

Afonso S. Bandeira (ETH Zürich)
Amit Singer (Princeton University)
Thomas Strohmer (UC Davis)

Topics in Mathematics of Data Science

Preprint: version 1.0-

September 18, 2025

Contents

0	Notes on this Version and Current Status	1
1	Introduction	3
1.1	Origins	3
1.2	The data science revolution	3
1.3	Mathematics as a foundation of data science	4
1.4	Whom is this book for?	6
1.5	Organization	7
2	Curses, Blessings, and Surprises in High Dimensions	9
2.1	Geometry of spheres and balls in high dimension	10
2.2	Basic concepts from probability	17
2.3	Concentration of measure, a “blessing of dimensionality”	22
3	Singular Value Decomposition and Principal Component Analysis	33
3.1	Singular Value Decomposition	33
3.2	Principal Component Analysis and dimension reduction	42
3.3	PCA in high dimensions and Marčenko-Pastur law	48
4	Graphs, Networks, and Clustering	57
4.1	PageRank	57
4.2	Graph theory	60
4.3	Clustering	62
5	Nonlinear Dimension Reduction and Diffusion Maps	77
5.1	Diffusion Maps	77
5.2	Connections between Diffusion Maps and Spectral Clustering	85
5.3	Semi-supervised learning	88

6	Linear Dimension Reduction via Random Projections	95
6.1	The Johnson-Lindenstrauss Lemma	95
6.2	Randomized SVD	103
7	Community Detection and the Power of Convex Relaxation	111
7.1	Max-Cut, lifting, and approximation algorithms	111
7.2	Community Detection	118
8	Concentration of Measure and Gaussian Analysis	135
8.1	Gaussian integration by parts and Wick's formula	135
8.2	Gaussian interpolation, Poincaré, and concentration	138
8.3	Gaussian comparison principles	143
8.4	Gordon's Theorem	146
9	Matrix Concentration Inequalities	153
9.1	Non-commutative Khintchine inequality	153
9.2	Matrix concentration inequalities	159
9.3	Improvements leveraging intrinsic freeness	162
10	Compressive Sensing and Sparsity	169
10.1	Sparse recovery	169
10.2	Null space property and exact recovery	173
10.3	The Restricted Isometry Property	176
10.4	Duality and exact recovery	181
10.5	Compressive sensing: from theory to practice	184
	Bibliography	193

Notes on this Version and Current Status

This is a preprint of a book in preparation by the authors.

We anticipate the focus and content of this manuscript not to change drastically and it can already be used for a graduate course in Mathematics of Data Science; it has been used as such by the authors at their home institutions.

We welcome suggestions and comments, and would like to learn about any possible errors and typos.

Please contact the authors at bandeira@math.ethz.ch, strohmer@math.ucdavis.edu, or amits@math.princeton.edu. In fact, this version was improved using many suggestions sent from, and typos found by, many students and readers. Thank you!

Thank you,
Afonso, Thomas, and Amit.

Introduction

1.1 Origins

The idea of this book emerged from the classes on “mathematics of data science” that the authors have been teaching at their home institutions Princeton University (AS), UC Davis (TS), and MIT, NYU, and ETH (ASB)¹. The lecture notes eventually became the basis to this book. While there are specialized books and many advanced research papers on each and every topic covered in this book, we believe it fills in the gap of being a single source that provides an accessible entry point at an advanced undergraduate level to mathematics of data science. The book can serve readers as a segue to those more advanced and specialized topics. In that light, one of the main purposes of this book is to be used as a textbook at both the (advanced) undergraduate and graduate levels. This book also differs from most texts in the way that it interweaves mathematical theory with some of its applications in data science.

1.2 The data science revolution

The data science revolution is one of the most profound shifts in modern history, reshaping industries, economies, and societies with unprecedented speed and scale. At its core, the data science revolution is driven by the exponential growth in data, advances in computing power, and the development of sophisticated analytical techniques.

Data science and its cousin artificial intelligence have become a cornerstone of modern decision-making and innovation, penetrating almost all aspects of our life. In healthcare, for example, data science is driving personalized medicine and predictive analytics. By analyzing patient data, healthcare

¹Lecture notes from these classes also contain a list of mathematical open problems [29], which has now helped inspire a mathematical blog of open problems at <https://randomstrasse101.math.ethz.ch/>. The open problems of 2024 are available at [31].

providers can tailor treatments to individual needs, predict potential health issues before they arise, and improve patient outcomes. In finance, data science is revolutionizing risk management, fraud detection, and investment strategies. Algorithms analyze market trends, assess creditworthiness, and identify unusual transactions to mitigate risks and enhance decision-making.

Retailers leverage data science to understand consumer behavior, optimize supply chains, and personalize marketing efforts. The transportation industry benefits from data science through advancements in autonomous vehicles, route optimization, and predictive maintenance. Companies use data to enhance logistics, reduce costs, and improve safety. The energy sector uses data science to optimize resource allocation, predict equipment failures, and enhance energy efficiency. Data-driven insights help in managing renewable energy sources and reducing environmental impact.

From the statistical techniques that form the basis of data analysis to the complex algorithms that power machine learning and artificial intelligence, mathematics is at the heart of data science. Understanding the connection between data science and mathematics is crucial for appreciating the depth and breadth of this revolution.

1.3 Mathematics as a foundation of data science

Data science is an inherently interdisciplinary field that merges techniques from mathematics, statistics, computer science, signal processing, and domain-specific knowledge to extract insights from data. However, at its core, data science is fundamentally about understanding and leveraging patterns in data. This understanding is built on a complex and nuanced foundation of mathematical principles that enable the development of models and algorithms capable of making sense of complex datasets. This book aims to unravel many of these mathematical principles and provide a thorough grounding in the essentials of data science from a rigorous, mathematical perspective.

Mathematics provides the language and framework for formalizing the processes involved in data analysis, from data collection and preprocessing to modeling and evaluation. It offers the tools necessary for designing effective algorithms, making precise predictions, visualizing complex datasets, quantifying uncertainty, and interpreting the results of data-driven experiments. As the field of data science keeps evolving at a rapid speed, new algorithms and techniques are constantly being developed. A strong mathematical foundation enables practitioners to quickly understand and adopt these innovations because they can grasp the underlying principles rather than just following black-box implementations. Moreover, with a solid grasp of mathematics, learning new methods becomes easier.

Also, when algorithms fail to produce correct results in real world applications, we would like to know why they failed. Is it because of mistakes in the experimental setup, did we gather the wrong kind of data, do we have

insufficient data, is it because of corrupted measurements, calibration errors, or incorrect modeling assumptions, or perhaps due to fundamentally unrealistic expectations, or is it due to a deficiency of the algorithm itself? If it is the latter, can it be fixed by a better initialization, a more careful tuning of the parameters, or by choosing a different algorithm? Or is a more fundamental modification required, such as developing a different model, including additional prior information, taking more measurements, or a better compensation of calibration errors? How reliable are our algorithms to small changes in the input data and can we explain why an algorithm arrives at a certain prediction?

As we are dealing with larger datasets and we aim for faster and faster throughput, it becomes increasingly important to address the aforementioned challenges in a systematic and principled manner. Thus, a rigorous and thorough study of computational algorithms both from a theoretical and numerical viewpoint is not a luxury, but emerges as an imperative ingredient towards effective data-driven discovery.

While a thorough mathematical framework can never completely avoid some guessing and trial-and-error, it dramatically can reduce the time and resources we spend on trial-and-error. Good mathematical theory can help in predicting if a problem is solvable at all in feasible time. Moreover, fundamental mathematical analysis often leads to new discoveries such as faster or more robust algorithms, and it may reveal hitherto unknown connections to other, already well studied problems.

At the heart of data science lies a rich tapestry of mathematical concepts. Key areas include statistics and probability, linear algebra, optimization, numerical algorithms, analysis, and graph theory, but this list is by no means exhaustive. This book emphasizes the geometric viewpoint of data sets. Specifically, it is a very powerful way to think about a dataset as a point cloud in a high dimensional space for a wide range of analysis tasks such as clustering, classification, dimension reduction, prediction, denoising, outlier detection, and more. In this geometric viewpoint, linear models of data are interpreted as linear subspaces in a Euclidean space, while non-linear models are viewed as manifolds embedded in the ambient space. Equally powerful is the representation of data points and their interconnections and similarities as a graph and learning using graph algorithms and analysis. In particular, the graph representation provides a powerful way to extend classical methods for function approximation and interpolation to the setting of complex and high dimensional setting of modern data sets. The computational machinery provided by classical harmonic analysis on Euclidean spaces such as Fourier analysis can be extended in this way to more complex geometries and networks.

As data science continues to evolve, new mathematical techniques and approaches will emerge, expanding the possibilities for analyzing and interpreting data. This book provides a solid foundation upon which readers can build their expertise, adapt to new developments, and contribute to the advancement of the field.

Data science also opens a new playground for Mathematics, motivating new questions and new areas of mathematical inquiry. For this reason, another goal of this book is to motivate mathematicians (and, in particular, mathematics students) to engage with these topics and to create new mathematics motivated by questions arising from data and computation.

1.4 Whom is this book for?

Many of the concepts, algorithms, and methods presented in this book do not only belong in the toolbox of a well trained data scientist, but are also found in the standard repertoire of researchers and practitioners of machine learning and AI.

The purpose of this book is twofold. On the one hand, the book is designed to guide readers through many of the mathematical foundations of data science in a structured and accessible manner. It assumes a basic familiarity with mathematical concepts and aims to build upon this foundation to explore more advanced topics relevant to data science. On the other hand, the selection of data science topics for this book is also guided by our desire to demonstrate how data science leads to deep and exciting mathematical problems, and how (sometimes seemingly unrelated) mathematical concepts can provide key insight into data science tasks. This mutual enrichment is akin to how physics and mathematics have had a vibrant cross-fertilization over many centuries.

What prerequisites should a reader have? An undergraduate level knowledge of linear algebra and probability is assumed.

What is covered in this book and what not? This book does not cover everything related to mathematics of data science. Indeed, there are many topics not included to keep the book at a manageable length. Thus, while this book is not an encyclopedia of data science, it does contain the material that one might want to teach in either a year-long course or a one-semester course on the mathematical foundations of data science.

The book covers both linear and non-linear dimension reduction methods. The linear methods include the singular value decomposition (SVD) and principal component analysis (PCA) as well as random projections. The latter also serve as an important ingredient to the computationally efficient randomized SVD algorithm. Furthermore, the book covers the essential background in probability, in particular, concentration of measure and large deviation theory that are used to explain the effectiveness of convex relaxations in data science methodologies such as clustering in the stochastic block model and compressive sensing. In addition, the book focuses on spectral methods in data analysis, ranging from spectral clustering to Laplacian eigenmaps and diffusion maps for non-linear dimension reduction to the PageRank algorithm for ranking.

There are many mathematical topics related to data science that are not covered in this book. Perhaps the most notable omissions are relevant topics from statistics such as parameter estimation (e.g., maximum likelihood estimation, Bayesian inference, etc.) and optimization (first order methods, second order methods, stochastic optimization). There are already excellent textbooks dedicated to these topics, which are also being typically taught in separate courses. Such courses in optimization and statistics are of course recommended to those interested in data science, and can be taken either before, after, or in parallel to a course based on this book. Currently, the book also does not cover Deep Learning, although we plan to add an introductory chapter on Deep Learning in a future version.

1.5 Organization

The book is separated in ten chapters, an effort was made to intertwine mathematical techniques and motivating applications. It starts with a description of the curious phenomena of high dimensional geometry, taking the opportunity to introduce and recall some probability theory in Chapter 2. It proceeds to discussing linear dimension reduction techniques (such as PCA and truncated SVD) and doing a recap of several linear algebra concepts in Chapter 3. It moves on to graph theory (including spectral graph theory, and spectral clustering in Chapter 4) and non-linear dimension reduction (including diffusion maps in Chapter 5), continuing to highlight how central the spectrum of matrices formed from data are to understanding the data geometry. of said data. Chapter 6 focus on dimension reduction via randomized methods, including a brief introduction to randomized linear algebra. The concept and power of convex relaxations is illustrated via the community detection problem in Chapter 7. Chapters 8 and 9 are devoted to probability, the first to Gaussian analysis and concentration of measure, and the second to matrix concentration inequalities. Chapter 10 is dedicated to Compressive Sensing.

Acknowledgement

We want to thank our students, colleagues, and readers of earlier drafts of this book, who provided valuable feedback, alerted us to errors, and pointed out typos. Please keep doing so! We are grateful to Junda (Albie) Sheng for generating some figures and to Yuan Ni for providing code for some numerical simulations. Amit Singer acknowledges support from the Simons Foundation Math+X Investigator Award, AFOSR FA9550-20-1-0266, AFOSR FA9550-23-1-0249, NSF DMS 2009753, NSF DMS 2510039, and NIH/NIGMS R01GM136780-01. Thomas Strohmer acknowledges support from NSF DMS-2208356, NIH R01HL16351, P41EB032840, and DE-SC0023490.

Curses, Blessings, and Surprises in High Dimensions

Most of the material in this book will naturally “live” in the context of high dimensions, either motivated by the analysis of high dimensional data or by associating with networks a space whose dimension is the number of network nodes, to name a few. This chapter is about High Dimensions. It discusses the curse of dimensionality, but also many of its blessings. The first is caused by the exponential increase in volume associated with adding extra dimensions to Euclidean space. The latter is a manifestation of an intriguing phenomenon called the concentration of measure. This concentration phenomenon will give rise to many surprising facts about high dimensional geometry that we will discuss. Since several of the results discussed in this chapter require basic tools from probability, we will also review some fundamental probabilistic concepts.

The curse of dimensionality

The *curse of dimensionality* refers to the fact that many algorithmic approaches to problems in $\mathbb{R}^d$ become *exponentially* more difficult as the dimension d grows. The expression “curse of dimensionality” was coined by Richard Bellman to describe the problem caused by the exponential increase in volume associated with adding extra dimensions to Euclidean space [35].

For instance, if we want to sample the unit interval such that the distance between adjacent points is at most 0.01, then 100 evenly-spaced sample points suffice; an equivalent sampling of a five-dimensional unit hypercube with a grid with a spacing of 0.01 between adjacent points would require 10^{10} sample points. Thus, a modest increase in dimensions results in a dramatic increase in required data points to cover the space at the same density.

Intimately connected to the curse of dimensionality is the problem of *overfitting* and *underfitting*. Here, overfitting refers to the issue that an algorithm may show good performance on the training data, but poor generalization to other data. Underfitting in turn, corresponds to poor performance on the training data (and poor generalization to other data). This problem manifests itself in many machine learning algorithms.

We will discuss a toy example from image classification in more detail to illustrate the underlying issues. Assume we want to classify images into two groups, cars and bicycles, say. From the vast number of images depicting cars or bicycles, we are only able to obtain a small number of training images, say five images of cars and five images of bicycles. We want to train a simple linear classifier based on these ten labeled training images to correctly classify the remaining unlabeled car/bicycle images. We start with a simple feature, e.g. the amount of red pixels in each image. However, this is unlikely to give a linear separation of the training data. We add more features and eventually the training images become linearly separable. This might suggest that increasing the number of features until perfect classification of the training data is achieved, is a sound strategy. However, as we *linearly increase* the dimension of the feature space, the density of our training data *decreases exponentially* with the feature dimension.

In other words, to maintain a comparable density of our training data, we would need to increase the size of the dataset exponentially – the curse of dimensionality. Thus, we risk producing a model that could be very good at predicting the target class on the training set, but it may fail miserably when faced with new data. This means that our model does not *generalize* from the training data to the test data.

2.1 Geometry of spheres and balls in high dimension

When we peel an orange, then after having removed the rind we are still left with the majority of the orange. Suppose now we peel a d -dimensional orange for large d , then after removing the orange peel we would be left with essentially nothing. The reason for this – from a healthy nutrition viewpoint discouraging – fact is that for a d -dimensional unit ball almost all of its volume is concentrated near the boundary sphere. This is just one of many surprising phenomena in high dimensions. Many of these surprises are actually a manifestation of some form of concentration of measure that we will analyze in more detail in the next section (and then these surprises are no longer so startling).

When introducing data analysis concepts, we typically use only a few dimensions in order to facilitate visualization. However, our intuition about space, which is naturally based on two and three dimensions, can often be misleading in high dimensions. Many properties of even very basic objects become counterintuitive in higher dimensions. Understanding these paradoxical properties is essential in data analysis as it allows us to avoid pitfalls in the design of algorithms and statistical methods for high-dimensional data. It is therefore instructive to analyze the shape and properties of some basic geometric forms that we understand very well in dimensions two and three, in high dimensions.

To that end, we will look at some of the properties of the sphere and the cube as the dimension increases. The d -dimensional hyperball of radius R is defined by

$$B^d(R) = \{x \in \mathbb{R}^d : x_1^2 + \cdots + x_d^2 \leq R^2\},$$

the d -dimensional hypersphere (or d -sphere) of radius R is given by

$$S^{d-1}(R) = \{x \in \mathbb{R}^d : x_1^2 + \cdots + x_d^2 = R^2\},$$

and the d -dimensional hypercube with side length $2R$ is the subset of $\mathbb{R}^d$ defined as the d -fold product of intervals $[-R, R]$:

$$C^d(R) = \underbrace{[-R, R] \times \cdots \times [-R, R]}_{d \text{ times}}.$$

If there is no danger of confusion, we may write B^d for the unit ball $B^d(1)$, S^{d-1} for the unit sphere $S^{d-1}(1)$, and C^d for the unit hypercube $C^d(\frac{1}{2})$.

High Dimensional Geometry and High Dimensional Probability

In the rest of this section we will establish a variety of (at first) surprising properties of high dimensional geometry. To illustrate the connections between geometry and probability we will derive some of the results first using analytical tools (computing volumes and integrals) and then using tools from Probability, namely concentration inequalities (which the reader might agree gives more elegant arguments). The goal here is not to get the best bounds, but to illustrate these properties and showcase the duality between geometry and probability.

2.1.1 Volume of the hyperball

Proposition 2.1. *The volume of $B^d(R)$ is given by*

$$\text{Vol}(B^d(R)) = \frac{\pi^{\frac{d}{2}} R^d}{\frac{d}{2} \Gamma(\frac{d}{2})}. \quad (2.1)$$

Proof. The volume of $B^d(R)$ is given by

$$\text{Vol}(B^d(R)) = \int_0^R s_d r^{d-1} dr = \frac{s_d R^d}{d}, \quad (2.2)$$

where s_d denotes the (hyper-)surface area of a unit d -sphere. A unit d -sphere must satisfy

$$s_d \int_0^\infty e^{-r^2} r^{d-1} dr = \underbrace{\int_{-\infty}^\infty \cdots \int_{-\infty}^\infty}_{d \text{ times}} e^{-(x_1^2 + \cdots + x_d^2)} dx_1 \cdots dx_d = \left(\int_{-\infty}^\infty e^{-x^2} dx \right)^d.$$

Recall that the Gamma function is given by

$$\Gamma(n) = \int_0^\infty r^{n-1} e^{-r} dr = 2 \int_0^\infty e^{-r^2} r^{2n-1} dr,$$

hence $\frac{1}{2}s_d\Gamma(\frac{d}{2}) = \left[\Gamma(\frac{1}{2})\right]^d = (\pi^{\frac{1}{2}})^d$, and thus $s_d = \frac{2\pi^{\frac{d}{2}}}{\Gamma(\frac{d}{2})}$. Plugging this expression into (2.2) gives

$$\text{Vol}(B^d(R)) = \frac{\pi^{\frac{d}{2}} R^d}{\frac{d}{2} \Gamma(\frac{d}{2})}. \quad (2.3)$$

□

Recall that for even d , we have $\Gamma(d/2) = (d/2 - 1)!$. To understand the asymptotic behavior $\text{Vol}(B^d(R))$ we can use the celebrated Stirling's Formula,

$$\Gamma(n) \approx \sqrt{\frac{2\pi}{n}} \left(\frac{n}{e}\right)^n$$

we obtain an approximation for the volume of the unit d -ball for large d

$$\text{Vol}(B^d) \approx \frac{1}{\sqrt{d\pi}} \left(\frac{2\pi e}{d}\right)^{\frac{d}{2}}, \quad (2.4)$$

where $\approx$ means that the quotient goes to 1 as the relevant parameter, in this case d , goes to ∞ .

Since the denominator in the parenthesis of equation (2.4) goes to infinity, the volume of the unit d -sphere goes rapidly to 0 as the dimension d increases to infinity, see also Figure 2.1. Notice that this would still be the case for any fixed radius R .

Thus, unit spheres in high dimensions have almost no volume—compare this to the unit cube, which has volume 1 in any dimension. For $B^d(R)$ to have volume equal to 1, its radius R must be approximately (asymptotically) equal to $\sqrt{\frac{d}{2\pi e}}$.

2.1.2 Concentration of the volume of a ball near its equator

If we take an orange and cut it into slices, then the slices near the center are larger since the sphere is wider there. This effect increases dramatically (exponentially with the dimension) with increasing dimension. Assume we want to cut off a slab around the “equator”¹ of the d -unit ball such that 99% of its volume is contained inside the slab. In two dimensions the width of the

¹To define the “equator” of a the d -dimensional ball, we need to pick a “north pole” as reference. Without loss of generality we could pick the unit vector in the x_1 -direction as defining “north”.

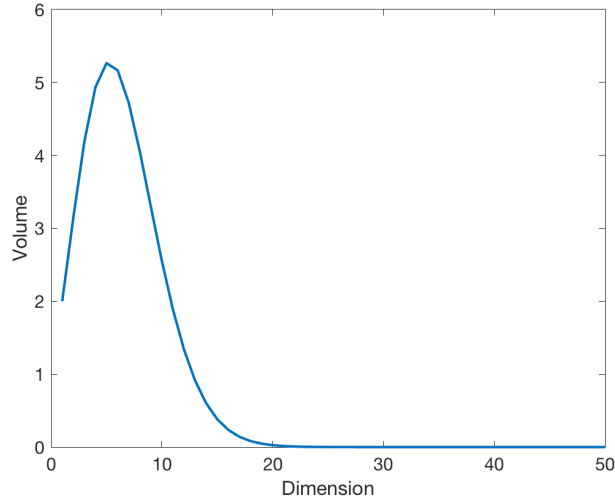


Fig. 2.1: The volume of the unit d -ball using the exact formula in equation (2.3). The volume reaches its maximum for $d = 5$ and decreases rapidly to zero with increasing dimension d .

slab has to be almost 2, so that 99% of the volume are captured by the slab. But as the dimension increases the width of the slab gets rapidly smaller. Indeed, in high dimensions only a very thin slab is required, since nearly all the volume of the unit ball lies a very small distance away from the equator. The following theorem makes the considerations above precise.

Fact 2.2 *Almost all the volume of $B^d(R)$ lies near its equator.*

Proof. It suffices to prove the result for the unit d -ball. Without loss of generality we pick as “north” the direction x_1 . The intersection of the sphere with the plane $x_1 = 0$ forms our equator, which is formally given by the $d - 1$ -dimensional region $\{x : \|x\| \leq 1, x_1 = 0\}$. This intersection is a sphere of dimension $d - 1$ with volume $\text{Vol}(B^{d-1})$ given by the $(d - 1)$ -analogue of formula (2.3) with $R = 1$.

We now compute the volume of B^d that lies between $x_1 = 0$ and $x_1 = p_0$. Let $P_0 = \{x : \|x\| \leq 1, x_1 \geq p_0\}$ be the “polar cap”, i.e., part of the sphere above the slab of width $2p_0$ around the equator. To compute the volume of the cap P we will integrate over all slices of the cap from p_0 to 1. Each such slice will be a sphere of dimension $d - 1$ and radius $\sqrt{1 - p^2}$, hence its volume is $(1 - p^2)^{\frac{d-1}{2}} \text{Vol}(B^{d-1})$. Therefore

$$\text{Vol}(P) = \int_{p_0}^1 (1 - p^2)^{\frac{d-1}{2}} \text{Vol}(B^{d-1}) dp = \text{Vol}(B^{d-1}) \int_{p_0}^1 (1 - p^2)^{\frac{d-1}{2}} dp.$$

Using $e^x \geq 1 + x$ for all x we can upper bound this integral by

$$\begin{aligned} \text{Vol}(P) &\leq \text{Vol}(B^{d-1}) \int_{p_0}^{\infty} e^{-\frac{d-1}{2}p^2} dp = \text{Vol}(B^{d-1}) \sqrt{\frac{2}{d-1}} \int_{p_0\sqrt{\frac{d-1}{2}}}^{\infty} e^{-u^2} du \\ &= \text{Vol}(B^{d-1}) \sqrt{\frac{\pi}{2(d-1)}} \text{erfc}\left(p_0\sqrt{\frac{d-1}{2}}\right), \end{aligned}$$

where $\text{erfc}(x) = \frac{2}{\sqrt{\pi}} \int_x^{\infty} e^{-u^2} du$ is the complementary error function. The upper bound $\text{erfc}(x) \leq \frac{e^{-x^2}}{\sqrt{\pi}x}$ gives

$$\text{Vol}(P) \leq \text{Vol}(B^{d-1}) \sqrt{\frac{\pi}{2(d-1)}} \frac{e^{-\frac{d-1}{2}p_0^2}}{\sqrt{\pi}p_0\sqrt{\frac{d-1}{2}}} = \frac{\text{Vol}(B^{d-1})}{d-1} \frac{e^{-\frac{d-1}{2}p_0^2}}{p_0}$$

Recall, from (2.3) that $\text{Vol}(B^d) = \frac{\pi^{\frac{d}{2}}}{\frac{d}{2} \Gamma(\frac{d}{2})}$, so, for d large enough (since $\frac{\Gamma(\frac{d}{2})}{\Gamma(\frac{d-1}{2})} \approx \sqrt{\frac{d}{2}}$),

$$\text{Vol}(B^{d-1}) = \frac{\pi^{-1/2}}{\frac{d-1}{d}} \frac{\Gamma(\frac{d}{2})}{\Gamma(\frac{d-1}{2})} \text{Vol}(B^d) \leq \frac{d-1}{2} \text{Vol}(B^d).$$

Finally, a simple calculation shows that the ratio between the volume of the polar caps and the entire hypersphere is bounded by

$$\frac{2 \text{Vol}(P)}{\text{Vol}(B^d)} \leq \frac{2 \text{Vol}(P)}{\text{Vol}(B^{d-1})} \leq \frac{1}{p_0} \exp\left(-\frac{d-1}{2}p_0^2\right).$$

The expression above shows that this ratio decreases exponentially as both d and p increase, proving our claim that the volume of the sphere concentrates strongly around its equator. $\square$

2.1.3 Concentration of the volume of a ball on shells

We consider two concentric balls $B^d(1)$ and $B^d(1-\varepsilon)$. Using equation (2.3), the ratio of their volumes is

$$\frac{\text{Vol}(B^d(1-\varepsilon))}{\text{Vol}(B^d(1))} = (1-\varepsilon)^d. \quad (2.5)$$

Clearly, for every ε this ratio tends to zero as $d \rightarrow \infty$. This implies that the spherical shell given by the region between $B^d(1)$ and $B^d(1-\varepsilon)$ will contain most of the volume of $B^d(1)$ for large enough d even if ε is very small. How quickly does the volume concentrate at the surface of $B^d(1)$? We choose ε as a function of d , e.g. $\varepsilon = \frac{t}{d}$, then

$$\frac{\text{Vol}(B^d(1-\varepsilon))}{\text{Vol}(B^d(1))} = \left(1 - \frac{t}{d}\right)^d \rightarrow e^{-t}.$$

Thus, almost all the volume of $B^d(R)$ is contained in an annulus of width R/d .

Therefore, if we peel a d -dimensional orange and even if we peel it very carefully so that we remove only a very thin layer of its peel, we will have removed most of the orange and are left with almost nothing.

2.1.4 Geometry of the hypercube

We have seen that most of the volume of the hypersphere is concentrated near its surface. A similar result also holds for the hypercube, and it is indeed a tendency for high-dimensional geometric objects. Yet, the hypercube exhibits an even more interesting volume concentration behavior, which we will establish below.

We start with a basic observation.

Proposition 2.3. *The hypercube C^d has volume 1 and diameter $\sqrt{d}$.*

The above proposition, although mathematically trivial, hints already at a somewhat counterintuitive behavior of the cube in high dimensions. Its corners seem to get “stretched out” more and more, while the rest of the cube must “shrink” to keep the volume constant. This property becomes even more striking when we compare the cube with the sphere as the dimension increases.

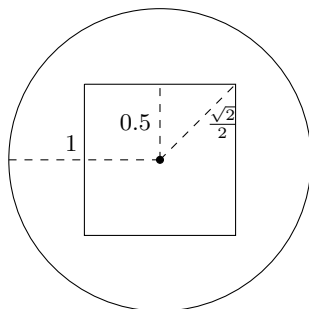


Fig. 2.2: 2-dimensional unit sphere and unit cube, centered at the origin.

In two dimensions (Figure 2.2), the unit square is completely contained in the unit sphere. The distance from the center to a vertex (radius of the circumscribed sphere) is $\frac{\sqrt{2}}{2}$ and the apothem (radius of the inscribed sphere) is $\frac{1}{2}$. In four dimensions (Figure 2.3), the distance from the center to a vertex is 1, so the vertices of the cube touch the surface of the sphere. However, the apothem is still $\frac{1}{2}$. The result, when projected in two dimensions no longer appears

convex, however all hypercubes are convex. This is part of the strangeness of higher dimensions - hypercubes are both convex and “pointy.” In dimensions greater than 4 the distance from the center to a vertex is $\frac{\sqrt{d}}{2} > 1$, and thus the vertices of the hypercube extend far outside the sphere, cf. Figure 2.4.

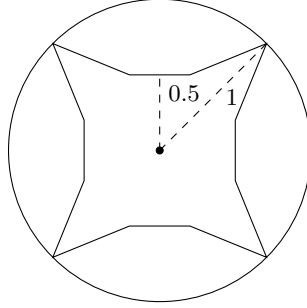


Fig. 2.3: Projections of the 4-dimensional unit sphere and unit cube, centered at the origin (4 of the 16 vertices of the hypercube are shown).

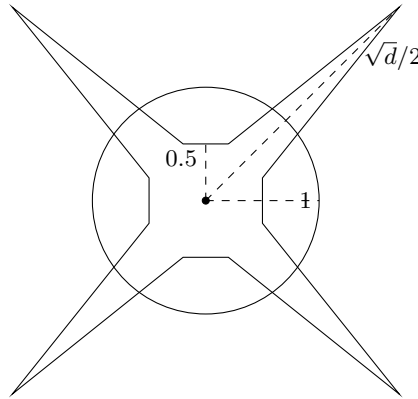


Fig. 2.4: Projections of the d -dimensional unit sphere and unit cube, centered at the origin (4 of the 2^d vertices of the hypercube are shown).

The considerations above suggest the following observation:

Fact 2.4 *Most of the volume of the high-dimensional cube is located in its corners.*

Proof. Let C^d denote the unit hypercube. The subset Q_R of C^d of points $x \in \mathbb{R}^d$ whose distance to the center is smaller than R ($Q_R = \{x \in C^d : \|x\| \leq R\}$) is given by $Q = C^d \cap B^d(R)$. Thus

$$\frac{\text{vol}(Q_R)}{\text{vol}(C^d)} = \text{vol}(Q) \leq B^d(R) \approx \frac{1}{\sqrt{d\pi}} \left(\frac{2\pi e R^2}{d} \right)^{\frac{d}{2}},$$

which goes to zero as long as $R < \sqrt{\frac{d}{2\pi e}}$. This shows that most of the volume of C^d lies in points at distance of order $\sqrt{d}$ of the center. $\square$

2.2 Basic concepts from probability

We briefly review some fundamental concepts from probability theory, which are helpful or necessary to understand the blessings of dimensionality and some of the surprises encountered in high dimensions. More advanced probabilistic concepts will be presented in Chapter ???. We assume that the reader is familiar with elementary probability as is covered in introductory probability courses (see, for example [71, 152]).

The two most basic concepts in probability associated with a random variable X are *expectation* (or *mean*) and *variance*, denoted by

$$\mathbb{E}[X] \quad \text{and} \quad \text{Var}(X) := \mathbb{E}[(X - \mathbb{E}[X])^2],$$

respectively. An important tool to describe probability distributions is the *moment generating function* of X , defined by

$$M_X(t) = \mathbb{E}[e^{tX}], \quad t \in \mathbb{R}, \quad (2.6)$$

the choice of nomenclature can be easily justified by expanding $M_X(t)$ in a series. The p -th moment of X is defined by $\mathbb{E}[X^p]$ for $p > 0$ and the p -th absolute moment is $\mathbb{E}[|X|^p]$.

We can introduce L^p -norms of random variables by taking the p -th root of moments, i.e.,

$$\|X\|_{L^p} := (\mathbb{E}[|X|^p])^{\frac{1}{p}}, \quad p \in [0, \infty],$$

with the usual extension to $p = \infty$ by setting

$$\|X\|_{\infty} := \text{ess sup } |X|.$$

Let $(\Omega, \Sigma, \mathbb{P})$ be a probability space, where Σ denotes a σ -algebra on the sample space Ω and $\mathbb{P}$ is a probability measure on (Ω, Σ) . For fixed p the vector space $L^p(\Omega, \Sigma, \mathbb{P})$ consists of all random variables X on Ω with finite L^p -norm, i.e.,

$$L^p(\Omega, \Sigma, \mathbb{P}) = \{X : \|X\|_{L^p} < \infty\}.$$

We will usually not mention the underlying probability space. For example, we will often simply write L^p for $L^p(\Omega, \Sigma, \mathbb{P})$.

The case $p = 2$ deserves special attention since L^2 is a Hilbert space with inner product and norm

$$\langle X, Y \rangle_{L^2} = \mathbb{E}[XY], \quad \|X\|_{L^2} = (\mathbb{E}[X^2])^{\frac{1}{2}},$$

respectively. Note that the *standard deviation* $\sigma(X) := \sqrt{\text{Var}(X)}$ of X can be written as

$$\sigma(X) = \|X - \mathbb{E}[X]\|_{L^2}.$$

The *covariance* of the random variables X and Y is

$$\text{cov}(X, Y) = \mathbb{E}[(X - \mathbb{E}[X])(Y - \mathbb{E}[Y])] = \langle X - \mathbb{E}[X], Y - \mathbb{E}[Y] \rangle_{L^2}. \quad (2.7)$$

We recall a few classical inequalities for random variables. *Hölder's inequality* states that for random variables X and Y on a common probability space and $p, q \geq 1$ with $1/p + 1/q = 1$, there holds

$$|\mathbb{E}[XY]| \leq \|X\|_{L^p} \|Y\|_{L^q}. \quad (2.8)$$

The special case $p = q = 2$ is the *Cauchy-Schwarz inequality*

$$|\mathbb{E}[XY]| \leq \sqrt{\mathbb{E}[X^2] \mathbb{E}[Y^2]}. \quad (2.9)$$

Jenssen's inequality states that for any random variable X and a convex function $\varphi : \mathbb{R} \rightarrow \mathbb{R}$, we have

$$\varphi(\mathbb{E}[X]) \leq \mathbb{E}[\varphi(X)]. \quad (2.10)$$

Since $\varphi(x) = x^{q/p}$ is a convex function for $q \geq p \geq 0$, it follows immediately from Jenssen's inequality that

$$\|X\|_{L^p} \leq \|X\|_{L^q} \quad \text{for } 0 \leq p \leq q < \infty.$$

Minkovskii's inequality states that for any $p \in [0, \infty]$ and any random variables X, Y , we have

$$\|X + Y\|_{L^p} \leq \|X\|_{L^p} + \|Y\|_{L^p}, \quad (2.11)$$

which can be viewed as the *triangle inequality*.

The *cumulative distribution function* of X is defined by

$$F_X(t) = \mathbb{P}(X \leq t), \quad t \in \mathbb{R}.$$

We have $\mathbb{P}\{X > t\} = 1 - F_X(t)$, where the function $t \mapsto \mathbb{P}\{|X| \geq t\}$ is called the *tail* of X . The following lemma establishes a close connection between expectation and tails.

Proposition 2.5 (Integral identity). *Let X be a non-negative random variable. Then*

$$\mathbb{E}[X] = \int_0^\infty \mathbb{P}\{X > t\} dt.$$

The two sides of this identity are either finite or infinite simultaneously.

Given an event E with non-zero probability, $\mathbb{P}(\cdot|E)$ denotes conditional probability, furthermore for a random variable X we use $\mathbb{E}[X|E]$ to denote the conditional expectation.

2.2.1 Tail bounds

Markov's inequality is a fundamental tool to bound the tail of a random variable in terms of its expectation.

Proposition 2.6. *For any non-negative random variable $X : S \rightarrow \mathbb{R}$ we have*

$$\mathbb{P}\{X \geq t\} \leq \frac{\mathbb{E}[X]}{t} \quad \text{for all } t > 0. \quad (2.12)$$

We provide two versions of the same proof, one using the language of conditional expectations.

Proof. Let $\mathcal{I}$ denote the event $\{X \geq t\}$. Then

$$\mathbb{E}[X] = \sum_{s \in S} p(s)X(s) = \sum_{s \in \mathcal{I}} p(s)X(s) + \sum_{s \in \mathcal{I}^c} p(s)X(s),$$

where $p(s)$ denotes the probability of s ; in case of continuous variables this should be replaced with the density function and $\sum$ with an integral.

Since X is non-negative, it holds $\sum_{s \in \mathcal{I}^c} p(s)X(s) \geq 0$ and

$$\mathbb{E}[X] \geq \sum_{s \in \mathcal{I}} p(s)X(s) \geq t \sum_{s \in \mathcal{I}} p(s) = t\mathbb{P}\{\mathcal{I}\}.$$

Proof (Using the language of conditional expectation).

$$\mathbb{E}[X] = \mathbb{P}(X < t)\mathbb{E}[X|X < t] + \mathbb{P}(X \geq t)\mathbb{E}[X|X \geq t],$$

where we take the product to be zero if the probability is zero.

Since X is non-negative, it holds $\mathbb{P}(X < t)\mathbb{E}[X|X < t] \geq 0$. Also, $\mathbb{E}[X|X \geq t] \geq t$. Hence,

$$\mathbb{E}[X] \geq \mathbb{P}(X \geq t)\mathbb{E}[X|X \geq t] \geq t\mathbb{P}(X \geq t).$$

An important consequence of Markov's inequality is *Chebyshev's inequality*.

Corollary 2.7. *Let X be a random variable with mean μ and variance σ^2 . Then, for any $t > 0$*

$$\mathbb{P}\{|X - \mu| \geq t\} \leq \frac{\sigma^2}{t^2}. \quad (2.13)$$

Chebyshev's inequality, which follows by applying Markov's inequality to the non-negative random variable $Y = (X - \mathbb{E}[X])^2$, is a form of concentration inequality, as it guarantees that X must be close to its mean μ whenever the variance of X is small. Both, Markov's and Chebyshev's inequality are sharp, i.e., in general they cannot be improved.

Markov's inequality only requires the existence of the first moment. We can say a bit more if in addition the random variable X has a moment generating

function in a neighborhood around zero, that is, there is a constant $b > 0$ such that $\mathbb{E}[e^{\lambda(X-\mu)}]$ exists for all $\lambda \in [0, b]$. In this case we can apply Markov's inequality to the random variable $Y = e^{\lambda(X-\mu)}$ and obtain the generic *Chernoff bound*

$$\mathbb{P}\{X - \mu \geq t\} = \mathbb{P}\{e^{\lambda(X-\mu)} \geq e^{\lambda t}\} \leq \frac{\mathbb{E}[e^{\lambda(X-\mu)}]}{e^{\lambda t}}. \quad (2.14)$$

In particular, optimizing over λ in order to obtain the tightest bound in (2.14) gives

$$\log \mathbb{P}\{X - \mu \geq t\} \leq - \sup_{\lambda \in [0, b]} \{\lambda t - \log \mathbb{E}[e^{\lambda(X-\mu)}]\}.$$

Gaussian random variables are among the most important random variables. A Gaussian random variable X with mean μ and standard deviation σ has a probability density function given by

$$\psi(t) = \frac{1}{\sqrt{2\pi}\sigma^2} \exp\left(-\frac{(t-\mu)^2}{2\sigma^2}\right). \quad (2.15)$$

We write $X \sim \mathcal{N}(\mu, \sigma^2)$. We call a Gaussian random variable X with $\mathbb{E}[X] = 0$ and $\mathbb{E}[X^2] = 1$ a *standard Gaussian* or *standard normal* (random variable). In this case we have the following tail bound.

Proposition 2.8 (Gaussian tail bounds). *Let $X \sim \mathcal{N}(\mu, \sigma^2)$. Then for all $t > 0$*

$$\mathbb{P}(X \geq \mu + t) \leq e^{-t^2/2\sigma^2}. \quad (2.16)$$

Proof. Without loss of generality we consider $\mu = 0$ and $\sigma = 1$, as it is straightforward to extend to the general case with a change of variables. We use the moment-generating function $\lambda \mapsto \mathbb{E}[e^{\lambda X}]$. A simple calculation gives

$$\mathbb{E}[e^{\lambda X}] = \frac{1}{\sqrt{2\pi}} \int_{-\infty}^{\infty} e^{\lambda x - x^2/2} dx = \frac{1}{\sqrt{2\pi}} e^{\lambda^2/2} \int_{-\infty}^{\infty} e^{-(x-\lambda)^2/2} dx = e^{\lambda^2/2},$$

where we have used the fact that $\int_{-\infty}^{\infty} e^{-(x-\lambda)^2/2} dx$ is just the entire Gaussian integral shifted and therefore its value is $\sqrt{2\pi}$. We now apply Chernoff's bound (2.14) and obtain $\mathbb{P}(X > t) \leq \mathbb{E}[e^{\lambda X}] e^{-\lambda t}$. Minimizing this expression over λ gives $\lambda = t$ and thus $\mathbb{P}(X > t) \leq e^{-t^2/2}$.

2.2.2 Sub-gaussian random variables

Definition 2.9. *A random variable X with mean $\mu = \mathbb{E}[X]$ is called sub-Gaussian if there is a positive number σ such that*

$$\mathbb{E}[e^{\lambda(X-\mu)}] \leq e^{\sigma^2 \lambda^2/2}, \quad \text{for all } \lambda \in \mathbb{R}.$$

Note that σ^2 is not necessarily the variance of X . If X satisfies the above definition, we also say that X is sub-Gaussian with parameter σ , or X is (μ, σ) sub-Gaussian in case we want to emphasize μ as well. Clearly, owing to the symmetry in the definition, $-X$ is sub-Gaussian if and only if X is sub-Gaussian. Obviously, any Gaussian random variable with variance σ^2 is sub-Gaussian with parameter σ . We refer to [185] for other, equivalent, definitions of sub-Gaussian random variables.

Combining the moment condition in Definition 2.9 with calculations similar to those that lead us to the Gaussian tail bounds in 2.8, yields the following concentration inequality for sub-Gaussian random variables.

Proposition 2.10 (Sub-Gaussian tail bounds). *Assume X is sub-Gaussian with parameter σ . Then for all $t > 0$*

$$\mathbb{P}(|X - \mu| \geq t) \leq 2e^{-t^2/2\sigma^2} \quad \text{for all } t \in \mathbb{R}. \quad (2.17)$$

An important example of non-Gaussian, but sub-Gaussian random variables are *Rademacher random variables*. A Rademacher random variable ε takes on the values ± 1 with equal probability and is sub-Gaussian with parameter σ . Indeed, any bounded random variable is sub-Gaussian (see Remark 2.16).

While many important random variables have a sub-Gaussian distribution, this class does not include several frequently occurring distributions with heavier tails. A classical example is the *chi-squared distribution*, which we will discuss at the end of this chapter.

2.2.3 Sub-exponential random variables

Relaxing slightly the condition on the moment-generating function in Definition 2.9 leads to the class of *sub-exponential* random variables.

Definition 2.11. *A random variable X with mean $\mu = \mathbb{E}[X]$ is called sub-exponential if there are parameters ν, b such that*

$$\mathbb{E}[e^{\lambda(X-\mu)}] \leq e^{\nu^2\lambda^2/2}, \quad \text{for all } |\lambda| \leq \frac{1}{b}.$$

Clearly, a sub-Gaussian random variable is sub-exponential (set $\nu = \sigma$ and $b = 0$, where $1/b$ is interpreted as $+\infty$). However, the converse is not true. Take for example $X \sim \mathcal{N}(0, 1)$ and consider the random variable $Z = X^2$. For $\lambda < \frac{1}{2}$ it holds that

$$\mathbb{E}[e^{\lambda(Z-1)}] = \frac{1}{\sqrt{2\pi}} \int_{-\infty}^{\infty} e^{\lambda(x^2-1)} e^{-x^2/2} dx = \frac{e^{-\lambda}}{\sqrt{1-2\lambda}}. \quad (2.18)$$

However, for $\lambda \geq \frac{1}{2}$ the moment-generating function does not exist, which implies that X^2 is not sub-Gaussian. But X^2 is sub-exponential. Indeed, a brief computation shows that

$$\frac{e^{-\lambda}}{\sqrt{1-2\lambda}} \leq e^{2\lambda^2} = e^{4\lambda^2/2}, \quad \text{for all } |\lambda| \leq 1/4,$$

which in turn implies that X^2 is sub-exponential with parameters $(\nu, b) = (2, 4)$.

Following a similar procedure that yielded sub-Gaussian tail bounds produces concentration inequalities for sub-exponential random variables. However, in this case we see two different types of concentration emerging, depending on the value of t .

Proposition 2.12 (Sub-exponential tail bounds). *Assume X is sub-exponential with parameters (ν, b) . Then*

$$\mathbb{P}(X \geq \mu + t) \leq \begin{cases} e^{-t^2/2\nu^2} & \text{if } 0 \leq t \leq \frac{\nu^2}{b}, \\ e^{-t/2b} & \text{if } t > \frac{\nu^2}{b}. \end{cases} \quad (2.19)$$

Both the sub-Gaussian property and the sub-exponential property is preserved under summation for independent random variables, and the associated parameters transform in a simple manner (see Remark 2.16).

2.3 Concentration of measure, a “blessing of dimensionality”

Suppose we wish to predict the outcome of an event of interest. One natural approach would be to compute the expected value of the object. However, how can we tell how good the expected value is to the actual outcome of the event? Without further information of how well the actual outcome concentrates around its expectation, the expected value is of little use. We would like to have an estimate for the probability that the actual outcome deviates from its expectation by a certain amount. This is exactly the role that *concentration inequalities* play in probability and statistics.

The concentration of measure phenomenon was put forward by Vitali Milman in the asymptotic geometry of Banach spaces regarding probabilities on product spaces in high dimensions [125, 112].

The celebrated law of large numbers of classical probability theory is the most well known form of *concentration of measure*; it states that sums of independent random variables are, under very mild conditions, close to their expectation with a large probability. We will see various quantitative versions of such concentration inequalities throughout this course. Some deal with sums of scalar random variables, others with sums of random vectors or sums of random matrices. Such concentration inequalities are instances of what is sometimes called *Blessings of dimensionality* (cf. [68]). This expression refers to the fact that certain random fluctuations can be well controlled in high dimensions, while it would be very complicated to make such predictive statements in moderate dimensions.

2.3.1 Large deviation inequalities

Concentration and large deviations inequalities are among the most useful tools when understanding the performance of some algorithms. We start with two of the most fundamental results in probability. We refer to Sections 1.7 and 2.4 in [71] for the proofs and variations.

Theorem 2.13 (Strong law of large numbers). *Let $X_1, X_2, \dots$ be a sequence of i.i.d. random variables with mean μ . Denote*

$$S_n := X_1 + \dots + X_n.$$

Then, as $n \rightarrow \infty$

$$\frac{S_n}{n} \rightarrow \mu \quad \text{almost surely.} \quad (2.20)$$

The celebrated *central limit theorem* tells us that the limiting distribution of a sum of i.i.d. random variables is always Gaussian. The best known version is probably due to Lindeberg-Lévy.

Theorem 2.14 (Lindeberg-Lévy Central limit theorem). *Let $X_1, X_2, \dots$ be a sequence of i.i.d. random variables with mean μ and variance σ^2 . Denote*

$$S_n := X_1 + \dots + X_n,$$

and consider the normalized random variable Z_n with mean zero and variance one, given by

$$Z_n := \frac{S_n - \mathbb{E}[S_n]}{\sqrt{\text{Var } S_n}} = \frac{1}{\sigma\sqrt{n}} \sum_{i=1}^n (X_i - \mu).$$

Then, as $n \rightarrow \infty$

$$Z_n \rightarrow \mathcal{N}(0, 1) \quad \text{in distribution.} \quad (2.21)$$

The strong law of large numbers and the central limit theorem give us qualitative statements about the behavior of a sum of i.i.d. random variables. In many applications it is desirable to be able to quantify how such a sum deviates around its mean. This is where concentration inequalities come into play.

The intuitive idea is that if we have a sum of independent random variables

$$X = X_1 + \dots + X_n,$$

where X_i are i.i.d. centered random variables, then while the value of X can be of order $\mathcal{O}(n)$ it will very likely be of order $\mathcal{O}(\sqrt{n})$ (note that this is the order of its standard deviation). The inequalities that follow are ways of very precisely controlling the probability of X being larger (or smaller) than $\mathcal{O}(\sqrt{n})$. While we could use, for example, Chebyshev’s inequality for this, in the inequalities that follow the probabilities will be exponentially small, rather

than just quadratically small, which will be crucial in many applications to come. Moreover, unlike the classical central limit theorem, the concentration inequalities below are *non-asymptotic* in the sense that they hold for all fixed n and not just for $n \rightarrow \infty$ (but the larger the n , the stronger the inequalities become).

Theorem 2.15 (Hoeffding's Inequality). *Let $X_1, X_2, \dots, X_n$ be independent bounded random variables, i.e., $|X_i| \leq a_i$ and $\mathbb{E}[X_i] = 0$. Then,*

$$\mathbb{P} \left\{ \left| \sum_{i=1}^n X_i \right| \geq t \right\} \leq 2 \exp \left(-\frac{t^2}{2 \sum_{i=1}^n a_i^2} \right).$$

The inequality implies that fluctuations larger than $\mathcal{O}(\sqrt{n})$ have small probability. For example, if $a_i = a$ for all i , setting $t = a\sqrt{2n \log n}$ yields that the probability is at most $\frac{2}{n}$.

Proof. We prove the result for the case $|X_i| \leq a$, the extension to the case $|X_i| \leq a_i$ is straightforward. We first get a probability bound for the event $\sum_{i=1}^n X_i > t$. The proof, again, will follow from Markov. Since we want an exponentially small probability, we use a classical trick that involves exponentiating with any $\lambda > 0$ and then choosing the optimal λ .

$$\begin{aligned} \mathbb{P} \left\{ \sum_{i=1}^n X_i \geq t \right\} &= \mathbb{P} \left\{ e^{\lambda \sum_{i=1}^n X_i} \geq e^{\lambda t} \right\} \\ &\leq \frac{\mathbb{E}[e^{\lambda \sum_{i=1}^n X_i}]}{e^{t\lambda}} \\ &= e^{-t\lambda} \prod_{i=1}^n \mathbb{E}[e^{\lambda X_i}], \end{aligned} \tag{2.22}$$

where the penultimate step follows from Markov's inequality and the last equality follows from independence of the X_i 's.

We now use the fact that $|X_i| \leq a$ to bound $\mathbb{E}[e^{\lambda X_i}]$. Because the function $f(x) = e^{\lambda x}$ is convex,

$$e^{\lambda x} \leq \frac{a+x}{2a} e^{\lambda a} + \frac{a-x}{2a} e^{-\lambda a},$$

for all $x \in [-a, a]$.

Since, for all i , $\mathbb{E}[X_i] = 0$ we get

$$\mathbb{E}[e^{\lambda X_i}] \leq \mathbb{E} \left[\frac{a+X_i}{2a} e^{\lambda a} + \frac{a-X_i}{2a} e^{-\lambda a} \right] = \frac{1}{2} (e^{\lambda a} + e^{-\lambda a}) = \cosh(\lambda a)$$

Note that²

²This follows immediately from the Taylor expansions: $\cosh(x) = \sum_{n=0}^{\infty} \frac{x^{2n}}{(2n)!}$, $e^{x^2/2} = \sum_{n=0}^{\infty} \frac{x^{2n}}{2^n n!}$, and $(2n)! \geq 2^n n!$.

$$\cosh(x) \leq e^{x^2/2}, \quad \text{for all } x \in \mathbb{R}$$

Hence,

$$\mathbb{E}[e^{\lambda X_i}] \leq e^{(\lambda a)^2/2}. \quad (2.23)$$

Together with (2.22), this gives

$$\begin{aligned} \mathbb{P}\left\{\sum_{i=1}^n X_i \geq t\right\} &\leq e^{-t\lambda} \prod_{i=1}^n e^{(\lambda a)^2/2} \\ &= e^{-t\lambda} e^{n(\lambda a)^2/2} \end{aligned}$$

This inequality holds for any choice of $\lambda \geq 0$, so we choose the value of λ that minimizes

$$\min_{\lambda} \left\{ n \frac{(\lambda a)^2}{2} - t\lambda \right\}$$

Differentiating readily shows that the minimizer is given by

$$\lambda = \frac{t}{na^2},$$

which satisfies $\lambda > 0$. For this choice of λ ,

$$n(\lambda a)^2/2 - t\lambda = \frac{1}{n} \left(\frac{t^2}{2a^2} - \frac{t^2}{a^2} \right) = -\frac{t^2}{2na^2}$$

Thus,

$$\mathbb{P}\left\{\sum_{i=1}^n X_i \geq t\right\} \leq e^{-\frac{t^2}{2na^2}}$$

By using the same argument on $\sum_{i=1}^n (-X_i)$, and union bounding over the two events we get,

$$\mathbb{P}\left\{\left|\sum_{i=1}^n X_i\right| \geq t\right\} \leq 2e^{-\frac{t^2}{2na^2}}$$

□

Remark 2.16 (Hoeffding’s inequality and sub-Gaussianity). It is useful to dissect the proof of Hoeffding’s inequality, and notice that it was done by showing two things: (i) that the sum of independent sub-Gaussian random variables is sub-Gaussian with parameter whose square is the sum of the squares of the sub-Gaussian parameters of the original random variables (this is effectively what is proven in (2.22)), and (ii) that bounded random variables are sub-Gaussian (shown in (2.23)).

Remark 2.17. Hoeffding's inequality is suboptimal in a sense we now describe. Let's say that we have random variables $r_1, \dots, r_n$ i.i.d. distributed as

$$r_i = \begin{cases} -1 & \text{with probability } p/2 \\ 0 & \text{with probability } 1-p \\ 1 & \text{with probability } p/2. \end{cases}$$

Then, $\mathbb{E}(r_i) = 0$ and $|r_i| \leq 1$ so Hoeffding's inequality gives:

$$\mathbb{P} \left\{ \left| \sum_{i=1}^n r_i \right| > t \right\} \leq 2 \exp \left(-\frac{t^2}{2n} \right).$$

Intuitively, the smaller p is, the more concentrated $|\sum_{i=1}^n r_i|$ should be, however Hoeffding's inequality does not capture this behaviour.

A natural way to attempt to capture this behaviour is by noting that the variance of $\sum_{i=1}^n r_i$ depends on p as $\text{Var}(r_i) = p$. The inequality that follows, Bernstein's inequality, uses the variance of the summands to improve over Hoeffding's inequality.

The way this is going to be achieved is by strengthening the proof above, more specifically in step (2.22) we will use the bound on the variance to get a better estimate on $\mathbb{E}[e^{\lambda X_i}]$ essentially by realizing that if X_i is centered, $\mathbb{E}X_i^2 = \sigma^2$, and $|X_i| \leq a$ then, for $k \geq 2$, $\mathbb{E}X_i^k \leq \mathbb{E}|X_i|^k \leq a^{k-2}\mathbb{E}|X_i|^2 \leq a^{k-2}\sigma^2 = \left(\frac{\sigma^2}{a^2}\right)a^k$.

Theorem 2.18 (Bernstein's Inequality).

Let $X_1, X_2, \dots, X_n$ be independent centered bounded random variables satisfying $|X_i| \leq a$ and $\mathbb{E}[X_i^2] = \sigma^2$. Then,

$$\mathbb{P} \left\{ \left| \sum_{i=1}^n X_i \right| \geq t \right\} \leq 2 \exp \left(-\frac{t^2}{2n\sigma^2 + \frac{2}{3}at} \right).$$

Remark 2.19. Before proving Bernstein's inequality, note that on the example of Remark 2.17 we get

$$\mathbb{P} \left\{ \left| \sum_{i=1}^n r_i \right| \geq t \right\} \leq 2 \exp \left(-\frac{t^2}{2np + \frac{2}{3}t} \right),$$

which exhibits a dependence on p and, for small values of p is considerably smaller than what Hoeffding's inequality gives.

Proof.

As before, we will prove

$$\mathbb{P} \left\{ \sum_{i=1}^n X_i \geq t \right\} \leq \exp \left(-\frac{t^2}{2n\sigma^2 + \frac{2}{3}at} \right),$$

and then union bound with the same result for $-\sum_{i=1}^n X_i$, to prove the Theorem.

For any $\lambda > 0$ we have

$$\begin{aligned}\mathbb{P}\left\{\sum_{i=1}^n X_i \geq t\right\} &= \mathbb{P}\{e^{\lambda \sum X_i} \geq e^{\lambda t}\} \\ &\leq \frac{\mathbb{E}[e^{\lambda \sum X_i}]}{e^{\lambda t}} \\ &= e^{-\lambda t} \prod_{i=1}^n \mathbb{E}[e^{\lambda X_i}]\end{aligned}$$

The following calculations reveal the source of the improvement over Hoeffding’s inequality.

$$\begin{aligned}\mathbb{E}[e^{\lambda X_i}] &= \mathbb{E}\left[1 + \lambda X_i + \sum_{m=2}^{\infty} \frac{\lambda^m X_i^m}{m!}\right] \\ &\leq 1 + \sum_{m=2}^{\infty} \frac{\lambda^m a^{m-2} \sigma^2}{m!} \\ &= 1 + \frac{\sigma^2}{a^2} \sum_{m=2}^{\infty} \frac{(\lambda a)^m}{m!} \\ &= 1 + \frac{\sigma^2}{a^2} (e^{\lambda a} - 1 - \lambda a)\end{aligned}$$

Therefore,

$$\mathbb{P}\left\{\sum_{i=1}^n X_i \geq t\right\} \leq e^{-\lambda t} \left[1 + \frac{\sigma^2}{a^2} (e^{\lambda a} - 1 - \lambda a)\right]^n$$

We will use a few simple inequalities (that can be easily proved with calculus) such as³ $1 + x \leq e^x$, for all $x \in \mathbb{R}$.

This means that,

$$1 + \frac{\sigma^2}{a^2} (e^{\lambda a} - 1 - \lambda a) \leq e^{\frac{\sigma^2}{a^2} (e^{\lambda a} - 1 - \lambda a)},$$

which readily implies

$$\mathbb{P}\left\{\sum_{i=1}^n X_i \geq t\right\} \leq e^{-\lambda t} e^{\frac{n\sigma^2}{a^2} (e^{\lambda a} - 1 - \lambda a)}.$$

As before, we try to find the value of $\lambda > 0$ that minimizes

³In fact $y = 1 + x$ is a tangent line to the graph of $f(x) = e^x$.

$$\min_{\lambda} \left\{ -\lambda t + \frac{n\sigma^2}{a^2} (e^{\lambda a} - 1 - \lambda a) \right\}$$

Differentiation gives

$$-t + \frac{n\sigma^2}{a^2} (ae^{\lambda a} - a) = 0$$

which implies that the optimal choice of λ is given by

$$\lambda^* = \frac{1}{a} \log \left(1 + \frac{at}{n\sigma^2} \right)$$

If we set

$$u = \frac{at}{n\sigma^2}, \quad (2.24)$$

then $\lambda^* = \frac{1}{a} \log(1 + u)$.

Now, the value of the minimum is given by

$$-\lambda^* t + \frac{n\sigma^2}{a^2} (e^{\lambda^* a} - 1 - \lambda^* a) = -\frac{n\sigma^2}{a^2} [(1 + u) \log(1 + u) - u].$$

This means that

$$\mathbb{P} \left\{ \sum_{i=1}^n X_i \geq t \right\} \leq \exp \left(-\frac{n\sigma^2}{a^2} \{ (1 + u) \log(1 + u) - u \} \right)$$

The rest of the proof follows by noting that, for every $u > 0$,

$$(1 + u) \log(1 + u) - u \geq \frac{u}{\frac{2}{u} + \frac{2}{3}}, \quad (2.25)$$

which implies:

$$\begin{aligned} \mathbb{P} \left\{ \sum_{i=1}^n X_i \geq t \right\} &\leq \exp \left(-\frac{n\sigma^2}{a^2} \frac{u}{\frac{2}{u} + \frac{2}{3}} \right) \\ &= \exp \left(-\frac{t^2}{2n\sigma^2 + \frac{2}{3}at} \right). \end{aligned}$$

□

We refer to [185] for several useful variations of Bernstein's inequality.

2.3.2 The geometry of the hypercube revisited

Equipped with the probabilistic tools from the previous sections, we will give an alternative proof for Fact 2.4.

Theorem 2.20. [See Fact 2.4] Let x be a random, uniformly drawn, point in C^d . We have that $\|x\| \geq \frac{\sqrt{d}}{4}$ with high probability (meaning with probability $1 - o(1)$).

The proof of this statement will be based on a probabilistic argument, thereby illustrating (again) the nice and fruitful connection between geometry and probability in high dimension. Pick a point at random in C^d , which corresponds to $[-\frac{1}{2}, \frac{1}{2}]^d$. We want to upperbound the probability that the point is also in the sphere of radius R .

Let $x = (x_1, \dots, x_d) \in \mathbb{R}^d$ and each $x_i \in [-\frac{1}{2}, \frac{1}{2}]$ is chosen uniformly at random. The event that x also lies in the sphere of radius R corresponds to the inequality

$$\|x\|_2 = \sqrt{\sum_{i=1}^d x_i^2} \leq R.$$

Let $z_i = x_i^2$ and note that

$$\mathbb{E}[z_i] = \frac{1}{2} \int_{-1/2}^{1/2} t^2 dt = 2 \frac{(1/2)^3}{3} = \frac{1}{12},$$

which implies $\mathbb{E}[\|x\|_2^2] = \frac{d}{24}$. Since the random variables $z_i - \mathbb{E}[z_i] = z_i - \frac{1}{12}$ are independent, centered, and bounded ($|z_i - \mathbb{E}[z_i]| \leq \max\{\frac{1}{12}, \frac{1}{4} - \frac{1}{12}\} = \frac{2}{12}$), we can use Hoeffding’s inequality (Theorem 2.15). This gives

$$\begin{aligned} \mathbb{P}(\|x\|_2^2 \leq R^2) &= \mathbb{P}\left(\sum_{i=1}^d x_i^2 \leq R^2\right) = \mathbb{P}\left(\sum_{i=1}^d (z_i - \mathbb{E}[z_i]) \leq R^2 - \frac{d}{12}\right) \\ &\leq \exp\left[-\frac{\left(\frac{d}{12} - R^2\right)^2}{2d\left(\frac{2}{12}\right)^2}\right] = \exp\left[-\frac{(d - 12R^2)^2}{8d}\right]. \end{aligned}$$

Taking $R = \frac{\sqrt{d}}{4}$ gives $\mathbb{P}\left(\|x\|_2 \leq \frac{\sqrt{d}}{4}\right) \leq \exp[-c_0 d]$ for a universal constant c_0

(which can be computed in this case to be $c_0 = \frac{(1 - \frac{16}{8})^2}{8} = \frac{1}{128}$, but that’s not important).

Since we now have gained a better understanding of the properties of the cube in high dimensions, we can use this knowledge to our advantage. For instance, this “pointiness” of the hypercube (in form of the ℓ_1 -ball) turns out to be very useful in the areas of compressive sensing and sparse recovery, see Chapter 10.

2.3.3 Tail bounds for the chi-squared distribution

If $X_1, \dots, X_n$ are independent, standard normal random variables, then the sum of their squares, $Z = \sum_{k=1}^n X_k^2$ is distributed according to the *chi-squared*

distribution with n degrees of freedom. We denote this by $Z \sim \chi^2(n)$. Its probability density function is

$$\varphi(t) = \begin{cases} \frac{t^{\frac{n}{2}-1} e^{-\frac{t}{2}}}{2^{\frac{n}{2}} \Gamma(\frac{n}{2})}, & t > 0. \\ 0, & \text{else.} \end{cases}$$

Since the $X_k^2, k = 1, \dots, n$ are subexponential with parameters $(2, 4)$ and independent, $Z = \sum_{k=1}^n X_k^2$ is subexponential with parameters $(2\sqrt{n}, 4)$ (analogously to Remark 2.16). Therefore, using (2.19), we obtain the χ^2 tail bound

$$\mathbb{P} \left(\left| \frac{1}{n} \sum_{k=1}^n X_k^2 - 1 \right| \geq t \right) \leq \begin{cases} 2e^{-nt^2/8} & \text{for } t \in (0, 1). \\ 2e^{-nt/8} & \text{if } t \geq 1. \end{cases} \quad (2.26)$$

A very useful variation of this bound, commonly known as “Lemma 1 in Laurent and Massart [111]” is the following Theorem.

Theorem 2.21 ([111]). *Let $X_1, \dots, X_n$ be i.i.d. standard Gaussian random variables ($\mathcal{N}(0, 1)$), and $a_1, \dots, a_n$ non-negative numbers and not all zero. Let*

$$Z = \sum_{k=1}^n a_k (X_k^2 - 1).$$

The following inequalities hold for any $x > 0$:

- $\mathbb{P} \{ Z \geq 2\|a\|_2 \sqrt{x} + 2\|a\|_\infty x \} \leq \exp(-x),$
- $\mathbb{P} \{ Z \leq -2\|a\|_2 \sqrt{x} \} \leq \exp(-x),$

where $\|a\|_2^2 = \sum_{k=1}^n a_k^2$ and $\|a\|_\infty = \max_{1 \leq k \leq n} |a_k|$.

Note that if $a_k = 1$, for all k , then Z is a χ_2 random variable with n degrees of freedom, so this theorem immediately gives a deviation inequality for χ_2 random variables, readily comparable to (2.26).

2.3.4 How to generate random points on a sphere

How can we sample a point uniformly at random from S^{d-1} ? The first approach that may come to mind is the following method to generate random points on a unit circle. Independently generate each coordinate uniformly at random from the interval $[-1, 1]$. This yields points that are distributed uniformly at random in a square that contains the unit ball. We could now project all points onto the unit ball. However, the resulting distribution will not be uniform since more points fall on a line from the origin to a vertex of the square, than fall on a line from the origin to the midpoint of an edge due to the difference in length of the diagonal of the square to its side length.

To remedy this problem, we could discard all points outside the unit ball and project the remaining points onto the sphere. However, if we generalize

this technique to higher dimensions, the analysis in the previous section has shown that the ratio of the volume of $S^{d-1}(1)$ to the volume of $C^d(1)$ decreases rapidly. This makes this process not practical, since almost all the generated points will be discarded in this process and we end up with essentially no points inside (and thus, after projection, on) the sphere.

Instead we can proceed as follows: Recall that the multivariate Gaussian distribution is rotation invariant about the origin. This means that we can simply construct a vector in $\mathbb{R}^d$ whose entries are independently drawn from a univariate Gaussian distribution, and then normalize it to lie on the sphere. This gives a distribution of points that is uniform over the sphere.

Having a method of generating points uniformly at random on S^{d-1} at our disposal, we can now give a probabilistic proof that points on S^{d-1} concentrate near its equator.

Theorem 2.22. *Let x be a point randomly, and uniformly, drawn in the sphere in d dimensions S^{d-1} . For any $\rho > 0$ and d large enough, $\|x\| \geq \sqrt{\frac{\rho}{d}}$ has probability $\geq 1 - e^{-\frac{\rho}{2}}$.*

Without loss of generality, we pick the first canonical basis vector e_1 to represent the “north pole”, and the intersection of the sphere with the plane $x_1 = 0$ forms our equator. We create a random vector by sampling $(z_1, \dots, z_d) \sim \mathcal{N}(0, I_d)$ and normalize the vector to get $x = (x_1, \dots, x_d) = \frac{1}{\sqrt{\sum_{k=1}^d z_k^2}}(z_1, \dots, z_d)$. Because x is on the sphere, it holds that $\|x\|_2 = 1$.

We are interested in understanding the random variable x_1^2 . Notice that $1 = \sum_{k=1}^d x_k^2 = \mathbb{E} \left[\sum_{k=1}^d x_k^2 \right] = d\mathbb{E}x_1^2$ since all coordinates are identically distributed (but not independent), this means that $\mathbb{E}x_k^2 = \frac{1}{d}$ for all coordinates x_k . Our goal is to upperbound the probability that $x_1 = \frac{z_1}{\sqrt{\sum_{k=1}^d z_k^2}}$ is large, we will do that by union bounding the events that z_1 is large (in absolute value) and that $\sum_{k=1}^d z_k^2$ is small (which can be done with (2.26)). We will not attempt to derive the strongest possible bound, but rather to showcase an useful argument. Let $\rho > 0$ and $0 < \varepsilon < 1$ be parameters that we will set later (perhaps depending on d). We have

$$\begin{aligned} \mathbb{P} \left(|x_1| \geq \sqrt{\frac{\rho}{d}} \right) &\leq \mathbb{P} \left(|z_1| \geq \sqrt{(1-\varepsilon)\rho} \text{ and } \sum_{k=1}^d z_k^2 \leq (1-\varepsilon)d \right) \\ &\leq \mathbb{P} \left(|z_1| \geq \sqrt{(1-\varepsilon)\rho} \right) + \mathbb{P} \left(\frac{1}{d} \sum_{k=1}^d z_k^2 \leq 1-\varepsilon \right) \\ &\leq 2 \exp \left[-\frac{(1-\varepsilon)\rho}{2} \right] + 2 \exp \left[-\frac{d\varepsilon^2}{2} \right], \end{aligned}$$

where in the last inequality we used the Gaussian tail bound (Proposition 2.8) in the first term and the χ^2 bound (2.26) in the second.

For d large enough, taking ε sufficiently small we have

$$\mathbb{P}\left(|x_1| \geq \sqrt{\frac{\rho}{d}}\right) \leq \exp\left[-\frac{\rho}{3}\right],$$

where the 3 could be replaced by any real larger than 2. $\square$

Remark 2.23 (Random vectors in high dimensions). The results shown in this chapter give us valuable insight about the behavior of random vectors in high dimensions. The results showing that most of the ball is near the boundary, together with the concentration of the norm of the gaussian vector $z \sim \mathcal{N}(0, I)$ give us the crucial insight that lengths of random high dimensional vectors tend to concentrate. The theorems that establish that most of the sphere is near the equator give us a perhaps even more valuable insight: two random vectors in high dimensions are almost always almost orthogonal. Indeed, we can take one of them to be the “north pole” without loss of generality and (in the case in which they are uniformly drawn in the sphere) Theorem 2.22 states that their inner product is in the order of $\frac{1}{\sqrt{d}}$ smaller than the product of their norms, which corresponds to an angle very close to perpendicular.

We will dive deeper into the topic of probability in high dimensions in Chapters 8 and 9. There are several excellent texts about different aspects of high dimensional probability, two notable examples are [180] and [185].

Singular Value Decomposition and Principal Component Analysis

Data is most often represented as a matrix, even network data and graphs are often naturally represented by their adjacency matrix. For this reason Linear Algebra is one of the key tools in data analysis. Perhaps more surprising is the fact that spectral properties of matrices representing data play a crucial role in data analysis. We illustrate this importance with a discussion of the Singular Value Decomposition (SVD) and Principal Component Analysis (PCA) that are often used for data compression, denoising, and dimension reduction. Tools from random matrix theory are then utilized to better understand the performance of SVD and PCA in the high dimensional regime.

3.1 Singular Value Decomposition

We recommend the reader [93] and [81] as base references in Linear Algebra.

The SVD is one of the most useful tools for analyzing data. Given a matrix $A \in \mathbb{R}^{m \times n}$, the SVD of A is given by

$$A = U \Sigma V^T, \quad (3.1)$$

where $U \in O(m)$, $V \in O(n)$ are orthogonal matrices (meaning that $U^T U = U U^T = I_{m \times m}$ and $V^T V = V V^T = I_{n \times n}$) and $\Sigma \in \mathbb{R}^{m \times n}$ is a matrix with non-negative entries on its diagonal and otherwise zero entries.

The columns of U and V are referred to, respectively, as the left and right singular vectors of A and the diagonal elements of Σ as the singular values of A . Through the SVD, any matrix can be written as a sum of rank-1 matrices

$$A = \sum_{k=1}^r \sigma_k u_k v_k^T, \quad (3.2)$$

where $\sigma_1 \geq \sigma_2 \geq \dots \geq \sigma_r > 0$ are the non-zero singular values of A , and u_k and v_k are the corresponding left and right singular vectors. In particular, $\text{rank}(A) = r$, that is, the number of non-zero singular values r is the rank of A .

Remark 3.1. The SVD of a complex valued matrix $A \in \mathbb{C}^{m \times n}$ is given by $A = U\Sigma V^*$, where $U \in \mathbb{C}^{m \times m}$ and $V \in \mathbb{C}^{n \times n}$ are unitary matrices ($UU^* = U^*U = I_{m \times m}$ and $VV^* = V^*V = I_{n \times n}$) and Σ is the same as in the real case (i.e., $\Sigma \in \mathbb{R}^{m \times n}$ is a matrix with non-negative entries on its diagonal and otherwise zero entries). We focus our presentation on the SVD of real valued matrices as they are more common in data science, but everything can be straightforwardly extended to the complex case.

Remark 3.2. Say $m \leq n$, it is easy to see that we can also think of the SVD as having $U \in \mathbb{R}^{m \times n}$ where $UU^T = I$, $\Sigma \in \mathbb{R}^{n \times n}$ a diagonal matrix with non-negative entries and $V \in O(n)$.

3.1.1 Matrix norms

A very powerful modelling tool in data science is low rank matrices. As already suggested in the expansion (3.2) the SVD will play an important role in this, being used to provide low rank approximation of data matrices.

In order to be able to talk about low rank approximations of matrices, we need a notion of distance between matrices. Just like with vectors, the distance between matrices can be measured using a suitable norm of the difference. One popular norm is the Frobenius norm, or the Hilbert-Schmidt norm, defined as

$$\|A\|_F = \sqrt{\sum_{i,j} a_{ij}^2}, \quad (3.3)$$

where a_{ij} are the entries of the matrix A . This norm is simply the Euclidean norm of a vector of length mn comprised of the matrix elements. The Frobenius norm can also be expressed in terms of the singular values. To see this, first express the Frobenius norm in terms of the trace of $A^T A$ as

$$\|A\|_F^2 = \sum_{i,j} a_{ij}^2 = \text{Tr}(A^T A), \quad (3.4)$$

where we recall that the trace of a square matrix A is defined as

$$\text{Tr}(A) = \sum_i a_{ii}. \quad (3.5)$$

A particularly important property of the trace is that for any A of size $m \times n$ and B of size $n \times m$

$$\text{Tr}(AB) = \text{Tr}(BA). \quad (3.6)$$

Note that this implies that, e.g., $\text{Tr}(ABC) = \text{Tr}(CAB)$, but it does not imply that, e.g., $\text{Tr}(ABC) = \text{Tr}(ACB)$ which is not true in general. Now, plugging the SVD (3.1) into (3.4) gives

$$\|A\|_F^2 = \text{Tr}(A^T A) = \text{Tr}(V \Sigma^T U^T U \Sigma V^T) = \text{Tr}(\Sigma^T \Sigma) = \sum_{k=1}^r \sigma_k^2, \quad (3.7)$$

where we used the orthogonality of U and V and the trace property (3.6). We conclude that the Frobenius norm of a matrix equals the Euclidean norm of its vector of singular values.

A different way to define the size of a matrix is by viewing it as an operator and measuring by how much it can dilate vectors. For example, the operator 2-norm is defined as

$$\|A\|_2 = \sup_{\|x\|=1} \|Ax\|. \quad (3.8)$$

This operator norm can also be succinctly expressed in terms of the singular values. Indeed, for any $x \in \mathbb{R}^n$

$$Ax = \sum_{k=1}^r \sigma_k u_k (v_k^T x). \quad (3.9)$$

Using the orthogonality of the left singular vectors u_k we get

$$\|Ax\|^2 = \sum_{k=1}^r \sigma_k^2 \langle v_k, x \rangle^2 \leq \sigma_1^2 \sum_{k=1}^r \langle v_k, x \rangle^2 \leq \sigma_1^2 \sum_{k=1}^n \langle v_k, x \rangle^2 = \sigma_1^2 \|x\|^2, \quad (3.10)$$

where the last equality is due to the orthogonality of the right singular vectors v_k . Moreover, we get equality by choosing $x = v_1$. We conclude that the 2-norm is simply the largest singular value

$$\|A\|_2 = \sigma_1. \quad (3.11)$$

Before proceeding we briefly recall Hölder's inequality

Proposition 3.3. *For $p, q \in (0, \infty)$ such that $\frac{1}{p} + \frac{1}{q} = 1$, and $x, y \in \mathbb{R}^n$, we have*

$$\sum_{k=1}^n |x_k y_k| \leq \left(\sum_{k=1}^n |x_k|^p \right)^{\frac{1}{p}} \left(\sum_{k=1}^n |y_k|^q \right)^{\frac{1}{q}}.$$

The following useful form is equivalent:

$$\left(\sum_{k=1}^n |c_k| |a_k|^r |b_k|^s \right)^{r+s} \leq \left(\sum_{k=1}^n |c_k| |a_k|^{r+s} \right)^r \left(\sum_{k=1}^n |c_k| |b_k|^{r+s} \right)^s, \quad (3.12)$$

for non-negative r and s .

3.1.2 Existence of the SVD

The matrix 2-norm plays an important role in proving that any matrix has an SVD.

Theorem 3.4. *If A is a real $m \times n$ matrix, then there exist orthogonal matrices*

$$U = [u_1, \dots, u_m] \in O(m) \text{ and } V = [v_1, \dots, v_n] \in O(n)$$

such that

$$U^T A V = \text{diag}(\sigma_1, \dots, \sigma_p) \in \mathbb{R}^{m \times n}, \quad p = \min\{m, n\}$$

where $\sigma_1 \geq \sigma_2 \geq \dots \geq \sigma_p \geq 0$.

Proof. We will use induction to prove the existence of SVD. Let $x \in \mathbb{R}^n$ and $y \in \mathbb{R}^m$ be unit 2-norm vectors that satisfy $Ax = \sigma y$ where $\sigma = \|A\|_2$, the 2-norm of A . There exist matrices $V_2 \in \mathbb{R}^{n \times (n-1)}$ and $U_2 \in \mathbb{R}^{m \times (m-1)}$ such that $V = [x, V_2] \in \mathbb{R}^{n \times n}$ and $U = [y, U_2] \in \mathbb{R}^{m \times m}$ are orthogonal matrices (this follows from a simple result in linear algebra that given a basis of a subspace, it is possible to extend it to a basis of the whole space). Consider now the matrix $U^T A V$ which has the following structure

$$U^T A V = \begin{bmatrix} \sigma & w^T \\ 0 & B \end{bmatrix} = A_1.$$

A short explanation: we can write $U^T A V = [y, U_2]^T A [x, V_2]$. Thus, the top left element in the resulting matrix is $y^T A x = y^T \sigma y = \sigma$ (because $Ax = \sigma y$ and $\|y\| = 1$). The leftmost column underneath σ is all 0's, as we have $U_2^T A x = \sigma U_2^T y = 0$ by orthogonality of U_2 and y . The topmost column to the right of σ is some vector w^T , and the bottom right of the matrix is represented by a smaller matrix $B \in \mathbb{R}^{(m-1) \times (n-1)}$. Since

$$\left\| A_1 \begin{bmatrix} \sigma \\ w \end{bmatrix} \right\|^2 = \left\| \begin{bmatrix} \sigma^2 + \|w\|^2 \\ Bw \end{bmatrix} \right\|^2 \geq (\sigma^2 + \|w\|^2)^2$$

we have $\|A_1\|_2^2 \geq (\sigma^2 + \|w\|^2)$. However, $\|A_1\|^2 = \|A\|_2^2 = \sigma^2$, because the matrix 2-norm is invariant under orthogonal transformations. Therefore, $\|w\|^2 = 0$ and

$$U^T A V = \begin{bmatrix} \sigma & 0 \\ 0 & B \end{bmatrix}.$$

The proof is completed by applying induction on the smaller matrix B , resulting in the next largest σ on the diagonal in each iteration, and we eventually end up with $U^T A V = \text{diag}(\sigma_1, \dots, \sigma_p)$ with $\sigma_1 \geq \dots \geq \sigma_p \geq 0$.

3.1.3 Low rank matrix approximation

A very important property of the SVD is that it provides the best low rank approximation of a matrix, when the approximation error is measured in terms of either the operator norm or the Frobenius norm. Specifically, for any $0 \leq s \leq r$ consider the rank- s matrix $A_s = \sum_{k=1}^s \sigma_k u_k v_k^T$. Then, among all matrices of

rank s , A_s best approximates A in terms of the operator norm (or Frobenius norm) error. When measuring error using the operator norm, the approximation error equals σ_{s+1} , which is the largest singular value among the remaining $r - s$ smallest singular values

$$\|A - A_s\|_2 = \inf_{B \in \mathbb{R}^{m \times n}, \text{rank}(B) \leq s} \|A - B\|_2 = \sigma_{s+1}. \quad (3.13)$$

A similar result holds for the Frobenius norm, with the approximation error given in terms of the remaining $r - s$ smallest singular values as

$$\|A - A_s\|_F = \inf_{B \in \mathbb{R}^{m \times n}, \text{rank}(B) \leq s} \|A - B\|_F = \sqrt{\sum_{k=s+1}^r \sigma_k^2}. \quad (3.14)$$

In fact, A_s is the best low rank approximation for any univariate matrix norm satisfying $\|UAV\| = \|A\|$ for any $U \in O(m), V \in O(n)$, that is, norms that are invariant to multiplication by orthogonal matrices [126].

We prove the low rank approximation property for the operator norm case.

Theorem 3.5. *Let the SVD of a matrix $A \in \mathbb{R}^{m \times n}$ be $A = U\Sigma V^T$. If $s < r = \text{rank}(A)$ and we have the rank- s matrix A_s ,*

$$A_s = \sum_{k=1}^s \sigma_k u_k v_k^T$$

then

$$\min_{\text{rank}(B) \leq s} \|A - B\|_2 = \|A - A_s\|_2 = \sigma_{s+1}.$$

Proof. Since A_s and A have the same left and right singular vectors and the same s leading singular values, it follows that $U^T A_s V = \text{diag}(\sigma_1, \dots, \sigma_s, 0, \dots, 0)$ and $U^T (A - A_s) V = \text{diag}(0, \dots, 0, \sigma_{s+1}, \dots, \sigma_p)$, where $p = \min\{m, n\}$. This implies $\|A - A_s\|_2 = \sigma_{s+1}$, the largest remaining singular value. Now, suppose we have a matrix $B \in \mathbb{R}^{m \times n}$ with $\text{rank}(B) = s$. We can find an orthonormal basis for the null space of B , i.e. orthonormal vectors $x_1, \dots, x_{n-s}$ such that $\text{null}(B) = \text{span}\{x_1, \dots, x_{n-s}\}$. By a dimension argument, we have that

$$\text{span}\{x_1, \dots, x_{n-s}\} \cap \text{span}\{v_1, \dots, v_{s+1}\} \neq \{0\},$$

since these are two subspaces of $\mathbb{R}^n$ of dimension $n - s$ and $s + 1$. Take a unit vector z in this intersection; for this vector, we know that $Bz = 0$ and $Az = \sum_{k=1}^{s+1} \sigma_k u_k (v_k^T z)$. Therefore,

$$\|A - B\|_2^2 \geq \|(A - B)z\|^2 = \|Az\|^2 = \sum_{k=1}^{s+1} \sigma_k^2 (v_k^T z)^2 \geq \sigma_{s+1}^2 \sum_{k=1}^{s+1} (v_k^T z)^2 = \sigma_{s+1}^2$$

which completes the proof.

For the Frobenius norm case, the low rank approximation property follows from von Neumann's trace inequality [187]:

Theorem 3.6. *If A and B are complex-valued $n \times n$ matrices with singular values $\sigma_1(A), \dots, \sigma_n(A)$ and $\sigma_1(B), \dots, \sigma_n(B)$, then*

$$|\operatorname{Tr}(AB)| \leq \sum_{k=1}^n \sigma_k(A) \sigma_k(B).$$

The reader is referred to [127] for a proof of Theorem 3.6. We are now ready to prove that the SVD provides the best low rank approximation in the Frobenius norm.

Theorem 3.7. *Let the SVD of a matrix $A \in \mathbb{R}^{m \times n}$ be $A = U \Sigma V^T$. If $s < r = \operatorname{rank}(A)$ and we have the rank- s matrix A_s ,*

$$A_s = \sum_{k=1}^s \sigma_k u_k v_k^T$$

then

$$\min_{\operatorname{rank}(B) \leq s} \|A - B\|_F = \|A - A_s\|_F = \sqrt{\sum_{k=s+1}^r \sigma_k^2}.$$

Proof. Let B be any matrix with $\operatorname{rank}(B) = s < r = \operatorname{rank}(A)$. Notice that

$$\|A - B\|_F^2 = \|A\|_F^2 + \|B\|_F^2 - 2 \operatorname{Tr}(A^T B) = \sum_{k=1}^r \sigma_k^2(A) + \sum_{k=1}^s \sigma_k^2(B) - 2 \operatorname{Tr}(A^T B).$$

By von Neumann's trace inequality (whose statement for square matrices implies the general case of rectangular matrices)

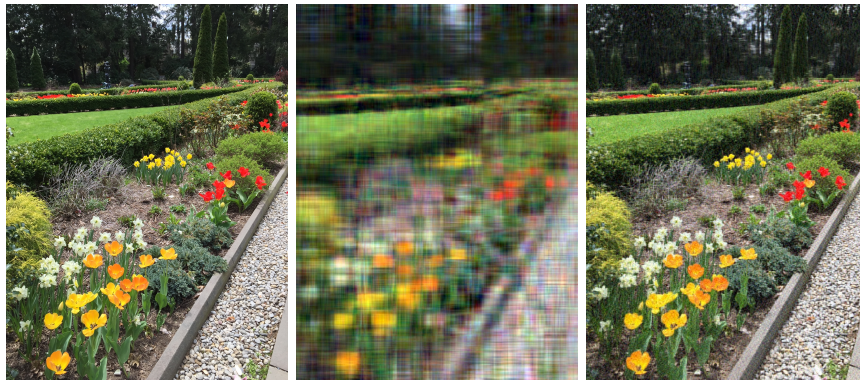
$$\operatorname{Tr}(A^T B) \leq \sum_{k=1}^s \sigma_k(A) \sigma_k(B).$$

Therefore,

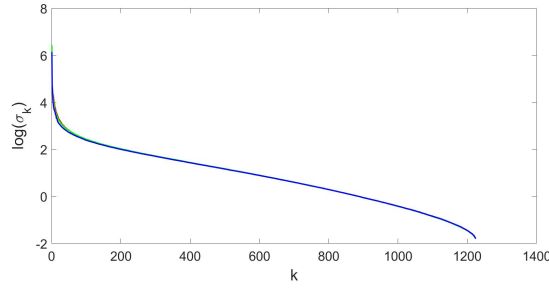
$$\begin{aligned} \|A - B\|_F^2 &\geq \sum_{k=1}^r \sigma_k^2(A) + \sum_{k=1}^s \sigma_k^2(B) - 2 \sum_{k=1}^s \sigma_k(A) \sigma_k(B) \\ &= \sum_{k=s+1}^r \sigma_k^2(A) + \sum_{k=1}^s (\sigma_k(A) - \sigma_k(B))^2 \\ &\geq \sum_{k=s+1}^r \sigma_k^2(A), \end{aligned}$$

which establishes the lower bound. Finally, the lower bound is clearly attained for $B = A_s$.

The low rank approximation property has a wide ranging implication on data compression. The storage size of an $m \times n$ data matrix is mn . If that matrix is of rank r , then storage size reduces from mn to $(n + m + 1)r$ (for storing r left and right singular vectors and values). For $r \ll \min\{n, m\}$ this reduction can be quite dramatic. For example, if $r = 10$ and $n = m = 10^6$, then storage reduces from 10^{12} entries to just $2 \cdot 10^7$. But even if the matrix is not precisely of rank r , but only approximately, in the sense that $\sigma_{r+1} \ll \sigma_1$, then we are guaranteed by the above approximation results to incur only a small approximation error due to compression using the top r singular vectors and values. In many cases, the singular values of large data matrices decay very quickly, motivating this type of low rank approximation which oftentimes is the only way to handle massive data sets that otherwise cannot be stored and/or manipulated efficiently. Remarkably, even treating an image as a matrix of pixel intensity values and compressing it this way yields good image compression and de-noising algorithms (as it mitigates the noise corresponding to singular values that are truncated). This is illustrated using a photo of size 1224×1632 pixels taken at Prospect Garden in Princeton University on April 28, 2015 by Amit Singer.



The original photo (left panel) is considered as three matrices (one for each RGB color channel). The main features of the image appear already in the rank-10 approximation (center panel), although many fine details are clearly missing and there are also some visible artifacts (such as striping, smearing, etc.). The rank-100 approximation (right panel) has exceptional resemblance with the original photo. Of course there are still differences between the two images (such as artifacts on the curbstone) but these are less visible to the naked eye on first inspection. Note that compression using $r = 10$ and $r = 100$ require less than 1% and 10%, respectively, storage of the original image. The decay of the singular values of the three matrices is illustrated below (in red, green, and blue). We use logarithmic scale due to the rapid decay of singular values. Notice that the singular values of the three matrices are very similar.



Remark 3.8. The computational complexity of computing the SVD of a matrix of size $m \times n$ with $m \geq n$ is $\mathcal{O}(mn^2)$. This cubic scaling could be prohibitive for massive data matrices, in a future version of this manuscript we plan to discuss numerical algorithms that use randomization to efficiently compute the low rank approximation of such large matrices more efficiently.

Remark 3.9. The SVD plays a critical role in sensitivity analysis of linear systems of the form $Ax = b$. In particular, the error is bounded in terms of the condition number $\kappa(A)$ which is the ratio of the largest singular value and the smallest singular value. A rigorous treatment of this important topic is beyond the scope of this book and can be found in other texts on numerical linear algebra [82].

3.1.4 Spectral Decomposition

If $M \in \mathbb{R}^{n \times n}$ is symmetric then it admits a spectral decomposition

$$M = V\Lambda V^T,$$

where $V \in O(n)$ is a matrix whose columns v_k are the eigenvectors of M and Λ is a diagonal matrix whose diagonal elements λ_k are the eigenvalues of M . Similarly, we can write

$$M = \sum_{k=1}^n \lambda_k v_k v_k^T.$$

When all of the eigenvalues of M are non-negative we say that M is positive semidefinite and write $M \succeq 0$. In that case we can write

$$M = \left(V\Lambda^{1/2}\right) \left(V\Lambda^{1/2}\right)^T.$$

A decomposition of M of the form $M = UU^T$ (such as the one above) is called a Cholesky decomposition (oftentimes, in a Cholesky decomposition, the matrix U is required to be triangular, here we do not make this requirement).

For symmetric matrices, the operator 2-norm is also known as the spectral norm, given by

$$\|M\| = \max_k |\lambda_k(M)|.$$

3.1.5 Quadratic forms

In both this and following chapters, we will be interested in solving problems of the type

$$\max_{\substack{V \in \mathbb{R}^{n \times d} \\ V^T V = I_{d \times d}}} \text{Tr}(V^T M V),$$

where M is a symmetric $n \times n$ matrix.

Note that this is equivalent to

$$\max_{\substack{v_1, \dots, v_d \in \mathbb{R}^n \\ v_i^T v_j = \delta_{ij}}} \sum_{k=1}^d v_k^T M v_k, \quad (3.15)$$

where δ is the Kronecker delta ($\delta_{ij} = 1$ for $i = j$ and $\delta_{ij} = 0$ otherwise).

When $d = 1$ this reduces to the more familiar

$$\max_{\substack{v \in \mathbb{R}^n \\ \|v\|_2=1}} v^T M v. \quad (3.16)$$

It is easy to see (for example, using the spectral decomposition of M) that (3.16) is maximized by the leading eigenvector of M and

$$\max_{\substack{v \in \mathbb{R}^n \\ \|v\|_2=1}} v^T M v = \lambda_{\max}(M).$$

Furthermore (3.15) is maximized by taking $v_1, \dots, v_d$ to be the d leading eigenvectors of M and its value is simply the sum of the d largest eigenvalues of M . This follows, for example, from a Theorem of Fan (see page 3 of [131]). A fortunate consequence is that the solution to (3.15) can be computed sequentially: we can first solve for $d = 1$, computing v_1 , then update the solution for $d = 2$ by simply computing v_2 without changing the first vector.

Remark 3.10. All of the tools and results above have natural analogues when the matrices have complex entries (and are Hermitian instead of symmetric).

3.1.6 SVD and the Moore-Penrose Pseudoinverse

The SVD is intimately connected to the Moore-Penrose pseudoinverse of a matrix [36]. The Moore-Penrose pseudoinverse (named after mathematicians E. H. Moore and Roger Penrose) of A , denoted as $A^\dagger$, is a generalization of the matrix inverse to matrices that may not be square, or if square, may not have full rank. For an invertible matrix, the pseudoinverse is identical to the standard inverse. However, for a rectangular or rank-deficient square matrix, the pseudoinverse provides a unique solution that satisfies a set of four specific conditions known as the Penrose conditions, that uniquely define the pseudoinverse $A^\dagger$ for any matrix A . These conditions are a generalization of

the properties of a standard matrix inverse. A matrix $A^\dagger$ is the pseudoinverse of A if and only if it satisfies all four of the following equations:

$$\begin{aligned} AA^\dagger A &= A, \\ A^\dagger AA^\dagger &= A^\dagger, \\ (AA^\dagger)^T &= AA^\dagger, \\ (A^\dagger A)^T &= A^\dagger A. \end{aligned}$$

These conditions ensure that the pseudoinverse acts “as much as possible” as an inverse for matrices that are not invertible, including rectangular or rank-deficient square matrices. The above four conditions are equivalent to the matrices $AA^\dagger$ and $A^\dagger A$ being orthogonal projectors onto the range of A and the range of A^T respectively.

The primary purpose of the pseudoinverse is to find “best-fit” solution to a system of linear equations $Ax = b$ when a unique solution does not exist. In cases where the system has multiple solutions, the pseudoinverse provides the one with the minimum Euclidean norm. For systems with no exact solution (an overdetermined system), it provides the least-squares solution, which minimizes the error $\|Ax - b\|_2$. This makes the pseudoinverse a tremendously useful tool in fields like numerical mathematics, statistics, machine learning, and signal processing.

The SVD provides a straightforward and computationally stable way to calculate $A^\dagger$. That is, the pseudoinverse of A is given by the formula.

$$A^\dagger = V \Sigma^\dagger U^T, \quad (3.17)$$

where the matrix $\Sigma^\dagger$ is found by taking the reciprocal of every non-zero singular value on the diagonal of Σ . This method is particularly powerful because it works for any matrix—square or rectangular, full rank or rank-deficient—and is numerically robust, avoiding the issues with ill-conditioned matrices that can plague other inversion methods. Indeed, in the practical computation of $\Sigma^\dagger$ one usually chooses a threshold τ , such that all singular values less than τ are treated as zero.

3.2 Principal Component Analysis and dimension reduction

When faced with a high dimensional dataset, a natural approach is to attempt to reduce its dimension, either by projecting it to a lower dimensional space or by finding a better representation for the data using a small number of meaningful features. Beyond data compression and visualization, dimension reduction facilitates downstream analysis such as clustering and regression that perform significantly better in lower dimensions. We will explore a few different ways of reducing the dimension, both linearly and non-linearly.

We start with the classical Principal Component Analysis (PCA). PCA continues to be one of the most effective and simplest tools for exploratory data analysis. It dates back to a 1901 paper by Karl Pearson [143] and was rediscovered several times by different research communities where it is also known under different names such as the proper orthogonal decomposition (POD) and the Karhunen–Loève transform (KLT) among others.

Suppose we have n data points $x_1, \dots, x_n$ in $\mathbb{R}^p$, and we are interested in (linearly) projecting the data to $d < p$ dimensions. This is particularly useful if, say, one wants to visualize the data in two or three dimensions ($d = 2, 3$). There are a couple of seemingly different criteria we can use to choose this projection:

1. Finding the d -dimensional affine subspace for which the projections of $x_1, \dots, x_n$ on it best approximate the original points $x_1, \dots, x_n$.
2. Finding the d -dimensional projection of $x_1, \dots, x_n$ that preserves as much variance of the data as possible.

As we will see below, these two approaches are equivalent and they correspond to Principal Component Analysis.

Before proceeding, we recall a couple of simple statistical quantities associated with $x_1, \dots, x_n$, that will reappear below.

Given $x_1, \dots, x_n$ we define its sample mean as

$$\mu_n = \frac{1}{n} \sum_{k=1}^n x_k, \quad (3.18)$$

and its sample covariance as

$$\Sigma_n = \frac{1}{n-1} \sum_{k=1}^n (x_k - \mu_n)(x_k - \mu_n)^T. \quad (3.19)$$

Remark 3.11. If $x_1, \dots, x_n$ are independently sampled from a distribution, μ_n and Σ_n are unbiased estimators for, respectively, the mean and covariance of the distribution.

PCA as the best d -dimensional affine fit

We start with the first interpretation of PCA and then show that it is equivalent to the second. We are trying to approximate each x_k by

$$x_k \approx \mu + \sum_{i=1}^d (\beta_k)_i v_i, \quad (3.20)$$

where $v_1, \dots, v_d$ is an orthonormal basis for the d -dimensional subspace, $\mu \in \mathbb{R}^p$ represents the translation, and $\beta_k \in \mathbb{R}^d$ corresponds to the coefficients of x_k . Without loss of generality we take

$$\sum_{k=1}^n \beta_k = 0, \quad (3.21)$$

as any joint translation of β_k can be absorbed into μ .

If we represent the subspace by $V = [v_1 \cdots v_d] \in \mathbb{R}^{p \times d}$ then we can rewrite (3.22) as

$$x_k \approx \mu + V\beta_k, \quad (3.22)$$

where $V^T V = I_{d \times d}$, because the vectors v_i are orthonormal.

We will measure goodness of fit in terms of least squares and attempt to solve

$$\min_{\substack{\mu, V, \beta_k \\ V^T V = I}} \sum_{k=1}^n \|x_k - (\mu + V\beta_k)\|_2^2 \quad (3.23)$$

We start by optimizing for μ . It is easy to see that the first order condition for μ corresponds to

$$\nabla_{\mu} \sum_{k=1}^n \|x_k - (\mu + V\beta_k)\|_2^2 = 0 \iff \sum_{k=1}^n (x_k - (\mu + V\beta_k)) = 0.$$

Thus, the optimal value μ^* of μ satisfies

$$\left(\sum_{k=1}^n x_k \right) - n\mu^* - V \left(\sum_{k=1}^n \beta_k \right) = 0.$$

Since we assumed in (3.21) that $\sum_{k=1}^n \beta_k = 0$, we have that the optimal μ is given by

$$\mu^* = \frac{1}{n} \sum_{k=1}^n x_k = \mu_n,$$

the sample mean.

We can then proceed to finding the solution for (3.23) by solving

$$\min_{\substack{V, \beta_k \\ V^T V = I}} \sum_{k=1}^n \|x_k - \mu_n - V\beta_k\|_2^2. \quad (3.24)$$

Let us proceed by optimizing for β_k . The problem almost fully decouples in each k , the only constraint coupling them being (3.21). We will ignore this constraint, solve the decoupled problems, and verify that it is automatically satisfied. Hence we focus on, for each k ,

$$\min_{\beta_k} \|x_k - \mu_n - V\beta_k\|_2^2 = \min_{\beta_k} \left\| x_k - \mu_n - \sum_{i=1}^d (\beta_k)_i v_i \right\|_2^2. \quad (3.25)$$

Since $v_1, \dots, v_d$ are orthonormal, it is easy to see that the solution is given by $(\beta_k^*)_i = v_i^T (x_k - \mu_n)$ which can be succinctly written as $\beta_k^* = V^T (x_k - \mu_n)$, which satisfies (3.21). Thus, (3.24) is equivalent to

$$\min_{V^T V = I} \sum_{k=1}^n \left\| (x_k - \mu_n) - V V^T (x_k - \mu_n) \right\|_2^2. \quad (3.26)$$

Note that

$$\begin{aligned} \left\| (x_k - \mu_n) - V V^T (x_k - \mu_n) \right\|_2^2 &= (x_k - \mu_n)^T (x_k - \mu_n) \\ &\quad - 2 (x_k - \mu_n)^T V V^T (x_k - \mu_n) \\ &\quad + (x_k - \mu_n)^T V (V^T V) V^T (x_k - \mu_n) \\ &= (x_k - \mu_n)^T (x_k - \mu_n) \\ &\quad - (x_k - \mu_n)^T V V^T (x_k - \mu_n). \end{aligned}$$

Since $(x_k - \mu_n)^T (x_k - \mu_n)$ does not depend on V , minimizing (3.26) is equivalent to

$$\max_{V^T V = I} \sum_{k=1}^n (x_k - \mu_n)^T V V^T (x_k - \mu_n). \quad (3.27)$$

A few algebraic manipulations using properties of the trace yields:

$$\begin{aligned} \sum_{k=1}^n (x_k - \mu_n)^T V V^T (x_k - \mu_n) &= \sum_{k=1}^n \text{Tr} \left[(x_k - \mu_n)^T V V^T (x_k - \mu_n) \right] \\ &= \sum_{k=1}^n \text{Tr} \left[V^T (x_k - \mu_n) (x_k - \mu_n)^T V \right] \\ &= \text{Tr} \left[V^T \sum_{k=1}^n (x_k - \mu_n) (x_k - \mu_n)^T V \right] \\ &= (n-1) \text{Tr} [V^T \Sigma_n V]. \end{aligned}$$

This means that the solution to (3.27) is given by the solution of

$$\max_{V^T V = I} \text{Tr} [V^T \Sigma_n V]. \quad (3.28)$$

As we saw above (recall (3.15)) the solution is given by $V = [v_1, \dots, v_d]$ where $v_1, \dots, v_d$ correspond to the d leading eigenvectors of Σ_n .

PCA as the d -dimensional projection that preserves the most variance

We now show that the alternative interpretation of PCA, of finding the d -dimensional projection of $x_1, \dots, x_n$ that preserves the most variance, also

arrives to the optimization problem (3.28). We aim to find an orthonormal basis $v_1, \dots, v_d$ (organized as $V = [v_1, \dots, v_d]$ with $V^T V = I_{d \times d}$) of a d -dimensional space such that the projection of $x_1, \dots, x_n$ onto this subspace has the most variance. Equivalently we can ask for the points

$$\left\{ \begin{bmatrix} v_1^T x_k \\ \vdots \\ v_d^T x_k \end{bmatrix} \right\}_{k=1}^n,$$

to have as much variance as possible. Hence, we are interested in solving

$$\max_{V^T V = I} \sum_{k=1}^n \left\| V^T x_k - \frac{1}{n} \sum_{r=1}^n V^T x_r \right\|^2. \quad (3.29)$$

Note that

$$\sum_{k=1}^n \left\| V^T x_k - \frac{1}{n} \sum_{r=1}^n V^T x_r \right\|^2 = \sum_{k=1}^n \|V^T (x_k - \mu_n)\|^2 = (n-1) \operatorname{Tr}(V^T \Sigma_n V),$$

showing that (3.29) is equivalent to (3.28) and that the two interpretations of PCA are indeed equivalent.

Finding the Principal Components

When given a dataset $x_1, \dots, x_n \in \mathbb{R}^p$, in order to compute the Principal Components one needs to compute the leading eigenvectors of

$$\Sigma_n = \frac{1}{n-1} \sum_{k=1}^n (x_k - \mu_n)(x_k - \mu_n)^T.$$

A naïve way of doing this is to construct Σ_n (which takes $\mathcal{O}(np^2)$ work) and then finding its spectral decomposition (which takes $\mathcal{O}(p^3)$ work). This means that the computational complexity of this procedure is $\mathcal{O}(\max\{np^2, p^3\})$ (see [93] or [81]).

An alternative is to use the Singular Value Decomposition (3.1). Let $X = [x_1 \cdots x_n]$, and recall that

$$\Sigma_n = \frac{1}{n-1} (X - \mu_n \mathbf{1}^T) (X - \mu_n \mathbf{1}^T)^T.$$

Let us take the SVD of $\frac{1}{\sqrt{n-1}} (X - \mu_n \mathbf{1}^T) = U_L D U_R^T$ with $U_L \in O(p)$, D diagonal, and $U_R^T U_R = I$. Then,

$$\Sigma_n = \frac{1}{n-1} (X - \mu_n \mathbf{1}^T) (X - \mu_n \mathbf{1}^T)^T = U_L D U_R^T U_R D U_L^T = U_L D^2 U_L^T,$$

meaning that U_L corresponds to the eigenvectors of Σ_n . Computing the SVD of $X - \mu_n \mathbf{1}^T$ takes $\mathcal{O}(\min\{n^2p, p^2n\})$ work but if one is interested in simply computing the top d eigenvectors then this computational cost reduces to $\mathcal{O}(dnp)$. This can be further improved with randomized algorithms. There are randomized algorithms that compute an approximate solution in $\mathcal{O}(pn \log d + (p+n)d^2)$ time (We plan to discuss this in more detail in a future version of this manuscript. Nonetheless, you can read more about this in [88, 150, 135]).

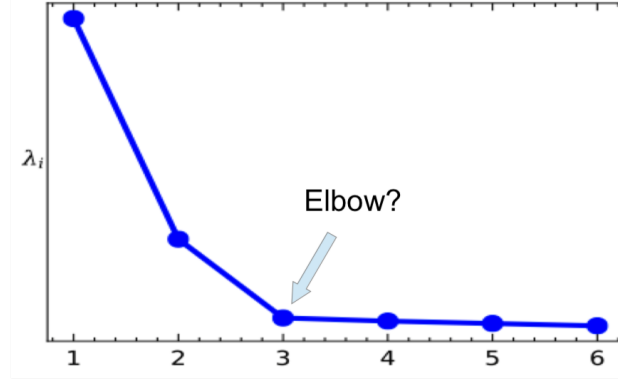
Numerical stability is another important reason why computing the principal components using the SVD is preferable. Since the eigenvalues of Σ_n are proportional to the squares of the singular values of $X - \mu_n \mathbf{1}^T$, problems arise when the ratio of singular values exceeds 10^8 , causing the ratio of the corresponding eigenvalues of Σ_n to be larger than 10^{16} . In this case, the smaller eigenvalue would be rounded to zero (due to machine precision), which is certainly not desirable.

Which d should we pick?

Given a dataset, if the objective is to visualize it then picking $d = 2$ or $d = 3$ might make the most sense. However, PCA is useful for many other purposes, for example:

1. Denoising: often times the data belongs to a lower dimensional space but is corrupted by high dimensional noise. In such cases, PCA helps reduce the noise while keeping the signal.
2. Downstream analysis: One may be interested in running an algorithm (clustering, regression, etc.) that would be too computationally expensive or too statistically insignificant to run in high dimensions. Dimension reduction using PCA may help there.

In these applications (and many others) it is not clear how to pick d . A fairly popular heuristic is to try to choose the cut-off at a component that has significantly more variance than the one immediately after. Since the total variance is $\text{Tr}(\Sigma_n) = \sum_{k=1}^p \lambda_k$, the proportion of variance in the i 'th component is nothing but $\frac{\lambda_i}{\text{Tr}(\Sigma_n)}$. A plot of the values of the ordered eigenvalues, also known as a scree plot, helps identify a reasonable choice of d . Here is an example:



It is common to then try to identify an “elbow” on the scree plot to choose the cut-off. In the next Section we will look into Random Matrix Theory to better understand the behavior of the eigenvalues of Σ_n and gain insight into choosing cut-off values.

3.3 PCA in high dimensions and Marčenko-Pastur law

Let us assume that the data points $x_1, \dots, x_n \in \mathbb{R}^p$ are independent draws of a zero mean Gaussian random variable $g \sim \mathcal{N}(0, \Sigma)$ with some covariance matrix $\Sigma \in \mathbb{R}^{p \times p}$. In this case, when we use PCA we are hoping to find a low dimensional structure in the distribution, which should correspond to the large eigenvalues of Σ (and their corresponding eigenvectors). For that reason, and since PCA depends on the spectral properties of Σ_n , we would like to understand whether the spectral properties of the sample covariance matrix Σ_n (eigenvalues and eigenvectors) are close to those of Σ , also known as the population covariance.

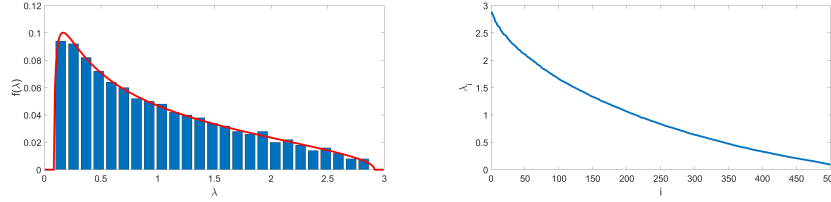
Since $\mathbb{E}\Sigma_n = \Sigma$, if p is fixed and $n \rightarrow \infty$ the law of large numbers guarantees that indeed $\Sigma_n \rightarrow \Sigma$. However, in many modern applications it is not uncommon to have p of the same order as n (or, sometimes, even larger). For example, if our dataset is composed of images then n is the number of images and p the number of pixels per image; it is conceivable that the number of pixels and the number of images are comparable. Unfortunately, in that case, it is no longer clear that $\Sigma_n \rightarrow \Sigma$. Dealing with this type of difficulties is an important goal of high dimensional statistics.

For simplicity we will try to understand the spectral properties of

$$S_n = \frac{1}{n} X X^T,$$

where $x_1, \dots, x_n$ are the columns of X . Since $x \sim \mathcal{N}(0, \Sigma)$ we know that $\mu_n \rightarrow 0$ (and, clearly, $\frac{n}{n-1} \rightarrow 1$), hence the spectral properties of S_n will be essentially the same as Σ_n .¹

Let us start by looking into a simple example, $\Sigma = I$. In that case, the distribution has no low dimensional structure, as the distribution is rotation invariant. The following is a normalized histogram (left panel) and a scree plot (right panel) of the eigenvalues of a realization of S_n (when $\Sigma = I$) for $p = 500$ and $n = 1000$. The red line is the eigenvalue distribution predicted by the Marčenko-Pastur distribution (3.30), that we will discuss below.



As one can see in the figure, there are many eigenvalues considerably larger than 1, as well as many eigenvalues significantly smaller than 1. Notice that, if given this profile of eigenvalues of Σ_n one could potentially be led to believe that the data has low dimensional structure, when in truth the distribution it was drawn from is isotropic.

Understanding the distribution of eigenvalues of random matrices is at the core of Random Matrix Theory (there are many good books on Random Matrix Theory, e.g. [17, 169, 13]). This particular limiting distribution was first established in 1967 by Marčenko and Pastur [122] and is now referred to as the Marčenko-Pastur distribution. They showed that, if p and n both grow indefinitely to ∞ with their ratio fixed $p/n = \gamma \leq 1$, the sample distribution of the eigenvalues of S_n (like the histogram above), in the limit, will be

$$dF_\gamma(\lambda) = \frac{1}{2\pi} \frac{\sqrt{(\gamma_+ - \lambda)(\lambda - \gamma_-)}}{\gamma\lambda} 1_{[\gamma_-, \gamma_+]}(\lambda) d\lambda, \quad (3.30)$$

with support $[\gamma_-, \gamma_+]$, where $\gamma_- = (1 - \sqrt{\gamma})^2$, $\gamma_+ = (1 + \sqrt{\gamma})^2$, and $\gamma = p/n$. This is plotted as the red line in the figure above.

Remark 3.12. We will not provide the proof of the Marčenko-Pastur law here (you can see, for example, [16] for several different proofs of it), but an approach to a proof is using the so-called moment method. The central idea is to note that one can compute moments of the eigenvalue distribution in two different ways and note that (in the limit) for any k ,

$$\frac{1}{p} \mathbb{E} \operatorname{Tr} \left[\left(\frac{1}{n} X X^T \right)^k \right] = \frac{1}{p} \mathbb{E} \operatorname{Tr} (S_n^k) = \mathbb{E} \frac{1}{p} \sum_{i=1}^p \lambda_i^k(S_n) = \int_{\gamma_-}^{\gamma_+} \lambda^k dF_\gamma(\lambda),$$

¹In this case, S_n is actually the maximum likelihood estimator for Σ .

and that the quantities $\frac{1}{p}\mathbb{E}\text{Tr}\left[\left(\frac{1}{n}XX^T\right)^k\right]$ can be estimated (these estimates rely essentially on combinatorics). The distribution $dF_\gamma(\lambda)$ can then be computed from its moments.

3.3.1 Spike models and the BBP phase transition

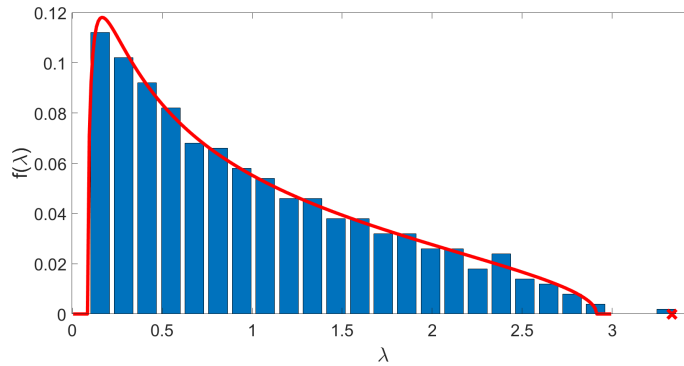
What if there actually is some (linear) low dimensional structure in the data? When can we expect to capture it with PCA? A particularly simple, yet relevant, example to analyze is when the covariance matrix Σ is a rank-1 perturbation of the identity matrix and is of the form $\Sigma = I + \beta uu^T$, where u is a unit norm vector and $\beta > 0$. This is a particular case of a spike model.

One way to think about this spike model is that each data point x consists of a signal part $\sqrt{\beta}g_0u$ where g_0 is a scalar standard Gaussian $\mathcal{N}(0, 1)$ (i.e. a normally distributed multiple of a fixed vector $\sqrt{\beta}u$) and a noise part $g \sim \mathcal{N}(0, I)$ (independent of g_0). Then $x = g + \sqrt{\beta}g_0u$ is a Gaussian random variable

$$x \sim \mathcal{N}(0, I + \beta uu^T).$$

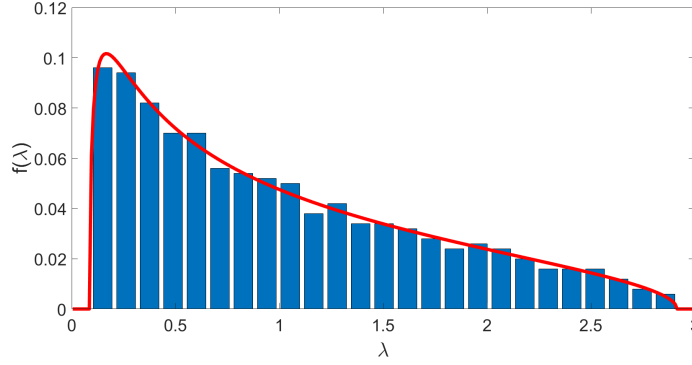
Whereas the signal part $\sqrt{\beta}g_0u$ resides on a central line in the direction of u , the noise part is high dimensional and isotropic. We therefore refer to β as the signal-to-noise ratio (SNR). Indeed, β is the ratio of the signal variance (in the u -direction) to the noise variance (in each direction).

A natural question is whether this rank-1 perturbation can be seen in S_n . In other words, is it possible to detect the direction of the line u from noisy measurements in high dimension? Let us build some intuition with an example. The following is the histogram of the eigenvalues of a random realization of S_n for $p = 500$, $n = 1000$, u is the first element of the canonical basis $u = e_1$, and $\beta = 1.5$:



The histogram suggests that there is an eigenvalue of S_n that “pops out” of the support of the Marčenko-Pastur distribution (below we will estimate the location of this eigenvalue, and that estimate corresponds to the red “x”).

It is worth noting that the largest eigenvalues of Σ is simply $1 + \beta = 2.5$ while the largest eigenvalue of S_n appears considerably larger than that. Let us try now the same experiment with $\beta = 0.5$:



It appears that, for $\beta = 0.5$, the histogram of the eigenvalues is indistinguishable from when $\Sigma = I$. In particular, no eigenvalue is separated from the Marčenko-Pastur distribution.

This motivates the following question:

Question 3.13. For which values of γ and β do we expect to see an eigenvalue of S_n popping out of the support of the Marčenko-Pastur distribution, and what is the limiting value that we expect it to take?

As we will see below, there is a critical value of β , denoted β_c , below which we do not expect to see a change in the distribution of eigenvalues and above which we expect one of the eigenvalues to pop outside of the support. This phenomenon is known as the BBP phase transition (after Baik, Ben Arous, and P\'ech\'e [18]). There are many very nice papers about this and similar phenomena, including [141, 97, 18, 142, 19, 98, 37, 38].²

In what follows we will find the critical value β_c and estimate the location of the largest eigenvalue of S_n for any β . While the argument we will use can be made precise (and is borrowed from [141]) we will be ignoring a few details for the sake of exposition. In other words, the argument below can be transformed into a rigorous proof, but it is not one at the present form.

We want to understand the behavior of the leading eigenvalue of the sample covariance matrix

$$S_n = \frac{1}{n} \sum_{i=1}^n x_i x_i^T.$$

²Notice that the Marčenko-Pastur theorem does not imply that all eigenvalues are actually in the support of the Marčenko-Pastur distribution, it just rules out that a non-vanishing proportion are. However, it is possible to show that indeed, in the limit, all eigenvalues will be in the support (see, for example, [141]).

Since $x \sim \mathcal{N}(0, I + \beta uu^T)$ we can write $x = (I + \beta uu^T)^{1/2} z$ where $z \sim \mathcal{N}(0, I)$ is an isotropic Gaussian. Then,

$$S_n = \frac{1}{n} \sum_{i=1}^n (I + \beta uu^T)^{1/2} z_i z_i^T (I + \beta uu^T)^{1/2} = (I + \beta uu^T)^{1/2} Z_n (I + \beta uu^T)^{1/2},$$

where $Z_n = \frac{1}{n} \sum_{i=1}^n z_i z_i^T$ is the sample covariance matrix of independent isotropic Gaussians. The matrices $S_n = (I + \beta uu^T)^{1/2} Z_n (I + \beta uu^T)^{1/2}$ and $Z_n (I + \beta uu^T)$ are related by a similarity transformation, and therefore have exactly the same eigenvalues. Hence, it suffices to find the leading eigenvalue of the matrix $Z_n (I + \beta uu^T)$, which is a rank-1 perturbation of Z_n (indeed, $Z_n (I + \beta uu^T) = Z_n + \beta Z_n uu^T$). We already know that the eigenvalues of Z_n follow the Marčenko-Pastur distribution, so we are left to understand the effect of a rank-1 perturbation on its eigenvalues.

To find the leading eigenvalue λ of $Z_n (I + \beta uu^T)$, let v be the corresponding eigenvector, that is,

$$Z_n (I + \beta uu^T) v = \lambda v.$$

Subtract $Z_n v$ from both sides to get

$$\beta Z_n uu^T v = (\lambda I - Z_n) v.$$

Assuming λ is not an eigenvalue of Z_n , we can multiply by $(\lambda I - Z_n)^{-1}$ to get³

$$\beta (\lambda I - Z_n)^{-1} Z_n uu^T v = v.$$

Our assumption also implies that $u^T v \neq 0$, for otherwise $v = 0$. Multiplying by u^T gives

$$\beta u^T (\lambda I - Z_n)^{-1} Z_n u (u^T v) = u^T v.$$

Dividing by $\beta u^T v$ (which is not 0 as explained above) yields

$$u^T (\lambda I - Z_n)^{-1} Z_n u = \frac{1}{\beta}. \quad (3.31)$$

Suppose $w_1, \dots, w_p$ are orthonormal eigenvectors of Z_n (with corresponding eigenvalues $\lambda_1, \dots, \lambda_p$), and expand u in that basis:

$$u = \sum_{i=1}^p \alpha_i w_i.$$

Plugging this expansion in (3.31) gives

$$\sum_{i=1}^p \frac{\lambda_i}{\lambda - \lambda_i} \alpha_i^2 = \frac{1}{\beta} \quad (3.32)$$

³Intuitively, λ is larger than all the eigenvalues of Z_n , because it corresponds to a perturbation of Z_n by a positive definite matrix βuu^T ; yet, a formal justification is beyond the present discussion.

For large p , each α_i^2 concentrates around its mean value $\mathbb{E}[\alpha_i^2] = \frac{1}{p}$ (again, this statement can be made rigorous), and (3.32) becomes

$$\lim_{p \rightarrow \infty} \frac{1}{p} \sum_{i=1}^p \frac{\lambda_i}{\lambda - \lambda_i} = \frac{1}{\beta} \quad (3.33)$$

Since the eigenvalues $\lambda_1, \dots, \lambda_p$ follow the Marčenko-Pastur distribution, the limit on the left hand side can be replaced by the integral

$$\int_{\gamma_-}^{\gamma_+} \frac{t}{\lambda - t} dF_\gamma(t) = \frac{1}{\beta} \quad (3.34)$$

Using an integral table (or an integral software), we find that

$$\frac{1}{\beta} = \int_{\gamma_-}^{\gamma_+} \frac{t}{\lambda - t} dF_\gamma(t) = \frac{1}{4\gamma} \left[2\lambda - (\gamma_- + \gamma_+) - 2\sqrt{(\lambda - \gamma_-)(\lambda - \gamma_+)} \right]. \quad (3.35)$$

For $\lambda = \gamma_+$, that is, when the top eigenvalue touches the right edge of the Marčenko-Pastur distribution, (3.35) becomes $\frac{1}{4\gamma}(\gamma_+ - \gamma_-)$. This is the critical point that one gets the pop out of the top eigenvalue from the bulk of the Marčenko-Pastur distribution. To calculate the critical value β_c , we recall that $\gamma_- = (1 - \sqrt{\gamma})^2$ and $\gamma_+ = (1 + \sqrt{\gamma})^2$, hence

$$\frac{1}{\beta_c} = \frac{1}{4\gamma} \left((1 + \sqrt{\gamma})^2 - (1 - \sqrt{\gamma})^2 \right). \quad (3.36)$$

Therefore, the critical SNR is

$$\beta_c = \sqrt{\gamma} = \sqrt{\frac{p}{n}}. \quad (3.37)$$

When $\beta > \sqrt{\frac{p}{n}}$ one can observe the pop out of the top eigenvalue from the bulk.

Eq. (3.37) illustrates the interplay of the SNR β , the number of samples n , and the dimension p . Low SNR, small sample size, and high dimensionality are all obstacles for detecting linear structure in noisy high dimensional data.

More generally, inverting the relationship between β and λ given by (3.35) (which simply amounts to solving a quadratic), we find that the largest eigenvalue λ of the sample covariance matrix S_n has the limiting value

$$\lambda \rightarrow \begin{cases} (\beta + 1) \left(1 + \frac{\gamma}{\beta} \right) & \text{for } \beta \geq \sqrt{\gamma}, \\ (1 + \sqrt{\gamma})^2 & \text{for } \beta < \sqrt{\gamma}. \end{cases} \quad (3.38)$$

In the finite sample case λ will be fluctuating around that value.

Notice that the critical SNR value, $\beta_c = \sqrt{\gamma}$ is buried deep inside the support of the Marčenko-Pastur distribution, because $\sqrt{\gamma} < \gamma_+ = (1 + \sqrt{\gamma})^2$.

In other words, the SNR does not have to be greater than the operator norm of the noise matrix in order for it to pop out. We see that the noise effectively pushes the eigenvalue to the right (indeed, $\lambda > \beta$).

The asymptotic squared correlation $|\langle u, v \rangle|^2$ between the top eigenvector v of the sample covariance matrix and true signal vector u can be calculated in a similar fashion. The limiting correlation value turns out to be

$$|\langle v, u \rangle|^2 \rightarrow \begin{cases} \frac{1 - \frac{\gamma}{\beta^2}}{1 + \frac{\gamma}{\beta^2}} & \text{for } \beta \geq \sqrt{\gamma} \\ 0 & \text{for } \beta < \sqrt{\gamma} \end{cases} \quad (3.39)$$

Notice that the correlation value tends to 1 as $\beta \rightarrow \infty$, but is strictly less than 1 for any finite SNR.

3.3.2 Wigner matrices

Another very important random matrix model is the Wigner matrix (and it will make appearances in Chapters 8, 9 and 7). Given an integer n , a standard Gaussian Wigner matrix $W \in \mathbb{R}^{n \times n}$ is a symmetric matrix with independent $\mathcal{N}(0, 1)$ off-diagonal entries (except for the fact that $W_{ij} = W_{ji}$) and jointly independent $\mathcal{N}(0, 2)$ diagonal entries. In the limit, the eigenvalues of $\frac{1}{\sqrt{n}}W$ are distributed according to the so-called semi-circular law

$$dSC(x) = \frac{1}{2\pi} \sqrt{4 - x^2} 1_{[-2, 2]}(x) dx, \quad (3.40)$$

and there is also a BBP like transition for this matrix ensemble [75]. More precisely, if v is a unit-norm vector in $\mathbb{R}^n$ and $\xi \geq 0$ then the largest eigenvalue of $\frac{1}{\sqrt{n}}W + \xi vv^T$ satisfies

- If $\xi \leq 1$ then

$$\lambda_{\max} \left(\frac{1}{\sqrt{n}}W + \xi vv^T \right) \rightarrow 2,$$

- and if $\xi > 1$ then

$$\lambda_{\max} \left(\frac{1}{\sqrt{n}}W + \xi vv^T \right) \rightarrow \xi + \frac{1}{\xi}. \quad (3.41)$$

The typical correlation, with v , of the leading eigenvector $v_{\max}$ of $\frac{1}{\sqrt{n}}W + \xi vv^T$ is also known:

- If $\xi \leq 1$ then

$$|\langle v_{\max}, v \rangle|^2 \rightarrow 0,$$

- and if $\xi > 1$ then

$$|\langle v_{\max}, v \rangle|^2 \rightarrow 1 - \frac{1}{\xi^2}.$$

From a statistical viewpoint, a central question is to understand for different distributions of matrices, when it is possible to detect and estimate a spike in a random matrix [144]. When the underlying random matrix corresponds to the adjacency matrix of a random graph and the spike to a bias on the distribution of the graph edges, corresponding to structural properties of the graph, the estimates above are able to predict important phase transitions in community detection in networks, as we will see in Chapter 7.

3.3.3 Rank and covariance estimation

The spike model and random matrix theory thus offer a principled way for determining the number of principal components, or equivalently of the rank of the hidden linear structure: simply count the number of eigenvalues to the right of the Marčenko-Pastur distribution. In practice, this approach for rank estimation is often too simplistic for several reasons. First, for actual datasets, n and p are finite, and one needs to take into account non-asymptotic corrections and finite sample fluctuations [104, 105]. Second, the noise may be heteroskedastic (that is, noise variance is different in different directions). Moreover, the noise statistics could also be unknown and it can be non-Gaussian [117]. In some situations it might be possible to estimate the noise statistics directly from the data and to homogenize the noise (a procedure sometimes known as “whitening”) [115]. These situations call for careful analysis, and many open problems remain in the field.

Another popular method for rank estimation is using permutation methods. In permutation methods, each column of the data matrix is randomly permuted, so that the low-rank linear structure in the data is destroyed through scrambling, while only the noise is preserved. The process can be repeated multiple times, and the statistics of the singular values of the scrambled data matrices are then used to determine the rank. In particular, only singular values of the original (unscrambled) data matrix that are larger than the largest singular value of the scrambled matrices (taking fluctuations into account of course) are considered as corresponding to signal and are counted towards the rank. The mathematical analysis of permutation methods is another active field of research [64, 65].

In some applications, the objective is to estimate the low rank covariance matrix of the clean signal Σ from the noisy measurements. We saw that in the spike model, the eigenvalues of the sample covariance matrix are inflated due to noise. It is therefore required to shrink the computed eigenvalues of S_n in order to obtain a better estimate of the eigenvalues of Σ . That is, if

$$S_n = \sum_{i=1}^p \lambda_i v_i v_i^T$$

is the spectral decomposition of S_n , then we seek an estimator of Σ , denoted $\hat{\Sigma}$ of the form

$$\hat{\Sigma} = \sum_{i=1}^p \eta(\lambda_i) v_i v_i^T.$$

The scalar nonlinearity $\eta : \mathbb{R}^+ \rightarrow \mathbb{R}^+$ is known as the shrinkage function. An obvious shrinkage procedure is to estimate $\beta = \eta(\lambda)$ from the computed λ by inverting (3.38) (and setting $\beta = 0$ for $\lambda < \gamma_+$). It turns out that this particular shrinker is optimal in terms of the operator norm loss. However, for other loss functions (such as the Frobenius norm loss), the optimal shrinkage function takes a different form [69]. The reason why the shrinker depends on the loss function is that the eigenvectors of S_n are not perfectly correlated with those of Σ but rather make some non-trivial angle, as in (3.39). In other words, the eigenvectors are noisy, and it may require more aggressive shrinkage to account for that error in the eigenvector. It can be shown that the eigenvector v of the sample covariance is uniformly distributed in a cone around u whose opening angle is given by (3.39). While we can improve the estimation of the eigenvalue via shrinkage, it is however unclear how to improve the estimation of the eigenvector (without any a priori knowledge about it). Finally, we remark that eigenvalue shrinkage also plays an important role in denoising.

Graphs, Networks, and Clustering

A crucial part of data science consists of the studying of networks. Network science, or graph theory, unifies the study of diverse types of networks, such as social networks, protein-protein interaction networks, gene-regulation networks, and the internet. In this chapter we introduce graph theory and treat the problem of clustering, to identify similar data points, or vertices, in (network) data.

4.1 PageRank

Before we introduce the formalism of graph theory, we describe the celebrated PageRank algorithm. This algorithm is a principal component¹ behind the web search algorithms, in particular in Google. The goal of PageRank is to quantitatively rate the importance of each page on the web, allowing the search algorithm to rank the pages and thereby present to the user the more important pages first. Search engines such as Google have to carry out three basic steps:²

- Crawl the web and locate all, or as many as possible, accessible webpages.
- Index the data of the webpages from step 1, so that they can be searched efficiently for relevant key words or phrases.
- Rate the importance of each page in the database, so that when a user does a search and the subset of pages in the database with the desired information has been found, the more important pages can be presented first.

Here, we will focus on the third step. We follow mainly the derivation in [43]. We aim to develop a score of importance for each webpage. A score will be a

¹It is difficult to resist using this pun.

²Another important component of modern search engines is personalization, which we do not discuss here.

non-negative number. A key idea in assigning a score to any given webpage is that the page's score is derived from the links made to that page from other webpages — “A person is important not if it knows a lot of people, but if a lot of people know that person”.

Suppose the web of interest contains n pages, each page indexed by an integer k , $1 \leq k \leq n$. A typical example is illustrated in Figure 4.1, in which an arrow from page k to page j indicates a link from page k to page j . Such a web is an example of a directed graph. The links to a given page are called the backlinks for that page. We will use x_k to denote the importance score of page k in the web. x_k is nonnegative and $x_j > x_k$ indicates that page j is more important than page k .

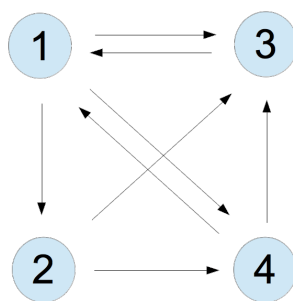


Fig. 4.1: A toy example of the Internet

A very simple approach is to take x_k as the number of backlinks for page k . In the example in Figure 4.1, we have $x_1 = 2$, $x_2 = 1$, $x_3 = 3$, and $x_4 = 2$, so that page 3 is the most important, pages 1 and 4 tie for second, and page 2 is least important. A link to page k becomes a vote for page k 's importance. This approach ignores an important feature one would expect a ranking algorithm to have, namely, that a link to page k from an important page should boost page k 's importance score more than a link from an unimportant page. In the web of Figure 4.1, pages 1 and 4 both have two backlinks: each links to the other, but the second backlink from page 1 is from the seemingly important page 3, while the second backlink for page 4 is from the relatively unimportant page 2. As such, perhaps the algorithm should rate the importance of page 1 higher than that of page 4.

As a first attempt at incorporating this idea, let us compute the score of page j as the sum of the scores of all pages linking to page j . For example, consider the web in our toy example. The score of page 1 would be determined by the relation $x_1 = x_3 + x_4$. However, since x_3 and x_4 will depend on x_1 , this seems like a circular definition, since it is self-referential (it is exactly

this self-referential property that will establish a connection to eigenvector problems!).

We also seek a scheme in which a webpage does not gain extra influence simply by linking to lots of other pages. We can do this by reducing the impact of each link, as more and more outgoing links are added to a webpage. If page j contains n_j links, one of which links to page k , then we will boost page k 's score by x_j/n_j , rather than by x_j . In this scheme, each webpage gets a total of one vote, weighted by that web page's score, that is evenly divided up among all of its outgoing links. To quantify this for a web of n pages, let $L_k \subset \{1, 2, \dots, n\}$ denote the set of pages with a link to page k , that is, L_k is the set of page k 's backlinks. For each k we require

$$x_k = \sum_{j \in L_k} \frac{x_j}{n_j},$$

where n_j is the number of outgoing links from page j .

If we apply these scheme to the toy example in Figure 4.1, then for page 1 we have $x_1 = x_3/1 + x_4/2$, since pages 3 and 4 are backlinks for page 1 and page 3 contains only one link, while page 4 contains two links (splitting its vote in half). Similarly, $x_2 = x_1/3$, $x_3 = x_1/3 + x_2/2 + x_4/2$, and $x_4 = x_1/3 + x_2/2$. These conditions can be expressed as linear system of equations $Ax = x$, where $x = [x_1, x_2, x_3, x_4]^T$ and

$$A = \begin{bmatrix} 0 & 0 & 1 & \frac{1}{2} \\ \frac{1}{3} & 0 & 0 & 0 \\ \frac{1}{3} & \frac{1}{2} & 0 & \frac{1}{2} \\ \frac{1}{3} & \frac{1}{2} & 0 & 0 \end{bmatrix}$$

Thus, we end up with an eigenvalue/eigenvector problem: Find the eigenvector x of the matrix A , associated with the eigenvalue 1. We note that A is a column-stochastic matrix, since it is a square matrix for which all of its entries are nonnegative and the entries in each column sum to 1. If we build a random walk on the internet where each link is clicked with equal probability then A_{ij} is the probability that a random walked in page j goes to page i . Later, in Chapter 5, we will use random walks on (undirected) graphs in order to embed their nodes in euclidean space, we note that, in that context, we will mostly work with $M = A^T$ the matrix for which M_{ij} corresponds to the probability of taking step from i to j . Stochastic matrices arise in the study of Markov chains and in a variety of modeling problems in economics and operations research. See e.g. [94] for more details on stochastic matrices. The fact that 1 is an eigenvalue of A is not just coincidence in this example, but holds true in general for stochastic matrices.

Theorem 4.1. *A column-stochastic matrix A has an eigenvalue equal to 1 and 1 is also its largest eigenvalue.*

Proof. Let A be an $n \times n$ column-stochastic matrix. We first note that A and A^T have the same eigenvalues (their eigenvector will usually be different though). Let $\mathbf{1} = [1, 1, \dots, 1]^T$ be the vector of length n which has all ones as entries. Since A is column-stochastic, we have $A^T \mathbf{1} = \mathbf{1}$ (since all columns of A sum up to 1). Hence $\mathbf{1}$ is an eigenvector of A^T (but not of A) with eigenvalue 1. Thus 1 is also an eigenvalue of A .

To show that 1 is the largest eigenvalue of A we apply the Gershgorin Circle Theorem [94] to A^T . Consider row k of A^T . Let us call the diagonal element $a_{k,k}$ and the radius will be $\sum_{i \neq k} |a_{k,i}| = \sum_{i \neq k} a_{k,i}$ since $a_{k,i} \geq 0$. This is a circle with its center at $a_{k,k} \in [0, 1]$ and with radius $\sum_{i \neq k} a_{k,i} = 1 - a_{k,k}$. Hence, this circle has 1 on its perimeter. This holds for all Gershgorin circles for this matrix. Thus, since all eigenvalues lie in the union of the Gershgorin circles, all eigenvalues λ_i satisfy $|\lambda_i| \leq 1$.

In our example, we obtain as eigenvector x of A associated with eigenvalue 1 the vector $x = [x_1, x_2, x_3, x_4]^T$ with entries $x_1 = \frac{12}{31}, x_2 = \frac{4}{31}, x_3 = \frac{9}{31}$, and $x_4 = \frac{6}{31}$. Hence, perhaps somewhat surprisingly, page 3 is no longer the most important one, but page 1. This can be explained by the fact, that the in principle quite important page 3 (which has three webpages linking to it) has only one outgoing link, which gets all its “voting power”, and that link points to page 1.

In reality, A can easily be of size a billion times a billion. Fortunately, we do not need compute all eigenvectors of A , only the eigenvector associated with the eigenvalue 1, which, as we know, is also the largest eigenvalue of A . This in turn means we can resort to standard *power iteration* to compute x fairly efficiently (and we can also make use of the fact that A will be a sparse matrix, i.e., many of its entries will be zero). The actual PageRank algorithms adds some minor modifications, but the essential idea is as described above.

The idea to use eigenvectors for ranking dates back to the late 1800s to the work of Edmund Landau in the context of ranking in chess tournaments, it was actually Landau’s first mathematical paper when he was 18 years old. You can read about this nice story in [158].

4.2 Graph theory

We now introduce the formalism for undirected³ graphs, one of the main objects of study in what follows. A graph $G = (V, E)$ contains a set of nodes $V = \{v_1, \dots, v_n\}$ and edges $E \subseteq \binom{V}{2}$. An edge $(i, j) \in E$ if v_i and v_j are connected. Figure 4.2 depicts one of the graph theorists’ favorite examples, the Petersen graph⁴.

³The previous Chapter featured directed graphs, in which edges (links) have a meaningful direction. In what follows we will focus in undirected graphs in which an edge represents a connection, without meaningful direction.

⁴The Petersen graph is often used as a counter-example in graph theory.

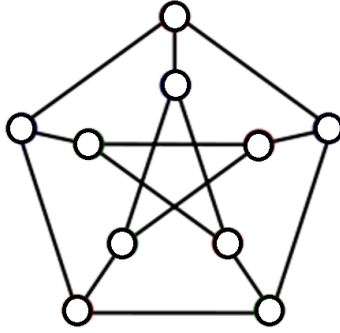


Fig. 4.2: The Petersen graph

Let us recall some concepts about graphs that we will need.

- A graph is *connected* if, for all pairs of vertices, there is a path between these vertices in the graph. The number of connected components is simply the size of the smallest partition of the nodes into connected subgraphs. The Petersen graph is connected (and thus it has only 1 connected component).
- A *clique* of a graph G is a subset S of its nodes such that the subgraph corresponding to it is complete. In other words S is a clique if all pairs of vertices in S share an edge. The clique number $c(G)$ of G is the size of the largest clique of G . The Petersen graph has a clique number of 2.
- An *independent set* of a graph G is a subset S of its nodes such that no two nodes in S share an edge. Equivalently it is a clique of the complement graph $G^c := (V, E^c)$. The independence number of G is simply the clique number of S^c . The Petersen graph has an independence number of 4.

A particularly useful way to represent a graph is through its adjacency matrix. Given a graph $G = (V, E)$ on n nodes ($|V| = n$), we define its adjacency matrix $A \in \mathbb{R}^{n \times n}$ as the symmetric matrix with entries

$$A_{ij} = \begin{cases} 1 & \text{if } (i, j) \in E, \\ 0 & \text{otherwise.} \end{cases}$$

Sometimes, we will consider weighted graphs $G = (V, E, W)$, where edges may have weights w_{ij} that are non-negative $w_{ij} \geq 0$ and symmetric $w_{ij} = w_{ji}$.

Much of the sequel will deal with graphs. Chapter 5 will treat (network) data visualization, dimension reduction, and embeddings of graphs in Euclidean space. Chapter 7 will introduce and study important random graph models. The rest of this Chapter will be devoted to spectral graph theory, clustering, and graph metrics.

4.3 Clustering

Clustering is one of the central tasks in machine learning. Given a set of data points, or nodes of a graph, the purpose of clustering is to partition the data into a set of clusters where data points assigned to the same cluster correspond to similar data points (depending on the context, it could be for example having small distance to each other if the points are in Euclidean space, or having high connectivity if on a graph). We will start with an example of clustering points in Euclidean space, and later move back to graphs.

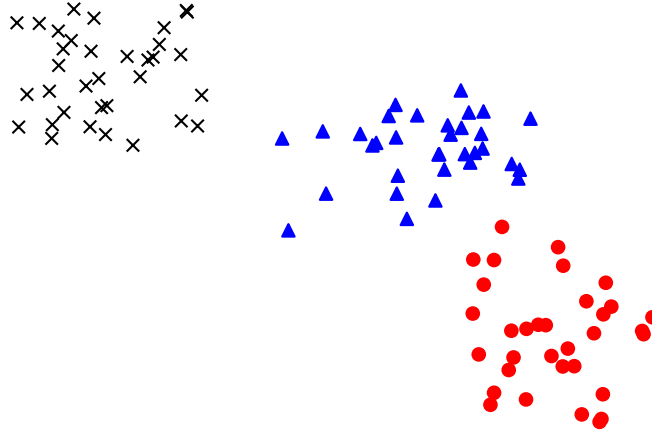


Fig. 4.3: Examples of points separated into clusters.

4.3.1 k -means Clustering

One of the most popular methods used for clustering is k -means clustering. Given $x_1, \dots, x_n \in \mathbb{R}^p$ the k -means clustering partitions the data points into clusters $S_1 \cup \dots \cup S_k$ with centers $\mu_1, \dots, \mu_k \in \mathbb{R}^p$ as the solution to:

$$\min_{\text{partition } \substack{S_1, \dots, S_k \\ \mu_1, \dots, \mu_k}} \sum_{l=1}^k \sum_{i \in S_l} \|x_i - \mu_l\|^2. \quad (4.1)$$

Note that, given the partition, the optimal centers are given by

$$\mu_l = \frac{1}{|S_l|} \sum_{i \in S_l} x_i.$$

Lloyd's algorithm [118] (also sometimes known as the k -means algorithm), is an iterative algorithm that alternates between

- Given centers $\mu_1, \dots, \mu_k$, assign each point x_i to the cluster S_l with nearest center

$$l = \operatorname{argmin}_{l=1, \dots, k} \|x_i - \mu_l\|.$$

- Update the centers $\mu_l = \frac{1}{|S_l|} \sum_{i \in S_l} x_i$.

Unfortunately, Lloyd's algorithm is not guaranteed to converge to the solution of (4.1). Indeed, Lloyd's algorithm oftentimes gets stuck in local optima of (4.1). In the sequel we will discuss convex relaxations for clustering, which can be used as an alternative algorithmic approach to Lloyd's algorithm, but since optimizing (4.1) is NP -hard there is no polynomial time algorithm that works in the worst-case (assuming the widely believed conjecture $P \neq NP$, see also Chapter 7.1)

While popular, k -means clustering has some potential issues:

- One needs to set the number of clusters a priori (a typical way to overcome this issue is by trying the algorithm for different number of clusters).
- The way (4.1) is defined it needs the points to be defined in an Euclidean space, oftentimes we are interested in clustering data for which we only have some measure of affinity between different data points, but not necessarily an embedding in $\mathbb{R}^p$ (this issue can be overcome by reformulating (4.1) in terms of distances only).
- The formulation is computationally hard, so algorithms may produce sub-optimal instances.
- The solutions of k -means are always convex clusters. This means that k -means may have difficulty in finding cluster such as in Figure 4.4.

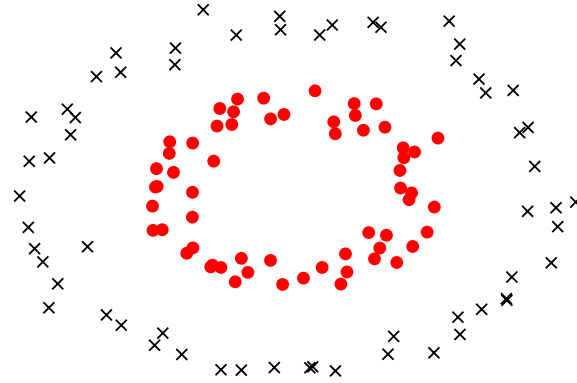


Fig. 4.4: Because the solutions of k -means are always convex clusters, it is not able to handle some cluster structures.

4.3.2 Spectral Clustering

A natural way to try to overcome the issues of k -means depicted in Figure 4.4 is by transforming the data into a graph and cluster the graph: Given the data points we can construct a weighted graph $G = (V, E, W)$ using a similarity kernel K_ε , such as $K_\varepsilon(u) = \exp(-\frac{1}{2\varepsilon}u^2)$, by associating each point to a vertex and, for which pair of nodes, set the edge weight as

$$w_{ij} = K_\varepsilon(\|x_i - x_j\|).$$

Another popular procedure to transform data into a graph is by constructing a graph where data points are connected if they correspond to the nearest neighbors. We note that this procedure only needs a measure of distance, or similarity, of data points and not necessarily that they lie in an Euclidean space. Given this motivation, and the prevalence of network data, we will now address the problem of clustering the nodes of a graph.

Normalized Cut

Given a graph $G = (V, E, W)$, the goal is to partition the graph into clusters in a way that keeps as many of the edges, or connections, within the clusters and has as few edges as possible across clusters. We will focus on the case of two clusters, and briefly address extensions at the end of this Chapter. A natural way to measure a vertex partition (S, S^c) is

$$\text{cut}(S) = \sum_{i \in S} \sum_{j \in S^c} w_{ij}.$$

If we represent the partition by a vector $y \in \{\pm 1\}^n$ where $y_i = 1$ if $i \in S$, and $y_i = -1$ otherwise, then the cut is a quadratic form on the Graph Laplacian.

Definition 4.2 (Graph Laplacian and Degree Matrix). *Let $G = (V, E, W)$ be a graph and W the matrix of weights (or adjacency matrix if the graph is unweighted). The degree matrix D is a diagonal matrix with diagonal entries*

$$D_{ii} = \deg(i) = \sum_j w_{ij}.$$

The graph Laplacian of G is given by

$$L_G = D - W.$$

Equivalently

$$L_G := \sum_{i < j} w_{ij} (e_i - e_j) (e_i - e_j)^T.$$

Note that the entries of L_G are given by

$$(L_G)_{ij} = \begin{cases} -w_{ij} & \text{if } i \neq j \\ \deg(i) & \text{if } i = j. \end{cases}$$

If $S \subset V$ and $y \in \{\pm 1\}^n$ such that $y_i = 1$ if $i \in S$, and $y_i = -1$ otherwise, then it is easy to see that

$$\text{cut}(S) = \frac{1}{4} \sum_{i < j} w_{ij} (y_i - y_j)^2.$$

The following proposition establishes

$$\text{cut}(S) = \frac{1}{4} y^T L_G y, \quad (4.2)$$

for $y \in \{\pm 1\}^n$ such that $y_i = 1$ if and only if $i \in S$.

Proposition 4.3. *Let $G = (V, E, W)$ be a graph and L_G its graph Laplacian. For any $x \in \mathbb{R}^n$*

$$x^T L_G x = \sum_{i < j} w_{ij} (x_i - x_j)^2.$$

Proof.

$$\begin{aligned} \sum_{i < j} w_{ij} (x_i - x_j)^2 &= \sum_{i < j} w_{ij} [x^T (e_i - e_j)] [(e_i - e_j)^T x] \\ &= \sum_{i < j} w_{ij} x^T (e_i - e_j) (e_i - e_j)^T x \\ &= x^T \left[\sum_{i < j} w_{ij} (e_i - e_j) (e_i - e_j)^T \right] x. \end{aligned}$$

□

Note that Proposition 4.3 implies that the graph Laplacian L_G is a positive semidefinite (PSD) matrix, since $x^T L_G x \geq 0$ for all x (assuming non-negative weights $w_{ij} \geq 0$). This property can also be deduced directly from Definition 4.2 of L_G as the (non-negative) weighted sum of rank-1 PSD matrices.

While $\text{cut}(S)$ is a good way of measuring the fit of a partition, it has a major drawback: the minimum cut is achieved for $S = \emptyset$ (since $\text{cut}(\emptyset) = 0$) which is a rather meaningless choice of partition. Constraining the partition to be non-trivial would not overcome this drawback, because very unbalanced partitions would still be favored (e.g., partitions with $|S| = 1$ containing just a single node). Below we discuss how to promote more balanced partitions.

Remark 4.4. One simple way to address this is to simply ask for an exactly balanced partition, $|S| = |S^c|$ (let us assume the number of vertices $n = |V|$ is

even). We can then identify a partition with a label vector $y \in \{\pm 1\}^n$ where $y_i = 1$ if $i \in S$, and $y_i = -1$ otherwise. Also, the balanced condition can be written as $\sum_{i=1}^n y_i = 0$. This means that we can write the minimum balanced cut as

$$\min_{\substack{S \subset V \\ |S|=|S^c|}} \text{cut}(S) = \frac{1}{4} \min_{\substack{y \in \{-1,1\}^n \\ \mathbf{1}^T y = 0}} y^T L_G y,$$

which is suggestive of the connection between clustering and spectral properties of L_G . This connection will be made precise below.

Asking for the partition to be exactly balanced is too restrictive in many cases. There are several ways to evaluate a partition that are variations of $\text{cut}(S)$ that take into account the intuition that one wants both S and S^c not to be too small (although not necessarily equal to $|V|/2$). A prime example is Cheeger's cut.

Definition 4.5 (Cheeger's cut). *Given a graph and a vertex partition (S, S^c) , the Cheeger cut (also known as conductance, or expansion) of S is given by*

$$h(S) = \frac{\text{cut}(S)}{\min\{\text{vol}(S), \text{vol}(S^c)\}},$$

where $\text{vol}(S) = \sum_{i \in S} \deg(i)$.

Also, the Cheeger's constant of G is given by

$$h_G = \min_{S \subset V} h(S).$$

A similar object is the Normalized Cut, Ncut , which is given by

$$\text{Ncut}(S) = \frac{\text{cut}(S)}{\text{vol}(S)} + \frac{\text{cut}(S^c)}{\text{vol}(S^c)}.$$

Note that $\text{Ncut}(S)$ and $h(S)$ are tightly related, in fact it is easy to see that:

$$h(S) \leq \text{Ncut}(S) \leq 2h(S).$$

Normalized Cut as a spectral relaxation

Below we will show that Ncut can be written in terms of a minimization of a quadratic form involving the graph Laplacian L_G , analogously to the balanced partition as described in Remark 4.4.

Recall that the cut value of a balanced partition can be written as

$$\frac{1}{4} \min_{\substack{y \in \{-1,1\}^n \\ \mathbf{1}^T y = 0}} y^T L_G y.$$

An intuitive way to relax the balanced condition is to allow the labels y to take values in two different real values a and b (e.g. $y_i = a$ if $i \in S$ and $y_j = b$ if $i \notin S$) but not necessarily ± 1 . We can then use the notion of volume of a set to ensure a less restrictive notion of balanced by asking that

$$a \operatorname{vol}(S) + b \operatorname{vol}(S^c) = 0, \quad (4.3)$$

where

$$\operatorname{vol}(S) = \sum_{i \in S} \deg(i). \quad (4.4)$$

Thus (4.3) corresponds to $\mathbf{1}^T D y = 0$.

We also need to fix a scale for a and b :

$$a^2 \operatorname{vol}(S) + b^2 \operatorname{vol}(S^c) = 1,$$

which corresponds to $y^T D y = 1$.

This suggests considering

$$\min_{\substack{y \in \{a,b\}^n \\ \mathbf{1}^T D y = 0, y^T D y = 1}} y^T L_G y.$$

As we will see below, this corresponds precisely to Ncut.

Proposition 4.6. *For a and b to satisfy $a \operatorname{vol}(S) + b \operatorname{vol}(S^c) = 0$ and $a^2 \operatorname{vol}(S) + b^2 \operatorname{vol}(S^c) = 1$ it must be that*

$$a = \left(\frac{\operatorname{vol}(S^c)}{\operatorname{vol}(S) \operatorname{vol}(G)} \right)^{\frac{1}{2}} \quad \text{and} \quad b = - \left(\frac{\operatorname{vol}(S)}{\operatorname{vol}(S^c) \operatorname{vol}(G)} \right)^{\frac{1}{2}},$$

corresponding to

$$y_i = \begin{cases} \left(\frac{\operatorname{vol}(S^c)}{\operatorname{vol}(S) \operatorname{vol}(G)} \right)^{\frac{1}{2}} & \text{if } i \in S \\ - \left(\frac{\operatorname{vol}(S)}{\operatorname{vol}(S^c) \operatorname{vol}(G)} \right)^{\frac{1}{2}} & \text{if } i \in S^c. \end{cases}$$

Note that vol is defined in (4.4).

Proof. The proof involves only doing simple algebraic manipulations together with noticing that $\operatorname{vol}(S) + \operatorname{vol}(S^c) = \operatorname{vol}(G)$. $\square$

Proposition 4.7.

$$\operatorname{Ncut}(S) = y^T L_G y,$$

where y is given by

$$y_i = \begin{cases} \left(\frac{\operatorname{vol}(S^c)}{\operatorname{vol}(S) \operatorname{vol}(G)} \right)^{\frac{1}{2}} & \text{if } i \in S \\ - \left(\frac{\operatorname{vol}(S)}{\operatorname{vol}(S^c) \operatorname{vol}(G)} \right)^{\frac{1}{2}} & \text{if } i \in S^c. \end{cases}$$

Proof.

$$\begin{aligned}
y^T L_G y &= \frac{1}{2} \sum_{i,j} w_{ij} (y_i - y_j)^2 \\
&= \sum_{i \in S} \sum_{j \in S^c} w_{ij} (y_i - y_j)^2 \\
&= \sum_{i \in S} \sum_{j \in S^c} w_{ij} \left[\left(\frac{\text{vol}(S^c)}{\text{vol}(S) \text{vol}(G)} \right)^{\frac{1}{2}} + \left(\frac{\text{vol}(S)}{\text{vol}(S^c) \text{vol}(G)} \right)^{\frac{1}{2}} \right]^2 \\
&= \sum_{i \in S} \sum_{j \in S^c} w_{ij} \frac{1}{\text{vol}(G)} \left[\frac{\text{vol}(S^c)}{\text{vol}(S)} + \frac{\text{vol}(S)}{\text{vol}(S^c)} + 2 \right] \\
&= \sum_{i \in S} \sum_{j \in S^c} w_{ij} \frac{1}{\text{vol}(G)} \left[\frac{\text{vol}(S^c)}{\text{vol}(S)} + \frac{\text{vol}(S)}{\text{vol}(S^c)} + \frac{\text{vol}(S)}{\text{vol}(S)} + \frac{\text{vol}(S^c)}{\text{vol}(S^c)} \right] \\
&= \sum_{i \in S} \sum_{j \in S^c} w_{ij} \left[\frac{1}{\text{vol}(S)} + \frac{1}{\text{vol}(S^c)} \right] \\
&= \text{cut}(S) \left[\frac{1}{\text{vol}(S)} + \frac{1}{\text{vol}(S^c)} \right] \\
&= \text{Ncut}(S).
\end{aligned}$$

□

This means that finding the minimum Ncut corresponds to solving

$$\begin{aligned}
&\min y^T L_G y \\
&\text{s. t. } y \in \{a, b\}^n \text{ for some } a \text{ and } b \\
&\quad y^T D y = 1 \\
&\quad y^T D \mathbf{1} = 0.
\end{aligned} \tag{4.5}$$

Since solving (4.5) is, in general, NP-hard, we consider a similar problem where the constraint that y can only take two values is removed:

$$\begin{aligned}
&\min y^T L_G y \\
&\text{s. t. } y \in \mathbb{R}^n \\
&\quad y^T D y = 1 \\
&\quad y^T D \mathbf{1} = 0.
\end{aligned} \tag{4.6}$$

Given a solution of (4.6) we can *round* it to a partition by setting a threshold τ and taking $S = \{i \in V : y_i \leq \tau\}$. We will see below that (4.6) is an eigenvector problem (for this reason we call (4.6) a spectral relaxation).

In order to better see that (4.6) is an eigenvector problem (and thus computationally tractable), set $z = D^{\frac{1}{2}} y$ and

$$\mathcal{L}_G = D^{-\frac{1}{2}} L_G D^{-\frac{1}{2}}, \tag{4.7}$$

then (4.6) is equivalent

$$\begin{aligned} \min \quad & z^T \mathcal{L}_G z \\ \text{s. t. } \quad & z \in \mathbb{R}^n \\ & \|z\|^2 = 1 \\ & \left(D^{\frac{1}{2}} \mathbf{1}\right)^T z = 0. \end{aligned} \tag{4.8}$$

Note that $\mathcal{L}_G = I - D^{-\frac{1}{2}} W D^{-\frac{1}{2}}$. We order its eigenvalues in increasing order $0 = \lambda_1(\mathcal{L}_G) \leq \lambda_2(\mathcal{L}_G) \leq \dots \leq \lambda_n(\mathcal{L}_G)$. The eigenvector associated with the smallest eigenvalue is given by $D^{\frac{1}{2}} \mathbf{1}$. By the variational interpretation of the eigenvalues, it follows that the minimum of (4.8) is $\lambda_2(\mathcal{L}_G)$ and the minimizer is given by the second smallest eigenvector of $\mathcal{L}_G = I - D^{-\frac{1}{2}} W D^{-\frac{1}{2}}$, which we call v_2 . Note that this corresponds also to the second largest eigenvector of $D^{-\frac{1}{2}} W D^{-\frac{1}{2}}$. This means that the optimal y in (4.6) is given by $\varphi_2 = D^{-\frac{1}{2}} v_2$. This motivates Algorithm 4.1, which is often referred to as Spectral Clustering:

Algorithm 4.1 Spectral Clustering

Given a graph $G = (V, E, W)$, let v_2 be the eigenvector corresponding to the second smallest eigenvalue of the normalized Laplacian $\mathcal{L}_G$, as defined in (4.7). Let $\varphi_2 = D^{-\frac{1}{2}} v_2$. Given a threshold τ (one can try all different possibilities, or run k -means for $k = 2$), set

$$S = \{i \in V : \varphi_2(i) \leq \tau\}.$$

4.3.3 Cheeger's Inequality

Because the relaxation (4.6) is obtained from (4.5) by removing a constraint we immediately have that

$$\lambda_2(\mathcal{L}_G) \leq \min_{S \subset V} \text{Ncut}(S).$$

This means that

$$\frac{1}{2} \lambda_2(\mathcal{L}_G) \leq h_G.$$

In what follows we will show a performance guarantee for Algorithm 4.1.

Lemma 4.8. *There is a threshold τ producing a partition S such that*

$$h(S) \leq \sqrt{2\lambda_2(\mathcal{L}_G)}.$$

This implies in particular that

$$h(S) \leq \sqrt{4h_G},$$

meaning that Algorithm 4.1 is suboptimal at most by a square-root factor. Note that this also directly implies the famous Cheeger's Inequality, stated next.

Theorem 4.9 (Cheeger's Inequality). *Recall the definitions above. The following holds:*

$$\frac{1}{2}\lambda_2(\mathcal{L}_G) \leq h_G \leq \sqrt{2\lambda_2(\mathcal{L}_G)}.$$

Cheeger's inequality was first established for manifolds by Jeff Cheeger in 1970 [54], and the graph version is due to Noga Alon and Vitaly Milman [8, 10] in the mid 80s. Cheeger's inequality for a closed manifold M gives a bound on the area of a hypersurface E , that partitions M into two disjoint parts.

Theorem 4.10 (Cheeger's Inequality (Cheeger 1970)). *Let h_M be the Cheeger isoperimetric constant of M , defined as*

$$h_M = \inf_E \frac{\text{area of } E}{\min\{\text{vol}(A), \text{vol}(B)\}}$$

where E is a smooth submanifold of M of $n-1$ dimensions that divides it into two disjoint submanifolds A and B . Let λ_M be the smallest positive eigenvalue of the Laplacian, or Laplace-Bertrami operator, on M . Then

$$\lambda_M \geq \frac{h_M^2}{4}$$

Bounding the diameter of a graph using $\lambda_2(\mathcal{L}_G)$

Before proving Cheeger's inequality for graphs, we first discuss another eigenvalue inequality, that gives a flavor of the connection between the second eigenvalue of the normalized graph Laplacian and the geometry of the graph.

Associate a cost $c(u, v)$ for each pair of nodes inversely proportional to the weight between them:

$$c(u, v) = \begin{cases} \frac{1}{w_{uv}} & \text{for } (u, v) \in E, \\ +\infty & \text{for } (u, v) \notin E. \end{cases} \quad (4.9)$$

A *path* of length k connecting u and v is of the form

$$P = (u = v_1, v_2, \dots, v_k = v).$$

The cost associated with this path is

$$c(P) = \sum_{i=1}^{k-1} c(v_i, v_{i+1}).$$

The *geodesic distance* or *shortest path* between u and v is the minimum of the cost over all possible paths connecting them:

$$d_g(u, v) = \min\{c(P) \mid P(1) = u, P(k) = v, |P| = k\}$$

The *diameter* of the graph G is the largest geodesic distance

$$\text{diam}(G) = \max_{u, v} d_g(u, v).$$

We demonstrate that for $\lambda_2(\mathcal{L}_G)$, the second smallest eigenvalue of the normalized graph Laplacian $\mathcal{L}_G = D^{-1/2}L_G D^{-1/2}$,

$$\text{diam}(G) \geq \frac{1}{\text{vol}(G)\lambda_2(\mathcal{L}_G)}. \quad (4.10)$$

That is, a small second eigenvalue implies that the diameter of the graph must be large. We already know that $\lambda_2(\mathcal{L}_G) = 0$ implies that the graph is disconnected and the diameter is infinite. The bound (4.10) is more quantitative from that standpoint.

Proof: Taking f to be an eigenfunction achieving the Rayleigh quotient characterization

$$\lambda_2 = \inf_{f^T D \mathbf{1} = 0} \frac{\frac{1}{2} \sum_{u, v} w_{uv} (f(v) - f(u))^2}{\sum_v f(v)^2 d_v}.$$

Choose v_0, u_0 such that

$$f(v_0) = \max_v |f(v)|,$$

and

$$f(u_0) < 0,$$

which exist since

$$\sum_v f(v) d_v = 0.$$

Now, suppose

$$P = (u_0 = v_1, v_2, \dots, v_k = v_0)$$

is the shortest path from u_0 to v_0 . We have

$$\begin{aligned}
\lambda_2 &= \frac{\frac{1}{2} \sum_{u,v} w_{uv} (f(v) - f(u))^2}{\sum_v f(v)^2 d_v} \\
&\geq \frac{\sum_{i=1}^{k-1} w_{v_i, v_{i+1}} (f(v_i) - f(v_{i+1}))^2}{f(v_0)^2 \sum_v d_v} \\
&\geq \frac{\sum_{i=1}^{k-1} w_{v_i, v_{i+1}} (f(v_i) - f(v_{i+1}))^2}{f(v_0)^2 \text{vol}(G)} \frac{d_g(u_0, v_0)}{\text{diam}(G)} \\
&= \frac{\sum_{i=1}^{k-1} w_{v_i, v_{i+1}} (f(v_i) - f(v_{i+1}))^2}{f(v_0)^2 \text{vol}(G)} \frac{\sum_{i=1}^{k-1} w_{v_i, v_{i+1}}^{-1}}{\text{diam}(G)} \\
&\geq \frac{\left(\sum_{i=1}^{k-1} f(v_i) - f(v_{i+1}) \right)^2}{f(v_0)^2 \text{vol}(G) \text{diam}(G)} \\
&= \frac{(f(v_0) - f(u_0))^2}{f(v_0)^2 \text{vol}(G) \text{diam}(G)} \\
&\geq \frac{1}{\text{vol}(G) \text{diam}(G)},
\end{aligned}$$

there the second-to-last inequality is an application of Cauchy-Schwarz inequality in dimension $k - 1$ (and relies on the fact that the weights $w_{v_i, v_{i+1}}$ are positive).

The upper bound in Cheeger's inequality (corresponding to Lemma 4.8) is more interesting but more difficult to prove, and it is often referred to as the “the difficult part” of Cheeger's inequality. We will prove this Lemma in what follows. There are several proofs of this inequality (see [56] for four different proofs!). The proof that follows is an adaptation of the proof in this blog post [172] for the case of weighted graphs.

Proof. [of Lemma 4.8]

We will show that given $y \in \mathbb{R}^n$ satisfying

$$\mathcal{R}(y) := \frac{y^T L_G y}{y^T D y} \leq \delta,$$

and $y^T D \mathbf{1} = 0$, there is a “rounding of it”, meaning a threshold τ and a corresponding choice of partition

$$S = \{i \in V : y_i \leq \tau\}$$

such that

$$h(S) \leq \sqrt{2\delta}.$$

Since $y = \varphi_2$ satisfies the conditions with $\delta = \lambda_2(\mathcal{L}_G)$, this proves the Lemma.

We will pick this threshold at random and use the probabilistic method to show that at least one of the thresholds works.

First we can, without loss of generality, assume that $y_1 \leq \dots \leq y_n$ (we can simply relabel the vertices). Also, note that scaling of y does not change the

value of $\mathcal{R}(y)$. Also, if $y^T D \mathbf{1} = 0$ adding a multiple of $\mathbf{1}$ to y can only decrease the value of $\mathcal{R}(y)$: the numerator does not change and the denominator $(y + c\mathbf{1})^T D(y + c\mathbf{1}) = y^T D y + c^2 \mathbf{1}^T D \mathbf{1} \geq y^T D y$.

This means that we can construct (from y by adding a multiple of $\mathbf{1}$ and scaling) a vector x such that

$$x_1 \leq \dots \leq x_n, \quad x_m = 0, \quad \text{and} \quad x_1^2 + x_n^2 = 1,$$

and

$$\frac{x^T L_G x}{x^T D x} \leq \delta,$$

where m be the index for which $\text{vol}(\{1, \dots, m-1\}) \leq \text{vol}(\{m, \dots, n\})$ but $\text{vol}(\{1, \dots, m\}) > \text{vol}(\{m+1, \dots, n\})$.

We consider a random construction of S with the following distribution. $S = \{i \in V : x_i \leq \tau\}$ where $\tau \in [x_1, x_n]$ is drawn at random with the distribution

$$\mathbb{P}\{\tau \in [a, b]\} = \int_a^b 2|\tau| d\tau,$$

where $x_1 \leq a \leq b \leq x_n$.

It is not difficult to check that

$$\mathbb{P}\{\tau \in [a, b]\} = \begin{cases} |b^2 - a^2| & \text{if } a \text{ and } b \text{ have the same sign} \\ a^2 + b^2 & \text{if } a \text{ and } b \text{ have different signs} \end{cases}$$

Let us start by estimating $\mathbb{E} \text{cut}(S)$.

$$\begin{aligned} \mathbb{E} \text{cut}(S) &= \mathbb{E} \frac{1}{2} \sum_{i \in V} \sum_{j \in V} w_{ij} \mathbf{1}_{(S, S^c) \text{ cuts the edge } (i, j)} \\ &= \frac{1}{2} \sum_{i \in V} \sum_{j \in V} w_{ij} \mathbb{P}\{(S, S^c) \text{ cuts the edge } (i, j)\} \end{aligned}$$

Note that $\mathbb{P}\{(S, S^c) \text{ cuts the edge } (i, j)\}$ is $|x_i^2 - x_j^2|$ if x_i and x_j have the same sign and $x_i^2 + x_j^2$ otherwise. Both cases can be conveniently upper bounded by $|x_i - x_j|(|x_i| + |x_j|)$. This means that

$$\begin{aligned} \mathbb{E} \text{cut}(S) &\leq \frac{1}{2} \sum_{i, j} w_{ij} |x_i - x_j| (|x_i| + |x_j|) \\ &\leq \frac{1}{2} \sqrt{\sum_{ij} w_{ij} (x_i - x_j)^2} \sqrt{\sum_{ij} w_{ij} (|x_i| + |x_j|)^2}, \end{aligned}$$

where the second inequality follows from the Cauchy-Schwarz inequality.

From the construction of x we know that

$$\sum_{ij} w_{ij} (x_i - x_j)^2 = 2x^T L_G x \leq 2\delta x^T D x.$$

Also,

$$\sum_{ij} w_{ij}(|x_i| + |x_j|)^2 \leq \sum_{ij} w_{ij}(2x_i^2 + 2x_j^2) = 2 \left(\sum_i \deg(i) x_i^2 \right) + 2 \left(\sum_j \deg(j) x_j^2 \right) = 4x^T D x.$$

This means that

$$\mathbb{E} \text{cut}(S) \leq \frac{1}{2} \sqrt{2\delta x^T D x} \sqrt{4x^T D x} = \sqrt{2\delta} x^T D x.$$

On the other hand,

$$\mathbb{E} \min\{\text{vol } S, \text{vol } S^c\} = \sum_{i=1}^n \deg(i) \mathbb{P}\{x_i \text{ is in the smallest set (in terms of volume)}\},$$

to break ties, if $\text{vol}(S) = \text{vol}(S^c)$ we take the “smallest” set to be the one with the first indices.

Note that m is always in the largest set. Any vertex $j < m$ is in the smallest set if $x_j \leq \tau \leq x_m = 0$ and any $j > m$ is in the smallest set if $0 = x_m \leq \tau \leq x_j$. This means that,

$$\mathbb{P}\{x_j \text{ is in the smallest set (in terms of volume)}\} = x_j^2.$$

Which means that

$$\mathbb{E} \min\{\text{vol } S, \text{vol } S^c\} = \sum_{i=1}^n \deg(i) x_i^2 = x^T D x.$$

Hence,

$$\frac{\mathbb{E} \text{cut}(S)}{\mathbb{E} \min\{\text{vol } S, \text{vol } S^c\}} \leq \sqrt{2\delta}.$$

Note, however, that $\frac{\mathbb{E} \text{cut}(S)}{\mathbb{E} \min\{\text{vol } S, \text{vol } S^c\}}$ is not necessarily the same as $\mathbb{E} \frac{\text{cut}(S)}{\min\{\text{vol } S, \text{vol } S^c\}}$. Therefore, we do not necessarily have

$$\mathbb{E} \frac{\text{cut}(S)}{\min\{\text{vol } S, \text{vol } S^c\}} \leq \sqrt{2\delta}.$$

However, since both random variables are positive,

$$\mathbb{E} \text{cut}(S) \leq \mathbb{E} \min\{\text{vol } S, \text{vol } S^c\} \sqrt{2\delta},$$

or equivalently

$$\mathbb{E} \left[\text{cut}(S) - \min\{\text{vol } S, \text{vol } S^c\} \sqrt{2\delta} \right] \leq 0,$$

which guarantees, by the probabilistic method, the existence of S such that

$$\text{cut}(S) \leq \min\{\text{vol } S, \text{vol } S^c\} \sqrt{2\delta},$$

which is equivalent to

$$h(S) = \frac{\text{cut}(S)}{\min\{\text{vol } S, \text{vol } S^c\}} \leq \sqrt{2\delta},$$

which concludes the proof of the Lemma. $\square$

Remark 4.11. Spectral clustering can also be viewed from the perspective of random walks. In fact, Proposition 5.7 shows that spectral clustering can be viewed as attempting to find clusters so as to minimize the probability of the random walker to jump from one cluster to the other.

Multiple Clusters

Much of the above can be easily adapted to multiple clusters. Algorithm 4.2 is a natural extension of spectral clustering to multiple clusters.⁵

Algorithm 4.2 Spectral Clustering

Given a graph $G = (V, E, W)$, let $v_2, \dots, v_k$ be the eigenvectors corresponding to the second through k 'th eigenvalues of the normalized Laplacian $\mathcal{L}_G$, as defined in (4.7). Let $\varphi_m = D^{-\frac{1}{2}}v_m$. Consider the map $\varphi : V \rightarrow \mathbb{R}^{k-1}$ defined as

$$\varphi(v_i) = \begin{bmatrix} \varphi_2(i) \\ \vdots \\ \varphi_k(i) \end{bmatrix}.$$

Cluster the n points in $k - 1$ dimensions into k clusters using k -means.

There is also an analogue of Cheeger's inequality. A natural way of evaluating k -way clustering is via the k -way expansion constant (see [114]):

$$\rho_G(k) = \min_{S_1, \dots, S_k} \max_{l=1, \dots, k} \left\{ \frac{\text{cut}(S_l)}{\text{vol}(S_l)} \right\},$$

where the maximum is over all choice of k disjoint subsets of V (but not necessarily forming a partition).

Another natural definition is

$$\varphi_G(k) = \min_{S: \text{vol } S \leq \frac{1}{k} \text{vol}(G)} \frac{\text{cut}(S)}{\text{vol}(S)}.$$

⁵We will see in Chapter 5 that the map $\varphi : V \rightarrow \mathbb{R}^{k-1}$ defined in Algorithm 4.2 can also be used for data visualization, not just clustering.

It is easy to see that

$$\varphi_G(k) \leq \rho_G(k).$$

The following are analogues of Cheeger's inequality for multiple clusters.

Theorem 4.12 ([114]). *Let $G = (V, E, W)$ be a graph and k a positive integer*

$$\rho_G(k) \leq \mathcal{O}(k^2) \sqrt{\lambda_k}. \quad (4.11)$$

Also,

$$\rho_G(k) \leq \mathcal{O}\left(\sqrt{\lambda_{2k} \log k}\right).$$

Nonlinear Dimension Reduction and Diffusion Maps

In Chapter 3 we discussed dimension reduction via Principal Component Analysis. Many datasets however have low dimensional structure that is not linear. In this chapter we will discuss nonlinear dimension reduction techniques. Just as with Spectral Clustering in Chapter 4 we will focus on graph data while noting that most types of data can be transformed into a weighted graph by means of a similarity kernel (Section 5.1.1). The goal of this chapter is to embed the nodes of a graph in Euclidean space in a way that best preserves the intrinsic geometry of the graph (or the data that gave rise to the graph).

5.1 Diffusion Maps

Diffusion Maps will allow us to represent (weighted) graphs $G = (V, E, W)$ in $\mathbb{R}^d$, i.e. associating, to each node, a point in $\mathbb{R}^d$. Before presenting Diffusion Maps, we'll introduce a few important notions. The reader may notice the similarities with the objects described in the context of PageRank in Chapter 4, the main difference is that here the connections between graphs have no direction, meaning that the weight matrix W is symmetric; this will be crucial in the derivations below.

Given $G = (V, E, W)$ we consider a random walk (with independent steps) on the vertices of V with transition probabilities:

$$\mathbb{P}\{X(t+1) = j | X(t) = i\} = \frac{w_{ij}}{\deg(i)},$$

where $\deg(i) = \sum_j w_{ij}$. Let M be the matrix of these probabilities,

$$M_{ij} = \frac{w_{ij}}{\deg(i)}. \quad (5.1)$$

It is easy to see that $M_{ij} \geq 0$ and $M\mathbf{1} = \mathbf{1}$ (indeed, M is a transition probability matrix). Recalling that D is the diagonal matrix with diagonal entries

$D_{ii} = \deg(i)$ we have

$$M = D^{-1}W.$$

If we start a random walker at node i ($X(0) = i$) then the probability that, at step t , is at node j is given by

$$\mathbb{P}\{X(t) = j | X(0) = i\} = (M^t)_{ij}.$$

In other words, the probability cloud of the random walker at point t , given that it started at node i is given by the row vector

$$\mathbb{P}\{X(t) | X(0) = i\} = e_i^T M^t = M^t[i, :].$$

Remark 5.1. A natural representation of the graph would be to associate each vertex to the probability cloud above, meaning

$$i \rightarrow M^t[i, :].$$

This would place nodes i_1 and i_2 for which the random walkers starting at i_1 and i_2 have, after t steps, very similar distribution of locations. However, this would require $d = n$. In what follows we will construct a similar mapping but for considerably smaller d .

M is not symmetric, but a matrix similar to M , $S = D^{\frac{1}{2}} M D^{-\frac{1}{2}}$ is, indeed $S = D^{-\frac{1}{2}} W D^{-\frac{1}{2}}$. We consider the spectral decomposition of S

$$S = V \Lambda V^T, \quad (5.2)$$

where $V = [v_1, \dots, v_n]$ satisfies $V^T V = I_{n \times n}$ and Λ is diagonal with diagonal elements $\Lambda_{kk} = \lambda_k$ (and we organize them as $\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_n$). Note that $S v_k = \lambda_k v_k$. Also,

$$M = D^{-\frac{1}{2}} S D^{\frac{1}{2}} = D^{-\frac{1}{2}} V \Lambda V^T D^{\frac{1}{2}} = \left(D^{-\frac{1}{2}} V\right) \Lambda \left(D^{\frac{1}{2}} V\right)^T.$$

We define $\Phi = D^{-\frac{1}{2}} V$ with columns $\Phi = [\varphi_1, \dots, \varphi_n]$ and $\Psi = D^{\frac{1}{2}} V$ with columns $\Psi = [\psi_1, \dots, \psi_n]$. Then

$$M = \Phi \Lambda \Psi^T,$$

and Φ, Ψ form a biorthogonal system in the sense that $\Phi^T \Psi = I_{n \times n}$ or, equivalently, $\varphi_j^T \psi_k = \delta_{jk}$. Note that φ_k and ψ_k are, respectively right and left eigenvectors of M , indeed, for all $1 \leq k \leq n$:

$$M \varphi_k = \lambda_k \varphi_k \quad \text{and} \quad \psi_k^T M = \lambda_k \psi_k^T.$$

Also, we can rewrite this decomposition as

$$M = \sum_{k=1}^n \lambda_k \varphi_k \psi_k^T.$$

and it is easy to see that

$$M^t = \sum_{k=1}^n \lambda_k^t \varphi_k \psi_k^T. \quad (5.3)$$

Let's revisit the embedding suggested on Remark 5.1. It would correspond to

$$v_i \rightarrow M^t[i, :] = \sum_{k=1}^n \lambda_k^t \varphi_k(i) \psi_k^T,$$

it is written in terms of the basis ψ_k . The Diffusion Map will essentially consist of the representing a node i by the coefficients of the above map

$$v_i \rightarrow \begin{bmatrix} \lambda_1^t \varphi_1(i) \\ \lambda_2^t \varphi_2(i) \\ \vdots \\ \lambda_n^t \varphi_n(i) \end{bmatrix}, \quad (5.4)$$

Note that $M\mathbf{1} = \mathbf{1}$, meaning that one of the right eigenvectors φ_k is simply a multiple of $\mathbf{1}$ and so it does not distinguish the different nodes of the graph. We will show that this indeed corresponds to the the first eigenvalue.

Proposition 5.2. *All eigenvalues λ_k of M satisfy $|\lambda_k| \leq 1$.*

Proof.

Let φ_k be a right eigenvector associated with λ_k whose largest entry in magnitude is positive $\varphi_k(i_{\max})$. Then,

$$\lambda_k \varphi_k(i_{\max}) = M \varphi_k(i_{\max}) = \sum_{j=1}^n M_{i_{\max}, j} \varphi_k(j).$$

This means, by triangular inequality that, that

$$|\lambda_k| = \sum_{j=1}^n |M_{i_{\max}, j}| \frac{|\varphi_k(j)|}{|\varphi_k(i_{\max})|} \leq \sum_{j=1}^n |M_{i_{\max}, j}| = 1.$$

□

Remark 5.3. It is possible that there are other eigenvalues with magnitude 1 but only if G is disconnected or if G is bipartite. Provided that G is a connected graph, a natural way to remove potential periodicity issues (like the graph being bipartite) is to make the walk lazy, i.e. to add a certain probability of the walker to stay in the current node. This can be conveniently achieved by taking, e.g.,

$$M' = \frac{1}{2}M + \frac{1}{2}I.$$

By the proposition above we can take $\varphi_1 = \mathbf{1}$, meaning that the first coordinate of (5.4) does not help differentiate points on the graph. This suggests removing that coordinate:

Definition 5.4 (Diffusion Map). *Given a graph $G = (V, E, W)$ construct M and its decomposition $M = \Phi\Lambda\Psi^T$ as described above. The Diffusion Map is a map $\varphi_t : V \rightarrow \mathbb{R}^{n-1}$ given by*

$$\varphi_t(v_i) = \begin{bmatrix} \lambda_2^t \varphi_2(i) \\ \lambda_3^t \varphi_3(i) \\ \vdots \\ \lambda_n^t \varphi_n(i) \end{bmatrix}.$$

This map is still a map to $n - 1$ dimensions. But note now that each coordinate has a factor of λ_k^t which, if λ_k is small will be rather small for moderate values of t . This motivates truncating the Diffusion Map by taking only the first d coefficients.

Definition 5.5 (Truncated Diffusion Map). *Given a graph $G = (V, E, W)$ and dimension d , construct M and its decomposition $M = \Phi\Lambda\Psi^T$ as described above. The Diffusion Map truncated to d dimensions is a map $\varphi_t : V \rightarrow \mathbb{R}^d$ given by*

$$\varphi_t^{(d)}(v_i) = \begin{bmatrix} \lambda_2^t \varphi_2(i) \\ \lambda_3^t \varphi_3(i) \\ \vdots \\ \lambda_{d+1}^t \varphi_{d+1}(i) \end{bmatrix}.$$

In the following theorem we show that the euclidean distance in the diffusion map coordinates (called diffusion distance) meaningfully measures distance between the probability clouds after t iterations.

Theorem 5.6. *For any pair of nodes v_{i_1}, v_{i_2} we have*

$$\|\varphi_t(v_{i_1}) - \varphi_t(v_{i_2})\|^2 = \sum_{j=1}^n \frac{1}{\deg(j)} [\mathbb{P}\{X(t) = j | X(0) = i_1\} - \mathbb{P}\{X(t) = j | X(0) = i_2\}]^2.$$

Proof. Note that $\sum_{j=1}^n \frac{1}{\deg(j)} [\mathbb{P}\{X(t) = j | X(0) = i_1\} - \mathbb{P}\{X(t) = j | X(0) = i_2\}]^2$ can be rewritten as

$$\begin{aligned} \sum_{j=1}^n \frac{1}{\deg(j)} \left[\sum_{k=1}^n \lambda_k^t \varphi_k(i_1) \psi_k(j) - \sum_{k=1}^n \lambda_k^t \varphi_k(i_2) \psi_k(j) \right]^2 &= \\ &= \sum_{j=1}^n \frac{1}{\deg(j)} \left[\sum_{k=1}^n \lambda_k^t (\varphi_k(i_1) - \varphi_k(i_2)) \psi_k(j) \right]^2 \end{aligned}$$

and

$$\begin{aligned} \sum_{j=1}^n \frac{1}{\deg(j)} \left[\sum_{k=1}^n \lambda_k^t (\varphi_k(i_1) - \varphi_k(i_2)) \psi_k(j) \right]^2 &= \sum_{j=1}^n \left[\sum_{k=1}^n \lambda_k^t (\varphi_k(i_1) - \varphi_k(i_2)) \frac{\psi_k(j)}{\sqrt{\deg(j)}} \right]^2 \\ &= \left\| \sum_{k=1}^n \lambda_k^t (\varphi_k(i_1) - \varphi_k(i_2)) D^{-\frac{1}{2}} \psi_k \right\|^2. \end{aligned}$$

Note that $D^{-\frac{1}{2}} \psi_k = v_k$ which forms an orthonormal basis, meaning that

$$\begin{aligned} \left\| \sum_{k=1}^n \lambda_k^t (\varphi_k(i_1) - \varphi_k(i_2)) D^{-\frac{1}{2}} \psi_k \right\|^2 &= \sum_{k=1}^n (\lambda_k^t (\varphi_k(i_1) - \varphi_k(i_2)))^2 \\ &= \sum_{k=2}^n (\lambda_k^t \varphi_k(i_1) - \lambda_k^t \varphi_k(i_2))^2, \end{aligned}$$

where the last equality follows from the fact that $\varphi_1 = \mathbf{1}$. The proof of the theorem is complete. $\square$

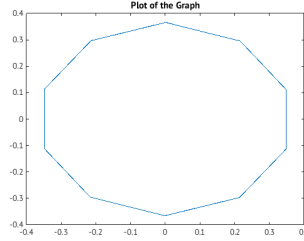


Fig. 5.1: The Diffusion Map of the ring graph gives a very natural way of displaying (indeed, if one is asked to draw the ring graph, this is probably the drawing that most people would do). It is actually not difficult to analytically compute the Diffusion Map of this graph and confirm that it displays the points in a circle.

5.1.1 Diffusion Maps of point clouds

Very often we are interested in embedding in $\mathbb{R}^d$ a point cloud of points $x_1, \dots, x_n \in \mathbb{R}^p$ and not necessarily a graph. One option is to use Principal Component Analysis (PCA), but PCA is only designed to find linear structure of the data and the low dimensionality of the dataset may be non-linear. For

example, let us say that our dataset is images of the face of someone taken from different angles and lighting conditions, for example, the dimensionality of this dataset is limited by the amount of muscles in the head and neck and by the degrees of freedom of the lighting conditions (see Figure 5.5) but it is not clear that this low dimensional structure is linearly apparent on the pixel values of the images.

Let us consider a point cloud that is sampled from a two dimensional swiss roll embedded in three dimension (see Figure 5.2). In order to learn the two dimensional structure of this object we need to differentiate points that are near each other because they are close by in the manifold and not simply because the manifold is curved and the points appear nearby even when they really are distant in the manifold (see Figure 5.2 for an example). We will achieve this by creating a graph from the data points.

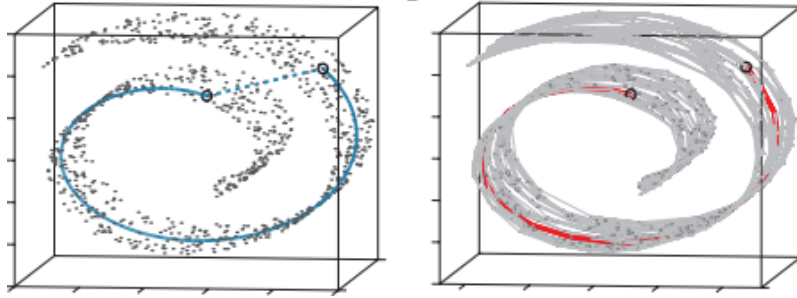


Fig. 5.2: A swiss roll point cloud (see, for example, [170]). The points are sampled from a two dimensional manifold curved in $\mathbb{R}^3$ and then a graph is constructed where nodes correspond to points.

Our goal is for the graph to capture the structure of the manifold. To each data point we will associate a node. For this we should only connect points that are close in the manifold and not points that maybe appear close in Euclidean space simply because of the curvature of the manifold. This is achieved by picking a small scale and linking nodes if they correspond to points whose distance is smaller than that scale. This is usually done smoothly via a kernel K_ε , and to each edge (i, j) associating a weight

$$w_{ij} = K_\varepsilon(\|x_i - x_j\|_2),$$

a common example of a Kernel is $K_\varepsilon(u) = \exp(-\frac{1}{2\varepsilon}u^2)$, that gives essentially zero weight to edges corresponding to pairs of nodes for which $\|x_i - x_j\|_2 \gg \sqrt{\varepsilon}$. We can then take the the Diffusion Maps of the resulting graph.

5.1.2 An illustrative simple example

A simple and illustrative example is to take images of a blob on a background in different positions (imagine a white square on a black background and each data point corresponds to the same white square in different positions). This dataset is clearly intrinsically two dimensional, as each image can be described by the (two-dimensional) position of the square. However, we don't expect this two-dimensional structure to be directly apparent from the vectors of pixel values of each image; in particular we don't expect these vectors to lie in a two dimensional affine subspace!

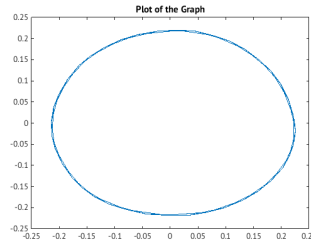


Fig. 5.3: The two-dimensional diffusion map of the dataset of the dataset where each data point is an image with the same vertical strip in different positions in the x-axis, the circular structure is apparent.

Let's start by experimenting with the above example for one dimension. In that case the blob is a vertical stripe and simply moves left and right. We think of our space as the one in many arcade games, if the square or stripe moves to the right all the way to the end of the screen, it shows up on the left side (and same for up-down in the two-dimensional case). Not only should this point cloud have a one dimensional structure but it should also exhibit a circular structure. Remarkably, this structure is completely apparent when taking the two-dimensional Diffusion Map of this dataset, see Figure 5.3.

For the two dimensional example, we expect the structure of the underlying manifold to be a two-dimensional torus. Indeed, Figure 5.4 shows that the three-dimensional diffusion map captures the toroidal structure of the data.

5.1.3 Similar non-linear dimensional reduction techniques

There are several other similar non-linear dimensional reduction methods. A particularly popular one is ISOMAP [170]. The idea is to find an embedding in $\mathbb{R}_d$ for which euclidean distances in the embedding correspond as much as possible to geodesic distances in the graph. This can be achieved by, between

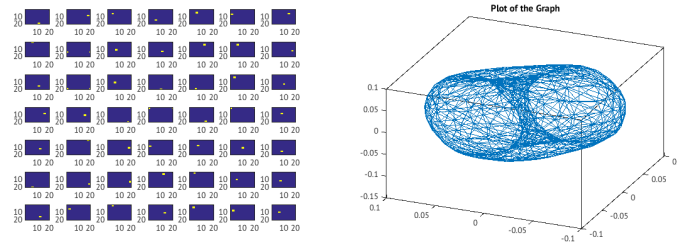


Fig. 5.4: On the left the data set considered and on the right its three dimensional diffusion map, the fact that the manifold is a torus is remarkably captured by the embedding.

pairs of nodes v_i, v_j finding their geodesic distance and then using, for example, Multidimensional Scaling to find points $y_i \in \mathbb{R}^d$ that minimize (for example)

$$\min_{y_1, \dots, y_n \in \mathbb{R}^d} \sum_{i,j} (\|y_i - y_j\|^2 - \delta_{ij}^2)^2,$$

which can be done with spectral methods (it is a good exercise to compute the optimal solution to the above optimization problem).

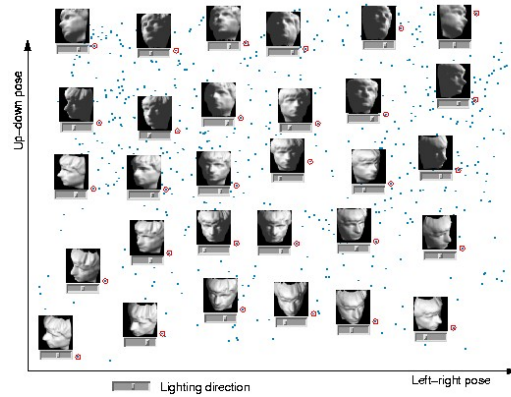


Fig. 5.5: The two dimensional representation of a data set of images of faces as obtained in [170] using ISOMAP. Remarkably, the two dimensionals are interpretable

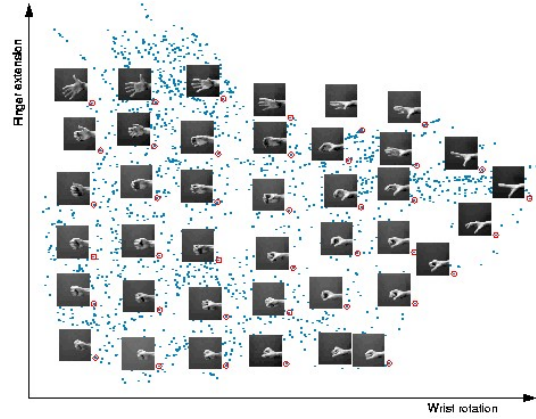


Fig. 5.6: The two dimensional representation of a data set of images of human hand as obtained in [170] using ISOMAP. Remarkably, the two dimensionals are interpretable

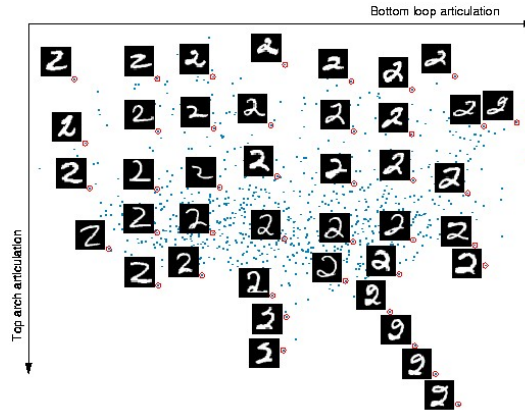


Fig. 5.7: The two dimensional representation of a data set of handwritten digits as obtained in [170] using ISOMAP. Remarkably, the two dimensionals are interpretable

5.2 Connections between Diffusion Maps and Spectral Clustering

Diffusion maps are tightly connected to Spectral Clustering (described in Chapter 4). In fact, Spectral Clustering can be understood as simply performing k -means on the embedding given by Diffusion Maps truncated to $k - 1$ dimensions.

A natural way to try to overcome the issues of k -means depicted in Figure 4.4 is by using Diffusion Maps: Given the data points we construct a weighted graph $G = (V, E, W)$ using a kernel K_ε , such as $K_\varepsilon(u) = \exp(-\frac{1}{2\varepsilon}u^2)$, by associating each point to a vertex and, for which pair of nodes, set the edge weight as

$$w_{ij} = K_\varepsilon(\|x_i - x_j\|).$$

Recall the construction of a matrix $M = D^{-1}W$ as the transition matrix of a random walk

$$\mathbb{P}\{X(t+1) = j | X(t) = i\} = \frac{w_{ij}}{\deg(i)} = M_{ij},$$

where D is the diagonal with $D_{ii} = \deg(i)$. The d -dimensional Diffusion Map is given by

$$\varphi_t^{(d)}(i) = \begin{bmatrix} \lambda_2^t \varphi_2(i) \\ \vdots \\ \lambda_{d+1}^t \varphi_{d+1}(i) \end{bmatrix},$$

where $M = \Phi \Lambda \Psi^T$ where Λ is the diagonal matrix with the eigenvalues of M and Φ and Ψ are, respectively, the right and left eigenvectors of M (note that they form a bi-orthogonal system, $\Phi^T \Psi = I$).

If we want to cluster the vertices of the graph in k clusters, then it is natural to truncate the Diffusion Map to have $k-1$ dimensions (since in $k-1$ dimensions we can have k linearly separable sets). If indeed the clusters were linearly separable after embedding then one could attempt to use k -means on the embedding to find the clusters, this is precisely the motivation for Spectral Clustering.

Algorithm 5.1 Spectral Clustering described using Diffusion Maps.

Spectral Clustering: Given a graph $G = (V, E, W)$ and a number of clusters k (and t), Spectral Clustering consists in taking a $(k-1)$ dimensional Diffusion Map

$$\varphi_t^{(k-1)}(i) = \begin{bmatrix} \lambda_2^t \varphi_2(i) \\ \vdots \\ \lambda_k^t \varphi_k(i) \end{bmatrix}$$

and clustering the points $\varphi_t^{(k-1)}(1), \varphi_t^{(k-1)}(2), \dots, \varphi_t^{(k-1)}(n) \in \mathbb{R}^{k-1}$ using, for example, k -means clustering. Usually, the scaling of λ_m^t is ignored (corresponding to $t = 0$).

In order to show that this indeed coincides with Algorithm 4.2, it is enough to show that $\varphi_m = D^{-\frac{1}{2}} v_m$ where v_m is the eigenvector associated with the m -th smallest eigenvalue of $\mathcal{L}_G$. This follows from the fact that $S = D^{-\frac{1}{2}} W D^{-\frac{1}{2}}$ as defined in (5.2) is related to $\mathcal{L}_G$ by

$$\mathcal{L}_G = I - S,$$

and $\Phi = D^{-1/2}V$.

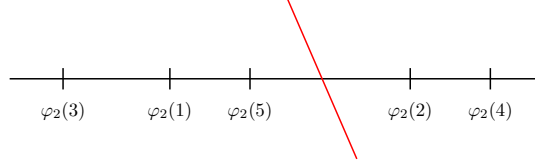


Fig. 5.8: For two clusters, spectral clustering consists in assigning to each vertex i a real number $\varphi_2(i)$, then setting a threshold τ and taking $S = \{i \in V : \varphi_2(i) \leq \tau\}$. This real number can both be interpreted through the spectrum of $\mathcal{L}_G$ as in Algorithm 4.1 or as the Diffusion Map embedding as in Algorithm 5.1.

Proposition 5.7 below establishes a connection between Ncut (as described in Chapter 4) and the random walks introduced above. Let M as defined in (5.1) denote the matrix of transition probabilities. Recall that $M\mathbf{1} = \mathbf{1}$, corresponding to $M\varphi_1 = \varphi_1$, which means that $\psi_1^T M = \psi_1^T$, where

$$\psi_1 = D^{\frac{1}{2}}v_1 = D\varphi_1 = [\deg(i)]_{1 \leq i \leq n}.$$

This means that $\left[\frac{\deg(i)}{\text{vol}(G)}\right]_{1 \leq i \leq n}$ is the stationary distribution of this random walk. Indeed it is easy to check that, if $X(t)$ has a certain distribution p_t then $X(t+1)$ has a distribution p_{t+1} given by $p_{t+1}^T = p_t^T M$.

Proposition 5.7. *Given a graph $G = (V, E, W)$ and a partition (S, S^c) of V , $\text{Ncut}(S)$ corresponds to the probability, in the random walk associated with G , that a random walker in the stationary distribution goes to S^c conditioned on being in S plus the probability of going to S condition on being in S^c , more explicitly:*

$$\text{Ncut}(S) = \mathbb{P}\{X(t+1) \in S^c | X(t) \in S\} + \mathbb{P}\{X(t+1) \in S | X(t) \in S^c\},$$

where $\mathbb{P}\{X(t) = i\} = \frac{\deg(i)}{\text{vol}(G)}$.

Proof. Without loss of generality we can take $t = 0$. Also, the second term in the sum corresponds to the first with S replaced by S^c and vice-versa, so we'll focus on the first one. We have:

$$\begin{aligned}
\mathbb{P}\{X(1) \in S^c | X(0) \in S\} &= \frac{\mathbb{P}\{X(1) \in S^c \cap X(0) \in S\}}{\mathbb{P}\{X(0) \in S\}} \\
&= \frac{\sum_{i \in S} \sum_{j \in S^c} \mathbb{P}\{X(1) \in j \cap X(0) \in i\}}{\sum_{i \in S} \mathbb{P}\{X(0) = i\}} \\
&= \frac{\sum_{i \in S} \sum_{j \in S^c} \frac{\deg(i)}{\text{vol}(G)} \frac{w_{ij}}{\deg(i)}}{\sum_{i \in S} \frac{\deg(i)}{\text{vol}(G)}} \\
&= \frac{\sum_{i \in S} \sum_{j \in S^c} w_{ij}}{\sum_{i \in S} \deg(i)} \\
&= \frac{\text{cut}(S)}{\text{vol}(S)}.
\end{aligned}$$

Analogously,

$$\mathbb{P}\{X(t+1) \in S | X(t) \in S^c\} = \frac{\text{cut}(S)}{\text{vol}(S^c)},$$

which concludes the proof. $\square$

5.3 Semi-supervised learning

Classification is a central task in machine learning. In a supervised learning setting we are given many labelled examples and want to use them to infer the label of a new, unlabeled example. For simplicity, let us focus on the case of two labels, $\{-1, +1\}$.

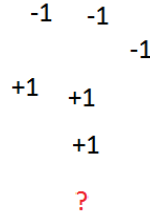


Fig. 5.9: Given a few labeled points, the task is to label an unlabeled point.

Let us consider the task of labelling the point “?” in Figure 5.9 given the labeled points. The natural label to give to the unlabeled point would be 1.

However, if we are given not just one unlabeled point, but many, as in Figure 5.10; then it starts being apparent that -1 is a more reasonable guess.

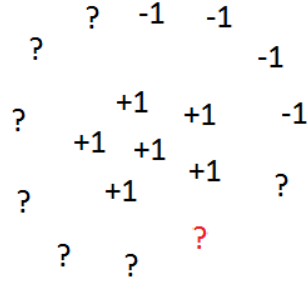


Fig. 5.10: In this example we are given many unlabeled points, the unlabeled points help us learn the geometry of the data.

Intuitively, the unlabeled data points allowed us to better learn the intrinsic geometry of the dataset. That is the idea behind Semi-Supervised Learning (SSL), to make use of the fact that often one has access to many unlabeled data points in order to improve classification.

Just as above, we will use the data points to construct (via a kernel K_ε) a graph $G = (V, E, W)$ where nodes correspond to points. More precisely, let l denote the number of labeled points with labels $f_1, \dots, f_l$, and u the number of unlabeled points (with $n = l + u$), the first l nodes $v_1, \dots, v_l$ correspond to labeled points and the rest $v_{l+1}, \dots, v_n$ are unlabeled. We want to find a function $f : V \rightarrow \{-1, 1\}$ that agrees on labeled points: $f(i) = f_i$ for $i = 1, \dots, l$ and that is “as smooth as possible” on the graph. A way to pose this is the following

$$\min_{f: V \rightarrow \{-1, 1\}: f(i) = f_i \text{ } i=1, \dots, l} \sum_{i < j} w_{ij} (f(i) - f(j))^2.$$

Instead of restricting ourselves to giving $\{-1, 1\}$ we allow ourselves to give real valued labels, with the intuition that we can “round” later by, e.g., assigning the sign of $f(i)$ to node i .

We thus are interested in solving

$$\min_{f: V \rightarrow \mathbb{R}: f(i) = f_i \text{ } i=1, \dots, l} \sum_{i < j} w_{ij} (f(i) - f(j))^2.$$

If we denote by f the vector (in $\mathbb{R}^n$ with the function values) then, recalling Proposition 4.3, we can rewrite the problem as

$$\sum_{i < j} w_{ij} (f(i) - f(j))^2 = f^T L_G f.$$

Remark 5.8. Consider an analogous example on the real line, where one would want to minimize

$$\int f'(x)^2 dx.$$

Integrating by parts

$$\int f'(x)^2 dx = \text{Boundary Terms} - \int f(x)f''(x)dx.$$

Analogously, in $\mathbb{R}^d$:

$$\begin{aligned} \int \|\nabla f(x)\|^2 dx &= \int \sum_{k=1}^d \left(\frac{\partial f}{\partial x_k}(x) \right)^2 dx = \text{B. T.} - \int f(x) \sum_{k=1}^d \frac{\partial^2 f}{\partial x_k^2}(x) dx = \\ &= \text{B. T.} - \int f(x) \Delta f(x) dx, \end{aligned}$$

which helps motivate the use of the term graph Laplacian for L_G .

Let us consider our problem

$$\min_{f: V \rightarrow \mathbb{R}: f(i)=f_i \ i=1,\dots,l} f^T L_G f.$$

We can write

$$D = \begin{bmatrix} D_L & 0 \\ 0 & D_U \end{bmatrix}, \quad W = \begin{bmatrix} W_{LL} & W_{LU} \\ W_{UL} & W_{UU} \end{bmatrix}, \quad L_G = \begin{bmatrix} D_L - W_{LL} & -W_{LU} \\ -W_{UL} & D_U - W_{UU} \end{bmatrix},$$

and $f = \begin{bmatrix} f_L \\ f_U \end{bmatrix}$, there L and U correspond to the set of indices of respectively labeled and unlabeled nodes (not to be confused with the Laplacian matrix L_G).

Then we want to find (recall that $W_{UL} = W_{LU}^T$)

$$\min_{f_U \in \mathbb{R}^u} f_L^T [D_L - W_{LL}] f_L - 2f_U^T W_{UL} f_L + f_U^T [D_U - W_{UU}] f_U.$$

by first-order optimality conditions, it is easy to see that the optimal satisfies

$$(D_U - W_{UU}) f_U = W_{UL} f_L.$$

If $D_U - W_{UU}$ is invertible¹ then

$$f_U^* = (D_U - W_{UU})^{-1} W_{UL} f_L.$$

Remark 5.9. The function f we constructed is called a harmonic extension. Indeed, it shares properties with harmonic functions in euclidean space such

¹It is not difficult to see that unless the problem is in some form degenerate, such as the unlabeled part of the graph being disconnected from the labeled one, then this matrix will indeed be invertible.

as the mean value property and maximum principles; if v_i is an unlabeled point then

$$f(i) = [D_U^{-1}(W_{UL}f_l + W_{UU}f_u)]_i = \frac{1}{\deg(i)} \sum_{j=1}^n w_{ij}f(j),$$

which immediately implies that the maximum and minimum value of f needs to be attained at a labeled point.

An interesting experience and the Sobolev Embedding Theorem

Let us try a simple experiment. Let's say we have a grid on $[-1, 1]^d$ dimensions (with say m^d points for some large m) and we label the center as $+1$ and every node that is at distance larger or equal to 1 to the center, as -1 . We are interested in understanding how the above algorithm will label the remaining points, hoping that it will assign small numbers to points far away from the center (and close to the boundary of the labeled points) and large numbers to points close to the center.

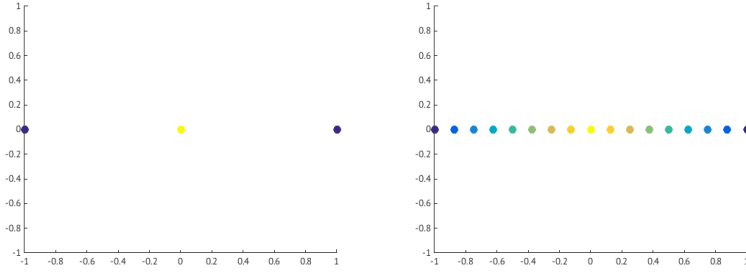


Fig. 5.11: The $d = 1$ example of the use of this method to the example described above, the value of the nodes is given by color coding. For $d = 1$ it appears to smoothly interpolate between the labeled points.

See the results for $d = 1$ in Figure 5.11, $d = 2$ in Figure 5.12, and $d = 3$ in Figure 5.13. While for $d \leq 2$ it appears to be smoothly interpolating between the labels, for $d = 3$ it seems that the method simply learns essentially -1 on all points, thus not being very meaningful. Let us turn to $\mathbb{R}^d$ for intuition:

Let's say that we want to find a function in $\mathbb{R}^d$ that takes the value 1 at zero and -1 at the unit sphere, that minimizes $\int_{B_0(1)} \|\nabla f(x)\|^2 dx$. Let us consider the following function on $B_0(1)$ (the ball centered at 0 with unit radius)

$$f_\varepsilon(x) = \begin{cases} 1 - 2\frac{|x|}{\varepsilon} & \text{if } |x| \leq \varepsilon \\ -1 & \text{otherwise.} \end{cases}$$

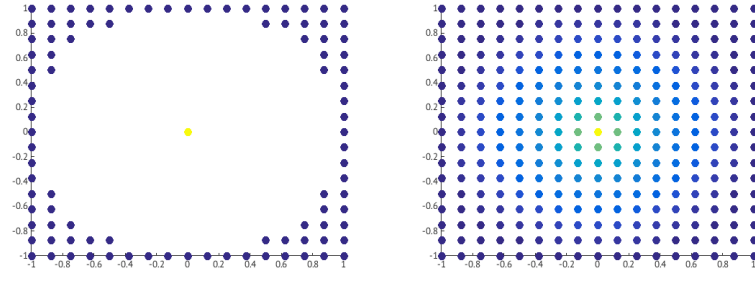


Fig. 5.12: The $d = 2$ example of the use of this method to the example described above, the value of the nodes is given by color coding. For $d = 2$ it appears to smoothly interpolate between the labeled points.

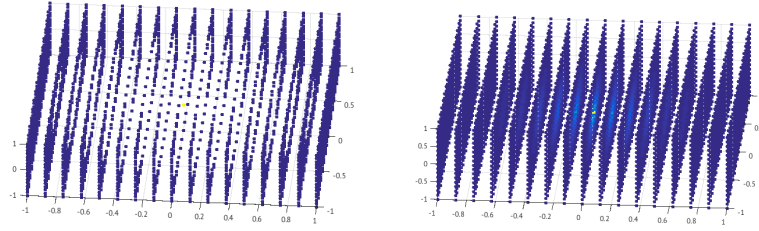


Fig. 5.13: The $d = 3$ example of the use of this method to the example described above, the value of the nodes is given by color coding. For $d = 3$ the solution appears to only learn the label -1 .

A quick calculation suggest that

$$\int_{B_0(1)} \|\nabla f_\varepsilon(x)\|^2 dx = \int_{B_0(\varepsilon)} \frac{1}{\varepsilon^2} dx = \text{vol}(B_0(\varepsilon)) \frac{1}{\varepsilon^2} dx \approx \varepsilon^{d-2},$$

meaning that, if $d > 2$, the performance of this function is improving as $\varepsilon \rightarrow 0$, explaining the results in Figure 5.13.

One way of thinking about what is going on is through the Sobolev Embedding Theorem. $H^m(\mathbb{R}^d)$ is the space of function whose derivatives up to order m are square-integrable in $\mathbb{R}^d$, Sobolev Embedding Theorem says that if $m > \frac{d}{2}$ then, if $f \in H^m(\mathbb{R}^d)$ then f must be continuous, which would rule out the behavior observed in Figure 5.13. It also suggests that if we are able to control also second derivates of f then this phenomenon should disappear (since $2 > \frac{3}{2}$). While we will not describe it here in detail, there is, in fact, a way of doing this by minimizing not $f^T L f$ but $f^T L^2 f$ instead, Figure 5.14 shows the outcome of the same experiment with the $f^T L f$ replaced by $f^T L^2 f$

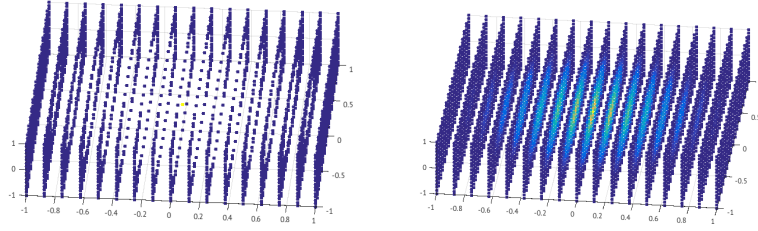


Fig. 5.14: The $d = 3$ example of the use of this method with the extra regularization $f^T L^2 f$ to the example described above, the value of the nodes is given by color coding. The extra regularization seems to fix the issue of discontinuities.

and confirms our intuition that the discontinuity issue should disappear (see, e.g., [136] for more on this phenomenon).

Linear Dimension Reduction via Random Projections

In Chapters 3 and 5 we saw both linear and non-linear methods for dimension reduction. In this chapter we will see one of the most fascinating consequences of the phenomenon of concentration of measure in high dimensions, one of the *blessings of high dimensions* described in Chapter 2. When given a data set in high dimensions, we will see that the projection to a lower dimensional space, taken at random, can preserve (under conditions made precise in the next section) certain geometric features of the data set. The remarkable aspect here is that this “lower” dimension can be strikingly lower. This gives rise to significant computational savings in many data processing tasks by including a random projection as a pre-processing step. Similar ideas are also the driving force behind many algorithms in *randomized linear algebra*, as we will demonstrate in this chapter via the *randomized SVD*. There is however another less obvious implication of this phenomenon with important practical implications: since the projection is agnostic of the data, it can be leveraged even when the data set is not explicit, such as the set of all *natural images* or the set of all “possible” *brain scans*; this is at the heart of *Compressive Sensing*, which we will explore in Chapter 10

6.1 The Johnson-Lindenstrauss Lemma

Suppose one has n points, $X = \{x_1, \dots, x_n\}$, in $\mathbb{R}^p$ (with p large). If $p > n$, the points actually lie in a subspace of dimension n , so the projection $f : \mathbb{R}^p \rightarrow \mathbb{R}^n$ of the points to that subspace acts without distorting the geometry of X . In particular, for every x_i and x_j , $\|f(x_i) - f(x_j)\|^2 = \|x_i - x_j\|^2$, meaning that f is an isometry in X . Suppose instead we allow a bit of distortion, and look for a map $f : \mathbb{R}^p \rightarrow \mathbb{R}^d$ that is an ε -isometry¹, meaning that

¹Another possible convention is to ask for $(1-\varepsilon)\|x_i - x_j\| \leq \|f(x_i) - f(x_j)\| \leq (1+\varepsilon)\|x_i - x_j\|$, this is equivalent up to a multiplicative constant on ε since for $0 \leq \varepsilon \leq 1$ we have $1+\varepsilon \leq (1+\varepsilon)^2 \leq 1+3\varepsilon$ and $1-2\varepsilon \leq (1-\varepsilon)^2 \leq 1-\varepsilon$.

$$(1 - \varepsilon)\|x_i - x_j\|^2 \leq \|f(x_i) - f(x_j)\|^2 \leq (1 + \varepsilon)\|x_i - x_j\|^2. \quad (6.1)$$

Can we do better than $d = n$?

In 1984, Johnson and Lindenstrauss [96] showed a remarkable lemma that answers this question affirmatively.

Theorem 6.1 (Johnson-Lindenstrauss Lemma [96]). *For any $0 < \varepsilon < 1$ and for any integer n , let d be such that*

$$d \geq 4 \frac{1}{\varepsilon^2/2 - \varepsilon^3/3} \log n. \quad (6.2)$$

Then, for any set X of n points in $\mathbb{R}^p$, there is a linear map $f : \mathbb{R}^p \rightarrow \mathbb{R}^d$ that is an ε -isometry for X (see (6.1)). This map can be found in randomized polynomial time².

We follow [58] for an elementary proof for Theorem 6.1. We need a few concentration of measure bounds. We will omit the proof of those but they are available in [58] and are essentially the same ideas as those used to show the concentration inequalities in Chapter 2.

Lemma 6.2 (see [58]). *Let $y_1, \dots, y_p$ be i.i.d standard Gaussian random variables and $Y = (y_1, \dots, y_p)$. Let $g : \mathbb{R}^p \rightarrow \mathbb{R}^d$ be the projection into the first d coordinates and $Z = g\left(\frac{Y}{\|Y\|}\right) = \frac{1}{\|Y\|}(y_1, \dots, y_d)$ and $L = \|Z\|^2$. It is clear that $\mathbb{E}L = \frac{d}{p}$. In fact, L is very concentrated around its mean*

- If $\beta < 1$,

$$\Pr \left[L \leq \beta \frac{d}{p} \right] \leq \exp \left(\frac{d}{2} (1 - \beta + \log \beta) \right).$$

- If $\beta > 1$,

$$\Pr \left[L \geq \beta \frac{d}{p} \right] \leq \exp \left(\frac{d}{2} (1 - \beta + \log \beta) \right).$$

Proof. [of Johnson-Lindenstrauss Lemma]

We will start by showing that, given a pair x_i, x_j a projection onto a random subspace of dimension d will satisfy (after appropriate scaling) property (6.1) with high probability. Without loss of generality we can assume that $u = x_i - x_j$ has unit norm. Understanding what is the norm of the projection of u on a random subspace of dimension d is the same as understanding the norm of the projection of a (uniformly) random point on S^{p-1} the unit sphere in $\mathbb{R}^p$ on a specific d -dimensional subspace—let us say the one generated by the first d canonical basis vectors.

²Randomized polynomial time refers to the complexity class of decision problems for which a randomized algorithm can solve the problem in polynomial time with a high probability of correctness.

This means that we are interested in the distribution of the norm of the first d entries of a random vector drawn from the uniform distribution over S^{p-1} – this distribution is the same as taking a standard Gaussian vector in $\mathbb{R}^p$ and normalizing it to the unit sphere.

Let $g : \mathbb{R}^p \rightarrow \mathbb{R}^d$ be the projection on a random d -dimensional subspace and let $f : \mathbb{R}^p \rightarrow \mathbb{R}^d$ defined as $f = \sqrt{\frac{p}{d}}g$. Then (by the above discussion), given a pair of distinct points x_i and x_j , $\frac{\|f(x_i) - f(x_j)\|^2}{\|x_i - x_j\|^2}$ has the same distribution as $\frac{p}{d}L$, as defined in Lemma 6.2. Using Lemma 6.2, we have, given a pair x_i, x_j ,

$$\Pr \left[\frac{\|f(x_i) - f(x_j)\|^2}{\|x_i - x_j\|^2} \leq (1 - \varepsilon) \right] \leq \exp \left(\frac{d}{2} (1 - (1 - \varepsilon) + \log(1 - \varepsilon)) \right),$$

since for $\varepsilon \geq 0$, $\log(1 - \varepsilon) \leq -\varepsilon - \varepsilon^2/2$ we have

$$\begin{aligned} \Pr \left[\frac{\|f(x_i) - f(x_j)\|^2}{\|x_i - x_j\|^2} \leq (1 - \varepsilon) \right] &\leq \exp \left(-\frac{d\varepsilon^2}{4} \right) \\ &\leq \exp(-2 \log n) = \frac{1}{n^2}. \end{aligned}$$

On the other hand,

$$\Pr \left[\frac{\|f(x_i) - f(x_j)\|^2}{\|x_i - x_j\|^2} \geq (1 + \varepsilon) \right] \leq \exp \left(\frac{d}{2} (1 - (1 + \varepsilon) + \log(1 + \varepsilon)) \right).$$

since for $\varepsilon \geq 0$, $\log(1 + \varepsilon) \leq \varepsilon - \varepsilon^2/2 + \varepsilon^3/3$ we have

$$\begin{aligned} \mathbb{P} \left[\frac{\|f(x_i) - f(x_j)\|^2}{\|x_i - x_j\|^2} \geq (1 + \varepsilon) \right] &\leq \exp \left(-\frac{d(\varepsilon^2 - 2\varepsilon^3/3)}{4} \right) \\ &\leq \exp(-2 \log n) = \frac{1}{n^2}. \end{aligned}$$

By the union bound it follows that

$$\Pr \left[\frac{\|f(x_i) - f(x_j)\|^2}{\|x_i - x_j\|^2} \notin [1 - \varepsilon, 1 + \varepsilon] \right] \leq \frac{2}{n^2}.$$

Since there exist $\binom{n}{2}$ such pairs, again, a simple union bound gives

$$\Pr \left[\exists_{i,j} : \frac{\|f(x_i) - f(x_j)\|^2}{\|x_i - x_j\|^2} \notin [1 - \varepsilon, 1 + \varepsilon] \right] \leq \frac{2}{n^2} \frac{n(n-1)}{2} = 1 - \frac{1}{n}.$$

Therefore, choosing f as a properly scaled projection onto a random d -dimensional subspace gives an ε -isometry on X (see (6.1)) with probability at least $\frac{1}{n}$. We can achieve any desirable constant probability of success by

trying $\mathcal{O}(n)$ such random projections, meaning we can find an ε -isometry in randomized polynomial time. $\square$

Note that by considering d slightly larger one can get a good projection on the first random attempt with high confidence. In fact, it is trivial to adapt the proof above to obtain the following lemma:

Proposition 6.3. *For any $0 < \varepsilon < 1$, $\tau > 0$, and for any integer n , let d be such that*

$$d \geq \frac{2(2 + \tau)}{\varepsilon^2/2 - \varepsilon^3/3} \log n.$$

Then, for any set X of n points in $\mathbb{R}^p$, take $f : \mathbb{R}^p \rightarrow \mathbb{R}^d$ to be a suitably scaled projection on a random subspace of dimension d , then f is an ε -isometry for X (see (6.1)) with probability at least $1 - \frac{1}{n^\tau}$.

Proposition 6.3 is quite remarkable. Consider the situation where we are given a high-dimensional data set in a streaming fashion – meaning that we get each data point at a time, consecutively. To run a dimension-reduction technique like PCA or Diffusion maps we would need to wait until we received the last data point and then compute the dimension reduction map (both PCA and Diffusion Maps are, in some sense, data adaptive). Using Lemma 6.3 one can just choose a projection at random in the beginning of the process (all one needs to know is an estimate of the logarithm of the size of the data set) and just map each point using this projection matrix which can be done online – we do not need to see the next point to compute the projection of the current data point. Proposition 6.3 ensures that this (seemingly naïve) procedure will, with high probability, not distort the data by more than ε .

One might wonder if such low-dimensional embeddings as provided by the Johnson-Lindenstrauss Lemma also extend to other norms besides the Euclidean norm. For the ℓ_1 -norm there exist lower bounds which prevent such a dramatic dimension reduction (see [113]), and for the ℓ_∞ -norm one can easily construct examples that demonstrate the impossibility of dimension reduction. Hence, the Johnson-Lindenstrauss Lemma seems to be a subtle result about the Euclidean norm.

6.1.1 The Fast Johnson-Lindenstrauss transform and optimality

Let us continue thinking about our example of high-dimensional streaming data. After we draw the random projection matrix³, say M , for each data point x , we still have to compute Mx which has a computational cost of

³An orthogonal projection P must satisfy $P = P^*$ and $P^2 = P$. Here, it is not M that represents a projection, but M^*M , yet for our purposes of approximate norm-preserving dimension reduction it suffices to apply M instead of M^*M . However, with a slight abuse of terminology, we still refer to M as projection.

$\mathcal{O}(\varepsilon^{-2} \log(n)p)$ since M has $\mathcal{O}(\varepsilon^{-2} \log(n)p)$ entries (since M is a random matrix, generically it will be a dense matrix). In some applications this might be too expensive, raising the natural question of whether one can do better. Moreover, storing a large-scale dense matrix M is not very desirable either. There is no hope of significantly reducing the number of rows in general, as it is known that the Johnson-Lindenstrauss Lemma is orderwise optimal [9, 107].

We might hope to replace the dense random matrix M by a sparse matrix S to speed up the matrix-vector multiplication and to reduce the storage requirements. This method was proposed and analyzed in [6]. Here we discuss a slightly simplified version, see also [57].

We let S be a very sparse $d \times p$ matrix, where each row of S has just one single non-zero entry of value $\sqrt{p/d}$ at a location drawn uniformly at random. Then, for any vector $x \in \mathbb{R}^p$

$$\mathbb{E}_i[(Sx)_i^2] = \sum_{j=1}^p \mathbb{P}(S_{ij} \neq 0) \cdot \frac{p}{d} \cdot x_j^2 = \frac{1}{d} \|x\|_2^2,$$

hence $\mathbb{E}[\|Sx\|_2^2] = \mathbb{E}[\sum_{i=1}^d (Sx_i)^2] = \|x\|_2^2$. This result is satisfactory with respect to expectation (even for $d = 1$), but not with respect to the variance of $\|Sx\|_2^2$. For instance, if x has only one non-zero entry we need $d \sim \mathcal{O}(p)$ to ensure that $\|Sx\|_2^2 \neq 0$ with non-negligible probability. More generally, if one coordinate of x is much larger (in absolute value) than all its other coordinates, then we will need a rather large value for d to guarantee that $\|Sx\|_2 \approx \|x\|_2$.

A natural way to quantify the “peakiness” of a vector x is via the *peak-to-average ratio*⁴ measured by the quantity $\|x\|_\infty / \|x\|_2$. It is easy to see that we have (assuming x is not the zero-vector)

$$\frac{1}{\sqrt{p}} \leq \frac{\|x\|_\infty}{\|x\|_2} \leq 1.$$

The upper bound is achieved by vectors with only one non-zero entry, while the lower bound is met by constant-modulus vectors. Thus, if

$$\frac{\|x\|_\infty}{\|x\|_2} \approx \frac{1}{\sqrt{p}}, \quad (6.3)$$

we can hope that sparse subsampling of x will still preserve its Euclidean norm.

⁴This quantity also plays an important role in wireless communications. There, one tries to avoid transmitting signals with a large peak-to-average ratio, since such signals would suffer from nonlinear distortions when they are passing through the power amplifiers that are usually installed in cell phones. The potentially large peak-to-average ratio of OFDM signals is one of the alleged reasons why CDMA was dominant over OFDM for such a long time.

Thus, this suggests to include a preprocessing step by applying a rotation so that sparse vectors become non-sparse in the new basis, thereby reducing their ∞ -norm (while their 2-norm remains invariant under rotation). Two natural choices for such a rotation are the Discrete Fourier transform (which maps unit-vectors into constant modulus vectors) and its $\mathbb{Z}_2$ -cousin, the Walsh-Hadamard matrix⁵. But since the Fast Johnson-Lindenstrauss Transform (FJLT) has to work for all vectors, we need to avoid that this rotation maps dense vectors into sparse vectors. We can address this issue by “randomizing” the rotation, thereby ensuring with overwhelming probability that dense vectors are not mapped into sparse vectors. This can be accomplished in a numerically efficient manner (thus maintaining our overall goal of numerical efficiency) by first randomizing the signs of x before applying the rotation. Putting these steps together we arrive at the following map.

Definition 6.4. *The Fast Johnson-Lindenstrauss Transform is the map $\Psi : \mathbb{C}^p \rightarrow \mathbb{C}^d$, defined by $\Psi := SFD$, where S and D are random matrices and F is a deterministic matrix. In particular,*

- S is a $d \times p$ matrix, where each row of S has just one single non-zero entry of value $\sqrt{p/d}$ at a location drawn uniformly at random.
- F is either the $p \times p$ DFT matrix or the $p \times p$ Hadamard matrix (if it exists), in each case normalized by $1/\sqrt{p}$ to obtain a unitary matrix.
- D is a $p \times p$ diagonal matrix whose entries are drawn independently from $\{-1, +1\}$ with probability $1/2$.

We can carry out the matrix-vector multiplication by the DFT matrix via the Fast Fourier Transform (FFT) in $\mathcal{O}(p \log p)$ operations; a similar algorithm exists for the Walsh-Hadamard matrix. The FJLT allows for a dimension reduction that is competitive with the Johnson-Lindenstrauss Lemma as manifested by the following theorem.

Theorem 6.5 (Fast Johnson-Lindenstrauss Transform). *For $0 < \varepsilon < 1$ and $0 < \delta < 1$, there is a random matrix Ψ of size $d \times p$ with $d = \mathcal{O}(\log(p/\delta) \log(1/\delta)/\varepsilon^2)$ such that, for each $x \in \mathbb{C}^p$,*

$$\|\Psi x\|_2 \in [1 - \varepsilon, 1 + \varepsilon] \cdot \|x\|_2$$

holds with probability at least $1 - \delta$. Matrix-vector multiplication with Ψ takes $\mathcal{O}(p \log p + d)$ operations.

For convenience we will prove Theorem 6.5 for the case when F is the Walsh-Hadamard matrix so that the random variables are real-valued and the exposition is lighter. But it is straightforward to adapt the argument to the case when F is the DFT matrix. Keeping this minor simplification in mind, the proof of Theorem 6.5 follows from the two lemmas below. We first show that with high probability the random rotation FD produces vectors with a sufficiently low peak-to-average ratio.

⁵Hadamard matrices do not exist for all dimension d . But we can always pad x with zeroes to achieve the desired length.

Lemma 6.6. *Let $y = FDx$, where F and D are as in Definition 6.4. Then*

$$\mathbb{P}_D \left(\frac{\|y\|_\infty}{\|y\|_2} \geq \sqrt{\frac{2 \log(4p/\delta)}{p}} \right) \leq \frac{\delta}{2}. \quad (6.4)$$

Proof. Since FD is unitary, the quantity $\|FDx\|_\infty/\|FDx\|_2$ is invariant under rescaling of x and therefore we can assume $\|x\|_2 = 1$.

Let $\xi_i = \pm 1$ be the i -th diagonal entry of D . We have $y_i = \sum_{j=1}^p \xi_j F_{ij} x_j$ and note that the terms of this sum are i.i.d. bounded random variables. We thus can apply Hoeffding's inequality. In the notation of Theorem 2.15, let $X_j = \xi_j F_{ij} x_j$. We note that $X_j = \pm F_{ij} x_j$, hence $\mathbb{E}[X_j] = 0$ and $|X_j| \leq a_j$, where $a_j = |F_{ij} x_j|$. It holds that

$$\sum_{j=1}^p a_j^2 = \sum_{j=1}^p |F_{ij}|^2 x_j^2 = \sum_{j=1}^p \frac{1}{p} x_j^2 = \frac{\|x\|_2^2}{p} = \frac{1}{p}.$$

We can now use Theorem 2.15 with $t = \sqrt{2 \log(4p/\delta)/p}$ and obtain

$$\mathbb{P} \left(|y_i| > \sqrt{\frac{2 \log(4p/\delta)}{p}} \right) \leq 2 \exp \left(-\frac{2 \log(4p/\delta)/p}{2/p} \right) = \frac{\delta}{2p}.$$

Applying the union bound finishes the proof. $\square$

Lemma 6.7. *Let $0 < \varepsilon < 1$ and $y \in \mathbb{R}^p$ satisfy $\|y\|_\infty^2 < \frac{2 \log(4p/\delta)}{p}$ and $\|y\|_2 = 1$, then it possible to pick $d \sim 1/\varepsilon^2 \log(p/\delta) \log(1/\delta)$ such that*

$$\mathbb{P}(\|Sy\|_2^2 - 1 \leq \varepsilon) \leq 1 - \frac{\delta}{2}.$$

Proof. The idea is to use Bernstein's Inequality (Theorem 2.18). We write $S_{ji} = \sqrt{p/d} \delta_{ji}$, where $\delta_{ji} \in \{0, 1\}$ is our random sample of the columns for row j . Hence for all j , $\sum_{i=1}^p \delta_{ji} = 1$. We write $z := Sy$ and compute

$$\begin{aligned} q_j &:= z_j^2 \\ &= \frac{p}{d} \left(\sum_{i=1}^p \delta_{ji} y_i \right)^2 \\ &= \frac{p}{d} \left(\sum_i \delta_{ji} y_i^2 + \sum_{i \neq \ell} \delta_{ji} \delta_{j\ell} y_i y_\ell \right) \\ &= \frac{p}{d} \sum_i \delta_{ji} y_i^2. \end{aligned}$$

We want to bound the random variable $\|z\|^2 - 1 = \sum_j (q_j - \frac{1}{d})$. Since the q_j 's are independent, and $\mathbb{E}[q_j - \frac{1}{d}] = 0$, we can apply Bernstein's Inequality (Theorem 2.18, provided we bound σ^2 and a . Firstly,

$$a \leq \frac{p}{d} \|y\|_\infty^2 - \frac{1}{d} < \frac{p}{d} \|y\|_\infty^2 \leq \frac{2 \log(4p/\delta)}{d}.$$

and secondly,

$$\begin{aligned} d\sigma^2 &\leq d\mathbb{E}[q_1^2] - 1 \leq d\mathbb{E}[q_1^2] = d\mathbb{E}\left[\frac{p^2}{d^2} \sum_i \delta_{ji} y_i^4\right] \\ &= \frac{p^2}{d} \frac{1}{p} \sum_i y_i^4 \\ &\leq \|y\|_\infty^2 \frac{p}{d} \sum_i y_i^2 \\ &= \frac{p}{d} \|y\|_\infty^2 \|y\|^2 \\ &= \frac{p}{d} \|y\|_\infty^2 \\ &\leq \frac{2 \log(4p/\delta)}{d}. \end{aligned}$$

Plugging these terms into Bernstein's Inequality (Theorem 2.18) gives

$$\begin{aligned} \mathbb{P}(|\|Sy\|^2 - 1| > \varepsilon) &= \mathbb{P}\left(\left|\sum_{j=1}^d q_j - 1\right| > \varepsilon\right) \\ &\leq 2 \exp\left(\frac{-\varepsilon^2}{2d \left(\frac{2 \log(4p/\delta)}{d^2}\right) + \frac{2}{3} \left(\frac{2 \log(4p/\delta)}{d}\right) \varepsilon}\right) \\ &\leq 2 \exp\left(\frac{-\varepsilon^2}{\frac{8}{3} \left(\frac{2 \log(4p/\delta)}{d}\right)}\right), \\ &= 2 \exp\left(\frac{-3\varepsilon^2 d}{16 \log(4p/\delta)}\right), \end{aligned}$$

where the last inequality uses $\varepsilon < 1$. Picking $d = \frac{16}{3\varepsilon^2} \log(4p/\delta) \log(4/\delta)$ gives the desired $\delta/2$ -bound. $\square$

Combining Lemmas 6.6 and 6.7, union bounding over the two events, and recalling that if $\|x\| = 1$ then $\|y\| = 1$ since both D and F are unitary, establishes Theorem 6.5.

Besides the potential speedup, another advantage of the FJLT is that it requires significantly less memory compared to storing an unstructured random projection matrix as is the case for the standard Johnson-Lindenstrauss approach.

There are many applications where one wants to perform dimension reduction while preserving the geometry of certain sets with infinite points, such as subspaces. Via an ε -net argument it is possible to show that random matrices

that are good as Johnson-Lindenstrauss maps for N points, also approximately preserve the geometry of a $\log N$ dimensional subspace. These go by the name of Oblivious Subspace Embeddings (see, e.g., Section 5.2.2. in [102] where the fast JL construction we show above is shown to be a Oblivious Subspace Embedding). A similar idea will be explored below (Section 6.2) where a random projection will allow us to still approximate the subspace spanned by the singular vectors corresponding to the largest singular values, resulting in a randomized SVD algorithm that is workhorse of large scale linear algebra. In Chapter 8 we will develop analogues of Johnson-Lindenstrauss' lemma for more general sets, through the celebrated Gordon's espace through a mesh theorem.

6.2 Randomized SVD

Randomized linear algebra has gained significant relevance in recent years due to its ability to efficiently handle large-scale data problems that traditional methods struggle with. The Johnson-Lindenstrauss dimension reduction discussed in the previous section is one prominent example. The techniques encountered in the Johnson-Lindenstrauss transform naturally lend themselves as key ingredient of the *randomized SVD*, a method to compute an approximation of the SVD of a matrix. It is especially useful when the matrix is large and dense and can be well approximated by a low-rank matrix.

In this section we will give a detailed description of the randomized SVD. It is based on using random projections to reduce the dimensionality of the matrix before performing a more traditional SVD on the reduced matrix.

Recall from Chapter 3 that the SVD of a matrix $A \in \mathbb{R}^{m \times n}$ is given by

$$A = U \Sigma V^*, \quad (6.5)$$

where $U \in \mathbb{R}^{m \times m}$ is a unitary matrix containing the left singular vectors, $\Sigma \in \mathbb{R}^{m \times n}$ is a diagonal matrix with singular values (in decreasing order) on the diagonal, and $V \in \mathbb{R}^{n \times n}$ is a unitary matrix containing the right singular vectors.

In practice A is often large, and we are interested in only a small number s of the largest singular values and corresponding singular vectors, giving a rank- k approximation of A :

$$A \approx U_k \Sigma_k V_k^*$$

where $U_k \in \mathbb{R}^{m \times k}$, $\Sigma_k \in \mathbb{R}^{k \times k}$, and $V_k \in \mathbb{R}^{n \times k}$.

We define the projection P_k onto the the span of the leading k left singular vectors of A via

$$P_k = \sum_{i=1}^k u_i u_i^*.$$

We know from Theorem 3.7 that the best rank- k approximation of A with respect to the Frobenius norm is given by

$$\min_{\text{rank}(B) \leq k} \|A - B\|_F = \|A - P_k A\|_F = \sqrt{\sum_{i>k} \sigma_i^2}.$$

Thus, we can compute the best approximation by projecting the target matrix onto the k -dimensional subspace that captures “most of the action” of the matrix A .

In many applications it suffices to compute the best approximation with moderate accuracy. Hence, following [89], we let $s = k + p$ for small p , and seek a rank- s approximation $\hat{A}_s$ that is comparable to the best rank- k approximation, i.e.,

$$\|A - \hat{A}_s\|_F^2 \leq (1 + \varepsilon) \|A - P_k A\|_F^2 = (1 + \varepsilon) \sum_{i>k} \sigma_i^2$$

for some chosen tolerance $\varepsilon > 0$.

The aim of the Randomized SVD (RSVD) is essentially to generate an efficient approximation for the best rank- k approximation of the matrix $P_k A$. We proceed in two steps. In the first step we compute a low-dimensional subspace that captures the range of A . In the second step we project A onto the subspace and compute the (standard SVD) of the resulting smaller matrix. Regarding the first step, recall that P_k is the orthogonal projection onto the k leading left singular vectors of A . But since we do not know the subspace spanned by P_k , a sensible choice, as suggested by the Johnson-Lindenstrauss Lemma, is to draw some *probing* random vectors from a rotationally invariant distribution to estimate the directions of the vectors u_i that make up P_k . We collect these probing random vectors in a matrix Ψ .

Hence, let $\Psi \in \mathbb{R}^{n \times (k+p)}$ be a random matrix, where k is the desired number of singular values/vectors, and p is an oversampling parameter (typically p is small, like 5 or 10). A promising choice for Ψ is to select each entry of Ψ as a sample from a standard normal distribution $\mathcal{N}(0, 1)$. We then form the matrix $Y \in \mathbb{R}^{m \times (k+p)}$ as

$$Y = A\Psi.$$

This matrix Y captures the action of A on a random subspace of dimension $k + p$. We will see (in Theorem 6.8) that, with high probability, the subspace spanned by the columns of Y contains a good approximation of the subspace spanned by the top k left singular vectors of A .⁶ Next, we find an orthonormal basis for the range of Y . Using the standard QR decomposition, we compute

$$Y = QR,$$

⁶In many applications the matrix A is implicit and one has access to it via matrix-vector products, if this is the case computing Y can be done with $k + p$ matrix-vector products.

where $Q \in \mathbb{R}^{m \times (k+p)}$ has orthonormal columns, and $R \in \mathbb{R}^{(k+p) \times (k+p)}$ is an upper triangular matrix. Now, we project A onto the lower-dimensional subspace spanned by the columns of Q by computing

$$B = Q^* A.$$

Here, $B \in \mathbb{R}^{(k+p) \times n}$ is a much smaller matrix that approximates A (in the sense that QB approximates A).

Next, we compute the exact SVD of the smaller matrix B

$$B = U_B \Sigma_B V_B^*.$$

where $U_B \in \mathbb{R}^{(k+p) \times (k+p)}$, $\Sigma_B \in \mathbb{R}^{(k+p) \times n}$, and $V_B \in \mathbb{R}^{n \times (k+p)}$. Finally, the approximate SVD of A is given by

$$A \approx \tilde{U} \tilde{\Sigma} \tilde{V}^*$$

where $\tilde{U} = QU_B \in \mathbb{R}^{m \times (k+p)}$, $\tilde{\Sigma} = \Sigma_B \in \mathbb{R}^{(k+p) \times (k+p)}$, and $\tilde{V} = V_B \in \mathbb{R}^{n \times (k+p)}$.

We summarize these steps in Algorithm 6.1.

Algorithm 6.1 Randomized SVD

Input: Matrix $A \in \mathbb{R}^{m \times n}$, target rank $k > 0$, oversampling factor $p > 0$; set $s = k + p$.

1. Generate a random matrix $\Psi \in \mathbb{R}^{N \times (k+p)}$
2. Compute the product $Y = A\Psi$
3. Compute the economy-sized QR factorization $Y = QR$, with $Q \in \mathbb{R}^{m \times (k+p)}$ and $R \in \mathbb{R}^{(k+p) \times (k+p)}$
4. Compute the product $B = Q^* A$
5. Compute the economy-sized SVD $B = U_B \Sigma_B V_B^*$, with $U_B \in \mathbb{R}^{m \times (k+p)}$, $\Sigma_B \in \mathbb{R}^{(k+p) \times (k+p)}$ and $V_B \in \mathbb{R}^{n \times (k+p)}$
6. Set $\tilde{U} = QU_B$, $\tilde{\Sigma} = \Sigma_B$, $\tilde{V} = V_B$.

Output: Orthogonal $\tilde{U} \in \mathbb{R}^{m \times (k+p)}$, orthogonal $\tilde{V} \in \mathbb{R}^{n \times (k+p)}$, and diagonal $\tilde{\Sigma} \in \mathbb{R}^{(k+p) \times (k+p)}$ such that $A \sim \hat{A}_s := \tilde{U} \tilde{\Sigma} \tilde{V}^*$

The random matrix Ψ used in the random projection step ensures that with high probability, the subspace spanned by the columns of $Y = A\Psi$ contains a good approximation to the dominant singular subspace of A .

Theorem 6.8 ([89]). *Consider a matrix $A \in \mathbb{R}^{m \times n}$ with singular values $\sigma_1 \geq \sigma_2 \geq \sigma_3 \geq \dots$. Fix the target rank $k \leq \min(m, n)$. When $s := k + p \geq k + 2$, the randomized SVD method produces a random rank- s approximation $\hat{A}_s$ that satisfies*

$$\mathbb{E}[\|A - \hat{A}_s\|_F] \leq \left(1 + \frac{k}{p-1}\right)^{\frac{1}{2}} \left(\sum_{j>k} \sigma_j^2\right)^{\frac{1}{2}}. \quad (6.6)$$

The error on the right hand side consists of the best rank- k approximation error and an additional factor due to using a randomized projection instead of the optimal projection P_k .

For the proof of Theorem 6.8 we need some preparation. In what follows it will be convenient to rewrite the SVD of A , cf. (6.5), in block matrix form:

$$A = \begin{bmatrix} U_k & U_\perp \end{bmatrix} \begin{bmatrix} \Sigma_k & 0 \\ 0 & \Sigma_\perp \end{bmatrix} \begin{bmatrix} V_k^* \\ V_\perp^* \end{bmatrix}, \quad (6.7)$$

where $U_\perp \in \mathbb{R}^{m \times m-k}$ denotes the orthogonal complement of $U_k \in \mathbb{R}^{m \times k}$ with respect to $U \in \mathbb{R}^{m \times m}$. Furthermore, we set

$$\Psi_k := V_k^* \Psi \quad \text{and} \quad \Psi_\perp := V_\perp^* \Psi.$$

Here, Ψ_k captures the alignment of Ψ with the leading right singular vectors of A , while $\Psi_\perp$ captures the alignment with the trailing right singular vectors of A . The matrix P_Y denotes the orthogonal projection onto the range of Y . Since this projection is unique, we have that $P_Y = QQ^*$.

The following deterministic error bound (the bound does not require Ψ to be Gaussian) will be instrumental:

Lemma 6.9. *We use the same notation as in Theorem 6.8. Let $\Psi \in \mathbb{R}^{n \times (k+p)}$ be an arbitrary matrix with $\text{rank}(\Psi) = k$. Then the rank- s approximation $\hat{A}_s$ computed by Algorithm 6.1 satisfies*

$$\|A - \hat{A}_s\|_F^2 \leq \|\Sigma_\perp\|_F^2 + \|\Sigma_k \Psi_k \Psi_\perp^\dagger\|_F^2. \quad (6.8)$$

Proof. We first fix the target rank $k \leq \min\{m, n\}$ and the oversampling factor p and set $s = k + p$. Denoting $\tilde{A} = U^* A$ and $\tilde{Y} = \tilde{A} \Psi$, we obtain

$$\|A - \hat{A}_s\|_F^2 = \|(I - P_Y)A\|_F^2 = \|U^*(I - P_Y)U\tilde{A}\|_F^2 = \|(I - P_{\tilde{Y}})\tilde{A}\|_F^2. \quad (6.9)$$

Here, the first equality uses the fact that the Frobenius norm is unitarily invariant, whereas in the second equality we have substituted $P_{\tilde{Y}} = U^* P_Y U$. Next, we set $Z = \tilde{Y} \Psi_k^\dagger \Sigma_k^{-1} =: [I \ F]^*$ where $F = \Sigma_\perp \Psi_\perp \Psi_k^\dagger \Sigma_k^{-1}$. Note that $\text{range}(Z) \subseteq \text{range}(\tilde{Y})$. This in turn yields

$$\|(I - P_{\tilde{Y}})\tilde{A}\|_F^2 \leq \|(I - P_Z)\tilde{A}\|_F^2 = \text{Tr}(\tilde{A}^*(I - P_Z)\tilde{A}) = \text{Tr}(\Sigma(I - P_Z)\Sigma). \quad (6.10)$$

Note that P_Z has the expansion

$$P_Z = \begin{bmatrix} I \\ F \end{bmatrix} [I + F^* F]^{-1} \begin{bmatrix} I \\ F \end{bmatrix}^*.$$

Since P_Z is a projection, we have that $I - P_Z \succ 0$. After some linear algebra calculations (left to the reader), we obtain

$$0 \prec I - P_Z \preceq \begin{bmatrix} F^* F & B \\ B^* & I \end{bmatrix},$$

where $B = -(I + F^*F)^{-1}F^*$. Multiplying both sides by Σ gives

$$\Sigma(I - P_Z)\Sigma \preceq \begin{bmatrix} \Sigma_k F^* F \Sigma_k & \Sigma_k B \Sigma_\perp \\ \Sigma_\perp B^* \Sigma_k & \Sigma_\perp \Sigma_\perp \end{bmatrix}. \quad (6.11)$$

Combining (6.9), (6.10) and (6.11) gives

$$\begin{aligned} \|A - \hat{A}_s\|_F^2 &= \|(I - P_Y)\tilde{A}\|_F^2 \\ &\leq \text{Tr}(\Sigma(I - P_Z)\Sigma) \\ &\leq \text{Tr}(\Sigma_k F^* F \Sigma_k) + \text{Tr}(\Sigma_\perp) \\ &\leq \|\Sigma_\perp \Psi_\perp \Psi_k^\dagger\|_F^2 + \|\Sigma_\perp\|_F^2, \end{aligned}$$

where in the penultimate inequality we have used the fact that the trace is monotone on the cone of positive semidefinite matrices. $\square$

Lemma 6.9 holds for *any* rank- k matrix Ψ . As pointed out in [102], the first term on the right-hand side of (6.8) captures the trailing singular values of the matrix, which is independent of the specific choice of Ψ , since this loss occurs for every rank- k approximation. The second term is the error caused by the dimension reducing matrix Ψ . We want the component Ψ_k in the leading directions to be well conditioned (ideally, an identity matrix). We want the component $\Psi_\perp$ in the trailing directions to be as small as possible.

The bound (6.8) suggests that in order to establish Theorem 6.8 we need to get a handle on $\Psi_k \Psi_\perp^\dagger$ when $\Psi_{ij} \stackrel{\text{i.i.d.}}{\sim} \mathcal{N}(0, 1)$. To that end, the following lemma will be helpful.

Lemma 6.10. *Let $G \in \mathbb{R}^{m \times n}$ be a standard Gaussian random matrix. Assume $m \geq 2$ and $n - m \geq 2$. Then*

$$\mathbb{E}[\|G^\dagger\|_F^2] = \frac{m}{m - n - 1}. \quad (6.12)$$

Proof. We follow the proof of the corresponding result in the appendix of [89]. First, note that (GG^*) is invertible almost surely. We compute

$$\|G^\dagger\|_F^2 = \text{trace}[(G^\dagger)^* G^\dagger] = \text{trace}[(GG^*)^{-1}].$$

The random matrix $(GG^*)^{-1}$ follows the inverted Wishart distribution. Hence, we can use formula (12) on page 97 of [134] for the expected trace and obtain (6.12). $\square$

With these lemmata in place, we are now ready to prove Theorem 6.8.

Proof (of Theorem 6.8). Since Ψ is chosen to be a standard Gaussian matrix and the Gaussian distribution is invariant under unitary transformations, it follows that $\Psi_k = V_k^* \Psi$ and $\Psi_\perp = V_\perp^* \Psi$ are also standard Gaussian. Moreover,

note that matrix $\Psi_\perp$ has full rank with probability one, since the rows of a (fat) Gaussian matrix are almost surely in general position.

We consider the bound established in Lemma 6.9, apply expectation to (6.8) and use Hölder's inequality to obtain

$$\mathbb{E}[\|A - \hat{A}_s\|_F] \leq (\mathbb{E}[\|(I - P_Y)A\|_F^2])^{1/2} \leq \left(\|\Sigma_\perp\|_F^2 + \mathbb{E}[\|\Sigma_k \Psi_k \Psi_\perp^\dagger\|_F^2] \right)^{1/2}. \quad (6.13)$$

Using the law of total expectation as well as the independence of Ψ_k and $\Psi_\perp$ we estimate

$$\begin{aligned} \mathbb{E}[\|\Sigma_k \Psi_k \Psi_\perp^\dagger\|_F^2] &= \mathbb{E}_{\Psi_k} [\mathbb{E}[\|\Sigma_\perp \Psi_\perp \Psi_k^\dagger\|_F^2 | \Psi_k]] \\ &\leq \mathbb{E}[\|\Sigma_\perp\|_F^2 \|\Psi_k^\dagger\|_F^2] \\ &\leq \|\Sigma_\perp\|_F^2 \mathbb{E}[\|\Psi_k^\dagger\|_F^2] \\ &\leq \frac{k}{p-1} \sum_{j>k} \sigma_j^2, \end{aligned} \quad (6.14)$$

where the second inequality follows from Lemma 6.10.

Combining now (6.13) with (6.14) establishes the bound (6.6). $\square$

6.2.1 Computational complexity of the Randomized SVD

What is the computational complexity of the RSVD? The dominant costs for the RSVD are:

1. $\mathcal{O}(mn(k+p))$ for the matrix-matrix multiplication $A\Psi$, assuming that Ψ is an unstructured, non-sparse matrix (like a Gaussian matrix).
2. $\mathcal{O}(m(k+p)^2)$ for the QR decomposition of Y .
3. $\mathcal{O}(mn(k+p))$ for computing Q^*A .
4. $\mathcal{O}((k+p)^2n)$ for computing the SVD of the small matrix B .

Thus, the overall complexity of the randomized SVD is $\mathcal{O}(mn(k+p))$, which is significantly smaller than the $\mathcal{O}(mn^2)$ cost of the classical SVD when k is much smaller than n .

We can further reduce the cost of the matrix-matrix product $A\Psi$ by introducing structure or sparsity into Ψ . We already encountered one such choice in Definition 6.4 in the form of the Fast Johnson-Lindenstrauss transform. The improved computational efficiency comes at the cost of a log-factor in terms of required samples to guarantee the same accuracy.

When the singular values of A decay slowly, Algorithm 6.1 is not very accurate in the regime of interest, namely when $k \ll n$. The reason is that to achieve a small approximation error the matrix Q must align well with the first $k+p$ left singular vectors of A . As a consequence, the product $A\Psi$ must reveal the singular vectors, which in turn would require the singular vectors associated with small singular values not to interfere with the calculation.

How can we remedy this issue? We can take inspiration from the power iteration. By raising our matrix to a power q via $(AA^*)^q A$ prior to applying it to Ψ we can enhance the decay of the spectrum of our matrix without effecting the direction of the eigenvectors. Hence, we replace in Algorithm 6.1 the step $Y = A\Psi$ by $Y = (AA^*)^q A\Psi$. A small power of q (such as $q = 1$ or 2) is often sufficient. In this modified RSVD it is also advisable to increase the oversampling factor p . We refer to [89, Corollary 10.10] for an approximation error bound for this modified RSVD as well as to [179] for a more detailed discussion.

The key ideas behind the RSVD also form the main ingredients for some other algorithms, such as randomized least squares and randomized Choleksy factorization, see [151, 89, 102].

Community Detection and the Power of Convex Relaxation

This chapter presents the idea of a convex relaxation, illustrated via problems on graphs. We start with the introduction of the Goemans-Williamson Semidefinite relaxation for the Max-Cut problem, and move on to the problem of community detection in the Stochastic Block Model.

7.1 Max-Cut, lifting, and approximation algorithms

Many data analysis tasks include in them a step consisting of solving a computational problem, oftentimes in the form of finding a hidden parameter that best explains the data, or model specifications that provide best-fits. Many such problems, including examples in previous chapters, are computationally intractable. In complexity theory this is often captured by *NP*-hardness. Unless the widely believed $P \neq NP$ conjecture is false, there is no polynomial algorithm that can solve all instances of an NP-hard problem. Thus, when faced with an NP-hard problem (such as the *Max-Cut* problem discussed below) one has three natural options: to use an exponential type algorithm that solves exactly the problem in all instances, to design polynomial time algorithms that only work for some of the instances (hopefully relevant ones), or to design polynomial algorithms that, in all instances, produce guaranteed approximate solutions. This section is about the third option, another example of this approach is the earlier discussion on Spectral Clustering and Cheeger's inequality. The second option, of designing algorithms that work in many, rather than all, instances is discussed later in this chapter, notably these goals are often achieved by the same algorithms.

The *Max-Cut* problem is defined as follows: Given a graph $G = (V, E, W)$ with non-negative weights w_{ij} on the edges, find a set $S \subset V$ for which $\text{cut}(S)$ is maximal.¹ Goemans and Williamson [80] introduced an approximation algorithm that runs in polynomial time, has a randomized component in it, and

¹The problems of maximizing and minimizing the cut are related if one considers the complement graph. As we saw in Chapter 4, if we want to cluster the nodes of a

is able to obtain a cut whose expected value is guaranteed to be no smaller than a particular constant α_{GW} times the optimum cut. The constant α_{GW} is referred to as the approximation ratio.

Let $V = \{1, \dots, n\}$. One can restate **Max-Cut** as

$$\begin{aligned} \max \quad & \frac{1}{2} \sum_{i < j} w_{ij} (1 - y_i y_j) \\ \text{s.t.} \quad & |y_i| = 1 \end{aligned} \quad (7.1)$$

The y_i 's are binary variables that indicate set membership, i.e., $y_i = 1$ if $i \in S$ and $y_i = -1$ otherwise.

Consider the following relaxation of this problem:

$$\begin{aligned} \max \quad & \frac{1}{2} \sum_{i < j} w_{ij} (1 - u_i^T u_j) \\ \text{s.t.} \quad & u_i \in \mathbb{R}^n, \|u_i\| = 1. \end{aligned} \quad (7.2)$$

This is a relaxation because if we restrict u_i to be a multiple of e_1 , the first element of the canonical basis, then (7.2) is equivalent to (7.1). For this to be a useful approach, the following two properties should hold:

- (a) Problem (7.2) is easy to solve.
- (b) The solution of (7.2) is, in some way, related to the solution of (7.1).

Definition 7.1. *Given a graph G , we define $\text{MaxCut}(G)$ as the optimal value of (7.1) and $\mathcal{R}\text{MaxCut}(G)$ as the optimal value of (7.2).*

We start with property (a). Set X to be the Gram matrix of $u_1, \dots, u_n$, that is, $X = U^T U$ where the i 'th column of U is u_i . We can rewrite the objective function of the relaxed problem as

$$\frac{1}{2} \sum_{i < j} w_{ij} (1 - X_{ij})$$

One can exploit the fact that X having a decomposition of the form $X = Y^T Y$ is equivalent to being positive semidefinite, denoted $X \succeq 0$. The set of PSD matrices is a convex set. Also, the constraint $\|u_i\| = 1$ can be expressed as $X_{ii} = 1$. This means that the relaxed problem is equivalent to the following semidefinite program (SDP):

$$\begin{aligned} \max \quad & \frac{1}{2} \sum_{i < j} w_{ij} (1 - X_{ij}) \\ \text{s.t.} \quad & X \succeq 0 \text{ and } X_{ii} = 1, i = 1, \dots, n. \end{aligned} \quad (7.3)$$

graph (by minimizing the cut) we have to add a mechanism that balances the size of clusters. The Max-Cut problem does not suffer from this issue, as the partitions that maximize the cut tend to be well balanced. In fact, this makes the Max-Cut problem conceptionally closer to clustering (and community detection) than the Min-Cut problem, making it an excellent theoretical testbed for algorithms and analysis. Furthermore, it also enjoys several direct applications.

This SDP can be solved (up to ε accuracy) in time polynomial on the input size and $\log(\varepsilon^{-1})$ [182].

There is an alternative way of viewing (7.3) as a relaxation of (7.1). By taking $X = yy^T$ one can formulate a problem equivalent to (7.1)

$$\begin{aligned} \max \quad & \frac{1}{2} \sum_{i < j} w_{ij} (1 - X_{ij}) \\ \text{s.t.} \quad & X \succeq 0, \quad X_{ii} = 1, \quad i = 1, \dots, n, \quad \text{and} \quad \text{rank}(X) = 1. \end{aligned} \quad (7.4)$$

The SDP (7.3) can be regarded as a relaxation of (7.4) obtained by removing the non-convex rank constraint. In fact, this is how we will later formulate a similar relaxation for the minimum bisection problem, in Section 7.2.

We now turn to property (b), and consider the problem of forming a solution to (7.1) from a solution to (7.3). From the solution $\{u_i\}_{i=1,\dots,n}$ of the relaxed problem (7.3), we produce a cut subset S' by randomly picking a vector $r \in \mathbb{R}^n$ from the uniform distribution on the unit sphere and setting

$$S' = \{i \mid r^T u_i \geq 0\}$$

In other words, we separate the vectors $u_1, \dots, u_n$ by a random hyperplane (perpendicular to r). We will show that the cut given by the set S' is comparable to the optimal one.

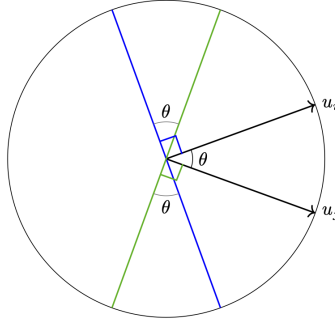


Fig. 7.1: Illustration of the relationship between the angle between vectors and their inner product, $\theta = \arccos(u_i^T u_j)$

Let W be the value of the cut produced by the procedure described above. Note that W is a random variable, whose expectation is easily seen (see Figure 7.1) to be given by

$$\begin{aligned}\mathbb{E}[W] &= \sum_{i < j} w_{ij} \Pr \{ \text{sign}(r^T u_i) \neq \text{sign}(r^T u_j) \} \\ &= \sum_{i < j} w_{ij} \frac{1}{\pi} \arccos(u_i^T u_j).\end{aligned}$$

If we define α_{GW} as

$$\alpha_{GW} = \min_{-1 \leq x \leq 1} \frac{\frac{1}{\pi} \arccos(x)}{\frac{1}{2}(1-x)},$$

it can be shown that $\alpha_{GW} > 0.87$ (see, for example [80]).

By linearity of expectation

$$\mathbb{E}[W] = \sum_{i < j} w_{ij} \frac{1}{\pi} \arccos(u_i^T u_j) \geq \alpha_{GW} \frac{1}{2} \sum_{i < j} w_{ij} (1 - u_i^T u_j). \quad (7.5)$$

Let $\text{MaxCut}(G)$ be the maximum cut of G , meaning the maximum of the original problem (7.1). Naturally, the optimal value of (7.2) is larger or equal than $\text{MaxCut}(G)$. Hence, an algorithm that solves (7.2) and uses the random rounding procedure described above produces a cut W satisfying

$$\text{MaxCut}(G) \geq \mathbb{E}[W] \geq \alpha_{GW} \frac{1}{2} \sum_{i < j} w_{ij} (1 - u_i^T u_j) \geq \alpha_{GW} \text{MaxCut}(G), \quad (7.6)$$

thus having an approximation ratio (in expectation) of α_{GW} . Note that one can run the randomized rounding procedure several times and select the best cut.² We thus have

$$\text{MaxCut}(G) \geq \mathbb{E}[W] \geq \alpha_{GW} \mathcal{R}\text{MaxCut}(G) \geq \alpha_{GW} \text{MaxCut}(G)$$

Can α_{GW} be improved?

A natural question is to ask whether there exists a polynomial time algorithm that has an approximation ratio better than α_{GW} .

The unique games problem (as depicted in Figure 7.2) is the following: Given a graph and a set of k colors, and, for each edge, a matching between the colors, the goal in the unique games problem is to color the vertices as to agree with as high of a fraction of the edge matchings as possible. For example, in Figure 7.2 the coloring agrees with $\frac{3}{4}$ of the edge constraints, and it is easy to see that one cannot do better.

The Unique Games Conjecture of Khot [99], has played a major role in hardness of approximation results.

²It is worth noting that one is only guaranteed to solve (7.2) up to an approximation of ε from its optimum value. However, since this ε can be made arbitrarily small, one can get the approximation ratio arbitrarily close to α_{GW} .

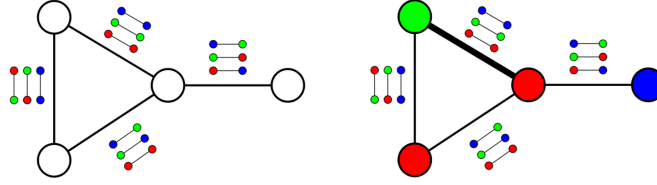


Fig. 7.2: The Unique Games Problem

Conjecture 7.2. For any $\varepsilon > 0$, the problem of distinguishing whether an instance of the Unique Games Problem is such that it is possible to agree with a $\geq 1 - \varepsilon$ fraction of the constraints or it is not possible to even agree with a ε fraction of them, is NP-hard.

There is a sub-exponential time algorithm capable of distinguishing such instances of the unique games problem [15], however no polynomial time algorithm has been found so far. At the moment one of the strongest candidates to break the Unique Games Conjecture is a relaxation based on the Sum-of-squares hierarchy that we will discuss below.

Remarkably, approximating **Max-Cut** with an approximation ratio better than α_{GW} is as hard as refuting the Unique Games Conjecture (UG-hard) [100]. More generally, if the Unique Games Conjecture is true, the semidefinite programming approach described above produces optimal approximation ratios for a large class of problems [147].

Not depending on the Unique Games Conjecture, there is a NP-hardness of approximation of $\frac{16}{17}$ for **Max-Cut** [90].

Remark 7.3. Note that a simple greedy method that assigns membership to each vertex as to maximize the number of edges cut involving vertices already assigned achieves an approximation ratio of $\frac{1}{2}$ (even of $\frac{1}{2}$ of the total number of edges, not just of the optimal cut).

7.1.1 A Sums-of-Squares interpretation

We now give a different interpretation to the approximation ratio obtained above. Let us first slightly reformulate the problem (recall that $w_{ii} = 0$).

Recall from Proposition 4.3 that a cut can be rewritten as a quadratic form involving the graph Laplacian. We can rewrite (7.1) as

$$\max_{y_i = \pm 1, i = 1, \dots, n} \frac{1}{4} y^T L_G y \quad (7.7)$$

Similarly, (7.3) can be written (by taking $X = yy^T$) as

$$\begin{aligned}
& \max \frac{1}{4} \text{Tr}(L_G X) \\
& \text{s.t. } X \succeq 0 \\
& \quad X_{ii} = 1, \quad i = 1, \dots, n.
\end{aligned} \tag{7.8}$$

In Section 7.2 we will derive the dual program to (7.8) in the context of recovery in the Stochastic Block Model. Here we will simply state it, and show weak duality as it will be important for the argument that follows.

$$\begin{aligned}
& \min \text{Tr}(D) \\
& \text{s.t. } D \text{ is a diagonal matrix} \\
& \quad D - \frac{1}{4} L_G \succeq 0.
\end{aligned} \tag{7.9}$$

Indeed, if X is a feasible solution to (7.8) and D a feasible solution to (7.9) then, since X and $D - \frac{1}{4} L_G$ are both positive semidefinite $\text{Tr}[X(D - \frac{1}{4} L_G)] \geq 0$ which gives

$$0 \leq \text{Tr}\left[X\left(D - \frac{1}{4} L_G\right)\right] = \text{Tr}(XD) - \frac{1}{4} \text{Tr}(L_G X) = \text{Tr}(D) - \frac{1}{4} \text{Tr}(L_G X),$$

since D is diagonal and $X_{ii} = 1$. This shows weak duality, the fact that the value of (7.9) is larger than the one of (7.8).

If certain conditions, the so called Slater conditions [183, 182], are satisfied then the optimal values of both programs are known to coincide, this is known as strong duality. In this case, the Slater conditions ask whether there is a matrix strictly positive definite that is feasible for (7.8), and the identity is such a matrix. This means that there exists $D^\natural$ feasible for (7.9) such that

$$\text{Tr}(D^\natural) = \mathcal{R}\text{MaxCut}.$$

Hence, for any $y \in \mathbb{R}^n$ we have

$$\frac{1}{4} y^T L_G y = \mathcal{R}\text{MaxCut} - y^T \left(D^\natural - \frac{1}{4} L_G\right)^T y + \sum_{i=1}^n D_{ii}^\natural (y_i^2 - 1). \tag{7.10}$$

Since $D^\natural - \frac{1}{4} L_G \succeq 0$, there exists V such that $D^\natural - \frac{1}{4} L_G = VV^T$ with the columns of V denoted by $v_1, \dots, v_n$, meaning that $y^T (D^\natural - \frac{1}{4} L_G)^T y = \|V^T y\|^2 = \sum_{k=1}^n (v_k^T y)^2$. Hence, for any $y \in \mathbb{R}^n$,

$$\mathcal{R}\text{MaxCut} - \frac{1}{4} y^T L_G y = \sum_{k=1}^n (v_k^T y)^2 + \sum_{i=1}^n D_{ii}^\natural (y_i^2 - 1). \tag{7.11}$$

Note that (7.11) certifies that no cut of G is larger than $\mathcal{R}\text{MaxCut}$. Indeed, if $y \in \{\pm 1\}^n$ then $y_i^2 = 1$ and so $\sum_{i=1}^n D_{ii}^\natural (y_i^2 - 1) = 0$ while $\sum_{k=1}^n (v_k^T y)^2 \geq 0$ since it is a sum of squares (of degree 2). This is known as a sum-of-squares certificate [34, 33, 140, 108, 157, 137]; by following this argument in reverse one can show that $\mathcal{R}\text{MaxCut}$ is precisely the smallest real number for which

a certificate of the type of (7.11) exists (while MaxCut is the smallest real number for which $\text{MaxCut} - \frac{1}{4}y^T L_G y$ is non-negative in the hypercube. This gap shrinks if one allows sum-of-square certificates of higher degree³.

The remarkable fact is that sum-of-squares certificates of at most a specified degree can be found using Semidefinite programming [140, 108] (in fact, the SDP (7.9) is finding the smallest real number Λ for which a certificate such as (7.11) (of degree 2) exists. On the other hand, the primal is in some sense constraining the degree 2 moments of y , $X_{ij} = y_i y_j$. Many natural questions remain open towards a precise understanding of the power of SDPs corresponding to higher degree sum-of-squares certificates. There are several very nice lecture notes, surveys, and courses on Sum-of-Squares, see for example [33, 76, 148]⁴

Remark 7.4 (triangular inequalities and Sum of squares level 4). A natural follow-up question is whether the relaxation of degree 4 is actually strictly tighter than the one of degree 2 for Max-Cut (in the sense of forcing extra constraints). What follows is an interesting set of inequalities that degree 4 enforces and that degree 2 does not, known as triangular inequalities. This example helps illustrate the differences between Sum-of-Squares certificates of different degree.

Since $y_i \in \{\pm 1\}$ we naturally have that, for all i, j, k

$$y_i y_j + y_j y_k + y_k y_i \geq -1,$$

this would mean that, for $X_{ij} = y_i y_j$ we would have,

$$X_{ij} + X_{jk} + X_{ik} \geq -1,$$

however it is not difficult to see that the SDP (7.8) of degree 2 is only able to constraint

$$X_{ij} + X_{jk} + X_{ik} \geq -\frac{3}{2},$$

which is considerably weaker. There are a few different ways of thinking about this, one is that the three vector u_i, u_j, u_k in the relaxation may be at an angle of 120 degrees with each other. Another way of thinking about this is that the inequality $y_i y_j + y_j y_k + y_k y_i \geq -\frac{3}{2}$ can be proven using sum-of-squares proof with degree 2:

$$(y_i + y_j + y_k)^2 \geq 0 \quad \Rightarrow \quad y_i y_j + y_j y_k + y_k y_i \geq -\frac{3}{2}$$

³This is related to Hilbert's 17th problem [154] and Stengle's Positivstellensatz [160]

⁴There are also very nice courses and notes on the topic, such as courses by Boaz Barak and David Steurer (<https://www.sumofsquares.org/public/index.html>), by Tselil Schramm (<https://tselilshchramm.org/sos-paradigm/winter21.html>), by Sam Hopkins (<http://www.samuelbhopskins.com/teaching/sos-fall-24/sos-fall-24.html>), Tim Kunisky (<http://www.kunisky.com/teaching/2022spring-sos/>), and Aaron Potechin (<https://canvas.uchicago.edu/courses/17604>).

However, the stronger constraint cannot.

On the other hand, if degree 4 monomials are involved, (let's say $X_S = \prod_{s \in S} y_s$, note that $X_\emptyset = 1$ and $X_{ij}X_{ik} = X_{jk}$) then the constraint

$$\begin{bmatrix} X_\emptyset \\ X_{ij} \\ X_{jk} \\ X_{ki} \end{bmatrix} \begin{bmatrix} X_\emptyset \\ X_{ij} \\ X_{jk} \\ X_{ki} \end{bmatrix}^T = \begin{bmatrix} 1 & X_{ij} & X_{jk} & X_{ki} \\ X_{ij} & 1 & X_{ik} & X_{jk} \\ X_{jk} & X_{ik} & 1 & X_{ij} \\ X_{ki} & X_{jk} & X_{ij} & 1 \end{bmatrix} \succeq 0$$

implies $X_{ij} + X_{jk} + X_{ik} \geq -1$ just by taking

$$\mathbf{1}^T \begin{bmatrix} 1 & X_{ij} & X_{jk} & X_{ki} \\ X_{ij} & 1 & X_{ik} & X_{jk} \\ X_{jk} & X_{ik} & 1 & X_{ij} \\ X_{ki} & X_{jk} & X_{ij} & 1 \end{bmatrix} \mathbf{1} \geq 0.$$

Also, note that the inequality $y_i y_j + y_j y_k + y_k y_i \geq -1$ can indeed be proven using sum-of-squares proof with degree 4 (recall that $y_i^2 = 1$):

$$(1 + y_i y_j + y_j y_k + y_k y_i)^2 \geq 0 \quad \Rightarrow \quad y_i y_j + y_j y_k + y_k y_i \geq -1.$$

Interestingly, it is known [101] that these extra inequalities alone will not increase the approximation power (in the worst case) of (7.3).

7.2 Community Detection

The problem of detecting communities in network data is a central problem in data science, examples of interest include social networks, the internet, or biological and ecological networks. In Chapter 4 we discussed clustering in the context of graphs, and described performance guarantees for spectral clustering (based on Cheeger's Inequality) that made no assumptions on the underlying graph. While these guarantees are remarkable, they are worst-case and hence pessimistic in nature. In an effort to understand the performance of some of these approaches on more realistic models of data, we will now analyze a generative model for graphs with community structure, the stochastic block model. On the methodology side, we will focus on convex relaxations, based on semidefinite programming (as in Chapter 7.1), and will show that this approach achieves exact recovery of the communities on graphs drawn from this model. The techniques developed to prove these guarantees mirror the ones used to prove analogous guarantees for a variety of other problems where convex relaxations yield exact recovery.

7.2.1 The Stochastic Block Model

The Stochastic Block Model is a random graph model that produces graphs with a community structure. While, as with any model, we do not expect it to

capture all properties of a real world network (examples include network hubs, power-law degree distributions, and other structures) the goal is to study a simple graph model that produces community structure, as a test bed for understanding fundamental limits of community detection and analyzing the performance of recovery algorithms.

Definition 7.5 (Stochastic Block Model). *Let n and k be positive integers representing respectively the number of nodes and communities, $c \in [k]^n$ be the vector of community memberships for the different nodes, and $P \in [0, 1]^{k \times k}$ a symmetric matrix of connectivity probabilities. A graph G is said to be drawn from the Stochastic Block Model on n nodes, when for each pair of nodes (i, j) the probability that $(i, j) \in E$ is independent from all other edges and given by P_{c_i, c_j} .*

We will focus on the special case of the two communities ($k = 2$) balanced symmetric block model where n is even, both communities are of the same size, and

$$P = \begin{bmatrix} p & q \\ q & p \end{bmatrix},$$

where $p, q \in [0, 1]$ are constants, cf. Figure 7.3. Furthermore, we will focus on the associative case ($p > q$), while noting that all that follows can be easily adapted to the disassociate case ($q > p$). We note also that when $p = q$ this model reduces to the classical Erdős-Renyí model described in Chapter 4. Since there are only two communities we will identify their membership labels with $+1$ and -1 .

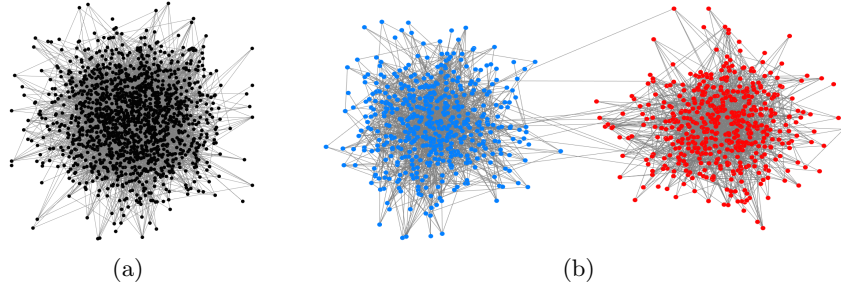


Fig. 7.3: A graph generated from the stochastic block model with 600 nodes and 2 communities, scrambled in Fig. 7.3(a), clustered and color-coded in Fig. 7.3(b). Nodes in this graph connect with probability $p = 6/600$ within communities and $q = 0.1/600$ across communities. (Image courtesy of Emmanuel Abbe.)

Many fascinating questions can be asked in the context of this model. Natural questions include to characterize statistics of the model, such as number

of triangles or larger cliques. In this chapter, motivated by the problem of community detection, we are interested in understanding when is it possible to reconstruct, or estimate, the community memberships from an observation of the graph, and what efficient algorithms succeed at this inference task.

Before proceeding we note that the difficulty of this problem should certainly depend on the value of p and q . As illustrative examples, this problem is trivial when $p = 1$ and $q = 0$ and hopeless when $p = q$ (notice that even in the easy case the actual membership can only be determined up to a re-labeling of the communities). As $p > q$, we will attempt to recover the original partition by trying to compute the minimum bisection of the graph; while related to the Max-Cut problem described in Section 7.1, notice how the objective here is to produce the minimum balanced cut.

7.2.2 Spike model prediction

A natural approach is to draw motivation from Chapter 4 and to use a form of spectral clustering to attempt to partition the graph.

Let A be the adjacency matrix of G ,

$$A_{ij} = \begin{cases} 1 & \text{if } (i, j) \in E(G) \\ 0 & \text{otherwise.} \end{cases} \quad (7.12)$$

Note that in our model, A is a random matrix. We would like to solve

$$\begin{aligned} \max \quad & \sum_{i,j} A_{ij} x_i x_j \\ \text{s.t.} \quad & x_i = \pm 1, \forall i \\ & \sum_j x_j = 0, \end{aligned} \quad (7.13)$$

The optimal solution x of (7.13) takes the value $+1$ on one side of a partition and -1 on the other side, where the partition is balanced and achieves the minimum cut between the resulting clusters.

Relaxing the condition $x_i = \pm 1, \forall i$ to $\|x\|_2^2 = n$ would yield a spectral method

$$\begin{aligned} \max \quad & \sum_{i,j} A_{ij} x_i x_j \\ \text{s.t.} \quad & \|x\|_2 = \sqrt{n} \\ & \mathbf{1}^T x = 0 \end{aligned} \quad (7.14)$$

The solution of (7.14) corresponds to the leading eigenvector of the matrix obtained by projecting A on the orthogonal complement of the all-ones vector $\mathbf{1}$.

The matrix A is a random matrix whose expectation is given⁵ by

$$\mathbb{E}[A_{ij}] = \begin{cases} p & \text{if } i \text{ and } j \text{ are in the same community} \\ q & \text{otherwise.} \end{cases}$$

Let g denote the vector corresponding to the true community memberships, with entries $+1$ and -1 ; note that this is the vector we want to recover.⁶ We can write

$$\mathbb{E}[A] = \frac{p+q}{2}\mathbf{1}\mathbf{1}^T + \frac{p-q}{2}gg^T,$$

and

$$A = (A - \mathbb{E}[A]) + \frac{p+q}{2}\mathbf{1}\mathbf{1}^T + \frac{p-q}{2}gg^T.$$

In order to remove the term $\frac{p+q}{2}\mathbf{1}\mathbf{1}^T$ we consider the random matrix

$$\mathcal{A} = A - \frac{p+q}{2}\mathbf{1}\mathbf{1}^T.$$

It is easy to see that

$$\mathcal{A} = (\mathcal{A} - \mathbb{E}[\mathcal{A}]) + \frac{p-q}{2}gg^T.$$

This means that $\mathcal{A}$ is the sum of a random matrix whose expected value is zero and a rank-1 matrix, i.e.

$$\mathcal{A} = W + \lambda vv^T$$

where $W = (\mathcal{A} - \mathbb{E}[\mathcal{A}])$ and $\lambda vv^T = \frac{p-q}{2}n \left(\frac{g}{\sqrt{n}}\right) \left(\frac{g}{\sqrt{n}}\right)^T$. In Chapter 3 we saw that for a large enough rank-1 additive perturbation to a Wigner matrix, there is an eigenvalue associated with the perturbation that pops outside of the distribution of eigenvalues of a Wigner Gaussian matrix W . Moreover, whenever this happens, we saw that the leading eigenvector has a non-trivial correlation with g .

Since $\mathcal{A}$ is simply A minus a multiple of $\mathbf{1}\mathbf{1}^T$, problem (7.14) is equivalent to

$$\begin{aligned} \max \quad & \sum_{i,j} \mathcal{A}_{ij} x_i x_j \\ \text{s.t.} \quad & \|x\|_2 = \sqrt{n} \\ & \mathbf{1}^T x = 0 \end{aligned} \tag{7.15}$$

⁵For simplicity we assume that self-loops also have probability p . This does not affect any of the conclusions, as it does not give information about the community memberships.

⁶We want to recover either g or $-g$, as they correspond to different labelings of the same community structure.

Since we have subtracted a multiple of $\mathbf{1}\mathbf{1}^T$, we will drop the constraint $\mathbf{1}^T x = 0$. Notice how a deviation from $\mathbf{1}^T x = 0$ would be penalized in the new objective, the fact that the multiple we subtracted is sufficient for us to drop the constraint will be confirmed by the success of the new optimization problem, now given by

$$\begin{aligned} \max \quad & \sum_{i,j} \mathcal{A}_{ij} x_i x_j \\ \text{s.t.} \quad & \|x\|_2 = \sqrt{n}, \end{aligned} \quad (7.16)$$

whose solution corresponds to the leading eigenvector of $\mathcal{A}$.

Recall that if $\mathcal{A} - \mathbb{E}[\mathcal{A}]$ is a Wigner matrix with i.i.d. entries with zero mean and variance σ^2 then its empirical spectral density follows the semicircle law and it is essentially supported in $[-2\sigma\sqrt{n}, 2\sigma\sqrt{n}]$. We would then expect the top eigenvector of $\mathcal{A}$ to correlate with g as long as

$$\frac{p-q}{2}n > \frac{2\sigma\sqrt{n}}{2}. \quad (7.17)$$

Unfortunately $\mathcal{A} - \mathbb{E}[\mathcal{A}]$ is not a Wigner matrix in general. In fact, half of its entries have variance $p(1-p)$ while the variance of the other half is $q(1-q)$.

Putting rigor aside for a second, if we were to take $\sigma^2 = \frac{p(1-p)+q(1-q)}{2}$ then (7.17) would suggest that the leading eigenvector of $\mathcal{A}$ correlates with the true partition vector g as long as

$$\frac{p-q}{2} > \frac{1}{\sqrt{n}} \sqrt{\frac{p(1-p)+q(1-q)}{2}}, \quad (7.18)$$

This argument is of course not valid, because the matrix in question is not a Wigner matrix. The special case $q = 1-p$ can be easily salvaged, since all entries of $\mathcal{A} - \mathbb{E}[\mathcal{A}]$ have the same variance and they can be made to be identically distributed by conjugating with $\text{diag}(g)$. This is still an impressive result, it says that if $p = 1-q$ then $p-q$ needs only to be around $\frac{1}{\sqrt{n}}$ to be able to make an estimate that correlates with the original partitioning!

An interesting regime (motivated, for example, by friendship networks in social sciences) is when the average degree of each node is constant. This can be achieved by taking $p = \frac{a}{n}$ and $q = \frac{b}{n}$ for constants a and b . While the argument presented to justify condition (7.18) is not valid in this setting, it nevertheless suggests that the condition on a and b needed to be able to make an estimate that correlates with the original partition, often referred to as partial recovery, is

$$(a-b)^2 > 2(a+b). \quad (7.19)$$

Remarkably this was posed as a conjecture by Decelle et al. [60] and proved in a series of works by Mossel et al. [133, 132] and Massoulié [124]. While

describing the proof of this conjecture is outside the scope of this book, we note that the conjectures were obtained by studying fixed points of a certain linearization of belief propagation using techniques from statistical physics. The lower bound can be proven by showing contiguity between the two models below the phase transition, and the upper bound is obtained by analyzing an algorithm that is an adaptation of belief propagation and studying the so-called non-backtracking operator. We refer the reader to the excellent survey of Abbe [4] and references therein for further reading.

Remark 7.6 (More than three communities). The balanced symmetric stochastic block model with $k > 3$ communities is conjectured to have a fascinating statistical-to-computational gap. In the sparse regime of inner probability $p = \frac{a}{n}$ and outer probability $q = \frac{b}{n}$ it is believed that, for $k > 3$ there is a regime of the parameters a and b such that the problem of partially recovering the community memberships is statistically, or information-theoretically, possible but that there does not exist a polynomial-time algorithm to perform this task. These conjectures are based on insight obtained with tools from statistical physics. We refer the reader to [60, 192, 79, 3] for further reading.

7.2.3 Exact recovery

We now turn our attention to the problem of recovering the cluster membership of every single node correctly, not simply having an estimate that correlates with the true labels. We will keep our focus on the balanced, symmetric, two communities setting. This will serve as an illustration for two phenomena: (i) the fact that convex relaxations often find optimal solutions, not just approximations, when the input are “typical instances”; (ii) convex relaxations such as semidefinite programs over random instances can often be readily understood via matrix concentration inequalities (such as the ones in Chapter 9). If the inner-probability is $p = \frac{a}{n}$ then it is not hard to show that each cluster will (with high probability) have isolated nodes, making it impossible to recover the membership of every possible node correctly. In fact this is the case whenever $p \leq \frac{(2-\varepsilon) \log n}{n}$, for some $\varepsilon > 0$. For that reason we focus on the regime

$$p = \frac{\alpha \log(n)}{n} \text{ and } q = \frac{\beta \log(n)}{n}, \quad (7.20)$$

for some constants $\alpha > \beta$.

A natural algorithm would be to compute the minimum bisection (7.13) which corresponds to the Maximum Likelihood Estimator, and also the Maximum a Posteriori Estimator when the community memberships are drawn uniformly at random. In fact, it is known (see [1] for a proof) that if

$$\sqrt{\alpha} - \sqrt{\beta} > \sqrt{2}, \quad (7.21)$$

then, with high probability, (7.13) recovers the true partition. Moreover, if

$$\sqrt{\alpha} - \sqrt{\beta} < \sqrt{2},$$

no algorithm can, with high probability, recover the true partition.

In this section we will analyze a semidefinite programming relaxation, analogous to the ones described in Section 7.1 for Max-Cut. By making use of convex duality, we will derive conditions for exact recovery with this particular algorithm, reducing the problem to a problem in random matrix theory. We will present a solution to the resulting random matrix question, using the matrix concentration tools developed in Chapter 9. While not providing the strongest known guarantee, this approach is extremely adaptable and can be used to solve a large number of similar questions.⁷

7.2.4 A semidefinite relaxation

Let $x \in \mathbb{R}^n$ with $x_i = \pm 1$ represent a partition of the nodes (recall that there is an ambiguity in the sense that x and $-x$ represent the same partition). Note that if we remove the constraint $\sum_j x_j = 0$ in (7.13) then the optimal solution becomes $x = \mathbf{1}$. Let us define $B = 2A - (\mathbf{1}\mathbf{1}^T - I)$, meaning that

$$B_{ij} = \begin{cases} 0 & \text{if } i = j \\ 1 & \text{if } (i, j) \in E(G) \\ -1 & \text{otherwise} \end{cases} \quad (7.22)$$

It is clear that the problem

$$\begin{aligned} \max \quad & \sum_{i,j} B_{ij} x_i x_j \\ \text{s.t. } & x_i = \pm 1, \forall_i \\ & \sum_j x_j = 0 \end{aligned} \quad (7.23)$$

has the same solution as (7.13). However, when the constraint is dropped,

$$\begin{aligned} \max \quad & \sum_{i,j} B_{ij} x_i x_j \\ \text{s.t. } & x_i = \pm 1, \forall_i, \end{aligned} \quad (7.24)$$

$x = \mathbf{1}$ is no longer an optimal solution. As with (7.16) above, the penalty created by subtracting a large multiple of $\mathbf{1}\mathbf{1}^T$ will be enough to discourage

⁷A tight guarantee can be obtained by adapting the random matrix theory part of the argument, the convex geometry part can remain the same; this is briefly discussed below.

unbalanced partitions (the reader may notice the connection with Lagrangian duality). In fact, (7.24) is the problem we will set ourselves to solve.

Unfortunately (7.24) is in general NP-hard (one can encode, for example, *Max-Cut* by picking an appropriate B). We will relax it to an easier problem by the same technique used to approximate the Max-Cut problem in the previous section (this technique is often known as *matrix lifting*). If we write $X = xx^T$ then we can formulate the objective of (7.24) as

$$\sum_{i,j} B_{ij} x_i x_j = x^T B x = \text{Tr}(x^T B x) = \text{Tr}(B x x^T) = \text{Tr}(B X)$$

Also, the condition $x_i = \pm 1$ implies $X_{ii} = x_i^2 = 1$. This means that (7.24) is equivalent to

$$\begin{aligned} \max \quad & \text{Tr}(B X) \\ \text{s.t.} \quad & X_{ii} = 1, \forall_i \\ & X = x x^T \text{ for some } x \in \mathbb{R}^n. \end{aligned} \tag{7.25}$$

The fact that $X = x x^T$ for some $x \in \mathbb{R}^n$ is equivalent to $\text{rank}(X) = 1$ and $X \succeq 0$. This means that (7.24) is equivalent to

$$\begin{aligned} \max \quad & \text{Tr}(B X) \\ \text{s.t.} \quad & X_{ii} = 1, \forall_i \\ & X \succeq 0 \\ & \text{rank}(X) = 1. \end{aligned} \tag{7.26}$$

We now relax the problem by removing the non-convex rank constraint

$$\begin{aligned} \max \quad & \text{Tr}(B X) \\ \text{s.t.} \quad & X_{ii} = 1, \forall_i \\ & X \succeq 0. \end{aligned} \tag{7.27}$$

This is an SDP that can be solved (up to arbitrary precision) in polynomial time [182] (note that it is the same SDP as (7.8), with a different coefficient matrix).

Since we removed the rank constraint, the solution to (7.27) is no longer guaranteed to be rank-1. We will take a different approach from the one we used in Section 7.1 to obtain an approximation ratio for **Max-Cut**, which was a worst-case approximation ratio guarantee. In this section we will show that, for some values of α and β , with high probability, the solution to (7.27) not only satisfies the rank constraint but it coincides with $X = g g^T$ where g corresponds to the true partition. From X one can compute g by simply calculating its leading eigenvector. Besides being a good approximation algorithm, SDP relaxations such as (7.8) provide optimal solutions in many instances (but not always).

7.2.5 Convex duality

A standard technique to show that a candidate solution is the optimal one for a convex problem is to use convex duality.

We will describe duality with a game theoretical intuition in mind. The idea will be to rewrite (7.27) without imposing constraints on X but rather have the constraints be implicitly enforced. Consider the following optimization problem.

$$\max_X \min_{\substack{Z, Q \\ Z \text{ is diagonal} \\ Q \succeq 0}} \text{Tr}(BX) + \text{Tr}(QX) + \text{Tr}(Z(I_{n \times n} - X)). \quad (7.28)$$

Let us provide a game theoretical interpretation for (7.28). Suppose that there is a *primal player* (picking X) whose objective is to maximize the objective and a *dual player* (picking Z and Q after seeing X) trying to make the objective as small as possible. If the primal player does not pick X satisfying the constraints of (7.27) then we claim that the dual player is capable of driving the objective to $-\infty$. Indeed, if there is an i for which $X_{ii} \neq 1$ then the dual player can simply pick $Z_{ii} = -c \frac{1}{1-X_{ii}}$ and make the objective as small as desired by taking a large enough c . Similarly, if X is not positive semidefinite, then the dual player can take $Q = cvv^T$ where v is such that $v^T X v < 0$. If, on the other hand, X satisfies the constraints of (7.27) then

$$\text{Tr}(BX) \leq \min_{\substack{Z, Q \\ Z \text{ is diagonal} \\ Q \succeq 0}} \text{Tr}(BX) + \text{Tr}(QX) + \text{Tr}(Z(I_{n \times n} - X)).$$

Since equality can be achieved if for example the dual player picks $Q = 0_{n \times n}$, then it is evident that the values of (7.27) and (7.28) are the same:

$$\max_{\substack{X, \\ X_{ii} \forall_i \\ X \succeq 0}} \text{Tr}(BX) = \max_X \min_{\substack{Z, Q \\ Z \text{ is diagonal} \\ Q \succeq 0}} \text{Tr}(BX) + \text{Tr}(QX) + \text{Tr}(Z(I_{n \times n} - X))$$

With this game theoretical intuition in mind, it is clear that if we change the “rules of the game” and have the dual player decide their variables before the primal player (meaning that the primal player can pick X knowing the values of Z and Q) then it is clear that the objective can only increase, which means that:

$$\max_{\substack{X, \\ X_{ii} \forall_i \\ X \succeq 0}} \text{Tr}(BX) \leq \min_{\substack{Z, Q \\ Z \text{ is diagonal} \\ Q \succeq 0}} \max_X \text{Tr}(BX) + \text{Tr}(QX) + \text{Tr}(Z(I_{n \times n} - X)).$$

Note that we can rewrite

$$\text{Tr}(BX) + \text{Tr}(QX) + \text{Tr}(Z(I_{n \times n} - X)) = \text{Tr}((B + Q - Z)X) + \text{Tr}(Z).$$

When playing:

$$\min_{\substack{Z, Q \\ Z \text{ is diagonal} \\ Q \succeq 0}} \max_X \text{Tr}((B + Q - Z)X) + \text{Tr}(Z),$$

if the dual player does not set $B + Q - Z = 0_{n \times n}$ then the primal player can drive the objective value to $+\infty$, this means that the dual player is forced to chose $Q = Z - B$ and so we can write

$$\min_{\substack{Z, Q \\ Z \text{ is diagonal} \\ Q \succeq 0}} \max_X \text{Tr}((B + Q - Z)X) + \text{Tr}(Z) = \min_{\substack{Z \\ Z \text{ is diagonal} \\ Z - B \succeq 0}} \max_X \text{Tr}(Z),$$

which clearly does not depend on the choices of the primal player. This means that

$$\max_{\substack{X \\ X_{ii} \leq 1 \\ X \succeq 0}} \text{Tr}(BX) \leq \min_{\substack{Z \\ Z \text{ is diagonal} \\ Z - B \succeq 0}} \text{Tr}(Z).$$

This is known as weak duality (strong duality says that, under some conditions the two optimal values actually match, see for example [182], recall that we used strong duality when giving a sum-of-squares interpretation to the Max-Cut approximation ratio, a similar interpretation can be given in this problem, see [21]).

Also, the problem

$$\begin{aligned} \min \quad & \text{Tr}(Z) \\ \text{s.t.} \quad & Z \text{ is diagonal} \\ & Z - B \succeq 0 \end{aligned} \tag{7.29}$$

is called the dual problem of (7.27).

The derivation above explains why the objective value of the dual problem is always greater or equal to the primal problem. Nevertheless, there is a much simpler proof (although not as enlightening): let X, Z be a feasible point of (7.27) and (7.29), respectively. Since Z is diagonal and $X_{ii} = 1$, it follows that $\text{Tr}(ZX) = \text{Tr}(Z)$. Also, $Z - B \succeq 0$ and $X \succeq 0$, therefore $\text{Tr}[(Z - B)X] \geq 0$. Altogether,

$$\text{Tr}(Z) - \text{Tr}(BX) = \text{Tr}[(Z - B)X] \geq 0,$$

as stated.

Recall that we want to show that gg^T is the optimal solution of (7.27). Then, if we find Z diagonal, such that $Z - B \succeq 0$ and

$$\text{Tr}[(Z - B)gg^T] = 0, \quad (\text{this condition is known as complementary slackness})$$

then $X = gg^T$ must be an optimal solution of (7.27). To ensure that gg^T is the unique solution we just have to ensure that the nullspace of $Z - B$

only has dimension 1 (which corresponds to multiples of g). Essentially, if this is the case, then for any other possible solution X one could not satisfy complementary slackness.

This means that if we can find Z with the following properties:

- (1) Z is diagonal
- (2) $\text{Tr}[(Z - B)gg^T] = 0$
- (3) $Z - B \succeq 0$
- (4) $\lambda_2(Z - B) > 0$,

then gg^T is the unique optimum of (7.27) and so recovery of the true partition is possible (with an efficient algorithm). Z is known as the dual certificate, or dual witness.

7.2.6 Building the dual certificate

The idea to build Z is to construct it to satisfy properties (1) and (2) and try to show that it satisfies (3) and (4) using concentration. In fact, since Z is diagonal this design problem has n free variables. If $Z - B \succeq 0$ then condition (2) is equivalent to $(Z - B)g = 0$ which provides n equations, as the resulting linear system is non-singular, the candidate arising from using conditions (1) and (2) will be unique.

A couple of preliminary definitions will be useful before writing out the candidate Z . Recall that the degree matrix D of a graph G is a diagonal matrix where each diagonal coefficient D_{ii} corresponds to the number of neighbors of vertex i and that $\lambda_2(M)$ is the second smallest eigenvalue of a symmetric matrix M .

Definition 7.7. Let $\mathcal{G}_+$ (resp. $\mathcal{G}_-$) be the subgraph of G that includes the edges that link two nodes in the same community (resp. in different communities) and A the adjacency matrix of G . We denote by $D_{\mathcal{G}}^+$ (resp. $D_{\mathcal{G}}^-$) the degree matrix of $\mathcal{G}_+$ (resp. $\mathcal{G}_-$) and define the Stochastic Block Model Laplacian to be

$$L_{SBM} = D_{\mathcal{G}}^+ - D_{\mathcal{G}}^- - A.$$

Note that the inclusion of self loops does not change L_{SBM} . Also, we point out that L_{SBM} is not in general positive-semidefinite.

Now we are ready to construct the candidate Z . Condition (2) implies that we need $Z_{ii} = \frac{1}{g_i} B[i, :]g$. Since $B = 2A - (\mathbf{1}\mathbf{1}^T - I)$ we have

$$Z_{ii} = \frac{1}{g_i} (2A - (\mathbf{1}\mathbf{1}^T - I))[i, :]g = 2\frac{1}{g_i} (Ag)_i + 1,$$

meaning that

$$Z = 2(D_{\mathcal{G}}^+ - D_{\mathcal{G}}^-) + I.$$

This is our candidate dual witness. As a result

$$Z - B = 2(D_G^+ - D_G^-) + I - [2A - (\mathbf{1}\mathbf{1}^T - I)] = 2L_{SBM} + \mathbf{1}\mathbf{1}^T.$$

It trivially follows (by construction) that

$$(Z - B)g = 0.$$

This immediately gives the following lemma.

Lemma 7.8. *Let L_{SBM} denote the Stochastic Block Model Laplacian as defined in Definition 7.7. If*

$$\lambda_2(2L_{SBM} + \mathbf{1}\mathbf{1}^T) > 0, \quad (7.30)$$

then the relaxation (7.27) recovers the true partition.

Note that $2L_{SBM} + \mathbf{1}\mathbf{1}^T$ is a random matrix and so this reduces to “an exercise” in random matrix theory.

7.2.7 Matrix concentration and Matrix Bernstein Inequality

We will dedicate Chapter 9 to random matrix theory and matrix concentration inequalities that aim to control the largest eigenvalue or spectral norm of random matrices. In this section we present a general use concentration inequality for sums of independent random matrices, while noting that, as with scalars, many random variables can be written as sums of independent random variables even when it’s not trivially apparent (and we will see below that this is the case for the matrix appearing in Lemma 7.8).

Let us recall Bernstein’s inequality (Theorem 2.18) copied here with slightly different notation, and with only one of the tails: If $X_1, X_2, \dots, X_n$ are independent centered random variables satisfying $|X_i| \leq r$ and $\mathbb{E}[X_i^2] = \frac{1}{n}\nu^2$. Then,

$$\mathbb{P}\left\{\sum_{i=1}^n X_i > t\right\} \leq \exp\left(-\frac{t^2}{2\nu^2 + \frac{2}{3}rt}\right). \quad (7.31)$$

A very useful generalization of this inequality exists for bounding the largest eigenvalue of the sum of independent random matrices

Theorem 7.9 (Theorem 1.4 in [175]). *Let $\{X_k\}_{k=1}^n$ be a sequence of independent random symmetric $d \times d$ matrices. Assume that each X_k satisfies:*

$$\mathbb{E}X_k = 0 \text{ and } \lambda_{\max}(X_k) \leq R \text{ almost surely.}$$

Then, for all $t \geq 0$,

$$\mathbb{P}\left\{\lambda_{\max}\left(\sum_{k=1}^n X_k\right) \geq t\right\} \leq d \cdot \exp\left(\frac{-t^2}{2\sigma^2 + \frac{2}{3}Rt}\right) \text{ where } \sigma^2 = \left\|\sum_{k=1}^n \mathbb{E}(X_k^2)\right\|.$$

Note that $\|A\|$ denotes the spectral norm of A . Comparing with (7.31) the attentive reader will notice an extra dimensional factor of d ; a simple change of variables shows that this corresponds to a poly-logarithmic factor on the random variable, a factor that will be discussed later in this chapter where we will also include an improved inequality (Theorem 9.17).

In Chapter 9 we will discuss matrix concentration inequalities in detail and prove inequalities similar to Theorem 7.9 (for the proof of the precise bound in Theorem 7.9, which is based on a different toolkit than what we will explore in Chapter 9, we point the reader to [175]). For now we will use this inequality to show that the SDP approach presented above achieves exact recovery of communities in the stochastic block model.

7.2.8 Using matrix concentration

In this section we show how the resulting question amounts to controlling the largest eigenvalue of a random matrix, which can be done with the matrix concentration tools described above.

Let us start by noting that

$$\mathbb{E}[2L_{\text{SBM}} + 11^T] = 2\mathbb{E}L_{\text{SBM}} + 11^T = 2\mathbb{E}D_{\mathcal{G}}^+ - 2\mathbb{E}D_{\mathcal{G}}^- - 2\mathbb{E}A + 11^T,$$

and $\mathbb{E}D_{\mathcal{G}}^+ = \frac{n}{2} \frac{\alpha \log(n)}{n} I$, $\mathbb{E}D_{\mathcal{G}}^- = \frac{n}{2} \frac{\beta \log(n)}{n} I$. Moreover, $\mathbb{E}A$ is a matrix with four $\frac{n}{2} \times \frac{n}{2}$ blocks where the diagonal blocks have entries $\frac{\alpha \log(n)}{n}$ and the off-diagonal blocks have entries $\frac{\beta \log(n)}{n}$.⁸ In other words

$$\mathbb{E}A = \frac{1}{2} \left(\frac{\alpha \log(n)}{n} + \frac{\beta \log(n)}{n} \right) 11^T + \frac{1}{2} \left(\frac{\alpha \log(n)}{n} - \frac{\beta \log(n)}{n} \right) gg^T.$$

This means that

$$\mathbb{E}[2L_{\text{SBM}} + 11^T] = ((\alpha - \beta) \log n) I + \left(1 - (\alpha + \beta) \frac{\log n}{n} \right) 11^T - (\alpha - \beta) \frac{\log n}{n} gg^T.$$

Since $L_{\text{SBM}}g = 0$ and $11^Tg = 0$, we can safely ignore what happens in the span of g , and it is not hard to see that

$$\lambda_2(\mathbb{E}[2L_{\text{SBM}} + 11^T]) = (\alpha - \beta) \log n.$$

Thus, it is enough to show that

$$\|L_{\text{SBM}} - \mathbb{E}[L_{\text{SBM}}]\| < \frac{\alpha - \beta}{2} \log n, \quad (7.32)$$

which is a large deviation inequality; recall that $\|\cdot\|$ denotes operator norm.

⁸For simplicity we assume the possibility of self-loops; notice however that the matrix in question does not depend on this, only its decomposition in the degree matrices and A .

The idea is to write $L_{\text{SBM}} - \mathbb{E}[L_{\text{SBM}}]$ as a sum of independent random matrices and use the Matrix Bernstein Inequality (Theorem 7.9). This gives an illustrative example of the applicability of matrix concentration tools, as many random matrices of interest can be rewritten as sums of independent matrices.

Let us start by defining, for i and j in the same community (i.e. $g_i = g_j$),

$$\gamma_{ij}^+ = \begin{cases} 1 & \text{if } (i, j) \in E \\ 0 & \text{otherwise,} \end{cases}$$

and

$$\Delta_{ij}^+ = (e_i - e_j)(e_i - e_j)^T,$$

where e_i (resp. e_j) is the vector of all zeros except the i^{th} (resp. j^{th}) coefficient which is 1.

For i and j in different communities (i.e. $g_i \neq g_j$), define

$$\gamma_{ij}^- = \begin{cases} 1 & \text{if } (i, j) \in E \\ 0 & \text{otherwise,} \end{cases}$$

and

$$\Delta_{ij}^- = -(e_i + e_j)(e_i + e_j)^T.$$

We have

$$L_{\text{SBM}} = \sum_{i < j: g_i = g_j} \gamma_{ij}^+ \Delta_{ij}^+ + \sum_{i < j: g_i \neq g_j} \gamma_{ij}^- \Delta_{ij}^-.$$

We note how $(\gamma_{ij}^+)_{i,j}$ and $(\gamma_{ij}^-)_{i,j}$ are jointly independent random variables with expectations $\mathbb{E}(\gamma_{ij}^+) = \frac{\alpha \log n}{n}$ and $\mathbb{E}(\gamma_{ij}^-) = \frac{\beta \log n}{n}$. Δ_{ij}^+ and Δ_{ij}^- are deterministic matrices. This means that

$$L_{\text{SBM}} - \mathbb{E}L_{\text{SBM}} = \sum_{\substack{i < j: \\ g_i = g_j}} \left(\gamma_{ij}^+ - \frac{\alpha \log n}{n} \right) \Delta_{ij}^+ + \sum_{i < j: g_i \neq g_j} \left(\gamma_{ij}^- - \frac{\beta \log n}{n} \right) \Delta_{ij}^-.$$

We can then use Theorem 7.9 by setting

$$\sigma^2 = \left\| \text{Var}[\gamma^+] \sum_{i < j: g_i = g_j} (\Delta_{ij}^+)^2 + \text{Var}[\gamma^-] \sum_{i < j: g_i \neq g_j} (\Delta_{ij}^-)^2 \right\|, \quad (7.33)$$

and $R = 2$, since $\|\Delta_{ij}^+\| = \|\Delta_{ij}^-\| = 2$ and both $(\gamma_{ij}^+)_{i,j}$ and $(\gamma_{ij}^-)_{i,j}$ take values in $[-1, 1]$. Note how this bound is for the spectral norm of the summands, not just the largest eigenvalue, as our goal is to bound the spectral norm of the random matrix. In order to compute σ^2 , we write

$$\sum_{i < j: g_i = g_j} (\Delta_{ij}^+)^2 = nI - (\mathbf{1}\mathbf{1}^T + gg^T),$$

and

$$\sum_{i < j: g_i \neq g_j} (\Delta_{ij}^-)^2 = nI + (\mathbf{1}\mathbf{1}^T - gg^T)$$

Since $\text{Var}[\gamma^+] \leq \frac{\alpha \log n}{n}$, $\text{Var}[\gamma^-] \leq \frac{\beta \log n}{n}$, and all the summands are positive semidefinite we have

$$\sigma^2 \leq \left\| \frac{(\alpha + \beta) \log n}{n} (nI - gg^T) - \frac{(\alpha - \beta) \log n}{n} \mathbf{1}\mathbf{1}^T \right\| = (\alpha + \beta) \log n.$$

Using Theorem 7.9 for $t = \frac{\alpha - \beta}{2} \log n$ on both the largest and smallest eigenvalue gives

$$\begin{aligned} \mathbb{P} \left\{ \|L_{\text{SBM}} - \mathbb{E}[L_{\text{SBM}}]\| \geq \frac{\alpha - \beta}{2} \log n \right\} &\leq \\ &\leq 2n \cdot \exp \left(\frac{-\left(\frac{\alpha - \beta}{2} \log n\right)^2}{2(\alpha + \beta) \log n + \frac{4}{3} \left(\frac{\alpha - \beta}{2} \log n\right)} \right) \\ &= 2 \cdot \exp \left(-\frac{(\alpha - \beta)^2 \log n}{8(\alpha + \beta) + \frac{8}{3}(\alpha - \beta)} + \log n \right) \\ &= 2n^{-\left(\frac{(\alpha - \beta)^2}{8(\alpha + \beta) + \frac{8}{3}(\alpha - \beta)} - 1\right)}. \end{aligned}$$

Together with Lemma 7.8, this implies that as long as

$$(\alpha - \beta)^2 > 8(\alpha + \beta) + \frac{8}{3}(\alpha - \beta), \quad (7.34)$$

the semidefinite programming relaxation (7.27) recovers the true partition, with probability tending to 1 as n increases.

While it is possible to obtain a stronger guarantee for this relaxation, the derivation above illustrates the matrix concentration technique in a simple, yet powerful, instance. In fact, the analysis in [1] uses the same technique. In order to obtain a sharp guarantee (Theorem 7.10 below) one needs more specialized tools. We refer the interested reader to [20] or [87] for a discussion and proof of Theorem 7.10; the main idea is to separate the diagonal from the non-diagonal part of $L_{\text{SBM}} - \mathbb{E}[L_{\text{SBM}}]$, treating the former with scalar concentration inequalities, and the latter with specialized matrix concentration inequalities such as the ones in [28].

Theorem 7.10. *Let G be a random graph with n nodes drawn according to the stochastic block model on two communities with edge probabilities p and q . Let $p = \frac{\alpha \log n}{n}$ and $q = \frac{\beta \log n}{n}$, where $\alpha > \beta$ are constants. Then, as long as*

$$\sqrt{\alpha} - \sqrt{\beta} > \sqrt{2}, \quad (7.35)$$

the semidefinite program considered above coincides with the true partition with high probability.

Note that, if

$$\sqrt{\alpha} - \sqrt{\beta} < \sqrt{2}, \quad (7.36)$$

then exact recovery of the communities is impossible, meaning that the SDP algorithm is optimal. Furthermore, in this regime (7.36), one can show that there will be a node on each community that is more connected to the other community than to its own, meaning that a partition that swaps them would have more likelihood. The fact that the SDP will start working essentially when this starts happening appears naturally in the analysis in [20]. Later, it was proven that the spectral method (7.14), followed by a simple thresholding step, also gives exact recovery of the communities [2]. An analogous analysis has recently been obtained for the (normalized or unnormalized) graph Laplacian in place of the adjacency matrix, see [62]. However, the proof techniques for the graph Laplacian are different and a bit more involved, since—unlike the adjacency matrix—the graph Laplacian does not exhibit row/column-wise independence.

Remark 7.11. An important advantage of semidefinite relaxations is that they are often able to produce certificates of optimality. Indeed, if the solution of the relaxation (7.27) is rank 1 then the user is sure that it must be the solution of (7.24). These advantages, and ways of producing such certificates while bypassing the need to solve the semidefinite program are discussed in [21]. Semidefinite relaxation has also been effectively used to show rigorously that spectral clustering can indeed outperform clustering via k -means [116, 40, 63]

Concentration of Measure and Gaussian Analysis

In this chapter we significantly expand on the concepts presented in Chapter 2, showcasing several other instances of the *Concentration of Measure* phenomena, and focusing on Gaussian Analysis to develop a toolkit that is used throughout the book, including the celebrated Gordon's Escape through a Mesh Theorem. Our treatment draws inspiration from the excellent texts [180, 181].

Notation: In this chapter, we try to reserve g for an isotropic Gaussian vector and use X and Y to denote Gaussian vectors with other covariances.

8.1 Gaussian integration by parts and Wick's formula

We start by covering some basics in Gaussian analysis, starting with a very useful fact about Gaussian random variables: Gaussian integration by parts.

Lemma 8.1 (Gaussian integration by parts). *Let g be an n dimensional vector of iid $\mathcal{N}(0, 1)$ random variables ($g \sim \mathcal{N}(0, I_{n \times n})$). Let $\psi \in \mathbb{R}^n \rightarrow \mathbb{R}$ be any differentiable function with at most exponential growth,¹ then for any $i \in [n]$,*

$$\mathbb{E}[g_i \psi(g)] = \mathbb{E} \left[\frac{\partial \psi}{\partial x_i}(g) \right].$$

Proof. For φ a univariate function, integration by parts gives

¹We will not make an effort to describe the minimal regularity assumptions as to not deviate from the main goal of this book. When we say differentiable with at most exponential growth we mean that the derivatives also are of at most exponential growth. The interested reader will realize that this is tightly connected to Sobolev spaces on the Gaussian measure.

$$\begin{aligned}
\mathbb{E}[g_i \varphi(g_i)] &= \int_{-\infty}^{\infty} x \varphi(x) \frac{\exp\left(-\frac{x^2}{2}\right)}{\sqrt{2\pi}} dx \\
&= - \int_{-\infty}^{\infty} \varphi(x) \frac{d}{dx} \left(\frac{\exp\left(-\frac{x^2}{2}\right)}{\sqrt{2\pi}} \right) dx \\
&= \int_{-\infty}^{\infty} \frac{d}{dx} \varphi(x) \frac{\exp\left(-\frac{x^2}{2}\right)}{\sqrt{2\pi}} dx = \mathbb{E}[\varphi'(g_i)],
\end{aligned}$$

as long as φ is differential and has at most exponential growth. The proof follows by taking $\varphi(x) = F(g_1, \dots, g_{i-1}, x, g_{i+1}, \dots, g_n)$ and conditioning on $g_1, \dots, g_{i-1}, g_{i+1}, \dots, g_n$. $\square$

It is straightforward to extend this formula to non-isotropic Gaussian vectors, and we record this here as a corollary.

Corollary 8.2. *Let X be a centered n -dimensional gaussian vector with covariance Σ , i.e. $X \sim \mathcal{N}(0, \Sigma)$. Let $\varphi \in \mathbb{R}^n \rightarrow \mathbb{R}$ be any differentiable function with at most exponential growth, then for any $j \in [n]$,*

$$\mathbb{E}[X_j \varphi(X)] = \sum_{k=1}^n \Sigma_{jk} \mathbb{E} \left[\frac{\partial \varphi}{\partial x_k}(X) \right].$$

Proof. We can write $X = \Sigma^{\frac{1}{2}} g$ where g has iid $\mathcal{N}(0, 1)$ entries (note that $\Sigma \succeq 0$). Defining $\psi(g) = \varphi(\Sigma^{\frac{1}{2}} g)$ and using Lemma 8.1 yields

$$\begin{aligned}
\mathbb{E}[X_j \varphi(X)] &= \mathbb{E} \left[\left(\Sigma^{\frac{1}{2}} g \right)_j \varphi(\Sigma^{\frac{1}{2}} g) \right] = \mathbb{E} \left[\left(\Sigma^{\frac{1}{2}} g \right)_j \psi(g) \right] \\
&= \mathbb{E} \left[\left(\sum_{i=1}^n \Sigma_{ji}^{\frac{1}{2}} g_i \right) \psi(g) \right] = \sum_{i=1}^n \Sigma_{ji}^{\frac{1}{2}} \mathbb{E} [g_i \psi(g)] \\
&= \sum_{i=1}^n \Sigma_{ji}^{\frac{1}{2}} \mathbb{E} \left[\frac{\partial \psi}{\partial x_i}(g) \right] = \sum_{i=1}^n \Sigma_{ji}^{\frac{1}{2}} \mathbb{E} \left[\sum_{k=1}^n \frac{\partial \varphi}{\partial x_k}(X) \Sigma_{ki}^{\frac{1}{2}} \right] \\
&= \sum_{k=1}^n \left(\sum_{i=1}^n \Sigma_{ji}^{\frac{1}{2}} \Sigma_{ki}^{\frac{1}{2}} \right) \mathbb{E} \left[\frac{\partial \varphi}{\partial x_k}(X) \right] = \sum_{k=1}^n \Sigma_{jk} \mathbb{E} \left[\frac{\partial \varphi}{\partial x_k}(X) \right],
\end{aligned}$$

where in the first equality in the third line we used Lemma 8.1. The last equality follows from matrix multiplication and the fact that $\Sigma^{\frac{1}{2}}$ is a symmetric matrix. $\square$

Lemma 8.3 (Gaussian moments). *Let g be a standard univariate Gaussian. We have*

$$\mathbb{E} g^{2p} = (2p - 1)!!,$$

where $(2p-1)!! = (2p-1)(2p-3)\cdots 3\cdot 1$. Also,

$$[\mathbb{E}g^{2p}]^{\frac{1}{2p}} \leq \sqrt{2p}.$$

Proof. This is a straightforward application of Gaussian integration by parts (Lemma 8.1). Let $\varphi(g) = g^{2p-1}$ then $\mathbb{E}[g \cdot g^{2p-1}] = \mathbb{E}[(2p-1)g^{2p-2}]$. Iterating and using $\mathbb{E}g^2 = 1$ gives $\mathbb{E}g^{2p} = (2p-1)!!$. While it is easy to see that $(2p-1)!! \leq (2p)^p$, which proves the intended inequality, it is instructive to see a different, direct, proof:

Since $\mathbb{E}[g^{2p}] = \mathbb{E}[(2p-1)g^{2p-2}]$ and, by Jensen, $\mathbb{E}[g^{2p-2}] \leq (\mathbb{E}[g^{2p}])^{\frac{2p-2}{2p}}$, we have

$$\mathbb{E}[g^{2p}] \leq (2p-1) (\mathbb{E}[g^{2p}])^{\frac{2p-2}{2p}},$$

which can be rewritten as $(\mathbb{E}[g^{2p}])^{\frac{1}{p}} \leq (2p-1)$. $\square$

We will be interested in computing trace moments of Gaussian matrices, which will involve computing expected values of polynomials of Gaussians. It is clear that the expected value of any polynomial of i.i.d. standard Gaussians can be computed using Gaussian integration by parts for each monomial (just as in Lemma 8.3). Wick's formula is a very useful way of organizing this calculation. Before stating Wick's formula, we need to define the notion of pair partition.

Definition 8.4 (Pair Partition). *Given k a positive integer, we define $\mathbb{P}_2[k]$ as the set of partitions of $[k]$ into subsets of size 2 each. If k is odd then $\mathbb{P}_2[k]$ is empty. Given a function u on $[k]$ and a pair partition $\nu \in \mathbb{P}_2[k]$ we say that u is compatible with ν , and write $u \sim \nu$ if for all sets $\{i, j\} \in \nu$ we have $u(i) = u(j)$.*

Lemma 8.5 (Wick's formula). *Let $g_1, \dots, g_n$ be iid standard gaussian random variables and let $u : [p] \rightarrow [n]$ then*

$$\mathbb{E}[g_{u(1)} \cdots g_{u(p)}] = \sum_{\nu \in \mathbb{P}_2[p]} 1_{u \sim \nu},$$

where $\mathbb{P}_2[p]$ denotes a set of pair partitions and $u \sim \nu$ means that the function u is compatible with ν (see Definition 8.4)

Proof. Since both sides are zero when p is odd we can focus on the case when p is even. We can then prove the identity by induction. It is trivially true for $p = 2$, let $p \geq 4$ an even number. Using Gaussian integration by parts

$$\begin{aligned}
\mathbb{E}[g_{u(1)} \cdot g_{u(2)} \cdots g_{u(p)}] &= \mathbb{E} \left[\frac{\partial}{\partial g_{u(1)}} g_{u(2)} \cdots g_{u(p)} \right] \\
&= \mathbb{E} \left[\sum_{i=2}^p g_{u(2)} \cdots g_{u(i-1)} 1_{u(1)=u(i)} g_{u(i+1)} \cdots g_{u(p)} \right] \\
&= \sum_{i=2}^p 1_{u(1)=u(i)} \mathbb{E} [g_{u(2)} \cdots g_{u(i-1)} g_{u(i+1)} \cdots g_{u(p)}].
\end{aligned}$$

Using the induction hypothesis for $\mathbb{E} [g_{u(2)} \cdots g_{u(i-1)} g_{u(i+1)} \cdots g_{u(p)}]$ (which is a product of $p-2$ factors) we get

$$\begin{aligned}
\mathbb{E}[g_{u(1)} \cdot g_{u(2)} \cdots g_{u(p)}] &= \sum_{i=2}^p 1_{u(1)=u(i)} \sum_{\nu' \in \mathbb{P}_2([p] \setminus \{1, i\})} 1_{u|_{[p] \setminus \{1, i\}} \sim \nu'} \\
&= \sum_{i=2}^p \sum_{\nu' \in \mathbb{P}_2([p] \setminus \{1, i\})} 1_{u \sim (\nu' \cup \{1, i\})} = \sum_{\nu \in \mathbb{P}_2[p]} 1_{u \sim \nu},
\end{aligned}$$

where $u|_{[p] \setminus \{1, i\}}$ denotes the restriction of u to the set $[p] \setminus \{1, i\}$. □

8.2 Gaussian interpolation, Poincaré, and concentration

In this section, we start with a classical tool in Gaussian analysis that we will use several times, namely the idea of Gaussian Interpolation.

Lemma 8.6 (Gaussian Interpolation). *Let X and Y be two centered independent n dimensional Gaussian vectors with covariances, respectively Σ^X and Σ^Y . Define, for $t \in [0, 1]$,*

$$Z_t = \sqrt{t}X + \sqrt{1-t}Y, \quad (8.1)$$

and let $f : \mathbb{R}^n \rightarrow \mathbb{R}$ be a twice differentiable function with at most exponential growth. Then

$$\frac{d}{dt} \mathbb{E} [f(Z_t)] = \frac{1}{2} \sum_{i,j=1}^n (\Sigma_{ij}^X - \Sigma_{ij}^Y) \mathbb{E} \left[\frac{\partial^2 f}{\partial x_i \partial x_j} (Z_t) \right]. \quad (8.2)$$

Before proving this lemma, we note that the parameterization of the path, although perhaps peculiar at first, is the one that in the case that X and Y are identically distributed, renders Z_t identically distributed to them for any t ; it is also the parameterization for which $\text{Cov}(Z_t) = t \text{Cov}(X) + (1-t) \text{Cov}(Y)$, where $\text{Cov}(X) = \mathbb{E}XX^T$ (since $\mathbb{E}X = 0$).

Proof. The idea will be to use Gaussian integration by parts (Corollary 8.2). By the chain rule (and swapping the order of differentiating and taking the expectation)

$$\begin{aligned} \frac{d}{dt} \mathbb{E}[f(Z_t)] &= \mathbb{E} \left[\sum_{i=1}^n \frac{\partial f}{\partial x_i}(Z_t) \frac{d}{dt}(Z_t)_i \right] \\ &= \mathbb{E} \left[\sum_{i=1}^n \frac{\partial f}{\partial x_i}(Z_t) \left(\frac{1}{2\sqrt{t}} X_i - \frac{1}{2\sqrt{1-t}} Y_i \right) \right] \end{aligned} \quad (8.3)$$

$$= \mathbb{E} \left[\sum_{i=1}^n \frac{\partial f}{\partial x_i}(Z_t) \frac{1}{2\sqrt{t}} X_i \right] - \mathbb{E} \left[\sum_{i=1}^n \frac{\partial f}{\partial x_i}(Z_t) \frac{1}{2\sqrt{1-t}} Y_i \right] \quad (8.4)$$

For the first term, conditioning on Y and using Corollary 8.2 on the random variable $\sqrt{t}X$ gives

$$\begin{aligned} \mathbb{E} \left[\sum_{i=1}^n \frac{\partial f}{\partial x_i}(Z_t) \frac{1}{2\sqrt{t}} X_i \right] &= \frac{1}{2t} \mathbb{E} \left[\sum_{i=1}^n \frac{\partial f}{\partial x_i}(\sqrt{t}X + \sqrt{1-t}Y) \sqrt{t} X_i \right] \\ &= \frac{1}{2t} \sum_{j=1}^n (t \Sigma^X)_{ij} \mathbb{E} \left[\sum_{i=1}^n \frac{\partial^2 f}{\partial x_j \partial x_i}(\sqrt{t}X + \sqrt{1-t}Y) \right] \\ &= \frac{1}{2} \sum_{j=1}^n \Sigma_{ij}^X \mathbb{E} \left[\sum_{i=1}^n \frac{\partial^2 f}{\partial x_j \partial x_i}(Z_t) \right]. \end{aligned}$$

An analogous calculation for the second terms completes the proof. $\square$

Gaussian interpolation will play (at least) two important roles in the following: (i) together with the fundamental theorem of calculus it will allow us to compare $\mathbb{E}f(X)$ with $\mathbb{E}f(Y)$ for various functions f and Gaussian random variables X and Y , this is the key tool behind the comparison inequalities of Section 8.3; (ii) another important role of Gaussian interpolation is to provide concentration inequalities, roughly speaking by interpolating between (X, X) containing two exact copies of a random variable and (X, X') containing two independent copies, allows to quantitatively say how much two independent copies of a random variable behave similarly, resulting in concentration of measure. The first instance of this idea is in the so-called Gaussian covariance identity, but we will see it again in the Gaussian Poincaré inequality and in Gaussian concentration.

Lemma 8.7 (Gaussian covariance identity). *Let φ and ψ be two differentiable functions (from $\mathbb{R}^n$ to $\mathbb{R}$) of at most exponential growth. Let g and g' be two iid n dimensional $\mathcal{N}(0, I)$ vectors. Define*

$$g_{(t)} = tg + \sqrt{1-t^2}g'. \quad (8.5)$$

We have

$$\text{cov}(\varphi(g), \psi(g)) = \int_0^1 \mathbb{E} \langle \nabla \varphi(g), \nabla \psi(g(t)) \rangle dt.$$

Before proving this lemma, let us remark on how the interpolation in (8.5) is related to the one in (8.1). The interpolation used in (8.1) is an interpolation between two distributions of random variables linearly parameterized on the covariance of these random variables. The interpolation in the covariance identity (8.5) creates a path between a random variable and an i.i.d. copy of it, linearly parameterized by the covariance between g and $g(t)$. As we will see in the proof below, we can study (8.5) by taking $Z_t = \sqrt{t}X + \sqrt{1-t}Y$ where X and Y are two $2n$ -dimensional Gaussian random variables with covariances

$$X \sim \mathcal{N}\left(0, \begin{bmatrix} I_{n \times n} & 0_{n \times n} \\ 0_{n \times n} & I_{n \times n} \end{bmatrix}\right) \text{ and } Y \sim \mathcal{N}\left(0, \begin{bmatrix} I_{n \times n} & I_{n \times n} \\ I_{n \times n} & I_{n \times n} \end{bmatrix}\right), \quad (8.6)$$

since, for each $t \in [0, 1]$, $Z_t \sim \begin{bmatrix} g \\ g(t) \end{bmatrix}$ (in the sense of having the same probability law).

Proof. Let us start by defining $\ell(t) = \mathbb{E} [\varphi(g)\psi(g(t))]$. One has

$$\ell(0) = (\mathbb{E}\varphi(g))(\mathbb{E}\psi(g)) \text{ and } \ell(1) = \mathbb{E} [\varphi(g)\psi(g)],$$

thus

$$\text{cov}(\varphi(g), \psi(g)) = \ell(1) - \ell(0) = \int_0^1 \ell'(t) dt.$$

The idea now is to compute $\ell'(t)$ using Lemma 8.6. We leverage the construction above and consider $Z_t = \sqrt{t}X + \sqrt{1-t}Y$ where X and Y are independent gaussian vectors distributed accordingly to (8.6). Since, for each $t \in [0, 1]$, Z_t has the same law as $\begin{bmatrix} g \\ g(t) \end{bmatrix}$ we have that $f(Z_t)$ has the same law as $\varphi(g)\psi(g(t))$ for $f : \mathbb{R}^{2n} \rightarrow \mathbb{R}$ given by $f\left(\begin{bmatrix} y \\ z \end{bmatrix}\right) = \varphi(y)\psi(z)$ for $y, z \in \mathbb{R}^n$, and so $\ell(t) = \mathbb{E} [f(Z_t)]$

The rest of the proof is a straightforward computation using Lemma 8.6:

$$\begin{aligned} \ell'(t) &= \frac{d}{dt} \mathbb{E} [f(Z_t)] = \frac{1}{2} \sum_{a,b=1}^{2n} \left(\Sigma_{ab}^{Z_1} - \Sigma_{ab}^{Z_0} \right) \mathbb{E} \left[\frac{\partial^2 f}{\partial x_a \partial x_b} (Z_t) \right] \\ &= \frac{1}{2} \left(2 \sum_{i=1}^n \mathbb{E} \frac{\partial}{\partial x_i} \varphi(g) \frac{\partial}{\partial x_i} \psi(g(t)) \right) \\ &= \mathbb{E} \langle \nabla \varphi(g), \nabla \psi(g(t)) \rangle. \end{aligned}$$

□

An elegant consequence of this lemma is the celebrated Gaussian Poincaré inequality, which bounds the variance of a function of a gaussian vector by its gradient yielding a form concentration of measure: “If a function f has small gradient, $f(X)$ has small variance”.

Proposition 8.8 (Gaussian Poincaré Inequality). *Let $f : \mathbb{R}^n \rightarrow \mathbb{R}$ be a differentiable function with at most exponential growth. Let $g \sim \mathcal{N}(0, I_{n \times n})$. Then*

$$\text{Var}(f(g)) \leq \mathbb{E} \|\nabla f(g)\|^2.$$

We note that this inequality (as the others in this chapter) can be effortlessly extended to Lipschitz functions by a standard approximation argument.

Proof. By the Gaussian covariance identity (Lemma 8.7) we have

$$\begin{aligned} \text{Var}(f(g)) &= \text{cov}(f(g), f(g)) = \int_0^1 \mathbb{E} \langle \nabla f(g), \nabla f(g_t) \rangle dt \\ &\leq \int_0^1 (\mathbb{E} \|\nabla f(g)\|^2)^{\frac{1}{2}} (\mathbb{E} \|\nabla f(g_t)\|^2)^{\frac{1}{2}} dt \\ &= \int_0^1 \mathbb{E} \|\nabla f(g)\|^2 dt = \mathbb{E} \|\nabla f(g)\|^2, \end{aligned}$$

where the inequality follows from Cauchy-Schwarz on both the inner-product and the expectation and the penultimate equality follows from the fact that $g \sim g_t$. $\square$

We will now prove one of the most important results in concentration of measure, Gaussian concentration. Although being a concentration result specific for normally distributed random variables, it will be very useful throughout this book. Much like Poincaré’s inequality above, it intuitively it says that if $F : \mathbb{R}^n \rightarrow \mathbb{R}$ is a function that is stable in terms of its input then $F(X)$ is well concentrated around its mean, where $X \in \mathcal{N}(0, I)$. Unlike Poincaré’s inequality above, Gaussian concentration provides Gaussian tail bounds. More precisely:

Theorem 8.9 (Gaussian Concentration). *Let $g \sim \mathcal{N}(0, I_{n \times n})$ be a vector with i.i.d. standard Gaussian entries and $f : \mathbb{R}^n \rightarrow \mathbb{R}$ a σ -Lipchitz function (i.e.: $|f(x) - f(y)| \leq \sigma \|x - y\|$, for all $x, y \in \mathbb{R}^n$). Then, for every $t \geq 0$*

$$\mathbb{P} \{|f(g) - \mathbb{E}f(g)| \geq t\} \leq 2 \exp \left(-\frac{t^2}{2\sigma^2} \right).$$

Proof. We will assume that the function f is smooth as a limiting argument can generalize the result from smooth functions to general Lipschitz functions. The existence of the integrals we use follows from the constraints that the Lipschitz property forces on growth of f . We also assume $\mathbb{E}f(g) = 0$ for ease of notation, as we can simply consider $\varphi(g) = f(g) - \mathbb{E}f(g)$.

We consider the moment generating function (as in (2.6))

$$M_f(\lambda) = \mathbb{E}[\exp(\lambda f(g))],$$

and note that

$$M'_f(\lambda) = \mathbb{E}[f(g) \exp(\lambda f(g))] = \text{cov}(f(g), \exp(\lambda f(g))),$$

since $\mathbb{E}f(g) = 0$.

The Gaussian covariance identity (Lemma 8.7) then gives

$$\begin{aligned} M'_f(\lambda) &= \int_0^1 \mathbb{E} \langle \nabla f(g), \nabla \exp(\lambda f(g_t)) \rangle dt \\ &= \int_0^1 \mathbb{E} [\langle \nabla f(g), \nabla f(g_t) \rangle \lambda \exp(\lambda f(g_t))] dt. \end{aligned}$$

The Lipchitz condition implies that $|\langle \nabla f(g), \nabla f(g_t) \rangle| \leq \sigma^2$ which together with the fact that $M_f(\lambda) \geq 0$ gives

$$|M'_f(\lambda)| \leq \int_0^1 \sigma^2 \mathbb{E} |\lambda \exp(\lambda f(g_t))| dt = \sigma^2 |\lambda| M_f(\lambda).$$

Since $M_f(0) = 1$ this implies (one can see this, e.g., by noting that $\log M_f(0) = 0$ and $|\frac{d}{d\lambda} \log M_f(\lambda)| \leq |\lambda| \sigma^2$) that

$$M_f(\lambda) \leq \exp\left(\frac{\lambda^2 \sigma^2}{2}\right).$$

The rest of the proof follows from the standard Chernoff trick and Markov Inequality (the same we did in Proposition 2.8). Formally, we showed that $f(g)$ is a sub-gaussian random variable (in the sense of Definition 2.9) and the proof can be finished by using Proposition 2.10. $\square$

There is a direct argument for a weaker version of this inequality that does not require the covariance identity. We refer the interested reader to Theorem 2.1.12 in [169] (the original argument is due to Maurey and Pisier).

8.2.1 Talagrand's concentration inequality

A remarkable result by Talagrand [166], Talangrad's concentration inequality, provides an analogue of Gaussian concentration for bounded random variables.

Theorem 8.10 (Talangrand concentration, Thm. 2.1.13 [169]). *Let $K > 0$, and let $X_1, \dots, X_n$ be independent bounded random variables with $|X_i| \leq K$ for all $1 \leq i \leq n$. Let $F : \mathbb{R}^n \rightarrow \mathbb{R}$ be a σ -Lipschitz and convex function. Then, for any $t \geq 0$,*

$$\mathbb{P}\{|F(X) - \mathbb{E}[F(X)]| \geq tK\} \leq c_1 \exp\left(-c_2 \frac{t^2}{\sigma^2}\right),$$

for positive constants c_1 , and c_2 .

Other useful similar inequalities (with explicit constants) are available in [123].

8.3 Gaussian comparison principles

We will start this section with the celebrated Slepian's Comparison Lemma, also known as the Sudakov-Fernique inequality,² are crucial tools to compare Gaussian processes. A Gaussian process is a family of Gaussian random variables indexed by some set T , $\{X_t\}_{t \in T}$ (if T is finite this is simply a Gaussian vector). Given a Gaussian process X_t , a particular quantity of interest is $\mathbb{E}[\sup_{t \in T} X_t]$.

We will focus on Gaussian Processes that are whose maximum is well approximated by maximum on finite sets, we will call these *separable Gaussian processes*.

Definition 8.11. We say a Gaussian process $\{X_t\}_{t \in T}$ is separable if

$$\mathbb{E}\left[\sup_{t \in T} X_t\right] = \sup_{\substack{T' \subseteq T \\ |T'| < \infty}} \mathbb{E}\left[\max_{t \in T'} X_t\right].$$

The Gaussian processes we will consider are either finite, which case they are trivially separable, or they are defined over a separable set T and have path continuity on the process metric $d(t, u) = \sqrt{\mathbb{E}(X_t - X_u)^2}$ which also ensures that they satisfy the condition in Definition 8.11. In most cases, in fact, one can simply think of the approximating subsets as ε -nets of T .

We will be particularly interested in comparing two different Gaussian processes, usually one of interest with one that is simpler to understand. Intuitively, if we have two Gaussian processes X_t and Y_t with mean zero $\mathbb{E}[X_t] = \mathbb{E}[Y_t] = 0$, for all $t \in T$, and the same variance, then the process that has the “least correlations” should have a larger maximum (think the maximum entry of vector with i.i.d. Gaussian entries versus one always with the same Gaussian entry). Theorem 8.12 makes this intuition precise and extends it to processes with different variances.³

²They are essentially the same statement, on being a slight generalization of the other, we state here the more general form: Theorem 8.12.

³Although intuitive in some sense, this turns out to be a delicate statement about Gaussian random variables, as it does not hold in general for other distributions.

Theorem 8.12 (Slepian/Sudakov-Fernique inequality).

Let $\{X_u\}_{u \in U}$ and $\{Y_u\}_{u \in U}$ be two (almost surely bounded) separable centered Gaussian processes indexed by the same (compact) set U (see Definition 8.11). If, for every $u, v \in U$:

$$\mathbb{E}[X_u - X_v]^2 \leq \mathbb{E}[Y_u - Y_v]^2, \quad (8.7)$$

then

$$\mathbb{E} \left[\max_{u \in U} X_u \right] \leq \mathbb{E} \left[\max_{u \in U} Y_u \right].$$

Proof. Since X and Y are separable processes, we can reduce to the case in which U is finite, by showing the inequality for all finite subsets of the indexing set (see Definition 8.11).

The idea is to use Gaussian Interpolation (Lemma 8.6) on a soft-max. We take

$$Z(t) = \sqrt{t}X + \sqrt{1-t}Y.$$

For $\beta > 0$ we define

$$F_\beta(Z(t)) = \frac{1}{\beta} \log \left(\sum_{u \in U} e^{\beta Z_u(t)} \right).$$

Gaussian Interpolation (Lemma 8.6) implies that

$$\frac{d}{dt} \mathbb{E} F_\beta(Z(t)) = \frac{1}{2} \sum_{u,v \in U} (\Sigma_{uv}^X - \Sigma_{uv}^Y) \mathbb{E} \left[\frac{\partial^2 F}{\partial x_u \partial x_v}(Z(t)) \right].$$

Setting $p_u(Z) = e^{\beta Z_u(t)} / \sum_{v \in U} e^{\beta Z_v(t)}$ (we dropped the t from $p_u(Z(t))$ to ease notation), we have

$$\begin{aligned} \frac{\partial F}{\partial x_u}(Z(t)) &= \frac{1}{\beta} \frac{\beta e^{\beta Z_u(t)}}{\sum_{v \in U} e^{\beta Z_v(t)}} = p_u(Z), \\ \frac{\partial^2 F}{\partial x_u \partial x_v}(Z(t)) &= \delta_{uv} \frac{\beta e^{\beta Z_u(t)}}{\sum_{w \in U} e^{\beta Z_w(t)}} - e^{\beta Z_u(t)} \frac{\beta e^{\beta Z_v(t)}}{(\sum_{w \in U} e^{\beta Z_w(t)})^2} \\ &= \beta (\delta_{uv} p_u(Z) - p_u(Z) p_v(Z)). \end{aligned}$$

Our goal is to show that $\frac{d}{dt} \mathbb{E} F_\beta(Z(t))$ is non-positive, implying that $\mathbb{E} F_\beta(X) \leq \mathbb{E} F_\beta(Y)$ which by taking $\beta \rightarrow \infty$, and recalling that

$$\max_{u \in U} Z_u(t) \leq F_\beta(Z(t)) \leq \frac{\log n}{\beta} + \max_{u \in U} Z_u(t),$$

finishes the proof.

Showing that $\frac{d}{dt} \mathbb{E} F_\beta(Z(t))$ is non-positive is mechanical computation (using $\sum_{u \in U} p_u(Z) = 1$ and $\mathbb{E}[X_u - X_v]^2 = \Sigma_{uu}^X - 2\Sigma_{uv}^X + \Sigma_{vv}^X$). Indeed,

$$\frac{d}{dt} \mathbb{E} F_\beta(Z(t)) = \frac{\beta}{2} \sum_{u \in U} (\Sigma_{uu}^X - \Sigma_{uu}^Y) (p_u(Z) - p_u(Z)p_u(Z)) \quad (8.8)$$

$$- \frac{\beta}{2} \sum_{u \neq v \in U} (\Sigma_{uv}^X - \Sigma_{uv}^Y) p_u(Z)p_v(Z), \quad (8.9)$$

and

$$\begin{aligned} 0 &\geq \frac{1}{2} \sum_{u \neq v \in U} \left(\mathbb{E}[X_u - X_v]^2 - \mathbb{E}[Y_u - Y_v]^2 \right) p_u(Z)p_v(Z) \\ &= \frac{1}{2} \sum_{u \neq v \in U} (\Sigma_{uu}^X - 2\Sigma_{uv}^X + \Sigma_{vv}^X - \Sigma_{uu}^Y + 2\Sigma_{uv}^Y - \Sigma_{vv}^Y) p_u(Z)p_v(Z) \\ &= \sum_{u \neq v \in U} (\Sigma_{uu}^X - \Sigma_{uu}^Y) p_u(Z)p_v(Z) - \sum_{u \neq v \in U} (\Sigma_{uv}^X - \Sigma_{uv}^Y) p_u(Z)p_v(Z) \\ &= \sum_{u \in U} (\Sigma_{uu}^X - \Sigma_{uu}^Y) p_u(Z) \left(\sum_{v \in U \setminus \{u\}} p_v(Z) \right) - \sum_{u \neq v \in U} (\Sigma_{uv}^X - \Sigma_{uv}^Y) p_u(Z)p_v(Z) \\ &= \frac{2}{\beta} \frac{d}{dt} \mathbb{E} F_\beta(Z(t)), \end{aligned}$$

where the last equality follows from $\sum_{v \in U \setminus \{u\}} p_v(Z) = 1 - p_u(Z)$. $\square$

We will also use in the sequel an extension due to Gordon [83, 84]. While we omit the proof here, we note that it can be shown by a similar argument as above, by taking a soft min-max, see, e.g., [84] or Problem 6.2. in [180]).

Theorem 8.13. *[Theorem A in [84]] Let $\{X_{t,u}\}_{(t,u) \in T \times U}$ and $\{Y_{t,u}\}_{(t,u) \in T \times U}$ be two (almost surely bounded) separable centered Gaussian processes⁴ indexed by the same (compact) sets T and U .*

If, for every $t_1, t_2 \in T$ and $u_1, u_2 \in U$:

$$\mathbb{E}[X_{t_1, u_1} - X_{t_1, u_2}]^2 \leq \mathbb{E}[Y_{t_1, u_1} - Y_{t_1, u_2}]^2, \quad (8.10)$$

and, for $t_1 \neq t_2$,

$$\mathbb{E}[X_{t_1, u_1} - X_{t_2, u_2}]^2 \geq \mathbb{E}[Y_{t_1, u_1} - Y_{t_2, u_2}]^2, \quad (8.11)$$

then

$$\mathbb{E} \left[\min_{t \in T} \max_{u \in U} X_{t,u} \right] \leq \mathbb{E} \left[\min_{t \in T} \max_{u \in U} Y_{t,u} \right].$$

Note that Theorem 8.12 easily follows by setting $|T| = 1$.

⁴Here we need the min-max to be well approximated by finite subsets of T and U in the sense of Definition 8.11, as before the processes on which we use this result will be well approximated by ε -nets.

8.3.1 Spectral norm of a Wigner Matrix

Before proceeding with Gordon's Escape Through a Mesh Theorem we take a small detour to showcase the usefulness of what we already developed in this chapter.

Let $W \in \mathbb{R}^{n \times n}$ be a standard Gaussian Wigner matrix, a symmetric matrix with (otherwise) independent Gaussian entries, the off-diagonal entries have unit variance and the diagonal entries have variance 2 (recall (3.40)). $\lambda_{\max}(W)$ depends on $\frac{n(n+1)}{2}$ independent (standard) Gaussian random variables and it is easy to see that it is a $\sqrt{2}$ -Lipschitz function of these variables, since

$$\left| \lambda_{\max}(W^{(1)}) - \lambda_{\max}(W^{(2)}) \right| \leq \lambda_{\max}(W^{(1)} - W^{(2)}) \leq \|W^{(1)} - W^{(2)}\|_F.$$

The symmetry of the matrix and the variance 2 of the diagonal entries are responsible for an extra factor of $\sqrt{2}$.

Using Gaussian Concentration (Theorem 8.9) we immediately get

$$\mathbb{P}\{\lambda_{\max}(W) \geq \mathbb{E}\lambda_{\max}(W) + t\} \leq 2 \exp\left(-\frac{t^2}{4}\right).$$

On the other hand, one can prove $\mathbb{E}\lambda_{\max}(W) \leq 2\sqrt{n}$ using Slepian's inequality (Theorem 8.12) by comparing the gaussian process $X_u = u^W u$ with the gaussian process $Y_u = \sqrt{2}g^T u$ for $g \in \mathcal{N}(0, I_{n \times n})$, this is an excellent exercises that we leave to the reader. Combining these we get the following.

Proposition 8.14. *Let $W \in \mathbb{R}^{n \times n}$ be a standard Gaussian Wigner matrix, a symmetric matrix with (otherwise) independent Gaussian entries, the off-diagonal entries have unit variance and the diagonal entries have variance 2. Then,*

$$\mathbb{P}\{\lambda_{\max}(W) \geq 2\sqrt{n} + t\} \leq 2 \exp\left(-\frac{t^2}{4}\right).$$

Note that this gives an extremely precise control of the fluctuations of $\lambda_{\max}(W)$. In fact, for $t = 2\sqrt{\log n}$ this gives

$$\mathbb{P}\{\lambda_{\max}(W) \geq 2\sqrt{n} + 2\sqrt{\log n}\} \leq 2 \exp\left(-\frac{4 \log n}{4}\right) = \frac{2}{n}.$$

This illustrated the level of concentration we can expect from the spectral norm, or largest eigenvalues, of random matrices. As we will see in Chapter 9 the main challenge in many cases is in controlling the expected value of the spectral norm, or largest eigenvalue.

8.4 Gordon's Theorem

In Section 6.1 we showed that in order to approximately preserve the distances (up to $1 \pm \varepsilon$) between n points, it suffices to randomly project them

to $\Theta(\varepsilon^{-2} \log n)$ dimensions. The key argument was that a random projection approximately preserves the norm of every point in a set S , in this case the set of differences between pairs of n points. What we showed is that in order to approximately preserve the norm of every point in S , it is enough to project to $\Theta(\varepsilon^{-2} \log |S|)$ dimensions. The question this section is meant to answer is: can this be improved if S has a special structure? Given a set S , what is the measure of complexity of S that explains how many dimensions one needs to project on to still approximately preserve the norms of points in S . We shall see below that this will be captured—via Gordon's Theorem—by the so called *Gaussian width* of S .

Definition 8.15 (Gaussian width). *Given a closed set $S \subset \mathbb{R}^p$, its Gaussian width $\omega(S)$ is defined as:*

$$\omega(S) = \mathbb{E} \max_{x \in S} [g_p^T x],$$

where $g_p \sim \mathcal{N}(0, I_{p \times p})$.

Similarly to the proof of Theorem 6.1 we will restrict our attention to sets S of unit norm vectors, meaning that $S \subseteq \mathbb{S}^{p-1}$ (which in particular implies it is compact).

Also, we will focus our attention not in random projections but in the similar model of random linear maps $G : \mathbb{R}^p \rightarrow \mathbb{R}^d$ that are given by matrices with i.i.d. Gaussian entries. For this reason the following proposition will be useful:

Proposition 8.16. *Let $g_d \sim \mathcal{N}(0, I_{d \times d})$, and define*

$$a_d := \mathbb{E} \|g_d\|.$$

Then $\sqrt{d-1} \leq \sqrt{\frac{d}{d+1}} \sqrt{d} \leq a_d \leq \sqrt{d}$.

We will prove here the weaker lower bound $\sqrt{d-1} \leq a_d \leq \sqrt{d}$ which is sufficient for our purposes and has a very elegant proof.⁵ *Proof.* The upper bound is a straightforward consequence of Jensen's inequality

$$a_d^2 = (\mathbb{E} \|g_d\|)^2 \leq \mathbb{E} \|g_d\|^2 = d.$$

The lower bound $\sqrt{d-1} \leq a_d$ follows from the Gaussian Poincaré Inequality (Proposition 8.8): let $f(g_d) = \|g_d\|$ (and recall that since f is Lipschitz it has a gradient a.e. and we can apply Proposition 8.8. Thus $\text{Var}(f(g_d)) \leq \mathbb{E} \|\nabla f(g_d)\|^2$ which gives:

$$d - (\mathbb{E} \|g\|)^2 = \mathbb{E} \|g\|^2 - (\mathbb{E} \|g\|)^2 = \text{Var}(f(g)) \leq \mathbb{E} \|\nabla f(g)\|^2 = \mathbb{E} \|g/\|g\|\|^2 = 1$$

□

We are now ready to present Gordon's Theorem.

⁵There is a very nice discussion of this lower bound in this blog post [73].

Theorem 8.17 (Gordon's Theorem [84]). *Let $G \in \mathbb{R}^{d \times p}$ be a random matrix with independent $\mathcal{N}(0, 1)$ entries and $S \subset \mathbb{S}^{p-1}$ be a closed subset of the unit sphere in p dimensions. Then*

$$\mathbb{E} \max_{x \in S} \left\| \frac{1}{a_d} Gx \right\| \leq 1 + \frac{\omega(S)}{a_d},$$

and

$$\mathbb{E} \min_{x \in S} \left\| \frac{1}{a_d} Gx \right\| \geq 1 - \frac{\omega(S)}{a_d},$$

where $a_d = \mathbb{E} \|g_d\|$ and $\omega(S)$ is the Gaussian width of S . Recall that $\sqrt{\frac{d}{d+1}} \sqrt{d} \leq a_d \leq \sqrt{d}$.

Before proving Gordon's Theorem we will note some of its direct implications. The theorem suggests that $\frac{1}{a_d} G$ preserves the norm of the points in S up to $1 \pm \frac{\omega(S)}{a_d}$, indeed we can make this precise with Gaussian concentration (Theorem 8.9).

Note that the function $F(G) = \max_{x \in S} \|Gx\|$ is 1-Lipschitz. Indeed

$$\begin{aligned} \left| \max_{x_1 \in S} \|G_1 x_1\| - \max_{x_2 \in S} \|G_2 x_2\| \right| &\leq \max_{x \in S} \left| \|G_1 x\| - \|G_2 x\| \right| \leq \max_{x \in S} \|(G_1 - G_2)x\| \\ &\leq \|G_1 - G_2\| \leq \|G_1 - G_2\|_F. \end{aligned}$$

Similarly, one can show that $F(G) = \min_{x \in S} \|Gx\|$ is 1-Lipschitz. Thus, one can use Gaussian concentration to get

$$\mathbb{P} \left\{ \max_{x \in S} \|Gx\| \geq a_d + \omega(S) + t \right\} \leq \exp \left(-\frac{t^2}{2} \right), \quad (8.12)$$

and

$$\mathbb{P} \left\{ \min_{x \in S} \|Gx\| \leq a_d - \omega(S) - t \right\} \leq \exp \left(-\frac{t^2}{2} \right). \quad (8.13)$$

This gives us the following theorem.

Theorem 8.18. *Let $G \in \mathbb{R}^{d \times p}$ a random matrix with independent $\mathcal{N}(0, 1)$ entries and $S \subset \mathbb{S}^{p-1}$ be a closed subset of the unit sphere in p dimensions.*

Then, for $\varepsilon > \frac{\omega(S)}{a_d}$, with probability $\geq 1 - 2 \exp \left[-\frac{a_d^2}{2} \left(\varepsilon - \frac{\omega(S)}{a_d} \right)^2 \right]$:

$$(1 - \varepsilon) \|x\| \leq \left\| \frac{1}{a_d} Gx \right\| \leq (1 + \varepsilon) \|x\|,$$

for all $x \in S$.

Recall that $d \frac{d}{d+1} \leq a_d^2 \leq d$.

Proof. This is readily obtained by taking $\varepsilon = \frac{\omega(S)+t}{a_d}$, using (8.12) and (8.13). $\square$

Remark 8.19. Note that a simple use of a union bound⁶ shows that $\omega(S) \lesssim \sqrt{2 \log |S|}$, which means that taking d to be of the order of $\log |S|$ suffices to ensure that $\frac{1}{a_d}G$ to have the Johnson Lindenstrauss property. This observation shows that Theorem 8.18 essentially directly implies Theorem 6.1 (although not exactly, since $\frac{1}{a_d}G$ is not a projection).

8.4.1 Gordon's Escape Through a Mesh Theorem

Theorem 8.18 suggests that, if $\omega(S) \leq a_d$, a uniformly chosen random subspace of $\mathbb{R}^n$ of dimension $(n-d)$ (which can be seen as the nullspace of G) avoids a set S with high probability. This is indeed the case and is known as Gordon's Escape Through a Mesh Theorem (Corollary 3.4. in Gordon's original paper [84]). See also [130] for a description of the proof. We include the Theorem below for the sake of completeness.

Theorem 8.20 (Corollary 3.4. in [84]). *Let $S \subset \mathbb{S}^{p-1}$ be a closed subset of the unit sphere in p dimensions. If $\omega(S) < a_d$, then for a $(p-d)$ -dimensional subspace Λ drawn uniformly from the Grassmanian manifold we have*

$$\mathbb{P}\{\Lambda \cap S \neq \emptyset\} \leq \frac{7}{2} \exp\left(-\frac{1}{18} (a_d - \omega(S))^2\right),$$

where $\omega(S)$ is the Gaussian width of S and $a_d = \mathbb{E}\|g_d\|$ where $g_d \sim \mathcal{N}(0, I_{d \times d})$.

8.4.2 Proof of Gordon's Theorem

We are now ready to prove Gordon's Theorem.

Proof. [of Theorem 8.17]

Let $G \in \mathbb{R}^{d \times p}$ with i.i.d. $\mathcal{N}(0, 1)$ entries. We define two Gaussian processes: For $v \in S \subset \mathbb{S}^{p-1}$ and $u \in \mathbb{S}^{d-1}$ let $g \sim \mathcal{N}(0, I_{d \times d})$ and $h \sim \mathcal{N}(0, I_{p \times p})$ and define the following processes:

$$A_{u,v} = g^T u + h^T v,$$

and

$$B_{u,v} = u^T G v.$$

For all $v, v' \in S \subset \mathbb{S}^{p-1}$ and $u, u' \in \mathbb{S}^{d-1}$,

⁶This follows from the fact that the maximum of n standard Gaussian random variables is $\lesssim \sqrt{2 \log n}$.

$$\begin{aligned}
\mathbb{E} |A_{v,u} - A_{v',u'}|^2 - \mathbb{E} |B_{v,u} - B_{v',u'}|^2 &= 4 - 2(u^T u' + v^T v') - \sum_{ij} (v_i u_j - v'_i u'_j)^2 \\
&= 4 - 2(u^T u' + v^T v') - [2 - 2(v^T v')(u^T u')] \\
&= 2 - 2(u^T u' + v^T v' - u^T u' v^T v') \\
&= 2(1 - u^T u')(1 - v^T v').
\end{aligned}$$

This means that $\mathbb{E} |A_{v,u} - A_{v',u'}|^2 - \mathbb{E} |B_{v,u} - B_{v',u'}|^2 \geq 0$ and $\mathbb{E} |A_{v,u} - A_{v',u'}|^2 - \mathbb{E} |B_{v,u} - B_{v',u'}|^2 = 0$ if $v = v'$. This implies that we can use Theorem 8.13 with $X = A$ and $Y = B$ (these processes are well approximated by ε -nets and separable, see Definition 8.11), to get

$$\mathbb{E} \min_{v \in S} \max_{u \in \mathbb{S}^{d-1}} A_{v,u} \leq \mathbb{E} \min_{v \in S} \max_{u \in \mathbb{S}^{d-1}} B_{v,u}.$$

Noting that

$$\mathbb{E} \min_{v \in S} \max_{u \in \mathbb{S}^{d-1}} B_{v,u} = \mathbb{E} \min_{v \in S} \max_{u \in \mathbb{S}^{d-1}} u^T G v = \mathbb{E} \min_{v \in S} \|G v\|,$$

and

$$\begin{aligned}
\mathbb{E} \left[\min_{v \in S} \max_{u \in \mathbb{S}^{d-1}} A_{v,u} \right] &= \mathbb{E} \max_{u \in \mathbb{S}^{d-1}} g^T u + \mathbb{E} \min_{v \in S} h^T v \\
&= \mathbb{E} \max_{u \in \mathbb{S}^{d-1}} g^T u - \mathbb{E} \max_{v \in S} (-h^T v) = a_d - \omega(S),
\end{aligned}$$

gives the second part of the theorem.

On the other hand, since $\mathbb{E} |A_{v,u} - A_{v',u'}|^2 - \mathbb{E} |B_{v,u} - B_{v',u'}|^2 \geq 0$ then we can similarly use Theorem 8.12 with $X = B$ and $Y = A$, to get

$$\mathbb{E} \max_{v \in S} \max_{u \in \mathbb{S}^{d-1}} A_{v,u} \geq \mathbb{E} \max_{v \in S} \max_{u \in \mathbb{S}^{d-1}} B_{v,u}.$$

Noting that

$$\mathbb{E} \max_{v \in S} \max_{u \in \mathbb{S}^{d-1}} B_{v,u} = \mathbb{E} \max_{v \in S} \max_{u \in \mathbb{S}^{d-1}} u^T G v = \mathbb{E} \max_{v \in S} \|G v\|,$$

and

$$\mathbb{E} \left[\max_{v \in S} \max_{u \in \mathbb{S}^{d-1}} A_{v,u} \right] = \mathbb{E} \max_{u \in \mathbb{S}^{d-1}} g^T u + \mathbb{E} \max_{v \in S} h^T v = a_d + \omega(S),$$

concludes the proof of the theorem. $\square$

A remarkable application of Gordon's Theorem is that one can use it for abstract sets S such as the set of all *natural images* or the set of all plausible *user-product* ranking matrices. In these cases Gordon's Theorem suggests that a measurements corresponding just to a random projection may be enough

to keep geometric properties of the data set in question, in particular, it may allow for reconstruction of the data point from just the projection. These phenomenon and the sensing savings that arises from it is at the heart of compressive sensing and several recommendation system algorithms, among many other data processing techniques. Motivated by these two applications we will focus in this section on understanding which projections are expected to preserve the norm of sparse vectors and low-rank matrices. Compressive sensing will be discussed in detail in Chapter 10.

Remark 8.21 (Ornstein-Uhlenbeck process). There is an important side of Gaussian Analysis that we do not cover, related to the theory of Markov semigroups, and in particular, the Ornstein-Uhlenbeck process. In fact, this approach is arguably the most natural way of proving Poincaré's inequality, and also yields Log-Sobolev and hypercontractivity inequalities. We point the reader to Chapter 2 of [180] for an excellent pedagogical introduction to these ideas.

Remark 8.22 (Hermite Polynomials). Another central set of ideas in Gaussian analysis that we do not cover in this book is the theory of Hermite polynomials, which are a basis of orthogonal polynomials in the Gaussian measure, forming an analogue of Fourier analysis in the Gaussian setting. Writing a function $f \in L^2(\mathcal{N}(0, 1))$, or in $L^2(\mathcal{N}(0, I))$, in its Hermite polynomial expansion is a classical idea, often known as a Wiener Chaos expansion. Due to Gaussian integration by parts (Lemma 8.1) taking the gradient of such a function yields simple transformations on the expansion coefficients. In fact, this yields a very simple proof of Poincaré's inequality (try it!).⁷ The Hermite polynomials are the eigenfunctions of the Ornstein-Uhlenbeck Operator (mentioned in Remark 8.21). We point the reader to the classical reference [165] for more on Hermite polynomials and an introduction to the beautiful theory of orthogonal polynomials.

⁷Hermite polynomials also play a crucial role on the development of low degree method for computational hardness of hypothesis testing (see, e.g., Appendix B of [106]).

Matrix Concentration Inequalities

There are many applications where one needs to control the spectrum of random matrices. Depending on the context, these matrices may represent the noise whose effect in a spectral algorithm is controlled by its spectral norm, or the size of a dual variable that needs to be controlled to show the exactness of a convex relaxation (such as in Chapter 7). While some of the tools we developed in Chapter 2 could be used to control the size of the entries of random matrices, which could translate to spectral bounds, this would likely introduce many suboptimal dimensional factors.

In what follows we will state and prove various matrix concentration results, somewhat similar to Theorem 7.9 (in Section 7.2.7). We will focus on understanding, and bounding, the typical value of the spectral norm of random matrices by upper bounding $\mathbb{E}\|X\|$, as these tend to be high dimensional objects themselves they often have enough concentration that tail bounds are then easy to obtain (as it was illustrated in Chapter 8 with Proposition 8.14). Our presentation will rely on Gaussian analysis (developed in Chapter 8). Our treatment of the Non-commutative Khintchine inequality, and the Matrix Concentration inequality we will prove follows parts of [181] and [178].¹

9.1 Non-commutative Khintchine inequality

We start with a particularly important inequality involving the expected value of a random matrix. It is intimately related to the non-commutative Khintchine inequality [145], and for that reason we will often refer to it as Non-commutative Khintchine (see, for example, (4.9) in [175]). We will prove this inequality in Section 9.1.2.

Theorem 9.1 (Non-commutative Khintchine (NCK)). *Let $A_1, \dots, A_n \in \mathbb{R}^{d \times d}$ be symmetric matrices and $g_1, \dots, g_n \sim \mathcal{N}(0, 1)$ i.i.d., then:*

¹For an approach to matrix concentration that includes a direct proof of Theorem 7.9 we recommend Tropp's excellent monograph [176].

$$\mathbb{E} \left\| \sum_{k=1}^n g_k A_k \right\| \leq \left(2 \lceil \log d \rceil + 1 \right)^{\frac{1}{2}} \sigma, \quad (9.1)$$

where

$$\sigma^2 = \left\| \sum_{k=1}^n A_k^2 \right\|. \quad (9.2)$$

Note that, akin to Proposition 8.14, we can also use Gaussian Concentration to get a tail bound on $\left\| \sum_{k=1}^n g_k A_k \right\|$. We consider the function

$$F : \mathbb{R}^n \rightarrow \left\| \sum_{k=1}^n g_k A_k \right\|.$$

We now estimate its Lipschitz constant; let $g, h \in \mathbb{R}^n$ then

$$\begin{aligned} \left| \left\| \sum_{k=1}^n g_k A_k \right\| - \left\| \sum_{k=1}^n h_k A_k \right\| \right| &\leq \left\| \left(\sum_{k=1}^n g_k A_k \right) - \left(\sum_{k=1}^n h_k A_k \right) \right\| \\ &= \left\| \sum_{k=1}^n (g_k - h_k) A_k \right\| \\ &= \max_{v: \|v\|=1} \left| v^T \left(\sum_{k=1}^n (g_k - h_k) A_k \right) v \right| \\ &= \max_{v: \|v\|=1} \left| \sum_{k=1}^n (g_k - h_k) (v^T A_k v) \right| \\ &\leq \max_{v: \|v\|=1} \sqrt{\sum_{k=1}^n (g_k - h_k)^2} \sqrt{\sum_{k=1}^n (v^T A_k v)^2} \\ &= \sqrt{\max_{v: \|v\|=1} \sum_{k=1}^n (v^T A_k v)^2} \|g - h\|_2, \end{aligned}$$

where in the first inequality we made use of the triangular inequality and in the last one we used Cauchy-Schwarz.

This motivates us to define a new parameter, the weak variance σ_* .

Definition 9.2 (Weak Variance (see, for example, [176])). *Given symmetric matrices $A_1, \dots, A_n \in \mathbb{R}^{d \times d}$. We define the weak variance parameter as*

$$\sigma_*^2 = \max_{v: \|v\|=1} \sum_{k=1}^n (v^T A_k v)^2.$$

This means that, using Gaussian concentration (Theorem 8.9), and setting $t = u\sigma_*$, we have

$$\mathbb{P} \left\{ \left\| \sum_{k=1}^n g_k A_k \right\| \geq \left(2 + 2 \log(2d) \right)^{\frac{1}{2}} \sigma + u \sigma_* \right\} \leq \exp \left(-\frac{1}{2} u^2 \right). \quad (9.3)$$

Thus, although the expected value of $\|\sum_{k=1}^n g_k A_k\|$ is controlled by the parameter σ , its fluctuations seem to be controlled by σ_* . We compare the two quantities in the following proposition.

Proposition 9.3. *Given $A_1, \dots, A_n \in \mathbb{R}^{d \times d}$ symmetric matrices, recall that*

$$\sigma = \sqrt{\left\| \sum_{k=1}^n A_k^2 \right\|^2} \text{ and } \sigma_* = \sqrt{\max_{v: \|v\|=1} \sum_{k=1}^n (v^T A_k v)^2}.$$

We have

$$\sigma_* \leq \sigma.$$

Proof. Using the Cauchy-Schwarz inequality,

$$\begin{aligned} \sigma_*^2 &= \max_{v: \|v\|=1} \sum_{k=1}^n (v^T A_k v)^2 = \max_{v: \|v\|=1} \sum_{k=1}^n (v^T [A_k v])^2 \\ &\leq \max_{v: \|v\|=1} \sum_{k=1}^n (\|v\| \|A_k v\|)^2 = \max_{v: \|v\|=1} \sum_{k=1}^n \|A_k v\|^2 \\ &= \max_{v: \|v\|=1} \sum_{k=1}^n v^T A_k^2 v = \left\| \sum_{k=1}^n A_k^2 \right\| = \sigma^2. \end{aligned}$$

□

9.1.1 How tight is the Non-commutative Khintchine inequality?

The following simple calculation is suggestive that the parameter σ in Theorem 9.1 is indeed the correct parameter to understand $\mathbb{E} \|\sum_{k=1}^n g_k A_k\|$.

$$\begin{aligned} \mathbb{E} \left\| \sum_{k=1}^n g_k A_k \right\|^2 &= \mathbb{E} \left\| \left(\sum_{k=1}^n g_k A_k \right)^2 \right\| = \mathbb{E} \max_{v: \|v\|=1} v^T \left(\sum_{k=1}^n g_k A_k \right)^2 v \\ &\geq \max_{v: \|v\|=1} \mathbb{E} v^T \left(\sum_{k=1}^n g_k A_k \right)^2 v = \max_{v: \|v\|=1} v^T \left(\sum_{k=1}^n A_k^2 \right) v = \sigma^2. \end{aligned}$$

But a natural question is whether the logarithmic term in (9.1) is needed. Motivated by this question we will explore a couple of examples.

Example 9.4. We can write a $d \times d$ Wigner matrix W (recall Section 3.3.2) as a gaussian series, by taking A_{ij} for $i \leq j$ defined as

$$A_{ij} = e_i e_j^T + e_j e_i^T,$$

if $i \neq j$, and

$$A_{ii} = \sqrt{2} e_i e_i^T.$$

It is not difficult to see that, in this case, $\sum_{i \leq j} A_{ij}^2 = (d+1)I_{d \times d}$, meaning that $\sigma = \sqrt{d+1}$. This implies that Theorem 9.1 gives us

$$\mathbb{E}\|W\| \lesssim \sqrt{d \log d},$$

however, we know that $\mathbb{E}\|W\| \asymp \sqrt{d}$, meaning that the bound given by NCK (Theorem 9.1) is, in this case, suboptimal by a logarithmic factor.²

The next example will show that the logarithmic factor is in fact needed in some examples

Example 9.5. Consider $A_k = e_k e_k^T \in \mathbb{R}^{d \times d}$ for $k = 1, \dots, d$. The matrix $\sum_{k=1}^n g_k A_k$ corresponds to a diagonal matrix with independent standard gaussian random variables as diagonal entries, and so its spectral norm is given by $\max_k |g_k|$. It is known that $\max_{1 \leq k \leq d} |g_k| \asymp \sqrt{\log d}$. On the other hand, a direct calculation shows that $\sigma = 1$. This shows that the logarithmic factor cannot, in general, be removed.

This motivates the question of trying to understand when is it that the extra dimensional factor is needed. For both these examples, the resulting matrix $X = \sum_{k=1}^n g_k A_k$ has independent entries (except for the fact that it is symmetric). The case of independent entries [149, 156, 109, 28, 110] is now somewhat understood:³

Theorem 9.6 ([28]). *If X is a $d \times d$ random symmetric matrix with gaussian independent entries (except for the symmetry constraint) whose entry i, j has variance b_{ij}^2 , then*

$$\mathbb{E}\|X\| \lesssim \sqrt{\max_{1 \leq i \leq d} \sum_{j=1}^d b_{ij}^2 + \max_{i,j} |b_{ij}| \sqrt{\log d}}.$$

Remark 9.7. X in the theorem above can be written in terms of a Gaussian series by taking

$$A_{ij} = b_{ij} (e_i e_j^T + e_j e_i^T),$$

for $i < j$ and $A_{ii} = b_{ii} e_i e_i^T$. One can then compute σ and σ_* :

$$\sigma^2 = \max_{1 \leq i \leq d} \sum_{j=1}^d b_{ij}^2 \text{ and } \sigma_*^2 \asymp \max_{i,j} b_{ij}^2.$$

²By $a \asymp b$ we mean $a \lesssim b$ and $a \gtrsim b$.

³A notable improvement of Theorem 9.6 is a dimension free bound obtained in [110].

This means that, when the random matrix in NCK (Theorem 9.1) has independent entries (modulo symmetry) then

$$\mathbb{E}\|X\| \lesssim \sigma + \sqrt{\log d} \sigma_*. \quad (9.4)$$

Recently an improvement, using ideas from Free Probability to remove the dimensional factor in some situations, was obtained in [22] (see [42] for non-gaussian extensions). Interestingly, the same work shows that (9.4) does not hold in general, disproving a conjecture that was included in an earlier version of this manuscript. See Section 9.3 for a more thorough discussion on these improvements.

9.1.2 Proof of the Non-commutative Khintchine inequality

We are now ready to prove Theorem 9.1, the Non-commutative Khintchine inequality (NCK). See Section 7.2. in [177, Section 7.2] and [181] for presentations of the same argument.

Proof. [of Theorem 9.1]

Let p be a positive integer. Let $X = \sum_{k=1}^n g_k A_k$, we have

$$(\mathbb{E}\|X\|)^{2p} \leq \mathbb{E}\|X\|^{2p} = \mathbb{E}\|X^{2p}\| \leq \mathbb{E}\operatorname{Tr} X^{2p},$$

where the first inequality follows from Jensen and the second from the fact that the trace of a positive semidefinite matrix dominates its spectral norm.

Using Gaussian integration by parts (Lemma 8.1) we get

$$\begin{aligned} \mathbb{E}\operatorname{Tr} X^{2p} &= \mathbb{E}\operatorname{Tr} \sum_{k=1}^n g_k A_k X^{2p-1} = \mathbb{E}\operatorname{Tr} \sum_{k=1}^n \sum_{q=0}^{2p-2} A_k X^q A_k X^{2p-2} \\ &= \mathbb{E} \sum_{k=1}^n \sum_{q=0}^{2p-2} \operatorname{Tr} (A_k X^q A_k X^{2p-2}). \end{aligned}$$

If the matrices $A_1, \dots, A_n$ were commutative, then we would have that $\operatorname{Tr} (A_k X^q A_k X^{2p-2}) = \operatorname{Tr} (A_k^2 X^{2p-2})$ for all q . It turns out that this is the worst possible situation and the commutative upper bound always dominates. We state and prove in Lemma 9.8 below that, for all $0 \leq q \leq 2p-2$, $\operatorname{Tr} (A_k X^q A_k X^{2p-2}) \leq \operatorname{Tr} (A_k^2 X^{2p-2})$. This implies that

$$\mathbb{E}\operatorname{Tr} X^{2p} \leq \mathbb{E} \sum_{k=1}^n (2p-1) \operatorname{Tr} (A_k^2 X^{2p-2}) = (2p-1) \mathbb{E} \sum_{k=1}^n \operatorname{Tr} \left(\left(\sum_{k=1}^n A_k^2 \right) X^{2p-2} \right).$$

Hölder's inequality on Schatten norms then gives

$$\mathbb{E}\operatorname{Tr} X^{2p} \leq (2p-1) \left\| \sum_{k=1}^n A_k^2 \right\| \mathbb{E} \sum_{k=1}^n \operatorname{Tr} (X^{2p-2}) = (2p-1) \sigma^2 \mathbb{E} \sum_{k=1}^n \operatorname{Tr} (X^{2p-2}).$$

Iterating this procedure gives

$$\mathbb{E} \operatorname{Tr} X^{2p} \leq (2p-1)!! \sigma^{2p} d,$$

which implies

$$\mathbb{E} \|X\| \leq (\operatorname{Tr} X^{2p})^{\frac{1}{2p}} \leq ((2p-1)!! \sigma^{2p} d)^{\frac{1}{2p}} \leq [(2p-1)!!]^{\frac{1}{2p}} \sigma d^{\frac{1}{2p}}. \quad (9.5)$$

Taking $p = \lceil \log d \rceil$ and using the fact that $(2p-1)!! \leq \left(\frac{2p+1}{e}\right)^p$ (see [178] for an elementary proof consisting essentially of taking logarithms and comparing the sum with an integral, note also that Lemma 8.3 would yield the same bound up to a constant of e) we get

$$\mathbb{E} \|X\| \leq \left(\frac{2\lceil \log d \rceil + 1}{e} \right)^{\frac{1}{2}} \sigma d^{\frac{1}{2\lceil \log d \rceil}} \leq (2\lceil \log d \rceil + 1)^{\frac{1}{2}} \sigma.$$

□

Lemma 9.8. *[Commutative is the worst-case I] Let A and X be symmetric matrices. For any p, q non-negative integers satisfying $0 \geq q \geq 2p-2$ we have*

$$\operatorname{Tr} (AX^q AX^{2p-2}) \leq \operatorname{Tr} (A^2 X^{2p-2}).$$

Proof. Consider the spectral decomposition $X = \sum_{i=1}^d \lambda_i v_i v_i^T$, then

$$\begin{aligned} \operatorname{Tr} (AX^q AX^{2p-2}) &= \sum_{i,j=1}^d \lambda_i^q \lambda_j^{2p-2-q} \operatorname{Tr} (A v_i v_i^T A v_j v_j^T) \\ &= \sum_{i,j=1}^d \lambda_i^q \lambda_j^{2p-2-q} (v_i^T A v_j)^2 \\ &\leq \left(\sum_{i,j=1}^d |\lambda_i|^{2p-2} (v_i^T A v_j)^2 \right)^{\frac{q}{2p-2}} \left(\sum_{i,j=1}^d |\lambda_j|^{2p-2} (v_i^T A v_j)^2 \right)^{\frac{2p-2-q}{2p-2}} \\ &= \sum_{i,j=1}^d |\lambda_i|^{2p-2} (v_i^T A v_j)^2 = \sum_{i,j=1}^d \lambda_i^{2p-2} (v_j^T A v_i) (v_i^T A v_j) \\ &= \sum_{i,j=1}^d \operatorname{Tr} (\lambda_i^{2p-2} v_j^T A v_i v_i^T A v_j) = \operatorname{Tr} (A^2 X^{2p-2}), \end{aligned}$$

where the inequality follows from Hölder's inequality (Proposition 3.3, in particular (3.12)). □

Remark 9.9. Theorem 9.1 is an upper bound on a Gaussian process, and thus an orderwise optimal bound can (in theory) be obtained by Talagrand's generic chaining, and up to a logarithmic factor by covering number arguments via a Dudley's (see [167]). Interestingly, there is no known proof of a bound that is optimal up to logarithmic factors based on these techniques (in other words, without using operator theoretical arguments). Such an argument would likely extend well beyond the setting of spectral norm in matrices, see [30] for a discussion on this and related problems.

9.2 Matrix concentration inequalities

In what follows, we closely follow [178] and present an elementary proof of a few useful matrix concentration inequalities using Theorem 9.1. We start with a Rademacher version of Theorem 9.1.

Theorem 9.10. *Let $H_1, \dots, H_n \in \mathbb{R}^{d \times d}$ be symmetric matrices and $\varepsilon_1, \dots, \varepsilon_n$ i.i.d. Rademacher random variables (meaning $= +1$ with probability $1/2$ and $= -1$ with probability $1/2$), then:*

$$\mathbb{E} \left\| \sum_{k=1}^n \varepsilon_k H_k \right\| \leq \left(\frac{\pi}{2} + \pi \lceil \log(d) \rceil \right)^{\frac{1}{2}} \sigma,$$

where

$$\sigma^2 = \left\| \sum_{k=1}^n H_k^2 \right\|. \quad (9.6)$$

Proof. Let $g_1, \dots, g_n$ be iid standard gaussians independent from the Rademacher random variables, since the sign and absolute value of g_k are independent and $\mathbb{E}|g_k| = \sqrt{\frac{2}{\pi}}$, Jensen implies

$$\begin{aligned} \mathbb{E} \left\| \sum_{k=1}^n \varepsilon_k H_k \right\| &= \frac{1}{\mathbb{E}|g_1|} \mathbb{E} \left\| \sum_{k=1}^n (\mathbb{E}|g_k|) \varepsilon_k H_k \right\| \\ &\leq \sqrt{\frac{\pi}{2}} \mathbb{E} \left\| \sum_{k=1}^n |g_k| \varepsilon_k H_k \right\| = \sqrt{\frac{\pi}{2}} \mathbb{E} \left\| \sum_{k=1}^n g_k H_k \right\|. \end{aligned}$$

The conclusion then follows by applying Theorem 9.1. $\square$

We note that Theorem 9.10 holds without the extra $\frac{\pi}{2}$ factor (see [178] for a proof that boils down to the same as argument as the one above to prove its Gaussian counterpart, Theorem 9.1).

Using Theorem 9.10, we will prove the following theorem.

Theorem 9.11. *Let $T_1, \dots, T_n \in \mathbb{R}^{d \times d}$ be random independent symmetric positive semidefinite matrices, then*

$$\mathbb{E} \left\| \sum_{i=1}^n T_i \right\| \leq \left[\left\| \sum_{i=1}^n \mathbb{E} T_i \right\|^{\frac{1}{2}} + \sqrt{C(d)} \left(\mathbb{E} \max_i \|T_i\| \right)^{\frac{1}{2}} \right]^2,$$

where

$$C(d) := 2\pi + 4\pi \lceil \log d \rceil. \quad (9.7)$$

A key step in the proof of Theorem 9.11 is an idea that is extremely useful in Probability, the trick of symmetrization. For this reason we isolate it in a lemma.

Lemma 9.12 (Symmetrization). *Let $T_1, \dots, T_n$ be independent random matrices (note that they do not necessarily need to be positive semidefinite, for the sake of this lemma) and $\varepsilon_1, \dots, \varepsilon_n$ random i.i.d. Rademacher random variables (independent also from the matrices). Then*

$$\mathbb{E} \left\| \sum_{i=1}^n T_i \right\| \leq \left\| \sum_{i=1}^n \mathbb{E} T_i \right\| + 2\mathbb{E} \left\| \sum_{i=1}^n \varepsilon_i T_i \right\|$$

Proof. The triangular inequality gives

$$\mathbb{E} \left\| \sum_{i=1}^n T_i \right\| \leq \left\| \sum_{i=1}^n \mathbb{E} T_i \right\| + \mathbb{E} \left\| \sum_{i=1}^n (T_i - \mathbb{E} T_i) \right\|.$$

Let us now introduce, for each i , a random matrix T'_i identically distributed to T_i and independent (all $2n$ matrices are independent). Then

$$\begin{aligned} \mathbb{E} \left\| \sum_{i=1}^n (T_i - \mathbb{E} T_i) \right\| &= \mathbb{E}_T \left\| \sum_{i=1}^n (T_i - \mathbb{E} T_i - \mathbb{E}_{T'_i} [T'_i - \mathbb{E}_{T'_i} T'_i]) \right\| \\ &= \mathbb{E}_T \left\| \mathbb{E}_{T'} \sum_{i=1}^n (T_i - T'_i) \right\| \leq \mathbb{E} \left\| \sum_{i=1}^n (T_i - T'_i) \right\|, \end{aligned}$$

where we use the notation $\mathbb{E}_a$ to mean that the expectation is taken with respect to the variable a and the last step follows from Jensen's inequality with respect to $\mathbb{E}_{T'}$.

Since $T_i - T'_i$ is a symmetric random variable,⁴ it is identically distributed to $\varepsilon_i (T_i - T'_i)$, which gives

⁴Note that we use the notation “symmetric random variable” to mean $X \sim -X$ and “symmetric matrix” to mean $X^T = X$

$$\begin{aligned}
\mathbb{E} \left\| \sum_{i=1}^n (T_i - T'_i) \right\| &= \mathbb{E} \left\| \sum_{i=1}^n \varepsilon_i (T_i - T'_i) \right\| \\
&\leq \mathbb{E} \left\| \sum_{i=1}^n \varepsilon_i T_i \right\| + \mathbb{E} \left\| \sum_{i=1}^n \varepsilon_i T'_i \right\| = 2\mathbb{E} \left\| \sum_{i=1}^n \varepsilon_i T_i \right\|,
\end{aligned}$$

concluding the proof. $\square$

Proof. [of Theorem 9.11]

Using Lemma 9.12 and Theorem 9.10 we get

$$\mathbb{E} \left\| \sum_{i=1}^n T_i \right\| \leq \left\| \sum_{i=1}^n \mathbb{E} T_i \right\| + \sqrt{C(d)} \mathbb{E} \left\| \sum_{i=1}^n T_i^2 \right\|^{\frac{1}{2}}$$

The trick now is to make a term like the one in the LHS appear in the RHS. For that we start by noting (you can see Fact 2.3 in [178] for an elementary proof) that, since $T_i \succeq 0$,

$$\left\| \sum_{i=1}^n T_i^2 \right\| \leq \max_i \|T_i\| \left\| \sum_{i=1}^n T_i \right\|.$$

This means that

$$\mathbb{E} \left\| \sum_{i=1}^n T_i \right\| \leq \left\| \sum_{i=1}^n \mathbb{E} T_i \right\| + \sqrt{C(d)} \mathbb{E} \left[\left(\max_i \|T_i\| \right)^{\frac{1}{2}} \left\| \sum_{i=1}^n T_i \right\|^{\frac{1}{2}} \right].$$

Furthermore, applying the Cauchy-Schwarz inequality for $\mathbb{E}$ gives,

$$\mathbb{E} \left\| \sum_{i=1}^n T_i \right\| \leq \left\| \sum_{i=1}^n \mathbb{E} T_i \right\| + \sqrt{C(d)} \left(\mathbb{E} \max_i \|T_i\| \right)^{\frac{1}{2}} \left(\mathbb{E} \left\| \sum_{i=1}^n T_i \right\| \right)^{\frac{1}{2}},$$

Now that the term $\mathbb{E} \left\| \sum_{i=1}^n T_i \right\|$ appears in the RHS, the proof can be finished with a simple application of the quadratic formula (see Section 6.1. in [178] for details). $\square$

We now show an inequality for general symmetric matrices

Theorem 9.13. *Let $Y_1, \dots, Y_n \in \mathbb{R}^{d \times d}$ be random independent symmetric matrices satisfying $\mathbb{E} Y_i = 0$, then*

$$\mathbb{E} \left\| \sum_{i=1}^n Y_i \right\| \leq \sqrt{C(d)} \sigma + C(d) L,$$

where,

$$\sigma^2 = \left\| \sum_{i=1}^n \mathbb{E} Y_i^2 \right\| \text{ and } L^2 = \mathbb{E} \max_i \|Y_i\|^2 \quad (9.8)$$

and, as in (9.7),

$$C(d) := 2\pi + 4\pi \lceil \log d \rceil.$$

Proof.

Using Symmetrization (Lemma 9.12) and Theorem 9.10, we get

$$\mathbb{E} \left\| \sum_{i=1}^n Y_i \right\| \leq 2\mathbb{E}_Y \left[\mathbb{E}_\varepsilon \left\| \sum_{i=1}^n \varepsilon_i Y_i \right\| \right] \leq \sqrt{C(d)} \mathbb{E} \left\| \sum_{i=1}^n Y_i^2 \right\|^{\frac{1}{2}}.$$

Jensen's inequality gives

$$\mathbb{E} \left\| \sum_{i=1}^n Y_i^2 \right\|^{\frac{1}{2}} \leq \left(\mathbb{E} \left\| \sum_{i=1}^n Y_i^2 \right\| \right)^{\frac{1}{2}},$$

and the proof can be concluded by noting that $Y_i^2 \succeq 0$ and using Theorem 9.11. $\square$

Remark 9.14 (The rectangular case). One can extend Theorem 9.13 to general rectangular matrices $S_1, \dots, S_n \in \mathbb{R}^{d_1 \times d_2}$ by setting

$$Y_i = \begin{bmatrix} 0 & S_i \\ S_i^T & 0 \end{bmatrix},$$

the so-called Hermitian dilation, and noting that

$$\|Y_i^2\| = \left\| \begin{bmatrix} 0 & S_i \\ S_i^T & 0 \end{bmatrix}^2 \right\| = \left\| \begin{bmatrix} S_i S_i^T & 0 \\ 0 & S_i^T S_i \end{bmatrix} \right\| = \max \{ \|S_i^T S_i\|, \|S_i S_i^T\| \}.$$

For details we refer to [178].

9.3 Improvements leveraging intrinsic freeness

In Section 9.1.1 we saw that the logarithmic factor in the Non-commutative Khintchine (Theorem 9.1) inequality is required in the worst-case but superfluous in some instances. If we pay close attention to the proof of Theorem 9.1, we notice that there is a potential loss in using Lemma 9.8 when $A_1, \dots, A_n$ are non-commutative, and that perhaps the $(2p-1)$ factor coming from bounding every summand with the $q=0$ one, could be (in some cases) replaced by a constant, which would remove the logarithmic factor in the final

bound.⁵ Tropp [177] took the first steps in connecting improvements in the Non-commutative Khintchine and ideas in Free Probability [139] showing an improvement of Theorem 9.1 in [177]. This is done by using Gaussian Integration by parts twice and controlling, with a (sometimes) smaller parameter ω – called the matrix alignment parameter – the summands where the factors from different applications of Gaussian Integration by parts “cross”. This allows for a more efficient iteration argument and replaces the $\theta((\log n)^{\frac{1}{2}})$ factor multiplying σ by a $\theta((\log n)^{\frac{1}{4}})$ factor.

We briefly describe below an improvement to Theorem 9.1 in [22] that often yields sharp bounds. To describe it, it is worth showing a slightly different proof of Theorem 9.1. Let $X = \sum_{k=1}^n g_k A_k$ and $p \geq 1$ as in Theorem 9.1. Wick’s formula (Lemma 8.5) gives

$$\begin{aligned} \mathbb{E}X^{2p} &= \sum_{u:[2p] \rightarrow [n]} \mathbb{E}[g_{u(1)} \cdots g_{u(2p)}] \operatorname{Tr}(A_{u(1)} \cdots A_{u(2p)}) \\ &= \sum_{u:[2p] \rightarrow [n]} \sum_{\nu \in \mathbb{P}_2[2p]} 1_{u \sim \nu} \operatorname{Tr}(A_{u(1)} \cdots A_{u(2p)}) \\ &= \sum_{\nu \in \mathbb{P}_2[2p]} \sum_{\substack{u:[2p] \rightarrow [n] \\ u \sim \nu}} \operatorname{Tr}(A_{u(1)} \cdots A_{u(2p)}). \end{aligned}$$

If the matrices A_k were commutative then the summands

$$\sum_{\substack{u:[2p] \rightarrow [n] \\ u \sim \nu}} \operatorname{Tr}(A_{u(1)} \cdots A_{u(2p)})$$

would coincide for all pair partitions, and in particular with the pair partition $\nu_0 := \{(1, 2), (3, 4), \dots, (2p-1, 2p)\}$ on which the summand is given by

$$\sum_{\substack{u:[2p] \rightarrow [n] \\ u \sim \nu_0}} \operatorname{Tr}(A_{u(1)} \cdots A_{u(2p)}) = \sum_{u:[p] \rightarrow [n]} \operatorname{Tr}(A_{u(1)}^2 \cdots A_{u(p)}^2) = \operatorname{Tr}\left(\sum_{k=1}^n A_k^2\right)^p.$$

The following Lemma is analogous to Lemma 9.8 and due to Buchholz [44].

Lemma 9.15. *[Commutative is the worst-case II [44]] For any $\nu \in \mathbb{P}[2p]$ and $A_1 \dots, A_n$ symmetric matrices*

$$\sum_{\substack{u:[2p] \rightarrow [n] \\ u \sim \nu}} \operatorname{Tr}(A_{u(1)} \cdots A_{u(2p)}) \leq \operatorname{Tr}\left(\sum_{k=1}^n A_k^2\right)^p.$$

⁵It is worth noting that both the term $q = 0$ and $q = 2p-2$ are equal to the upper bound in Lemma 9.8, however while it is tempting to hope to bound all other other terms by a smaller order parameter, this would yield a final bound of $\sqrt{2}\sigma$, which we know does not hold for Wigner matrices (which would need 2σ). The notion of crossing and non-crossing partitions (See Definition 9.16) is an elegant way of singling out which terms are substantial.

If $\sum_{k=1}^n A_k^2$ is a multiple of the identity matrix, there are many pair partitions that match the upper bound above: whenever there is an adjacent pair, it can be “peeled-off” in the sum and potentially make adjacent pairs that were not adjacent before; the partitions that can be fully “peeled-off” this way are precisely the so-called *non-crossing* partitions and they indeed match the upper bound above (see Lemma 9.21).

Definition 9.16 (Crossing and Non-crossing Partitions). *We say $\nu \in \mathbb{P}[2p]$ is a crossing partition when it has pairs $(i_1, i_2) \in \nu$ and $(j_1, j_2) \in \nu$ such that $i_1 < j_1 < i_2 < j_2$. Otherwise we say ν is non-crossing. The set of non-crossing partition is denoted by $\text{NC}[2p] \subset \mathbb{P}[2p]$.*

Bandeira, Boedihardjo, and van Handel [22] showed that, in many settings, the random matrix X exhibits intrinsic freeness in the sense that only the summand corresponding to non-crossing partitions are non-negligible.

This is done by interpolating (using Gaussian Interpolation, Lemma 8.6) the gaussian model X with a Free Probability model X_{free} where the gaussians are replaced by a *semi-circular family* – these can be viewed, in a sense, as “non-commutative random variables” for which Wick’s formula (Lemma 8.5) holds when summing only over $\text{NC}[2p]$ (this can be viewed as the limiting behavior of Lemma 9.19 below, and indeed the argument can be viewed as replacing Gaussians by ever larger Wigner matrices). While the matrix alignment parameter of Tropp is a key component in the argument, intrinsic freeness is controlled by $v^2 = \|\text{Cov} X\|$, the spectral norm of the covariance matrix of the entries of X (this covariance matrix is a $d^2 \times d^2$ matrix). When $\frac{\sigma}{v} \gg \text{polylog}(d)$ the improvement yields $\mathbb{E}\|X\| \leq (2 + o(1))\sigma$. Furthermore, Brailovskaya and van Handel [42] developed a universality principle that roughly states that for any random matrix $Y = \sum_{i=1}^n Y_i$ sum of independent matrices (such as in Theorem 9.13), as long as the summands are sufficiently small, the matrix Y behaves like a gaussian analogue Y_G where the entries of Y are replaced by gaussian random variables with the same mean and covariance, and for which the results in [22] can be used (notice that this is different than the argument, based on symmetrization, done to prove Theorem 9.13 from Theorem 9.1). Combining these two tools, one obtains an improvement over the matrix Bernstein inequality.

Theorem 9.17 ([22, 42]). *Let $Y_1, \dots, Y_n \in \mathbb{R}^{d \times d}$ be random independent symmetric matrices satisfying $\mathbb{E}Y_i = 0$, and such that $\|Y_i\| \leq R$, for all $i \in [n]$, almost surely. Then*

$$\mathbb{E} \left\| \sum_{i=1}^n Y_i \right\| \leq 2\sigma + C \left(v^{\frac{1}{2}} \sigma^{\frac{1}{2}} (\log d)^{\frac{3}{4}} + R^{\frac{1}{3}} \sigma^{\frac{2}{3}} (\log d)^{\frac{2}{3}} + R \log d \right),$$

and

$$\mathbb{P} \left[\left\| \sum_{i=1}^n Y_i \right\| \geq 2\sigma + C \left(v^{\frac{1}{2}} \sigma^{\frac{1}{2}} (\log d)^{\frac{3}{4}} + \sigma_* t^{\frac{1}{2}} + R^{\frac{1}{3}} \sigma^{\frac{2}{3}} t^{\frac{2}{3}} + Rt \right) \right] \leq de^{-t}$$

where, C is a universal constant,

$$\sigma^2 = \left\| \sum_{i=1}^n \mathbb{E} Y_i^2 \right\|, v^2 = \|\text{Cov}(Y)\| \text{ and } \sigma_*^2 = \sup_{\|u\|=\|w\|=1} \mathbb{E} |v^T Y w|^2. \quad (9.9)$$

Note that if $v, \sigma_*, R \ll \frac{1}{\text{polylog}(d)} \sigma$, which happens often in applications (see [22, 42]), then all terms multiplying the universal constant C are negligible and, in the tail bound, the tail parameter t only appears in low-order terms.

9.3.1 Crossings and Wick's formula for the GUE

An instructive pursuit is to compute an analogue of Wick's formula for a complex valued analogue of the Wigner matrix that has appeared in Section 3.3.2 and Example 3.3.2). This is an instance in which the phenomenon of cancellations arising from crossing partitions is particularly transparent (and it is tightly connected to the seminal work of Voiculescu [186] and Haagerup-Thorbjørnsen [86], and to the methods in [22] that give rise to the inequalities in Theorem 9.17).

Definition 9.18 (Gaussian Unitary Ensemble (GUE)). *A random $D \times D$ GUE matrix W is a random Hermitian (self-adjoint) matrix whose upper off-diagonal entries are iid $\mathbb{CN}(0, \frac{1}{D})$ and whose diagonal entries are iid $\mathcal{N}(0, \frac{1}{D})$ (also independent from the off-diagonal entries). $W_{ij} \sim \mathbb{CN}(0, \frac{1}{D})$ means that $\text{Re}(W_{ij}) \sim \mathcal{N}(0, \frac{1}{D})$, $\text{Im}(W_{ij}) \sim \mathcal{N}(0, \frac{1}{D})$, and that both are independent.*

The goal of this subsection is to show that, for large-dimensional GUEs, there is Wick's formula whose non-negligible terms only involve non-crossing partitions.

Lemma 9.19 (Wick's formula for GUEs (cf. Lemma 8.5)). *Let $W_1, \dots, W_n$ be iid $D \times D$ GUE matrices and let $u : [2p] \rightarrow [n]$ then*

$$\frac{1}{D} \mathbb{E} [\text{Tr } W_{u(1)} \cdots W_{u(2p)}] = \sum_{\nu \in \text{NC}_2[2p]} 1_{u \sim \nu} + \mathcal{O}\left(\frac{1}{D^2}\right).$$

where $\text{NC}_2[2p]$ denotes the set of non-crossing pair partitions, $u \sim \nu$ means the function u is compatible with ν , and the constant in $\mathcal{O}$ may depend on n, p , and u (see Definitions 8.4 and 9.16).

We start by showing that non-crossing terms do not have cancellations.

Definition 9.20. *Let $\nu \in \mathbb{P}[2p]$ (see Definition 8.4). We define $u_\nu : [2p] \rightarrow [p]$ as first surjective assignment $u : [2p] \rightarrow [p]$, in the lexicographic order, that satisfies $u \sim \nu$.⁶*

⁶For all our uses here, we could have picked any surjective assignment $u : [2p] \rightarrow [p]$, the one that is lexicographically first is an arbitrary choice.

Proposition 9.21. *Let $\nu \in \text{NC}[2p]$, and let $W_1, \dots, W_p$ be iid $D \times D$ GUE matrices, then*

$$\mathbb{E} [\text{Tr } W_{u_\nu(1)} \cdots W_{u_\nu(2p)}] = \text{Tr}(I) \quad (9.10)$$

Proof. Since ν is non-crossing there must exist $i \in [2p-1]$ such that $u_\nu(i) = u_\nu(i+1)$. For such as i we have that the LHS of (9.10) is equal to

$$\mathbb{E} \left[\text{Tr } W_{u_\nu(1)} \cdots W_{u_\nu(i-1)} \left(\mathbb{E}_{W_{u_\nu(i)}} W_{u_\nu(i)}^2 \right) W_{u_\nu(i+2)} \cdots W_{u_\nu(2p)} \right],$$

where the outside expectation is over the other random matrices. Furthermore, since $\left(\mathbb{E}_{W_{u_\nu(i)}} W_{u_\nu(i)}^2 \right) = I$, such pairs can be “peeled-off” one by one to finish the proof (formally, one can do induction by noting that $\nu \setminus \{i, i+1\}$ is also non-crossing). $\square$

We are now ready to examine crossing partitions.

Proposition 9.22. *Let B_1, B_2, B_3, B_4 be $D \times D$ (complex) matrices and let W, V be two independent $D \times D$ GUE matrices. We have*

$$\mathbb{E} [\text{Tr } W B_1 V B_2 W B_3 V B_4] = \frac{1}{D^2} \text{Tr} (B_3 B_2 B_1 B_4). \quad (9.11)$$

Proof. This follows from a computation, and noting that for $W \sim \text{GUE}$ we have $\mathbb{E} W_{i_1 j_1} W_{i_2 j_2} = \frac{1}{D} \delta_{i_1 j_2} \delta_{i_2 j_1}$ (where δ_{ij} is the indicator of $i = j$).⁷ Indeed, expanding the LHS of (9.11) we see it is equal to

$$\begin{aligned} \text{LHS of (9.11)} &= \sum_{i_1, \dots, i_8=1}^D \mathbb{E} [W_{i_1 i_2} (B_1)_{i_2 i_3} V_{i_3 i_4} (B_2)_{i_4 i_5} W_{i_5 i_6} (B_3)_{i_6 i_7} V_{i_7 i_8} (B_4)_{i_8 i_1}] \\ &= \sum_{i_1, \dots, i_8=1}^D \frac{1}{D^2} \delta_{i_1 i_6} \delta_{i_2 i_5} \delta_{i_3 i_8} \delta_{i_4 i_7} (B_1)_{i_2 i_3} (B_2)_{i_4 i_5} (B_3)_{i_6 i_7} (B_4)_{i_8 i_1} \\ &= \frac{1}{D^2} \sum_{i_1, \dots, i_4=1}^D (B_1)_{i_2 i_3} (B_2)_{i_4 i_2} (B_3)_{i_1 i_4} (B_4)_{i_3 i_1} \\ &= \frac{1}{D^2} \text{Tr} [B_3 B_2 B_1 B_4]. \end{aligned}$$

$\square$

Proposition 9.23. *Let $\nu \in \mathbb{P}[2p] \setminus \text{NC}[2p]$ be a crossing partition, and let $W_1, \dots, W_p$ be iid $D \times D$ GUE matrices, then*

$$0 \leq \mathbb{E} [\text{Tr } W_{u_\nu(1)} \cdots W_{u_\nu(2p)}] \leq \frac{1}{D^2} \text{Tr}(I) \quad (9.12)$$

⁷An analogous calculation can be done for real value Wigner matrices (Example 9.4), also called GOE, but there are more choices of i_1, j_1, i_2, j_2 that yield non-zero moments, making the calculation less transparent).

Proof. This follows by taking a crossing $\{i, j, i', j'\}$ such that $i < j < i' < j'$, $u_\nu(i) = u_\nu(i')$, and $u_\nu(j) = u_\nu(j')$, applying Proposition 9.22 on $W_{u_\nu(i)}$ and $W_{u_\nu(j)}$ and using induction on p . $\square$

Proof. [of Lemma 9.19]

With all the ingredients collected, this follows directly from the useful equality

$$\mathbb{E} [\text{Tr } W_{u(1)} \cdots W_{u(2p)}] = \sum_{\nu \in \mathbb{P}[2p]} 1_{u \sim \nu} \mathbb{E} [\text{Tr } W_{u_\nu(1)} \cdots W_{u_\nu(2p)}], \quad (9.13)$$

and Propositions 9.23 and 9.21.

The equality (9.13) is a consequence of Lemma 8.5, a formal proof follows by noting that W_ℓ are gaussian matrices and thus we can write $W_\ell = \sum_{j=1}^m g_{\ell;j} A_j$ for some A_j (complex valued). This means we have

$$\begin{aligned} \mathbb{E} [\text{Tr } W_{u(1)} \cdots W_{u(2p)}] &= \sum_{j_1, \dots, j_{2p}=1}^m \text{Tr} (A_{j_1} \cdots A_{j_{2p}}) \mathbb{E} [g_{u(1);j_1} \cdots g_{u(2p);j_{2p}}] \\ &= \sum_{j_1, \dots, j_{2p}=1}^m \text{Tr} (A_{j_1} \cdots A_{j_{2p}}) \sum_{\nu \in \mathbb{P}[2p]} 1_{u \sim \nu} 1_{j \sim \nu} \\ &= \sum_{\nu \in \mathbb{P}[2p]} 1_{u \sim \nu} \left(\sum_{j_1, \dots, j_{2p}=1}^m \text{Tr} (A_{j_1} \cdots A_{j_{2p}}) 1_{j \sim \nu} \right), \end{aligned}$$

and on the other hand,

$$\begin{aligned} \mathbb{E} [\text{Tr } W_{u_\nu(1)} \cdots W_{u_\nu(2p)}] &= \sum_{j_1, \dots, j_{2p}=1}^m \text{Tr} (A_{j_1} \cdots A_{j_{2p}}) \mathbb{E} [g_{u_\nu(1);j_1} \cdots g_{u_\nu(2p);j_{2p}}] \\ &= \sum_{j_1, \dots, j_{2p}=1}^m \text{Tr} (A_{j_1} \cdots A_{j_{2p}}) 1_{j \sim \nu}, \end{aligned}$$

where the second equality uses the fact that ν is the only pairing compatible with u_ν . $\square$

Compressive Sensing and Sparsity

Most of us have noticed how saving an image in JPEG dramatically reduces the space it occupies on our hard drives (as opposed to file types that save the value of each pixel in the image). The idea behind these compression methods is to exploit known structure in the images; although our cameras will record the value (even three values in RGB) for each pixel, it is clear that most collections of pixel values will not correspond to pictures that we would expect to see. Natural images do not correspond to arbitrary arrays of pixel values, but have some specific structure to them. It is this special structure one aims to exploit by choosing a proper representation of the image. Indeed, natural images are known to be approximately sparse in certain bases (such as the wavelet bases) and this is the core idea behind JPEG (actually, JPEG2000; JPEG uses a different basis).

10.1 Sparse recovery

Let us think of $x \in \mathbb{C}^p$ as the signal corresponding to the image already represented in the basis in which it is sparse. The modeling assumption is that x is s -sparse, or $\|x\|_0 \leq s$, meaning that x has at most s non-zero components and, usually, $s \ll p$. The ℓ_0 -norm¹ $\|x\|_0$ of a vector x is the number of non-zero entries of x . This means that when we take a picture, our camera makes p measurements (each corresponding to a pixel) but then, after an appropriate change of basis, it keeps only $s \ll p$ non-zero coefficients and drops the others. This seems a rather wasteful procedure and thus motivates the question: “If only a few degrees of freedom are kept after compression, why not in the first place measure in a more efficient way and take considerably less than p measurements?”.

The question whether we can carry out data acquisition and compression simultaneously is at the heart of *Compressive Sensing* [46, 47, 48, 49, 66, 77].

¹We recall that the ℓ_0 norm is not actually a norm, as it does not necessarily rescale linearly with a rescaling of x .

It is particularly important in MRI imaging [119, 74], as fewer measurements potentially means shorter data acquisition time. Indeed, current MRI technology based on concepts from compressive sensing can reduce the time needed to collect the data by a factor of six or more [119], which has significant benefits especially in pediatric MR imaging [184]. We recommend the book [77] as a great in-depth reference about compressive sensing.

In mathematical terms, the acquired measurements $y \in \mathbb{C}^m$ are connected to the signal of interest $x \in \mathbb{C}^p$, with $m \ll p$, via

$$\begin{bmatrix} y \end{bmatrix} = \begin{bmatrix} & A & \end{bmatrix} \begin{bmatrix} x \end{bmatrix}. \quad (10.1)$$

The matrix $A \in \mathbb{C}^{m \times p}$ models the linear measurement (information) process. Classical linear algebra tells us that if $m < p$, then the linear system (10.1) is underdetermined and that there are infinitely many solutions (assuming that there exists at least one solution). In other words, without additional information, it is impossible to recover x from y in the case $m < p$.

In this chapter, we assume that x is s -sparse with $s < m \ll p$. The goal is to recover x from this underdetermined system and do this in a computationally efficient manner. We emphasize that we *do not know* the location of the non-zero coefficients of x a priori², otherwise the task would be trivial.

10.1.1 Gaussian width of s -sparse vectors

Before discussing algorithms, let us discuss well-posedness of the problem at hand: In order to be able to reconstruct x from y we need at the very least that A is injective on sparse vectors. Furthermore, in order for reconstruction to be stable, one should ask not only that A is injective with respect to s -sparse vectors, but actually that it is almost an isometry, meaning that the ℓ_2 distance between Ax_1 and Ax_2 should be comparable to the distances between x_1 and x_2 , if they are s -sparse. Since the difference between two s -sparse vectors is in general a $2s$ -sparse vector, we can alternatively ask for A to approximately preserve the norm of $2s$ -sparse vectors. Gordon's Theorem (Theorem 8.17) suggests that we can take $A \in \mathbb{R}^{m \times p}$ to have i.i.d. Gaussian entries and to take $m \approx \omega^2(\mathcal{S}_{2s})$, where $\mathcal{S}_{2s} = \{x : x \in \mathbb{S}^{p-1}, \|x\|_0 \leq 2s\}$ is the set of $2s$ -sparse vectors, and $\omega(\mathcal{S}_{2s})$ denotes the Gaussian width of $\mathcal{S}_{2s}$ (see Definition 8.15).

²And therein lies the challenge, since s -sparse signals do not form a linear subspace of $\mathbb{R}^p$ (the sum of two s -sparse signals is in general no longer s -sparse but $2s$ -sparse).

Proposition 10.1. *If $s \leq p$, the Gaussian width $\omega(\mathcal{S}_s)$ of $\mathcal{S}_s$, the set of unit-norm vectors that are at most s -sparse, satisfies*

$$\omega(\mathcal{S}_s)^2 \lesssim s \log\left(\frac{p}{s}\right).$$

Proof.

$$\omega(\mathcal{S}_s) = \mathbb{E} \max_{v \in \mathbb{S}^{p-1}, \|v\|_0 \leq s} g^T v,$$

where $g \sim \mathcal{N}(0, I_{p \times p})$. We have

$$\omega(\mathcal{S}_s) = \mathbb{E} \max_{\Gamma \subset [p], |\Gamma|=s} \|g_\Gamma\|,$$

where g_Γ is the restriction of g to the set of indices Γ .

Given a set Γ , Theorem 2.21 yields

$$\mathbb{P} \left\{ \|g_\Gamma\|^2 \geq s + 2\sqrt{s}\sqrt{t} + 2t \right\} \leq \exp(-t).$$

Union bounding over all $\Gamma \subset [p]$, $|\Gamma| = s$ gives

$$\mathbb{P} \left\{ \max_{\Gamma \subset [p], |\Gamma|=s} \|g_\Gamma\|^2 \geq s + 2\sqrt{s}\sqrt{t} + 2t \right\} \leq \binom{p}{s} \exp(-t).$$

Taking u such that $t = su$, gives

$$\mathbb{P} \left\{ \max_{\Gamma \subset [p], |\Gamma|=s} \|g_\Gamma\|^2 \geq s(1 + 2\sqrt{u} + 2u) \right\} \leq \exp \left[-su + s \log \left(e \frac{p}{s} \right) \right]. \quad (10.2)$$

Taking $u > \log(e \frac{p}{s})$ it can be readily seen that the typical size of $\max_{\Gamma \subset [p], |\Gamma|=s} \|g_\Gamma\|^2$ is $\lesssim s \log(\frac{p}{s})$. The proof can be completed by integrating (10.2) in order to get a bound of the expectation of $\sqrt{\max_{\Gamma \subset [p], |\Gamma|=s} \|g_\Gamma\|^2}$. $\square$

This suggests that $\approx 2s \log(\frac{p}{2s})$ measurements suffice to stably identify a $2s$ -sparse vector. Indeed, we will see below that at this order of number of measurements it will actually be possible to efficiently recover an s -sparse vector.

10.1.2 Sparse recovery and ℓ_1 optimization

Since the system (10.1) is underdetermined and we know that x is sparse, the natural approach to try to recover x is to solve

$$\begin{aligned} \min \quad & \|z\|_0 \\ \text{s.t.} \quad & Az = y, \end{aligned} \quad (10.3)$$

and hope that the optimal solution z corresponds to the signal in question x . However the optimization problem (10.3) is NP-hard in general [77]. Instead,

the approach usually taken in sparse recovery is to consider a convex surrogate of the ℓ_0 norm, namely the ℓ_1 norm: $\|z\|_1 = \sum_{i=1}^p |z_i|$. Figure 10.1 depicts the ℓ_p balls and illustrates how the ℓ_1 norm can be seen as a convex surrogate of the ℓ_0 norm due to the *pointiness* of the ℓ_1 ball in the direction of the basis vectors, i.e. in “sparse” directions.

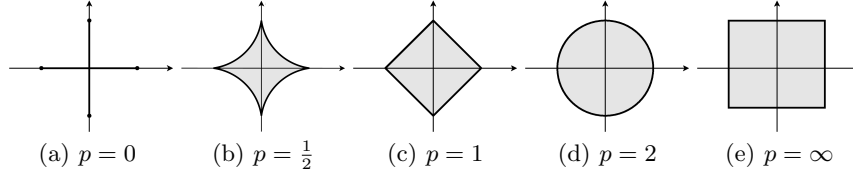


Fig. 10.1: ℓ_p norm unit balls with different values for p

The process of ℓ_p minimization can be understood as inflating (or deflating) the ℓ_p ball until one hits the affine subspace of interest. Figure 10.2 illustrates how ℓ_1 norm minimization promotes sparsity, while ℓ_2 norm minimization does not favor sparse solutions. We have seen in Chapter 2.1.4 that the ℓ_1 ball becomes “increasingly pointy” with increasing dimension. This behavior works in our favor in compressive sensing—another manifestation of the *blessings of dimensionality*.

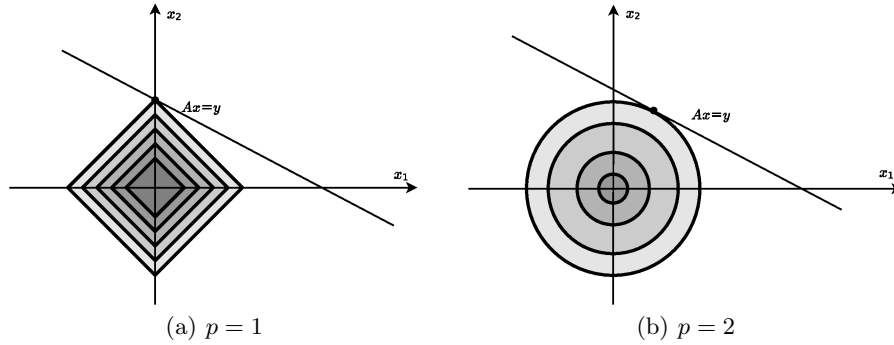


Fig. 10.2: A two-dimensional depiction of ℓ_1 and ℓ_2 minimization. In ℓ_p minimization, one inflates the ℓ_p ball until it hits the affine subspace of interest. This image conveys how the ℓ_1 norm (left) promotes sparsity due to the “pointiness” of the ℓ_1 ball. In contrast, ℓ_2 norm minimization (right) does not favor sparse solutions.

This motivates us to consider the following optimization problem (surrogate to (10.3)):

$$\begin{aligned} \min \quad & \|z\|_1 \\ \text{s.t.} \quad & Az = y, \end{aligned} \tag{10.4}$$

In order for (10.4) to be a good procedure for sparse recovery we need two things: for the solution of it to be meaningful (hopefully to coincide with x); and for (10.4) to be efficiently solvable.

Remark 10.2. We will consider for the moment the real-valued case $x \in \mathbb{R}^p$, $A \in \mathbb{R}^{m \times p}$ and formulate (10.4) as a linear program³ (and thus show that it is efficiently solvable). Let us think of ω^+ as the positive part of z and ω^- as the symmetric of the negative part of it, meaning that $z = \omega^+ - \omega^-$ and, for each i , either ω_i^- or ω_i^+ is zero. Note that, in that case,

$$\|z\|_1 = \sum_{i=1}^p \omega_i^+ + \omega_i^- = \mathbf{1}^T (\omega^+ + \omega^-).$$

Motivated by this, we consider:

$$\begin{aligned} \min \quad & \mathbf{1}^T (\omega^+ + \omega^-) \\ \text{s.t.} \quad & A(\omega^+ - \omega^-) = y \\ & \omega^+ \geq 0 \\ & \omega^- \geq 0, \end{aligned} \tag{10.5}$$

which is a linear program. It is not difficult to see that the optimal solution of (10.5) will indeed satisfy that, for each i , either ω_i^- or ω_i^+ is zero and it is indeed equivalent to (10.4); if both ω_i^- and ω_i^+ are non-zero, one can lower the objective while keep satisfying the constraints by reducing both variables. Since linear programs are efficiently solvable [183], this implies that the ℓ_1 -optimization problem (10.4) is efficiently solvable. This optimization program is also efficiently solvable over $\mathbb{C}$ despite not being a linear program.

In what follows we will discuss under which circumstances one can guarantee that the solution of (10.4) coincides with the sparse signal of interest. We will discuss a couple of different strategies to show this, as different strategies generalize better to other problems of interest. Later in this chapter we also discuss strategies for constructing sensing matrices.

10.2 Null space property and exact recovery

Given a s -sparse vector x , our goal is to show that under certain conditions x is the unique optimal solution to

$$\begin{aligned} \min \quad & \|z\|_1 \\ \text{s.t.} \quad & Az = y, \end{aligned} \tag{10.6}$$

³In the complex case, we are dealing with a quadratic program.

Let $S = \text{supp}(x)$, with $|S| = s$.⁴ If x is not the unique optimal solution of the ℓ_1 minimization problem, there exists $z \neq x$ as optimal solution. Let $v = z - x$, it satisfies

$$\|v + x\|_1 \leq \|x\|_1 \quad \text{and} \quad A(v + x) = Ax,$$

this means that $Av = 0$. Also,

$$\|x\|_S = \|x\|_1 \geq \|v + x\|_1 = \|(v + x)_S\|_1 + \|v_{S^c}\|_1 \geq \|x_S\|_1 - \|v_S\|_1 + \|v_{S^c}\|_1,$$

where the last inequality follows by applying the triangle inequality. This means that $\|v_S\|_1 \geq \|v_{S^c}\|_1$, but since $|S| \ll p$ it is unlikely for A to have vectors in its nullspace that are so concentrated on such few entries. This motivates the following definition.

Definition 10.3 (Null Space Property). *A is said to satisfy the s -Null Space Property ($A \in s\text{-NSP}$) if, for all v in $\ker(A)$ (the nullspace of A) and all $|S| = s$, we have*

$$\|v_S\|_1 < \|v_{S^c}\|_1.$$

From the argument above, it is clear that if A satisfies the Null Space Property for s , then x will indeed be the unique optimal solution to (10.4). In fact, as the property is described in terms of any set S of size s , it implies recovery for any s -sparse vector.

Theorem 10.4. *Let x be an s -sparse vector. If $A \in s\text{-NSP}$ then x is the unique solution to the ℓ_1 optimization problem (10.4) with $y = Ax$.*

The Null Space Property is a statement about certain vectors not belonging to the null space of A , thus we can again resort to Gordon's Theorem (Theorem 8.17) to establish recovery guarantees for Gaussian sensing matrices. Let us define the intersection with the unit-sphere of the cone of such vectors

$$C_s := \{v \in \mathbb{S}^{p-1} : \exists_{S \subset [p], |S|=s} \|v_S\|_1 \geq \|v_{S^c}\|_1\}. \quad (10.7)$$

10.2.1 Gordon's Theorem and the Null Space Property

Since for an $m \times p$ matrix A , $A \in s\text{-NSP}$ is equivalent to $\ker(A) \cap C_s = \emptyset$, Gordon's Theorem, or more specifically Gordon's Escape Through a Mesh Theorem (Theorem 8.20), implies that there exists a universal $C > 0$ such that if A is drawn with iid Gaussian entries, it will satisfy the $s\text{-NSP}$ with high probability provided that $m \geq C \omega^2(C_s)$, where $\omega(C_s)$ is the Gaussian width of C_s (as defined in (10.7)).

⁴If x has support size strictly smaller than s , in what follows, we can simply take a superset of it with size s

Proposition 10.5. *Let $s \leq p$ and $C_s \subset \mathbb{S}^{p-1}$ defined in (10.7). There exists a universal constant C such that*

$$\omega(C_s) \leq C \sqrt{s \log\left(\frac{p}{s}\right)},$$

where $\omega(C_s)$ is the Gaussian width of C_s (as defined in (10.7)).

Proof. The goal is to upper bound

$$\omega(C_s) = \mathbb{E} \max_{v \in C_s} v^T g,$$

for $g \sim \mathcal{N}(0, I)$. Note that C_s is invariant under permutations of the indices. Thus, the maximizer $v \in C_s$ will have its largest entries (in absolute value) in the coordinates where g has its largest entries (in absolute value). Let S be the set of the s coordinates with largest absolute value of g . We have

$$\mathbb{E} \max_{v \in C_s} v^T g = \mathbb{E} \max_{v: \|v_S\|_1 \geq \|v_{S^c}\|_1, \|v\|_2=1} v_S^T g_S + v_{S^c}^T g_{S^c}.$$

The key idea is to notice that the condition $\|v_S\|_1 \geq \|v_{S^c}\|_1$ imposes a strong bound on the ℓ_1 norm of v_{S^c} via $\|v_{S^c}\|_1 \leq \|v_S\|_1 \leq \sqrt{s} \|v_S\|_2 \leq \sqrt{s}$. This can be leveraged by noticing that

$$v_S^T g_S + v_{S^c}^T g_{S^c} \leq \|v_S\|_2 \|g_S\|_2 + \|v_{S^c}\|_1 \|g_{S^c}\|_\infty.$$

This gives

$$\omega(C_s) \leq \mathbb{E} \|g_S\|_2 + \sqrt{s} \|g_{S^c}\|_\infty,$$

where S corresponds to the set of the s coordinates with largest absolute value of g .

We saw in the proof of Proposition 10.1, in the context of computing the Gaussian width of the set of sparse vectors, that $\mathbb{E} \|g_S\|_2 \lesssim \sqrt{s \log\left(\frac{p}{s}\right)}$. Since all entries of g_{S^c} are smaller, in absolute value than any entry in g_S we have that $\|g_{S^c}\|_\infty^2 \leq \frac{1}{s} \|g_S\|_2^2$. This implies that $\mathbb{E} \|g_{S^c}\|_\infty \lesssim \sqrt{\log\left(\frac{p}{s}\right)}$, concluding the proof. $\square$

Together with Theorem 10.4 this implies the following recovery guarantee, matching the order of number of measurements suggested by the Gaussian width of sparse vectors.

Theorem 10.6. *There exists a universal constant $C \geq 0$ such that if A is a $m \times p$ matrix with i.i.d. Gaussian entries, and $m \geq Cs \log\left(\frac{p}{s}\right)$, the following holds with high probability: For any x an s -sparse vector, x is the unique solution to the ℓ_1 -optimization problem (10.4) with $y = Ax$.*

10.3 The Restricted Isometry Property

An alternative (and more classical) approach to establishing exact recovery via ℓ_1 -minimization is through the Restricted Isometry Property (RIP), which corresponds precisely with the property of approximately preserving the length of sparse vectors.

Definition 10.7 (Restricted Isometry Property (RIP)). *An $m \times p$ matrix A (with either real or complex valued entries) is said to satisfy the (s, δ) -Restricted Isometry Property (RIP),*

$$(1 - \delta)\|x\|^2 \leq \|Ax\|^2 \leq (1 + \delta)\|x\|^2,$$

for all s -sparse x .

If A satisfies the RIP for sparsity $2s$, it means that it approximately preserves distances between s -sparse vectors (hence the name *RIP*). This can be leveraged to show that A satisfies the NSP.

Theorem 10.8 ([51]). *Let $y = Ax$ where x is an s -sparse vector. Assume that A satisfies the RIP property with $\delta_{2s} < \frac{1}{3}$, then the solution x_* to the ℓ_1 -minimization problem*

$$\min_z \|z\|_1, \quad \text{subject to } Az = y = Ax \quad (10.8)$$

becomes x exactly, i.e., $x_* = x$

To prove this theorem we need the following lemma.

Lemma 10.9 ([51]). *We have*

$$|\langle Ax, Ax' \rangle| \leq \delta_{s+s'} \|x\|_2 \|x'\|_2$$

for all x, x' supported on disjoint subsets $S, S' \subseteq [1, \dots, p]$, $x, x' \in \mathbb{R}^p$, and $|S| \leq s$, $|S'| \leq s'$

Proof.

Without loss of generality, we can assume $\|x\|_2 = \|x'\|_2 = 1$, so that the right hand side of the inequality becomes just $\delta_{s+s'}$. Since A satisfies the RIP property, we have

$$(1 - \delta_{s+s'}) \|x \pm x'\|_2^2 \leq \|A(x \pm x')\|_2^2 \leq (1 + \delta_{s+s'}) \|x \pm x'\|_2^2.$$

Since x and x' have disjoint support, $\|x \pm x'\|_2^2 = \|x\|_2^2 + \|x'\|_2^2 = 2$; the RIP property then becomes

$$2(1 - \delta_{s+s'}) \leq \|Ax \pm Ax'\|_2^2 \leq 2(1 + \delta_{s+s'})$$

The polarization identity implies:

$$\begin{aligned}
|\langle Ax, Ax' \rangle| &= \frac{1}{4} \left| \|Ax + Ax'\|_2^2 - \|Ax - Ax'\|_2^2 \right| \\
&\leq \frac{1}{4} \left| 2(1 + \delta_{s+s'}) - 2(1 - \delta_{s+s'}) \right| \\
&= \delta_{s+s'}.
\end{aligned}$$

□

To prove Theorem 10.8, we simply need to show that the Null Space Property holds for the given conditions.

Proof (of Theorem 10.8). Take $h \in \ker(A) \setminus 0$. Let index set S_0 be the set of indices of s largest entries (by modulus) of h . Let index sets $S_1, S_2, \dots$ be index sets corresponding to the next s to $2s$, $2s$ to $3s$, $\dots$ largest entries of h .

Since A satisfies the RIP, we have

$$\|h_{S_0}\|_2^2 \leq \frac{1}{1 - \delta_s} \|Ah_{S_0}\|_2^2 \quad (10.9)$$

$$= \frac{1}{1 - \delta_s} \sum_{j \geq 1} \langle Ah_{S_0}, A(-h_{S_j}) \rangle \quad (\text{because } h_{S_0} = \sum_{j \geq 1} (-h_{S_j})) \quad (10.10)$$

$$\leq \frac{1}{1 - \delta_s} \sum_{j \geq 1} \delta_{2s} \|h_{S_0}\|_2 \|h_{S_j}\|_2 \quad (\text{by Lemma 10.9}) \quad (10.11)$$

$$\leq \frac{\delta_{2s}}{1 - \delta_s} \|h_{S_0}\|_2 \sum_{j \geq 1} \|h_{S_j}\|_2 \quad (10.12)$$

$$\|h_{S_0}\|_2 \leq \frac{\delta_{2s}}{1 - \delta_s} \sum_{j \geq 1} \|h_{S_j}\|_2. \quad (10.13)$$

Note that

$$\|h_{S_j}\|_2 \leq s^{\frac{1}{2}} \|h_{S_j}\|_\infty \leq s^{-\frac{1}{2}} \|h_{S_{j-1}}\|_1.$$

We can rewrite (10.13) as

$$\|h_{S_0}\|_2 \leq \frac{\delta_{2s}}{1 - \delta_s} s^{-\frac{1}{2}} \sum_{j \geq 1} \|h_{S_{j-1}}\|_1 \quad (10.14)$$

$$= \frac{\delta_{2s}}{1 - \delta_s} s^{-\frac{1}{2}} \|h\|_1. \quad (10.15)$$

Also, by the Cauchy-Schwarz inequality,

$$\|h_{S_0}\|_1 = \sum_{i \in S_0} 1 \times |h_i| \leq \sqrt{\sum_{i \in S_0} 1^2} \sqrt{\sum_{i \in S_0} h_i^2} = \sqrt{s} \|h_{S_0}\|_2. \quad (10.16)$$

We have $\delta_{2s} < \frac{1}{3}$ as a condition, so

$$\frac{\delta_{2s}}{1 - \delta_s} < \frac{\delta_{2s}}{1 - \delta_{2s}} < \frac{1}{2} \quad \text{for } \delta_{2s} < \frac{1}{3}. \quad (10.17)$$

Combining (10.15), (10.16), and (10.17), we get

$$\|h_{S_0}\|_1 < \frac{1}{2}\|h\|_1. \quad (10.18)$$

Now we show that (10.18) is equivalent to $\|h_S\|_1 < \|h_{S^c}\|_1$:

$$\begin{aligned} & \|h_S\|_1 < \|h_{S^c}\|_1 \\ \Leftrightarrow & 2\|h_S\|_1 < \|h_{S^c}\|_1 + \|h_S\|_1 \\ \Leftrightarrow & 2\|h_S\|_1 < \|h\|_1 \\ \Leftrightarrow & \|h_S\|_1 < \frac{1}{2}\|h\|_1. \end{aligned}$$

Thus, we have shown that $\|h_{S_0}\|_1 < \|h_{S^c}\|_1$, which is the Null Space Property and by virtue of Theorem 10.4 our proof is complete. $\square$

Many results in compressive sensing (such as Theorem 10.8) can be extended with little extra effort to the case where x is not exactly s -sparse, but only approximately s -sparse, a property that is sometimes referred to as *compressible*. See [50, 77] for a detailed discussion.

10.3.1 Random matrices and the Restricted Isometry Property

Theorem 10.6 (and its proof) showed conditions under which a random gaussian matrices satisfy the NSP. To show the same for RIP is straightforward with the mathematical machinery we have now developed. Indeed, using Proposition 10.1 and Theorem 8.18, one can readily show⁵ that matrices with Gaussian entries satisfy the RIP with $m \approx s \log\left(\frac{p}{s}\right)$.

Theorem 10.10. *Let A be an $m \times p$ matrix with i.i.d. standard Gaussian entries, there exists a constant C such that, if*

$$m \geq Cs \log\left(\frac{p}{s}\right), \quad (10.19)$$

then $\frac{1}{\sqrt{m}}A$ satisfies the $(s, \frac{1}{4})$ -RIP with high probability.

We point out an important aspect in this context. Theorems 10.6 and 10.10 combined with Theorem 10.8 yield a *uniform recovery guarantee* for sparse vectors with Gaussian sensing matrices. Once a Gaussian matrix satisfies the RIP or NSP (which it will for certain parameters with high probability), then

⁵Note that the $1 \pm \delta$ term in the RIP property corresponds to $(1 \pm \varepsilon)^2$ in Gordon's Theorem. Since the RIP is a stronger property when δ is smaller, one can simply use $\varepsilon = \frac{1}{3}\delta$.

exact recovery via ℓ_1 -minimization holds uniformly for *all* sufficiently sparse vectors.

Figure 10.3 illustrates the phase transition phenomenon in compressive sensing: The plot shows that for a given sparsity level s , ℓ_1 minimization almost always succeeds in finding the sparsest solution if the number of measurements is above a certain threshold, while it almost always fails if the number of measurements is below that threshold. The very sharp transition between failure and success visible in the plot is supported by theoretical analysis [67, 12]. The sensing matrix in this experiment is a Gaussian random matrix of dimension $m \times p$ where the ambient dimension $p = 100$ is fixed. The number of measurements, m varies from 1 to 100, and the sparsity s varies from 1 to m . For each choice of s and m we construct a sparse vector x with non-zero random coefficients at s locations chosen uniformly at random from $1, \dots, 100$. We compute $y = Ax$ and solve for x via (10.8). For each choice of s and m this experiment (randomly chosen A and x) is repeated 50 times. We plot the empirical rate of success (here, success means that the relative reconstruction error is less than 10^{-5}), where black means complete failure and white means complete success. A detailed analysis of this phase transition phenomenon can be found in [67, 12].

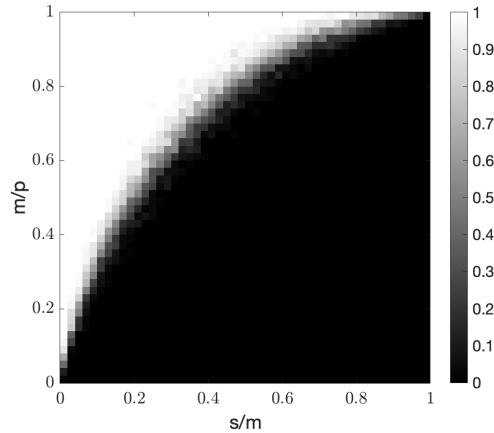


Fig. 10.3: The phase transition phenomenon in compressive sensing: if the number of measurements is above a certain threshold, ℓ_1 minimization almost always succeeds in finding the sparsest solution, while it almost always fails if the number of measurements is below that threshold. We plot the empirical probability for different sparsity levels and different number of measurements at fixed ambient dimension, where black represents 100% failure and white represents 100% success. Note the sharp transition between the two extremes.

While there are obvious similarities between Johnson-Lindenstrauss projections and sensing matrices that satisfy the RIP, there are also important differences.⁶ We note that for JL dimension reduction to be applicable (an upper estimate of) the number of vectors must be known a priori (and this number if finite). JL projection preserves (up to ϵ) pairwise distances between these vectors, but the vectors do not have to be sparse. As a consequence, JL projections P are, in general, a one-way street, as in general one cannot recover x from $y = Px$. In contrast, a matrix that satisfies the RIP works for infinitely many vectors, however with the caveat that these vectors must be sparse. Moreover, one can recover such sparse vectors x from $y = Ax$ (and can do so numerically efficiently).

As a consequence of these considerations, a matrix that satisfies the RIP does not necessarily have to satisfy the Johnson-Lindenstrauss Lemma. While a Gaussian random matrix does indeed satisfy both, RIP and the Johnson-Lindenstrauss Lemma, other matrices do not satisfy both simultaneously. For example, take a randomly subsampled Fourier matrix A of dimensions $m \times p$. In the notation of the definition of the Fast Johnson-Lindenstrauss transform, this matrix A would correspond to $A = SF$, but without the diagonal matrix D that randomizes phases (or signs) of x . This matrix A will not meet the Johnson-Lindenstrauss properties of Theorem 6.1. But the absence of the phase randomization matrix D is not a hurdle for $A = SF$ to satisfy the RIP under appropriate conditions on the matrix dimensions.

Indeed, it is known [49] that if $m = \Omega_\delta(s \text{ polylog } p)$, then the partial Fourier matrix satisfies the RIP with high probability. The exact number of logarithmic factors needed is the object of much research with the best known upper bound due to Haviv and Regev [91], giving an upper bound of $m = \Omega_\delta(s \log^2 s \log p)$. On the side of lower bounds it is known that the asymptotics established for Gaussian matrices of $m = \Theta_\delta(s \log(p/s))$ are not achievable in general [26].

Checking whether a matrix satisfies the RIP or not is in general NP-hard [23, 171]. While Theorem 10.10 suggests that RIP matrices are abundant for $s \approx \frac{m}{\log(p)}$, it appears to be very difficult to deterministically construct matrices that satisfy RIP for $s \gg \sqrt{m}$, known as the square bottleneck [168, 25, 24, 27, 41, 129]. The only known unconditional construction that is able to break this bottleneck is due to Bourgain et al. [41]; their construction achieves $s \approx m^{\frac{1}{2} + \epsilon}$ for a small, but positive, ϵ . There is also a conditional construction, based on the Paley Equiangular Tight Frame [25, 27].

In Section 10.5.1 we will consider more practical conditions for designing sensing matrices. These conditions, which are better suited for applications, are based on the concept of the *coherence* of a matrix. Interestingly, the phase randomization of x that is notably absent in the partial Fourier matrix mentioned above, will reappear in this context in connection with *nonuniform recovery guarantees*.

⁶Interestingly, there are known relationships between the two objects [103].

Remark 10.11. If one is interested in understanding the probability of exact recovery of a specific sparse vector, and not a uniform guarantee on all sparse vectors simultaneously, then it is possible to do a more refined version of the arguments above that are able to predict the exact asymptotics of the number of measurements required; see [53] for an approach based on Gaussian widths and [11] for an approach based on Integral Geometry [11].

10.4 Duality and exact recovery

In this section we describe yet another approach to show exact recovery of sparse vectors via (10.4). In this section we take an approach based on duality, the same strategy we took in Chapter 7 to show exact recovery in the Stochastic Block Model. The approach presented here is essentially the same as the one followed in [78] for the real case, and in [173] for the complex case.

Let us start by presenting duality in Linear Programming with a game theoretic view point, similarly to how we did for Semidefinite Programming in Chapter 7. The idea is again to reformulate (10.5) without constraints, by adding a dual player that wants to maximize the objective and would exploit a deviation from the original constraints (by, for example, giving the dual player a variable u and adding to the objective $u^T (y - A(\omega^+ - \omega^-))$). More precisely consider the following

$$\min_{\substack{\omega^+ \\ \omega^-}} \max_{\substack{u \\ v^+ \geq 0 \\ v^- \geq 0}} \mathbf{1}^T (\omega^+ + \omega^-) - (v^+)^T \omega^+ - (v^-)^T \omega^- + u^T (y - A(\omega^+ - \omega^-)). \quad (10.20)$$

Indeed, if the primal player (picking ω^+ and ω^- and attempting to minimize the objective) picks variables that do not satisfy the original constraints, then the dual player (picking u, v^+ , and v^- and trying to maximize the objective) will be able to make the objective value as large as possible. It is then clear that (10.5) = (10.20).

If the order with which the players choose variable values, this can only benefit the primal player, that now gets to see the value of the dual variables before picking the primal variables, meaning that (10.20) $\geq$ (10.21), where (10.21) is given by:

$$\max_{\substack{u \\ v^+ \geq 0 \\ v^- \geq 0}} \min_{\substack{\omega^+ \\ \omega^-}} \mathbf{1}^T (\omega^+ + \omega^-) - (v^+)^T \omega^+ - (v^-)^T \omega^- + u^T (y - A(\omega^+ - \omega^-)). \quad (10.21)$$

Rewriting

$$\max_{\substack{u \\ v^+ \geq 0 \\ v^- \geq 0}} \min_{\substack{\omega^+ \\ \omega^-}} (\mathbf{1} - v^+ - A^T u)^T \omega^+ + (\mathbf{1} - v^- + A^T u)^T \omega^- + u^T y \quad (10.22)$$

With this formulation, it becomes clear that the dual players needs to set $\mathbf{1} - v^+ - A^T u = 0$, $\mathbf{1} - v^- + A^T u = 0$ and thus (10.22) is equivalent to

$$\begin{aligned} \max_u \quad & u^T y \\ \text{s.t.} \quad & v^+ \geq 0 \\ & v^- \geq 0 \\ & \mathbf{1} - v^+ - A^T u = 0 \\ & \mathbf{1} - v^- + A^T u = 0 \end{aligned}$$

or equivalently,

$$\begin{aligned} \max_u \quad & u^T y \\ \text{s.t.} \quad & -\mathbf{1} \leq A^T u \leq \mathbf{1}. \end{aligned} \tag{10.23}$$

The linear program (10.23) is known as the dual program to (10.5). The discussion above shows that (10.23) $\leq$ (10.5) which is known as weak duality. More remarkably, strong duality guarantees that the optimal values of the two programs match.

There is a considerably easier way to show weak duality (although not as enlightening as the one above). If ω^- and ω^+ are primal feasible and u is dual feasible, then

$$\begin{aligned} 0 &\leq (\mathbf{1}^T - u^T A) \omega^+ + (\mathbf{1}^T + u^T A) \omega^- \\ &= \mathbf{1}^T (\omega^+ + \omega^-) - u^T [A (\omega^+ - \omega^-)] = \mathbf{1}^T (\omega^+ + \omega^-) - u^T y, \end{aligned} \tag{10.24}$$

showing that (10.23) $\leq$ (10.5).

10.4.1 Finding a dual certificate

In order to show that $\omega^+ - \omega^- = x$ is an optimal solution⁷ to (10.5), we will find a dual feasible point u for which the dual matches the value of $\omega^+ - \omega^- = x$ in the primal, u is known as a *dual certificate* or *dual witness*.

From (10.24) it is clear that u must satisfy $(\mathbf{1}^T - u^T A) \omega^+ = 0$ and $(\mathbf{1}^T + u^T A) \omega^- = 0$, this is known as *complementary slackness*. This means that we must take the entries of $A^T u$ be +1 or -1 when x is non-zero (and be +1 when it is positive and -1 when it is negative), in other words

$$(A^T u)_S = \text{sign}(x_S),$$

where $S = \text{supp}(x)$, and $\|A^T u\|_\infty \leq 1$ (in order to be dual feasible).

Remark 10.12. It is not difficult to see that if we further ask $\|(A^T u)_{S^c}\|_\infty < 1$, any optimal primal solution would have to have its support contained in the support of x . This observation gives us the following lemma.

⁷For now we will focus on showing that it is an optimal solution, see Remark 10.12 for a brief discussion of how to strengthen the argument to show uniqueness.

Lemma 10.13. *Consider the problem of sparse recovery discussed above. Let $S = \text{supp}(x)$, if A_S is injective and there exists $u \in \mathbb{R}^M$ such that*

$$(A^T u)_S = \text{sign}(x_S),$$

and

$$\|(A^T u)_{S^c}\|_\infty < 1,$$

then x is the unique optimal solution to the ℓ_1 -minimization problem (10.4).

Since we know that $(A^T u)_S = \text{sign}(x_S)$ (and that A_S is injective), we try to construct⁸ u by least squares and hope that it satisfies $\|(A^T u)_{S^c}\|_\infty < 1$. More precisely, we take

$$u = (A_S^T)^\dagger \text{sign}(x_S),$$

where $(A_S^T)^\dagger = A_S (A_S^T A_S)^{-1}$ is the Moore-Penrose pseudo-inverse of A_S^T . This gives the following corollary.

Corollary 10.14. *Consider the problem of sparse recovery. Let $S = \text{supp}(x)$. If A_S is injective and*

$$\|A_{S^c}^T A_S (A_S^T A_S)^{-1} \text{sign}(x_S)\|_\infty < 1,$$

then x is the unique optimal solution to the ℓ_1 -minimization problem (10.4).

Theorem 10.10 establishes that if $m \geq Cs \log(\frac{p}{s})$, for a universal constant C , and $A \in \mathbb{R}^{m \times p}$ is drawn with i.i.d. Gaussian entries $\mathcal{N}(0, \frac{1}{m})$ then⁹ it will, with high probability, satisfy the $(s, 1/3)$ -RIP. Note that, if A satisfies the $(s, 1/3)$ -RIP then, for any $|S| \leq s$ one has $\|A_S\| \leq \sqrt{1 + \frac{1}{3}}$ and $\|(A_S^T A_S)^{-1}\| \leq (1 - \frac{1}{3})^{-1} = \frac{3}{2}$, where $\|\cdot\|$ denotes the operator norm $\|B\| = \max_{\|x\|=1} \|Bx\|$.

This means that if we take A random with i.i.d. $\mathcal{N}(0, \frac{1}{m})$ entries then, for any $|S| \leq s$ we have that

$$\|A_S (A_S^T A_S)^{-1} \text{sign}(x_S)\| \leq \sqrt{1 + \frac{1}{3}} \frac{3}{2} \sqrt{s} = \sqrt{3} \sqrt{s},$$

and because of the independency among the entries of A , A_{S^c} is independent of this vector and so for each $j \in S^c$ we have

$$\mathbb{P}\left(\left|A_j^T A_S (A_S^T A_S)^{-1} \text{sign}(x_S)\right| \geq \frac{1}{\sqrt{M}} \sqrt{3} \sqrt{st}\right) \leq 2 \exp\left(-\frac{t^2}{2}\right),$$

⁸Note how this differs from the situation in Chapter 7 where the linear inequalities were enough to determine a unique candidate for a dual certificate.

⁹Note that the normalization here is taken slightly differently: entries are normalized by $\frac{1}{\sqrt{m}}$, rather than $\frac{1}{a_m}$, but the difference is negligible for our purposes.

where A_j is the j -th column of A .

An application of the union bound gives

$$\mathbb{P} \left(\left\| A_{S^c}^T A_S (A_S^T A_S)^{-1} \text{sign}(x_S) \right\|_\infty \geq \frac{1}{\sqrt{M}} \sqrt{3} \sqrt{st} \right) \leq 2N \exp \left(-\frac{t^2}{2} \right),$$

which implies

$$\begin{aligned} \mathbb{P} \left(\left\| A_{S^c}^T A_S (A_S^T A_S)^{-1} \text{sign}(x_S) \right\|_\infty \geq 1 \right) &\leq 2p \exp \left(-\frac{\left(\frac{\sqrt{m}}{\sqrt{3s}} \right)^2}{2} \right) \\ &= \exp \left(-\frac{1}{2} \left[\frac{m}{3s} - 2 \log(2p) \right] \right), \end{aligned}$$

which means that we expect to exactly recover x via ℓ_1 minimization when $m \gg s \log(p)$. While this can be asymptotically worse than the bound of $m \gtrsim s \log \left(\frac{p}{s} \right)$, and this guarantee is not uniformly obtained for all sparse vectors, the technique in this section is generalizable to many circumstances and illustrates the flexibility of approaches based in construction of dual witnesses. We will discuss nonuniform guarantees in more detail in the next section.

10.5 Compressive sensing: from theory to practice

10.5.1 Sensing matrices and incoherence

In applications, we usually cannot completely freely choose the sensing matrix to our liking. This means that Gaussian random matrices play an important role as benchmark, but from a practical viewpoint they play a marginal role. Clearly, randomness in the sensing matrix seems to be very beneficial for compressive sensing. However, in practice, there are many design constraints on the sensing matrix A , as in many applications one only has access to structured measurement systems. For example, we may still have the freedom to choose, say the positions of the antennas in radar systems that employ multiple antennas [163, 164], the position of sensors in MRI [119, 5, 39, 138], or the sampling locations in digital signal acquisition [128]. By choosing these randomly, we can still introduce randomness in our system. Or, we can transmit random waveforms in sonar and radar systems [92, 146]. Yet, in all these cases the overall structure of A is still dictated by the physics of wave propagation. In other applications, it will be other physical constraints or design limitations that will dictate how much randomness we can introduce into the sensing matrix.

While establishing the RIP for Gaussian or Bernoulli random matrices is not too difficult, it is already significantly harder to do so for the partial Fourier matrix [49, 153, 91] and time-frequency matrices [70], and even harder for more specific sensing matrices.

A useful concept to overcome the practical limitations of the RIP is via the concept of the (in)coherence of a matrix. This concept has proven to be widely applicable in practice. While we want to avoid the constraints of the RIP, we nevertheless take it as our point of departure. Recall that the RIP (Definition 10.7) asks that any $S \subset [p]$, $|S| \leq s$ satisfies:

$$(1 - \delta)\|x\|^2 \leq \|A_S x\|^2 \leq (1 + \delta)\|x\|^2,$$

for all $x \in \mathbb{R}^{|S|}$. This is equivalent to

$$\max_x \frac{x^T (A_S^T A_S - I) x}{x^T x} \leq \delta,$$

or equivalently

$$\|A_S^T A_S - I\| \leq \delta.$$

If the columns of A are unit-norm vectors (in $\mathbb{R}^m$), then the diagonal of $A_S^T A_S$ is all-ones, this means that $A_S^T A_S - I$ consists only of the non-diagonal elements of $A_S^T A_S$. If, moreover, for any two columns a_i, a_j , of A we have $|a_i^T a_j| \leq \mu$ for some μ then, Gershgorin's circle theorem tells us that $\|A_S^T A_S - I\| \leq \mu(s - 1)$.

More precisely, given a symmetric matrix B , the Gershgorin's Circle Theorem [93] states that all of the eigenvalues of B are contained in the so called Gershgorin discs (for each i , the Gershgorin disc corresponds to $\{\lambda : |\lambda - B_{ii}| \leq \sum_{j \neq i} |B_{ij}|\}$). If B has zero diagonal, then this reads: $\|B\| \leq \sum_{j \neq i} |B_{ij}|$.

Given a set of p unit-norm vectors $a_1, \dots, a_p \in \mathbb{R}^m$ we define its worst-case coherence μ as

$$\mu(a_1, \dots, a_p) = \max_{i \neq j} |\langle a_i, a_j \rangle|. \quad (10.25)$$

Given a set of unit-norm vectors $a_1, \dots, a_p \in \mathbb{R}^m$ with worst-case coherence μ , if we form a matrix with these vectors as columns, then it will be $(s, \mu(s - 1))$ -RIP, meaning that it will be $(s, \frac{1}{3})$ -RIP for $s \leq \frac{1}{\frac{1}{3} - \mu}$.

This motivates the problem of designing sets of vectors $a_1, \dots, a_p \in \mathbb{R}^m$ with smallest possible worst-case coherence. This is a central problem in Frame Theory [162, 55, 52, 188]. Recall that in finite dimensions, a set of vectors $a_1, \dots, a_p \in \mathbb{H}^m$ (where $\mathbb{H} = \mathbb{R}$ or $\mathbb{C}$) is called a frame for $\mathbb{H}^m$ if there exist constants (frame bounds) $0 < A \leq B < \infty$ such that¹⁰

$$A\|x\|^2 \leq \sum_{k=1}^p |\langle x, a_k \rangle|^2 \leq B\|x\|^2 \quad (10.26)$$

for every $x \in \mathbb{H}^m$. The associated *frame operator* S is defined by

¹⁰Since here we are dealing with a finite dimensional Hilbert space, the upper frame bound is always trivially fulfilled.

$$Sx = \sum_{i=1}^p \langle x, a_i \rangle a_i. \quad (10.27)$$

The frame definition (10.26) can be equivalently expressed as requiring that $AI \preceq S \preceq BI$ holds, where $\preceq$ denotes the positive-semidefinite order. A frame is called *tight* if $A = B$, in which case $S = AI$. If $\|a_i\| = 1$ for all $i = 1, \dots, p$ then $\{a_k\}_{k=1}^p$ is called a *unit norm frame*. We call a unit norm frame $\{a_k\}_{k=1}^p$ *equiangular* if

$$|\langle a_i, a_j \rangle| = c \quad \text{for all } i, j \text{ with } i \neq j, \quad (10.28)$$

for some constant $c \geq 0$. Obviously, any orthonormal basis is equiangular.

We can now provide an elucidating answer to the question of a lower bound on the coherence of a set of vectors $a_1, \dots, a_p$ via the following theorem from frame theory.

Theorem 10.15 ([162]). *Let $\{a_k\}_{k=1}^p$ be a unit-norm frame for $\mathbb{H}^m$, where $\mathbb{H} = \mathbb{R}$ or $\mathbb{C}$. Then*

$$\mu(a_1, \dots, a_p) \geq \sqrt{\frac{p-m}{m(p-1)}}. \quad (10.29)$$

Equality holds in (10.29) if and only if $\{a_k\}_{k=1}^p$ is an equiangular tight frame.

Proof. A simple calculation shows that

$$\text{trace}(S) = \|S^{\frac{1}{2}}\|_F^2 = \sum_{k=1}^p \|a_k\|^2 = p, \text{ and } \text{trace}(S^2) = \|S\|_F^2 = \sum_{k,l=1}^p |\langle a_k, a_l \rangle|^2,$$

Let $\lambda_1 \geq \dots \geq \lambda_m$ be the eigenvalues of S . Let $\mathbf{1}$ denote the all-one vector. By the Cauchy-Schwarz inequality

$$\text{trace}(S)^2 = \left(\sum_{k=1}^m \lambda_k \right)^2 = \langle \mathbf{1}, \{\lambda_k\} \rangle^2 \leq \|\mathbf{1}\|^2 \|\{\lambda_k\}\|^2 = m \sum_{k=1}^m \lambda_k^2 = m \text{trace } S^2,$$

which can be stated equivalently as

$$\sum_k \sum_{l=1}^p |\langle a_k, a_l \rangle|^2 \geq \frac{p^2}{m}. \quad (10.30)$$

Equality holds in (10.30) if and only if $\lambda_1 = \dots = \lambda_p = p/m$, i.e., if and only if $\{a_k\}_{k=1}^p$ is a tight frame

We now consider

$$\max_{k \neq l} |\langle a_k, a_l \rangle|^2 \geq \frac{1}{p^2 - p} \sum_{k \neq l} |\langle a_k, a_l \rangle|^2 = \frac{1}{p^2 - p} \left(\sum_{k=1}^p \sum_{l=1}^p |\langle a_k, a_l \rangle|^2 - p \right), \quad (10.31)$$

with equality if and only if $\{a_k\}_{k=1}^p$ is equiangular. By the bound (10.30) we have

$$\frac{1}{p^2 - p} \left(\sum_{k=1}^p \sum_{l=1}^p |\langle a_k, a_l \rangle|^2 - p \right) \geq \frac{1}{p^2 - p} \left(\frac{p^2}{m} - p \right) = \frac{p - m}{m(p - 1)}, \quad (10.32)$$

with equality if and only if $\{a_k\}_{k=1}^p$ is a tight frame. Inequalities (10.31) and (10.32) together give (10.29) and the proof is complete. $\square$

We note that if $\mathbb{H} = \mathbb{R}$ equality in (10.29) can only hold if $p \leq m(m+1)/2$, while if $\mathbb{H} = \mathbb{C}$ then equality in (10.29) can only hold if $p \leq m^2$. These statements follow from the bounds in Table II of [61].

The coherence bound (10.29) of a set of p unit-norm vectors in m dimensions is also known as Welch bound [189]. Due to this limitation implied by Theorem 10.15, deterministic constructions based on coherence cannot yield matrices that satisfy the RIP for $s \gg \sqrt{m}$, known as the square-root bottleneck [25, 168].

However, as we will see in Section 10.5.2, if we are willing to accept slightly weaker recovery guarantees than provided by the RIP, matrices with low coherence do give rise to sensing matrices that come with appealing theoretical guarantees while also being useful in practice.

Theorem 10.15 suggests that in order to design sensing vectors of minimal coherence, we should look for equiangular tight frames (ETFs). The existence and construction of ETFs is a rather challenging problem that intersects with a wide range of areas such as harmonic analysis, algebraic, combinatorics, finite geometry, group theory and representation theory, and coding theory [162, 188]. Applications include signal processing, wireless communications, numerical linear algebra, and quantum physics. Indeed, ETFs are deeply connected to Zauner's conjecture, particularly through their role in quantum information theory and the concept of SIC-POVMs (Symmetric Informationally Complete Positive Operator-Valued Measures) [191, 155, 14].

Another interesting and very useful construction of sensing vectors with low coherence is given by so-called *mutually unbiased bases* [190, 72, 32, 45, 162]. These bases are widely used in quantum information theory as well as in radar and signal processing. Two orthonormal bases $\{\varphi_j\}_{j=0}^{m-1}, \{\psi_j\}_{j=0}^{m-1} \in \mathbb{C}^m$ are called *mutually unbiased* if and only if

$$\forall i, j \quad |\langle \varphi_i, \psi_j \rangle| = \frac{1}{\sqrt{m}}. \quad (10.33)$$

The classical example of an MUB consists of the identity matrix and the Discrete Fourier transform matrix. But, amazingly, for certain values of m , there exist MUBs in $\mathbb{C}^m$ consisting of $m+1$ orthonormal bases $\{\varphi_j^{(k)}\}_{j=0}^{m-1}, k = 0, \dots, m$ such that condition (10.33) is satisfied for any two orthonormal bases $\{\varphi_j^{(i)}\}_{j=0}^{m-1}, \{\varphi_j^{(k)}\}_{j=0}^{m-1}$, see for example [190, 32, 45, 162, 7]. Thus, MUBs give rise to structured sensing matrices of dimension $m \times m^2$ with coherence $\mu = 1/\sqrt{m}$, which thus almost achieve the Welch bound. Here is a simple example for an MUB in $\mathbb{C}^m$. Let m be a prime number ≥ 5 and set $g(j) = \frac{1}{\sqrt{m}} e^{2\pi i j^3/m}$ for $j = 0, \dots, m-1$. Then the vectors $\{g_{k,l}\}_{k,l=0}^{m-1}$ defined via

$$g_{k,l}(j) = g(j-k)e^{2\pi ilj/m}, \quad k, l = 0, \dots, m-1, \quad (10.34)$$

satisfy $|\langle g_{k,l}, g_{k',l'} \rangle| \in \{0, 1/\sqrt{m}\}$ for all $g_{k,l} \neq g_{k',l'}$, which follows from basic properties of Gaussian sums (cf. Theorem 2 in [7]). We can add the $m \times m$ identity matrix and end up with $m^2 + m$ vectors (split into $m+1$ ONBs) in $\mathbb{C}^m$ which form a MUB of maximal size. For other, nonequivalent constructions, of MUBs (which, however, still revolve around translations and modulations in form of the finite Heisenberg group) see [45]. Remarkably, for m not a power of a prime number it is not known how many MUB exist; this is an open problem even for $m = 6$ [31]. For $\mathbb{R}^m$ the construction and existence of MUBs faces more constraints. For example, when m is of the form $m = 4^k$, there exist $\frac{m+1}{2}$ MUBs in $\mathbb{R}^m$, see e.g. [45].

10.5.2 Nonuniform recovery guarantees and coherence

As mentioned before, it follows from Theorem 10.15 that deterministic constructions based on coherence cannot yield matrices that satisfy the RIP for $s \gg \sqrt{M}$, known as the square-root bottleneck [25, 168]. To overcome this square root bottleneck something has to give. One fruitful direction is to sacrifice the *uniform recovery* granted by the RIP. Namely, once a matrix satisfies the RIP, it is guaranteed that the solution to the ℓ_1 -optimization problem is identical to the solution to the ℓ_0 -optimization problem (the sparsest solution) for all s -sparse vectors. In contrast, we can consider scenarios in which this equivalence between ℓ_0 -optimization and ℓ_1 optimization holds “only” for *most* s -sparse vectors. This leads to *nonuniform recovery* results, which we will pursue below. The benefits are worth the sacrifice, since we end up with theoretical guarantees that are much more practical.

Recall that we consider a general linear system of equations $Ax = y$, where $A \in \mathbb{C}^{m \times p}$, $x \in \mathbb{C}^p$ and $m \leq p$. We introduce the following generic s -sparse model:

- (i) The support $S \subset \{1, \dots, p\}$ of the s nonzero coefficients of x is selected uniformly at random.
- (ii) The non-zero entries of $\text{sign}(x)$ form a Steinhaus sequence, i.e., $\text{sign}(x_k) := x_k/|x_k|$, $k \in I_s$, is a complex random variable that is uniformly distributed on the unit circle, and independent from each other (and from I_s).

As an example for a theoretical nonuniform guarantee using coherence-based sensing matrices we state (without proof) the following theorem, and refer to [77] for a proof.

Theorem 10.16 (Theorem 14.5 in [77]). *Let $A \in \mathbb{C}^{m \times p}$, $m \leq p$, be a matrix with ℓ_2 -normalized columns and coherence μ . Let S be a subset of $\{1, \dots, p\}$ selected at random according to the uniform model with $\text{card}(S) = s$. Let $x \in \mathbb{C}^p$ be a vector supported on S for which $\text{sign}(x_S)$ is a Steinhaus sequence independent of S . Assume that, for $\eta, \varepsilon \in (0, 1)$*

$$\eta \leq \frac{c}{\ln(p/\varepsilon)} \quad (10.35)$$

$$\frac{s}{p} \|A\|^2 \leq \frac{c}{\ln(p/\varepsilon)} \quad (10.36)$$

for an appropriate constant $c > 0$. Then, with probability at least, the vector x is the unique minimizer of $\|z\|_1$ subject to $Az = Ax$.

A simple calculation shows that for a unit-norm tight frame the condition 10.35 is satisfied if

$$m \geq Cs \ln(p/\varepsilon),$$

which is in the ballpark of condition (10.19) for Gaussian sensing matrices. Theorem 10.16 demonstrates that by allowing for nonuniform recovery guarantees one can successfully reconstruct sparse signals under much milder coherence conditions than imposed by the aforementioned square root bottleneck. That being said, establishing that the support set of real-world signals can be modeled as random is, in general, difficult to justify rigorously. For instance, the wavelet coefficients of natural images are typically concentrated along edges, giving rise to structured, tree-like patterns in the locations of significant coefficients. This insight is utilized in practical compressive sensing systems as, for example, in MRI [119, 5]. The fact that the support set of many natural signals does not conform to a uniform distribution highlights the relevance of the recovery results established in Sections 10.1–10.3, which hold for *all* sufficiently sparse signals. Furthermore, as pointed out in [77], the stability guarantees currently available for models based on random supports are notably weaker than those derived from the restricted isometry property.

Various other versions of *nonuniform recovery* results can be found e.g., in [174, 50, 77]. See [161, 95, 5] for some theoretical results geared towards applications.

10.5.3 Further practical considerations

When moving the ideas of compressive sensing to move from theory into practice we encounter some challenges: (i) the choice of the sensing matrix is often limited by practical considerations and physical constraints (both usually rule out a Gaussian sensing matrix); (ii) signals are often not sparse in the measurement domain, but only with respect to some properly chosen transform (and even then, signals are often only approximately sparse); (iii) for various reasons (speed being one of them) we may need alternative sparse solvers instead of vanilla ℓ_1 -minimization. We already have discussed (i) earlier in this section and will discuss items (ii) and (iii) succinctly below. We recommend [5] for a detailed treatment of how to span the gap between theory and practice of compressive sensing.

As indicated above, signals are often not sparse in the canonical basis, but they are (approximately) sparse after a suitable transformation. For examples

audio signals are sparse with respect to a localized Fourier transform such as a Gabor transform [59] or some other local trigonometric transform [121, 120]. JPEG exploits the sparsity of images in the discrete cosine basis or wavelet basis [120] while curvelet transforms are effective sparsifiers for images arising in astronomy [159]. Radar signals are sparse when (properly) transformed into the time-frequency domain [92]. Some signals may even require a redundant dictionary (perhaps consisting of a combination of a wavelet basis, a curvelet basis, and a Fourier basis), see e.g. [120, Chapter 12]. But how do such transform-sparse signals fit into our framework?

As before, let $x \in \mathbb{C}^p$ be the signal (or image) of interest and we observe $y = Ax$, where $A \in \mathbb{C}^{m \times p}$ is a sensing matrix, where $m < p$. We assume that x is sparse with respect to some basis or frame $U \in \mathbb{C}^{p \times n}$, where $n \geq p$ (and $n = p$ if we restrict U to be a basis). Denote this sparse representation by $w \in \mathbb{C}^n$, i.e., $x = Uw$. We write $B := AU$ and can now solve for w via

$$\min \|z\|_1 \text{ s.t. } Bz = y. \quad (10.37)$$

Having solved this problem (assuming it does indeed yield w), we can recover x from w by simply computing $x = Uw$.

Thus, we can (try to) apply the theoretical results developed in the previous sections to the matrix $B = AU$ in place of A . In a nutshell, the sensing matrix A has to be *incoherent* with respect to the transform in which the signal is sparse. An important example arises in MRI. There, the wavelet transform is often chosen as sparsifying transform and the sensing matrix is derived from randomly subsampling the two-dimensional Discrete Fourier Transform or the Radon transform, see [119, 5]. In practice, it has been observed that even better results are obtained when the random subsampling of the Fourier coefficients is combined with a deterministic sampling of the Fourier coefficients corresponding to the low frequencies, a strategy that has already been suggested in one of the original papers on compressive sensing [66]. We refer to [119, 5] for many more details and modifications.

Figure 10.4 shows an example of compressive sensing applied to MRI. The test image is the so-called GLPU phantom (introduced by Guerquin-Kern, Lejeune, Pruessmann and Unser in [85]). It is a continuous, piecewise constant MRI-type image with an analytic expression for its Fourier transform. The reconstruction of the discrete GLPU phantom from 30% randomly chosen discrete Fourier measurements. The wavelet transform with a Haar wavelet has been used as sparsifying transform U . The image x is recovered by first solving (10.37) for w , followed by computing $x = Uw$.

There are various other efficient and rigorous methods to recover sparse vectors from underdetermined systems besides ℓ_1 -minimization and the lasso. For example, homotopy methods, greedy algorithms or methods based on approximate message passing. We refer to [77, 5] for a comprehensive discussion of these techniques. Furthermore, [5] contains a detailed discussion of some subtle potential numerical stability issues one should be aware of.

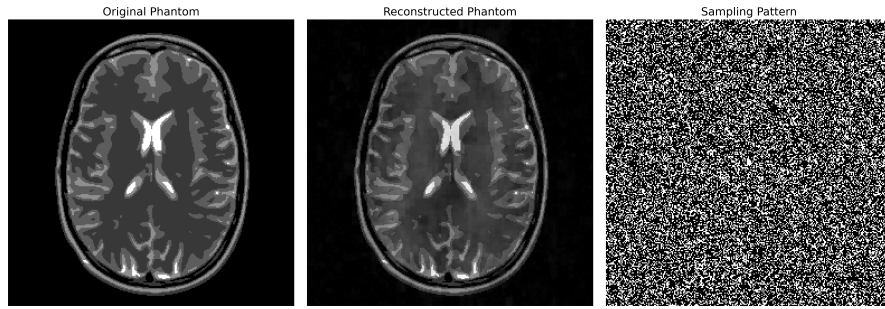


Fig. 10.4: The GLPU MRI phantom (left), its reconstruction (middle) via ℓ_1 minimization from 30% discrete randomly chosen Fourier measurements using the wavelet transform as sparsifying transform U . The sampling pattern is depicted on the right. [The authors want to thank Yuan Ni for providing the code to run this numerical simulation.]

Bibliography

1. E. Abbe, A. S. Bandeira, and G. Hall. Exact recovery in the stochastic block model. *IEEE Transactions on Information Theory*, 62(1):417–487, 2016. (Cited on pp. 123, 132.)
2. E. Abbe, J. Fan, and Y. Wang, K. and Zhong. Entrywise eigenvector analysis of random matrices with low expected rank. *Available online at arXiv:1709.09565*, 2017. (Cited on p. 133.)
3. E. Abbe and C. Sandon. Detection in the stochastic block model with multiple clusters: proof of the achievability conjectures, acyclic bp, and the information-computation gap. *Available online at arXiv:1512.09080 [math.PR]*, 2015. (Cited on p. 123.)
4. Emmanuel Abbe. Community detection and stochastic block models. *Foundations and Trends® in Communications and Information Theory*, 14(1-2):1–162, 2018. (Cited on p. 123.)
5. Ben Adcock and Anders C Hansen. *Compressive imaging: structure, sampling, learning*. Cambridge University Press, 2021. (Cited on pp. 184, 189, 189, 189, 190, 190, 190, 190.)
6. Nir Ailon and Bernard Chazelle. The fast Johnson-Lindenstrauss transform and approximate nearest neighbors. *SIAM J. Comput*, pages 302–322, 2009. (Cited on p. 99.)
7. W Alltop. Complex sequences with low periodic correlations (Corresp.). *IEEE Transactions on Information Theory*, 26(3):350–354, 1980. (Cited on pp. 187, 188.)
8. N. Alon. Eigenvalues and expanders. *Combinatorica*, 6:83–96, 1986. (Cited on p. 70.)
9. N. Alon. Problems and results in extremal combinatorics i. *Discrete Mathematics*, 273(1–3):31–53, 2003. (Cited on p. 99.)
10. N. Alon and V. Milman. Isoperimetric inequalities for graphs, and superconcentrators. *Journal of Combinatorial Theory*, 38:73–88, 1985. (Cited on p. 70.)
11. D. Amelunxen, M. Lotz, M. B. McCoy, and J. A. Tropp. Living on the edge: phase transitions in convex programs with random data. *Information and Inference, available online*, 2014. (Cited on pp. 181, 181.)
12. Dennis Amelunxen, Martin Lotz, Michael B McCoy, and Joel A Tropp. Living on the edge: Phase transitions in convex programs with random data. *Infor-*

- mation and Inference: A Journal of the IMA*, 3(3):224–294, 2014. (Cited on pp. 179, 179.)
13. G. W. Anderson, A. Guionnet, and O. Zeitouni. *An introduction to random matrices*. Cambridge studies in advanced mathematics. Cambridge University Press, Cambridge, New York, Melbourne, 2010. (Cited on p. 49.)
 14. Marcus Appleby, Ingemar Bengtsson, Steven Flammia, and Dardo Goyeneche. Tight frames, Hadamard matrices and Zauner’s conjecture. *Journal of Physics A: Mathematical and Theoretical*, 52(29):295301, 2019. (Cited on p. 187.)
 15. S. Arora, B. Barak, and D. Steurer. Subexponential algorithms for unique games related problems. 2010. (Cited on p. 115.)
 16. Z. D. Bai. Methodologies in spectral analysis of large dimensional random matrices, a review. *Statistics Sinica*, 9:611–677, 1999. (Cited on p. 49.)
 17. Zhidong Bai and Jack W Silverstein. *Spectral analysis of large dimensional random matrices*, volume 20. Springer, 2010. (Cited on p. 49.)
 18. J. Baik, G. Ben-Arous, and S. Péché. Phase transition of the largest eigenvalue for nonnull complex sample covariance matrices. *The Annals of Probability*, 33(5):1643–1697, 2005. (Cited on pp. 51, 51.)
 19. J. Baik and J. W. Silverstein. Eigenvalues of large sample covariance matrices of spiked population models. 2005. (Cited on p. 51.)
 20. A. S. Bandeira. Random Laplacian matrices and convex relaxations. *Available online at arXiv:1504.03987 [math.PR]*, 2015. (Cited on pp. 132, 133.)
 21. A. S. Bandeira. A note on probably certifiably correct algorithms. *Comptes Rendus Mathématique*, to appear, 2016. (Cited on pp. 127, 133.)
 22. A. S. Bandeira, M. Boedihardjo, and R. van Handel. Matrix concentration inequalities and free probability. *Inventiones Mathematicae*, (234):419–487, 2023. (Cited on pp. 157, 163, 164, 164, 164, 165, 165.)
 23. A. S. Bandeira, E. Dobriban, D.G. Mixon, and W.F. Sawin. Certifying the restricted isometry property is hard. *IEEE Trans. Inform. Theory*, 59(6):3448–3450, 2013. (Cited on p. 180.)
 24. A. S. Bandeira, M. Fickus, D. G. Mixon, and J. Moreira. Derandomizing restricted isometries via the Legendre symbol. *Available online at arXiv:1406.4089 [math.CO]*, 2014. (Cited on p. 180.)
 25. A. S. Bandeira, M. Fickus, D. G. Mixon, and P. Wong. The road to deterministic matrices with the restricted isometry property. *Journal of Fourier Analysis and Applications*, 19(6):1123–1149, 2013. (Cited on pp. 180, 180, 187, 188.)
 26. A. S. Bandeira, M. E. Lewis, and D. G. Mixon. Discrete uncertainty principles and sparse signal processing. *Available online at arXiv:1504.01014 [cs.IT]*, 2015. (Cited on p. 180.)
 27. A. S. Bandeira, D. G. Mixon, and J. Moreira. A conditional construction of restricted isometries. *Available online at arXiv:1410.6457 [math.FA]*, 2014. (Cited on pp. 180, 180.)
 28. A. S. Bandeira and R. v. Handel. Sharp nonasymptotic bounds on the norm of random matrices with independent entries. *Annals of Probability*, to appear, 2015. (Cited on pp. 132, 156, 156.)
 29. Afonso S. Bandeira. Ten lectures and forty-two open problems in the mathematics of data science. *Available online at: <https://people.math.ethz.ch/~abandeira/TenLecturesFortyTwoProblems.pdf>*, 2016. (Cited on p. 3.)

30. Afonso S. Bandeira, Sivakanth Gopi, Haotian Jiang, Kevin Lucca, and Thomas Rothvoss. A geometric perspective on the injective norm of sums of random tensors. *Available online at arXiv:2411.10633 [math.PR]*, 2024. (Cited on p. 159.)
31. Afonso S. Bandeira, Anastasia Kireeva, Antoine Maillard, and Almut Rödder. Randomstrasse101: Open problems of 2024. *arXiv preprint arXiv:2504.20539*, 2025. (Cited on pp. 3, 188.)
32. Bandyopadhyay, Boykin, Roychowdhury, and Vatan. A new proof for the existence of mutually unbiased bases. *Algorithmica*, 34:512–528, 2002. (Cited on pp. 187, 187.)
33. B. Barak. Sum of squares upper bounds, lower bounds, and open questions. *Available online at <http://www.boazbarak.org/sos/files/all-notes.pdf>*, 2014. (Cited on pp. 116, 117.)
34. B. Barak and D. Steurer. Sum-of-squares proofs and the quest toward optimal algorithms. *Survey, ICM 2014*, 2014. (Cited on p. 116.)
35. Richard Bellman. *Dynamic programming*. Princeton University Press, Princeton, 1957. (Cited on p. 9.)
36. Adi Ben-Israel and Thomas NE Greville. *Generalized inverses: theory and applications*. Springer, 2003. (Cited on p. 41.)
37. F. Benaych-Georges and R. R. Nadakuditi. The eigenvalues and eigenvectors of finite, low rank perturbations of large random matrices. *Advances in Mathematics*, 2011. (Cited on p. 51.)
38. F. Benaych-Georges and R. R. Nadakuditi. The singular values and vectors of low rank perturbations of large rectangular random matrices. *Journal of Multivariate Analysis*, 2012. (Cited on p. 51.)
39. Kai Tobias Block, Martin Uecker, and Jens Frahm. Undersampled radial MRI with multiple coils. iterative image reconstruction using a total variation constraint. *Magnetic Resonance in Medicine: An Official Journal of the International Society for Magnetic Resonance in Medicine*, 57(6):1086–1098, 2007. (Cited on p. 184.)
40. March Boedihardjo, Shaofeng Deng, and Thomas Strohmer. A performance guarantee for spectral clustering. *SIAM Journal on Mathematics of Data Science*, 3(1):369–387, 2021. (Cited on p. 133.)
41. J. Bourgain et al. Explicit constructions of RIP matrices and related problems. *Duke Mathematical Journal*, 159(1), 2011. (Cited on pp. 180, 180.)
42. Tatiana Brailovskaya and Ramon van Handel. Universality and sharp matrix concentration inequalities. *Geometric and Functional Analysis*, 34(6):1734–1838, 2024. (Cited on pp. 157, 164, 164, 165.)
43. K. Bryan and T. Leise. The \$25,000,000,000 eigenvector: The linear algebra behind google. *Siam Review*, 48(3):569–581, 2006. (Cited on p. 57.)
44. Artur Buchholz. Operator khintchine inequality in non-commutative probability. *Mathematische Annalen*, 319, 2001. (Cited on pp. 163, 163.)
45. A Robert Calderbank, Peter J Cameron, William M Kantor, and Jaap J Seidel. F_4 -Kerdock codes, orthogonal spreads, and extremal Euclidean line-sets. *Proceedings of the London Mathematical Society*, 75(2):436–480, 1997. (Cited on pp. 187, 187, 188, 188.)
46. E. J. Candès, J. Romberg, and T. Tao. Robust uncertainty principles: exact signal reconstruction from highly incomplete frequency information. *IEEE Trans. Inform. Theory*, 52:489–509, 2006. (Cited on p. 169.)

47. E. J. Candès, J. Romberg, and T. Tao. Stable signal recovery from incomplete and inaccurate measurements. *Comm. Pure Appl. Math.*, 59:1207–1223, 2006. (Cited on p. 169.)
48. E. J. Candès and T. Tao. Decoding by linear programming. *IEEE Trans. Inform. Theory*, 51:4203–4215, 2005. (Cited on p. 169.)
49. E. J. Candès and T. Tao. Near optimal signal recovery from random projections: universal encoding strategies? *IEEE Trans. Inform. Theory*, 52:5406–5425, 2006. (Cited on pp. 169, 180, 184.)
50. E.J. Candès and Y. Plan. Near-ideal model selection by ℓ_1 minimization. *Annals of Statistics*, 37(5A):2145–2177, 2009. (Cited on pp. 178, 189.)
51. E.J. Candès. The restricted isometry property and its implications for compressed sensing. *Comptes Rendus Mathématique*, 346(9):589–592, 2008. (Cited on pp. 176, 176.)
52. P. G. Casazza and G. Kutyniok. *Finite Frames: Theory and Applications*. 2012. (Cited on p. 185.)
53. V. Chandrasekaran, B. Recht, P.A. Parrilo, and A.S. Willsky. The convex geometry of linear inverse problems. *Foundations of Computational Mathematics*, 12(6):805–849, 2012. (Cited on p. 181.)
54. J. Cheeger. A lower bound for the smallest eigenvalue of the Laplacian. *Problems in analysis (Papers dedicated to Salomon Bochner, 1969)*, pp. 195–199. Princeton Univ. Press, 1970. (Cited on p. 70.)
55. O. Christensen. *An introduction to frames and Riesz bases*. Birkhäuser, Boston, 2003. (Cited on p. 185.)
56. F. Chung. Four proofs for the cheeger inequality and graph partition algorithms. *Fourth International Congress of Chinese Mathematicians*, pp. 331–349, 2010. (Cited on p. 72.)
57. Michael B Cohen, Jelani Nelson, and David P Woodruff. Optimal approximate matrix product in terms of stable rank. *ICALP 2016*, pages 11:1–11:14, 2016. (Cited on p. 99.)
58. S. Dasgupta and A. Gupta. An elementary proof of the johnson-lindenstrauss lemma. Technical report, 2002. (Cited on pp. 96, 96, 96.)
59. Laurent Daudet and Bruno Torrèsani. Sparse adaptive representations for musical signals. In *Signal processing methods for music transcription*, pages 65–98. Springer, 2006. (Cited on p. 190.)
60. A. Decelle, F. Krzakala, C. Moore, and L. Zdeborová. Asymptotic analysis of the stochastic block model for modular networks and its algorithmic applications. *Phys. Rev. E*, 84, December 2011. (Cited on pp. 122, 123.)
61. P. Delsarte, J. M. Goethals, and J. J. Seidel. Bounds for systems of lines and Jacobi polynomials. *Philips Res. Repts*, 30(3):91*–105*, 1975. Issue in honour of C.J. Bouwkamp. (Cited on p. 187.)
62. Shaofeng Deng, Shuyang Ling, and Thomas Strohmer. Strong consistency, graph Laplacians, and the Stochastic Block Model. *arXiv preprint arXiv:2004.09780*, 2020. (Cited on p. 133.)
63. Shaofeng Deng, Shuyang Ling, and Thomas Strohmer. Strong consistency, graph laplacians, and the stochastic block model. *Journal of Machine Learning Research*, 22(117):1–44, 2021. (Cited on p. 133.)
64. E Dobriban. Permutation methods for factor analysis and pca. *arXiv preprint arXiv:1710.00479*, 2017. (Cited on p. 55.)

65. Edgar Dobriban and Art B Owen. Deterministic parallel analysis: an improved method for selecting factors and principal components. *Journal of the Royal Statistical Society: Series B (Statistical Methodology)*, 2018. (Cited on p. 55.)
66. D. L. Donoho. Compressed sensing. *IEEE Trans. Inform. Theory*, 52:1289–1306, 2006. (Cited on pp. 169, 190.)
67. David Donoho and Jared Tanner. Observed universality of phase transitions in high-dimensional geometry, with implications for modern data analysis and signal processing. *Philosophical Transactions of the Royal Society A: Mathematical, Physical and Engineering Sciences*, 367(1906):4273–4293, 2009. (Cited on pp. 179, 179.)
68. David L Donoho et al. High-dimensional data analysis: The curses and blessings of dimensionality. *AMS math challenges lecture*, 1(32):375, 2000. (Cited on p. 22.)
69. David L Donoho, Matan Gavish, and Iain M Johnstone. Optimal shrinkage of eigenvalues in the spiked covariance model. *Annals of statistics*, 46(4):1742, 2018. (Cited on p. 56.)
70. Dominik Dorsch and Holger Rauhut. Refined analysis of sparse mimo radar. *Journal of Fourier Analysis and Applications*, 23(3):485–529, 2017. (Cited on p. 184.)
71. Rick Durrett. *Probability: theory and examples*, volume 49. Cambridge university press, 2019. (Cited on pp. 17, 23.)
72. Thomas Durt, Berthold-Georg Englert, Ingemar Bengtsson, and Karol Życzkowski. On mutually unbiased bases. *International journal of quantum information*, 8(04):535–640, 2010. (Cited on p. 187.)
73. E. N. Epperly. Blog post: “note to self: Norm of a gaussian random vector”. 2023. (Cited on p. 147.)
74. Li Feng, Thomas Benkert, Kai Tobias Block, Daniel K Sodickson, Ricardo Otazo, and Hersh Chandarana. Compressed sensing for body MRI. *Journal of Magnetic Resonance Imaging*, 45(4):966–987, 2017. (Cited on p. 170.)
75. D. Féral and S. Péché. The largest eigenvalue of rank one deformation of large wigner matrices. *Communications in Mathematical Physics*, 272(1):185–228, 2006. (Cited on p. 54.)
76. Noah Fleming, Pravesh Kothari, and Toniann Pitassi. Semialgebraic proofs and efficient algorithm design. *Foundations and Trends in Theoretical Computer Science*, 14(1-2):1–221, 2019. (Cited on p. 117.)
77. S. Foucart and H. Rauhut. *A Mathematical Introduction to Compressive Sensing*. Birkhauser, 2013. (Cited on pp. 169, 170, 171, 178, 188, 188, 189, 189, 190.)
78. J. J. Fuchs. On sparse representations in arbitrary redundant bases. *Information Theory, IEEE Transactions on*, 50(6):1341–1344, 2004. (Cited on p. 181.)
79. Amir Ghasemian, Pan Zhang, Aaron Clauset, Cristopher Moore, and Leto Peel. Detectability thresholds and optimal algorithms for community structure in dynamic networks. *Available online at arXiv:1506.06179 [stat.ML]*, 2015. (Cited on p. 123.)
80. M. X. Goemans and D. P. Williamson. Improved approximation algorithms for maximum cut and satisfiability problems using semidefinite programming. *Journal of the Association for Computing Machinery*, 42:1115–1145, 1995. (Cited on pp. 111, 114.)
81. G. H. Golub. *Matrix Computations*. Johns Hopkins University Press, third edition, 1996. (Cited on pp. 33, 46.)

82. Gene H Golub and Charles F Van Loan. *Matrix computations*. JHU press, 2013. (Cited on p. 40.)
83. Y. Gordon. Some inequalities for gaussian processes and applications. *Israel J. Math*, 50:109–110, 1985. (Cited on p. 145.)
84. Y. Gordon. On milnan’s inequality and random subspaces which escape through a mesh in $\mathbb{R}^n$. 1988. (Cited on pp. 145, 145, 145, 148, 149, 149.)
85. Matthieu Guerquin-Kern, Laurent Lejeune, Klaas Paul Pruessmann, and Michael Unser. Realistic analytical phantoms for parallel magnetic resonance imaging. *IEEE Transactions on Medical Imaging*, 31(3):626–636, 2011. (Cited on p. 190.)
86. Uffe Haagerup and Søren Thorbjørnsen. A new application of random matrices: $\text{ext}(c_{\text{red}}^*(f_2))$ is not a group. *Annals of Mathematics*, 162(2):711–775, 2005. (Cited on p. 165.)
87. B. Hajek, Y. Wu, and J. Xu. Achieving exact cluster recovery threshold via semidefinite programming. *Available online at arXiv:1412.6156*, 2014. (Cited on p. 132.)
88. N. Halko, P. G. Martinsson, and J. A. Tropp. Finding structure with randomness: probabilistic algorithms for constructing approximate matrix decompositions. *Available online at arXiv:0909.4061v2 [math.NA]*, 2009. (Cited on p. 47.)
89. Nathan Halko, Per-Gunnar Martinsson, and Joel A Tropp. Finding structure with randomness: Probabilistic algorithms for constructing approximate matrix decompositions. *SIAM review*, 53(2):217–288, 2011. (Cited on pp. 104, 105, 107, 109, 109.)
90. J. Hastad. Some optimal inapproximability results. 2002. (Cited on p. 115.)
91. I. Haviv and O. Regev. The restricted isometry property of subsampled fourier matrices. *SODA 2016*. (Cited on pp. 180, 184.)
92. Matthew A Herman and Thomas Strohmer. High-resolution radar via compressed sensing. *IEEE transactions on signal processing*, 57(6):2275–2284, 2009. (Cited on pp. 184, 190.)
93. R. A. Horn and C. R. Johnson. *Matrix Analysis*. Cambridge University Press, 1985. (Cited on pp. 33, 46, 185.)
94. R.A. Horn and C.R. Johnson. *Matrix analysis*. Cambridge University Press, Cambridge, 1990. Corrected reprint of the 1985 original. (Cited on pp. 59, 60.)
95. Max Hügel, Holger Rauhut, and Thomas Strohmer. Remote sensing via ℓ_1 -minimization. *Foundations of Computational Mathematics*, 14(1):115–150, 2014. (Cited on p. 189.)
96. W. Johnson and J. Lindenstrauss. Extensions of Lipschitz mappings into a Hilbert space. In *Conference in modern analysis and probability (New Haven, Conn., 1982)*, volume 26 of *Contemporary Mathematics*, pages 189–206. American Mathematical Society, 1984. (Cited on pp. 96, 96.)
97. I. M. Johnston. On the distribution of the largest eigenvalue in principal components analysis. *The Annals of Statistics*, 29(2):295–327, 2001. (Cited on p. 51.)
98. N. E. Karoui. Recent results about the largest eigenvalue of random covariance matrices and statistical application. *Acta Physica Polonica B*, 36(9), 2005. (Cited on p. 51.)
99. S. Khot. On the power of unique 2-prover 1-round games. *Thirty-fourth annual ACM symposium on Theory of computing*, 2002. (Cited on p. 114.)

100. S. Khot, G. Kindler, E. Mossel, and R. O'Donnell. Optimal inapproximability results for max-cut and other 2-variable csps? 2005. (Cited on p. 115.)
101. S. A. Khot and N. K. Vishnoi. The unique games conjecture, integrality gap for cut problems and embeddability of negative type metrics into l_1 . *Available online at arXiv:1305.4581 [cs.CC]*, 2013. (Cited on p. 118.)
102. Anastasia Kireeva and Joel A Tropp. Randomized matrix computations: Themes and variations. *arXiv preprint arXiv:2402.17873*, 2024. (Cited on pp. 103, 107, 109.)
103. Felix Krahmer and Rachel Ward. New and improved johnson–lindenstrauss embeddings via the restricted isometry property. *SIAM Journal on Mathematical Analysis*, 43(3):1269–1281, 2011. (Cited on p. 180.)
104. Shira Kritchman and Boaz Nadler. Determining the number of components in a factor model from limited noisy data. *Chemometrics and Intelligent Laboratory Systems*, 94(1):19–32, 2008. (Cited on p. 55.)
105. Shira Kritchman and Boaz Nadler. Non-parametric detection of the number of signals: Hypothesis testing and random matrix theory. *IEEE Transactions on Signal Processing*, 57(10):3930–3941, 2009. (Cited on p. 55.)
106. Dmitriy Kunisky, Alexander S. Wein, and Afonso S. Bandeira. Notes on computational hardness of hypothesis testing: Predictions using the low-degree likelihood ratio. In Paula Cerejeiras and Michael Reissig, editors, *Mathematical Analysis, its Applications and Computation: ISAAC 2019, Aveiro, Portugal, July 29–August 2*, volume 385 of *Springer Proceedings in Mathematics & Statistics*, pages 1–50. Springer, Cham, 2022. (Cited on p. 151.)
107. K. G. Larsen and J. Nelson. Optimality of the johnson-lindenstrauss lemma. *Available online at arXiv:1609.02094*, 2016. (Cited on p. 99.)
108. J. B. Lasserre. Global optimization with polynomials and the problem of moments. *SIAM Journal on Optimization*, 11(3):796–817, 2001. (Cited on pp. 116, 117.)
109. R. Latała. Some estimates of norms of random matrices. *Proc. Amer. Math. Soc.*, 133(5):1273–1282 (electronic), 2005. (Cited on p. 156.)
110. Rafał Latała, Ramon van Handel, and Pierre Youssef. The dimension-free structure of nonhomogeneous random matrices. *Inventiones mathematicae*, 214(3):1031–1080, 2018. (Cited on pp. 156, 156.)
111. B. Laurent and P. Massart. Adaptive estimation of a quadratic functional by model selection. *Ann. Statist.*, 2000. (Cited on pp. 30, 30.)
112. Michel Ledoux. *The concentration of measure phenomenon*. Number 89. American Mathematical Soc., 2001. (Cited on p. 22.)
113. James R Lee and Assaf Naor. Embedding the diamond graph in l_p and dimension reduction in l_1 . *Geometric & Functional Analysis GAFA*, 14(4):745–747, 2004. (Cited on p. 98.)
114. J.R. Lee, S.O. Gharan, and L. Trevisan. Multi-way spectral partitioning and higher-order cheeger inequalities. *STOC '12 Proceedings of the forty-fourth annual ACM symposium on Theory of computing*, 2012. (Cited on pp. 75, 76.)
115. William Leeb and Elad Romanov. Optimal singular value shrinkage with noise homogenization. *arXiv preprint arXiv:1811.02201*, 2018. (Cited on p. 55.)
116. Shuyang Ling and Thomas Strohmer. Certifying global optimality of graph cuts via semidefinite relaxation: A performance guarantee for spectral clustering. *Foundations of Computational Mathematics*, 20(3):367–421, 2020. (Cited on p. 133.)

117. Lydia T Liu, Edgar Dobriban, and Amit Singer. *e pca*: High dimensional exponential family *pca*. *The Annals of Applied Statistics*, 12(4):2121–2150, 2018. (Cited on p. 55.)
118. S. Lloyd. Least squares quantization in pcm. *IEEE Trans. Inf. Theor.*, 28(2):129–137, 1982. (Cited on p. 62.)
119. Michael Lustig, David Donoho, and John M Pauly. Sparse MRI: The application of compressed sensing for rapid MR imaging. *Magnetic Resonance in Medicine: An Official Journal of the International Society for Magnetic Resonance in Medicine*, 58(6):1182–1195, 2007. (Cited on pp. 170, 170, 184, 189, 190, 190.)
120. Stephane Mallat. *A wavelet tour of signal processing: The sparse way, 3rd edition*. Academic Press, 2009. (Cited on pp. 190, 190, 190.)
121. Henrique S Malvar. *Signal processing with lapped transforms*. Artech House, Inc., 1992. (Cited on p. 190.)
122. V. A. Marchenko and L. A. Pastur. Distribution of eigenvalues in certain sets of random matrices. *Mat. Sb. (N.S.)*, 72(114):507–536, 1967. (Cited on p. 49.)
123. P. Massart. About the constants in Talagrand’s concentration inequalities for empirical processes. *The Annals of Probability*, 28(2), 2000. (Cited on p. 143.)
124. L. Massoulié. Community detection thresholds and the weak ramanujan property. In *Proceedings of the 46th Annual ACM Symposium on Theory of Computing*, STOC ’14, pages 694–703, New York, NY, USA, 2014. ACM. (Cited on p. 122.)
125. Vitali D Milman and Gideon Schechtman. Asymptotic theory of finite-dimensional normed spaces, lecture notes in mathematics 1200, 1986. (Cited on p. 22.)
126. Leon Mirsky. Symmetric gauge functions and unitarily invariant norms. *The quarterly journal of mathematics*, 11(1):50–59, 1960. (Cited on p. 37.)
127. Leon Mirsky. A trace inequality of john von neumann. *Monatshefte für mathematik*, 79(4):303–306, 1975. (Cited on p. 38.)
128. Moshe Mishali and Yonina C Eldar. From theory to practice: Sub-nyquist sampling of sparse wideband analog signals. *IEEE Journal of selected topics in signal processing*, 4(2):375–391, 2010. (Cited on p. 184.)
129. D. G. Mixon. Explicit matrices with the restricted isometry property: Breaking the square-root bottleneck. *available online at arXiv:1403.3427 [math.FA]*, 2014. (Cited on p. 180.)
130. D. G. Mixon. Short, Fat matrices BLOG: Gordon’s escape through a mesh theorem. 2014. (Cited on p. 149.)
131. M. S. Moslehian. Ky Fan inequalities. *Available online at arXiv:1108.1467 [math.FA]*, 2011. (Cited on p. 41.)
132. E. Mossel, J. Neeman, and A. Sly. A proof of the block model threshold conjecture. *Available online at arXiv:1311.4115 [math.PR]*, January 2014. (Cited on p. 122.)
133. E. Mossel, J. Neeman, and A. Sly. Stochastic block models and reconstruction. *Probability Theory and Related Fields (to appear)*, 2014. (Cited on p. 122.)
134. Robb J Muirhead. *Aspects of multivariate statistical theory*. John Wiley & Sons, 2009. (Cited on p. 107.)
135. C. Musco and C. Musco. Stronger and faster approximate singular value decomposition via the block lanczos method. *Available at arXiv:1504.05477 [cs.DS]*, 2015. (Cited on p. 47.)

136. B. Nadler, N. Srebro, and X. Zhou. Semi-supervised learning with the graph laplacian: The limit of infinite unlabelled data. 2009. (Cited on p. 93.)
137. Y. Nesterov. Squared functional systems and optimization problems. *High performance optimization*, 13(405–440), 2000. (Cited on p. 116.)
138. Yuan Ni and Thomas Strohmer. Auto-calibration and biconvex compressive sensing with applications to parallel MRI. *arXiv preprint arXiv:2401.10400*, 2024. (Cited on p. 184.)
139. A. Nica and R. Speicher. Lectures on the combinatorics of free probability. *London Mathematical Society Lecture Note Series*, Cambridge University Press, 335, 2006. (Cited on p. 163.)
140. P. A. Parrilo. *Structured semidefinite programs and semialgebraic geometry methods in robustness and optimization*. PhD thesis, 2000. (Cited on pp. 116, 117.)
141. D. Paul. Asymptotics of the leading sample eigenvalues for a spiked covariance model. Available online at <http://anson.ucdavis.edu/~debashis/techrep/eigenlimit.pdf>. (Cited on pp. 51, 51, 51.)
142. D. Paul. Asymptotics of sample eigenstructure for a large dimensional spiked covariance model. *Statistics Sinica*, 17:1617–1642, 2007. (Cited on p. 51.)
143. K. Pearson. On lines and planes of closest fit to systems of points in space. *Philosophical Magazine, Series 6*, 2(11):559–572, 1901. (Cited on p. 43.)
144. A. Perry, A. S. Wein, A. S. Bandeira, and A. Moitra. Optimality and sub-optimality of pca for spiked random matrices and synchronization. Available online at [arXiv:1609.05573 \[math.ST\]](https://arxiv.org/abs/1609.05573), 2016. (Cited on p. 55.)
145. G. Pisier. *Introduction to operator space theory*, volume 294 of *London Mathematical Society Lecture Note Series*. Cambridge University Press, Cambridge, 2003. (Cited on p. 153.)
146. Lee C Potter, Emre Ertin, Jason T Parker, and Mijdat Cetin. Sparsity and compressed sensing in radar imaging. *Proceedings of the IEEE*, 98(6):1006–1020, 2010. (Cited on p. 184.)
147. P. Raghavendra. Optimal algorithms and inapproximability results for every CSP? In *Proceedings of the Fortieth Annual ACM Symposium on Theory of Computing*, STOC '08, pages 245–254. ACM, 2008. (Cited on p. 115.)
148. Prasad Raghavendra, Tselil Schramm, and David Steurer. High-dimensional estimation from sum-of-squares proofs. *ICM*, 2018. (Cited on p. 117.)
149. S. Riemer and C. Schütt. On the expectation of the norm of random matrices with non-identically distributed entries. *Electron. J. Probab.*, 18, 2013. (Cited on p. 156.)
150. V. Rokhlin, A. Szlam, and M. Tygert. A randomized algorithm for principal component analysis. Available at [arXiv:0809.2274 \[stat.CO\]](https://arxiv.org/abs/0809.2274), 2009. (Cited on p. 47.)
151. Vladimir Rokhlin and Mark Tygert. A fast randomized algorithm for overdetermined linear least-squares regression. *Proceedings of the National Academy of Sciences*, 105(36):13212–13217, 2008. (Cited on p. 109.)
152. Sheldon M Ross. *Introduction to probability models*. Academic press, 2014. (Cited on p. 17.)
153. M. Rudelson and R. Vershynin. On sparse reconstruction from Fourier and Gaussian measurements. *Comm. Pure Appl. Math.*, 61:1025–1045, 2008. (Cited on p. 184.)
154. K. Schmudgen. Around hilbert’s 17th problem. *Documenta Mathematica - Extra Volume ISMP*, pages 433–438, 2012. (Cited on p. 117.)

155. A. J. Scott and M. Grassl. Sic-povms: A new computer study. *J. Math. Phys.*, 2010. (Cited on p. 187.)
156. Y. Seginer. The expected norm of random matrices. *Combin. Probab. Comput.*, 9(2):149–166, 2000. (Cited on p. 156.)
157. N. Shor. An approach to obtaining global extremums in polynomial mathematical programming problems. *Cybernetics and Systems Analysis*, 23(5):695–700, 1987. (Cited on p. 116.)
158. Rainer Sinn and Günter M Ziegler. Landau on chess tournaments and Google’s PageRank. *arXiv preprint arXiv:2210.17300*, 2022. (Cited on p. 60.)
159. Jean-Luc Starck, David L Donoho, and Emmanuel J Candès. Astronomical image representation by the curvelet transform. *Astronomy & Astrophysics*, 398(2):785–800, 2003. (Cited on p. 190.)
160. G. Stengle. A nullstellensatz and a positivstellensatz in semialgebraic geometry. *Math. Ann.* 207, 207:87–97, 1974. (Cited on p. 117.)
161. T. Strohmer and B. Friedlander. Analysis of sparse MIMO radar. *Applied and Computational Harmonic Analysis*, 37:361–388, 2014. (Cited on p. 189.)
162. T. Strohmer and R. Heath. Grassmannian frames with applications to coding and communications. *Appl. Comput. Harmon. Anal.*, 14(3):257–275, 2003. (Cited on pp. 185, 186, 187, 187, 187.)
163. Thomas Strohmer and Benjamin Friedlander. Analysis of sparse MIMO radar. *Applied and Computational Harmonic Analysis*, 37(3):361–388, 2014. (Cited on p. 184.)
164. Thomas Strohmer and Haichao Wang. Accurate imaging of moving targets via random sensor arrays and Kerdock codes. *Inverse Problems*, 29(8):085001, 2013. (Cited on p. 184.)
165. Gábor Szegő. *Orthogonal Polynomials*, volume 23 of *American Mathematical Society Colloquium Publications*. American Mathematical Society, Providence, RI, 4th edition, 1975. (Cited on p. 151.)
166. M. Talagrand. Concentration of measure and isoperimetric inequalities in product spaces. *Inst. Hautes Etudes Sci. Publ. Math.*, (81):73–205, 1995. (Cited on p. 142.)
167. M. Talagrand. *Upper and Lower Bounds for Stochastic Processes: Modern Methods and Classical Problems*. Ergebnisse der Mathematik und ihrer Grenzgebiete. 3. Folge / A Series of Modern Surveys in Mathematics. Springer Berlin Heidelberg, 2014. (Cited on p. 159.)
168. T. Tao. What’s new blog: Open question: deterministic UUP matrices. 2007. (Cited on pp. 180, 187, 188.)
169. T. Tao. *Topics in Random Matrix Theory*. Graduate studies in mathematics. American Mathematical Soc., 2012. (Cited on pp. 49, 142, 142.)
170. J. B. Tenenbaum, V. de Silva, and J. C. Langford. A global geometric framework for nonlinear dimensionality reduction. *Science*, 290(5500):2319–2323, 2000. (Cited on pp. 82, 83, 84, 85, 85.)
171. A. M. Tillmann and M. E. Pfetsch. The computational complexity of the restricted isometry property, the nullspace property, and related concepts in compressed sensing. 2013. (Cited on p. 180.)
172. L. Trevisan. in theory BLOG: CS369G Lecture 4: Spectral Partitioning. 2011. (Cited on p. 72.)
173. J. A. Tropp. Recovery of short, complex linear combinations via ℓ_1 minimization. *IEEE Transactions on Information Theory*, 4:1568–1570, 2005. (Cited on p. 181.)

174. J. A. Tropp. On the conditioning of random subdictionaries. *Appl. Comput. Harmon. Anal.*, 25:1–24, 2008. (Cited on p. 189.)
175. J. A. Tropp. User-friendly tail bounds for sums of random matrices. *Foundations of Computational Mathematics*, 12(4):389–434, 2012. (Cited on pp. 129, 130, 153.)
176. J. A. Tropp. An introduction to matrix concentration inequalities. *Foundations and Trends in Machine Learning*, 2015. (Cited on pp. 153, 154.)
177. J. A. Tropp. Second-order matrix concentration inequalities. *In preparation*, 2015. (Cited on pp. 157, 163, 163.)
178. Joel A. Tropp. The expected norm of a sum of independent random matrices: An elementary approach. In Christian Houdré, David M. Mason, Patricia Reynaud-Bouret, and Jan Rosiński, editors, *High Dimensional Probability VII*, pages 173–202, Cham, 2016. Springer International Publishing. (Cited on pp. 153, 158, 159, 159, 161, 161, 162.)
179. Joel A Tropp and Robert J Webber. Randomized algorithms for low-rank matrix approximation: Design, analysis, and applications. *arXiv preprint arXiv:2306.12418*, 2023. (Cited on p. 109.)
180. R. van Handel. Probability in high dimensions. *ORF 570 Lecture Notes, Princeton University*, 2014. (Cited on pp. 32, 135, 145, 151.)
181. R. van Handel. Nonasymptotic random matrix theory. *St Flour Lecture Notes*, 2022. (Cited on pp. 135, 153, 157.)
182. L. Vanderberghe and S. Boyd. Semidefinite programming. *SIAM Review*, 38:49–95, 1996. (Cited on pp. 113, 116, 125, 127.)
183. L. Vanderberghe and S. Boyd. *Convex Optimization*. Cambridge University Press, 2004. (Cited on pp. 116, 173.)
184. Shreyas S Vasanawala, Marcus T Alley, Brian A Hargreaves, Richard A Barth, John M Pauly, and Michael Lustig. Improved pediatric MR imaging with compressed sensing. *Radiology*, 256(2):607–616, 2010. (Cited on p. 170.)
185. Roman Vershynin. *High-dimensional probability: An introduction with applications in data science*, volume 47. Cambridge University Press, 2018. (Cited on pp. 21, 28, 32.)
186. Dan Voiculescu. Limit laws for random matrices and free products. *Inventiones mathematicae*, 104(1):201–220, 1991. (Cited on p. 165.)
187. John von Neumann. Some matrix-inequalities and metrization of matrix-space. *Tomsk Univ. Rev.*, 1:286–300, 1937. (Cited on p. 38.)
188. Shayne FD Waldron. *An introduction to finite tight frames*. Springer, 2018. (Cited on pp. 185, 187.)
189. Lloyd Welch. Lower bounds on the maximum cross correlation of signals (corresp.). *IEEE Transactions on Information theory*, 20(3):397–399, 2003. (Cited on p. 187.)
190. William K Wootters and Brian D Fields. Optimal state-determination by mutually unbiased measurements. *Annals of Physics*, 191(2):363–381, 1989. (Cited on pp. 187, 187.)
191. G. Zauner. *Quantendesigns–Grundzüge einer nichtkommutativen Designtheorie*. PhD thesis, PhD Thesis Universität Wien, 1999. (Cited on p. 187.)
192. Pan Zhang, Cristopher Moore, and Lenka Zdeborova. Phase transitions in semisupervised clustering of sparse networks. *Phys. Rev. E*, 90, 2014. (Cited on p. 123.)