

Lecture Notes on Random Matrix Theory in Data Science and Statistics

Dmitriy (Tim) Kunisky

Fall 2026 (Last Updated: September 22, 2026)

Contents

1	Random Vector Theory	2
1.1	Independent Entries and Subgaussian Concentration	2
1.2	Gaussian Vectors and Subexponential Concentration	5
1.3	Orthogonal Invariance	8
1.4	Exercises	10
1.4.1	High-Dimensional Geometry	10
1.4.2	Concentration Inequalities	11
1.4.3	Orthogonal Invariance for Vectors	13
2	Geometry of Rectangular Random Matrices	14
2.1	Gaussian Random Field Viewpoint	14
2.2	The Geometric Method of Random Matrix Theory	17
2.3	Singular Values	18
2.3.1	Applying the Geometric Method	18
2.3.2	Nets, Coverings, and Packings	19
2.3.3	Coarse Non-Asymptotic Bound on Singular Values	21
2.4	Singular Vectors	24
2.4.1	Distinctness of Singular Values	25
2.4.2	Orthogonal Invariance for Matrices and Haar Measures	27
2.4.3	Connection to Random Projections	30
2.5	Application: Compressed Sensing	31
2.5.1	Null Space and Restricted Isometry Properties	31
2.5.2	Random Sensing Matrices	32
2.6	Application: Neural Network Initialization	34
2.6.1	General Setup	34
2.6.2	Initialization and Local Analysis	35
2.6.3	Jacobian Calculations and Kaiming Initialization	36
	Bibliography	38

1 | RANDOM VECTOR THEORY

Before we continue to random matrix theory, it will be useful to establish some preliminaries about random *vectors*, which will also lead us to a few important intuitions about high-dimensional probability and geometry.

As in the case of matrices, vectors have at least two natural interpretations: they are a **list of numbers**, as well as a geometric object with a **magnitude and direction** (related notions about vector and Banach spaces are abstractions of this latter viewpoint) with respect to a basis of the space to which they belong. Each interpretation gives rise to a natural model¹ of random vector. These are random vectors with *independent and identically distributed (i.i.d.) entries* and random vectors *uniformly random on the sphere*, respectively. In this chapter, we will look at these models and their behaviors, and alongside them introduce some of the basic tools and intuitions for probability and geometry in high dimension.

1.1 INDEPENDENT ENTRIES AND SUBGAUSSIAN CONCENTRATION

In the list-of-numbers interpretation, it is most logical to model a random vector where the entries are as statistically decoupled as possible. Thus, we consider

$$\mathbf{x} \sim \mu^{\otimes d}, \tag{1.1.1}$$

a vector with entries i.i.d. from some probability measure μ on $\mathbb{R}$. The symbol $\otimes$ here and throughout will denote forming the product measure, i.e., the above notation means $x_i \sim \mu$ are drawn i.i.d. The only exception is when it is applied to vectors or matrices, in which case it denotes the Kronecker or tensor product; which one we mean should always be clear from context.

The behavior of such vectors is closely linked to *high-dimensional geometry*. Consider, for example, the case $\mu = \text{Unif}([-1, 1])$. In this case, $\mathbf{x}$ is drawn uniformly at random from the solid box $[-1, 1]^d \subset \mathbb{R}^d$, obeying the measure $\text{Unif}([-1, 1]^d)$, and so properties that $\mathbf{x}$ has with high probability can equivalently be viewed as properties that “most” vectors in this box have, with respect to Lebesgue measure of volume.

We will see that this measure behaves rather paradoxically in high dimension. If you imagine a picture of the measures $\text{Unif}([-1, 1]^d)$ for $d = 1, 2, 3$ as densities, they look like solid convex “blobs” of mass. But, this intuition breaks down in high dimension: for large d , the mass of the box is instead “fragmented” and “hollow”, and in fact this follows from basic statements in scalar probability theory that you have likely encountered already.

¹We will sometimes use this word “model” instead of “law” or “distribution” to emphasize that such a choice includes some judgement of what distributional assumption is natural or reasonable—we make such a choice to try to *model* typical situations we might encounter in applications.

Consider the radius at which the random point $\mathbf{x}$ lies,

$$R^2 := \|\mathbf{x}\|^2 = \sum_{i=1}^d x_i^2 =: \sum_{i=1}^d Y_i.$$

This is a sum of i.i.d. random variables $Y_i = x_i^2$, precisely the situation to which the law of large numbers (LLN), central limit theorem (CLT), large deviations principles (LDP), and other tools of scalar probability apply. For instance, the law of large numbers gives

$$\frac{1}{d}R^2 = \frac{1}{d} \sum_{i=1}^d Y_i \xrightarrow[\mathbb{P}]{d \rightarrow \infty} \mathbb{E}Y_1 = \frac{1}{3},$$

the last equality following by a simple integral calculation. Thus already we see that, in high dimension, $\mathbf{x}$ usually lies close to a sphere of radius $\frac{1}{\sqrt{3}}\sqrt{d} \approx 0.577\sqrt{d}$.

The central limit theorem gives more detail, stating that

$$\frac{1}{\sqrt{d}} \left(R^2 - \frac{d}{3} \right) = \frac{1}{\sqrt{d}} \sum_{i=1}^d (Y_i - \mathbb{E}Y_i) \xrightarrow[\text{(law)}]{d \rightarrow \infty} \mathcal{N}(0, \text{Var}[Y_1]) = \mathcal{N}\left(0, \frac{4}{45}\right),$$

where $\frac{4}{45} = \frac{1}{5} - (\frac{1}{3})^2$ comes from similar integral calculations to before. We will not care that much about the detailed constants here; the important part is that the fluctuations of R^2 around $\frac{1}{3}d$ are of order $O(\sqrt{d})$. Thus heuristically we have $R^2 = \frac{1}{3}d + O(\sqrt{d})$, and so $R = \frac{1}{\sqrt{3}}\sqrt{d} + O(1)$.²

Most important for our purposes will be *non-asymptotic* concentration inequalities, which give concrete bounds on probabilities of deviation from the mean as a function of d . We develop some of the basic tools for working with such inequalities in this section, the next section, and the chapter exercises.

Our basic goal is to show bounds of the form

$$\mathbb{P} \left[\left| R^2 - \frac{d}{3} \right| \geq t \right] \leq \delta(t, d)$$

for some function $\delta(t, d)$ as small as possible. You could try using any inequalities you are already aware of for this purpose. For example, Markov's inequality gives

$$\mathbb{P} \left[R^2 \geq \frac{d}{3} + t \right] \leq \frac{\frac{d}{3}}{\frac{d}{3} + t} = \frac{1}{1 + \frac{3t}{d}}, \quad (1.1.2)$$

which only tells us that $R^2 = O(d)$ with high probability (say, $\mathbb{P}[R^2 \leq 100d] \geq 0.99$ and similar statements); the inequality “kicks in” and gives useful information only when $t \gtrsim d$. Chebyshev's inequality gives

$$\mathbb{P} \left[R^2 \geq \frac{d}{3} + t \right] = \mathbb{P}[R^2 \geq \mathbb{E}R^2 + t] \leq \frac{\text{Var}[R^2]}{t^2} = \frac{\frac{4}{45}d}{t^2}, \quad (1.1.3)$$

where we calculate $\text{Var}[R^2] = d \text{Var}[Y_1] = \frac{4}{45}d$. This gives the correct scaling, kicking in when $t \gtrsim \sqrt{d}$. However, as the following result shows, the actual dependence on t above is far from optimal.

Sharper such inequalities require some more tools. We say a random scalar $Y \in \mathbb{R}$ is *centered* if $\mathbb{E}Y = 0$ and similarly for random vectors.

²To see this, write $R = \frac{1}{\sqrt{3}}\sqrt{d} \cdot \sqrt{1 + \frac{3}{d}(R^2 - \frac{1}{3}d)}$ and take a Taylor expansion of the square root.

Definition 1.1.1 (Subgaussianity). A centered random scalar $Y \in \mathbb{R}$ is σ^2 -subgaussian if $\mathbb{E} \exp(\lambda Y) \leq \exp(\frac{\sigma^2}{2} \lambda^2)$ for all $\lambda \in \mathbb{R}$.

The function

$$\phi_Y(\lambda) := \mathbb{E} \exp(\lambda Y)$$

is called the *moment generating function* of Y . The terminology “subgaussian” is justified by the following.

Proposition 1.1.2. If $Y \sim \mathcal{N}(0, \sigma^2)$, then for all $\lambda \in \mathbb{R}$ we have $\phi_Y(\lambda) = \exp(\frac{\sigma^2}{2} \lambda^2)$.

Exercise 1.4.5 will develop some further tools for working with subgaussianity.

The following is a form of *Hoeffding’s inequality* giving concentration for sums of subgaussian random variables.

Lemma 1.1.3 (Subgaussian sum concentration). Suppose that $Y_1, \dots, Y_d$ are independent and each is centered and σ^2 -subgaussian. Then, for all $t \geq 0$,

$$\mathbb{P} \left[\left| \sum_{i=1}^d Y_i \right| \geq t \right] \leq 2 \exp \left(-\frac{1}{2\sigma^2} \cdot \frac{t^2}{d} \right).$$

Proof. We use a general proof technique that you can call the “Chernoff method” or the “nonlinear Markov inequality” that is behind the proofs of all Chebyshev-, Hoeffding-, and Chernoff-like inequalities. The idea is to “preprocess” a random variable with a nonlinearity before using Markov’s inequality. In our case, for any $\lambda > 0$ we have

$$\begin{aligned} \mathbb{P} \left[\sum_{i=1}^d Y_i \geq t \right] &= \mathbb{P} \left[\exp \left(\lambda \sum_{i=1}^d Y_i \right) \geq \exp(\lambda t) \right] \\ &\leq \frac{\mathbb{E}[\exp(\lambda \sum_{i=1}^d Y_i)]}{\exp(\lambda t)} && \text{(Markov's inequality)} \\ &= \frac{\prod_{i=1}^d \mathbb{E}[\exp(\lambda Y_i)]}{\exp(\lambda t)} && \text{(independence)} \\ &\leq \exp \left(\frac{\sigma^2 d}{2} \lambda^2 - t \lambda \right). && \text{(subgaussianity)} \end{aligned}$$

Optimizing over λ , we choose $\lambda = \frac{t}{\sigma^2 d}$, and obtain

$$\mathbb{P} \left[\sum_{i=1}^d Y_i \geq t \right] \leq \exp \left(-\frac{1}{2\sigma^2} \cdot \frac{t^2}{d} \right),$$

and applying the same argument to the lower tail and using the union bound gives the result. $\square$

In particular, if σ^2 is a constant as d grows, then this probability indeed becomes small once $t \gg \sqrt{d}$, as our softer argument with the central limit theorem above predicted. As a corollary, we find the following concrete concentration inequality about the uniform measure over the box.

Corollary 1.1.4 (Typical set of high-dimensional cube). Let $\mathbf{x} \sim \text{Unif}([-1, 1]^d)$ and $R^2 = \|\mathbf{x}\|^2$. Then, for all $t \geq 0$,

$$\mathbb{P} \left[\left| R^2 - \frac{d}{3} \right| \geq t \right] \leq 2 \exp \left(-\frac{2t^2}{d} \right).$$

This follows from Lemma 1.1.3 together with the following, which follows from a calculus argument that you can find in [RH17, Lemma 1.8].

Lemma 1.1.5 (Hoeffding's lemma). Suppose that $Y \in \mathbb{R}$ is centered with $Y \in [A, B]$ almost surely. Then, Y is σ^2 -subgaussian for $\sigma^2 = (\frac{B-A}{2})^2$.

Above we use this on $Y_i = x_i^2 - \frac{1}{3}$, so we can take $[A, B] = [-\frac{1}{3}, \frac{2}{3}]$, giving the specific constants above.

Thus, we see that most of the volume of the box $[-1, 1]^d$ lies around a thin spherical *shell* of thickness $O(1)$ around the sphere of radius $\frac{1}{\sqrt{3}}\sqrt{d}$. In particular, you should not think of this typical set³ as being convex anymore; further, you should think of it as being rather close to the corners of the cube, which lie at distance $\sqrt{d}$ from the origin, very far from the inscribed solid unit ball $\mathbb{B}^d := \{\mathbf{x} \in \mathbb{R}^d : \|\mathbf{x}\| \leq 1\}$. We give a schematic picture of this in two dimension in Figure 1.1.

Another way to phrase this phenomenon is that, for $d > 3$, the volume of $\mathbb{B}^d$ relative to its circumscribing box $[-1, +1]^d$ is bounded by

$$\begin{aligned} \frac{\text{vol}(\mathbb{B}^d)}{\text{vol}([-1, +1]^d)} &= \mathbb{P}_{\mathbf{x} \sim \text{Unif}([-1, +1]^d)} [\|\mathbf{x}\| \leq 1] \\ &\leq \mathbb{P}_{\mathbf{x} \sim \text{Unif}([-1, +1]^d)} [|\|\mathbf{x}\|^2 - d/3| \geq d/3 - 1] \\ &= \exp(-\Omega(d)). \end{aligned} \tag{1.1.4}$$

Thus, these same probabilistic tools show that the volume of the unit ball is exponentially smaller than the volume of its circumscribing box.⁴

1.2 GAUSSIAN VECTORS AND SUBEXPONENTIAL CONCENTRATION

The above model is useful for getting a sense of high-dimensional geometry, but the most important kind of random vector is the following, that you have likely seen before.

Definition 1.2.1 (Gaussian random vector). $\mathcal{N}(\boldsymbol{\mu}, \boldsymbol{\Sigma})$ for parameters $\boldsymbol{\mu} \in \mathbb{R}^d$ and $\boldsymbol{\Sigma} \in \mathbb{R}_{\geq 0}^{d \times d}$ is the probability measure on $\mathbb{R}^d$ with density

$$\frac{1}{\sqrt{\det^*(2\pi\boldsymbol{\Sigma})}} \exp \left(-\frac{1}{2} (\mathbf{x} - \boldsymbol{\mu})^\top \boldsymbol{\Sigma}^+ (\mathbf{x} - \boldsymbol{\mu}) \right) \tag{1.2.1}$$

relative to the Lebesgue measure on the row space of $\boldsymbol{\Sigma}$ translated by $\boldsymbol{\mu}$. Here, $\boldsymbol{\Sigma}^+$ is the

³Let us not try to define a precise notion of typical set; think of it as meaning the smallest set in which a random variable falls with high probability.

⁴Actually, as you may look up, the true volume is even smaller than that, but to show that requires a more careful calculation.

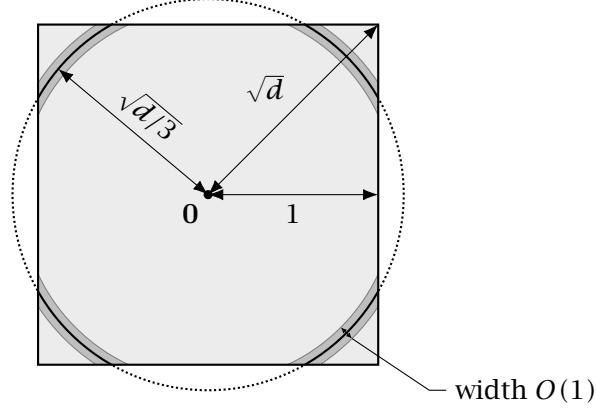


Figure 1.1: Schematic picture of the cube $[-1, 1]^d$ and the thin spherical shell giving the typical set of its uniform measure, per Corollary 1.1.4. Distances are labeled for large dimension d and are not drawn to scale.

Moore-Penrose pseudoinverse, and $\det^*$ is the product of all non-zero eigenvalues. When Σ is invertible (the most typical situation for us), then $\det^*$ above is the ordinary determinant, $\Sigma^+ = \Sigma^{-1}$, and the above density is with respect to the Lebesgue measure on all of $\mathbb{R}^d$. We call x a *Gaussian random vector* if $x \sim \mathcal{N}(\mu, \Sigma)$ for some μ and Σ , and a *standard Gaussian random vector* if $\mu = 0$ and $\Sigma = I_d$.

In particular, we will spend lots of time working with the standard Gaussian random vector, which is equivalently $\mathcal{N}(0, I_d) = \mathcal{N}(0, 1)^{\otimes d}$.

The following are the two fundamental properties of Gaussian random vectors.

Proposition 1.2.2 (Sufficiency of two moments). The parameters μ and Σ are the mean vector and covariance matrix of the random vector $x \sim \mathcal{N}(\mu, \Sigma)$:

$$\mu = \mathbb{E}x, \quad (1.2.2)$$

$$\Sigma = \mathbb{E}(x - \mu)(x - \mu)^\top = \mathbb{E}xx^\top - \mu\mu^\top \quad (1.2.3)$$

Thus, the law of a Gaussian random vector is determined by its mean and covariance, or equivalently by its linear and quadratic moments, $\mathbb{E}x$ and $\mathbb{E}xx^\top$.

Proposition 1.2.3 (Affine closure). If $x \in \mathbb{R}^d$ is a Gaussian random vector and $A \in \mathbb{R}^{d' \times d}$ and $b \in \mathbb{R}^{d'}$, then $Ax + b$ is also a Gaussian random vector. If $x \sim \mathcal{N}(\mu, \Sigma)$, then the corresponding parameters after this linear transformation are

$$\text{Law}(Ax + b) = \mathcal{N}(A\mu + b, A\Sigma A^\top). \quad (1.2.4)$$

These two facts taken together make it easy to do certain computations with Gaussian random vectors: provided we only make affine transformations, we both remain within the class of Gaussian

random vectors and can keep track of all distributions involved merely by linear algebra.⁵

We will see in the next section the key symmetry that makes the standard Gaussian random vector so fundamental. For now, let us take it for granted that this is an important distribution to understand and try to replay the concentration argument from before. The following calculation shows that something will go wrong:

Proposition 1.2.4. Suppose $X \sim \mathcal{N}(0, 1)$ and $Y = X^2$. Then,

$$\phi_Y(\lambda) = \mathbb{E} \exp(\lambda Y) = \begin{cases} (1 - 2\lambda)^{-1/2} & \text{if } \lambda < 1/2, \\ +\infty & \end{cases}.$$

In particular, Y is not σ^2 -subgaussian for any choice of $\sigma^2 > 0$.

We can still repeat the argument of Lemma 1.1.3, but we will have to be more careful to choose a value of λ where we can control the above expression. This situation, which happens often with sums of heavier-tailed random variables, motivates the following definition and argument. We use a two-parameter version of [RH17, Definition 1.11].

Definition 1.2.5 (Subexponentiality). Let $\sigma^2 > 0$ and $K > 0$. A centered random scalar $Y \in \mathbb{R}$ is (σ^2, K) -subexponential if

$$\mathbb{E}[\exp(\lambda Y)] \leq \exp\left(\frac{\sigma^2}{2} \lambda^2\right) \text{ for all } |\lambda| \leq K.$$

One way to think of this is as a form of *local subgaussianity*: σ^2 is the variance proxy, and K is the radius of the interval $[-K, K]$ on which the subgaussian bound holds. The corresponding bound for sums is a form of *Bernstein's inequality* (compare with [RH17, Theorem 1.13]).

Lemma 1.2.6 (Subexponential sum concentration). Suppose that $Y_1, \dots, Y_d$ are independent and each is (σ^2, K) -subexponential. Then, for all $t \geq 0$,

$$\mathbb{P}\left[\left|\sum_{i=1}^d Y_i\right| \geq t\right] \leq 2 \exp\left(-\min\left\{\frac{1}{2\sigma^2} \cdot \frac{t^2}{d}, \frac{K}{2} \cdot t\right\}\right).$$

Proof. The same argument as in Lemma 1.1.3 gives, for $0 < \lambda \leq K$,

$$\mathbb{P}\left[\sum_{i=1}^d Y_i \geq t\right] \leq \exp\left(\frac{d\sigma^2}{2} \lambda^2 - t\lambda\right).$$

Now, however, we cannot choose $\lambda > K$, so for $t > 0$ we instead optimize by taking

$$\lambda := \min\left\{\frac{t}{d\sigma^2}, K\right\}.$$

For this choice, $d\sigma^2\lambda \leq t$, so the exponent is at most $-t\lambda/2$. Applying the same argument to the lower tail and using the union bound gives the result. The case $t = 0$ is immediate. $\square$

⁵We will not use it in this class, but another important fact to be aware of is that the more complicated operation of *conditioning* a Gaussian random vector $\mathbf{x}$ on the values of some linear forms $\langle \mathbf{a}_i, \mathbf{x} \rangle$ also yields another Gaussian random vector.

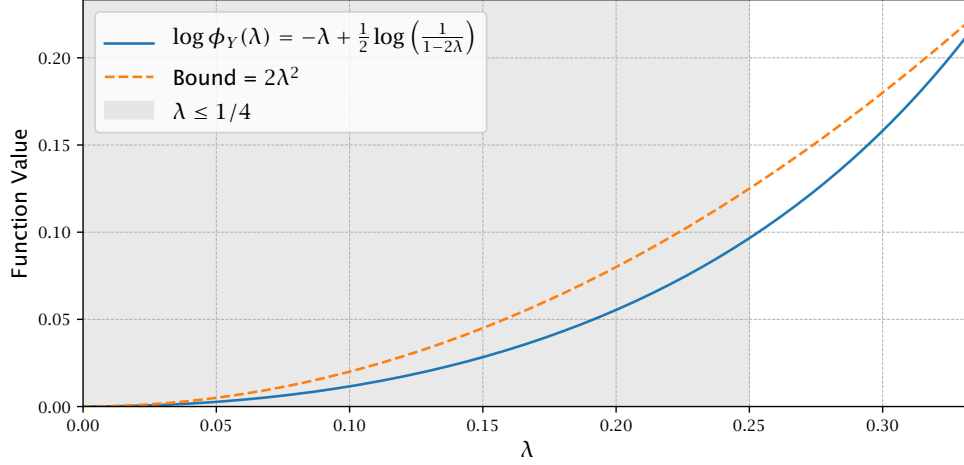


Figure 1.2: The cumulant generating function of the random variable Y appearing in Proposition 1.2.8, our bound on it, and the region where we apply this bound.

Thus the bound is Gaussian when $t \leq d\sigma^2 K$ and exponential when $t \geq d\sigma^2 K$.

The following is the corollary of this general result for specifically standard Gaussian vectors, which we will use several times.

Corollary 1.2.7. For $\mathbf{x} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_d)$ and any $t \geq 0$,

$$\begin{aligned} \mathbb{P}\left[\left|\|\mathbf{x}\|^2 - d\right| \geq t\right] &= \mathbb{P}\left[\left|\sum_{i=1}^d x_i^2 - d\right| \geq t\right] \\ &\leq 2 \begin{cases} \exp\left(-\frac{t^2}{8d}\right) & \text{if } t \leq d, \\ \exp\left(-\frac{t}{8}\right) & \text{if } t \geq d \end{cases} \end{aligned} \quad (1.2.5)$$

$$= 2 \exp\left(-\frac{1}{8} \min\left\{\frac{t^2}{d}, t\right\}\right). \quad (1.2.6)$$

This follows from Lemma 1.2.6 together with the following bound on the result of Proposition 1.2.4.

Proposition 1.2.8. For all $\lambda \leq \frac{1}{4}$, $\exp(-\lambda)(1 - 2\lambda)^{-1/2} \leq \exp(2\lambda^2)$. Thus, if $X \sim \mathcal{N}(0, 1)$, then $Y = X^2 - 1$ is $(4, \frac{1}{4})$ -subexponential.

We illustrate this estimate in Figure 1.2.

1.3 ORTHOGONAL INVARIANCE

Aside from independence of entries, the following is the other simple symmetry property a random vector can have.

Definition 1.3.1 (Orthogonal invariance). We say that a random vector $\mathbf{x} \in \mathbb{R}^d$ (or its law $\text{Law}(\mathbf{x})$) is *orthogonally invariant* if, for each (deterministic) orthogonal matrix $\mathbf{Q} \in \mathcal{O}(d)$, we have $\text{Law}(\mathbf{Q}\mathbf{x}) = \text{Law}(\mathbf{x})$.

Without going into more technical details of actually defining this probability measure explicitly, the following result states that there is a well-defined uniform measure on the unit sphere $\mathbb{S}^{d-1} := \{\mathbf{x} \in \mathbb{R}^d : \|\mathbf{x}\| = 1\}$ (so denoted because it is a $(d - 1)$ -dimensional manifold).

Proposition 1.3.2. There exists a unique probability measure supported on $\mathbb{S}^{d-1}$ that is orthogonally invariant, which we denote $\text{Unif}(\mathbb{S}^{d-1})$.

The following also shows that such distributions are nothing more than vectors distributed as $\text{Unif}(\mathbb{S}^{d-1})$ rescaled by some independent random scalar, and thus are precisely the more general class of models natural to the magnitude-and-direction interpretation that we indicated above.

Proposition 1.3.3. Suppose that $\mathbf{x}$ is orthogonally invariant and $\mathbf{x} \neq \mathbf{0}$ almost surely. Then, $\|\mathbf{x}\|$ is independent of $\mathbf{x}/\|\mathbf{x}\|$, and $\text{Law}(\mathbf{x}/\|\mathbf{x}\|) = \text{Unif}(\mathbb{S}^{d-1})$.

The standard Gaussian random vector turns out to be a useful bridge between independent entries and concentration inequalities on the one hand and these measures on the other.

Proposition 1.3.4. Suppose $\mathbf{x} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_d)$. For any $\mathbf{a} \in \mathbb{S}^{d-1}$, we have $\text{Law}(\langle \mathbf{a}, \mathbf{x} \rangle) = \mathcal{N}(0, 1)$. In particular, this law does not depend on $\mathbf{a}$; the distribution of $\mathbf{x}$ is “the same in any direction.” Further, $\mathbf{x}$ is orthogonally invariant.

Proof. Using Proposition 1.2.3, for $\mathbf{Q} \in \mathcal{O}(d)$,

$$\text{Law}(\mathbf{Q}\mathbf{x}) = \mathcal{N}(\mathbf{0}, \mathbf{Q}^\top \mathbf{I}_d \mathbf{Q}) = \mathcal{N}(\mathbf{0}, \mathbf{I}_d) = \text{Law}(\mathbf{x}), \quad (1.3.1)$$

giving the second result. The first follows from the second by considering a $\mathbf{Q}$ whose first row is $\mathbf{a}$. $\square$

So, a standard Gaussian random vector is *both* an i.i.d. random vector *and* an orthogonally invariant one. Explicitly, we have the following, which also has the useful practical interpretation of giving an explicit way to sample uniformly at random from spheres of any dimension.

Corollary 1.3.5. If $\mathbf{x} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_d)$, then $\text{Law}(\mathbf{x}/\|\mathbf{x}\|) = \text{Unif}(\mathbb{S}^{d-1})$.

Proof. The result follows directly from the above together with Proposition 1.3.3. $\square$

In fact, the following elegant result, which we will not prove here, shows that the Gaussian scalar measure is the *only* one having this property.

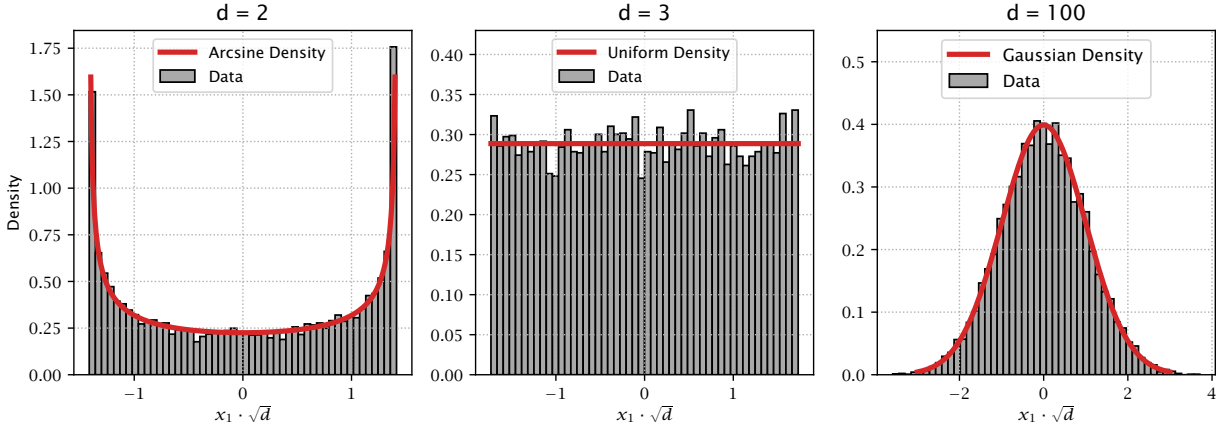


Figure 1.3: An illustration of Borel’s theorem, Theorem 1.3.7, showing two low-dimensional examples $d = 2, 3$ where the distribution of $x_1 \cdot \sqrt{d}$ is known exactly (and is not Gaussian) as well as the Gaussian limit for large d .

Theorem 1.3.6 (Maxwell). Suppose that $\mathbf{x} \sim \mu^{\otimes d}$ for some μ a probability measure on $\mathbb{R}$. If $d \geq 2$ and $\mathbf{x}$ is orthogonally invariant, then $\mu = \mathcal{N}(0, \sigma^2)$ for some $\sigma^2 \geq 0$.

These observations together give some of the fundamental intuitions of high-dimensional geometry. Analogously to random vectors in the box, Corollary 1.2.7 says that $\mathbf{x} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_d)$ has $\|\mathbf{x}\|^2 = d + O(\sqrt{d})$ with high probability, or $\|\mathbf{x}\| = \sqrt{d} + O(1)$. That is, a random standard Gaussian vector usually falls close to the spherical shell of width $O(1)$ around the sphere of radius $\sqrt{d}$. As in the case of the solid box, this typical set is non-convex and instead “hollow” and concentrated far from the origin.

Another way to look at our result is as implying the following approximation of laws, which we emphasize is an informal one:

$$\text{“ } \mathcal{N}(\mathbf{0}, \mathbf{I}_d) \approx \text{Unif}(\mathbb{S}^{d-1}(\sqrt{d})) \text{ ”} \quad (1.3.2)$$

We arrive at this by combining our observations that, when $\mathbf{x} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_d)$, then $\text{Law}(\mathbf{x}/\|\mathbf{x}\|) = \text{Unif}(\mathbb{S}^{d-1})$ (Corollary 1.3.5) and $\|\mathbf{x}\| \approx \sqrt{d}$ (Corollary 1.2.7). This idea can be made rigorous, for instance in the following result.

Theorem 1.3.7 (Borel). Let $\mathbf{x} \sim \text{Unif}(\mathbb{S}^{d-1})$. Then, for any fixed $k \geq 1$, $(x_1 \cdot \sqrt{d}, \dots, x_k \cdot \sqrt{d})$ converges in distribution as $d \rightarrow \infty$ to $\mathcal{N}(\mathbf{0}, \mathbf{I}_k)$.

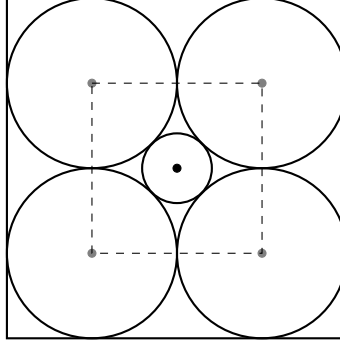
1.4 EXERCISES

1.4.1 HIGH-DIMENSIONAL GEOMETRY

Exercise 1.4.1. Consider the box B in $\mathbb{R}^d$ of side length 2, centered at the origin, with vertices at the points $(\pm 1, \dots, \pm 1)$. For each $\mathbf{s} \in \{\pm 1\}^d$, let $S_{\mathbf{s}}$ be the sphere of radius $\frac{1}{2}$ centered at the point $\frac{1}{2}\mathbf{s}$.

These are 2^d spheres packed on a cubic lattice into the box B . Consider the sphere S' centered at the origin that is tangent to every S_s . Find d_0 such that, if $d < d_0$, then S' is contained in B , but if $d \geq d_0$, then S' is *not* contained in B .

The case $d = 2$ looks as follows. The innermost circle is S' , and the four around it are $S_{(\pm 1, \pm 1)}$. The larger outermost square is B ; the smaller square in a dotted line is just for reference to show how the centers of the latter circles are arranged. Your task is to show that, in high dimension, the innermost circle is not contained in the outermost square (!).



Exercise 1.4.2. Show that there is a constant $c > 0$ such that, for all $\epsilon \in (0, 1)$, for d sufficiently large, there are at least $N = \exp(c\epsilon^2 d)$ unit vectors $\mathbf{v}_1, \dots, \mathbf{v}_N$ in $\mathbb{R}^d$ such that $|\langle \mathbf{v}_i, \mathbf{v}_j \rangle| \leq \epsilon$ for all $1 \leq i < j \leq N$ (i.e., such that the $\mathbf{v}_i$ are almost orthogonal).

(**HINT:** Consider random vectors. Choose a convenient distribution to work with. Note, though, that the question is not making a probabilistic statement.)

Exercise 1.4.3. Consider the shape $\Delta^d \subset \mathbb{R}^d$ that is the convex hull of the points $\mathbf{0}, \mathbf{e}_1, \mathbf{e}_1 + \mathbf{e}_2, \dots, \mathbf{e}_1 + \mathbf{e}_2 + \dots + \mathbf{e}_d$ (the origin plus the “partial sums” of the standard basis vectors). This is a *simplex* or high-dimensional tetrahedron, though not an equilateral one: edges have lengths varying among $1 = \sqrt{1}, \sqrt{2}, \dots, \sqrt{d}$. Compute the volume of Δ^d . What is the side length of a cube in $\mathbb{R}^d$ with the same volume? (Give an asymptotic approximation as $d \rightarrow \infty$.)

(**HINT:** For the volume computation, consider gluing several copies of Δ^d together to tile a more familiar object.)

Exercise 1.4.4. Suppose that $\mathbf{g} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_d)$. Show the non-asymptotic inequalities

$$\left(1 - \frac{2}{d}\right) \sqrt{d} \leq \mathbb{E} \|\mathbf{g}\| \leq \sqrt{d}.$$

(**HINT:** For the lower bound, that $\sqrt{x} \geq \frac{1}{2}(1+x-(x-1)^2)$ for all $x \geq 0$. Apply this with $x := \|\mathbf{g}\|^2/d$.)

1.4.2 CONCENTRATION INEQUALITIES

Exercise 1.4.5 (Subgaussianity). Recall that a centered random scalar $X \in \mathbb{R}$ is called σ^2 -*subgaussian* if $\mathbb{E} \exp(\lambda X) \leq \exp(\sigma^2 \lambda^2 / 2)$ for all $\lambda \in \mathbb{R}$. A centered random vector $\mathbf{x} \in \mathbb{R}^d$ is called the same if, for all $\mathbf{u} \in \mathbb{R}^d$ with $\|\mathbf{u}\| = 1$, $\langle \mathbf{u}, \mathbf{x} \rangle$ is a σ^2 -subgaussian random scalar.

1. Show that if X is a centered random scalar that is σ^2 -subgaussian, then aX is $a^2\sigma^2$ -subgaussian, and that if $X_1, \dots, X_n$ are independent and each X_i is σ_i^2 -subgaussian, then $\sum_{i=1}^n X_i$ is $(\sum_{i=1}^n \sigma_i^2)$ -subgaussian. Conclude that a random vector with independent entries that are each centered and σ^2 -subgaussian is itself σ^2 -subgaussian.

2. Show that the following three conditions on a centered random scalar X are equivalent:

- (a) There exists $\sigma^2 > 0$ such that X is σ^2 -subgaussian.
- (b) There exists $c > 0$ such that, for all $t > 0$, $\mathbb{P}[|X| \geq t] \leq 2 \exp(-ct^2)$.
- (c) There exists $\lambda > 0$ such that $\mathbb{E} \exp(\lambda X^2) < \infty$.

3. For X a centered subgaussian scalar, define

$$\|X\|_{\text{sg}} := \inf\{t > 0 : \mathbb{E} \exp(X^2/t^2) \leq 2\}.$$

Note that this is well-defined by Part 2. Show that $\|\cdot\|_{\text{sg}}$ is a norm on the space of centered subgaussian random variables (over a given probability space).

(**HINT:** Prove and use that $(X + Y)^2 \leq (1 + a)X^2 + (1 + \frac{1}{a})Y^2$ for any $a > 0$ (you might be more familiar with the case $a = 1$), together with Hölder's inequality.)

Exercise 1.4.6 (Subexponentiality). Recall that, for $\sigma^2, K > 0$, a centered random scalar $X \in \mathbb{R}$ is called (σ^2, K) -subexponential if

$$\mathbb{E}[\exp(\lambda X)] \leq \exp(\sigma^2 \lambda^2 / 2) \text{ for all } |\lambda| \leq K.$$

A centered random vector $\mathbf{x} \in \mathbb{R}^d$ is called the same if, for every $\mathbf{u} \in \mathbb{R}^d$ with $\|\mathbf{u}\| = 1$, $\langle \mathbf{u}, \mathbf{x} \rangle$ is a (σ^2, K) -subexponential random scalar.

1. Show that multiplying a (σ^2, K) -subexponential scalar X by $a \neq 0$ gives a subexponential scalar with parameters $(a^2 \sigma^2, K/|a|)$. If $X_1, \dots, X_n$ are independent and each X_i is (σ_i^2, K_i) -subexponential, show that $\sum_{i=1}^n X_i$ is subexponential with parameters

$$\left(\sum_{i=1}^n \sigma_i^2, \min_{1 \leq i \leq n} K_i \right).$$

Conclude that a random vector with independent (σ^2, K) -subexponential entries is (σ^2, K) -subexponential.

2. Show that the following three conditions on a centered random scalar X are equivalent:

- (a) There exist $\sigma^2, K > 0$ such that X is (σ^2, K) -subexponential.
- (b) There exists $c > 0$ such that, for all $t > 0$, $\mathbb{P}[|X| \geq t] \leq 2 \exp(-ct)$.
- (c) There exists $\lambda > 0$ such that $\mathbb{E}[\exp(\lambda |X|)] < \infty$.

3. For X a centered subexponential scalar, define

$$\|X\|_{\text{se}} := \inf\{t > 0 : \mathbb{E}[\exp(|X|/t)] \leq 2\}.$$

Note that this is well-defined by Part 2. Show that $\|\cdot\|_{\text{se}}$ is a norm on the space of centered subexponential random variables (over a given probability space). Show also that, for every centered subgaussian scalar X ,

$$\|X^2 - \mathbb{E}[X^2]\|_{\text{se}} \leq 2\|X\|_{\text{sg}}^2,$$

where $\|\cdot\|_{\text{sg}}$ is defined in Exercise 1.4.5.

(**HINT:** Use Hölder's inequality for the triangle inequality and Jensen's inequality for the bound on centered squares.)

1.4.3 ORTHOGONAL INVARIANCE FOR VECTORS

Exercise 1.4.7 (Factorization of distribution). Suppose $\mathbf{x} \in \mathbb{R}^d$ is an orthogonally invariant random vector, i.e., one having $\text{Law}(\mathbf{Q}\mathbf{x}) = \text{Law}(\mathbf{x})$ for all $\mathbf{Q} \in \mathcal{O}(d)$. Suppose also that $\mathbf{x} \neq \mathbf{0}$ almost surely. Show that the scalar $\|\mathbf{x}\|$ and the vector $\hat{\mathbf{x}} := \mathbf{x}/\|\mathbf{x}\|$ are independent random variables, i.e., that $\mathbb{P}[\|\mathbf{x}\| \in A, \hat{\mathbf{x}} \in B] = \mathbb{P}[\|\mathbf{x}\| \in A] \cdot \mathbb{P}[\hat{\mathbf{x}} \in B]$ for all pairs of Borel sets $A \subseteq \mathbb{R}_{\geq 0}$ and $B \subseteq \mathbb{S}^{d-1}$.

Exercise 1.4.8 (Moments). Suppose $\mathbf{x} \in \mathbb{R}^d$ is an orthogonally invariant random vector whose entries have finite first and second moments. Show that $\mathbb{E}\mathbf{x} = \mathbf{0}$ and $\mathbb{E}\mathbf{x}\mathbf{x}^\top = \sigma^2 \mathbf{I}_d$ for some $\sigma^2 \geq 0$.

2 | GEOMETRY OF RECTANGULAR RANDOM MATRICES

We now take our first steps into random matrix theory, using the results from the previous chapter to study the behavior of asymmetric or rectangular matrices with i.i.d. Gaussian entries. In particular, we will focus on the random matrix model

$$\mathbf{G} \sim \mathcal{N}(0, 1)^{\otimes d \times m},$$

a $d \times m$ random matrix with i.i.d. Gaussian entries.

2.1 GAUSSIAN RANDOM FIELD VIEWPOINT

Let us start by trying to gain some intuition about what $\mathbf{G}$ does to a single $\mathbf{y} \in \mathbb{R}^m$ as well as several given $\mathbf{y}_1, \dots, \mathbf{y}_n \in \mathbb{R}^m$. Immediately our observations about Gaussian random vectors become very useful, since we can easily completely characterize the laws of these images.

Proposition 2.1.1. $\text{Law}(\mathbf{G}\mathbf{y}) = \mathcal{N}(\mathbf{0}, \|\mathbf{y}\|^2 \mathbf{I}_d)$.

Proof. Since $\mathbf{G}\mathbf{y}$ is linear in $\mathbf{G}$, whose entries form a Gaussian random vector, $\mathbf{G}\mathbf{y}$ is itself a Gaussian random vector. So, it suffices to compute its mean and covariance:

$$\begin{aligned} \mathbb{E}(\mathbf{G}\mathbf{y})_i &= \mathbb{E} \sum_{a=1}^m G_{ia} y_a \\ &= 0, \\ \text{Cov}((\mathbf{G}\mathbf{y})_i, (\mathbf{G}\mathbf{y})_j) &= \mathbb{E}(\mathbf{G}\mathbf{y})_i (\mathbf{G}\mathbf{y})_j \\ &= \mathbb{E} \left(\sum_{a=1}^m G_{ia} y_a \right) \left(\sum_{b=1}^m G_{jb} y_b \right) \\ &= \sum_{a,b=1}^m y_a y_b \cdot \mathbb{E} G_{ia} G_{jb} \\ &= \begin{cases} 0 & \text{if } i \neq j, \\ \|\mathbf{y}\|^2 & \text{if } i = j \end{cases}, \end{aligned}$$

where we have used that the G_{ia} are i.i.d. and have mean zero.

Another approach is to work with matrices; it is good to get used to using matrices inside of expectations in the way written below. For the mean, the following operation is valid in general

when an expectation of a matrix product only has one random matrix involved:

$$\mathbb{E}[\mathbf{G}\mathbf{y}] = \mathbb{E}[\mathbf{G}]\mathbf{y} = \mathbf{0}. \quad (2.1.1)$$

You may check that this is just another form of the linearity of expectation. For the covariance, if we write $\mathbf{g}_1, \dots, \mathbf{g}_m \in \mathbb{R}^d$ for the columns of $\mathbf{G}$, so that $\mathbf{g}_a \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_d)$ are i.i.d. standard Gaussian random vectors, we may treat the entire covariance using similar manipulations:

$$\begin{aligned} \text{Cov}(\mathbf{G}\mathbf{y}) &= \mathbb{E}(\mathbf{G}\mathbf{y})(\mathbf{G}\mathbf{y})^\top \\ &= \mathbb{E} \left(\sum_{a=1}^m \mathcal{Y}_a \mathbf{g}_a \right) \left(\sum_{b=1}^m \mathcal{Y}_b \mathbf{g}_b^\top \right) \\ &= \sum_{a,b=1}^m \mathcal{Y}_a \mathcal{Y}_b \cdot \mathbb{E} \mathbf{g}_a \mathbf{g}_b^\top \\ &= \sum_{a=1}^m \mathcal{Y}_a^2 \cdot \mathbf{I}_d \\ &= \|\mathbf{y}\|^2 \mathbf{I}_d, \end{aligned}$$

now using that the $\mathbf{g}_a$ are i.i.d. and have mean zero and covariance identity. $\square$

By the same kinds of calculations, we may also deduce the joint law of any collection of finitely many $\mathbf{G}\mathbf{y}_1, \dots, \mathbf{G}\mathbf{y}_n$.

Proposition 2.1.2. Let $\mathbf{y}_1, \dots, \mathbf{y}_n \in \mathbb{R}^m$. Then,

$$\text{Law} \left(\begin{bmatrix} \mathbf{G}\mathbf{y}_1 \\ \vdots \\ \mathbf{G}\mathbf{y}_n \end{bmatrix} \right) = \mathcal{N} \left(\mathbf{0}, \begin{bmatrix} \|\mathbf{y}_1\|^2 \mathbf{I}_d & \langle \mathbf{y}_1, \mathbf{y}_2 \rangle \mathbf{I}_d & \cdots & \langle \mathbf{y}_1, \mathbf{y}_n \rangle \mathbf{I}_d \\ \langle \mathbf{y}_2, \mathbf{y}_1 \rangle \mathbf{I}_d & \|\mathbf{y}_2\|^2 \mathbf{I}_d & \cdots & \langle \mathbf{y}_2, \mathbf{y}_n \rangle \mathbf{I}_d \\ \vdots & \vdots & \ddots & \vdots \\ \langle \mathbf{y}_n, \mathbf{y}_1 \rangle \mathbf{I}_d & \langle \mathbf{y}_n, \mathbf{y}_2 \rangle \mathbf{I}_d & \cdots & \|\mathbf{y}_n\|^2 \mathbf{I}_d \end{bmatrix} \right). \quad (2.1.2)$$

Written more concisely, if $\mathbf{Y} \in \mathbb{R}^{m \times n}$ has the $\mathbf{y}_i$ as its columns, then the covariance matrix above is $(\mathbf{Y}^\top \mathbf{Y}) \otimes \mathbf{I}_d$.

Proof. The result follows from identical calculations to the previous ones in each block of the covariance matrix. $\square$

The previous result fully characterizes the *Gaussian random field* (a term for a higher-dimensional Gaussian process) that attaches the random vector $\mathbf{G}\mathbf{y} \in \mathbb{R}^d$ to each point $\mathbf{y} \in \mathbb{R}^m$, since we have characterized the joint law of any finitely many points of this field.

From this characterization, we may start to probe more particular geometric properties of the action of $\mathbf{G}$. We may, for example, compute what multiplication by $\mathbf{G}$ does to expected lengths:

$$\mathbb{E} \|\mathbf{G}\mathbf{y}\|^2 = \text{Tr}(\text{Cov}(\mathbf{G}\mathbf{y})) = d \|\mathbf{y}\|^2. \quad (2.1.3)$$

What about angles? By a similar calculation you can find

$$\mathbb{E} \langle \mathbf{G}\mathbf{y}_1, \mathbf{G}\mathbf{y}_2 \rangle = \text{Tr}(\text{Cov}(\mathbf{G}\mathbf{y}_1, \mathbf{G}\mathbf{y}_2)) = d \langle \mathbf{y}_1, \mathbf{y}_2 \rangle, \quad (2.1.4)$$

(Actually, this also follows directly from the calculation for distances by the polarization identity $\langle \mathbf{y}_1, \mathbf{y}_2 \rangle = \frac{1}{4} \|\mathbf{y}_1 + \mathbf{y}_2\|^2 - \frac{1}{4} \|\mathbf{y}_1 - \mathbf{y}_2\|^2$. Try working it out.) Indeed, $\mathbf{G}$ acts in the same way on the

entire *Gram matrix* of any finite collection of vectors in expectation: for $\mathbf{y}_1, \dots, \mathbf{y}_n \in \mathbb{R}^m$ organized into the columns of $\mathbf{Y}$, we have

$$\mathbb{E}(\mathbf{G}\mathbf{Y})^\top(\mathbf{G}\mathbf{Y}) = d\mathbf{Y}^\top\mathbf{Y}. \quad (2.1.5)$$

Remark 2.1.3 (Gram matrices). The Gram matrix is a fundamental geometric object that does not always get its due in a linear algebra class. For a set of vectors $\mathbf{y}_i$ as above, $\mathbf{Y}^\top\mathbf{Y}$ describes the entire relative geometry of this set. Looking at the entries, we see that the Gram matrix contains the lengths of the $\mathbf{y}_i$ (on the diagonal) and the angles between each pair (on the off-diagonal). Actually, this information fully specifies the geometry of this set of vectors: if any other $\mathbf{y}'_1, \dots, \mathbf{y}'_n \in \mathbb{R}^m$ have the same Gram matrix, then there must be an orthogonal $\mathbf{Q} \in \mathcal{O}(m)$ such that $\mathbf{Q}\mathbf{y}_i = \mathbf{y}'_i$; that is, it is possible to get from one point cloud to the other by rotations and reflections.

The above calculations suggest that we should normalize away the factors of d throughout by considering instead

$$\widehat{\mathbf{G}} := \frac{1}{\sqrt{d}}\mathbf{G}, \quad \text{Law}(\widehat{\mathbf{G}}) = \mathcal{N}\left(0, \frac{1}{d}\right)^{\otimes d \times m}. \quad (2.1.6)$$

This matrix will satisfy the appealing properties

$$\mathbb{E}\|\widehat{\mathbf{G}}\mathbf{y}\|^2 = \|\mathbf{y}\|^2, \quad (2.1.7)$$

$$\mathbb{E}\langle \widehat{\mathbf{G}}\mathbf{y}_1, \widehat{\mathbf{G}}\mathbf{y}_2 \rangle = \langle \mathbf{y}_1, \mathbf{y}_2 \rangle, \quad (2.1.8)$$

$$\mathbb{E}(\widehat{\mathbf{G}}\mathbf{Y})^\top(\widehat{\mathbf{G}}\mathbf{Y}) = \mathbf{Y}^\top\mathbf{Y}. \quad (2.1.9)$$

That is, in expectation, $\widehat{\mathbf{G}}$ preserves lengths, angles, and Gram matrices, and thus looks as if it preserves all of the geometry of $\mathbb{R}^m$! We will see in short order that this is true in a specific quantitative sense with high probability as well, making multiplication by $\widehat{\mathbf{G}}$ a very useful algorithmic and computational tool.

First, though, let us clarify an important caveat. Again, the above makes it look like $\widehat{\mathbf{G}}$ preserves the relative geometry of any given point cloud $\mathbf{y}_1, \dots, \mathbf{y}_n \in \mathbb{R}^m$ while mapping it to $\mathbb{R}^d$. But, this holds for any d , including $d \ll m$, and even $d = 1$! Thus, part of the above optimistic intuition is a kind of illusion created by taking expectations.

There are two important cautions about this intuition. First, once $d \ll m$, once we draw $\widehat{\mathbf{G}}$ at random, of course we can find $\mathbf{y} \neq \mathbf{0}$ such that $\widehat{\mathbf{G}}\mathbf{y} = \mathbf{0}$ (any $\mathbf{y}$ in the kernel of $\widehat{\mathbf{G}}$), and for such $\mathbf{y}$ we will have $0 = \|\widehat{\mathbf{G}}\mathbf{y}\| \ll \|\mathbf{y}\|$. Thus, $\widehat{\mathbf{G}}$ cannot and will not preserve the relative geometry of *any* vectors we multiply it by, and in particular we can easily find vectors *depending* on $\widehat{\mathbf{G}}$ whose relative geometry it badly distorts. On the other hand, if we *first* pick $\mathbf{y}$ and then draw a random $\widehat{\mathbf{G}}$, then we have reason to hope that $\|\widehat{\mathbf{G}}\mathbf{y}\|$ will concentrate around the value $\|\mathbf{y}\|$ suggested by the previous calculations; indeed, our calculations show that $\text{Law}(\widehat{\mathbf{G}}\mathbf{y}) = \mathcal{N}(0, \frac{1}{d}\|\mathbf{y}\|^2\mathbf{I}_d)$, to which Corollary 1.2.7 should apply.

Second, even provided we make sure $\widehat{\mathbf{G}}$ is independent of the $\mathbf{y}$ that we multiply by it, the extent to which the above properties hold with high probability rather than merely in expectation will depend on how large d is: when $d = 1$, then $\widehat{\mathbf{G}} = \mathbf{g}^\top$ for $\mathbf{g} \sim \mathcal{N}(0, \mathbf{I}_m)$, and we still have $\mathbb{E}\|\widehat{\mathbf{G}}\mathbf{y}_i\|^2 = \mathbb{E}\langle \mathbf{g}, \mathbf{y}_i \rangle^2 = \|\mathbf{y}_i\|^2$. But, it should be intuitively clear that we cannot actually have $\langle \mathbf{g}, \mathbf{y}_i \rangle^2 \approx \|\mathbf{y}_i\|^2$ simultaneously for very many $\mathbf{y}_i$, even if we choose them in advance of drawing $\widehat{\mathbf{G}}$. (Exercise: Try constructing concrete counterexamples.)

The next step is to understand the singular values of $\mathbf{G}$. This will require controlling a quadratic form uniformly over all unit vectors, using a finite approximation of the sphere or ball called an ϵ -net.

2.2 THE GEOMETRIC METHOD OF RANDOM MATRIX THEORY

We will now move from the action of G on fixed vectors to its singular values. This will give us a first spectral analysis of rectangular matrices with i.i.d. Gaussian entries.

We will see below that this amounts to controlling the operator norm of a certain symmetric “error” matrix M . The general approach we take to this task is one of the main methods of random matrix theory that we will study, so let us outline it separately here.

The chief trouble is that the operator norm is a quantity whose definition involves an uncountable quantification over test vectors. We adopt the following notation for closed balls in $\mathbb{R}^d$, which will be useful soon:

$$\mathbb{B}^d(\mathbf{x}, r) := \{\mathbf{y} \in \mathbb{R}^d : \|\mathbf{y} - \mathbf{x}\| \leq r\}, \quad (2.2.1)$$

$$\mathbb{B}^d := \mathbb{B}^d(\mathbf{0}, 1) \supset \mathbb{S}^{d-1}. \quad (2.2.2)$$

The operator norm can be formulated in either of the equivalent ways

$$\|M\| = \sup_{\mathbf{x} \in \mathbb{S}^{d-1}} |\mathbf{x}^\top M \mathbf{x}| = \sup_{\mathbf{x} \in \mathbb{B}^d} |\mathbf{x}^\top M \mathbf{x}|. \quad (2.2.3)$$

In either case, if $\mathbb{S}^{d-1}$ or $\mathbb{B}^d$ were replaced by a finite set, we could use the union bound (also called the first moment method), bounding $\mathbb{P}[|\mathbf{x}^\top M \mathbf{x}| > t]$ (which as we will see boils down to yet another application of Corollary 1.2.7 on the concentration of norms of Gaussian random vectors) and using a union bound to control $\mathbb{P}[\|M\| > t]$, but this is not the case.

But, it will turn out to be useful to take this hypothetical situation more seriously, so let us define notation for these notions.

Definition 2.2.1. For any nonempty bounded set $X \subset \mathbb{R}^d$ and $M \in \mathbb{R}_{\text{sym}}^{d \times d}$, we define

$$\|M\|_X := \sup_{\mathbf{x} \in X} |\mathbf{x}^\top M \mathbf{x}|. \quad (2.2.4)$$

This is not always actually a norm; depending on X , it may only be a *seminorm*, where $\|M\|_X = 0$ can hold for $M \neq \mathbf{0}$. The definition will be useful in any case, however. Applying the first moment method to a general random matrix gives the following:

Proposition 2.2.2. Let $M \in \mathbb{R}_{\text{sym}}^{d \times d}$ be a random matrix and $X \subset \mathbb{R}^d$ be a nonempty finite set. Then,

$$\mathbb{P}[\|M\|_X \geq t] \leq |X| \cdot \max_{\mathbf{x} \in X} \mathbb{P}[|\mathbf{x}^\top M \mathbf{x}| \geq t]. \quad (2.2.5)$$

As mentioned above, controlling the second factor above involving scalar probabilities will be possible in our case using Corollary 1.2.7; in general, it reduces to simpler questions involving scalar probability theory. Thus, the key point will be the tradeoff between the factor of $|X|$ in the bound above, which makes the bound deteriorate as $|X|$ grows, versus the error in approximating $\|M\|$ by $\|M\|_X$, which we will soon see decreases as $|X|$ increases. We will therefore reduce the matter of bounding $\|M\|$ to merely that of finding a sufficiently small set X that behaves like a sufficiently dense “grid” of $\mathbb{S}^{d-1}$ or $\mathbb{B}^d$ to approximate the operator norm accurately (the sets we work with will not actually be grids but will behave similarly). We call this general approach the **Geometric Method** of random matrix theory, because it reduces the matter of understanding norms of random matrices to (possibly more complicated variants of) the above purely geometric question.

2.3 SINGULAR VALUES

We focus on “short fat” matrices, that is, the setting $d \leq m$ (often we will also mention how our results specialize to the asymptotic regime $d \ll m$). We also return to working with $G \sim \mathcal{N}(0, 1)^{\otimes d \times m}$ again in order to simplify the notation.

For any fixed unit vector $x \in \mathbb{R}^m$, the vector Gx has law $\mathcal{N}(0, I_d)$. Thus Corollary 1.2.7 gives concentration of the quadratic form

$$x^\top (G^\top G - dI_m)x = \|Gx\|^2 - d. \quad (2.3.1)$$

When d is large, this is small compared with d with high probability. A union bound gives simultaneous control for a sufficiently small finite set of such test vectors. But, as we observed above, $G^\top G$ has rank at most d , so when $d < m$ it cannot be close to dI_m in operator norm relative to d . To understand the matrix itself, we will study its eigenvalues, or equivalently the singular values of G .

Note that $G^\top G \geq 0$, so its eigenvalues are non-negative. Call these $\lambda_1 \geq \dots \geq \lambda_m \geq 0$. Since $\text{rank}(G^\top G) \leq d$, we have $\lambda_{d+1} = \dots = \lambda_m = 0$. The remaining eigenvalues are related to the singular values of G by $\lambda_i = \sigma_i(G)^2$. We begin by studying these eigenvalues, or equivalently the singular values of G .

2.3.1 APPLYING THE GEOMETRIC METHOD

It will be more convenient to instead consider the eigenvalues of $GG^\top \in \mathbb{R}_{\text{sym}}^{d \times d}$. Since the eigenvalues of $GG^\top$ and $G^\top G$ are equal except that the latter has $m - d$ extra eigenvalues equal to 0, $GG^\top$ will provide us with the same information; at the end, we will return to $G^\top G$ and deduce some geometric intuitions about it. The probabilistic reason for working with $GG^\top$ instead is that, if $g_1, \dots, g_m \in \mathbb{R}^d$ are the columns of G (distributed i.i.d. as $\mathcal{N}(0, I_d)$), then we have

$$GG^\top = \sum_{i=1}^m g_i g_i^\top$$

which is a sum of $m \gg d$ i.i.d. random matrices of dimension d . As we will see, in such a situation it is justified to hope for a matrix-valued law of large numbers to hold, giving us

$$\begin{aligned} &\approx m \cdot \mathbb{E}[g_1 g_1^\top] \\ &= mI_d \\ &= \mathbb{E}[GG^\top]. \end{aligned} \quad (2.3.2)$$

Roughly speaking, this heuristic is more accurate than the same applied to $G^\top G$ where the roles of d and m are reversed because we want in our appeal to the law of large numbers to work with an expression having “as much averaging as possible.”

Remark 2.3.1. To emphasize the difference in how we are discussing these two matrices: above we observed that $G^\top G$ *looks like* dI_m from the point of view of a small enough number of quadratic forms, but we also observed that it *cannot actually* be close to this matrix (for instance by simple rank considerations). On the other hand, the above heuristic suggests that $GG^\top$ looks like mI_d , and we will show below that this is actually true, and use that to understand the trickier behavior of $G^\top G$.

Concretely, when $m \gg d$, we expect the operator norm of the difference $\|\mathbf{G}\mathbf{G}^\top - m\mathbf{I}_d\|$ to be small compared with m . This kind of statement that a sum of independent or weakly dependent random matrices is close in operator norm to its expectation is often called a *matrix concentration inequality*. Let us write $\Delta := \mathbf{G}\mathbf{G}^\top - m\mathbf{I}_d$.

Remark 2.3.2 (Statistical interpretation). Expanding in a sum of rank one terms again, we may also view

$$\frac{1}{m}\Delta = \left(\frac{1}{m} \sum_{i=1}^m \mathbf{g}_i \mathbf{g}_i^\top \right) - \mathbb{E}[\mathbf{g}_1 \mathbf{g}_1^\top] \quad (2.3.3)$$

as the difference between a *sample covariance* and the true covariance of several i.i.d. draws $\mathbf{g}_1, \dots, \mathbf{g}_m$ from $\mathcal{N}(\mathbf{0}, \mathbf{I}_d)$. Thus, our investigation will also address the natural statistical question of how good of a statistical estimator the sample covariance is (in this particular case).

To show that $\|\Delta\|$ is relatively small, we will apply the Geometric Method outlined in Section 2.2 with $\mathbf{M} = \Delta$.

2.3.2 NETS, COVERINGS, AND PACKINGS

To do this, our first task will be to develop some tools for relating $\|\mathbf{M}\|_{\mathcal{X}}$ for finite $\mathcal{X}$ to $\|\mathbf{M}\|$, *discretizing* the continuous optimization defining the operator norm.

Definition 2.3.3 (ϵ -net). A set $\mathcal{X} = \{\mathbf{x}_1, \dots, \mathbf{x}_N\} \subseteq \mathcal{Y} \subset \mathbb{R}^d$ is an ϵ -net or ϵ -covering of $\mathcal{Y}$ if

$$\mathcal{Y} \subseteq \bigcup_{i=1}^N \mathbb{B}^d(\mathbf{x}_i, \epsilon), \quad (2.3.4)$$

or equivalently if, for all $\mathbf{y} \in \mathcal{Y}$, there is $\mathbf{x}_i \in \mathcal{X}$ such that $\|\mathbf{y} - \mathbf{x}_i\| \leq \epsilon$.

The following key result shows that we may indeed discretize the operator norm and bound it by a maximum over a suitable net.

Lemma 2.3.4. Suppose that $\mathbf{M} \in \mathbb{R}_{\text{sym}}^{d \times d}$ and $\mathcal{X}$ is an ϵ -net of $\mathbb{S}^{d-1}$ for $0 < \epsilon < \frac{1}{2}$. Then,

$$\|\mathbf{M}\|_{\mathcal{X}} \leq \|\mathbf{M}\| \leq \frac{1}{1-2\epsilon} \|\mathbf{M}\|_{\mathcal{X}}. \quad (2.3.5)$$

The same is true if $\mathcal{X}$ is instead an ϵ -net of $\mathbb{B}^d$.

Proof. The first inequality is immediate. For the second, let $\mathbf{x} \in \mathbb{S}^{d-1}$ be such that $|\mathbf{x}^\top \mathbf{M} \mathbf{x}| = \|\mathbf{M}\|$, and let $\mathbf{x}_i \in \mathcal{X}$ be such that $\|\mathbf{x} - \mathbf{x}_i\| \leq \epsilon$. We then have

$$\begin{aligned} \|\mathbf{M}\|_{\mathcal{X}} &\geq |\mathbf{x}_i^\top \mathbf{M} \mathbf{x}_i| \\ &= |(\mathbf{x} + \mathbf{x}_i - \mathbf{x})^\top \mathbf{M}(\mathbf{x} + \mathbf{x}_i - \mathbf{x})| \\ &= |\mathbf{x}^\top \mathbf{M} \mathbf{x} + \mathbf{x}^\top \mathbf{M}(\mathbf{x}_i - \mathbf{x}) + (\mathbf{x}_i - \mathbf{x})^\top \mathbf{M} \mathbf{x}| \\ &\geq |\mathbf{x}^\top \mathbf{M} \mathbf{x}| - |\mathbf{x}^\top \mathbf{M}(\mathbf{x}_i - \mathbf{x})| - |(\mathbf{x}_i - \mathbf{x})^\top \mathbf{M} \mathbf{x}| \\ &\geq (1 - 2\epsilon) \|\mathbf{M}\|, \end{aligned}$$

and rearranging gives the result. The same argument applies verbatim in the case of a net of $\mathbb{B}^d$. $\square$

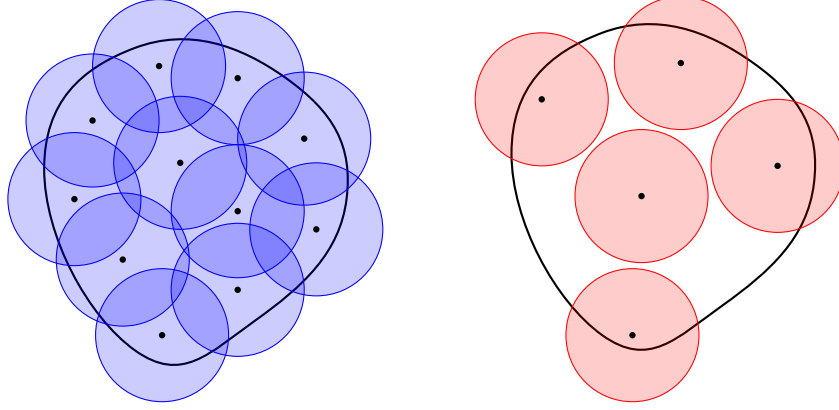


Figure 2.1: An illustration of a covering (left) and a packing (right) of a given set in $\mathbb{R}^2$ by balls of the same radius.

Next, we need to develop some further tools for constructing ϵ -nets in a way such that we have some control over their size. As you can convince yourself, it is rather tricky to construct explicit ϵ -nets by hand by completely describing the vectors of $\mathcal{X}$. Fortunately, there is a more abstract approach that lets us non-constructively show that small ϵ -nets exist. This proceeds by considering a complementary notion:

Definition 2.3.5 (ϵ -packing). A set $\mathcal{X} = \{x_1, \dots, x_N\} \subseteq \mathcal{Y} \subset \mathbb{R}^d$ is an ϵ -packing of $\mathcal{Y}$ if the balls $\mathbb{B}^d(x_i, \epsilon)$ are pairwise disjoint, or equivalently if $\|x_i - x_j\| > 2\epsilon$ for all $i, j \in [N]$ distinct.

Proposition 2.3.6. Let $\mathcal{X}$ be a maximal ϵ -packing of $\mathcal{Y}$, one that cannot be extended to a larger packing by adding more points. Then, $\mathcal{X}$ is a 2ϵ -net of $\mathcal{Y}$.

Proof. Let $y \in \mathcal{Y} \setminus \mathcal{X}$. Since adding y to $\mathcal{X}$ must yield a set that is not an ϵ -packing, there must be some $x \in \mathcal{X}$ such that $\mathbb{B}^d(x, \epsilon) \cap \mathbb{B}^d(y, \epsilon) \neq \emptyset$. Letting z be in the intersection, we have $\|x - y\| \leq \|x - z\| + \|y - z\| \leq 2\epsilon$. Thus, either $y \in \mathcal{X}$ or $y \in \mathbb{B}^d(x, 2\epsilon)$ for some $x \in \mathcal{X}$, giving the result. $\square$

Thus, to show that small nets exist, it suffices to show that small maximal packings exist. But, there is a fundamental *volumetric* bound on the size of arbitrary packings, including maximal ones: roughly speaking, the total volume of the union of the balls in a packing cannot be much bigger than the volume of the set being packed (at least for a nice full-dimensional convex set).

Proposition 2.3.7. Let $\mathcal{X}$ be any ϵ -packing of $\mathbb{B}^d$. Then, $|\mathcal{X}| \leq (1 + \frac{1}{\epsilon})^d$.

Proof. We have

$$\bigcup_{x \in \mathcal{X}} \mathbb{B}^d(x, \epsilon) \subseteq \mathbb{B}^d(\mathbf{0}, 1 + \epsilon), \quad (2.3.6)$$

and the union on the left-hand side is disjoint. Thus, taking volumes,

$$|\mathcal{X}| \cdot \text{vol}(\mathbb{B}^d(\mathbf{0}, \epsilon)) \leq \text{vol}(\mathbb{B}^d(\mathbf{0}, 1 + \epsilon)). \quad (2.3.7)$$

Rearranging,

$$|\mathcal{X}| \leq \frac{\text{vol}(\mathbb{B}^d(\mathbf{0}, 1 + \epsilon))}{\text{vol}(\mathbb{B}^d(\mathbf{0}, \epsilon))} = \left(\frac{1 + \epsilon}{\epsilon}\right)^d = \left(1 + \frac{1}{\epsilon}\right)^d, \quad (2.3.8)$$

by scaling considerations of the volumes of d -dimensional balls. $\square$

Corollary 2.3.8. For any $\epsilon > 0$, there exists an ϵ -net $\mathcal{X}$ of $\mathbb{B}^d$ of size $|\mathcal{X}| \leq (1 + \frac{2}{\epsilon})^d$.

Proof. Let $\mathcal{X}$ be a maximal $\frac{\epsilon}{2}$ -packing of $\mathbb{B}^d$. Such a packing exists: greedily add points while possible, and Proposition 2.3.7 guarantees that this process stops after finitely many steps. By Proposition 2.3.7, $|\mathcal{X}| \leq (1 + \frac{2}{\epsilon})^d$, while by Proposition 2.3.6, $\mathcal{X}$ is an ϵ -net. $\square$

2.3.3 COARSE NON-ASYMPTOTIC BOUND ON SINGULAR VALUES

Finally, we may put together the pieces to bound $\|\Delta\|$: we use the existence of a not-too-large ϵ -net by the above non-constructive method, approximate the operator norm by discretizing to test vectors in this net, and use the first moment method.

Theorem 2.3.9. If $d \leq m$, then

$$\mathbb{P}[\|\Delta\| > 64\sqrt{dm}] \leq 2 \exp(-d). \quad (2.3.9)$$

Proof. Let us fix $\epsilon := 1/4$. Let $\mathcal{X}$ be an ϵ -net of $\mathbb{B}^d$ as in Corollary 2.3.8, by which we can assume $|\mathcal{X}| \leq (1 + \frac{2}{\epsilon})^d = 9^d$. By Lemma 2.3.4, we have $\|\Delta\| \leq \frac{1}{1-2\epsilon} \|\Delta\|_{\mathcal{X}} = 2\|\Delta\|_{\mathcal{X}}$. Thus, for $t > 0$, we may bound

$$\begin{aligned} \mathbb{P}[\|\Delta\| \geq t] &\leq \mathbb{P}\left[\|\Delta\|_{\mathcal{X}} \geq \frac{t}{2}\right] && \text{(Lemma 2.3.4)} \\ &\leq 9^d \cdot \max_{x \in \mathcal{X}} \mathbb{P}\left[|x^\top \Delta x| \geq \frac{t}{2}\right]. && \text{(Proposition 2.2.2)} \end{aligned}$$

For any fixed $x \in \mathcal{X}$, the vector $G^\top x$ has law $\mathcal{N}(\mathbf{0}, \|x\|^2 I_m)$. Expanding the definition of Δ and using $\|x\| \leq 1$, we therefore have

$$\begin{aligned} \mathbb{P}\left[|x^\top \Delta x| \geq \frac{t}{2}\right] &= \mathbb{P}\left[|\|G^\top x\|^2 - m\|x\|^2| \geq \frac{t}{2}\right] \\ &= \mathbb{P}_{g \sim \mathcal{N}(\mathbf{0}, I_m)}\left[|\|g\|^2 - m| \geq \frac{t}{2}\right] \\ &\leq \mathbb{P}_{g \sim \mathcal{N}(\mathbf{0}, I_m)}\left[|\|g\|^2 - m| \geq \frac{t}{2}\right]. \end{aligned} \quad (2.3.10)$$

The last bound is independent of x , so we have

$$\mathbb{P}[\|\Delta\| \geq t] \leq 9^d \cdot \mathbb{P}_{g \sim \mathcal{N}(\mathbf{0}, I_m)}\left[|\|g\|^2 - m| \geq \frac{t}{2}\right]$$

To the remaining probability we can apply Corollary 1.2.7, which it will be convenient to do specifically in the form (1.2.6). We can also bound $9 \leq \exp(3)$. Together, these observations give

$$\leq 2 \exp\left(3d - \frac{1}{8} \min\left\{\frac{t^2}{4m}, \frac{t}{2}\right\}\right) \quad (2.3.11)$$

Now, let us take $t := C\sqrt{dm}$ for some C to be chosen momentarily. Noting that $\sqrt{dm} \geq d$ since $d \leq m$, we may further bound

$$\begin{aligned} &\leq 2 \exp \left(3d - \frac{1}{8} \min \left\{ \frac{C^2}{4}d, \frac{C}{2}d \right\} \right) \\ &= 2 \exp \left(d \cdot \left[3 - \min \left\{ \frac{C^2}{32}, \frac{C}{16} \right\} \right] \right) \end{aligned} \quad (2.3.12)$$

and finally we see that taking $C := 64$ gives the result. $\square$

To give a corollary in a form that is maybe easier and more intuitive to understand concerning the singular values of G , we will use the following basic eigenvalue perturbation inequalities.

Proposition 2.3.10. Let $A, \Delta \in \mathbb{R}_{\text{sym}}^{d \times d}$. Then,

$$\lambda_d(A) - \|\Delta\| \leq \lambda_d(A + \Delta) \leq \lambda_1(A + \Delta) \leq \lambda_1(A) + \|\Delta\|. \quad (2.3.13)$$

Proof. For both bounds, we use the variational description of the eigenvalues. For the upper bound, we have

$$\begin{aligned} \lambda_1(A + \Delta) &= \sup_{x \in \mathbb{S}^{d-1}} x^\top (A + \Delta)x \\ &\leq \sup_{x \in \mathbb{S}^{d-1}} (x^\top Ax + \|\Delta\|) \\ &= \left(\sup_{x \in \mathbb{S}^{d-1}} x^\top Ax \right) + \|\Delta\| \\ &= \lambda_1(A) + \|\Delta\|. \end{aligned}$$

Similarly, for the lower bound,

$$\begin{aligned} \lambda_d(A + \Delta) &= \inf_{x \in \mathbb{S}^{d-1}} x^\top (A + \Delta)x \\ &\geq \inf_{x \in \mathbb{S}^{d-1}} (x^\top Ax - \|\Delta\|) \\ &= \left(\inf_{x \in \mathbb{S}^{d-1}} x^\top Ax \right) - \|\Delta\| \\ &= \lambda_d(A) - \|\Delta\|. \end{aligned} \quad \square$$

These are special cases of *Weyl's inequalities*, which are a good more general tool to be aware of. The proof of the more general inequalities similarly uses a more general variational form of arbitrary eigenvalues given by the *Courant-Fischer min-max theorem*.

Corollary 2.3.11. For any $d \leq m$,

$$\mathbb{P} \left[\sqrt{m} - 64\sqrt{d} \leq \sigma_d(\mathbf{G}) \leq \sigma_1(\mathbf{G}) \leq \sqrt{m} + 32\sqrt{d} \right] \geq 1 - 2\exp(-d).$$

Proof. We will use the basic bounds

$$\sqrt{1+x} \leq 1 + \frac{x}{2} \text{ for all } x \geq 0, \quad (2.3.14)$$

$$\sqrt{1-x} \geq 1 - x \text{ for all } 0 \leq x \leq 1. \quad (2.3.15)$$

Both may be viewed as consequences of the concavity of the square root function: the first says that one affine transformation of the square root lies below a tangent line to it, while the second says that another affine transformation lies above a secant segment connecting two points on its graph. Then, on the event that $\|\Delta\| \leq 64\sqrt{dm}$ from Theorem 2.3.9, using that $\mathbf{G}\mathbf{G}^\top = m\mathbf{I}_d + \Delta$, we have

$$\begin{aligned} \sigma_1(\mathbf{G}) &= \sqrt{\lambda_1(\mathbf{G}\mathbf{G}^\top)} \\ &\leq \sqrt{m + \|\Delta\|} && \text{(Proposition 2.3.10)} \\ &\leq \sqrt{m + 64\sqrt{dm}} \\ &= \sqrt{m} \cdot \sqrt{1 + 64\sqrt{\frac{d}{m}}} \\ &\leq \sqrt{m} \cdot \left(1 + 32\sqrt{\frac{d}{m}} \right) && (2.3.14) \\ &= \sqrt{m} + 32\sqrt{d}. \end{aligned}$$

Likewise, for $\sigma_d(\mathbf{G})$ we have a similar argument. First, though, note that if $m - 64\sqrt{dm} < 0$, then $\sqrt{m} - 64\sqrt{d} < 0$ as well, so the result holds vacuously in this case. Otherwise,

$$\begin{aligned} \sigma_d(\mathbf{G}) &= \sqrt{\lambda_d(\mathbf{G}\mathbf{G}^\top)} \\ &\geq \sqrt{m - \|\Delta\|} && \text{(Proposition 2.3.10)} \\ &\geq \sqrt{m - 64\sqrt{dm}} \\ &= \sqrt{m} \cdot \sqrt{1 - 64\sqrt{\frac{d}{m}}} \\ &\geq \sqrt{m} \cdot \left(1 - 32\sqrt{\frac{d}{m}} \right) && (2.3.15) \\ &= \sqrt{m} - 32\sqrt{d}, \end{aligned}$$

completing the proof. □

Of course, what you should remember is just the softer statement that, when $d \leq m$, the singular values of $\mathbf{G}$ are, with high probability, all in the range

$\sqrt{m} - O(\sqrt{d}) \leq \sigma_d(\mathbf{G}) \leq \dots \leq \sigma_1(\mathbf{G}) \leq \sqrt{m} + O(\sqrt{d})$

(2.3.16)

We highlight this because it is a fundamental fact about the spectra of random matrices that you should try to commit to your intuition. We will see considerably more precise elaborations of this claim soon, too. Note that, indeed, once $m \gg d$, then all of these singular values are $\sqrt{m}$ to leading order and $\mathbf{G}\mathbf{G}^\top \approx m\mathbf{I}_d$, as we proposed earlier.

The top d eigenvalues of the larger matrix $\mathbf{G}^\top \mathbf{G} \in \mathbb{R}_{\text{sym}}^{m \times m}$ are also then approximately m ; however, for this matrix, there are $m - d$ more eigenvalues that are identically zero. Thus, the eigenvalues of $\frac{1}{m}\mathbf{G}^\top \mathbf{G}$ are all either zero or close to 1; that is, it is close to a *projection matrix*. It is then natural to ask: a projection to what subspace? To answer this question we must understand the singular vectors of $\mathbf{G}$, which we will study next.

2.4 SINGULAR VECTORS

While the ordered singular values $\sigma_1(\mathbf{G}) \geq \dots \geq \sigma_d(\mathbf{G})$ are well-defined for any $\mathbf{G}$, the individual singular vectors are not always well-defined. Suppose we take the singular value decomposition (SVD),

$$\mathbf{G} = \mathbf{U}\mathbf{\Sigma}\mathbf{V}^\top = \sum_{i=1}^d \sigma_i \mathbf{u}_i \mathbf{v}_i^\top,$$

for $\{\mathbf{u}_1, \dots, \mathbf{u}_d\} \subset \mathbb{R}^d$ and $\{\mathbf{v}_1, \dots, \mathbf{v}_d\} \subset \mathbb{R}^m$ orthonormal sets. When are we justified in viewing $\mathbf{u}_i(\mathbf{G})$ and $\mathbf{v}_i(\mathbf{G})$ as well-defined functions?

Actually, the answer is *never*, because the SVD does not disambiguate between the pair of singular vectors $(\mathbf{u}_i, \mathbf{v}_i)$ and the pair $(-\mathbf{u}_i, -\mathbf{v}_i)$; either pair would give rise to the same decomposition. To get around this minor issue, we may instead consider the rank-one matrices associated to the pairs of singular vectors, $\mathbf{Z}_i = \mathbf{u}_i \mathbf{v}_i^\top$. These matrices are unchanged by a simultaneous change of sign, and from them we may also recover the more geometrically meaningful projection matrices

$$\mathbf{L}_i := \mathbf{Z}_i \mathbf{Z}_i^\top = \mathbf{u}_i \mathbf{u}_i^\top, \quad (2.4.1)$$

$$\mathbf{R}_i := \mathbf{Z}_i^\top \mathbf{Z}_i = \mathbf{v}_i \mathbf{v}_i^\top. \quad (2.4.2)$$

It is sometimes useful for the sake of intuition to identify a projection matrix with the subspace it projects to; thus, $\mathbf{u}_i \mathbf{u}_i^\top$ and $\mathbf{v}_i \mathbf{v}_i^\top$ describe not the singular *vectors* but the entire *lines* in those directions. $\mathbf{Z}_i$ contains the information of this pair of lines, and its being one of the rank-one matrices in the SVD says that $\mathbf{G}$ maps the line spanned by $\mathbf{v}_i$ to that spanned by $\mathbf{u}_i$. $\mathbf{Z}_i$ also contains one more bit of information, saying whether that mapping maps $\mathbf{v}_i$ to a positive or negative multiple of $\mathbf{u}_i$.

However, there are more nuances. For instance, if $\sigma_i = 0$, then $\mathbf{Z}_i$ is equivalent to $-\mathbf{Z}_i$ in the SVD. And, more substantially, if $\sigma_i = \sigma_j$, even if neither are zero, then $\mathbf{Z}_i$ and $\mathbf{Z}_j$ are not uniquely determined. Indeed, writing

$$\tilde{\mathbf{U}} := \begin{bmatrix} | & | \\ \mathbf{u}_i & \mathbf{u}_j \\ | & | \end{bmatrix} \in \mathbb{R}^{d \times 2}, \quad \tilde{\mathbf{V}} := \begin{bmatrix} | & | \\ \mathbf{v}_i & \mathbf{v}_j \\ | & | \end{bmatrix} \in \mathbb{R}^{m \times 2},$$

we have that the sum of the two corresponding terms in the SVD is

$$\sigma_i(\mathbf{Z}_i + \mathbf{Z}_j) = \sigma_i(\mathbf{u}_i \mathbf{v}_i^\top + \mathbf{u}_j \mathbf{v}_j^\top) = \sigma_i \tilde{\mathbf{U}} \tilde{\mathbf{V}}^\top = \sigma_i (\tilde{\mathbf{U}} \mathbf{Q})(\tilde{\mathbf{V}} \mathbf{Q})^\top$$

for any $\mathbf{Q} \in \mathcal{O}(2)$. This gives rise to equivalent decompositions of $\mathbf{Z}_i + \mathbf{Z}_j$ as a sum of rank-one matrices. Since $\sigma_i = \sigma_j$, the expression $\sigma_i(\mathbf{Z}_i + \mathbf{Z}_j)$ is what appears in the SVD for these two

components, so these equivalent decompositions mean that Z_i and Z_j are not uniquely determined in this case.

As an extreme example, consider the case $d = m$ and $G = I_d$, the identity matrix. All d of its singular values are 1, and it admits many SVDs of the form $G = \sum_{i=1}^d u_i u_i^\top$ for *any* orthonormal basis $\{u_1, \dots, u_d\}$. Thus, the more repeated singular values there are, the more dire this issue becomes. However, as the following fact of linear algebra shows, these are the only issues.

Proposition 2.4.1. Suppose that $G \in \mathbb{R}^{d \times m}$ with $d \leq m$, $\sigma_d(G) > 0$, and the singular values $\sigma_1(G), \dots, \sigma_d(G)$ are distinct. Then, there exist unique $Z_i = Z_i(G) \in \mathbb{R}^{d \times m}$ such that each Z_i has rank one, has $\sigma_1(Z_i) = 1$, and

$$G = \sum_{i=1}^d \sigma_i(G) Z_i(G).$$

2.4.1 DISTINCTNESS OF SINGULAR VALUES

Remarkably, the conditions of the above Proposition 2.4.1 hold not merely with high probability but *almost surely* for Gaussian random matrices (and, though we will not discuss it here, for other matrices whose entries are independent random variables whose laws are absolutely continuous with respect to the Lebesgue measure).

Theorem 2.4.2. Let $G \sim \mathcal{N}(0, 1)^{\otimes d \times m}$ with $d \leq m$. Then, almost surely, the following hold:

1. $\sigma_d(G) > 0$.
2. The singular values $\sigma_1(G), \dots, \sigma_d(G)$ are distinct, i.e., we have the strict inequalities $\sigma_1(G) > \dots > \sigma_d(G)$.

We will use the following result to establish both parts, which should be intuitive but is a little bit clunky to prove. You can prove it as an exercise—one approach is to show the more geometric fact that any hypersurface has Lebesgue measure zero.

Lemma 2.4.3. Suppose $p \in \mathbb{R}[x_1, \dots, x_N]$ is a non-constant polynomial, and $g \sim \mathcal{N}(0, I_N)$. Then, for any $a \in \mathbb{R}$, $\mathbb{P}[p(g) = a] = 0$.

The first part of Theorem 2.4.2 can be proved quite directly using this result, as follows.

Proof of Theorem 2.4.2, Claim 1. Note that we have $\sigma_d(G) = 0$ if and only if G is rank-deficient, which occurs if and only if $P(G) := \det(GG^\top) = 0$. (Note that this is not the case if we take $\det(G^\top G)$ instead, which is zero whenever $d < m$.) It is easy to show that P is a non-constant polynomial of the dm entries of the matrix argument: it is non-constant because, for example, for rank-deficient matrix inputs it is zero while for full-rank inputs it is non-zero. Thus $\mathbb{P}[\sigma_d(G) = 0] = \mathbb{P}[P(G) = 0] = 0$ by Lemma 2.4.3. $\square$

For the second part we will need to establish some tools from the theory of symmetric polynomials as well as a relationship between polynomials of the *eigenvalues* of a matrix and polynomials of the *entries* of a matrix. First, the following are the basic notions of the general theory of symmetric polynomials.

Definition 2.4.4. A polynomial $f \in \mathbb{R}[t_1, \dots, t_d]$ is *symmetric* if, for all permutations $\pi \in S_d$, we have $f(t_{\pi(1)}, \dots, t_{\pi(d)}) = f(t_1, \dots, t_d)$.

Theorem 2.4.5 (Fundamental theorem of symmetric polynomials). Define the *power sum* polynomials

$$s_k(t_1, \dots, t_d) := \sum_{i=1}^d t_i^k.$$

Then, if $f \in \mathbb{R}[t_1, \dots, t_d]$ is a symmetric polynomial, then there exists some polynomial g of d variables such that

$$f(t_1, \dots, t_d) = g(s_1(t_1, \dots, t_d), \dots, s_d(t_1, \dots, t_d)).$$

In words, this result says that any symmetric polynomial is some polynomial of the power sum polynomials, or that the power sum polynomials *algebraically generate* the symmetric polynomials (which form a subring of the ring of all polynomials).

The power sum polynomials have an important role in random matrix theory, because they serve as a bridge between the properties of the entries of a matrix and the properties of the eigenvalues, a connection that we will see much more sophisticated uses of later. The basic underlying fact is quite simple, however. Suppose here and below that $\mathbf{X} \in \mathbb{R}_{\text{sym}}^{d \times d}$. Let us define the *power trace* polynomials of $\mathbf{X}$,

$$S_k(\mathbf{X}) := \text{Tr}(\mathbf{X}^k) = \sum_{a_1, \dots, a_k=1}^d X_{a_1 a_2} X_{a_2 a_3} \cdots X_{a_{k-1} a_k} X_{a_k a_1}.$$

The latter expression shows that $S_k(\mathbf{X})$ is a polynomial of the entries of $\mathbf{X}$.

Proposition 2.4.6. $s_k(\lambda_1(\mathbf{X}), \dots, \lambda_d(\mathbf{X})) = S_k(\mathbf{X})$.

This is just a simple consequence of the fact that the eigenvalues of $\mathbf{X}^k$ are the $\lambda_i(\mathbf{X})^k$. However, it leads to the following important corollary.

Corollary 2.4.7. Suppose that $F(\mathbf{X}) = f(\lambda_1(\mathbf{X}), \dots, \lambda_d(\mathbf{X}))$ for f a symmetric polynomial (i.e., $F(\mathbf{X})$ is a symmetric polynomial of the eigenvalues of $\mathbf{X}$). Then, $F(\mathbf{X})$ is also a (not necessarily symmetric) polynomial of the entries of $\mathbf{X}$.

Remark 2.4.8. The assumption that f is symmetric is crucial. For instance, you may check as an exercise that $F(\mathbf{X}) = \lambda_1(\mathbf{X})$ is not a polynomial of $\mathbf{X}$.

With these tools, we are now ready for the proof of the second part of the main Theorem.

Proof of Theorem 2.4.2, Claim 2. Consider first a function of $\mathbf{X} \in \mathbb{R}_{\text{sym}}^{d \times d}$, defined as:

$$\begin{aligned} f(t_1, \dots, t_d) &:= \prod_{1 \leq i < j \leq d} (t_i - t_j)^2, \\ F(\mathbf{X}) &:= f(\lambda_1(\mathbf{X}), \dots, \lambda_d(\mathbf{X})) \\ &= \prod_{1 \leq i < j \leq d} (\lambda_i(\mathbf{X}) - \lambda_j(\mathbf{X}))^2. \end{aligned}$$

Clearly we have that $\mathbf{X}$ has a repeated eigenvalue if and only if $F(\mathbf{X}) = 0$. Further, $F(\mathbf{X})$ satisfies the conditions of Corollary 2.4.7. Therefore, $F(\mathbf{X})$ is a polynomial of the entries of $\mathbf{X}$. (Note that, remarkably, we can draw this conclusion non-constructively from the tools above.)

Now, returning to the setting of the Theorem, we see that $\mathbf{G}$ has a repeated singular value if and only if $P(\mathbf{G}) = 0$, where $P(\mathbf{G}) := F(\mathbf{G}\mathbf{G}^\top)$. As before, P is non-constant merely because some matrices do have repeated singular values and some do not. Thus, we may conclude as before by Lemma 2.4.3 that

$$\mathbb{P}[\mathbf{G} \text{ has a repeated singular value}] = \mathbb{P}[P(\mathbf{G}) = 0] = 0,$$

as claimed. $\square$

Remark 2.4.9. The proof above shows that the set of symmetric matrices with a repeated eigenvalue is an algebraic set in the space $\mathbb{R}_{\text{sym}}^{d \times d}$, which may be identified with $\mathbb{R}^{d(d+1)/2}$. *A priori*, since we described this set as $\{\mathbf{X} : F(\mathbf{X}) = 0\}$ for a non-zero polynomial F , we only know that its codimension is at least 1 (i.e., the dimension of this set is at most $\frac{d(d+1)}{2} - 1$). Actually, its codimension is exactly 2. For example, the dimension of $\mathbb{R}_{\text{sym}}^{2 \times 2}$ is $\frac{2(2+1)}{2} = 3$, while the 2×2 matrices with a repeated eigenvalue are merely multiples of the identity, whose dimension is clearly $1 = 3 - 2$ (indeed, these matrices form a line in the space of all symmetric matrices). Note that one polynomial constraint can in fact encode several, for instance if $F(\mathbf{X}) = F_1(\mathbf{X})^2 + F_2(\mathbf{X})^2 = 0$, then $F_1(\mathbf{X}) = F_2(\mathbf{X}) = 0$. As discussed for instance by [BKL18], a more complicated version of this sum-of-squares structure is indeed what happens in the case of $F(\mathbf{X})$.

We have established that, on an almost sure event, we are justified in speaking of *the* singular vectors of a Gaussian random matrix, at least via the rank-one matrices $\mathbf{Z}_i(\mathbf{G})$ and the projections to the left and right singular vectors $\mathbf{L}_i(\mathbf{G})$ and $\mathbf{R}_i(\mathbf{G})$. We are now ready to study the laws of these objects.

2.4.2 ORTHOGONAL INVARIANCE FOR MATRICES AND HAAR MEASURES

We will see that orthogonal invariance, as we discussed for vectors earlier, is crucial to understanding these distributions. The appropriate definitions for random matrices are as follows:

Definition 2.4.10 (Orthogonal invariance of matrices). We say that a random matrix $\mathbf{X} \in \mathbb{R}^{d \times m}$ (or its law $\text{Law}(\mathbf{X})$) is:

- *Left-orthogonally invariant* if, for each (deterministic) orthogonal matrix $\mathbf{Q} \in \mathcal{O}(d)$, we have $\text{Law}(\mathbf{Q}\mathbf{X}) = \text{Law}(\mathbf{X})$.
- *Right-orthogonally invariant* if, for each (deterministic) orthogonal matrix $\mathbf{Q} \in \mathcal{O}(m)$, we have $\text{Law}(\mathbf{X}\mathbf{Q}) = \text{Law}(\mathbf{X})$.
- *Bi-orthogonally invariant* if it is both left- and right-orthogonally invariant.

Further, we say that a random symmetric matrix $\mathbf{X} \in \mathbb{R}_{\text{sym}}^{d \times d}$ is *symmetric-orthogonally invariant* if $\text{Law}(\mathbf{Q}\mathbf{X}\mathbf{Q}^\top) = \text{Law}(\mathbf{X})$ for each $\mathbf{Q} \in \mathcal{O}(d)$.

The following property is then simple to check using either explicit mean and covariance calculations or by using that the rows and columns of an i.i.d. Gaussian random matrix are i.i.d. standard Gaussian vectors, which are themselves orthogonally invariant.

Proposition 2.4.11. $G \sim \mathcal{N}(0, 1)^{\otimes d \times m}$ is bi-orthogonally invariant.

Corollary 2.4.12. If $G \sim \mathcal{N}(0, 1)^{\otimes d \times m}$, then $L_i(G)$ and $R_i(G)$ are both symmetric-orthogonally invariant, for each $i = 1, \dots, d$. Further, the stronger equalities of law hold that

$$\begin{aligned} \text{Law}\left((QL_1(G)Q^\top, \dots, QL_d(G)Q^\top)\right) &= \text{Law}\left((L_1(G), \dots, L_d(G))\right) \quad \text{for all } Q \in \mathcal{O}(d), \\ \text{Law}\left((QR_1(G)Q^\top, \dots, QR_d(G)Q^\top)\right) &= \text{Law}\left((R_1(G), \dots, R_d(G))\right) \quad \text{for all } Q \in \mathcal{O}(m). \end{aligned}$$

The following basic linear algebra observation is a useful preliminary to the proof.

Proposition 2.4.13 (Invariance and equivariance of SVD). Let $G \in \mathbb{R}^{d \times m}$ be a (deterministic) matrix such that $\sigma_1(G) > \dots > \sigma_d(G) > 0$. Let $Q_L \in \mathcal{O}(d)$ and $Q_R \in \mathcal{O}(m)$. Then,

$$\begin{aligned} \sigma_i(Q_L G Q_R^\top) &= \sigma_i(G), \\ Z_i(Q_L G Q_R^\top) &= Q_L Z_i(G) Q_R^\top, \\ L_i(Q_L G Q_R^\top) &= Q_L L_i(G) Q_L^\top, \\ R_i(Q_L G Q_R^\top) &= Q_R R_i(G) Q_R^\top. \end{aligned}$$

Proof of Corollary 2.4.12. Let us prove the simpler claim about the individual L_i : we have, by combining the above results,

$$\begin{aligned} \text{Law}(QL_i(G)Q^\top) &= \text{Law}(L_i(QG)) && \text{(Proposition 2.4.13)} \\ &= \text{Law}(L_i(G)). && \text{(Proposition 2.4.11)} \end{aligned}$$

The result for the individual R_i and the collections of all L_i and R_i follows similarly. $\square$

Let us comment on the intuition behind what we have derived so far. Consider the $L_i(G)$. Recall that, for u_i some choice of the left singular vectors in the SVD (which for G are almost surely well-defined up to a sign flip), $L_i(G) = u_i u_i^\top$. As we have discussed, we may identify this rank-one projection matrix with the line formed by $\text{span}(u_i)$. Thus the collection $(L_1(G), \dots, L_d(G))$ is a collection of d orthogonal lines, i.e., a coordinate system for $\mathbb{R}^d$ consisting of d orthogonal axes, although without a “positive direction” chosen for any of the axes. The above result says that this object for $G \sim \mathcal{N}(0, 1)^{\otimes d \times m}$ has an orthogonally invariant distribution. We will next try to formally explain in what sense this makes $(L_1(G), \dots, L_d(G))$ have the law of the unique “uniformly random” choice of coordinate system.

To do this, we will show that after all it is possible to give a meaning to the left singular vectors “ $u_i(G)$ ” (and similarly for the right singular vectors), provided we allow for some further randomization.

Definition 2.4.14. Suppose $G \sim \mathcal{N}(0, 1)^{\otimes d \times m}$ with $d \leq m$. $Z_i = Z_i(G) \in \mathbb{R}^{d \times m}$ is (almost surely) a well-defined rank-one matrix with $\sigma_1(Z_i) = 1$. There then exist exactly two pairs of unit vectors $(u_i^{(1)}, v_i^{(1)})$, $(u_i^{(2)}, v_i^{(2)})$ such that $u_i^{(a)} v_i^{(a)\top} = Z_i$, which are related by $u_i^{(1)} = -u_i^{(2)}$ and $v_i^{(1)} = -v_i^{(2)}$. Let us define further random variables $(\tilde{u}_i, \tilde{v}_i) \sim \text{Unif}(\{(u_i^{(1)}, v_i^{(1)}), (u_i^{(2)}, v_i^{(2)})\})$.

These random variables have a well-defined law, and may be viewed as generated randomly from $\mathbf{G}$ along with one more bit of randomness to make the above choice. Finally, we define associated random matrices

$$\widetilde{\mathbf{U}} := \begin{bmatrix} | & & | \\ \tilde{\mathbf{u}}_1 & \cdots & \tilde{\mathbf{u}}_d \\ | & & | \end{bmatrix} \in \mathbb{R}^{d \times d}, \quad \widetilde{\mathbf{V}} := \begin{bmatrix} | & & | \\ \tilde{\mathbf{v}}_1 & \cdots & \tilde{\mathbf{v}}_d \\ | & & | \end{bmatrix} \in \mathbb{R}^{m \times d}.$$

The above is just a slightly elaborate way of formalizing that we break the sign ambiguity of $\mathbf{u}_i$ by choosing a random sign, and making that choice compatible between $\mathbf{u}_i$ and $\mathbf{v}_i$. It is easy to check that

$$\mathbf{G} = \sum_{i=1}^d \sigma_i(\mathbf{G}) \tilde{\mathbf{u}}_i \tilde{\mathbf{v}}_i^\top = \widetilde{\mathbf{U}} \boldsymbol{\Sigma}(\mathbf{G}) \widetilde{\mathbf{V}}^\top,$$

and we may also view the above as choosing a random such SVD of $\mathbf{G}$ among the 2^d equivalent ones. At a heuristic level, you may just think of

$$\widetilde{\mathbf{U}} := \begin{bmatrix} | & & | \\ \pm \mathbf{u}_1(\mathbf{G}) & \cdots & \pm \mathbf{u}_d(\mathbf{G}) \\ | & & | \end{bmatrix}, \quad \widetilde{\mathbf{V}} := \begin{bmatrix} | & & | \\ \pm \mathbf{v}_1(\mathbf{G}) & \cdots & \pm \mathbf{v}_d(\mathbf{G}) \\ | & & | \end{bmatrix},$$

with the signs chosen to make the above SVD hold.

We will now see how the following fact determines the laws of these matrices.

Proposition 2.4.15. The random matrices $\widetilde{\mathbf{U}}$ and $\widetilde{\mathbf{V}}$ are left-orthogonally invariant.

For $\widetilde{\mathbf{U}}$, first note that $\widetilde{\mathbf{U}} \in \mathcal{O}(d)$, i.e., this is an orthogonal matrix. The characterization of the law then follows from the next result.

Theorem 2.4.16 (Haar). There exists a unique probability measure supported on $\mathcal{O}(d)$ that is left- or right-orthogonally invariant, called the *Haar measure* and which we denote $\text{Haar}(\mathcal{O}(d))$. It is the same as the law of the matrix whose columns are the vectors formed by applying the Gram-Schmidt orthonormalization procedure to d i.i.d. random vectors sampled from either $\mathcal{N}(\mathbf{0}, \mathbf{I}_d)$ or $\text{Unif}(\mathbb{S}^{d-1})$.

Corollary 2.4.17. $\text{Law}(\widetilde{\mathbf{U}}) = \text{Haar}(\mathcal{O}(d))$.

For $\widetilde{\mathbf{V}}$ consisting of the right singular vectors, the above cannot quite work verbatim since this matrix is not square. Indeed, it belongs to a different special set of matrices, which is worth knowing about in its own right.

Definition 2.4.18 (Stiefel manifold). The *Stiefel manifold* of dimensions $m \geq d$, denoted $\text{Stief}(m, d)$, is the set

$$\text{Stief}(m, d) := \{\mathbf{V} \in \mathbb{R}^{m \times d} : \mathbf{V}^\top \mathbf{V} = \mathbf{I}_d\},$$

or equivalently the set of $m \times d$ matrices with orthonormal columns.

This is a very natural geometric structure: it is the set to which the matrix of right singular vectors in the SVD belongs, but also it has an important relationship to parametrizing the set of d -dimensional subspaces of $\mathbb{R}^m$ (an object called the *Grassmannian manifold*), since it may be viewed as the set of all possible orthonormal bases of such subspaces.

One practical way to think of the Stiefel manifold is as the set of matrices that can be extended to orthogonal matrices in $\mathcal{O}(m)$ by adding $m - d$ more columns. Therefore, the following extension of Theorem 2.4.16 should not be surprising.

Corollary 2.4.19. For any $1 \leq d \leq m$, there exists a unique probability measure supported on $\text{Stief}(m, d)$ that is left-orthogonally invariant, again called the Haar measure and which we denote $\text{Haar}(\text{Stief}(m, d))$. It is the same as the law of a matrix formed by sampling from $\text{Haar}(\mathcal{O}(m))$ and keeping only the submatrix formed by the first d columns.

Remark 2.4.20. Note that $\text{Stief}(m, 1) = \mathbb{S}^{m-1}$, and $\text{Haar}(\text{Stief}(m, 1)) = \text{Unif}(\mathbb{S}^{m-1})$, the uniform measure we discussed in Chapter 1.

Corollary 2.4.21. $\text{Law}(\tilde{\mathbf{V}}) = \text{Haar}(\text{Stief}(m, d))$.

We have gotten as close as we might hope to defining “the laws of the singular vectors” of $\mathbf{G}$: though those singular vectors are not well-defined, we have described well-defined randomized versions $\tilde{\mathbf{U}}$ and $\tilde{\mathbf{V}}$ of the matrices formed by those vectors and identified their laws.

2.4.3 CONNECTION TO RANDOM PROJECTIONS

Previously, we raised the issue of the behavior of $\mathbf{G}^\top \mathbf{G} \in \mathbb{R}_{\text{sym}}^{m \times m}$: this is a matrix of rank d , yet which “looks like” $d\mathbf{I}_m$ from the point of view of a small number of quadratic form evaluations. Using our findings about the singular values and vectors of $\mathbf{G}$, we may clarify this situation. Let $\Sigma = \Sigma(\mathbf{G}) \in \mathbb{R}_{\text{sym}}^{d \times d}$ be the diagonal matrix of singular values. If $m \gg d$, then $\Sigma \approx \sqrt{m}\mathbf{I}_d$.

We have $\mathbf{G} = \tilde{\mathbf{U}}\Sigma\tilde{\mathbf{V}}^\top$, so

$$\mathbf{G}^\top \mathbf{G} = \tilde{\mathbf{V}}\Sigma^2\tilde{\mathbf{V}}^\top \approx m\tilde{\mathbf{V}}\tilde{\mathbf{V}}^\top.$$

Since $\tilde{\mathbf{V}} \in \text{Stief}(m, d)$, $\tilde{\mathbf{V}}\tilde{\mathbf{V}}^\top =: \mathbf{P}_V$ is a projection matrix to the d -dimensional column space $V := \text{col}(\tilde{\mathbf{V}})$. By the left-orthogonal invariance of $\tilde{\mathbf{V}}$, this random d -dimensional subspace has law invariant under orthogonal transformations. You should think of this law as, in a similar spirit to the Haar measures above, the unique law of a “uniformly random” d -dimensional subspace of $\mathbb{R}^m$.

Now, consider a quadratic form: for $\mathbf{x} \in \mathbb{R}^m$,

$$\|\mathbf{G}\mathbf{x}\|^2 = \mathbf{x}^\top \mathbf{G}^\top \mathbf{G} \mathbf{x} \approx m\|\mathbf{P}_V \mathbf{x}\|^2.$$

We may view $\mathbf{P}_V \mathbf{x}$ as first expressing $\mathbf{x}$ in a random coordinate system, and then zeroing out all but d coordinates of $\mathbf{x}$ in this coordinate system. Since in a random coordinate system $\mathbf{x}$ should have entries of roughly the same size, we should expect

$$\|\mathbf{P}_V \mathbf{x}\|^2 \approx \frac{d}{m} \|\mathbf{x}\|^2,$$

and thus

$$\|Gx\|^2 \approx d\|x\|^2,$$

compatible with Corollary 1.2.7 applied to Gx .

Thus the geometric mechanism behind the above observation about quadratic forms and the Johnson-Lindenstrauss transform can be seen as the principle that *suitably rescaled random projections preserve the norms of deterministic vectors*. As we have discussed, while for any given projection there will exist some vectors that “escape” this behavior and are either very close to or nearly orthogonal to the range of a random projection, one must fix a very large number of deterministic vectors in order to expect to find such a vector with respect to a random projection. (Indeed, you may reflect on how this is basically the same phenomenon as needing exponentially many vectors to produce an ϵ -net of $\mathbb{S}^{d-1}$.)

2.5 APPLICATION: COMPRESSED SENSING

We will now see a few applications of the ideas developed so far. The first is to the problem of *compressed sensing*, recovering $x \in \mathbb{R}^m$ from $y = Gx \in \mathbb{R}^d$ with $G \in \mathbb{R}^{d \times m}$ (the same dimensions as before).

For general $x \in \mathbb{R}^m$, we need G to be injective, which requires $d \geq m$. Compressed sensing concerns making d much smaller, provided we are promised that x is *sparse*. In this setting, while we will discuss taking G random, we view ourselves as having control over G (the so-called *sensing matrix*) in general, so the sparsity assumption should be seen as in a basis of our choosing. Thus to take advantage of compressed sensing we should encode any prior knowledge of x in a basis that makes x sparse. (For example, if x encodes an image, we might choose a wavelet basis.)

2.5.1 NULL SPACE AND RESTRICTED ISOMETRY PROPERTIES

Let us denote sparsity by

$$\|x\|_0 := \#\{i \in [m] : x_i \neq 0\}. \quad (2.5.1)$$

Note that this “ ℓ^0 -norm” is not actually a norm, because it is not homogeneous in x .

Assuming that $\|x\|_0 \leq k$ makes the task of recovering x easier. At the extreme, if $k = 1$, then $y = Gx$ is just (up to scaling) one of the columns of G . Thus provided the columns of G are well-separated, it will be easy to recover x . You can show that it is possible to construct $\exp(\Omega(d))$ unit vectors in $\mathbb{R}^d$ that have, say, pairwise inner products of magnitude each at most $1/2$, which we may use as the columns of G (and such a choice will be robust to a small amount of noise). Thus when $k = 1$ then we may take d as small as $O(\log m)$ while still being able to recover any 1-sparse x .

How does the situation change for larger k ? We will show below that it actually does not change that much.

First, let us specify what we mean by x being “recoverable.” The following definition is not standard but useful.

Definition 2.5.1. We say that G *distinguishes k -sparse vectors* if $Gx \neq Gx'$ for any $x \neq x'$ with $\|x\|_0, \|x'\|_0 \leq k$.

If G distinguishes k -sparse vectors, then compressed sensing is *information-theoretically* possible: by, say, brute force search over an ϵ -net of sparse vectors followed by a rounding procedure, we may exactly recover any k -sparse x from $y = Gx$. We will not go into the *computational* feasibility of compressed sensing here, which is a deep area in its own right. You may look up the role of ℓ^1 -norm minimization for such algorithms to get started with computational approaches.

The following more linear-algebraic definition is actually equivalent to distinguishing k -sparse vectors.

Definition 2.5.2 (Null space property). We say that G has the *k -null space property (k -NSP)* if, for all $x \neq 0$ with $\|x\|_0 \leq k$, $Gx \neq 0$.

Proposition 2.5.3. G distinguishes k -sparse vectors if and only if G has the $2k$ -NSP.

Proof. If G does not distinguish k -sparse vectors, then there exist $x \neq x'$ with $\|x\|_0, \|x'\|_0 \leq k$ such that $Gx = Gx'$. In particular, $G(x - x') = 0$, and $\|x - x'\|_0 \leq 2k$, so G does not have the $2k$ -NSP. Conversely, if G does distinguish k -sparse vectors and $x'' \neq 0$ has $\|x''\|_0 \leq 2k$, then we may write $x'' = x - x'$ for $x \neq x'$ and $\|x\|_0, \|x'\|_0 \leq k$ (by partitioning its at most $2k$ non-zero indices into two sets of size at most k). Reversing the above argument then shows that G has the $2k$ -NSP. $\square$

We will in fact be able to show the following, more quantitative property.

Definition 2.5.4 (Restricted isometry property). We say that G has the *(k, δ) -restricted isometry property $((k, \delta)$ -RIP)* if, for all $x \neq 0$ with $\|x\|_0 \leq k$,

$$(1 - \delta)\|x\|^2 \leq \|Gx\|^2 \leq (1 + \delta)\|x\|^2. \quad (2.5.2)$$

The following implication is immediate.

Proposition 2.5.5. If G has the (k, δ) -RIP for any $\delta < 1$, then G has the k -NSP.

The RIP is more useful than the NSP both for analyzing compressed sensing when some noise is further added to $y = Gx$, and for analyzing algorithms for recovering x , but here we will just view it as an analytic tool for establishing the NSP.

2.5.2 RANDOM SENSING MATRICES

The following is the main result that we will show.

Theorem 2.5.6. For any $1 \leq k \leq m$ and $\delta \in (0, 1)$, if

$$d \geq \frac{160}{\delta^2} k \left(1 + \log \left(\frac{m}{k} \right) \right)$$

and $\mathbf{G} \sim \mathcal{N}(0, 1/d)^{\otimes d \times m}$, then

$$\mathbb{P}[\mathbf{G} \text{ has the } k\text{-NSP}] \geq \mathbb{P}[\mathbf{G} \text{ has the } (k, \delta)\text{-RIP}] \geq 1 - \frac{2}{\binom{m}{k}} \geq 1 - 2 \left(\frac{k}{m} \right)^k. \quad (2.5.3)$$

Note that this scaling of d is consistent with our earlier observation that $d \gtrsim \log m$ was the right condition for $k = 1$. Remarkably, the result implies that, while

$$d \geq m$$

is needed for *arbitrary* $\mathbf{x}$ to be recoverable from $\mathbf{G}\mathbf{x}$, for the specific case of k -sparse vectors it suffices to have only the far smaller

$$d \gtrsim k \log \left(\frac{m}{k} \right) \quad (2.5.4)$$

Proof. For $S \subseteq [m]$ with $|S| = k$ and $\mathbf{x} \in \mathbb{R}^m$, let $\mathbf{x}^{(S)} \in \mathbb{R}^k$ be the restriction of $\mathbf{x}$ to the indices in S . Likewise, for $\mathbf{G} \in \mathbb{R}^{d \times m}$, let $\mathbf{G}^{(S)} \in \mathbb{R}^{d \times k}$ be the restriction of $\mathbf{G}$ to the columns whose indices are in S . If $\|\mathbf{x}\|_0 \leq k$ and the non-zero indices of $\mathbf{x}$ are contained in such S , then

$$\|\mathbf{x}\|^2 = \|\mathbf{x}^{(S)}\|^2, \quad (2.5.5)$$

$$\mathbf{G}\mathbf{x} = \mathbf{G}^{(S)}\mathbf{x}^{(S)}. \quad (2.5.6)$$

We may then view the (k, δ) -RIP as requiring that, for all $S \subseteq [m]$ with $|S| = k$ and all $\mathbf{x}^{(S)} \in \mathbb{R}^k$, we have

$$(1 - \delta)\|\mathbf{x}^{(S)}\|^2 \leq \|\mathbf{G}^{(S)}\mathbf{x}^{(S)}\|^2 \leq (1 + \delta)\|\mathbf{x}^{(S)}\|^2. \quad (2.5.7)$$

But this just amounts to asking that

$$1 - \delta \leq \lambda_k(\mathbf{G}^{(S)\top} \mathbf{G}^{(S)}) \leq \lambda_1(\mathbf{G}^{(S)\top} \mathbf{G}^{(S)}) \leq 1 + \delta, \quad (2.5.8)$$

or again equivalently that

$$\|\mathbf{G}^{(S)\top} \mathbf{G}^{(S)} - \mathbf{I}_k\| \leq \delta. \quad (2.5.9)$$

We will then proceed by union bounding,

$$\mathbb{P}[\mathbf{G} \text{ does not have the } (k, \delta)\text{-RIP}] \leq \sum_{S \in \binom{[m]}{k}} \mathbb{P}[\|\mathbf{G}^{(S)\top} \mathbf{G}^{(S)} - \mathbf{I}_k\| > \delta]. \quad (2.5.10)$$

This is almost exactly the kind of deviation that we bounded before in Theorem 2.3.9. First, let's decode our statement into that form. Introduce $\mathbf{H} \sim \mathcal{N}(0, 1)^{\otimes k \times d}$, the scaling and aspect ratio that Theorem 2.3.9 addresses. We have $\text{Law}(\mathbf{G}^{(S)}) = \text{Law}(d^{-1/2} \mathbf{H}^\top)$. Thus,

$$\begin{aligned} \mathbb{P}[\|\mathbf{G}^{(S)\top} \mathbf{G}^{(S)} - \mathbf{I}_k\| > \delta] &= \mathbb{P}\left[\left\|\frac{1}{d} \mathbf{H} \mathbf{H}^\top - \mathbf{I}_k\right\| > \delta\right] \\ &= \mathbb{P}[\|\mathbf{H} \mathbf{H}^\top - d \mathbf{I}_k\| > \delta d]. \end{aligned}$$

The only difference between this and what Theorem 2.3.9 covered is that there we were concerned with deviations of the order $O(\sqrt{kd})$, while here we are concerned with $O(d)$, which is much larger since $d \gg k$ under our assumption. But, we may still apply the intermediate result (2.3.11) from that proof, which gives

$$\begin{aligned} &\leq 2 \exp \left(3k - \frac{1}{8} \min \left\{ \frac{\delta^2 d^2}{4d}, \frac{\delta d}{2} \right\} \right) \\ &= 2 \exp \left(3k - \frac{\delta^2}{32} d \right), \end{aligned} \tag{2.5.11}$$

using that $\delta < 1$ so $\delta^2/4 = (\delta/2)^2 < \delta/2$. Returning to the union bound, we first recall the non-asymptotic bounds on binomial coefficients

$$\left(\frac{m}{k} \right)^k \leq \binom{m}{k} \leq \left(\frac{em}{k} \right)^k,$$

which we will use twice below. We may bound the terms in the union bound uniformly as

$$\begin{aligned} \mathbb{P}[\mathbf{G} \text{ does not have the } (k, \delta)\text{-RIP}] &\leq \binom{m}{k} \cdot 2 \exp \left(3k - \frac{\delta^2}{32} d \right) \\ &\leq \left(\frac{em}{k} \right)^k \cdot 2 \exp \left(3k - \frac{\delta^2}{32} d \right) \\ &\leq 2 \exp \left(k \left(1 + \log \left(\frac{m}{k} \right) \right) + 3k - \frac{\delta^2}{32} d \right) \\ &\leq 2 \exp \left(4k \left(1 + \log \left(\frac{m}{k} \right) \right) - \frac{\delta^2}{32} d \right) \\ &\leq 2 \exp \left(-k \left(1 + \log \left(\frac{m}{k} \right) \right) \right) \\ &\leq 2 \left(\frac{em}{k} \right)^{-k} \\ &\leq \frac{2}{\binom{m}{k}} \end{aligned} \tag{2.5.12}$$

completing the proof. The alternate form of the bound follows from the lower bound on binomial coefficients above. $\square$

2.6 APPLICATION: NEURAL NETWORK INITIALIZATION

2.6.1 GENERAL SETUP

We next give a quite different application to the training of *neural network* machine learning models. While this is a toy example, as we will mention below the result of our derivation is actually exactly what is used in some real model-training software.

A neural network forms a sequence of representations of its input by composing linear maps with nonlinear functions. In the simplest form without any constraints, this is a function $f : \mathbb{R}^{d_{\text{in}}} \rightarrow \mathbb{R}^{d_{\text{out}}}$ of the form

$$f(x) = \rho(\mathbf{W}_L(\rho(\mathbf{W}_{L-1} \cdots \rho(\mathbf{W}_1 x))))$$

for a sequence of rectangular matrices $\mathbf{W}_1 \in \mathbb{R}^{m_1 \times d_{\text{in}}}$, $\mathbf{W}_2 \in \mathbb{R}^{m_2 \times m_1}$, $\dots$, $\mathbf{W}_L \in \mathbb{R}^{d_{\text{out}} \times m_{L-1}}$ for some intermediate dimensions $m_1, \dots, m_{L-1}$. Here $\rho : \mathbb{R} \rightarrow \mathbb{R}$ is some nonlinear function, which is applied coordinatewise to vectors. A common choice that we focus on here is the *rectified linear unit (ReLU)*,

$$\rho(x) = \max\{0, x\}.$$

This kind of structure, with no constraints on the $\mathbf{W}_i$, is called a *multilayer perceptron (MLP)*. The function f is *parametrized* by the $\mathbf{W}_i$, which we may emphasize by writing $\boldsymbol{\theta} = (\mathbf{W}_1, \dots, \mathbf{W}_L)$ and writing $f(\mathbf{x}) = f(\mathbf{x}; \boldsymbol{\theta})$.

In *training* of such a network, we are given a dataset $(\mathbf{x}_1, \mathbf{y}_1), \dots, (\mathbf{x}_M, \mathbf{y}_M)$ of pairs of inputs and outputs, and wish to find $\boldsymbol{\theta}$ such that $f(\mathbf{x}_i; \boldsymbol{\theta}) \approx \mathbf{y}_i$. To do this, we define a *loss function*, a common choice being the ℓ^2 loss,

$$L(\boldsymbol{\theta}) = \sum_{i=1}^M \|f(\mathbf{x}_i; \boldsymbol{\theta}) - \mathbf{y}_i\|^2.$$

We then attempt to minimize this over $\boldsymbol{\theta}$ by some form of local optimization such as gradient descent, stochastic gradient descent, and so forth. This yields a sequence of iterates $\boldsymbol{\theta}^{(0)}, \boldsymbol{\theta}^{(1)}, \dots$, starting from some *initialization* $\boldsymbol{\theta}^{(0)}$ that we choose at the beginning.

There are many interesting questions about what happens during neural network training, why and when it works well, and how best to configure it. We will return to some others later, but for now we focus on the basic matter of how to choose the initialization $\boldsymbol{\theta}^{(0)}$.

2.6.2 INITIALIZATION AND LOCAL ANALYSIS

Often the initial weight matrices are chosen to have i.i.d. Gaussian (or some other entry distribution) entries of various scales, $(\mathbf{W}_i)_{a,b} \sim \mathcal{N}(0, \sigma_i^2)$. This reduces our question to that of how to choose the σ_i^2 . Let us consider a single nonlinear layer $f(\mathbf{x}) = \rho(\mathbf{W}\mathbf{x})$ for $\mathbf{x} \in \mathbb{R}^d$, $\mathbf{W} \in \mathbb{R}^{m \times d}$, and therefore $f(\mathbf{x}) \in \mathbb{R}^m$, but bearing in mind that we plan to compose several such layers. (Here the rows of $\mathbf{W}$ correspond to outputs, so the matrix dimensions are reversed from our earlier convention for rectangular matrices $\mathbf{G}$.) We then have just one parameter σ^2 to choose to set $\mathbf{W} \sim \mathcal{N}(0, \sigma^2)^{\otimes m \times d}$.

What is it that we want this choice of σ^2 to achieve? The following two criteria are reasonable:

1. First, we want the layer to roughly preserve norms, $\|f(\mathbf{x})\|^2 \approx \|\mathbf{x}\|^2$. If this fails, then composing many layers will either shrink or blow up the scale of the data, which may lead to issues with numerical stability at the beginning of gradient descent (among other issues).
2. Second, we want the layer to distinguish different points, mapping far apart $\mathbf{x}_1$ and $\mathbf{x}_2$ to far apart $f(\mathbf{x}_1)$ and $f(\mathbf{x}_2)$. If this fails, then we will “forget” some differences in the input data, and this will be impossible to undo by composing with further layers. If, say, we eventually want to solve a classification problem, then we risk “collapsing” two or more classes together in a way that they cannot be distinguished.

Let us focus on the *first* layer, where the second concern is most significant. In light of it, we should certainly choose $m \geq d$, or else we risk losing lots of information present in the data. (Choosing a very large m is also an instance of *overparametrization* of our model, which has many other effects and is an important aspect of most effective neural network architectures.)

Because $f(\mathbf{x})$ is nonlinear, these questions are tricky (though sometimes still possible) to address directly and globally. Instead, we consider a local version of the same question: we fix some $\mathbf{x} \in \mathbb{R}^d$ and consider small perturbations. Since f is differentiable almost everywhere, we have

$$f(\mathbf{x} + \mathbf{h}) \approx f(\mathbf{x}) + (\mathbf{J}f)(\mathbf{x})\mathbf{h},$$

where $(Jf)(x) \in \mathbb{R}^{m \times d}$ is the *Jacobian matrix*,

$$(Jf)(x)_{ab} = \frac{\partial f_a}{\partial x_b}(x).$$

A natural condition on the Jacobian that will achieve the above criteria is if it acts as an approximately isometric embedding, i.e., if its d singular values are all close to 1 with high probability for a given fixed x .

2.6.3 JACOBIAN CALCULATIONS AND KAIMING INITIALIZATION

Fix a deterministic input $x \in \mathbb{R}^d \setminus \{0\}$. Write $w_a^\top$ for row a of W , and let

$$b_a := \mathbb{1}\{\langle w_a, x \rangle > 0\}.$$

Each $\langle w_a, x \rangle$ is a non-degenerate Gaussian random variable, so almost surely all of them are non-zero. On this event, f is differentiable at x , with Jacobian

$$(Jf)(x) =: J(x) = D_x W, \\ D_x := \text{Diag}(b_1, \dots, b_m).$$

Thus, the derivative keeps the rows corresponding to positive outputs and replaces the other rows by zero. We call those positive outputs *active*.

For all sufficiently small h , none of these signs changes, and we actually have exactly

$$f(x + h) = f(x) + J(x)h.$$

The singular values therefore bound the change in the output by

$$\sigma_d(J(x))\|h\| \leq \|f(x + h) - f(x)\| \leq \sigma_1(J(x))\|h\|, \quad (2.6.1)$$

justifying the previous claims in case $\sigma_1(J(x)), \sigma_d(J(x)) \approx 1$.

To move towards analyzing this matrix, let $B := \sum_{a=1}^m b_a$ be the number of active outputs. The variables b_a are independent with law $\text{Ber}(1/2)$, so $B \sim \text{Bin}(m, 1/2)$. There is a dependence to handle carefully: a row of W is kept precisely when its inner product with x is positive. Conditional on being kept, that row is constrained to a half-space. However, the singular values depend on the row only through its outer product, which has a useful symmetry.

Proposition 2.6.1. Conditional on $B = k$, the matrix $J(x)^\top J(x)$ has the same law as $\sigma^2 H^\top H$, where $H \sim \mathcal{N}(0, 1)^{\otimes k \times d}$. For $k = 0$, the latter matrix is understood to be zero.

Proof. Write $w_a = \sigma g_a$, so the g_a are independent standard Gaussian vectors in $\mathbb{R}^d$. For one such vector g , the transformation $g \mapsto -g$ preserves its law and the outer product $gg^\top$, while reversing the sign of $\langle g, x \rangle$. Thus, for any measurable set $\mathcal{A}$ of symmetric matrices,

$$\begin{aligned} \mathbb{P}[gg^\top \in \mathcal{A}, \langle g, x \rangle > 0] &= \mathbb{P}[gg^\top \in \mathcal{A}, \langle g, x \rangle < 0] \\ &= \frac{1}{2} \mathbb{P}[gg^\top \in \mathcal{A}] \\ &= \mathbb{P}[\langle g, x \rangle > 0] \cdot \mathbb{P}[gg^\top \in \mathcal{A}]. \end{aligned}$$

This proves that $\text{sgn}(\langle g, x \rangle)$ is independent of $gg^\top$. Independence of the rows then shows that the collection $(g_a g_a^\top)_{a=1}^m$ is independent of the entire vector $(b_a)_{a=1}^m$. Since

$$\mathbf{J}(x)^\top \mathbf{J}(x) = \sigma^2 \sum_{a=1}^m b_a g_a g_a^\top,$$

conditioning on any set of k active rows gives a sum of k independent outer products of standard Gaussian vectors. This law depends only on k , which proves the claim. $\square$

The Jacobian nonlinear layer has therefore led us back to a rectangular Gaussian matrix $\mathbf{H}$, but now with a random number of rows B . Since B is usually close to $m/2$, the bounds from earlier in this chapter suggest that all singular values of $\mathbf{J}(x)$ should be close to $\sigma\sqrt{m}/2$ in relative terms when m is much larger than d . We also obtain the exact identity

$$\mathbb{E}[\mathbf{J}(x)^\top \mathbf{J}(x)] = \sigma^2 \mathbb{E}[B] \mathbf{I}_d = \frac{m\sigma^2}{2} \mathbf{I}_d. \quad (2.6.2)$$

Thus, the correct choice, most immediately the choice that preserves the squared length of every fixed perturbation in expectation, is

$$\sigma^2 = \frac{2}{m}$$

We will now show that sufficiently large m indeed makes this apply.

Theorem 2.6.2. There is a constant $C > 0$ such that the following holds. Let $\sigma^2 = 2/m$, $\epsilon, \delta \in (0, 1)$, fix $x \in \mathbb{R}^d \setminus \{0\}$, and suppose

$$m \geq \frac{C}{\epsilon^2} \left(d + \log \frac{1}{\delta} \right).$$

Then, with probability at least $1 - \delta$,

$$\|\mathbf{J}(x)^\top \mathbf{J}(x) - \mathbf{I}_d\| \leq \epsilon. \quad (2.6.3)$$

Proof. By Hoeffding's lemma (Lemma 1.1.5), each $b_a - 1/2$ is $1/4$ -subgaussian. Lemma 1.1.3 therefore gives

$$\mathbb{P} \left[\left| B - \frac{m}{2} \right| > \frac{\epsilon m}{4} \right] \leq 2 \exp \left(-\frac{\epsilon^2 m}{8} \right).$$

On the complementary event, $m/4 \leq B \leq m$ and $|2B/m - 1| \leq \epsilon/2$. Taking C large enough ensures $m \geq 4d$, so $B \geq d$ on this event.

Since $\sigma^2 = 2/m$, Proposition 2.6.1 says that, conditional on $B = k$, the matrix $\mathbf{J}(x)^\top \mathbf{J}(x)$ has the same law as $(2/m)\mathbf{H}^\top \mathbf{H}$, where $\mathbf{H}$ is a $k \times d$ standard Gaussian matrix. For every integer $k \in [m/4, m]$, applying (2.3.11) to $\mathbf{H}^\top$ with $t = \epsilon m/4$ gives

$$\mathbb{P} \left[\left\| \mathbf{J}(x)^\top \mathbf{J}(x) - \frac{2k}{m} \mathbf{I}_d \right\| > \frac{\epsilon}{2} \mid B = k \right] \leq 2 \exp \left(3d - \frac{\epsilon^2 m}{512} \right).$$

Here we used $t^2/(4k) \geq \epsilon^2 m/64$ and $t/2 \geq \epsilon^2 m/64$, since $k \leq m$ and $\epsilon < 1$. When both bounds hold, the triangle inequality gives

$$\begin{aligned} \|\mathbf{J}(x)^\top \mathbf{J}(x) - \mathbf{I}_d\| &\leq \left\| \mathbf{J}(x)^\top \mathbf{J}(x) - \frac{2B}{m} \mathbf{I}_d \right\| + \left| \frac{2B}{m} - 1 \right| \\ &\leq \frac{\epsilon}{2} + \frac{\epsilon}{2} = \epsilon. \end{aligned}$$

The probability of failure is at most $2 \exp(-\epsilon^2 m/8) + 2 \exp(3d - \epsilon^2 m/512)$, which is at most δ under the assumed lower bound on m , for a sufficiently large constant C . $\square$

The calculation gives two distinct roles to width m and variance σ^2 . Large width makes the singular values close together, while choosing the “right” variance $\sigma^2 = 2/m$ places the singular values near 1. We note that the condition that the width be large is important: if, say, $m = d$, then the typical rank of $\mathbf{J}(\mathbf{x})$ is about $d/2$, and thus $\sigma_d(\mathbf{J}(\mathbf{x})) = 0$ usually. Working in an asymptotic regime of large width is a common assumption in the neural network literature and leads to many such limiting simplifications.

Theorem 2.6.2 concerns one fixed input at initialization, with simultaneous control over all perturbation directions at that input. A random input independent of the weights is also allowed, by conditioning on it, provided it is non-zero almost surely. Uniform control over inputs, or control after training changes the weights, requires further, considerably more complicated arguments. The broader question of making Jacobian singular values close to 1 is studied under the name *dynamical isometry*. See, for example, [PSG17].

The choice $\sigma^2 = 2/m$ is called the *fan-out* version of the *He* or *Kaiming initialization*,¹ as derived in [HZRS15, Section 2.2]. Precisely this normalization is used in actual implementations in various utilities in PyTorch and other neural network libraries.²

Remark 2.6.3 (Another criterion for the variance). Symmetry of the Gaussian distribution also gives $\mathbb{E}[\|\mathbf{f}(\mathbf{x})\|^2] = (m\sigma^2/2)\|\mathbf{x}\|^2$. Thus our choice $\sigma^2 = 2/m$ preserves the total squared length of the input in expectation. One could instead ask to preserve the average squared coordinate, meaning $\mathbb{E}[\|\mathbf{f}(\mathbf{x})\|^2]/m = \|\mathbf{x}\|^2/d$. This criterion gives $\sigma^2 = 2/d$, the *fan-in* version of the Kaiming initialization. For a wide layer, the theorem then places the singular values near $\sqrt{m/d}$, and in particular the norm itself can be dramatically inflated over several such layers. The choice of variance therefore depends on which quantity we want to preserve.

¹The last and first names of the first author of the reference.

²See, e.g., [here](#) and [here](#).

BIBLIOGRAPHY

- [BKL18] Paul Breiding, Khazhgali Kozhasov, and Antonio Lerario. On the geometry of the set of symmetric matrices with repeated eigenvalues. *Arnold Mathematical Journal*, 4(3):423–443, 2018.
- [HZRS15] Kaiming He, Xiangyu Zhang, Shaoqing Ren, and Jian Sun. Delving deep into rectifiers: Surpassing human-level performance on ImageNet classification. In *Proceedings of the IEEE International Conference on Computer Vision (ICCV)*, pages 1026–1034. IEEE, 2015.
- [PSG17] Jeffrey Pennington, Samuel S. Schoenholz, and Surya Ganguli. Resurrecting the sigmoid in deep learning through dynamical isometry: theory and practice. *arXiv preprint arXiv:1711.04735*, Nov 2017. Appeared at NeurIPS 2017.
- [RH17] Philippe Rigollet and Jan-Christian Hütter. High dimensional statistics, 2017. Lecture notes.