A^\top). \quad (28)$$

Proof. The Frobenius norm separates over columns, so writing a_j and b_j for the columns of A and B gives $\|A - B\|_F^2 = \sum_j \|a_j - b_j\|^2$ with each $b_j \in C$. The closest point of C to a_j is the orthogonal projection $P_C a_j$, so the sum is minimized term by term at $b_j = P_C a_j$, that is, at $B = P_C A$. For the value, the Pythagorean identity $\|a_j\|^2 = \|P_C a_j\|^2 + \|a_j - P_C a_j\|^2$ summed over j yields $\|A - P_C A\|_F^2 = \|A\|_F^2 - \|P_C A\|_F^2$, and $\|P_C A\|_F^2 = \text{tr}(P_C A A^\top P_C) = \text{tr}(P_C A A^\top)$ by Lemma 2.6, the idempotence $P_C^2 = P_C$, and the cyclic invariance of the trace. $\square$

The second fact bounds how much of a positive semidefinite quadratic form a k -dimensional projection can capture. No k -dimensional slice sees more than the k largest eigendirections.

Lemma 4.5 (Trace maximum). *Let $M \in \mathbb{R}^{m \times m}$ be symmetric positive semidefinite with eigenvalues $\lambda_1 \geq \dots \geq \lambda_m \geq 0$. For every orthogonal projector P of rank k ,*

$$\text{tr}(PM) \leq \lambda_1 + \dots + \lambda_k. \quad (29)$$

Proof. Write $P = \sum_{j=1}^k q_j q_j^\top$ for an orthonormal set $q_1, \dots, q_k$, and let $M = \sum_{i=1}^m \lambda_i w_i w_i^\top$ be a spectral decomposition (Theorem 2.14) with $w_1, \dots, w_m$ orthonormal. Then

$$\text{tr}(PM) = \sum_{j=1}^k q_j^\top M q_j = \sum_{i=1}^m \lambda_i t_i, \quad t_i = \sum_{j=1}^k (w_i^\top q_j)^2. \quad (30)$$

Extending $q_1, \dots, q_k$ to an orthonormal basis of $\mathbb{R}^m$ shows $t_i \leq \sum_{j=1}^m (w_i^\top q_j)^2 = \|w_i\|^2 = 1$, while $\sum_{i=1}^m t_i = \sum_{j=1}^k \|q_j\|^2 = k$. Thus the weights t_i lie in $[0, 1]$ and sum to k , and since the λ_i are nonincreasing,

$$\sum_{i=1}^k \lambda_i - \sum_{i=1}^m \lambda_i t_i = \sum_{i=1}^k \lambda_i (1 - t_i) - \sum_{i>k} \lambda_i t_i \geq \lambda_k \left(\sum_{i=1}^k (1 - t_i) - \sum_{i>k} t_i \right) = 0, \quad (31)$$

the inequality using $\lambda_i \geq \lambda_k$ for $i \leq k$ and $\lambda_i \leq \lambda_k$ for $i > k$, and the final equality using $\sum_i t_i = k$. This is (29). $\square$

Theorem 4.6 (Eckart–Young). *Let $A \in \mathbb{R}^{m \times n}$ have rank r and singular values $\sigma_1 \geq \dots \geq \sigma_r > 0$. For every k with $1 \leq k \leq r$ and every $B \in \mathbb{R}^{m \times n}$ of rank at most k ,*

$$\|A - B\|_F \geq \|A - A_k\|_F = \left(\sum_{i=k+1}^r \sigma_i^2 \right)^{1/2}. \quad (32)$$

The truncated decomposition A_k is thus a closest rank- k matrix to A in the Frobenius norm [2].

Proof. The value $\|A - A_k\|_F = (\sum_{i>k} \sigma_i^2)^{1/2}$ is Lemma 4.3, so only the inequality remains. Let B have rank at most k and set $C = \text{Col}(B)$, a subspace of dimension at most k . Since $\text{Col}(B) \subseteq C$, Lemma 4.4 gives $\|A - B\|_F \geq \|A - P_C A\|_F$, and by (28) bounding this residual from below is the same as bounding $\text{tr}(P_C A A^\top)$ from above. The matrix $M = A A^\top$ is symmetric and positive semidefinite, and the singular value decomposition gives $A A^\top = U(\Sigma \Sigma^\top)U^\top$, a spectral decomposition whose eigenvalues are $\sigma_1^2 \geq \dots \geq \sigma_r^2 \geq 0$. Enlarging C to a k -dimensional superspace only increases $\text{tr}(P_C M)$, since M is positive semidefinite, so Lemma 4.5 applies and $\text{tr}(P_C A A^\top) \leq \sigma_1^2 + \dots + \sigma_k^2$. Using $\|A\|_F^2 = \|\Sigma\|_F^2 = \sum_{i=1}^r \sigma_i^2$ from Lemma 2.11, equation (28) now gives

$$\|A - B\|_F^2 \geq \|A\|_F^2 - \text{tr}(P_C A A^\top) \geq \sum_{i=1}^r \sigma_i^2 - \sum_{i=1}^k \sigma_i^2 = \sum_{i=k+1}^r \sigma_i^2, \quad (33)$$

and taking square roots yields (32). $\square$

Remark 4.7. The same matrix A_k is optimal in the spectral norm as well: $\|A - B\|_2 \geq \sigma_{k+1} = \|A - A_k\|_2$ for every B of rank at most k , a statement with a short proof by a dimension-counting argument. Mirsky's theorem subsumes both cases at once, establishing that A_k minimizes $\|A - B\|$ in *every* norm left unchanged by multiplication on either side by an orthogonal matrix, the *unitarily invariant* norms, of which the spectral and Frobenius norms are the two most familiar [7]. We do not prove the extension; the two norms treated above are the ones the rest of the paper uses.

4.3. A worked example and its geometry.

Example 4.8. Return to the matrix of Example 3.5,

$$A = \begin{pmatrix} 2 & 2 \\ 1 & 0 \\ 0 & 1 \end{pmatrix}, \quad (34)$$

whose decomposition produced $\sigma_1 = 3$ and $\sigma_2 = 1$ with $v_1 = \frac{1}{\sqrt{2}}(1, 1)^\top$ and $u_1 = \frac{1}{3\sqrt{2}}(4, 1, 1)^\top$. The best rank-one approximation keeps only the leading triple,

$$A_1 = \sigma_1 u_1 v_1^\top = \frac{1}{2} \begin{pmatrix} 4 \\ 1 \\ 1 \end{pmatrix} (1 \quad 1) = \begin{pmatrix} 2 & 2 \\ \frac{1}{2} & \frac{1}{2} \\ \frac{1}{2} & \frac{1}{2} \end{pmatrix}. \quad (35)$$

Subtracting,

$$A - A_1 = \begin{pmatrix} 0 & 0 \\ \frac{1}{2} & -\frac{1}{2} \\ -\frac{1}{2} & \frac{1}{2} \end{pmatrix}, \quad \|A - A_1\|_F^2 = 4 \cdot \frac{1}{4} = 1, \quad (36)$$

so $\|A - A_1\|_F = 1 = \sigma_2$, the value Theorem 4.6 predicts, since the only discarded singular value is σ_2 .

The geometry behind the theorem is clearest in three dimensions. A matrix acts on the unit sphere by stretching it into an ellipsoid whose semi-axes have lengths σ_i along the left singular vectors u_i (Remark 3.2). Truncating the decomposition removes the shortest axes, collapsing the ellipsoid onto the subspace spanned by the dominant singular directions, and the error of the approximation is the size of what was removed. Figure 3 shows this collapse.

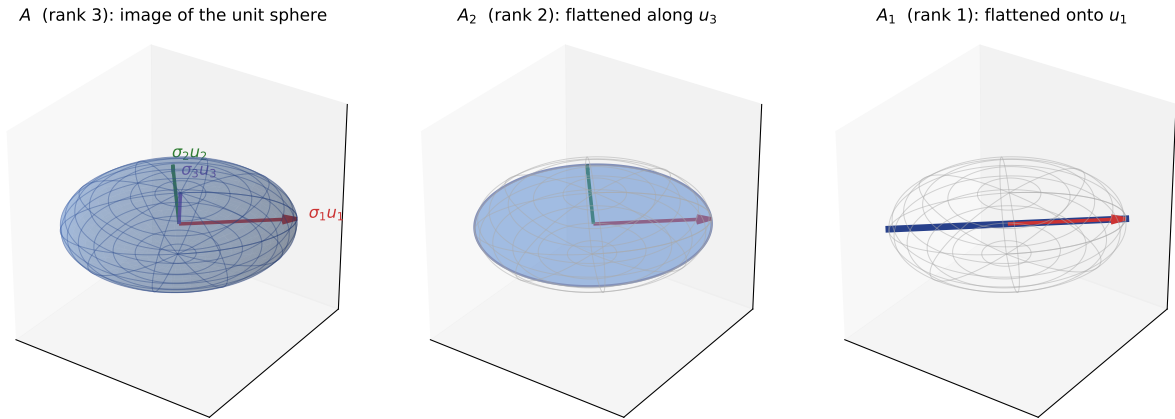


FIGURE 3. Low-rank approximation as geometric flattening. A matrix carries the unit sphere of $\mathbb{R}^3$ to an ellipsoid with semi-axes $\sigma_i u_i$ (left); the best rank-two and rank-one approximations flatten it onto $\text{span}\{u_1, u_2\}$ (center) and the line through u_1 (right), discarding the shortest axes. The error is governed by the discarded singular values, as Theorem 4.6 asserts.

Remark 4.9. The force of Theorem 4.6 lies in what it avoids. The matrices of rank at most k form neither a subspace nor a convex set, since the sum of two rank- k matrices generally has rank $2k$ (Remark 2.20). The feasible region is a curved variety, and minimizing a distance over such a set gives no a priori reason to expect a solution in closed form, still less one available without optimization. Eckart–Young says that for this objective the optimum is not merely closed form but already computed: it is a truncation of the decomposition of A . The singular values do more than describe A ; they pre-solve every low-rank approximation problem it poses, for all k simultaneously.

4.4. From the matrix to its perturbation. Theorem 4.6 concerns a matrix we possess in full. It names the best rank- k summary of A and the exact price of that summary in discarded singular values. In practice, though, A is rarely the matrix in hand. What we measure is a perturbed matrix $\tilde{A} = A + E$, corrupted by noise, and the decomposition we are able to compute is that of $\tilde{A}$ rather than of A . Whether the truncation of $\tilde{A}$ still recovers the structure of A depends on how far the perturbation can move the singular values and singular vectors to begin with. The next section takes up the first half of this question and bounds the displacement of the singular values through Weyl’s inequality. 🐥

5. PERTURBATION OF SINGULAR VALUES: WEYL’S INEQUALITY

Section 4 closed on the distance between the matrix we possess and the matrix we want. We now make that distance precise. Write

$$\tilde{A} = A + E, \quad (37)$$

where $A \in \mathbb{R}^{m \times n}$ is the matrix whose structure we seek and $E \in \mathbb{R}^{m \times n}$ is an unknown error, constrained only by a bound on its spectral norm $\|E\|_2$ and otherwise arbitrary. The Eckart–Young theorem (Theorem 4.6) approximates whatever matrix it is given, so when run on $\tilde{A}$ it returns the truncation built from the singular values and singular vectors of $\tilde{A}$ rather than those of A . Whether that truncation is a faithful stand-in for A_k depends entirely on how far the passage from A to $\tilde{A}$ moves the two ingredients of the decomposition.

Those ingredients, the singular values and the singular subspaces, can be disturbed to different degrees, and we separate the two questions. This section settles the singular values, and the answer is as strong as one could ask: every singular value of $\tilde{A}$ lies within $\|E\|_2$ of the corresponding singular value of A , so the values absorb the size of the perturbation and nothing more. The singular subspaces are the subtler matter, since directions can swing even when magnitudes hold steady, and Section 6 takes them up. The instrument behind the singular-value bound is a variational description of σ_j that exposes its dependence on A one vector at a time, and we develop it first.

5.1. A variational characterization. The singular values were defined through a global construction, as the square roots of the eigenvalues of $A^\top A$ (Section 3). For perturbation that definition is awkward, because it offers no direct way to compare $\sigma_j(A)$ with $\sigma_j(\tilde{A})$. What we want instead is a description of each singular value as the optimal value of a problem posed over vectors, so that changing A changes only the quantity being optimized. The Courant–Fischer characterization supplies exactly this.

Theorem 5.1 (Minimax characterization). *Let $A \in \mathbb{R}^{m \times n}$ have singular values $\sigma_1 \geq \dots \geq \sigma_n \geq 0$, with $\sigma_j = 0$ for $j > \text{rank}(A)$. For each $1 \leq j \leq n$,*

$$\sigma_j = \max_{\dim S=j} \min_{\substack{x \in S \\ \|x\|=1}} \|Ax\| = \min_{\dim S=n-j+1} \max_{\substack{x \in S \\ \|x\|=1}} \|Ax\|, \quad (38)$$

where S ranges over subspaces of $\mathbb{R}^n$ of the stated dimension [5].

Proof. We prove the first equality; the second follows from the same argument with the roles of the leading and trailing singular directions exchanged. Fix a singular value decomposition $A = U\Sigma V^\top$

(Theorem 3.1) and expand a vector $x \in \mathbb{R}^n$ in the right singular basis, $x = \sum_{i=1}^n c_i v_i$ with $c_i = v_i^\top x$. Then $Ax = \sum_i c_i \sigma_i u_i$, and because the left singular vectors are orthonormal,

$$\|Ax\|^2 = \sum_{i=1}^n \sigma_i^2 c_i^2, \quad \|x\|^2 = \sum_{i=1}^n c_i^2, \quad (39)$$

the terms with $\sigma_i = 0$ contributing nothing. Equation (39) reduces the geometry of A to a weighted sum, and the characterization is read off from it.

For the lower bound, take $S = \text{span}\{v_1, \dots, v_j\}$. A unit vector $x \in S$ has $c_i = 0$ for $i > j$, so (39) gives $\|Ax\|^2 = \sum_{i \leq j} \sigma_i^2 c_i^2 \geq \sigma_j^2 \sum_{i \leq j} c_i^2 = \sigma_j^2$, using $\sigma_i \geq \sigma_j$ for $i \leq j$. The minimum over unit $x \in S$ is therefore at least σ_j , attained at $x = v_j$, so the outer maximum is at least σ_j .

For the upper bound, let S be any j -dimensional subspace. Since $\dim S + \dim \text{span}\{v_j, \dots, v_n\} = j + (n - j + 1) = n + 1 > n$, the two subspaces share a nonzero vector, which we scale to a unit vector x . It has $c_i = 0$ for $i < j$, so $\|Ax\|^2 = \sum_{i \geq j} \sigma_i^2 c_i^2 \leq \sigma_j^2 \sum_{i \geq j} c_i^2 = \sigma_j^2$, using $\sigma_i \leq \sigma_j$ for $i \geq j$. Hence $\min_{x \in S} \|Ax\| \leq \sigma_j$ for every such S , and the outer maximum is at most σ_j . The two bounds coincide, giving the max-min form. $\square$

5.2. Weyl's inequality. The characterization (38) writes each singular value as a max-min of the single quantity $\|Ax\|$. Replacing A by $\tilde{A}$ changes $\|Ax\|$ by at most $\|E\|_2$ at every vector, and a quantity defined by maxima and minima cannot move by more than its integrand. This is the whole content of Weyl's inequality.

Theorem 5.2 (Weyl's inequality). *Let $A, \tilde{A} \in \mathbb{R}^{m \times n}$ have singular values $\sigma_1 \geq \dots \geq \sigma_n$ and $\tilde{\sigma}_1 \geq \dots \geq \tilde{\sigma}_n$. Then*

$$|\tilde{\sigma}_j - \sigma_j| \leq \left\| \tilde{A} - A \right\|_2 \quad \text{for every } j \quad (40)$$

[5].

Proof. Set $E = \tilde{A} - A$, so that $\|E\|_2 = \left\| \tilde{A} - A \right\|_2$. For every unit vector x ,

$$\left\| \tilde{A}x \right\| = \|Ax + Ex\| \leq \|Ax\| + \|Ex\| \leq \|Ax\| + \|E\|_2, \quad (41)$$

the first inequality from the triangle inequality and the second from the definition of the spectral norm (Definition 2.3). Fix a j -dimensional subspace S . The unit vector minimizing $\|Ax\|$ over S obeys (41), so $\min_{x \in S} \left\| \tilde{A}x \right\| \leq \min_{x \in S} \|Ax\| + \|E\|_2$. Taking the maximum over all j -dimensional S and reading each side through the max-min form of Theorem 5.1,

$$\tilde{\sigma}_j = \max_{\dim S=j} \min_{x \in S} \left\| \tilde{A}x \right\| \leq \max_{\dim S=j} \left(\min_{x \in S} \|Ax\| + \|E\|_2 \right) = \sigma_j + \|E\|_2. \quad (42)$$

The construction is symmetric in A and $\tilde{A}$, and $\left\| A - \tilde{A} \right\|_2 = \|E\|_2$, so exchanging the two matrices gives $\sigma_j \leq \tilde{\sigma}_j + \|E\|_2$. The two inequalities together are (40). $\square$

Remark 5.3. Weyl's inequality is a Lipschitz statement: each map $A \mapsto \sigma_j(A)$ is 1-Lipschitz in the spectral norm, and the constant 1 cannot be lowered. Equality is genuinely attained. Take any A with leading singular triple (σ_1, u_1, v_1) and perturb along that direction alone, $E = t u_1 v_1^\top$ with $t > 0$. Then $\tilde{A} = (\sigma_1 + t) u_1 v_1^\top + \sum_{i \geq 2} \sigma_i u_i v_i^\top$ is again a singular value decomposition, so $\tilde{\sigma}_1 = \sigma_1 + t$ while $\|E\|_2 = t$, and the first index of (40) holds with equality. No constant smaller than 1 survives this example.

Remark 5.4. The bound (40) is identical for every index, but its significance is not. Dividing by a positive σ_i , the *relative* shift satisfies

$$\frac{|\tilde{\sigma}_i - \sigma_i|}{\sigma_i} \leq \frac{\|E\|_2}{\sigma_i}, \quad (43)$$

which shrinks as σ_i grows. A perturbation of fixed size moves a large singular value by a small fraction of itself and a small one by a potentially large fraction. The dominant singular directions are thus the stable ones, which is exactly why the truncation of Section 4, retaining the largest singular values, stays trustworthy when the matrix is seen through noise: the part it keeps is the part the noise can least disturb.

5.3. A worked example.

Example 5.5. Take the matrix A of Examples 3.5 and 4.8, with $\sigma_1 = 3$ and $\sigma_2 = 1$, and perturb it along its leading singular direction by

$$E = u_1 v_1^\top = \frac{1}{6} \begin{pmatrix} 4 & 4 \\ 1 & 1 \\ 1 & 1 \end{pmatrix}, \quad \|E\|_2 = 1. \quad (44)$$

Because E adds to the leading term of the decomposition, $\tilde{A} = A + E = 4u_1 v_1^\top + u_2 v_2^\top$ is itself a singular value decomposition, with $\tilde{\sigma}_1 = 4$ and $\tilde{\sigma}_2 = 1$. The two shifts stand against the Weyl bound as follows.

i	σ_i	$\tilde{\sigma}_i$	$ \tilde{\sigma}_i - \sigma_i $	$\ E\ _2$
1	3	4	1	1
2	1	1	0	1

The bound holds at both indices, and the first row attains it. Since the perturbation lies entirely in the u_1, v_1 direction, it raises σ_1 by exactly $\|E\|_2$ and leaves σ_2 fixed, the equality case of Remark 5.3.

5.4. From values to subspaces. Weyl’s inequality closes the question of the singular values: they travel no farther than the perturbation that produced them. It is silent about the singular vectors. The values $\tilde{\sigma}_i$ may sit close to the σ_i while the directions $\tilde{u}_i$ and $\tilde{v}_i$ point well away from u_i and v_i , so a truncation that recovers the right magnitudes could still rest on the wrong subspace. Whether the singular subspaces are stable is a separate question, and its answer turns not on the size of the perturbation alone but on that size measured against the gaps between singular values. The next section takes it up through the Davis–Kahan theorem.

6. PERTURBATION OF SINGULAR SUBSPACES: THE DAVIS–KAHAN THEOREM

Section 5 left the singular vectors untreated, and before bounding their movement we must decide what movement even means. The question “how far has v_i rotated to $\tilde{v}_i$ ” presumes that v_i and $\tilde{v}_i$ are each determined by their matrices, and Proposition 3.3 showed they are not. When a singular value is repeated, any orthonormal basis of the corresponding eigenspace of $A^\top A$ serves as singular vectors, so an individual v_i is not a function of A at all. The trouble does not disappear when the singular values are merely close rather than equal: there the singular vectors are defined but ill-conditioned, and an arbitrarily small perturbation can swing them across the near-degenerate subspace while changing A almost not at all.

The resolution, anticipated in Remark 3.4, is to stop tracking vectors and start tracking the subspace a cluster of them spans. A subspace is determined by A even when no basis for it is, and two subspaces have a basis-independent distance, the principal angles of Section 2. This section bounds that distance. We first prove the governing inequality for the eigenvectors of a symmetric matrix, the setting in which it is cleanest, and then transfer it to singular subspaces through the symmetric matrix $A^\top A$, whose eigenvectors are exactly the right singular vectors of A .

6.1. Principal angles as the right measure. Recall from Definitions 2.21 and 2.23 that two d -dimensional subspaces, spanned by matrices V and $\tilde{V}$ with orthonormal columns, meet at principal angles $\theta_1, \dots, \theta_d$, and that their distance is $\|\sin \Theta(\tilde{V}, V)\|$ in the spectral or Frobenius norm. The

measure is exactly what Section 2 promised: it depends on the subspaces alone and not on the bases $V, \tilde{V}$ chosen to represent them, so it is immune to the indeterminacy that makes individual singular vectors meaningless. The one identity we need is the projection form of Remark 2.24: the numbers $\sin \theta_i$ are the singular values of $(I - VV^\top)\tilde{V}$, the projection of a basis of one subspace onto the orthogonal complement of the other, so that

$$\left\| \sin \Theta(\tilde{V}, V) \right\|_F^2 = \left\| (I - VV^\top)\tilde{V} \right\|_F^2. \quad (45)$$

Everything below is a bound on the right-hand side.

6.2. The Davis–Kahan theorem. The singular vectors of A are the eigenvectors of $A^\top A$, so we prove the perturbation bound first for the eigenvectors of a symmetric matrix and specialize afterward. To keep the singular-value matrix Σ unambiguous, symmetric matrices in this subsection are written H .

The idea of the proof is visible before any formula. An eigenvector belonging to an eigenvalue that is far from all the others is stable, because moving it would force the quadratic form $x \mapsto x^\top Hx$ to disagree with its near neighbours, and there are none. An eigenvector whose eigenvalue is crowded is fragile. The theorem turns “far” into the eigenvalue gap and makes the dependence quantitative.

Theorem 6.1 (Davis–Kahan, $\sin \Theta$ form). *Let $H, \tilde{H} \in \mathbb{R}^{p \times p}$ be symmetric, with eigenvalues $\lambda_1 \geq \dots \geq \lambda_p$ and $\tilde{\lambda}_1 \geq \dots \geq \tilde{\lambda}_p$ and orthonormal eigenvectors. Fix indices $r \leq s$, let $V = [v_r \ \dots \ v_s]$ and $\tilde{V} = [\tilde{v}_r \ \dots \ \tilde{v}_s]$ collect the corresponding eigenvectors of H and $\tilde{H}$, and set*

$$\delta = \min \left\{ \left| \tilde{\lambda}_j - \lambda_i \right| : r \leq j \leq s, i \notin [r, s] \right\}. \quad (46)$$

If $\delta > 0$, then

$$\left\| \sin \Theta(\tilde{V}, V) \right\|_F \leq \frac{\left\| \tilde{H} - H \right\|_F}{\delta} \quad (47)$$

[1].

Proof. Let $P = VV^\top$ be the orthogonal projector onto the eigenspace $\text{span}\{v_r, \dots, v_s\}$. Because this subspace is spanned by eigenvectors of H , it is H -invariant, and P commutes with H ; consequently so does $I - P$. By (45) the quantity to bound is $\left\| (I - P)\tilde{V} \right\|_F^2 = \sum_{j=r}^s \left\| (I - P)\tilde{v}_j \right\|^2$, so it suffices to control each $\left\| (I - P)\tilde{v}_j \right\|$.

Fix $j \in [r, s]$. The eigenvalue relation $\tilde{H}\tilde{v}_j = \tilde{\lambda}_j\tilde{v}_j$ lets the residual of $\tilde{v}_j$ against H be written through the perturbation alone:

$$(H - \tilde{\lambda}_j I)\tilde{v}_j = H\tilde{v}_j - \tilde{H}\tilde{v}_j = (H - \tilde{H})\tilde{v}_j. \quad (48)$$

Applying the commuting projector $I - P$ and using that it leaves $H - \tilde{\lambda}_j I$ invariant gives

$$(H - \tilde{\lambda}_j I)(I - P)\tilde{v}_j = (I - P)(H - \tilde{H})\tilde{v}_j. \quad (49)$$

Here the gap enters. On the range of $I - P$, the matrix H acts with the eigenvalues λ_i for $i \notin [r, s]$, so $H - \tilde{\lambda}_j I$ has the eigenvalues $\lambda_i - \tilde{\lambda}_j$ there, every one of absolute value at least δ by (46). A symmetric matrix whose eigenvalues on a subspace are bounded below in magnitude by δ expands every vector of that subspace by at least δ , so the left side of (49) has norm at least $\delta \left\| (I - P)\tilde{v}_j \right\|$. The right side has norm at most $\left\| (H - \tilde{H})\tilde{v}_j \right\|$, since $I - P$ is a projection. Therefore

$$\delta \left\| (I - P)\tilde{v}_j \right\| \leq \left\| (H - \tilde{H})\tilde{v}_j \right\|. \quad (50)$$

Squaring (50) and summing over $j \in [r, s]$,

$$\delta^2 \sum_{j=r}^s \|(I - P)\tilde{v}_j\|^2 \leq \sum_{j=r}^s \|(\tilde{H} - H)\tilde{v}_j\|^2 = \|(\tilde{H} - H)\tilde{V}\|_F^2 \leq \|\tilde{H} - H\|_F^2, \quad (51)$$

the last step because $\tilde{V}$ has orthonormal columns. Combining (51) with (45) gives (47). $\square$

The gap (46) measures the perturbed block eigenvalues against the unperturbed outer ones, which is natural in the proof but awkward in practice, where one knows the unperturbed spectrum and not its perturbation. Weyl's inequality removes the difficulty, and the next remark records the form of the theorem that results.

Remark 6.2. Write $\delta_0 = \min(\lambda_{r-1} - \lambda_r, \lambda_s - \lambda_{s+1})$ for the gap between the block and the rest of the *unperturbed* spectrum, with the conventions $\lambda_0 = +\infty$ and $\lambda_{p+1} = -\infty$. Weyl's inequality (Theorem 5.2, applied to the symmetric matrices H and $\tilde{H}$) gives $|\tilde{\lambda}_j - \lambda_j| \leq \|\tilde{H} - H\|_2$, so the perturbed block cannot stray far from the unperturbed one, and the separation (46) stays comparable to δ_0 . Yu, Wang, and Samworth turned this observation into a clean bound stated entirely through the unperturbed gap,

$$\|\sin \Theta(\tilde{V}, V)\|_F \leq \frac{2 \min(\sqrt{d} \|\tilde{H} - H\|_2, \|\tilde{H} - H\|_F)}{\delta_0}, \quad d = s - r + 1, \quad (52)$$

which is the version most often used in applications [9].

6.3. From eigenvectors to singular subspaces. The passage to singular subspaces costs nothing new. By Remark 2.15, the right singular vectors of $A \in \mathbb{R}^{m \times n}$ are orthonormal eigenvectors of the symmetric positive semidefinite matrix $A^\top A$, with eigenvalues σ_i^2 . Applying Theorem 6.1 to $H = A^\top A$ and $\tilde{H} = \tilde{A}^\top \tilde{A}$ therefore bounds the rotation of any block of right singular subspaces, once the perturbation $\tilde{H} - H$ is expressed through $E = \tilde{A} - A$. Expanding the square,

$$\tilde{A}^\top \tilde{A} - A^\top A = A^\top E + E^\top A + E^\top E, \quad (53)$$

and bounding each term with $\|XY\|_F \leq \|X\|_2 \|Y\|_F$ together with the transpose-invariance of the Frobenius norm gives

$$\|\tilde{A}^\top \tilde{A} - A^\top A\|_F \leq 2\sigma_1 \|E\|_F + \|E\|_2 \|E\|_F = (2\sigma_1 + \|E\|_2) \|E\|_F. \quad (54)$$

Corollary 6.3 (Davis–Kahan for singular subspaces). *Let $A, \tilde{A} \in \mathbb{R}^{m \times n}$ have right singular vectors collected, over a block of indices $r \leq s$, into V and $\tilde{V}$, and let δ be the separation (46) for the eigenvalues σ_i^2 of $A^\top A$ and $\tilde{\sigma}_i^2$ of $\tilde{A}^\top \tilde{A}$. Then*

$$\|\sin \Theta(\tilde{V}, V)\|_F \leq \frac{(2\sigma_1 + \|E\|_2) \|E\|_F}{\delta}. \quad (55)$$

The left singular subspaces obey the same bound, with $A^\top A$ replaced by $AA^\top$.

The gap in (55) is a separation of *squared* singular values, of the form $\sigma_{r-1}^2 - \sigma_r^2$, because the eigenvalues of $A^\top A$ are the σ_i^2 . A sharper account in terms of the singular values themselves, with gaps $\sigma_{r-1} - \sigma_r$, is available through a symmetric dilation of A rather than the Gram matrix; we do not pursue it here, and (55) is the bound this paper uses.

6.4. The single-vector case and a sharpness example. When a singular value σ_j is simple and isolated, the block reduces to one index and (55) bounds the rotation of a single direction. Writing the gap as $g_j = \min(\sigma_{j-1}^2 - \sigma_j^2, \sigma_j^2 - \sigma_{j+1}^2)$ and using the operator norm in (54),

$$\sin \Theta(\tilde{v}_j, v_j) \leq \frac{(2\sigma_1 + \|E\|_2) \|E\|_2}{g_j}. \quad (56)$$

Remark 6.4. The angle controls the vector when the vector is meaningful. A singular vector is fixed only up to sign, and choosing the sign of $\tilde{v}_j$ so that $\tilde{v}_j^\top v_j \geq 0$ gives $\|\tilde{v}_j - v_j\|^2 = 2(1 - \tilde{v}_j^\top v_j) \leq 2\sin^2 \Theta(\tilde{v}_j, v_j)$, since $1 - \cos \theta \leq \sin^2 \theta$ when $\cos \theta \geq 0$. Hence $\|\tilde{v}_j - v_j\| \leq \sqrt{2} \sin \Theta(\tilde{v}_j, v_j)$, so for a simple singular value the subspace bound translates back into an ordinary distance between vectors.

Example 6.5. The eigenvector bound (47) is tight up to the constant. Take the symmetric matrices

$$H = \begin{pmatrix} 1 & 0 \\ 0 & 0 \end{pmatrix}, \quad \tilde{H} = \begin{pmatrix} 1 & \varepsilon \\ \varepsilon & 0 \end{pmatrix}, \quad \varepsilon > 0, \quad (57)$$

with $\|\tilde{H} - H\|_2 = \varepsilon$ and eigenvalue gap $\lambda_1 - \lambda_2 = 1$, so the bound on the leading eigenvector reads $\sin \Theta(\tilde{v}_1, v_1) \leq \varepsilon$. The leading eigenvector of $\tilde{H}$ is proportional to $(\varepsilon, \tilde{\lambda}_1 - 1)$ with $\tilde{\lambda}_1 = \frac{1}{2}(1 + \sqrt{1 + 4\varepsilon^2})$, and a short computation gives $\sin \Theta(\tilde{v}_1, v_1) = 0.0985$ at $\varepsilon = 0.1$ and 0.0498 at $\varepsilon = 0.05$. The ratio of the true angle to the bound is 0.985 and 0.996 respectively, approaching 1 as $\varepsilon \rightarrow 0$. The rotation grows linearly in the perturbation at exactly the rate (47) predicts, so no constant below 1 can hold.

6.5. A worked example.

Example 6.6. Return once more to A of Examples 3.5, 4.8, and 5.5, with $\sigma_1 = 3$, $\sigma_2 = 1$, and leading right singular vector $v_1 = \frac{1}{\sqrt{2}}(1, 1)^\top$. The perturbation $E = u_1 v_1^\top$ of Example 5.5, which made Weyl's bound tight, rotates the singular subspace not at all: lying entirely along the leading singular direction, it leaves v_1 fixed, and $\sin \Theta(\tilde{v}_1, v_1) = 0$. The perturbation that maximally moves the singular value moves the subspace least, a clean illustration that the two questions of this part of the paper are genuinely independent.

A rotation appears only when the perturbation reaches across singular directions. Take instead

$$E' = \begin{pmatrix} 0 & \frac{1}{2} \\ 0 & 0 \\ 0 & 0 \end{pmatrix}, \quad \|E'\|_2 = \frac{1}{2}, \quad (58)$$

and set $\tilde{A} = A + E'$. Computing its singular value decomposition gives $\tilde{\sigma}_1 = 3.354$ and $\tilde{\sigma}_2 = 1.000$, and the leading right singular vector turns by $\sin \Theta(\tilde{v}_1, v_1) = 0.110$. Corollary 6.3 applies with the gap $\delta = \tilde{\sigma}_1^2 - \sigma_2^2 = 11.25 - 1 = 10.25$:

quantity	value
$\sin \Theta(\tilde{v}_1, v_1)$	0.110
$\ \tilde{A}^\top \tilde{A} - A^\top A\ _F$	2.658
gap $\delta = \tilde{\sigma}_1^2 - \sigma_2^2$	10.25
Davis–Kahan bound (55)	0.259

The bound holds with room to spare, exceeding the true rotation by a factor of about two. The slack traces to the large gap: with σ_1 and σ_2 far apart in square, the leading subspace is strongly anchored, and the Frobenius numerator charges the bound for all of the perturbation's energy rather than the part that actually tilts v_1 . A matrix with two close singular values would show the same bound much nearer to equality.

6.6. Toward entrywise localization. Weyl's inequality and the Davis–Kahan theorem together give a first-order account of the singular value decomposition under perturbation: the singular values move by at most the norm of the perturbation, and the singular subspaces rotate by at most that norm divided by a spectral gap. Both bounds are stated in the norm of $\tilde{A} - A$, and both therefore presume that this norm, or the spectral gap it is weighed against, can be estimated. Neither is always at hand, since a spectral norm is itself the solution of an optimization and a gap requires knowing the spectrum. The final tool of the paper asks for less. Gershgorin's circle theorem confines the eigenvalues of a matrix to discs whose centers and radii are read directly off the entries, with no norm to compute and no spectrum to know in advance, and we take it up next.

7. GERSHGORIN'S CIRCLE THEOREM AS A SUPPORTING TOOL

Section 6 closed on a practical gap. Weyl's inequality and the Davis–Kahan theorem are sharp, but each is stated in a norm of the perturbation, a single number distilled from the whole matrix and not, in general, available without work of its own. It is fair to ask for something cruder and cheaper: a statement about where the eigenvalues of a matrix lie that uses nothing but the entries, read off without solving any optimization. Gershgorin's theorem, which predates the perturbation results of this paper by decades, supplies exactly that, and for all its simplicity it remains a serviceable diagnostic, both as a quick enclosure of a spectrum and, as we will see, as a way to certify the hypotheses the earlier sections require.

7.1. The discs and the theorem. Throughout, $A = (a_{ij})$ is square. We state the theorem for $A \in \mathbb{C}^{n \times n}$, since eigenvalues are complex in general, and specialize to the real symmetric case when we return to the singular value decomposition. The construction attaches to each row a disc in the complex plane, centered at that row's diagonal entry and with radius the combined size of its off-diagonal entries.

Definition 7.1. For $A = (a_{ij}) \in \mathbb{C}^{n \times n}$, the i th *Gershgorin disc* is

$$D_i = \{ z \in \mathbb{C} : |z - a_{ii}| \leq R_i \}, \quad R_i = \sum_{j \neq i} |a_{ij}|, \quad (59)$$

the closed disc centered at the diagonal entry a_{ii} whose radius R_i is the absolute row sum off the diagonal.

A matrix with small off-diagonal entries is nearly diagonal, and its eigenvalues should sit near the diagonal entries; the radius R_i measures the departure from diagonality in row i and bounds how far an eigenvalue can drift from a_{ii} .

Theorem 7.2 (Gershgorin). *Every eigenvalue of $A \in \mathbb{C}^{n \times n}$ lies in the union $\bigcup_{i=1}^n D_i$ of its Gershgorin discs [?].*

Proof. Let λ be an eigenvalue with eigenvector $x \neq 0$, scaled so its largest coordinate has modulus one, $|x_k| = 1 \geq |x_j|$ for all j . This choice isolates the coordinate on which the eigenvalue equation bears hardest. The k th row of $Ax = \lambda x$ reads $\sum_j a_{kj}x_j = \lambda x_k$, and moving the diagonal term across gives

$$(\lambda - a_{kk})x_k = \sum_{j \neq k} a_{kj}x_j. \quad (60)$$

Taking absolute values, the left side is $|\lambda - a_{kk}|$ because $|x_k| = 1$, and the triangle inequality with $|x_j| \leq 1$ bounds the right,

$$|\lambda - a_{kk}| \leq \sum_{j \neq k} |a_{kj}| |x_j| \leq \sum_{j \neq k} |a_{kj}| = R_k. \quad (61)$$

Hence $\lambda \in D_k$, and in particular λ lies in the union. $\square$

Remark 7.3. Since A and $A^\top$ share the same eigenvalues, applying Theorem 7.2 to $A^\top$ places every eigenvalue in the union of the discs built from absolute *column* sums as well. An eigenvalue must lie in both unions, so intersecting the row and column families can sharpen the localization.

Corollary 7.4 (Isolated discs). *If a union of k Gershgorin discs is disjoint from the remaining $n - k$, then it contains exactly k eigenvalues of A , counted with multiplicity.*

The reason is continuity. Scale the off-diagonal entries of A by a parameter running from 1 down to 0, deforming A to its diagonal; the discs shrink to their centers a_{ii} , and the eigenvalues, which vary continuously along the way, end there. The discs of the isolated cluster never touch the others during the deformation, so an eigenvalue cannot pass between the groups without crossing the empty gap separating them, and the count within the cluster is preserved. The continuity of eigenvalues that makes this precise is standard, and we refer to [5] for it.

7.2. The supporting role in the singular value decomposition. The theorem earns its place here through two services to the rest of the paper, both purely entrywise.

The first is direct localization of the singular values. They are the square roots of the eigenvalues of the symmetric matrix $A^\top A$ (Remark 2.15), so the Gershgorin discs of $A^\top A$, formed from its entries alone, bound the σ_i^2 with no decomposition computed. For the matrix A of Example 3.5, where $A^\top A = \begin{pmatrix} 5 & 4 \\ 4 & 5 \end{pmatrix}$ by (20), both rows give the same disc, centered at 5 with radius 4, so every eigenvalue of $A^\top A$ lies in $[1, 9]$ and therefore $\sigma_i \in [1, 3]$. The true singular values $\sigma_1 = 3$ and $\sigma_2 = 1$ sit exactly at the endpoints $\sqrt{9}$ and $\sqrt{1}$. For this matrix the enclosure is not merely valid but tight, a consequence of $A^\top A$ being a symmetric 2×2 matrix with equal diagonal entries, whose eigenvalues are $a_{11} \pm |a_{12}|$. Figure 4 shows the disc with the two eigenvalues on its boundary.

The second service is a certificate, computed in advance, that the gap the Davis–Kahan bound requires actually exists. Corollary 6.3 controls the rotation of a singular subspace only when the relevant eigenvalues of $A^\top A$ stand apart from the rest by a positive gap δ , and the isolated-disc refinement turns disjointness of discs into exactly such a guarantee. If the discs surrounding $\sigma_r^2, \dots, \sigma_s^2$ are disjoint from the others, those squared singular values occupy a band separated from the remaining spectrum, and δ is at least the distance between the disc clusters, all without computing an eigenvalue. For the matrix above the certificate is vacuous, since its two discs coincide and cannot be separated; a matrix with more spread-out diagonal entries would have disjoint discs and a gap legible directly from (59). The structural point stands regardless: Gershgorin provides an entrywise precondition under which the perturbation theory of Section 6 is known to apply.

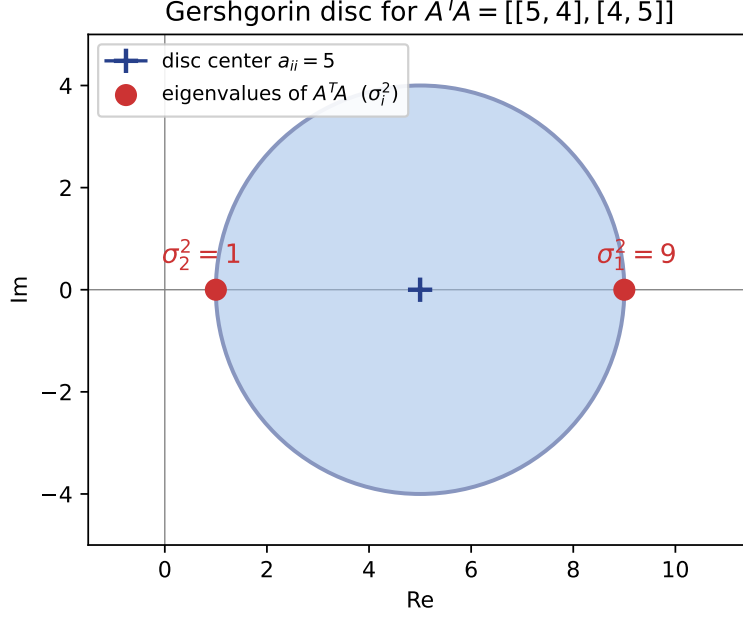


FIGURE 4. Gershgorin localization of the singular values of the matrix A of Example 3.5, read from $A^T A = \begin{pmatrix} 5 & 4 \\ 4 & 5 \end{pmatrix}$ without any decomposition. Both rows of $A^T A$ give the same disc (shaded), centered at the diagonal entry 5 (cross) with radius the off-diagonal row sum 4. The eigenvalues $\sigma_1^2 = 9$ and $\sigma_2^2 = 1$ (dots) lie in the disc, here exactly on its boundary, certifying $\sigma_i \in [1, 3]$.

7.3. Closing the toolkit. The instruments are now assembled. The singular value decomposition organizes a matrix into magnitudes and directions (Sections 3 and 4), Weyl’s inequality and the Davis–Kahan theorem bound how far each moves under perturbation (Sections 5 and 6), and Gershgorin’s theorem localizes the spectrum from the entries alone. The final section turns from the machinery to its uses, sketching how these results underwrite principal component analysis, covariance estimation, and low-rank recovery from noisy data, and setting the theorems in their historical place.

8. APPLICATIONS

The theorems of this paper are not confined to matrix analysis; they underwrite methods across the data sciences, and a single example from each of four fields shows how directly. Each rests on a result proved above.

Machine learning. Principal component analysis and dimensionality reduction summarize a data matrix by its leading singular directions and discard the rest. Weyl’s inequality (Theorem 5.2) is what makes the summary trustworthy: the large singular values, and with them the dominant components, move by at most $\|E\|_2$ under the sampling noise through which the matrix is always observed, so the retained directions are exactly the stable ones (Remark 5.4).

Artificial intelligence. Low-rank adaptation, now a standard method for fine-tuning large language models, freezes a pretrained weight matrix and learns only a low-rank update to it [6]. Its premise is the Eckart–Young theorem (Theorem 4.6): a matrix of rank k captures the dominant part of a change with a small fraction of the parameters, and the truncated decomposition is the optimal such approximation, so the update need never be full rank to be nearly as expressive.

Data science. Covariance estimation is the perturbation model itself. A sample covariance is the population covariance plus a finite-sample error, $\tilde{A} = A + E$, and the principal components

one extracts from it are its singular subspaces. Weyl’s inequality and the Davis–Kahan theorem (Theorems 5.2 and 6.1) bound how far the estimated eigenvalues and eigenvectors can lie from the population truth, and so make precise when an empirical covariance can be trusted [3].

Computer science. Spectral clustering and community detection recover a partition of a graph from the leading eigenvectors of its adjacency or Laplacian matrix. The clusters are encoded in a singular subspace, and the Davis–Kahan theorem (Corollary 6.3) bounds how far that subspace rotates under the noise separating the observed graph from its idealized block structure; a sufficient spectral gap therefore guarantees that the recovered partition is close to the correct one [9].

Across all four the governing principle is the one this paper has made precise. The recoverable signal is the large singular values and the well-separated subspaces, and everything else is noise that the perturbation bounds mark as untrustworthy.

REFERENCES

1. Chandler Davis and William M. Kahan, *The rotation of eigenvectors by a perturbation. III*, SIAM Journal on Numerical Analysis **7** (1970), no. 1, 1–46.
2. Carl Eckart and Gale Young, *The approximation of one matrix by another of lower rank*, Psychometrika **1** (1936), no. 3, 211–218.
3. Jianqing Fan, Weichen Wang, and Yiqiao Zhong, *An ℓ_∞ eigenvector perturbation bound and its application to robust covariance estimation*, Journal of Machine Learning Research **18** (2018), no. 207, 1–42, arXiv:1603.03516.
4. Gene H. Golub and Charles F. Van Loan, *Matrix computations*, 4 ed., Johns Hopkins University Press, Baltimore, MD, 2013.
5. Roger A. Horn and Charles R. Johnson, *Matrix analysis*, 2 ed., Cambridge University Press, Cambridge, 2013.
6. Edward J. Hu, Yelong Shen, Phillip Wallis, Zeyuan Allen-Zhu, Yanzhi Li, Shean Wang, Lu Wang, and Weizhu Chen, *LoRA: Low-rank adaptation of large language models*, 2022.
7. Leon Mirsky, *Symmetric gauge functions and unitarily invariant norms*, The Quarterly Journal of Mathematics **11** (1960), no. 1, 50–59.
8. G. W. Stewart and Ji-guang Sun, *Matrix perturbation theory*, Academic Press, Boston, 1990.
9. Yi Yu, Tengyao Wang, and Richard J. Samworth, *A useful variant of the Davis–Kahan theorem for statisticians*, Biometrika **102** (2015), no. 2, 315–323.

EULER CIRCLE