
**COURSE NOTES (MTH 415, FALL 2025):
APPLIED LINEAR ALGEBRA II**

Mark Iwen

Chapters 1, 2, and Sections 3.1 – 3.4 are based on notes and figures initially scribed by Eric Brodsky while he was an undergraduate student working with me at Michigan State University in 2024.

M.A.I.

Contents

1	Why We Should Care: Online Computing, Artificial Intelligence, and Data, Data, DATA!	3
1.1	Data, and What You Might Do with It	4
1.2	The Basics of Feed-forward Neural Networks (FNNs)	5
1.2.1	Affine Functions and Single Neurons	5
1.2.2	Layers of Neurons, and Some Helpful Matrix Notation	8
1.2.3	Feed-Forward Neural Networks (FNNs) in Full Generality	10
2	Linear Algebra over the Real and Complex Numbers	17
2.1	The Complex Numbers	17
2.1.1	Euler's Identity	20
2.1.2	The Polar Representation of a Complex Number	21
2.1.3	Complex Conjugation	22
2.1.4	The Roots of Unity	23
2.1.5	The Triangle Inequality for Complex Numbers	23
2.2	Basic Linear Algebra over $\mathbb{C}$ and $\mathbb{R}$	24
2.2.1	Some Inner Product Geometry for Real-valued Vectors $\mathbf{x}, \mathbf{y} \in \mathbb{R}^n$. .	26
2.2.2	The Cauchy–Schwarz Inequality	27
2.2.3	General Norms on $\mathbb{C}^{m \times n}$, and the Euclidean Vector Norm	27
2.3	Subspaces, Span, and Linear Independence	29
2.3.1	Bases, Orthonormal Bases, Dimension, and Rank	31
2.4	Orthonormal Bases and the Gram–Schmidt Algorithm	36
2.4.1	The QR Decomposition of a Matrix	41
2.5	Near-Optimal Compression of Low Rank Matrices	43
2.5.1	A Very Brief Review of Gaussian Elimination, and Some Useful Notation	43
2.6	Set Addition, Orthogonal Projections, and Perpendicular Subspaces	48
2.6.1	Representing Orthogonal Projections with Matrices	54
2.6.2	Least-Squares Theory for (Approximately) Solving Systems of Linear Equations	57
2.7	The Four Fundamental Linear Subspaces of a Matrix	58

2.8	The Spectral Theorem for Hermitian Matrices	61
2.9	Positive (Semi)Definite Matrices	65
2.9.1	The Cholesky Decomposition	67
2.10	A Review of the Trace and Determinant Functions	68
2.10.1	The Trace of a Matrix	68
2.10.2	The Determinant of a Matrix	70
3	Some More Advanced Topics in Linear Algebra	71
3.1	One Factorization to Rule Them All: The Singular Value Decomposition . .	71
3.1.1	The Relationship Between any Valid SVD of A and the Spectral Decompositions of A^*A and AA^*	76
3.1.2	The SVD and the Moore–Penrose Inverse of a Matrix	78
3.1.3	Singular Values, Matrix Norms, and Some Related Inequalities . . .	79
3.1.4	Optimal Low-Rank Approximation	83
3.2	Solving Ill-Conditioned and Noisy Linear Systems	85
3.2.1	Improving Conditioning by SVD Truncation	87
3.2.2	Tikhonov Regularization	89
3.3	Linear Least Squares Regression	91
3.4	Discrete Convolution and Fourier Transform Matrices	91
3.4.1	Circulant and Toeplitz Matrices	93
3.4.2	Discrete Fourier Transforms and Circular Convolutions	95
3.4.3	The Fast Fourier Transform (FFT)	103
3.4.4	The FFT for Vectors of Arbitrary Size	107
3.4.5	Fast Matrix Multiplication for Toeplitz Matrices	108
3.5	ℓ^p -norms and Associated Inequalities	111

Chapter 1

Why We Should Care: Online Computing, Artificial Intelligence, and Data, Data, DATA!

Artificial Intelligence, which began being generally useful in the 2020's, resulted from the combination of three crucial historical developments: *(i)* the exponential increase in available computing power from the 1950's until the 2020's,¹ *(ii)* the development of machine learning techniques beginning in the second half of the 20th century (Neural Network methods in particular), and *(iii)* the collection of super-massive data sets for training and learning.² This book is meant to give the reader a solid introduction to the mathematics necessary to begin understanding developments *(ii)* and *(iii)* above. In particular, you will learn about the mathematics needed to understand what a neural network is and how the algorithms work that one might use to compile, process, analyze, and store the types of extremely super-massive datasets needed to train one well. Many of the mathematical topics needed are covered beginning in Chapter 2.

In this chapter we simply aim to prepare you to understand why that material is so important, as well as to state some application problems in a mathematical way that makes them easier to begin understanding more rigorously. Our main contention is this: learning the mathematics first makes all the application problems below much easier to learn about and begin solving later! However, we do understand that mathematics is difficult, and that it helps to have some solid motivation going into a long hike to help keep you trekking uphill until you reach the beautiful views nearer to the top of the mountain. We hope the following sections will help give you that motivation.

¹Mainly due to steady innovations in integrated circuit manufacturing techniques over many decades – read up on Moore's law for a good time!

²Largely made possible by the development of modern communication infrastructure and the subsequent wide-scale adaptation of the internet beginning in the early 1990s.

1.1 Data, and What You Might Do with It

Let N be a positive integer. Herein we will let $[N]$ denote the first N non-negative integers from 0 to $N - 1$, $[N] := \{0, \dots, N - 1\}$ for any natural number $N \in \mathbb{N} = \{0, 1, 2, 3, \dots\}$. Our data herein will (almost always) be a vector of N numbers indexed by $[N]$. We will denote vectors with boldface letters. For example, $\mathbf{x} \in \mathbb{R}^N$ is a vector. We denote the entries of $\mathbf{x}$ by $x_j \in \mathbb{R}$. Pictorially, we have

$$\mathbf{x} = \begin{pmatrix} x_0 \\ x_1 \\ \vdots \\ x_{N-1} \end{pmatrix}.$$

In most settings we consider in this book vectors will be rich enough to represent the data we want to work with. This is primarily because, given the discrete and finite nature of digital computers, one can always simply vectorize other data one might have even if it isn't a vector to begin with. A related application example follows.

Example 1.1.1 (Image Classification Described with Vectors and Functions).

Suppose we want a model to separate pictures into two classes: pictures of cats and pictures of dogs. How can we describe this mathematically? Let's start with a picture of a cat. Assume this picture is 1000 pixels by 1000 pixels, and each pixel has some triple of color values associated to it (one for red, one for green, and one for blue), each a real number in the interval $[0, 1]$. Since a pixel is described by its three color values, each pixel in this image can be described a vector of length 3:

$$\begin{pmatrix} r \\ g \\ b \end{pmatrix} \in \mathbb{R}^3$$

where r denotes the red value of the pixel, and so on. Doing this for each pixel in the image, we attain $1000 \times 1000 = 10^6$ vectors of length 3. We can re-express this data as a single object by concatenating these vectors (in some arbitrary order, such as reading the pixel rows of the image left-to-right and top-to-bottom) into one large vector $\mathbf{x}_{\text{cat}} \in \mathbb{R}^{3 \times 10^6}$. Hence, our cat picture is now simply a big vector.

Now, let's focus on the question of classification. A classification model can be thought of as a function whose input is, e.g., a 1000×1000 picture of a cat or a dog, and whose output is either "cat" or "dog". If we assign the label 0 to cats, and 1 to dogs (or the other way around, if you prefer cats!), then our classification question boils down to finding a function $f : \mathbb{R}^{3 \times 10^6} \rightarrow \{0, 1\}$ such that, given a vectorized picture $\mathbf{x}_{\text{cat}}$ of a cat or a vectorized picture $\mathbf{x}_{\text{dog}}$ of a dog, we correctly get $f(\mathbf{x}_{\text{cat}}) = 0$ and $f(\mathbf{x}_{\text{dog}}) = 1$.

We can use a similar framework for other sorts of problems. For example, the problem of reducing noise in a 1000×1000 cat picture can be viewed as a problem of finding a function $f : \mathbb{R}^{3 \times 10^6} \rightarrow \mathbb{R}^{3 \times 10^6}$ such that $f(\mathbf{x}_{\text{cat}})$ is “less noisy” than the original picture $\mathbf{x}_{\text{cat}}$.

There are a lot of specific image processing methods and techniques built around processing images as two-dimensional objects. For simplicity herein, however, we will use the flexibility of discrete representations to allow us to turn any image, etc., into a vector as an excuse to ignore non-vector data (i.e., we will vectorize everything). Though this can always be done, we note that it certainly shouldn’t always be done... Nonetheless, it’s generally useful enough that we will do it here. It also will make understanding the mathematics involved much easier, which we will take as an additional reason to assume that our datasets are almost always collections of vectors herein.

1.2 The Basics of Feed-forward Neural Networks (FNNs)

Continuing for the moment in the spirit of our first example above, we will now briefly take a detour to discuss what kinds of functions f one might actually build and evaluate with a computer to, e.g., classify images as in Example 1.1.1. FNNs provide exactly one such “computer friendly” class of functions that are also expressive enough to be able to do many useful tasks quite well. Given their value in artificial intelligence applications we will now take some time to explain what they are and how they depend on, and utilize, ideas from, e.g., both linear algebra and optimization. To begin we will first discuss the atomic building block of every neural network – the neuron.

1.2.1 Affine Functions and Single Neurons

Let $\mathbf{x}$ and $\mathbf{y}$ be vectors in $\mathbb{R}^N$. We define the **inner product** of $\mathbf{x}$ and $\mathbf{y}$, denoted $\langle \mathbf{x}, \mathbf{y} \rangle$, to be the sum

$$\langle \mathbf{x}, \mathbf{y} \rangle = \sum_{j=0}^{N-1} x_j y_j \quad (1.1)$$

Definition 1.2.1 (Affine Functions). *Fix $\mathbf{w} \in \mathbb{R}^N$ and $a, b \in \mathbb{R}$. Then the **affine function** determined by $\mathbf{w}$ and b is the function $a_{\mathbf{w},b} : \mathbb{R}^N \rightarrow \mathbb{R}$ defined by*

$$a_{\mathbf{w},b}(\mathbf{x}) := \langle \mathbf{x}, \mathbf{w} \rangle + b$$

Here $\mathbf{w}$ is called the affine function’s **weight vector** and b is called its **bias**.

Note that we can also write the above as a single inner product of two vectors in $\mathbb{R}^{N+1}$,

$$\langle \mathbf{x}, \mathbf{w} \rangle + b = \left\langle \begin{pmatrix} \mathbf{x} \\ 1 \end{pmatrix}, \begin{pmatrix} \mathbf{w} \\ b \end{pmatrix} \right\rangle.$$

We can also represent an affine function $a_{\mathbf{w},b} : \mathbb{R}^N \rightarrow \mathbb{R}$ graphically as in Figure 1.1.

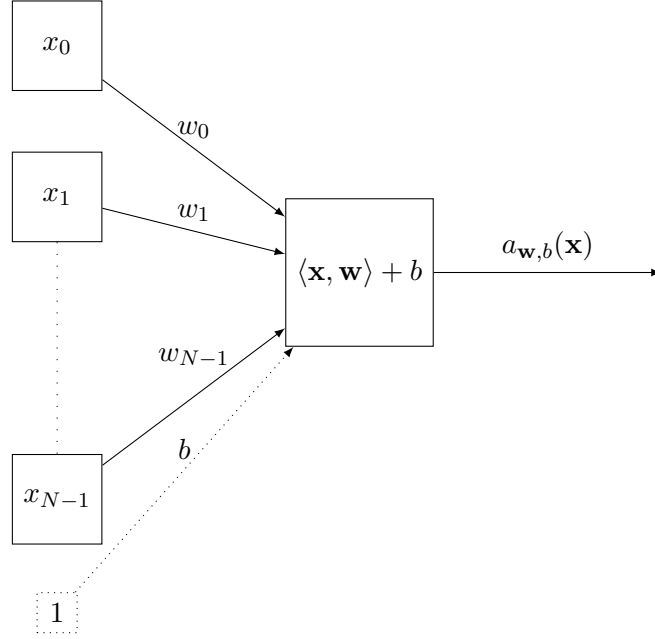


Figure 1.1: A graphical representation of an affine function $a_{\mathbf{w},b} : \mathbb{R}^N \rightarrow \mathbb{R}$. The first column of boxes represents the inputs (i.e., the entries of $\mathbf{x}$). The edge weights are the entry of the weight vector $\mathbf{w}$ that multiplies each corresponding input entry in the affine function's inner product (e.g., w_0 multiplies against x_0 , etc.). The dotted box around the constant input 1 used here to include the bias b as an edge weight is often omitted.

Definition 1.2.2 (Neurons). A **neuron** $\eta : \mathbb{R}^N \rightarrow \mathbb{R}$ is a composition of an affine function $a_{\mathbf{w},b}$ with a nonlinear function $\sigma : \mathbb{R} \rightarrow \mathbb{R}$ given by

$$\eta(\mathbf{x}) := \sigma(a_{\mathbf{w},b}(\mathbf{x})) = \sigma(\langle \mathbf{x}, \mathbf{w} \rangle + b).$$

Note that a neuron is determined by two choices: the parameters $\mathbf{w} \in \mathbb{R}^N$ and $b \in \mathbb{R}$, and the **activation function** σ .

A neuron also admits the commonly used graphical representation in Figure 1.2. In Figure 1.2 the first column of boxes is called the **input layer** and the circle is called a **node** or **neuron**. Some typical choices of activation functions σ include the

- Perceptron (or Heaviside, or step function): $\sigma(y) = \begin{cases} 0 & \text{if } y \leq 0 \\ 1, & \text{if } y > 0 \end{cases}$
- Sigmoid: $\sigma(y) = 1/(1 + e^{-y})$
- Hyperbolic tangent: $\sigma(y) = \tanh(y)$

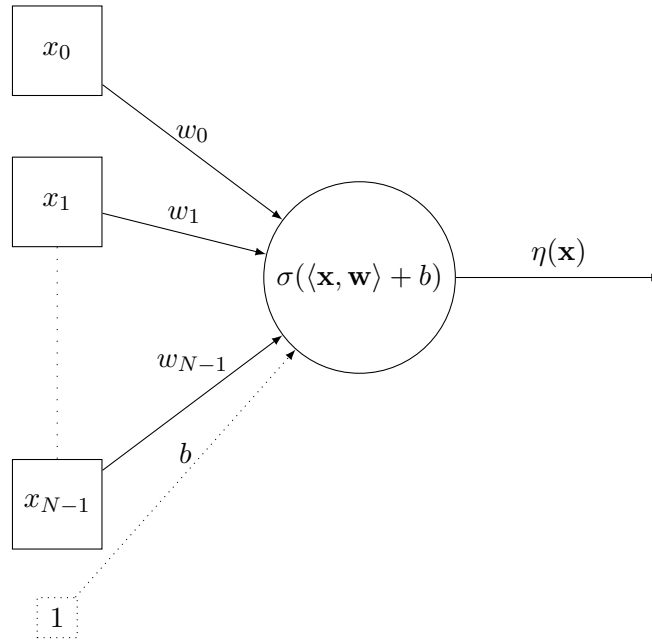


Figure 1.2: A graphical representation of a neuron. The first column of boxes represents the inputs (i.e., the entries of $\mathbf{x}$). The edge weights are the entry of the weight vector $\mathbf{w}$ that multiplies each corresponding input entry in the neuron's inner product (e.g., w_0 multiplies against x_0). Note in particular that a circle is used to represent a neuron here, as opposed to a box which is used to represent an affine function as per Figure 1.1. Again, the dotted box around the constant input 1 used here to include the bias b as an edge weight is often omitted.

- Rectified Linear Unit (ReLU): $\sigma(y) = \max(0, y)$
- Leaky ReLU: $\sigma_a(y) = \max(ay, y)$ with $0 \leq a < 1$.
- Absolute value (or modulus): $\sigma(y) = |y|$
- Smoothed versions of (leaky) ReLU to eliminate non-differentiability at $y = 0$.

Example 1.2.3 (A Simple Way to Smooth Non-Differentiability). *Fix $a \in [0, 1)$ and $\alpha \in \mathbb{R}^+$, and define the function $g : \mathbb{R} \rightarrow \mathbb{R}$ to be*

$$g(y) = \begin{cases} a, & y < -\alpha \\ 1, & y > \alpha \\ a + \frac{1-a}{2\alpha} \cdot (y + \alpha), & -\alpha \leq y \leq \alpha \end{cases}$$

A smoothed leaky ReLU function, $\tilde{\sigma}_{a,\alpha}(x)$, can be defined to be

$$\tilde{\sigma}_{a,\alpha}(x) = \int_0^x g(y) dy. \quad (1.2)$$

Exercise 1.2.1. The following problems concern the smoothed (leaky) ReLU function $\tilde{\sigma}_{a,\alpha} : \mathbb{R} \rightarrow \mathbb{R}$ defined in (1.2) with $a = 1/2$ and $\alpha = 1/4$.

- (a) Compute the integral in (1.2) and write down the resulting piecewise polynomial formula for $\tilde{\sigma}_{\frac{1}{2},\frac{1}{4}}(x)$. What is $\tilde{\sigma}_{\frac{1}{2},\frac{1}{4}}(1)$?
- (b) Plot $\tilde{\sigma}_{\frac{1}{2},\frac{1}{4}}$ together with the leaky ReLU function $\sigma_{\frac{1}{2}}$.

Given an activation function $\sigma : \mathbb{R} \rightarrow \mathbb{R}$ we will extend it to a function $\sigma : \mathbb{R}^N \rightarrow \mathbb{R}^N$ for any given $N \in \mathbb{N}$ entrywise by

$$\sigma(\mathbf{x}) := \begin{pmatrix} \sigma(x_0) \\ \sigma(x_1) \\ \vdots \\ \sigma(x_{N-1}) \end{pmatrix}.$$

We will now continue to build on this notation in order to help combine multiple neurons into more complicated (and useful!) functions.

1.2.2 Layers of Neurons, and Some Helpful Matrix Notation

A **matrix** $W \in \mathbb{R}^{N \times d}$ is a table of data with N rows and d columns. We denote the entry in the j^{th} row and k^{th} column of W by $W_{j,k} \in \mathbb{R}$ for all $j \in [N]$ and $k \in [d]$. We denote the j^{th} row of W , which is a vector in $\mathbb{R}^d$, by $W_{j,:} \in \mathbb{R}^d$. Similarly, we denote the j^{th} column of W , which is a vector in $\mathbb{R}^N$, by $W_{:,j} \in \mathbb{R}^N$. We can also build a matrix out of vectors. Given d vectors $\mathbf{w}_0, \dots, \mathbf{w}_{d-1} \in \mathbb{R}^N$, we can write the $N \times d$ matrix whose j^{th} column is $W_{:,j} = \mathbf{w}_j$ for all $j \in [d]$ as

$$\begin{pmatrix} | & & | \\ \mathbf{w}_0 & \cdots & \mathbf{w}_{d-1} \\ | & & | \end{pmatrix} \in \mathbb{R}^{N \times d}.$$

Given a matrix $W \in \mathbb{R}^{N \times d}$ and a vector $\mathbf{y} \in \mathbb{R}^N$ we will also denote by $(W|\mathbf{y}) \in \mathbb{R}^{N \times (d+1)}$ matrix whose first d columns are the columns of W , and whose $(d+1)^{\text{st}}$ column is $\mathbf{y}$.

The **transpose** of a matrix $W \in \mathbb{R}^{N \times d}$, denoted by $W^T \in \mathbb{R}^{d \times N}$, is the $d \times N$ matrix with entries given in terms of W by $(W^T)_{j,k} = W_{k,j}$ for all $j \in [d]$ and $k \in [N]$. That is, we swap the roles of rows and columns so that, e.g., $W_{j,:} = W_{:,j}^T$ for all $j \in [N]$.

Finally, a matrix $W \in \mathbb{R}^{N \times d}$ also always represents a linear function $W : \mathbb{R}^d \rightarrow \mathbb{R}^N$ where $W(\mathbf{x}) = W\mathbf{x} \in \mathbb{R}^N$ has entries given by

$$(W\mathbf{x})_j := \sum_{k \in [d]} W_{j,k} x_k$$

for all $j \in [N]$.

Exercise 1.2.2. Let $\mathbf{x} \in \mathbb{R}^N$, $\mathbf{y} \in \mathbb{R}^d$, and $W \in \mathbb{R}^{N \times d}$. Show that $\langle \mathbf{x}, W\mathbf{y} \rangle = \langle W^T \mathbf{x}, \mathbf{y} \rangle$.

We can also represent a matrix graphically as multiple affine functions. Let $W \in \mathbb{R}^{N \times d}$ and $\mathbf{x} \in \mathbb{R}^d$. Then we can express the matrix-vector product $W\mathbf{x} \in \mathbb{R}^N$ with the diagram in Figure 1.3

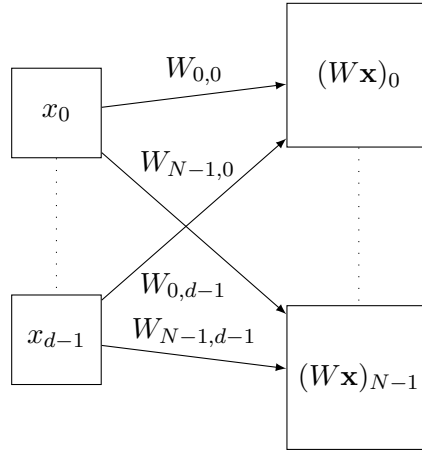


Figure 1.3: A graphical representation of a matrix $W \in \mathbb{R}^{N \times d}$ as an input layer of width d connected directly to a linear output layer of width N .

We now have enough notation to define and represent a single layer of neurons.

Definition 1.2.4 (A Layer of Neurons). A **layer of neurons** $\ell : \mathbb{R}^d \rightarrow \mathbb{R}^N$ is determined by a collection of d weight vectors $\mathbf{w}_0, \dots, \mathbf{w}_{d-1} \in \mathbb{R}^N$, d biases $b_0, \dots, b_{d-1}$, and a choice of activation function $\sigma : \mathbb{R} \rightarrow \mathbb{R}$. We call d the **width** of ℓ . The layer ℓ is defined using these parameters by

$$\ell(\mathbf{x}) := \begin{pmatrix} \sigma(\langle \mathbf{x}, \mathbf{w}_0 \rangle + b_0) \\ \vdots \\ \sigma(\langle \mathbf{x}, \mathbf{w}_{d-1} \rangle + b_{d-1}) \end{pmatrix} = \sigma \begin{pmatrix} \langle \mathbf{x}, \mathbf{w}_0 \rangle + b_0 \\ \vdots \\ \langle \mathbf{x}, \mathbf{w}_{d-1} \rangle + b_{d-1} \end{pmatrix} = \sigma(W\mathbf{x} + \mathbf{b}),$$

where $W^T = \begin{pmatrix} | & & | \\ \mathbf{w}_0 & \cdots & \mathbf{w}_{d-1} \\ | & & | \end{pmatrix} \in \mathbb{R}^{N \times d}$ and $\mathbf{b} = \begin{pmatrix} b_0 \\ \vdots \\ b_{d-1} \end{pmatrix}$. Note that $\ell : \mathbb{R}^N \rightarrow \mathbb{R}^d$ is effectively created by stacking d different neurons $\eta_0, \dots, \eta_{d-1} : \mathbb{R}^N \rightarrow \mathbb{R}$ into a vector. Here $W \in \mathbb{R}^{d \times N}$ is called the layer's **weight matrix** and $\mathbf{b}$ is called the layer's **bias vector**.

Above can also write $\ell(\mathbf{x}) = \sigma(W\mathbf{x} + \mathbf{b})$ as $\ell(\mathbf{x}) = \sigma\left(\tilde{A}\begin{pmatrix} \mathbf{x} \\ 1 \end{pmatrix}\right)$, where $\tilde{A} = (W|\mathbf{b}) \in \mathbb{R}^{d \times (N+1)}$. Thus, if we define the affine function $A : \mathbb{R}^N \rightarrow \mathbb{R}^d$ by $A(\mathbf{x}) := \tilde{A}\begin{pmatrix} \mathbf{x} \\ 1 \end{pmatrix} = W\mathbf{x} + \mathbf{b}$, we may further write ℓ compactly as a composition of σ and A , i.e., $\ell(\mathbf{x}) = \sigma(A(\mathbf{x})) = (\sigma \circ A)(\mathbf{x})$. This compositional form will be used below. Finally, we note that one can also represent a layer of d neurons graphically as per Figure 1.4.

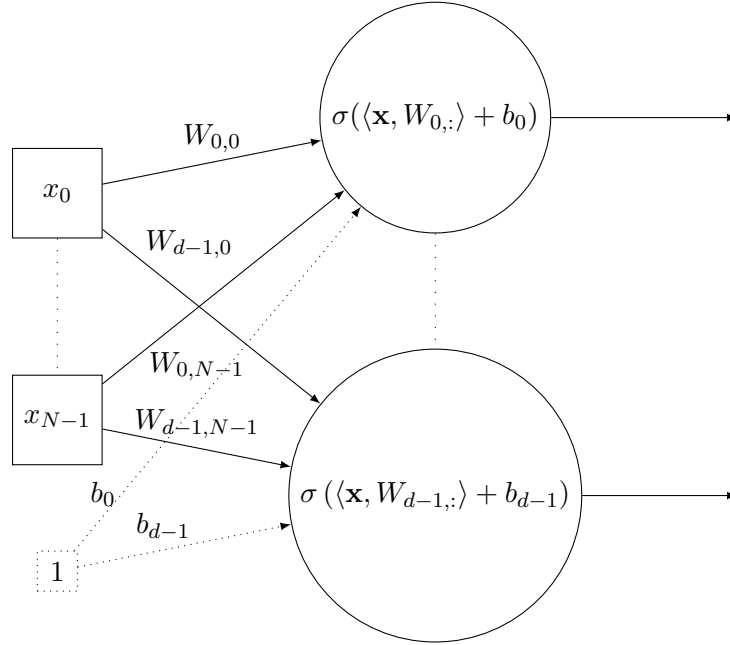


Figure 1.4: A graphical representation of a layer of neurons $\ell : \mathbb{R}^N \rightarrow \mathbb{R}^d$ defined by $\ell(\mathbf{x}) = \sigma(W\mathbf{x} + \mathbf{b})$ with weight matrix $W \in \mathbb{R}^{d \times N}$ and bias vector $\mathbf{b} \in \mathbb{R}^d$. Here the input layer of width N connects to a layer of d neurons.

1.2.3 Feed-Forward Neural Networks (FNNs) in Full Generality

Informally, a FNN is a series of layers of neurons with each layer feeding its outputs “forward” into the layer following it. From the discussion of neuron layers above, we can therefore

choose to describe an FNN as a concatenation of functions that alternate between affine functions and the activation function. More formally, one can define FNNs as follows.

Definition 1.2.5 (Feed-forward Neural Network (FNN)). A **Feed-forward Neural Network (FNN)** $f : \mathbb{R}^N \rightarrow \mathbb{R}^{d_L}$ is determined by an activation function $\sigma : \mathbb{R} \rightarrow \mathbb{R}$, a **depth** $L \in \mathbb{N}$, and layer widths $d_0, \dots, d_L$. It contains an input layer with N inputs, L layers of neurons (often called **hidden layers**), and a final linear output layer with d_L outputs. More specifically, let $\ell^0 : \mathbb{R}^N \rightarrow \mathbb{R}^{d_0}$ and $\ell^j : \mathbb{R}^{d_{j-1}} \rightarrow \mathbb{R}^{d_j} \forall j \in \{1, \dots, L-1\}$ be L layers of neurons, and let $A^L : \mathbb{R}^{d_{L-1}} \rightarrow \mathbb{R}^{d_L}$ be an affine function defined by $A^L(\mathbf{y}) := W^L \mathbf{y} + \mathbf{b}^L$ for $W^L \in \mathbb{R}^{d_L \times d_{L-1}}$, $\mathbf{b}^L \in \mathbb{R}^{d_L}$. The resulting FNN of depth L , f , is then given for all $\mathbf{x} \in \mathbb{R}^N$ by

$$\begin{aligned} f(\mathbf{x}) &= (A^L \circ \ell^{L-1} \circ \ell^{L-2} \circ \dots \circ \ell^1 \circ \ell^0)(\mathbf{x}) \\ &= (A^L \circ \sigma \circ A^{L-1} \circ \sigma \circ A^{L-2} \circ \dots \circ \sigma \circ A^0)(\mathbf{x}), \end{aligned}$$

where $A^{L-k}(\mathbf{y}) = W^{L-k}(\mathbf{y}) + \mathbf{b}^{L-k}$, with $W^{L-k} \in \mathbb{R}^{d_{L-k} \times d_{L-k-1}}$ and $\mathbf{b}^{L-k} \in \mathbb{R}^{d_{L-k}}$ for all $k \in [L]$, and $A^0(\mathbf{y}) = W^0 \mathbf{y} + \mathbf{b}^0$, with $W^0 \in \mathbb{R}^{d_0 \times N}$ and $\mathbf{b}^0 \in \mathbb{R}^{d_0}$.

Even after fixing the activation function σ we note that FNNs are functions that depend on a potentially huge number of parameters. Using our notation from above, the number of parameters in a FNN f is equal to the sum of the number of weights in the matrices W^j and the number of biases in the vectors $\mathbf{b}^j$ for all $j \in [L+1]$. Recall that when $j > 0$, $W^j \in \mathbb{R}^{d_j \times d_{j-1}}$ and $\mathbf{b}^j \in \mathbb{R}^{d_j}$, and when $j = 0$, $W^0 \in \mathbb{R}^{d_0 \times N}$ and $\mathbf{b}^0 \in \mathbb{R}^{d_0}$. Thus, the total number of parameters for a depth L FNN f with input layer width N and hidden layer widths $d_0, d_1, \dots, d_L$ is

$$\# \text{ FNN parameters} = d_0(N+1) + \sum_{j=1}^L d_j(d_{j-1}+1).$$

Finding a good way of choosing all of these parameters during training so that the resulting trained FNN is capable of, e.g., correctly classifying cat versus dog pictures is usually accomplished via optimization techniques. Techniques one can use to help reduce the number of these parameters in order to save space when storing a previously-trained FNN is something we will discuss more in, e.g., Sections 2.5 and 3.1. We urge you to keep reading to learn about these useful tricks, and more!

For now though, we will simply try to mitigate the fact that the general definition of a depth L FNN given above is rather complicated. In order to help digest it, let's consider some examples. Our first example will be that of a **shallow** FNN (that is, of an FNN of depth $L = 1$).

Example 1.2.6 (A Shallow FNN $f : \mathbb{R} \rightarrow \mathbb{R}$). A *shallow* (i.e., $L = 1$) FNN $f : \mathbb{R} \rightarrow \mathbb{R}$ will have the form

$$f(x) = b^1 + \sum_{j=0}^{d_0-1} w_j^1 \sigma(w_j^0 x + b_j^0). \quad (1.3)$$

where $b^1 \in \mathbb{R}$ is the single output layer bias (the output width is $d_1 = 1$), b_j^0 for all $j \in [d_0]$ are the biases of the single layer of neurons of width d_0 , and where the weights of the layer of neurons and the output layer are $w_j^0, w_j^1 \in \mathbb{R}$ for all $j \in [d_0]$, respectively.

Example 1.2.7 (The Graphical Representation of a Shallow FNN $f : \mathbb{R}^3 \rightarrow \mathbb{R}^2$). For a graphical representation of a shallow (i.e., depth $L = 1$) FNN $f : \mathbb{R}^3 \rightarrow \mathbb{R}^2$ with widths $d_0 = 2$ and $d_1 = 2$ see Figure 1.5. Note that such a network will be determined by two weight matrices $W^0 \in \mathbb{R}^{2 \times 3}$, $W^1 \in \mathbb{R}^{2 \times 2}$ and two bias vectors $\mathbf{b}^0, \mathbf{b}^1 \in \mathbb{R}^2$. Hence, it has a total of 14 parameters.

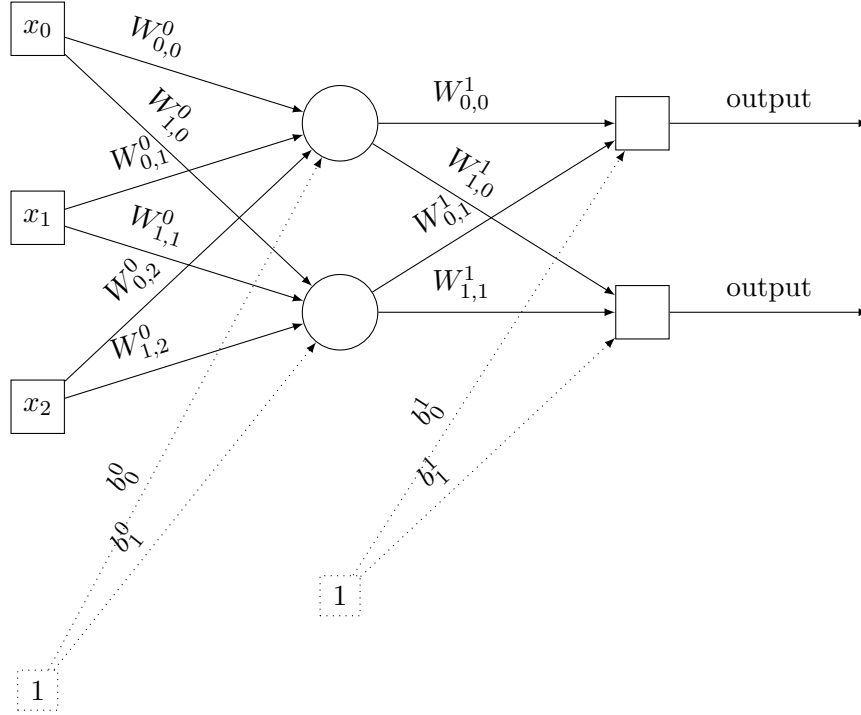


Figure 1.5: An example of a shallow neural network $f : \mathbb{R}^3 \rightarrow \mathbb{R}^2$. We call a depth 1 FNN **shallow**. The leftmost layer is the input layer with $N = 3$ inputs. The middle layer, which is the only nonlinear layer in this diagram, is a hidden layer of neurons with $d_0 = 2$ neurons. The right layer is the output layer with $d_L = 2$ outputs.

Example 1.2.8 (The Graphical Representation of a Depth $L = 2$ FNN $f : \mathbb{R}^2 \rightarrow \mathbb{R}^2$). For a graphical representation of a depth $L = 2$ FNN $f : \mathbb{R}^2 \rightarrow \mathbb{R}^2$ with widths $d_0 = 3$, $d_1 = 2$, and $d_2 = 2$ see Figure 1.6. Such a network will be determined by three weight matrices $W^0 \in \mathbb{R}^{3 \times 2}$, $W^1 \in \mathbb{R}^{2 \times 3}$, $W^2 \in \mathbb{R}^{2 \times 2}$ and three bias vectors $\mathbf{b}^0 \in \mathbb{R}^3$, $\mathbf{b}^1, \mathbf{b}^2 \in \mathbb{R}^2$. Hence, it has a total of 23 parameters.

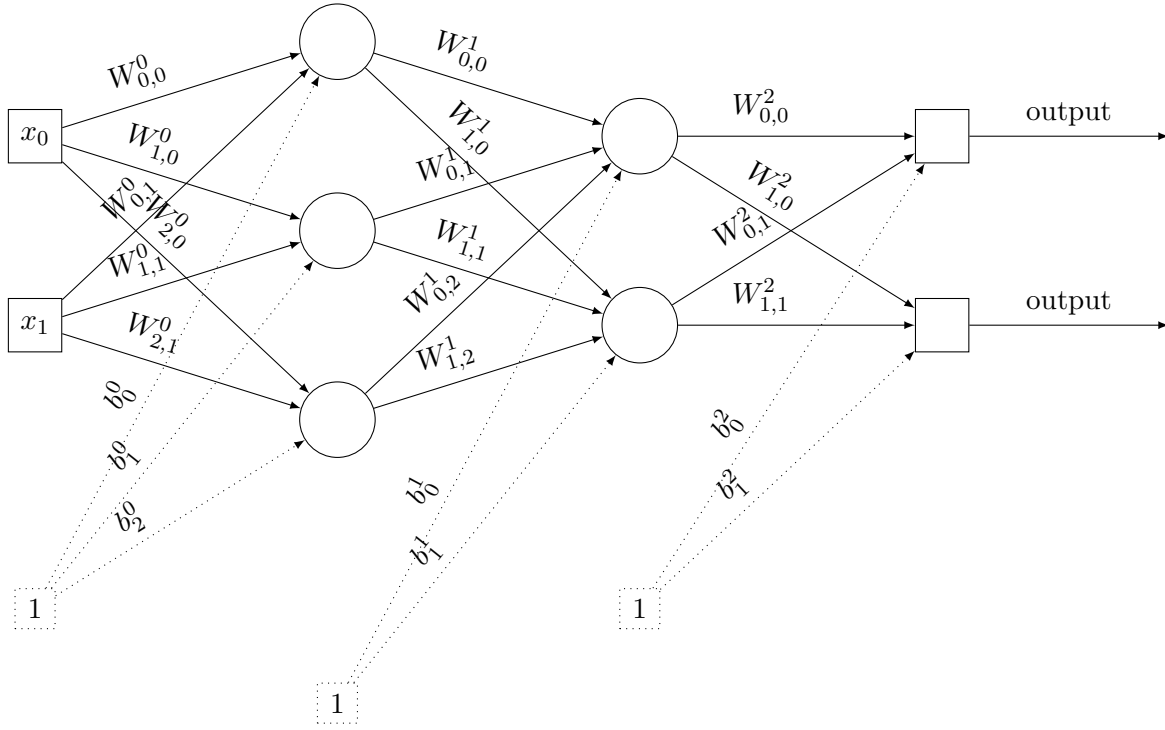


Figure 1.6: An example of a neural network $f : \mathbb{R}^2 \rightarrow \mathbb{R}^2$ of depth $L = 2$. The leftmost layer is the input layer with $N = 2$ inputs. The second layer from the left is the first hidden layer of neurons, which has $d_0 = 3$ neurons. The third layer from the left is the second hidden layer of neurons, which has width $d_1 = 2$, and the rightmost layer is the linear output layer, which has $d_L = d_2 = 2$ outputs.

Exercise 1.2.3. Draw the graphical representation of a shallow neural network $f : \mathbb{R} \rightarrow \mathbb{R}$ of width $d_0 = 5$. How many parameters does it have?

Exercise 1.2.4. Draw the graphical representation of a depth $L = 3$ neural network $f : \mathbb{R} \rightarrow \mathbb{R}$ with widths $d_0 = 2, d_1 = 2, d_2 = 2$. How many parameters does it have?

We will now briefly discuss why choosing, e.g., a greater value for its depth L might allow a FNN to “work better” at a variety of tasks. This is directly linked to the notion of the “expressivity” of an FNN.

Some Basics Concerning the Expressivity of FNNs

In practice the activation function $\sigma : \mathbb{R} \rightarrow \mathbb{R}$ is always chosen to be a nonlinear function. The reason why is directly linked to the notion of the “expressivity” of an FNN. Suppose

for example that we choose $\sigma : \mathbb{R} \rightarrow \mathbb{R}$ in (1.3) to be linear so that $\sigma(y) = ay + c$ for some $a, c \in \mathbb{R}$. Substituting this activation function into (1.3) we obtain

$$\begin{aligned} f(x) &= b^1 + \sum_{j=0}^{d_0-1} w_j^1 \sigma(w_j^0 x + b_j^0) = b^1 + \sum_{j=0}^{d_0-1} w_j^1 [a(w_j^0 x + b_j^0) + c] \\ &= \underbrace{\left(\sum_{j=0}^{d_0-1} w_j^1 a w_j^0 \right)}_{=: \tilde{a}} x + \underbrace{\left(b^1 + \sum_{j=0}^{d_0-1} w_j^1 (a b_j^0 + c) \right)}_{=: \tilde{c}} = \tilde{a}x + \tilde{c}, \end{aligned}$$

with the two new constants $\tilde{a}, \tilde{c} \in \mathbb{R}$ defined as above. That is, if we choose σ to be linear then the complicated shallow FNN $f : \mathbb{R} \rightarrow \mathbb{R}$ in (1.3) is just another linear function itself. All the weight and bias parameters used to define it were a total waste of time! Stated another way, choosing σ to be linear only allows shallow FNNs such as (1.3) to express simple linear functions.

As we shall see next, choosing σ to be something even “barely nonlinear” such as a ReLU function $\sigma(y) = \text{ReLU}(y) := \max(0, y)$ already allows shallow FNNs such as (1.3) to express/represent significantly more complicated functions than simple linear ones.³ The following Theorem is paraphrased from Foucart’s fantastic book on data science [12]. Informally, it tells us that choosing σ to be a ReLU function allows shallow FNNs such as (1.3) to express any continuous piecewise linear function you like. Note that this is a dramatically larger class of functions than the simple linear ones shallow FNNs such as (1.3) can express if σ is chosen to be linear. Hence, in this case *choosing σ to be nonlinear increases expressivity*.

Theorem 1.2.9 (See Theorem 24.1 in [12]). *Let $\sigma : \mathbb{R} \rightarrow \mathbb{R}$ be the ReLU function $\text{ReLU}(x) = \max\{0, x\}$. Then, every continuous piecewise linear function $f : \mathbb{R} \rightarrow \mathbb{R}$ as in (1.4) can be expressed by a shallow FNN whose single hidden layer contains $n + 2$ neurons. More specifically, let*

$$f(x) = \begin{cases} a_0 x + b_0 & x \leq \tau_1 \\ a_1 x + b_1 & \tau_1 \leq x \leq \tau_2 \\ \vdots & \\ a_{n-1} x + b_{n-1} & \tau_{n-1} \leq x \leq \tau_n \\ a_n x + b_n & \tau_n \leq x \end{cases} \quad (1.4)$$

where $\tau_1 < \tau_2 < \dots < \tau_n$ are real numbers, and $a_0, \dots, a_n$ and $b_0, \dots, b_n$ are real numbers such that the function f above is continuous (i.e., $a_j \tau_{j+1} + b_j = a_{j+1} \tau_{j+1} + b_{j+1}$ for all $j \in [N]$). In other words, f is a piecewise linear function whose slope changes finitely many (specifically, n) times. Any such function can be obtained via a shallow FNN of width $n + 2$.

³Note that the ReLU function itself is linear everywhere except at 0. Hence, I feel it is appropriate to label it as “barely nonlinear”.

Proof. We begin by noting two useful properties of the ReLU function:

$$\begin{aligned}\text{ReLU}(\gamma x) &= \gamma \text{ReLU}(x) \quad \forall x \in \mathbb{R}, \gamma > 0, \quad \text{and} \\ x &= \text{ReLU}(x) - \text{ReLU}(-x) \quad \forall x \in \mathbb{R}.\end{aligned}$$

Using these two properties, we can write f as the following linear combination of $n + 2$ ReLU functions as follows

$$\begin{aligned}f(x) &= a_0 x + b_0 + \sum_{j=1}^n (a_j - a_{j-1}) \text{ReLU}(x - \tau_j) \\ &= \text{ReLU}(a_0 x + b_0) - \text{ReLU}(-a_0 x - b_0) + \sum_{j=1}^n (a_j - a_{j-1}) \text{ReLU}(x - \tau_j).\end{aligned}$$

□

Note that the class of piecewise linear functions is actually quite powerful approximation-theoretically since one can, e.g., approximate any continuous function $\mathbb{R} \rightarrow \mathbb{R}$ within a bounded domain arbitrarily well using increasingly fine piecewise linear approximations. Thus, the theorem above tells us that even when using the most basic tools available to us (a straightforward *nonlinear* activation function within a FNN with just a single layer) we can already approximate a very general class of functions from $\mathbb{R} \rightarrow \mathbb{R}$ as well as we want.

When we consider functions of two variables, however, things become a bit more complicated. For example, [12] also shows that the bivariate piecewise linear function $g(x_0, x_1) = \min(0, \max(x_0, x_1))$ can *not* be exactly represented by a *shallow* ReLU FNN of any width. That said, as the next theorem demonstrates, g can in fact be exactly represented by a FNN of depth $L = 2$. This simple example is meant to demonstrate the following more general principal: *Increasing the depth of a FNN increases its expressivity*. In practice the depths (and widths) of modern neural networks are very large for this reason, leading to the necessity of practitioners to deal with many very large matrices. This is just one of the many many reasons it's crucial for the modern data scientist to know the linear algebra we will review in the next chapter. Hope to see you there!

Theorem 1.2.10 (Section 24.3 in [12]). *Define the function $g : \mathbb{R}^2 \rightarrow \mathbb{R}$ by $g(x_0, x_1) = \min\{0, \max\{x_0, x_1\}\}$. This function g cannot be generated by a shallow ReLU FNN, but g can be obtained as a depth $L = 2$ ReLU FNN.*

Proof. For a proof that g cannot be generated by a shallow ReLU FNN, consult [12, Theorem 24.1]. Below we show explicitly how g can be written as a depth 2 ReLU FNN.

$$\begin{aligned}g(x_0, x_1) &= \min\{0, \max\{x_0, x_1\}\} \\ &= -\text{ReLU}(-\max\{x_0, x_1\}) \\ &= -\text{ReLU}(-(x_0 + \text{ReLU}(x_1 - x_0))) \\ &= -\text{ReLU}(-\text{ReLU}(x_0) + \text{ReLU}(-x_0) - \text{ReLU}(x_1 - x_0))\end{aligned}$$

We can also draw this neural network as in Figure 1.7, omitting arrows with weight 0. $\square$

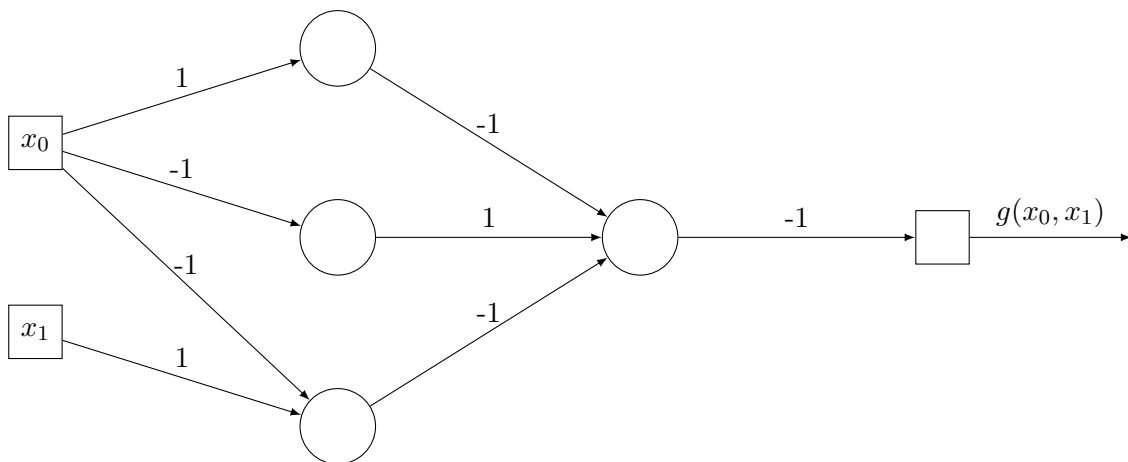


Figure 1.7: The graphical representation of the depth $L = 2$ ReLU FNN from the proof of Theorem 1.2.10 that computes $g(x_0, x_1) = \min\{0, \max\{x_0, x_1\}\}$.

Chapter 2

Linear Algebra over the Real and Complex Numbers

In this chapter we will introduce/review linear algebra over the complex numbers. We note immediately, however, that the real numbers are also complex numbers! **If the reader is intimidated by (or temporarily disinterested in) doing linear algebra over the complex numbers, they can simply skip down to Section 2.2 and replace the symbol “ $\mathbb{C}$ ” everywhere it appears there with an “ $\mathbb{R}$ ”. Doing so will not affect the correctness of anything in this chapter, or limit your understanding in an important way until Section 3.4.** We will also continue to use the matrix notation and conventions discussed in, e.g., Section 1.2.2 going forward. All of that material (where one restricts oneself to thinking about the reals $\mathbb{R} \subset \mathbb{C}$) also remains true in this chapter. In short, if you know how linear algebra works over $\mathbb{C}$, then you can reduce to linear algebra over $\mathbb{R}$ by simply replacing “ $\mathbb{C}$ ” everywhere it appears with an “ $\mathbb{R}$ ”. Doing linear algebra over the complex numbers instead of the reals in the first place does require a few minor adaptations, though (mainly, you need to use complex conjugation in a few crucial definitions). We will do that for you below. Before we begin, however, let’s review the complex numbers.

2.1 The Complex Numbers

In this book the letter $\mathbf{i}$ will be reserved for the imaginary number $\sqrt{-1}$. That is, $\mathbf{i}^2 := -1$. The imaginary number $\mathbf{i}$ satisfies all the properties you hope it would when interacting with elements of $\mathbb{R}$ including: $0\mathbf{i} = 0$ and $1\mathbf{i} = \mathbf{i}$, as well as all the usual associative, commutative, and distributive properties (e.g., $\mathbf{i}x = x\mathbf{i}$ and $\mathbf{i} + x = x + \mathbf{i} \ \forall x \in \mathbb{R}$). A **complex number** is an object of the form $z = x + \mathbf{i}y$ for $x, y \in \mathbb{R}$. The set of complex numbers is denoted

$$\mathbb{C} := \{x + \mathbf{i}y \mid x, y \in \mathbb{R}\}.$$

The number x in $z = x + iy$ is called the real part of z , and is denoted $\operatorname{Re}(z) \in \mathbb{R}$. Similarly, the number y is called the imaginary part of z , and is denoted $\operatorname{Im}(z) \in \mathbb{R}$. A real number is simply a complex number with a zero imaginary part. Hence, $\mathbb{R} \subset \mathbb{C}$. There is also a common geometric interpretation of a complex number as illustrated in Figure 2.1. In fact, the existence of this picture is why $\mathbb{C}$ is sometimes also referred to as “the complex plane”.

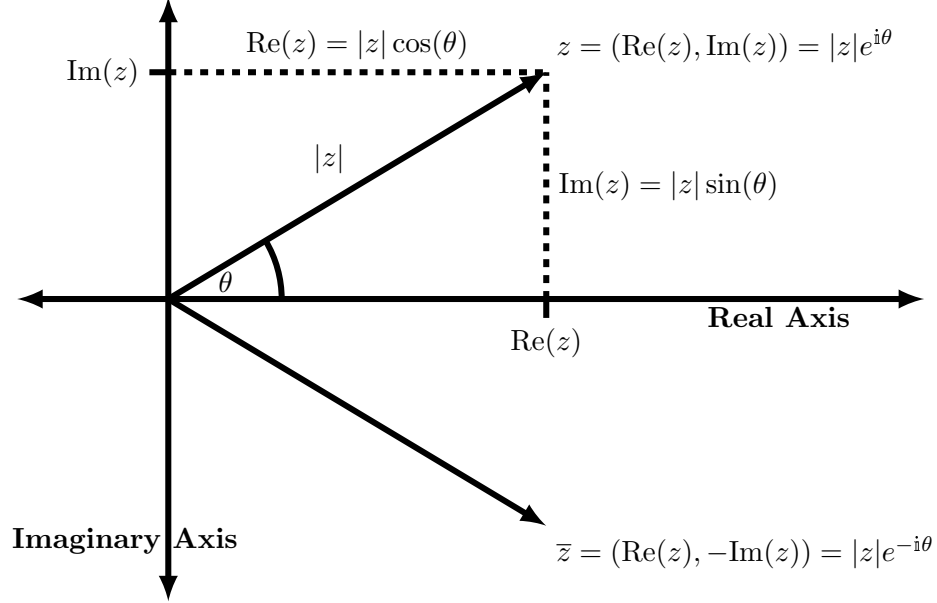


Figure 2.1: The geometry of a complex number $z \in \mathbb{C}$.

Figure 2.1 represents many of most important quantities related to a complex number $z = x + iy$ stemming from geometry. In particular, the **modulus**, **magnitude**, or **absolute value** of $z = x + iy$ is denoted by $|z|$. It is defined to be the Euclidean distance from the origin to $(\operatorname{Re}(z), \operatorname{Im}(z)) = (x, y)$ in the complex plane. It is therefore also the length of the hypotenuse of a right triangle whose other two sides have lengths $|\operatorname{Re}(z)|$ and $|\operatorname{Im}(z)|$, and so can be computed using the Pythagorean theorem to be

$$|z| = \sqrt{(\operatorname{Re}(z))^2 + (\operatorname{Im}(z))^2} = \sqrt{x^2 + y^2}.$$

Note that if $z \in \mathbb{R}$ so that $z = \operatorname{Re}(z)$ (i.e., if $y = 0$) then $|z| = |\operatorname{Re}(z)| = |x|$. That is, this definition extends the usual definition of absolute value over the real numbers $\mathbb{R}$ to all of $\mathbb{C}$.

Exercise 2.1.1. Let $z \in \mathbb{C}$. Prove that $|\operatorname{Re}(z)| \leq |z|$ and $|\operatorname{Im}(z)| \leq |z|$ always hold.

Another fundamental geometric quantity illustrated in Figure 2.1 related to $z = x + iy$ is its **phase angle** or **argument**, $\theta = \arg(z) \in [0, 2\pi)$, defined to be the angle between the real axis and the vector from the origin to $(\operatorname{Re}(z), \operatorname{Im}(z)) = (x, y)$ in the complex plane.

Using the geometric definitions of $\sin$ and $\cos$ involving right triangles one can immediately derive the formulas

$$x = \operatorname{Re}(z) = |z| \cos \theta \quad \text{and} \quad y = \operatorname{Im}(z) = |z| \sin \theta.$$

Similarly, one can appeal to trigonometry to see that, e.g, the phase angle θ of $z = x + \mathrm{i}y$ is

$$\theta = \arg(z) := \cos^{-1} \left(\frac{\operatorname{Re}(z)}{|z|} \right) = \cos^{-1} \left(\frac{x}{\sqrt{x^2 + y^2}} \right),$$

where one needs to remember to correct θ based on the quadrant of the complex plane z belongs to in the usual way. Note that positive real numbers (with $\operatorname{sign} 1 = \cos(0)$) always have the phase angle $\theta = 0$, and that negative real numbers (with $\operatorname{sign} -1 = \cos(\pi)$) always have the phase angle $\theta = \pi$. Hence, phase angles effectively extend the notion of “sign” from the real numbers $\mathbb{R}$ to all of $\mathbb{C}$ in a consistent fashion.

Two complex numbers $z_1 = x_1 + \mathrm{i}y_1$ and $z_2 = x_2 + \mathrm{i}y_2$ can be added component-wise (effectively as vectors) via the definition

$$z_1 + z_2 = (x_1 + \mathrm{i}y_1) + (x_2 + \mathrm{i}y_2) := (x_1 + x_2) + \mathrm{i}(y_1 + y_2) = \operatorname{Re}(z_1) + \operatorname{Re}(z_2) + \mathrm{i}(\operatorname{Im}(z_1) + \operatorname{Im}(z_2)).$$

Note again that the usual relationship between $\mathbb{R}$ and $\mathbb{C}$ holds: if $z_1, z_2 \in \mathbb{R}$ so that $\operatorname{Im}(z_1) = \operatorname{Im}(z_2) = 0$ then this definition of addition matches addition in $\mathbb{R}$. We have once again managed to extend the usual definition (of addition here) from $\mathbb{R}$ to all of $\mathbb{C}$ in a totally consistent way.

Similarly, two complex numbers $z_1, z_2 \in \mathbb{C}$ can be multiplied using the standard distributive law for the multiplication of two real numbers, but making sure to use the identity $\mathrm{i}^2 = -1$. Indeed, if $z_1 = x_1 + \mathrm{i}y_1$ and $z_2 = x_2 + \mathrm{i}y_2$, then

$$\begin{aligned} z_1 z_2 &= (x_1 + \mathrm{i}y_1)(x_2 + \mathrm{i}y_2) := x_1 x_2 + \mathrm{i}x_1 y_2 + \mathrm{i}y_1 x_2 + \mathrm{i}^2 y_1 y_2 \\ &= (x_1 x_2 - y_1 y_2) + \mathrm{i}(x_1 y_2 + y_1 x_2). \end{aligned}$$

Note once again that this definition of multiplication matches multiplication over the reals whenever $z_1, z_2 \in \mathbb{R}$ so that $y_1 = y_2 = 0 = \operatorname{Im}(z_1) = \operatorname{Im}(z_2)$. This, of course, allows us to compute powers of $z \in \mathbb{C}$, z^n , for any positive integer n in a way that is again a consistent extension of how one computes powers of real numbers.

Exercise 2.1.2. *Verify the following properties of complex number addition and multiplication.*

1. **Commutativity of addition and multiplication:** $z_1 + z_2 = z_2 + z_1$ and $z_1 z_2 = z_2 z_1$ for all $z_1, z_2 \in \mathbb{C}$.
2. **Associativity of addition and multiplication:** $(z_1 + z_2) + z_3 = z_1 + (z_2 + z_3)$ and $(z_1 z_2) z_3 = z_1 (z_2 z_3)$ for all $z_1, z_2, z_3 \in \mathbb{C}$.

3. **Distributivity:** $z_1(z_2 + z_3) = z_1z_2 + z_1z_3$ for all $z_1, z_2, z_3 \in \mathbb{C}$.

Exercise 2.1.3. Let $z_1, z_2 \in \mathbb{C}$. Show that $|z_1z_2| = |z_1||z_2|$.

Importantly, we can now see that more complicated functions that can be defined on $\mathbb{R}$ in terms of series expansions (like $\exp$, $\cos$, $\sin$, ...) should also believably extend in a consistent way to all of $\mathbb{C}$ since all of their basic building blocks (addition, multiplication, and integer powers) have been consistently extended from $\mathbb{R}$ to all of $\mathbb{C}$.¹ Recall the Taylor series for the exponential function centered at 0 is

$$\exp(x) = e^x = \sum_{n=0}^{\infty} \frac{x^n}{n!}.$$

This series converges absolutely for every $x \in \mathbb{R}$. The exponential function of any complex number $z \in \mathbb{C}$ can be defined analogously as

$$\exp(z) = e^z = \sum_{n=0}^{\infty} \frac{z^n}{n!}. \quad (2.1)$$

It matches the usual definition of $\exp$ on $\mathbb{R}$ (i.e., whenever $z \in \mathbb{R} \subset \mathbb{C}$) for all the reasons emphasized above. For more of these interesting properties on $\mathbb{C}$ we recommend taking a look at, e.g., [6].

2.1.1 Euler's Identity

We may now derive Euler's identity for complex exponentials. Consider the purely imaginary number $z = \mathbf{i}\theta$ for some $\theta \in \mathbb{R}$. In this case, we have

$$\exp(\mathbf{i}\theta) = e^{\mathbf{i}\theta} = \sum_{n=0}^{\infty} \frac{(\mathbf{i}\theta)^n}{n!}. \quad (2.2)$$

Before simplifying the above expression, we observe that since $\mathbf{i}^2 = -1$, we have

$$\mathbf{i}^{2n} = (-1)^n \quad \text{and} \quad \mathbf{i}^{2n+1} = (-1)^n \mathbf{i} \quad \text{for all } n \geq 0.$$

¹If you want to learn about how very nicely this idea ends up working out, I strongly recommend taking a class on complex analysis!

Breaking the sum (2.2) into the parts where n is even and n is odd we get that

$$\begin{aligned}
 e^{i\theta} &= \sum_{n \text{ even}} \frac{(i\theta)^n}{n!} + \sum_{n \text{ odd}} \frac{(i\theta)^n}{n!} \\
 &= \sum_{n=0}^{\infty} \frac{(i\theta)^{2n}}{(2n)!} + \sum_{n=0}^{\infty} \frac{(i\theta)^{2n+1}}{(2n+1)!} \\
 &= \sum_{n=0}^{\infty} \frac{i^{2n}\theta^{2n}}{(2n)!} + \sum_{n=0}^{\infty} \frac{i^{2n+1}\theta^{2n+1}}{(2n+1)!} \\
 &= \sum_{n=0}^{\infty} \frac{(-1)^n\theta^{2n}}{(2n)!} + \sum_{n=0}^{\infty} \frac{(-1)^n i\theta^{2n+1}}{(2n+1)!} \\
 &= \left(\sum_{n=0}^{\infty} \frac{(-1)^n\theta^{2n}}{(2n)!} \right) + i \left(\sum_{n=0}^{\infty} \frac{(-1)^n\theta^{2n+1}}{(2n+1)!} \right).
 \end{aligned}$$

Now recall from calculus that the Taylor series for $\cos \theta$ and $\sin \theta$ about 0 are

$$\cos \theta = \sum_{n=0}^{\infty} \frac{(-1)^n\theta^{2n}}{(2n)!} \quad \text{and} \quad \sin \theta = \sum_{n=0}^{\infty} \frac{(-1)^n\theta^{2n+1}}{(2n+1)!},$$

and that the series above converge absolutely for all $\theta \in \mathbb{R}$. Consequently, we obtain **Euler's identity**, i.e., that

$$e^{i\theta} = \cos \theta + i \sin \theta. \quad (2.3)$$

Exercise 2.1.4. Use Euler's identity and trigonometric identities involving sine and cosine to show that $e^{i\theta}e^{i\omega} = e^{i(\theta+\omega)}$ holds for all $\omega, \theta \in \mathbb{R}$.

Exercise 2.1.5. Use induction in addition to the last exercise to prove that $(e^{i\theta})^n = e^{in\theta}$ holds for all $n \in \mathbb{N}$.

2.1.2 The Polar Representation of a Complex Number

As illustrated in Figure 2.1, every $z = x + iy \in \mathbb{C}$ corresponds to a point $(x, y) = (\operatorname{Re}(z), \operatorname{Im}(z))$ in the complex plane. This suggests another way of representing a complex number using polar coordinates as done for $\mathbb{R}^2$. Specifically, if $x = r \cos \theta$ and $y = r \sin \theta$ for some $r \geq 0$ and $\theta \in [0, 2\pi)$, then a complex number $z = x + iy$ can be represented in terms of r and θ as follows:

$$z = x + iy = r \cos \theta + ir \sin \theta = r (\cos \theta + i \sin \theta). \quad (2.4)$$

Note that the identity $\cos^2(\theta) + \sin^2(\theta) = 1$ shows that $r = |z|$ in the polar representation above.

Using Euler's identity in (2.4) gives us the polar representation of a complex number in terms of complex exponentials:

$$z = r\mathrm{e}^{\mathrm{i}\theta}.$$

Stating the same formula another way, we have that

$$z = \operatorname{Re}(z) + \mathrm{i}\operatorname{Im}(z) = |z|\mathrm{e}^{\mathrm{i}\arg(z)}.$$

Exercise 2.1.6. *Prove the following useful identities involving the polar representation of complex numbers.*

1. Show that if $z = r\mathrm{e}^{\mathrm{i}\theta}$, then for any $n \in \mathbb{N}$ we have $z^n = r^n (\cos(n\theta) + \mathrm{i}\sin(n\theta))$.
2. Every complex number of unit modulus can be written as $\mathrm{e}^{\mathrm{i}\theta}$ for some $\theta \in [0, 2\pi)$.
3. If $z = r\mathrm{e}^{\mathrm{i}\theta}$ and $w = s\mathrm{e}^{\mathrm{i}\varphi}$, then $zw = rs\mathrm{e}^{\mathrm{i}(\theta+\varphi)}$.

2.1.3 Complex Conjugation

The **complex conjugate** of a complex number $z = x + \mathrm{i}y$ is the complex number $\bar{z} = x - \mathrm{i}y$. Geometrically, as illustrated in Figure 2.1, $\bar{z}$ is the reflection of z across the real axis. One can verify that

$$\operatorname{Re}(z) = \frac{z + \bar{z}}{2}, \quad \operatorname{Im}(z) = \frac{z - \bar{z}}{2\mathrm{i}}.$$

Consequently, a complex number z is a real number if and only if $z = \bar{z}$.

Exercise 2.1.7. *Prove the following useful identities involving complex conjugation. Let $z_1, z_2 \in \mathbb{C}$.*

1. Show that $\overline{z_1 + z_2} = \bar{z}_1 + \bar{z}_2$.
2. Show that $\overline{z_1 z_2} = \bar{z}_1 \bar{z}_2$.
3. Show that $|z_1|^2 = z_1 \bar{z}_1$.
4. Show that $|z_1 + z_2|^2 = |z_1|^2 + |z_2|^2 + 2\operatorname{Re}(z_1 \bar{z}_2)$.

Exercise 2.1.8. *Let $z \in \mathbb{C}$ have the polar representation $z = r\mathrm{e}^{\mathrm{i}\theta}$. Show that $\bar{z} = r\mathrm{e}^{-\mathrm{i}\theta}$.*

Exercise 2.1.9. *Let $z \in \mathbb{C}$ be nonzero. Show that*

$$z^{-1} := \frac{1}{z} = \frac{\bar{z}}{|z|^2}.$$

2.1.4 The Roots of Unity

Fix $n \in \mathbb{N}$ and consider the equation

$$z^n = 1, \quad z \in \mathbb{C}.$$

We wish to find all solutions $z \in \mathbb{C}$ of this equation. First, notice that necessarily $|z| = 1$. Hence, $z = e^{i\theta}$ for $\theta \in [0, 2\pi)$. By Euler's formula, this means that

$$1 = \cos(n\theta) + i \sin(n\theta),$$

which implies $\cos(n\theta) = 1$ and $\sin(n\theta) = 0$. This can only happen if $n\theta = 2\pi k$ for some $k \in \mathbb{Z}$. Thus, $\theta = \frac{2\pi k}{n}$ must hold. Note that we have n distinct values of $\theta \in [0, 2\pi)$ satisfying this formula, one for each $k \in [n]$. The set of these solutions is therefore $\{e^{\frac{2\pi k i}{n}} \mid k \in [n]\}$ are called the n^{th} **roots of unity**. Notice that they all satisfy $z^n = 1$ by design, and are placed in an equidistant fashion around the unit circle $|z| = 1$ in the complex plane. These values will be of special significance later in Section 3.4.

2.1.5 The Triangle Inequality for Complex Numbers

We will now prove the triangle inequality for complex numbers.

Lemma 2.1.1 (The Triangle Inequality for $\mathbb{C}$).

$$|z_1 + z_2| \leq |z_1| + |z_2| \quad \text{for all } z_1, z_2 \in \mathbb{C}. \quad (2.5)$$

Furthermore, equality holds in (2.5) if and only if $z_1 = cz_2$ for some real number $c \geq 0$.

Proof. Using the results of Exercises 2.1.7, 2.1.1, and 2.1.3 we can see that

$$\begin{aligned} |z_1 + z_2|^2 &= |z_1|^2 + |z_2|^2 + 2\operatorname{Re}(z_1 \overline{z_2}) \\ &\leq |z_1|^2 + |z_2|^2 + 2|z_1 \overline{z_2}| \\ &= |z_1|^2 + |z_2|^2 + 2|z_1||\overline{z_2}| \\ &= |z_1|^2 + |z_2|^2 + 2|z_1||z_2| \\ &= (|z_1| + |z_2|)^2. \end{aligned}$$

Taking square roots now gives us the desired inequality.

Now suppose that we have equality in (2.5). If either z_1 or z_2 is 0 we are finished. Thus, suppose that $z_1 = r_1 e^{i\theta_1}$ and $z_2 = r_2 e^{i\theta_2}$ with $r_1, r_2 > 0$. Notice that the only place we have an inequality in the argument above is in the estimate $\operatorname{Re}(z_1 \overline{z_2}) \leq |z_1 \overline{z_2}|$. If $|z_1 + z_2| = |z_1| + |z_2|$, then we must, in fact, have $\operatorname{Re}(z_1 \overline{z_2}) = |z_1 \overline{z_2}|$. That means $\operatorname{Im}(z_1 \overline{z_2}) = r_1 r_2 \sin(\theta_1 - \theta_2) = 0$. Given that $r_1, r_2 > 0$ we must therefore have $\sin(\theta_1 - \theta_2) = 0$, implying that $\theta_1 - \theta_2 = m\pi$ for some integer m .

If m were odd it would imply that

$$\operatorname{Re}(z_1 \overline{z_2}) = r_1 r_2 \cos(\theta_1 - \theta_2) = r_1 r_2 \cos(m\pi) = -r_1 r_2 < 0.$$

This is impossible here since we have $\operatorname{Re}(z_1 \overline{z_2}) = |z_1 z_2| = r_1 r_2 > 0$. Therefore, m must be even, and so $\theta_1 - \theta_2 = 2\pi n$ for some integer n . This implies that $\arg(z_1) = \arg(z_2)$, and so $z_1 = cz_2$ for some positive real number c . $\square$

We now have all the prerequisites we need to begin discussing linear algebra over $\mathbb{C}$.

2.2 Basic Linear Algebra over $\mathbb{C}$ and $\mathbb{R}$

A complex valued matrix $A \in \mathbb{C}^{m \times n}$ is a matrix of complex values with m rows and n columns whose entries are denoted by $A_{j,k} \in \mathbb{C}$ for all $j \in [m]$ and $k \in [n]$. A complex valued vector $\mathbf{x} \in \mathbb{C}^n$ of length n is also considered to be an $n \times 1$ matrix (i.e, vectors are “column vectors” by default). Its entries are denoted by $x_j \in \mathbb{C}$ for all $j \in [n]$, and can themselves be safely considered to be scalars, length 1 vectors, and 1×1 matrices as convenient. Matrices (and vectors) are always added entrywise, and scalar-vector/scalar-matrix multiplication is also always performed entrywise, as in your first linear algebra course.

Given a matrix $A \in \mathbb{C}^{m \times n}$ and a vector $\mathbf{x} \in \mathbb{C}^n$, their **matrix-vector product**, $A\mathbf{x} \in \mathbb{C}^m$, is a vector which can be defined in two equivalent ways. First, it can be defined entrywise via

$$(A\mathbf{x})_j := \sum_{k \in [n]} A_{j,k} x_k \in \mathbb{C} \quad \forall j \in [m]. \quad (2.6)$$

Alternatively, it can be defined as a weighted sum of the columns of A via the formula

$$A\mathbf{x} = \sum_{k \in [n]} x_k A_{:,k} \in \mathbb{C}^m. \quad (2.7)$$

Both equations are true and will be used often below.

Exercise 2.2.1. *Show that (2.6) holds if and only if (2.7) holds.*

We can use the matrix-vector product notation to describe the **product of two matrices**. For any natural numbers m, n, p , and two matrices $A \in \mathbb{C}^{m \times n}$ and $B \in \mathbb{C}^{n \times p}$, we can define their product columnwise by $(AB)_{:,k} := A(B_{:,k}) = A\mathbf{b}_k \quad \forall k \in [p]$, where $\mathbf{b}_k = B_{:,k}$ denotes the k^{th} column of B , and $A\mathbf{b}_k$ is a matrix-vector product as per (2.6). Equivalently we may write

$$AB = A \begin{pmatrix} | & & | \\ \mathbf{b}_0 & \cdots & \mathbf{b}_{p-1} \\ | & & | \end{pmatrix} = \begin{pmatrix} | & & | \\ A\mathbf{b}_0 & \cdots & A\mathbf{b}_{p-1} \\ | & & | \end{pmatrix}. \quad (2.8)$$

Similarly, we may also define matrix-matrix multiplication entrywise by

$$(AB)_{j,k} := \sum_{l=0}^{n-1} A_{j,l} B_{l,k}. \quad (2.9)$$

Note further that since we always consider a vector $\mathbf{v} \in \mathbb{C}^n$ to be an $n \times 1$ matrix, we should check that the resulting matrix-matrix product $A\mathbf{v}$ agrees with the matrix-vector product definition of $A\mathbf{v}$ above. It does – check!

Exercise 2.2.2. *Show that (2.8) holds if and only if (2.9) holds.*

Matrix-vector multiplication further allows us to view matrices as functions. Given $A \in \mathbb{C}^{m \times n}$, A acts on $\mathbb{C}^n$ by vector multiplication, and can therefore be viewed as a map $A : \mathbb{C}^n \rightarrow \mathbb{C}^m$ defined by $A(\mathbf{x}) = A\mathbf{x}$ for all $\mathbf{x} \in \mathbb{C}^n$. One can confirm that A is then a linear function (i.e., that $A(\alpha\mathbf{x} + \beta\mathbf{y}) = \alpha A\mathbf{x} + \beta A\mathbf{y}$ for all $\alpha, \beta \in \mathbb{C}$ and $\mathbf{x}, \mathbf{y} \in \mathbb{C}^n$), and that the range of A (as a function) is the column space of A (as a matrix). That is,

$$\begin{aligned} \text{Range}(A) &:= \{A\mathbf{x} \mid \mathbf{x} \in \mathbb{C}^n\} \subset \mathbb{C}^m \\ &= \left\{ \sum_{j \in [n]} \alpha_j A_{:,j} \mid \alpha_j \in \mathbb{C} \ \forall j \in [n] \right\} =: \text{span}(\{A_{:,j} \mid j \in [n]\}) \\ &=: \mathcal{C}(A) = \text{Column Space of } A. \end{aligned}$$

Let $A \in \mathbb{C}^{m \times n}$ be a matrix. The **adjoint** of A , denoted $A^* \in \mathbb{C}^{n \times m}$, is the conjugate transpose of A , i.e., the matrix produced by transposing A and taking the complex conjugate of each entry. It is defined entrywise by $(A^*)_{j,k} = \overline{A_{k,j}}$. Note that if $A \in \mathbb{R}^{m \times n}$ then $A^* = A^T$. We also note that $A = (A^*)^*$ always holds (check this!).

Exercise 2.2.3. *Let $A, B \in \mathbb{C}^{m \times n}$. Show that $(A + B)^* = A^* + B^*$.*

Exercise 2.2.4. *Let $A \in \mathbb{C}^{m \times n}$ and $B \in \mathbb{C}^{n \times p}$. Show that $(AB)^* = B^* A^*$.*

Given two vectors of the same length, $\mathbf{x}, \mathbf{y} \in \mathbb{C}^n$, we can define their Euclidean **inner product** to be

$$\langle \mathbf{x}, \mathbf{y} \rangle := \sum_{j=0}^{n-1} \overline{x_j} y_j = \mathbf{x}^* \mathbf{y} \in \mathbb{C}.$$

Also note that when two vectors $\mathbf{x}$ and $\mathbf{y}$ are real-valued, the complex inner product of $\mathbf{x}$ and $\mathbf{y}$ equals the real inner product of $\mathbf{x}$ and $\mathbf{y}$. Thus, we can view linear algebra over the complex numbers as a natural extension of linear algebra over the reals, where any statement about complex linear algebra still holds true when we restrict ourselves to the real numbers. This again supports my prior claim that you can simply “replace $\mathbb{C}$ everywhere

in this section with $\mathbb{R}$ ” and have a chapter on linear algebra over $\mathbb{R}$ as a result, should you desire to do so.

These next four exercises are highly recommended. As always, using the result of prior exercises to will help you complete subsequent ones more quickly is also always highly recommended.

Exercise 2.2.5. Let $\mathbf{x}, \mathbf{y} \in \mathbb{C}^n$. Show that $\langle \mathbf{x}, \mathbf{y} \rangle = \overline{\langle \mathbf{y}, \mathbf{x} \rangle}$.

Exercise 2.2.6. Show that the inner product is conjugate-linear in the first argument and linear in the second argument. That is, for $\alpha, \beta \in \mathbb{C}$ and $\mathbf{x}, \mathbf{y}, \mathbf{z} \in \mathbb{C}^n$ show that

$$1. \langle \alpha \mathbf{x} + \beta \mathbf{y}, \mathbf{z} \rangle = \overline{\alpha} \langle \mathbf{x}, \mathbf{z} \rangle + \overline{\beta} \langle \mathbf{y}, \mathbf{z} \rangle, \text{ and that}$$

$$2. \langle \mathbf{x}, \alpha \mathbf{y} + \beta \mathbf{z} \rangle = \alpha \langle \mathbf{x}, \mathbf{y} \rangle + \beta \langle \mathbf{x}, \mathbf{z} \rangle.$$

Exercise 2.2.7. Let $A \in \mathbb{C}^{m \times n}$ and $\mathbf{x} \in \mathbb{C}^n$. Show that $(A\mathbf{x})_j = \langle (A^*)_{:,j}, \mathbf{x} \rangle = \langle \overline{A_{j,:}}, \mathbf{x} \rangle$ for all $j \in [m]$.

Exercise 2.2.8. Let $A \in \mathbb{C}^{m \times n}$, $\mathbf{x} \in \mathbb{C}^n$, and $\mathbf{y} \in \mathbb{C}^m$. Show that both $\langle A\mathbf{x}, \mathbf{y} \rangle = \langle \mathbf{x}, A^*\mathbf{y} \rangle$ and $\langle A^*\mathbf{y}, \mathbf{x} \rangle = \langle \mathbf{y}, A\mathbf{x} \rangle$ hold.

Let’s now briefly review a geometric concept related to inner products that’s reserved for real-valued vectors.

2.2.1 Some Inner Product Geometry for Real-valued Vectors $\mathbf{x}, \mathbf{y} \in \mathbb{R}^n$

The inner product can be used to express the angle between two real vectors. Given two non-zero vectors $\mathbf{x}, \mathbf{y} \in \mathbb{R}^n$, the **angle** $\theta \in [0, \pi]$ **between** $\mathbf{x}, \mathbf{y} \in \mathbb{R}^n$ is

$$\theta = \cos^{-1} \left(\frac{\langle \mathbf{x}, \mathbf{y} \rangle}{\sqrt{\langle \mathbf{x}, \mathbf{x} \rangle \langle \mathbf{y}, \mathbf{y} \rangle}} \right). \quad (2.10)$$

Note that $\theta = \pi/2$ (or 90 degrees) whenever $\langle \mathbf{x}, \mathbf{y} \rangle = 0$, indicating that the two vectors are perpendicular, or orthogonal, to one another. Further note that the angle between $\mathbf{x}$ and $\mathbf{y}$ can always be reasoned about with regular two-dimensional plane geometry no matter how large n is here since $\mathbf{x}$ and $\mathbf{y}$ will always belong to the (at most) two-dimensional subspace $\text{span}\{\mathbf{x}, \mathbf{y}\} \subset \mathbb{R}^n$. Hence, all the pictures of right triangles you are tempted to draw on a piece of paper to better understand θ are 100% justified.²

Back to $\mathbb{C}^n$: Using the inner product geometry for real-valued vectors as motivation, we will also say that **two complex-valued vectors $\mathbf{x}, \mathbf{y} \in \mathbb{C}^n$ are orthogonal** if $\langle \mathbf{x}, \mathbf{y} \rangle = 0$. We will now recall an important inequality for inner products.

²Simply rewrite $\mathbf{x}$ and $\mathbf{y}$ in terms of an orthonormal basis of the $\text{span}\{\mathbf{x}, \mathbf{y}\}$, and then draw your pictures with axes in the directions of these orthonormal basis vectors. If this footnote is confusing I recommend you continue on, review orthogonality and orthogonal projections, and then come back here again for a rematch.

2.2.2 The Cauchy–Schwarz Inequality

Note that for any vector $\mathbf{x} \in \mathbb{C}^n$, $\langle \mathbf{x}, \mathbf{x} \rangle = \sum_{j \in [n]} x_j \overline{x_j} = \sum_{j \in [n]} |x_j|^2 \geq 0$ (this fact will become important later). Now let $t \in \mathbb{R}$, and $\mathbf{x}, \mathbf{y} \in \mathbb{C}^n$, and set $\alpha := \frac{\langle \mathbf{y}, \mathbf{x} \rangle}{\langle \mathbf{y}, \mathbf{x} \rangle}$. One can see that α is a complex number with magnitude 1. Finally, define the function $f : \mathbb{R} \rightarrow \mathbb{R}$ by

$$f(t) := \langle t\alpha\mathbf{x} + \mathbf{y}, t\alpha\mathbf{x} + \mathbf{y} \rangle$$

Recall that $f(t) \geq 0$ for all $t \in \mathbb{R}$ by the fact above.

Continuing, the following sequence of inequalities can be seen to hold using properties of the inner product together with the definition of α (check each step!). We have that

$$\begin{aligned} 0 \leq f(t) &= t\overline{\alpha}\langle \mathbf{x}, t\alpha\mathbf{x} + \mathbf{y} \rangle + \langle \mathbf{y}, t\alpha\mathbf{x} + \mathbf{y} \rangle \\ &= t^2\overline{\alpha}\alpha\langle \mathbf{x}, \mathbf{x} \rangle + t\overline{\alpha}\langle \mathbf{x}, \mathbf{y} \rangle + t\alpha\langle \mathbf{y}, \mathbf{x} \rangle + \langle \mathbf{y}, \mathbf{y} \rangle \\ &= t^2\langle \mathbf{x}, \mathbf{x} \rangle + 2\operatorname{Re}(t\alpha\langle \mathbf{y}, \mathbf{x} \rangle) + \langle \mathbf{y}, \mathbf{y} \rangle \\ &= \langle \mathbf{x}, \mathbf{x} \rangle t^2 + 2|\langle \mathbf{x}, \mathbf{y} \rangle|t + \langle \mathbf{y}, \mathbf{y} \rangle, \end{aligned}$$

which is a quadratic polynomial in t with real coefficients. Since the polynomial f above is ≥ 0 for all t , it must have at most one real root.

Recalling the quadratic equation for a generic polynomial $p(t) = at^2 + bt + c$, we note that its discriminant $b^2 - 4ac$ must be non-positive (i.e., ≤ 0) in order for the polynomial to have at most one real root. Applying this to our f above we learn that

$$\begin{aligned} (2|\langle \mathbf{x}, \mathbf{y} \rangle|)^2 &\leq 4\langle \mathbf{x}, \mathbf{x} \rangle \langle \mathbf{y}, \mathbf{y} \rangle \\ |\langle \mathbf{x}, \mathbf{y} \rangle| &\leq \sqrt{\langle \mathbf{x}, \mathbf{x} \rangle \langle \mathbf{y}, \mathbf{y} \rangle}. \end{aligned}$$

This inequality holds for all vectors $\mathbf{x}, \mathbf{y} \in \mathbb{C}^n$ since we chose them arbitrarily. It is known as the Cauchy-Schwarz Inequality (i.e., it has a name!) due to its importance.

Lemma 2.2.1 (The Cauchy-Schwarz Inequality). *For any two vectors $\mathbf{x}, \mathbf{y} \in \mathbb{C}^n$,*

$$|\langle \mathbf{x}, \mathbf{y} \rangle| \leq \sqrt{\langle \mathbf{x}, \mathbf{x} \rangle} \sqrt{\langle \mathbf{y}, \mathbf{y} \rangle}.$$

It is expressed here slightly differently than usual in Lemma 2.2.1, however. Usually it is stated like “ $|\langle \mathbf{x}, \mathbf{y} \rangle| \leq \|\mathbf{x}\|_2 \|\mathbf{y}\|_2$ ” where $\|\cdot\|_2$ denotes the ℓ^2 -vector norm, which we will recall next.

2.2.3 General Norms on $\mathbb{C}^{m \times n}$, and the Euclidean Vector Norm

A **matrix norm on $\mathbb{C}^{m \times n}$** is a function $f : \mathbb{C}^{m \times n} \rightarrow \mathbb{R}^+ := [0, \infty)$ satisfying all of the following properties:

1. (The triangle inequality): $f(A + B) \leq f(A) + f(B)$ for all $A, B \in \mathbb{C}^{m \times n}$,

2. $f(\alpha A) = |\alpha|f(A)$ for all $\alpha \in \mathbb{C}$ and $A \in \mathbb{C}^{m \times n}$, and
3. $f(A) = 0 \iff A = 0_{m \times n}$, where $0_{m \times n}$ denotes the $m \times n$ matrix of all zeros (i.e., the zero matrix).

Recall that we also view vectors in $\mathbb{C}^m$ as $m \times 1$ matrices. Thus, a norm on $m \times 1$ matrices (i.e., on vectors in $\mathbb{C}^m$) will also be called a **vector norm** for this reason.

We can now see that the **Euclidean**, or ℓ^2 -**norm**, of a vector $\mathbf{x} \in \mathbb{C}^n$ defined by $\|\mathbf{x}\|_2 := \sqrt{\langle \mathbf{x}, \mathbf{x} \rangle}$ is indeed a vector norm.

Lemma 2.2.2. *Let $f : \mathbb{C}^n \rightarrow \mathbb{R}^+$ be the ℓ^2 -norm so that $f(x) = \|\mathbf{x}\|_2 := \sqrt{\langle \mathbf{x}, \mathbf{x} \rangle}$. We claim that $f(x) = \|\mathbf{x}\|_2$ is a vector norm on $\mathbb{C}^n$.*

Proof. We will verify that each condition of a norm is satisfied. First, we will check that the triangle inequality holds. Note that the last inequality just below depends on the Cauchy-Schwarz inequality. Let $\mathbf{x}, \mathbf{y} \in \mathbb{C}^n$. Then

$$\begin{aligned}
 \|\mathbf{x} + \mathbf{y}\|_2 &= \sqrt{\langle \mathbf{x} + \mathbf{y}, \mathbf{x} + \mathbf{y} \rangle} \\
 &= \sqrt{\|\mathbf{x}\|_2^2 + \|\mathbf{y}\|_2^2 + 2\operatorname{Re}(\langle \mathbf{x}, \mathbf{y} \rangle)} \\
 &\leq \sqrt{\|\mathbf{x}\|_2^2 + \|\mathbf{y}\|_2^2 + 2|\langle \mathbf{x}, \mathbf{y} \rangle|} \\
 &\leq \sqrt{\|\mathbf{x}\|_2^2 + \|\mathbf{y}\|_2^2 + 2\|\mathbf{x}\|_2^2\|\mathbf{y}\|_2^2} \\
 &= \|\mathbf{x}\|_2 + \|\mathbf{y}\|_2.
 \end{aligned}$$

Next we verify that the norm scales correctly. Let $\alpha \in \mathbb{C}$ and $\mathbf{x} \in \mathbb{C}^n$. We have that

$$\begin{aligned}
 \|\alpha \mathbf{x}\|_2 &= \sqrt{\langle \alpha \mathbf{x}, \alpha \mathbf{x} \rangle} = \sqrt{\alpha \bar{\alpha} \langle \mathbf{x}, \mathbf{x} \rangle} = \sqrt{|\alpha|^2 \langle \mathbf{x}, \mathbf{x} \rangle} \\
 &= |\alpha| \|\mathbf{x}\|_2.
 \end{aligned}$$

Finally, we verify that the ℓ^2 -norm of a vector $\mathbf{x} \in \mathbb{C}^n$ can only be 0 if $\mathbf{x}$ is the vector of all zeros, $\mathbf{0}$. We have that

$$\|\mathbf{x}\|_2 = 0 \iff \|\mathbf{x}\|_2^2 = 0 \iff \sum_{j \in [n]} |x_j|^2 = 0 \iff |x_j| = 0 \ \forall j \in [n].$$

Having now shown that the ℓ^2 -norm satisfies all the properties of a norm, we may conclude that it indeed is one. $\square$

The following exercise demonstrates a useful property of the inner product which is perhaps most easily seen by using the properties of the ℓ^2 -norm.

Exercise 2.2.9. *Let $\mathbf{x} \in \mathbb{C}^n$. Show that if $\langle \mathbf{x}, \mathbf{y} \rangle = 0$ for all $\mathbf{y} \in \mathbb{C}^n$, then $\mathbf{x} = \mathbf{0}$ must hold.*

Though the ℓ^2 -norm is by far the most often used norm, all of the other norms in the following exercises are also commonly used. Even if you don't do each exercise (you should of course!), you should look at them for the norm definitions.

Exercise 2.2.10. Show that the ℓ^1 -norm defined by

$$\|A\|_1 := \sum_{j \in [m], k \in [n]} |A_{j,k}|$$

is indeed a norm on $\mathbb{C}^{m \times n}$.

Exercise 2.2.11. Show that the **Frobenius matrix norm** defined by

$$\|A\|_F := \sqrt{\sum_{j \in [m], k \in [n]} |A_{j,k}|^2}$$

is indeed a norm on $\mathbb{C}^{m \times n}$. HINT: Suppose you vectorize A . What does the Frobenius norm look like then?

Exercise 2.2.12. Show that the ℓ^∞ -norm defined by

$$\|A\|_\infty := \max_{j \in [m], k \in [n]} |A_{j,k}|$$

is indeed a norm on $\mathbb{C}^{m \times n}$.

Exercise 2.2.13. Show that the (ℓ^2, ℓ^2) -operator norm defined by

$$\|A\|_{2 \rightarrow 2} := \max_{\mathbf{x} \in \mathbb{C}^n \text{ s.t. } \|\mathbf{x}\|_2=1} \|A\mathbf{x}\|_2$$

is indeed a norm on $\mathbb{C}^{m \times n}$.

Exercise 2.2.14. Suppose that $f : \mathbb{C}^{m \times n} \rightarrow \mathbb{R}^+$ and $g : \mathbb{C}^{m \times n} \rightarrow \mathbb{R}^+$ are both norms on $\mathbb{C}^{m \times n}$. Let $\alpha, \beta \in \mathbb{R}^+ \setminus \{0\}$. Show that $h = \alpha f + \beta g$ will also be a norm on $\mathbb{C}^{m \times n}$.

With the aim in mind of recalling what the “rank” of a matrix really means, let's now briefly review linear independence and subspace basis properties.

2.3 Subspaces, Span, and Linear Independence

Let $S = \{\mathbf{v}_0, \dots, \mathbf{v}_{m-1}\} \subset \mathbb{C}^n$ be a finite and nonempty set of vectors in $\mathbb{C}^n$. The **span** of S , denoted $\text{span}(S)$, is the set

$$\text{span}(S) := \left\{ \sum_{j \in [m]} \alpha_j \mathbf{v}_j \mid \alpha_0, \dots, \alpha_{m-1} \in \mathbb{C} \right\} \subset \mathbb{C}^n.$$

If $S \subset \mathbb{C}^n$ is infinite, we instead define the span of S to be the set

$$\text{span}(S) := \bigcup_{A \subset S, A \text{ finite}} \text{span}(A) \subset \mathbb{C}^n.$$

Note that $S \subset \text{span}(S)$ always holds for any $S \subset \mathbb{C}^n$ since $\mathbf{x} \in S$ implies that $1 \cdot \mathbf{x} \in \text{span}(S)$. Furthermore, note that $\mathbf{0}$ is in the span of every nonempty set S since $\mathbf{0} = 0 \cdot \mathbf{x}$ for any $\mathbf{x} \in S$.

Exercise 2.3.1. Verify that if $A \subset S$, then $\text{span}(A) \subset \text{span}(S)$.

Exercise 2.3.2. Let $S, T \subset \mathbb{C}^n$. Verify that $\text{span}(T \cap S) \subset \text{span}(T) \cap \text{span}(S)$.

A subset $\mathcal{L} \subset \mathbb{C}^n$ is called a **linear subspace** of $\mathbb{C}^n$ if $\text{span}(\mathcal{L}) = \mathcal{L}$. That is, subspaces are sets that are closed under taking spans. Note that the so-called **trivial subspace** $\{\mathbf{0}\} \subset \mathbb{C}^n$ is always a subspace since $\text{span}(\{\mathbf{0}\}) = \{\mathbf{0}\}$. Similarly, $\mathbb{C}^n$ is a linear subspace because both of the following hold: (i) $\mathbb{C}^n \subset \text{span}(\mathbb{C}^n)$ (since $S \subset \text{span}(S)$ for any $S \subset \mathbb{C}^n$), and (ii) $\text{span}(\mathbb{C}^n) \subset \mathbb{C}^n$ (trivially by definition).

Lemma 2.3.1. The span of every nonempty subset $S \subset \mathbb{C}^n$ is a linear subspace of $\mathbb{C}^n$.

Proof. We need to show that

$$\text{span}(\text{span}(S)) = \text{span}(S).$$

As usual with set equalities of this type we will proceed by showing that both (i) $\text{span}(S) \subset \text{span}(\text{span}(S))$, and (ii) $\text{span}(\text{span}(S)) \subset \text{span}(S)$, hold. In fact (i) follows from the fact above that $S \subset \text{span}(S)$ holds for any $S \subset \mathbb{C}^n$. Hence, we only really need to verify (ii).

To verify that $\text{span}(\text{span}(S)) \subset \text{span}(S)$, let $\mathbf{y} \in \text{span}(\text{span}(S))$. By the definition of span , $\mathbf{y}$ must be the linear combination of a finite number $p \in \mathbb{N}$ of elements of $\text{span}(S)$. Hence, $\mathbf{y}$ will have the form

$$\mathbf{y} = \sum_{j=0}^{p-1} \beta_j \left(\sum_{k=0}^{q_j-1} \alpha_{j,k} \mathbf{x}_{j,k} \right) = \sum_{j=0}^{p-1} \sum_{k=0}^{q_j-1} \beta_j \alpha_{j,k} \mathbf{x}_{j,k},$$

where $\beta_j \in \mathbb{C}$ and $q_j \in \mathbb{N}$ for all $j \in [p]$, and where $\mathbf{x}_{j,k} \in S$ and $\alpha_{j,k} \in \mathbb{C}$ for all $k \in [q_j]$ for each $j \in [p]$. Thus, we can see that $\mathbf{y} \in \text{span}(S)$ too since it will be a linear combination of a finite number, $\min\left(\sum_{j \in [p]} q_j, |S|\right) \in \mathbb{N}$, of elements of S . $\square$

Given a matrix $A \in \mathbb{C}^{m \times n}$, recall that the **column space** of A , denoted by $\mathcal{C}(A)$, is defined to be $\mathcal{C}(A) := \text{span}(\{A_{:,j} \mid j \in [n]\}) \subset \mathbb{C}^m$. One important consequence of Lemma 2.3.1 is that the column space of every matrix $A \in \mathbb{C}^{m \times n}$ is a linear subspace of $\mathbb{C}^m$.

Exercise 2.3.3. Let $\mathcal{L}, \mathcal{K} \subset \mathbb{C}^n$ be two linear subspaces of $\mathbb{C}^n$. Show that $\mathcal{L} \cap \mathcal{K}$ is also a linear subspace of $\mathbb{C}^n$.

A finite set of vectors $\{\mathbf{v}_0, \dots, \mathbf{v}_{m-1}\} \subset \mathbb{C}^n$ is called **linearly independent** if $\sum_{j \in [m]} \alpha_j \mathbf{v}_j = \mathbf{0}$ if and only if $\alpha_j = 0$ for all j . In other words, no nontrivial (i.e., all zero) linear combination of the vectors can equal the zero vector. If a set of vectors is not linearly independent, we call it **linearly dependent**.

Exercise 2.3.4. Show that any set of vectors in $\mathbb{C}^n$ containing the zero vector is linearly dependent.

Exercise 2.3.5. Let $\mathbf{x} \in \mathbb{C}^n$ and suppose that $B \subset \mathbb{C}^n$ is a linearly independent set. Prove that if $\mathbf{x} \notin \text{span}(B)$, then $B' = B \cup \{\mathbf{x}\}$ is a new linearly independent set with $|B| + 1$ elements.

Definition 2.3.2 (The Standard Basis Vectors of $\mathbb{C}^n$). The **standard basis vectors** of $\mathbb{C}^n$ are the n vectors $\{\mathbf{e}_j\}_{j \in [n]} := \{\mathbf{e}_0, \mathbf{e}_1, \dots, \mathbf{e}_{n-1}\} \subset \mathbb{C}^n$ whose entries are given by

$$(\mathbf{e}_j)_k = \begin{cases} 1 & \text{if } j = k \\ 0 & \text{if } j \neq k \end{cases}$$

for all $k \in [n]$.

Example 2.3.3 (The Standard Basis Vectors are Linearly Independent). The standard basis vectors $\{\mathbf{e}_j\}_{j \in [n]} \subset \mathbb{C}^n$ are linearly independent because for any $\alpha_0, \alpha_1, \dots, \alpha_{n-1} \in \mathbb{C}$ we can see that

$$\mathbf{0} = \sum_{j \in [n]} \alpha_j \mathbf{e}_j = \begin{pmatrix} \alpha_0 \\ \alpha_1 \\ \vdots \\ \alpha_{n-1} \end{pmatrix} \iff \alpha_j = 0 \ \forall j \in [n].$$

Having just defined a set of vectors called the “standard basis”, it behooves us to briefly recall what a “basis” actually is. We do so next.

2.3.1 Bases, Orthonormal Bases, Dimension, and Rank

The following lemma ultimately guarantees that the notions of “dimension” and “rank” are well defined. Since these notions are inextricably linked to the notion of a “basis”, we will prepare the ground for them here.

Lemma 2.3.4 (The Exchange Lemma). Let $B_1, B_2 \subset \mathbb{C}^n$ be finite. Furthermore, suppose that B_2 is linearly independent, and that $\mathcal{L} := \text{span}(B_2) \subset \text{span}(B_1)$. Then $|B_2| \leq |B_1|$.

Proof. Suppose, towards a contradiction, that $|B_1| < |B_2|$. Let $B_1 = \{\mathbf{x}_0, \mathbf{x}_1, \dots, \mathbf{x}_{s-1}\} \subset \mathbb{C}^n$, and $B_2 = \{\mathbf{y}_0, \mathbf{y}_1, \dots, \mathbf{y}_{s+m-1}\} \subset \mathbb{C}^n$, where $m > 0$. Recall that the $\mathbf{y}_j$ vectors are linearly independent by assumption. Furthermore, we have the assumed inclusion $\mathcal{L} = \text{span}(B_2) \subset \text{span}(\{\mathbf{x}_0, \mathbf{x}_1, \dots, \mathbf{x}_{s-1}\})$.

Because $\mathbf{y}_0 \in \text{span}(B_1)$, there exist $\alpha_0, \dots, \alpha_{s-1} \in \mathbb{C}$ such that

$$\mathbf{y}_0 = \sum_{j \in [s]} \alpha_j \mathbf{x}_j.$$

Furthermore, because the $\mathbf{y}_j$ vectors are linearly independent, we recall that $\mathbf{y}_0$ can't be the zero vector. Hence, at least one of the α_j 's must be nonzero. Without loss of generality (w.l.g.), we may assume that $\alpha_0 \neq 0$. Thus, we can write $\mathbf{x}_0$ in terms of $\mathbf{y}_0$ and the other $\mathbf{x}_j$'s to see that

$$\mathbf{x}_0 = \frac{1}{\alpha_0} \left(\mathbf{y}_0 - \sum_{j=1}^{s-1} \alpha_j \mathbf{x}_j \right).$$

Hence, $\mathcal{L} \subset \text{span}(\{\mathbf{y}_0, \mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_{s-1}\})$ also holds. Note that we have effectively exchanged $\mathbf{x}_0$ for $\mathbf{y}_0$ in our initially assumed inclusion.

Now, we repeat this process to exchange $\mathbf{x}_1$ for $\mathbf{y}_1$ in the last inclusion just above: Since $\mathbf{y}_1 \in \mathcal{L}$, $\mathbf{y}_1 \in \text{span}(\{\mathbf{y}_0, \mathbf{x}_1, \dots, \mathbf{x}_{s-1}\})$. Thus, there exists $\beta_0 \in \mathbb{C}$ and $\gamma_1, \dots, \gamma_{s-1} \in \mathbb{C}$ such that

$$\mathbf{y}_1 = \beta_0 \mathbf{y}_0 + \sum_{j=1}^{s-1} \gamma_j \mathbf{x}_j.$$

Note that at least one $\gamma_j \in \mathbb{C}$ above must be nonzero (otherwise, we'd have $\mathbf{y}_1 = \beta_0 \mathbf{y}_0$, violating the assumed linear independence of the $\mathbf{y}_j$'s). Without loss of generality, we may assume that $\gamma_1 \neq 0$. Thus, we can write

$$\mathbf{x}_1 = \frac{1}{\gamma_1} \left(\mathbf{y}_1 - \beta_0 \mathbf{y}_0 - \sum_{j=2}^{s-1} \gamma_j \mathbf{x}_j \right).$$

As a result we have successfully exchanged $\mathbf{x}_1$ for $\mathbf{y}_1$ in our prior inclusion to see that $\mathcal{L} \subset \text{span}(\{\mathbf{y}_0, \mathbf{y}_1, \mathbf{x}_2, \dots, \mathbf{x}_{s-1}\})$ also holds.

Repeating this process $s - 2$ more times we find that $\mathcal{L} \subset \text{span}(\{\mathbf{y}_0, \mathbf{y}_1, \mathbf{y}_2, \dots, \mathbf{y}_{s-1}\})$ must hold. This generates a contradiction, however, because it implies that $\mathbf{y}_s \in B_2$ can be written as a linear combination of $\mathbf{y}_0, \dots, \mathbf{y}_{s-1}$, contradicting the fact that $\mathbf{y}_j$'s are linearly independent. Therefore, $|B_1| < |B_2|$ can't hold. $\square$

The following corollary of Lemma 2.3.4 guarantees that any two linearly independent sets that generate the same subspace have to have the same cardinality.

Corollary 2.3.5. *Let $B_1, B_2 \subset \mathbb{C}^n$ be finite sets that are both linearly independent. Furthermore, suppose that $\text{span}(B_1) = \text{span}(B_2)$. Then, $|B_1| = |B_2|$.*

Exercise 2.3.6. *Prove Corollary 2.3.5.*

We are now able to give a well defined definition of the dimension of a linear subspace. Let $\mathcal{L}$ be a linear subspace of $\mathbb{C}^n$. A **basis** of $\mathcal{L}$ is any linearly independent finite set B with $\mathcal{L} = \text{span}(B)$.³ Note that by Corollary 2.3.5 all bases of $\mathcal{L}$ must have the same cardinality. We call this cardinality the **dimension** of $\mathcal{L}$, and denote it by $\dim(\mathcal{L}) \in [n+1]$. If $\mathcal{L}$ is the subspace containing only the zero vector, we say that $\mathcal{L}$ is the **trivial subspace** and has dimension zero.

Example 2.3.6 (The Dimension of $\mathbb{C}^n$). *The n standard basis vectors $\{\mathbf{e}_j\}_{j \in [n]} \subset \mathbb{C}^n$ are indeed a basis of $\mathbb{C}^n$ because they are linearly independent and satisfy $\mathbb{C}^n = \text{span}(\{\mathbf{e}_j\}_{j \in [n]})$. As a result, we can see that the dimension of $\mathbb{C}^n$ is n .*

Exercise 2.3.7. *Let $\mathcal{L}$ be a linear subspace of $\mathbb{C}^n$. Use Lemma 2.3.4 to show that any linearly independent set of vectors $B \subset \mathcal{L}$ has cardinality $\leq n$.*

Exercise 2.3.8. *Let $\mathcal{L} \subset \mathbb{C}^n$ be a linear subspace. Prove that the dimension of $\mathcal{L}$ is at most n .*

Exercise 2.3.9. *Let $\mathcal{L} \subset \mathbb{C}^n$ be a linear subspace. Show that any linearly independent set of vectors $B \subset \mathcal{L}$ has cardinality $\leq$ the dimension of $\mathcal{L}$.*

Exercise 2.3.10. *Let $\mathcal{L}, \mathcal{K} \subset \mathbb{C}^n$ be two linear subspaces of $\mathbb{C}^n$ such that $\mathcal{L} \subset \mathcal{K}$. Prove that $\dim(\mathcal{L}) \leq \dim(\mathcal{K})$.*

The following lemma is crucial in several later arguments.

Lemma 2.3.7. *Let $\mathcal{L}, \mathcal{K} \subset \mathbb{C}^n$ be two linear subspaces of $\mathbb{C}^n$ with $\mathcal{L} \cap \mathcal{K} = \{\mathbf{0}\}$. Then $\mathcal{L} \cup \mathcal{K}$ contains $\dim(\mathcal{L}) + \dim(\mathcal{K})$ linearly independent vectors.*

Proof. Let $r = \dim(\mathcal{L})$ and $B = \{\mathbf{b}_j\}_{j \in [r]} \subset \mathcal{L}$ be a basis of $\mathcal{L}$. Similarly, let $s = \dim(\mathcal{K})$ and $A = \{\mathbf{a}_k\}_{k \in [s]} \subset \mathcal{K}$ be a basis of $\mathcal{K}$. We can see that $B \cup A$ must have cardinality $\dim(\mathcal{L}) + \dim(\mathcal{K})$ since $\mathcal{L} \cap \mathcal{K} = \{\mathbf{0}\}$, and neither B nor A can contain $\mathbf{0}$ (recall Exercise 2.3.4). Hence, we will be finished if we can show that $B \cup A$ is linearly independent.

Suppose for the sake of contradiction that $B \cup A$ is linearly dependent. Then, there exists a nonzero vector $\alpha \in \mathbb{C}^{r+s}$ such that

$$\begin{aligned} \sum_{j \in [r]} \alpha_j \mathbf{b}_j + \sum_{k \in [s]} \alpha_{k+r} \mathbf{a}_k = \mathbf{0} &\iff \mathcal{L} \ni \sum_{j \in [r]} \alpha_j \mathbf{b}_j = \sum_{k \in [s]} (-\alpha_{k+r}) \mathbf{a}_k \in \mathcal{K} \\ &\iff \sum_{j \in [r]} \alpha_j \mathbf{b}_j = \mathbf{0} \text{ and } \sum_{k \in [s]} (-\alpha_{k+r}) \mathbf{a}_k = \mathbf{0} \end{aligned}$$

³Note that by our definition of “linear subspace” it’s not immediately clear that every linear subspace of $\mathbb{C}^n$ has to have a basis. They do, and you can build a basis for any subspace of $\mathbb{C}^n$ in a finite number of steps using the Gram–Schmidt algorithm (see, e.g., Section 2.4).

since $\mathcal{L} \cap \mathcal{K} = \{\mathbf{0}\}$. Furthermore, at least one of $\sum_{j \in [r]} \alpha_j \mathbf{b}_j$ or $\sum_{k \in [s]} (-\alpha_{k+r}) \mathbf{a}_k$ is a nonzero sum since $\boldsymbol{\alpha} \in \mathbb{C}^{r+s}$ is nonzero. However, we then have a contradiction since both A and B are linearly independent. $\square$

Exercise 2.3.11. Let $\mathcal{L}, \mathcal{K} \subset \mathbb{C}^n$ be two linear subspaces of $\mathbb{C}^n$ with $\dim(\mathcal{L}) + \dim(\mathcal{K}) > n$. Prove that there exists a nonzero vector $\mathbf{x} \in \mathcal{L} \cap \mathcal{K}$.

Exercise 2.3.12. Let $\mathcal{L}, \mathcal{K} \subset \mathbb{C}^n$ be two linear subspaces of $\mathbb{C}^n$ with $\dim(\mathcal{L}) + \dim(\mathcal{K}) > n$. Prove that $\mathcal{L} \cap \mathcal{K}$ is a linear subspace of $\mathbb{C}^n$ with $\dim(\mathcal{L} \cap \mathcal{K}) \geq 1$.

Given a matrix $A \in \mathbb{C}^{m \times n}$, we define the **rank** of A to be the dimension of its column space $\mathcal{C}(A) \subset \mathbb{C}^m$ (which, as a reminder, is the span of the columns of A).

Exercise 2.3.13. Show that a rank r matrix $A \in \mathbb{C}^{m \times n}$ has exactly r linearly independent columns.

Exercise 2.3.14. Show that the rank of a matrix $A \in \mathbb{C}^{m \times n}$ is always $\leq \min\{m, n\}$.

We define a set of nonzero vectors $\{\mathbf{v}_j\}_{j \in [m]} \subset \mathbb{C}^n$ to be **mutually orthogonal** (or just **orthogonal**) if, for all $j \neq k$, $\langle \mathbf{v}_j, \mathbf{v}_k \rangle = 0$. We will also say that a set containing a single vector $\{\mathbf{v}\} \subset \mathbb{C}^n$ is **trivially orthogonal** since it contains nothing else for $\mathbf{v}$ to fail to be orthogonal with. The next lemma shows that orthogonal vectors are always linearly independent. Hence, they always form a basis of their span.

Lemma 2.3.8. An orthogonal set of nonzero vectors is always linearly independent.

Proof. Let $\{\mathbf{y}_j\}_{j \in [m]} \subset \mathbb{C}^n$ be orthogonal nonzero vectors. Suppose that there exist some $\alpha_0, \dots, \alpha_{m-1} \in \mathbb{C}$ such that

$$\sum_{j \in [m]} \alpha_j \mathbf{y}_j = \mathbf{0}.$$

Let $k \in [m]$. Since the inner product of the zero vector with any other vector is 0, we can see that

$$0 = \left\langle \mathbf{y}_k, \sum_{j \in [m]} \alpha_j \mathbf{y}_j \right\rangle = \sum_{j \in [m]} \alpha_j \langle \mathbf{y}_k, \mathbf{y}_j \rangle = \alpha_k \langle \mathbf{y}_k, \mathbf{y}_k \rangle = \alpha_k \|\mathbf{y}_k\|_2^2.$$

Recalling the properties of norms, we note that since $\mathbf{y}_k \neq \mathbf{0}$, $\|\mathbf{y}_k\|_2^2 > 0$. Hence, $\alpha_k = 0$ must hold for all $k \in [m]$. $\square$

A set of orthogonal vectors in $\mathbb{C}^n$ that all have norm 1 is called an **orthonormal set**. Note that given a set of orthogonal nonzero vectors, we can normalize each of them by replacing $\mathbf{y}_j$ with each $\frac{\mathbf{y}_j}{\|\mathbf{y}_j\|_2}$. This then guarantees that $\left\| \frac{\mathbf{y}_j}{\|\mathbf{y}_j\|_2} \right\|_2 = \frac{1}{\|\mathbf{y}_j\|_2} \|\mathbf{y}_j\|_2 = 1$. Thus, any orthogonal set of nonzero vectors can be turned into an orthonormal set. If a set of orthonormal vectors span a linear subspace $\mathcal{L} \subset \mathbb{C}^n$, we say that they are an **orthonormal basis for $\mathcal{L}$** .

Exercise 2.3.15. Show that the standard basis vectors $\{\mathbf{e}_j\}_{j \in [n]} \subset \mathbb{C}^n$ form an orthonormal basis of $\mathbb{C}^n$.

Orthonormal bases have several nice properties. For example, if we know that a vector $\mathbf{x} \in \mathbb{C}^n$ is in the span of an orthonormal basis $\{\mathbf{v}_j\}_{j \in [m]} \subset \mathbb{C}^n$, then we can find an expansion of $\mathbf{x}$ in terms of $\{\mathbf{v}_j\}_{j \in [m]}$ by noting that

$$\mathbf{x} = \sum_{j \in [m]} \alpha_j \mathbf{v}_j \iff \langle \mathbf{v}_k, \mathbf{x} \rangle = \sum_{j \in [m]} \alpha_j \langle \mathbf{v}_k, \mathbf{v}_j \rangle = \alpha_k \|\mathbf{v}_k\|_2^2 = \alpha_k \quad \forall k \in [m]. \quad (2.11)$$

Thus, we can easily recover the coefficients $\alpha_j \in \mathbb{C}$ of the linear combination making up $\mathbf{x}$ by taking the inner product of $\mathbf{x}$ with the orthonormal basis vectors. In addition, these coefficients will also satisfy the famous Pythagorean theorem.

Theorem 2.3.9 (The Pythagorean Theorem). Suppose $\{\mathbf{v}_j\}_{j \in [m]} = \{\mathbf{v}_0, \mathbf{v}_1, \dots, \mathbf{v}_{m-1}\} \subset \mathbb{C}^n$ is an orthonormal set of vectors. Then

$$\left\| \sum_{j \in [m]} \alpha_j \mathbf{v}_j \right\|_2^2 = \sum_{j \in [m]} |\alpha_j|^2$$

for all $\alpha_0, \dots, \alpha_{m-1} \in \mathbb{C}$. Equivalently, for any $\mathbf{x} \in \text{span}(\{\mathbf{v}_j\}_{j \in [m]})$,

$$\|\mathbf{x}\|_2^2 = \sum_{j \in [m]} |\langle \mathbf{x}, \mathbf{v}_j \rangle|^2$$

Proof. Note that the second equation follows immediately from the first since we have $\mathbf{x} = \sum_{j \in [m]} \langle \mathbf{v}_j, \mathbf{x} \rangle \mathbf{v}_j$. To show the first equation let $\alpha_0, \dots, \alpha_{m-1} \in \mathbb{C}$. Then,

$$\begin{aligned} \left\| \sum_{j \in [m]} \alpha_j \mathbf{v}_j \right\|_2^2 &= \left\langle \sum_{j \in [m]} \alpha_j \mathbf{v}_j, \sum_{k \in [m]} \alpha_k \mathbf{v}_k \right\rangle = \sum_{j \in [m]} \sum_{k \in [m]} \langle \alpha_j \mathbf{v}_j, \alpha_k \mathbf{v}_k \rangle \\ &= \sum_{j \in [m]} \sum_{k \in [m]} \overline{\alpha_j} \alpha_k \langle \mathbf{v}_j, \mathbf{v}_k \rangle = \sum_{j \in [m]} \overline{\alpha_j} \alpha_j \langle \mathbf{v}_j, \mathbf{v}_j \rangle = \sum_{j \in [m]} |\alpha_j|^2. \end{aligned}$$

□

Having hopefully reminded you why orthonormal bases are so great, we will now discuss how to generate one.

2.4 Orthonormal Bases and the Gram–Schmidt Algorithm

Algorithm 1 is an implementation of the Gram–Schmidt Algorithm which, when given a finite set $S \subset \mathbb{C}^n$ as input, outputs an orthonormal basis of $\text{span}(S)$. Before we analyze this algorithm to see that it works as intended we highly recommend that the reader take a close look at it. Here are some recommended exercises to help you pay close attention to how it works.

Exercise 2.4.1. Run Algorithm 1 on the set

$$S = \left\{ \begin{pmatrix} 1 \\ \mathbf{i} \end{pmatrix}, \begin{pmatrix} 2 + \mathbf{i} \\ 1 \end{pmatrix} \right\} \subset \mathbb{C}^2$$

by hand. Verify that the basis $B \subset \mathbb{C}^2$ it produces for $\text{span}(S)$ is indeed an orthonormal basis.

Exercise 2.4.2. Run Algorithm 1 on the set

$$S = \left\{ \begin{pmatrix} 1 \\ -1 \\ 1 \end{pmatrix}, \begin{pmatrix} 1 \\ 1 \\ 0 \end{pmatrix}, \begin{pmatrix} 4 \\ 0 \\ 2 \end{pmatrix} \right\} \subset \mathbb{C}^3$$

by hand. Verify that the basis $B \subset \mathbb{C}^3$ it produces for $\text{span}(S)$ is indeed an orthonormal basis.

When you are finished inspecting Algorithm 1 come back here and we prove a lemma which takes a step toward showing that the set B Algorithm 1 outputs is *always* an orthonormal basis for the span of the input set S .

Lemma 2.4.1. The set $B \subset \mathbb{C}^n$ output by Algorithm 1 is always orthonormal.

Proof. First, we observe from Lines 3 and 8 of Algorithm 1 that each $\mathbf{b}_j \in B$ will have norm 1. Thus, it only remains to show that B is orthogonal. To show orthogonality it suffices to show that the set $\{\mathbf{b}_\ell\}_{\ell=0}^j = \{\mathbf{b}_0, \mathbf{b}_1, \dots, \mathbf{b}_j\} \subset B$ is orthogonal for all $j \in [|B|]$. We will proceed by induction on j .

To begin, we note that $\{\mathbf{b}_\ell\}_{\ell=0}^0 = \{\mathbf{b}_0\}$ when $j = 0$ is trivially orthogonal as a singleton set. Now, as our induction hypothesis, assume that $\{\mathbf{b}_\ell\}_{\ell=0}^j$ is orthogonal. To show that $\{\mathbf{b}_\ell\}_{\ell=0}^{j+1}$ must also then be orthogonal it suffices to show that $\langle \mathbf{b}_k, \mathbf{b}_{j+1} \rangle$ will be 0 for all integers $k \in [0, j]$. Referring to Lines 6 and 8 of Algorithm 1, and noting that the vector

Algorithm 1 THE GRAM–SCHMIDT ALGORITHM FOR FINITE SETS

```

1: Input: A finite set  $S \subset \mathbb{C}^n$  with at least one nonzero element.
2: Output: An orthonormal set  $B \subset \mathbb{C}^n$  with  $\text{span}(B) = \text{span}(S)$ .
   # Initialize  $S$  and  $B$ .
3: Pick a nonzero  $\mathbf{x}_0 \in S$ , and set  $\mathbf{b}_0 := \mathbf{x}_0 / \|\mathbf{x}_0\|_2$  and  $B = \{\mathbf{b}_0\}$ .
4: Set  $S = S \setminus \{\mathbf{0}, \mathbf{x}_0\}$ , and initialize  $j = 1$ .
5: while  $S \neq \{\}$  do
6:   Pick  $\mathbf{x}_j \in S$ , and set  $\mathbf{y}_j = \mathbf{x}_j - \sum_{\ell=0}^{j-1} \langle \mathbf{b}_\ell, \mathbf{x}_j \rangle \mathbf{b}_\ell$ .
   # If  $\mathbf{y}_j = \mathbf{0}$  then  $\mathbf{x}_j \in \text{span}(B)$  already, so we'll immediately remove this  $\mathbf{x}_j$  from  $S$ 
   # in Line 11 and pick a new one. If  $\mathbf{y}_j \neq \mathbf{0}$  then  $\mathbf{x}_j \notin \text{span}(B)$ , so we will add a new
   # element to  $B$  so that  $\mathbf{x}_j$  will then belong to its new span.
7:   if  $\mathbf{y}_j \neq \mathbf{0}$  then
8:     Set  $\mathbf{b}_j := \mathbf{y}_j / \|\mathbf{y}_j\|_2$  and let  $B = B \cup \{\mathbf{b}_j\}$ .
9:     Set  $j = j+1$ .
10:  end if
11:  Set  $S = S \setminus \{\mathbf{x}_j\}$ .
12: end while
13: Return  $B$ .

```

permanently selected as $\mathbf{x}_j$ in Line 6 has its $\mathbf{y}_j \neq \mathbf{0}$ for all $j \in [|B|]$, we can see that indeed

$$\begin{aligned}
\langle \mathbf{b}_k, \mathbf{b}_{j+1} \rangle &= \left\langle \mathbf{b}_k, \frac{1}{\|\mathbf{y}_{j+1}\|_2} \left(\mathbf{x}_{j+1} - \sum_{\ell=0}^j \langle \mathbf{b}_\ell, \mathbf{x}_{j+1} \rangle \mathbf{b}_\ell \right) \right\rangle \\
&= \frac{1}{\|\mathbf{y}_{j+1}\|_2} \left(\langle \mathbf{b}_k, \mathbf{x}_{j+1} \rangle - \sum_{\ell=0}^j \langle \mathbf{b}_\ell, \mathbf{x}_{j+1} \rangle \langle \mathbf{b}_k, \mathbf{b}_\ell \rangle \right) \\
&= \frac{1}{\|\mathbf{y}_{j+1}\|_2} (\langle \mathbf{b}_k, \mathbf{x}_{j+1} \rangle - \langle \mathbf{b}_k, \mathbf{x}_{j+1} \rangle) = 0
\end{aligned}$$

for all $k \in [0, j]$, where we have used the inductive hypothesis that $\{\mathbf{b}_\ell\}_{\ell=0}^j$ is orthogonal in the last line. As a result we can see that $\{\mathbf{b}_\ell\}_{\ell=0}^{j+1}$ will also be orthogonal whenever $\{\mathbf{b}_\ell\}_{\ell=0}^j$ is for all $j = 0, 1, \dots, |B| - 2$, finishing our induction argument. $\square$

Lemma 2.4.1 guarantees that the output, $B \subset \mathbb{C}^n$, of Algorithm 1 is always orthonormal, but in order for it to be a basis of $\text{span}(S)$ we also need that $\text{span}(B) = \text{span}(S)$. This is established in our next lemma.

Lemma 2.4.2. *The set $B \subset \mathbb{C}^n$ output by Algorithm 1 always satisfies $\text{span}(B) = \text{span}(S)$.*

Proof. It suffices to show that $\text{span}(\{\mathbf{x}_\ell\}_{\ell=0}^j) = \text{span}(\{\mathbf{b}_\ell\}_{\ell=0}^j)$ for all $j \in [|B|]$ (think

about why!⁴). We will show this by induction on j . To begin, we note that when $j = 0$ we have $\text{span}(\{\mathbf{x}_0\}) = \text{span}(\{\mathbf{b}_0\})$ since $\mathbf{b}_0$ is a nonzero scalar multiple of $\mathbf{x}_0$ (see Line 3). Now, suppose for the sake of induction that $\text{span}(\{\mathbf{x}_\ell\}_{\ell=0}^j) = \text{span}(\{\mathbf{b}_\ell\}_{\ell=0}^j)$. We will prove that then $\text{span}(\{\mathbf{x}_\ell\}_{\ell=0}^{j+1}) = \text{span}(\{\mathbf{b}_\ell\}_{\ell=0}^{j+1})$ must also hold in the usual two steps.

$\text{span}(\{\mathbf{x}_\ell\}_{\ell=0}^{j+1}) \subset \text{span}(\{\mathbf{b}_\ell\}_{\ell=0}^{j+1})$: Let $\mathbf{x} \in \text{span}(\{\mathbf{x}_\ell\}_{\ell=0}^{j+1})$. Then, we can write

$$\mathbf{x} = \alpha_{j+1}\mathbf{x}_{j+1} + \mathbf{y}$$

where $\alpha_{j+1} \in \mathbb{C}$ and $\mathbf{y} \in \text{span}(\{\mathbf{x}_\ell\}_{\ell=0}^j) = \text{span}(\{\mathbf{b}_\ell\}_{\ell=0}^j)$. By Lines 6 and 8 of Algorithm 1 we also have that

$$\mathbf{x}_{j+1} = \|\mathbf{y}_{j+1}\|_2 \mathbf{b}_{j+1} + \sum_{\ell=0}^j \langle \mathbf{b}_\ell, \mathbf{x}_{j+1} \rangle \mathbf{b}_\ell$$

so $\mathbf{x}_{j+1} \in \text{span}(\{\mathbf{b}_\ell\}_{\ell=0}^{j+1})$. Therefore, $\mathbf{x} \in \text{span}(\{\mathbf{b}_\ell\}_{\ell=0}^{j+1})$.

$\text{span}(\{\mathbf{b}_\ell\}_{\ell=0}^{j+1}) \subset \text{span}(\{\mathbf{x}_\ell\}_{\ell=0}^{j+1})$: Let $\mathbf{z} \in \text{span}(\{\mathbf{b}_\ell\}_{\ell=0}^{j+1})$. Then we can write

$$\mathbf{z} = \beta_{j+1}\mathbf{b}_{j+1} + \mathbf{y}$$

where $\beta_{j+1} \in \mathbb{C}$ and $\mathbf{y} \in \text{span}(\{\mathbf{b}_\ell\}_{\ell=0}^j) = \text{span}(\{\mathbf{x}_\ell\}_{\ell=0}^j)$. Again, by Lines 6 and 8 of Algorithm 1 we also have that

$$\mathbf{b}_{j+1} = \frac{1}{\|\mathbf{y}_{j+1}\|_2} \left(\mathbf{x}_{j+1} - \sum_{\ell=0}^j \langle \mathbf{b}_\ell, \mathbf{x}_{j+1} \rangle \mathbf{b}_\ell \right),$$

where the sum $\sum_{\ell=0}^j \langle \mathbf{b}_\ell, \mathbf{x}_{j+1} \rangle \mathbf{b}_\ell$ above is in $\text{span}(\{\mathbf{b}_\ell\}_{\ell=0}^j) = \text{span}(\{\mathbf{x}_\ell\}_{\ell=0}^j)$. Thus, $\mathbf{b}_{j+1} \in \text{span}(\{\mathbf{x}_\ell\}_{\ell=0}^{j+1})$ which in turn implies that $\mathbf{z} \in \text{span}(\{\mathbf{x}_\ell\}_{\ell=0}^{j+1})$.

Having now shown that both $\text{span}(\{\mathbf{x}_\ell\}_{\ell=0}^{j+1}) \subset \text{span}(\{\mathbf{b}_\ell\}_{\ell=0}^{j+1})$ and $\text{span}(\{\mathbf{b}_\ell\}_{\ell=0}^{j+1}) \subset \text{span}(\{\mathbf{x}_\ell\}_{\ell=0}^{j+1})$ hold, we conclude that indeed $\text{span}(\{\mathbf{x}_\ell\}_{\ell=0}^{j+1}) = \text{span}(\{\mathbf{b}_\ell\}_{\ell=0}^{j+1})$. $\square$

Combining Lemmas 2.4.1 and 2.4.2 we obtain the following theorem guaranteeing that Algorithm 1 always produces an orthonormal basis of the span of its input set, as intended.

⁴Note that multiple candidates may be selected to be $\mathbf{x}_j$ in Line 6 before j is actually incremented in Line 9. All such temporarily-selected $\mathbf{x}_j$ candidates simply overwrite one another. However, only the final vector permanently selected to be $\mathbf{x}_j$ in Line 6 has its $\mathbf{y}_j \neq \mathbf{0}$. All other temporarily-selected $\mathbf{x}_j$ candidates must have $\mathbf{y}_j = \mathbf{0}$ in Line 7 in order to be overwritten, which means that they are already in the span of $\{\mathbf{b}_\ell\}_{\ell=0}^{j-1}$. Thus, only the span of the final permanently selected $\mathbf{x}_j$'s really matters.

Algorithm 2 THE GRAM–SCHMIDT ALGORITHM FOR SUBSPACES OF $\mathbb{C}^n$

```

1: Input: A nontrivial subspace  $\mathcal{L} \subset \mathbb{C}^n$  (i.e., an  $\mathcal{L} \neq \{\mathbf{0}\}$ ).
2: Output: An orthonormal set  $B \subset \mathbb{C}^n$  with  $\text{span}(B) = \mathcal{L}$ .
3: Pick  $\mathbf{x} \in \mathcal{L} \setminus \{\mathbf{0}\}$  and initialize  $B = \{\mathbf{x}/\|\mathbf{x}\|_2\}$ .
4: while  $\mathcal{L} \not\subset \text{span}(B)$  do
5:   Pick  $\mathbf{x} \in \mathcal{L} \setminus \text{span}(B)$ .
6:   Let  $\mathbf{y} = \mathbf{x} - \sum_{\mathbf{b} \in B} \langle \mathbf{b}, \mathbf{x} \rangle \mathbf{b}$ .
7:   Set  $B = B \cup \{\mathbf{y}/\|\mathbf{y}\|_2\}$ .
8: end while
9: Return  $B$ .

```

Theorem 2.4.3. *Algorithm 1 always returns an orthonormal basis B of $\text{span}(S) \subset \mathbb{C}^n$.*

The Gram–Schmidt algorithm is also useful for a lot of other theoretical reasons as well, which I would like to briefly mention here (please indulge me!). For example, based on our definition of what a linear subspace of $\mathbb{C}^n$ is, it’s not immediately clear that every such subspace has to have a basis. This can be established by, e.g., analyzing Algorithm 2 which is a variant of Algorithm 1 (except for subspaces). Please go and look it over.

Looking at Algorithm 2 we can see that a slightly modified version of Lemma 2.4.1 will again guarantee that B will remain orthonormal at all times. The main open question here is therefore whether the “while loop” in Line 4 of Algorithm 2 will ever terminate (subspaces are, after all, infinite sets ... there are many worst-case $\mathbf{x}$ values to pick from in each iteration!). We need not fear, however. The while loop must terminate after at most n iterations no matter what by the Exchange Lemma (Lemma 2.3.4) exactly because B will always be linearly independent (see, e.g., Exercise 2.3.7). More precisely, it will terminate after $\dim(\mathcal{L}) \leq n$ iterations (see, e.g., Exercise 2.3.9). Failing to do so would generate a contradiction. Formalizing this argument proves the following theorem.

Theorem 2.4.4. *Every nontrivial linear subspace $\mathcal{L} \subset \mathbb{C}^n$ has an orthonormal basis.*

As a final thought regarding Gram–Schmidt algorithm variants, we note that they can also be used to expand an orthonormal basis of a low-dimensional subspace of $\mathbb{C}^n$ into a larger orthonormal basis of all of $\mathbb{C}^n$. This fact comes in handy on many occasions. More precisely, suppose that we have an orthonormal basis B of a subspace $\mathcal{L} \subset \mathbb{C}^n$ with $|B| = \dim(\mathcal{L}) < n$ in our possession. Then, we can use Algorithm 3 to extend it to a larger basis $\tilde{B}$ of $\mathbb{C}^n$ with $B \subset \tilde{B}$. Please go take a look at Algorithm 3, paying special attention to its similarities and differences with Algorithm 2.

Looking at Algorithm 3 we can see that it is effectively a continuation of Algorithm 2. That is, Algorithm 3 effectively picks up where Algorithm 2 leaves off and then continues in the exact same way after substituting $\mathcal{L}$ with $\mathbb{C}^n$ everywhere in its “while loop”. As a consequence of this substitution, we can use essentially the same reasoning as above to

Algorithm 3 GRAM–SCHMIDT FOR EXTENDING AN ORTHONORMAL BASIS

```

1: Input: An orthonormal basis  $B$  of a subspace  $\mathcal{L} \subset \mathbb{C}^n$ .
2: Output: An orthonormal basis  $\tilde{B}$  of  $\mathbb{C}^n$  with  $B \subset \tilde{B}$ .
3: Initialize  $\tilde{B} = B$ .
4: while  $\mathbb{C}^n \not\subset \text{span}(\tilde{B})$  do
5:   Pick  $\mathbf{x} \in \mathbb{C}^n \setminus \text{span}(\tilde{B})$ .
6:   Let  $\mathbf{y} = \mathbf{x} - \sum_{\mathbf{b} \in \tilde{B}} \langle \mathbf{b}, \mathbf{x} \rangle \mathbf{b}$ .
7:   Set  $\tilde{B} = \tilde{B} \cup \{\mathbf{y}/\|\mathbf{y}\|_2\}$ .
8: end while
9: Return  $\tilde{B}$ .

```

see that Algorithm 3 will indeed output an orthonormal basis $\tilde{B}$ of $\mathbb{C}^n$. Furthermore, the fact that $B \subset \tilde{B}$ is entirely a result of how $\tilde{B}$ is initialized. Formalizing this line of thought proves the following theorem.

Theorem 2.4.5. *Let $B \subset \mathbb{C}^n$ be an orthonormal basis of a linear subspace $\mathcal{L} \subset \mathbb{C}^n$. Then there exists an orthonormal basis $\tilde{B}$ of $\mathbb{C}^n$ such that $B \subset \tilde{B}$.*

Exercise 2.4.3. *Implement a version of Algorithm 3 in the language of your choice⁵ and use it to complete the orthonormal set*

$$S = \left\{ \frac{1}{2} \begin{pmatrix} 1 \\ -1 \\ 1 \\ \mathbf{i} \end{pmatrix}, \frac{1}{2} \begin{pmatrix} 1 \\ 1 \\ 1 \\ -\mathbf{i} \end{pmatrix} \right\} \subset \mathbb{C}^4$$

to an orthonormal basis of all of $\mathbb{C}^4$. Verify that your resulting orthonormal basis set is indeed orthonormal.

Exercise 2.4.4. *Prove that every set of n orthonormal vectors in $\mathbb{C}^n$ is an orthonormal basis of $\mathbb{C}^n$.*

Exercise 2.4.5. *Let $\mathcal{L} \subset \mathbb{C}^n$ be a linear subspace of dimension $r \leq n$. Prove that every set of r orthonormal vectors in $\mathcal{L}$ is an orthonormal basis of $\mathcal{L}$.*

We will now explore yet another important consequence of the Gram–Schmidt algorithm – the existence of a QR factorization for any matrix $A \in \mathbb{C}^{m \times n}$.

⁵The language of your choice can also be “by hand” (recommended). If you do decide to write computer code, though, feel free to ask A.I. for help.

2.4.1 The QR Decomposition of a Matrix

Let's consider what happens when we apply Algorithm 1 to the columns of a matrix $A \in \mathbb{C}^{m \times n}$ so that its input is $S = \{A_{:,j}\}_{j \in [n]} \subset \mathbb{C}^m$. Even more specifically, suppose that we run Algorithm 1 with $\mathbf{x}_0 = A_{:,0}$, $\mathbf{x}_1 = A_{:,1}$ (or, more generally, = the first column after $A_{:,0}$ that isn't a multiple of $A_{:,0}$), $\mathbf{x}_2 = A_{:,2}$ (or, more generally, = the first column after $\mathbf{x}_1$ that isn't in the span of $\mathbf{x}_0$ and $\mathbf{x}_1$), etc.. First, we know that Algorithm 1 will output an orthonormal basis $B \subset \mathbb{C}^m$ of the column space, $\mathcal{C}(A)$, of A when it finishes. Second, by the definition of rank we also know that $|B| = \text{rank}(A)$. Denote the rank of A by r , and the elements of B by $\{\mathbf{b}_j\}_{j \in [r]}$.

By our analysis of Algorithm 1 we can further see that $A_{:,0} \in \text{span}(\{\mathbf{b}_0\})$, $A_{:,1} \in \text{span}(\{\mathbf{b}_0, \mathbf{b}_1\})$, etc.. More generally, $A_{:,j} \in \text{span}(\{\mathbf{b}_\ell\}_{\ell=0}^{\min\{j, r-1\}})$ for all $j \in [n]$. As a consequence, there exist complex numbers $R_{i,j} \in \mathbb{C}$, with $i, j \in [n]$ and $i \leq j$, such that

$$\begin{aligned} A_{:,0} &= R_{0,0} \mathbf{b}_0 \\ A_{:,1} &= R_{0,1} \mathbf{b}_0 + R_{1,1} \mathbf{b}_1 \\ &\vdots \\ A_{:,j} &= \sum_{\ell=0}^{\min\{j, r-1\}} R_{\ell,j} \mathbf{b}_\ell \text{ for all } j \in [n]. \end{aligned}$$

Now, define $Q \in \mathbb{C}^{m \times r}$ to be the matrix with the elements of B as its columns, and $R \in \mathbb{C}^{r \times n}$ to be the upper triangular matrix whose nonzero entries are defined above so that

$$Q = \begin{pmatrix} | & & | \\ \mathbf{b}_0 & \cdots & \mathbf{b}_{r-1} \\ | & & | \end{pmatrix} \text{ and } R = \begin{pmatrix} R_{0,0} & R_{0,1} & \cdots & R_{0,r-1} & \cdots & R_{0,n-1} \\ 0 & R_{1,1} & \cdots & R_{1,r-1} & \cdots & R_{1,n-1} \\ 0 & 0 & \ddots & \vdots & & \vdots \\ \vdots & \vdots & 0 & R_{r-1,r-1} & \cdots & R_{r-1,n-1} \end{pmatrix}.$$

Doing so we can see that

$$\begin{pmatrix} | & | & \cdots & | \\ A_{0,:} & A_{1,:} & \cdots & A_{:,n-1} \\ | & | & & | \end{pmatrix} = \begin{pmatrix} | & | & \cdots & | \\ \mathbf{b}_0 & \mathbf{b}_1 & \cdots & \mathbf{b}_{r-1} \\ | & | & & | \end{pmatrix} \begin{pmatrix} R_{0,0} & R_{0,1} & \cdots & R_{0,r-1} & \cdots \\ 0 & R_{1,1} & \cdots & R_{1,r-1} & \cdots \\ 0 & 0 & \ddots & \vdots & \\ \vdots & \vdots & 0 & R_{r-1,r-1} & \cdots \end{pmatrix}.$$

That is, $A = QR$. This is called a **QR decomposition** of A , and it is very useful computationally since both Q and R have special properties. Namely, Q has orthonormal columns, and R is upper triangular. By formalizing the discussion above one may prove the following theorem.

Theorem 2.4.6 (Every Matrix Has a QR Decomposition). *Let $A \in \mathbb{C}^{m \times n}$ be rank r . Then, there exists a matrix $Q \in \mathbb{C}^{m \times r}$ with orthonormal columns, and an upper triangular matrix $R \in \mathbb{C}^{r \times n}$, so that $A = QR$.*

Example 2.4.7. *The following is an example of a QR decomposition for a rank 2 matrix $A \in \mathbb{C}^{4 \times 4}$.*

$$A = \begin{pmatrix} 1 & 2 & 3 & 1 \\ 1 & 2 & 3 & -1 \\ 1 & 2 & 3 & 1 \\ 1 & 2 & 3 & -1 \end{pmatrix} = \begin{pmatrix} \frac{1}{2} & \frac{1}{2} \\ \frac{1}{2} & -\frac{1}{2} \\ \frac{1}{2} & \frac{1}{2} \\ \frac{1}{2} & -\frac{1}{2} \end{pmatrix} \begin{pmatrix} 2 & 4 & 6 & 0 \\ 0 & 0 & 0 & 2 \end{pmatrix} = QR.$$

Note that the matrix Q guaranteed by Theorem 2.4.6 is also clearly rank r since Q has orthonormal columns. In fact, with just a bit more work one can further see that R will always be rank r as well. We will save this final rank analysis for Section 2.7, however. For now, let us turn our attention to some implications of the QR decomposition with regard to low rank matrix compression.

Exercise 2.4.6. *Compute a QR decomposition of the matrix*

$$A = \begin{pmatrix} 1 & 2 \\ i & 1 \end{pmatrix}.$$

Verify that Q has orthonormal columns.

Exercise 2.4.7. *Compute a QR decomposition of the matrix*

$$A = \begin{pmatrix} 1 & 5 & -1 \\ -1 & 1 & -1 \\ 2 & 1 & 1 \end{pmatrix}.$$

Verify that Q has orthonormal columns.

Some Comments on Computing a QR Decomposition of a Matrix: I hope that this section has begun to convince you that the QR decomposition might be interesting. In fact, we will see going forward that the QR decomposition is also incredibly useful – useful enough that I am pretty certain that anyone reading this sentence will likely compute one at some point (probably using a preexisting software package like – these days – MATLAB, SciPy, LAPACK, or ... there are many!). When you do compute that QR decomposition it's important to point out that it won't be by (shouldn't be by!) running Algorithm 1 on the columns of the matrix. Theoretically Algorithm 1 is fantastic, but in practice a digital computer will likely turn a straightforward coding of Algorithm 1 into the inaccurate numerical equivalent of a reeking garbage scow (i.e., it'll be numerically unstable). In practice QR decompositions are instead computed using Householder reflections which, if interested, you can read about in standard numerical linear algebra texts such as, e.g., [34, 11].

2.5 Near-Optimal Compression of Low Rank Matrices

In this section we briefly consider the minimum number of complex values we need to store in order to fully represent a rank r matrix $A \in \mathbb{C}^{m \times n}$. Clearly, we can always do it by storing all mn entries in A , but can we do better? The answer is definitely “yes” if the matrix is low rank. To see why, consider a QR decomposition of $A \in \mathbb{C}^{m \times n}$, $A = QR$. Recalling that $Q^{m \times r}$ and $R \in \mathbb{C}^{r \times n}$, we immediately see that in fact we can completely represent A by instead storing the at most $mr + nr = r(m + n)$ entries of Q and R . And if, for example, $n = m$ and $r < n/2$, storing the at most $mr + nr = 2nr < n^2$ entries of Q and R will require less memory than directly storing the $mn = n^2$ entries of A .

In fact, however, we can do even better than this by taking full advantage of the structure that a QR decomposition guarantees us. Since, e.g., R is upper triangular we know that it will always have $\frac{(r-1)r}{2}$ zero entries below its main diagonal in predictable positions. Thus, there is no need to actually store those 0-valued entries of R . As a result, we can see that it really suffices to only store

$$mr + rn - \frac{(r-1)r}{2} = r \left(m + n - \frac{r-1}{2} \right) \quad (2.12)$$

complex numbers in order to fully represent both Q and R , and therefore A . *Note that this reduction in entries can have noticeable space-saving effects, especially when we need to store a large number of very large matrices. Further note that this is exactly the case one is in when, e.g., one wants to store the many large weight matrices needed to fully describe a trained deep neural network (recall Section 1.2.3)!*

Exercise 2.5.1. *Show that an upper triangular matrix $R \in \mathbb{C}^{r \times n}$ with $r \leq n$ will always have $\frac{(r-1)r}{2}$ zero entries below its main diagonal.*

The number of complex entries (2.12) one needs to store in order to represent a QR decomposition as described above is not quite optimal. To see why, we note that the dimension of the manifold of rank r matrices in $\mathbb{C}^{m \times n}$ is $(m + n - r)r$ (see, e.g., [16, Chapter 1]), so one should be able to represent any rank r matrix $A \in \mathbb{C}^{m \times n}$ by storing just $(m + n - r)r$ complex values. This means that storing a QR decomposition of A requires storing $\frac{r^2+r}{2}$ additional complex values beyond the theoretical minimum. As we will see, Gaussian elimination can help us reduce this number of additional values closer to 0. Using this as motivation we will now very briefly summarize Gaussian elimination while simultaneously introducing and reviewing a lot of other very useful notation.

2.5.1 A Very Brief Review of Gaussian Elimination, and Some Useful Notation

First let’s recall some notation. As mentioned above, we view vectors $\mathbf{u}, \mathbf{v} \in \mathbb{C}^n$ as $n \times 1$ matrices. As a result, the inner product $\langle \mathbf{u}, \mathbf{v} \rangle \in \mathbb{C}$ can be viewed as the matrix product of

a $1 \times n$ matrix with an $n \times 1$ matrix,

$$\mathbf{u}^* \mathbf{v} = (\overline{u_0}, \overline{u_1}, \dots, \overline{u_{n-1}}) \begin{pmatrix} v_0 \\ v_1 \\ \vdots \\ v_{n-1} \end{pmatrix} = \sum_{j \in [n]} \overline{u_j} v_j = \langle \mathbf{u}, \mathbf{v} \rangle,$$

the result of which is a 1×1 matrix (i.e., the scalar $\langle \mathbf{u}, \mathbf{v} \rangle$). We can similarly define the “outer product of two vectors” in $\mathbb{C}^n$ as the product of an $n \times 1$ matrix with a $1 \times n$ matrix. That is, given $\mathbf{u}, \mathbf{v} \in \mathbb{C}^n$, their **outer product** is

$$\mathbf{v} \mathbf{u}^* = \begin{pmatrix} v_0 \\ v_1 \\ \vdots \\ v_{n-1} \end{pmatrix} (\overline{u_0}, \overline{u_1}, \dots, \overline{u_{n-1}}) = \begin{pmatrix} v_0 \overline{u_0} & v_0 \overline{u_1} & \dots & v_0 \overline{u_{n-1}} \\ v_1 \overline{u_0} & v_1 \overline{u_1} & \dots & v_1 \overline{u_{n-1}} \\ \vdots & \vdots & \ddots & \vdots \\ v_{n-1} \overline{u_0} & v_{n-1} \overline{u_1} & \dots & v_{n-1} \overline{u_{n-1}} \end{pmatrix} \in \mathbb{C}^{n \times n}.$$

Note that this is an $n \times n$ matrix whose $(j, k)^{\text{th}}$ entry is $v_j \overline{u_k} \in \mathbb{C}$.

Next, the **standard basis of** $\mathbb{C}^{m \times n}$ consists of the mn matrices in $\mathbb{C}^{m \times n}$, denoted by $E^{(j,k)} \in \mathbb{C}^{m \times n}$, whose entries are given by

$$\left(E^{(j,k)} \right)_{\ell, h} = \begin{cases} 1 & \text{if } \ell = j \text{ and } h = k \\ 0 & \text{else} \end{cases} \quad \text{for all } j, \ell \in [m] \text{ and } k, h \in [n].$$

Note that we also have $E^{(j,k)} = \mathbf{e}_j \mathbf{e}_k^*$. We call these matrices the standard basis for $\mathbb{C}^{m \times n}$ because any matrix $A \in \mathbb{C}^{m \times n}$ can be expressed as the linear combination

$$A = \sum_{j \in [m]} \sum_{k \in [n]} A_{j,k} E^{(j,k)} = \sum_{j \in [m]} \sum_{k \in [n]} A_{j,k} \mathbf{e}_j \mathbf{e}_k^*.$$

Continuing, given a vector $\mathbf{v} \in \mathbb{C}^n$, we denote the diagonal matrix in $\mathbb{C}^{n \times n}$ with $\mathbf{v}$ on its diagonal by $\text{diag}(\mathbf{v}) \in \mathbb{C}^{n \times n}$. Equivalently, $\text{diag}(\mathbf{v})$ is the $n \times n$ matrix with entries given by

$$(\text{diag}(\mathbf{v}))_{j,k} = \begin{cases} v_j & \text{if } j = k \\ 0 & \text{else} \end{cases}.$$

Finally, we will denote the vector of all ones in $\mathbb{C}^n$ by $\mathbf{1} \in \mathbb{C}^n$. The following exercises will help you get more familiar with all of this notation.

Exercise 2.5.2. Let $\mathbf{v} \in \mathbb{C}^n$. Show that $\text{diag}(\mathbf{v}) = \sum_{j \in [n]} v_j \mathbf{e}_j \mathbf{e}_j^* = \sum_{j \in [n]} v_j E^{(j,j)}$.

Exercise 2.5.3. Let $\mathbf{v}, \mathbf{u} \in \mathbb{C}^n$. Show that $\text{diag}(\mathbf{v}) \mathbf{u} = \text{diag}(\mathbf{u}) \mathbf{v} \in \mathbb{C}^n$. As a consequence, show that $\text{diag}(\mathbf{1}) \mathbf{v} = \text{diag}(\mathbf{v}) \mathbf{1} = \mathbf{v}$ holds for all $\mathbf{v} \in \mathbb{C}^n$.

We can now see that the $n \times n$ **identity matrix**, denoted by $I_n \in \mathbb{C}^{n \times n}$, can be expressed in several equivalent forms. First, we know that its entries are $(I_n)_{j,k} = \begin{cases} 1 & \text{if } j = k \\ 0 & \text{else} \end{cases}$, for all $j, k \in [n]$. As a consequence we can see that

$$\begin{aligned} I_n &= \begin{pmatrix} | & | & & | \\ \mathbf{e}_0 & \mathbf{e}_1 & \cdots & \mathbf{e}_{n-1} \\ | & | & & | \end{pmatrix} = \text{diag}(\mathbf{1}) \\ &= \sum_{j \in [n]} E^{(j,j)} = \sum_{j \in [n]} \mathbf{e}_j \mathbf{e}_j^*. \end{aligned}$$

Additionally, we recall that the **inverse of a matrix** $A \in \mathbb{C}^{n \times n}$, if it exists, is the matrix $A^{-1} \in \mathbb{C}^{n \times n}$ satisfying $AA^{-1} = A^{-1}A = I_n$.

Having equipped ourselves with this new notation, we may now more easily and quickly review Gaussian elimination. In short, **Gaussian elimination** is the process of multiplying three types of invertible elementary matrices against a given matrix $A \in \mathbb{C}^{m \times n}$ in order to, usually, make A sparser (i.e., contain more zero entries). These three types of invertible **elementary matrices** are:

1. Rescaling Matrices: These $m \times m$ matrices multiply a given row and/or column of $A \in \mathbb{C}^{m \times n}$ by a scalar $\alpha \in \mathbb{C}$. We will denote them by

$$M(j, \alpha) := \text{diag}(\mathbf{1} + (\alpha - 1)\mathbf{e}_j) = I_m + (\alpha - 1)\mathbf{e}_j \mathbf{e}_j^*$$

for any given $\alpha \in \mathbb{C}$ and $j \in [m]$. If multiplied against $A \in \mathbb{C}^{m \times n}$ from the left, $M(j, \alpha) \in \mathbb{C}^{m \times m}$ will multiply the j^{th} row of A by α . If multiplied against $A \in \mathbb{C}^{m \times n}$ on the right, $M(j, \alpha) \in \mathbb{C}^{n \times n}$ will multiply the j^{th} column of A by α .

Example 2.5.1. Let $M(0, \alpha) = \begin{pmatrix} \alpha & 0 \\ 0 & 1 \end{pmatrix} \in \mathbb{C}^{2 \times 2}$. Then, we can see that both

$$\begin{aligned} M(0, \alpha) \begin{pmatrix} a & b \\ c & d \end{pmatrix} &= \begin{pmatrix} \alpha & 0 \\ 0 & 1 \end{pmatrix} \begin{pmatrix} a & b \\ c & d \end{pmatrix} = \begin{pmatrix} \alpha a & \alpha b \\ c & d \end{pmatrix}, \text{ and} \\ \begin{pmatrix} a & b \\ c & d \end{pmatrix} M(0, \alpha) &= \begin{pmatrix} a & b \\ c & d \end{pmatrix} \begin{pmatrix} \alpha & 0 \\ 0 & 1 \end{pmatrix} = \begin{pmatrix} \alpha a & b \\ \alpha c & d \end{pmatrix}. \end{aligned}$$

Exercise 2.5.4. Show that $(M(j, \alpha))^{-1} = I_m + (1/\alpha - 1)\mathbf{e}_j \mathbf{e}_j^* = M(j, \alpha^{-1})$ for all $\alpha \neq 0$ and $j \in [m]$.

2. Summing Matrices: These $m \times m$ matrices add a multiple of one row/column to another row/column. We will denote them by

$$S(j, k, \alpha) := I_m + \alpha E^{(j,k)} = I_m + \alpha \mathbf{e}_j \mathbf{e}_k^*$$

for any given $\alpha \in \mathbb{C}$ and $j, k \in [m]$ with $j \neq k$. Given $A \in \mathbb{C}^{m \times n}$, the product $S(j, k, \alpha)A$ effectively adds $\alpha(\text{row } k \text{ of } A)$ to row j of A , and then stores the result back in row j . Similarly, if $S(j, k, \alpha) \in \mathbb{C}^{n \times n}$ is multiplied against A from the right it will add $\alpha(\text{column } j \text{ of } A)$ to column k of A , and then store the result back in column k .

Example 2.5.2. Let $S(0, 1, \alpha) = \begin{pmatrix} 1 & \alpha \\ 0 & 1 \end{pmatrix} \in \mathbb{C}^{2 \times 2}$. Then, we can see that both

$$S(0, 1, \alpha) \begin{pmatrix} a & b \\ c & d \end{pmatrix} = \begin{pmatrix} 1 & \alpha \\ 0 & 1 \end{pmatrix} \begin{pmatrix} a & b \\ c & d \end{pmatrix} = \begin{pmatrix} a + \alpha c & b + \alpha d \\ c & d \end{pmatrix}, \text{ and}$$

$$\begin{pmatrix} a & b \\ c & d \end{pmatrix} S(0, 1, \alpha) = \begin{pmatrix} a & b \\ c & d \end{pmatrix} \begin{pmatrix} 1 & \alpha \\ 0 & 1 \end{pmatrix} = \begin{pmatrix} a & b + \alpha a \\ c & d + \alpha c \end{pmatrix}.$$

Exercise 2.5.5. Show that $(S(j, k, \alpha))^{-1} = I_m - \alpha \mathbf{e}_j \mathbf{e}_k^* = S(j, k, -\alpha)$ for all $\alpha \in \mathbb{C}$ and $j, k \in [m]$ with $j \neq k$.

3. Atomic Permutation Matrices: These $m \times m$ matrices swap two rows/columns of a given matrix. We will denote them by

$$P(j, k) := I_m - \mathbf{e}_j \mathbf{e}_j^* - \mathbf{e}_k \mathbf{e}_k^* + \mathbf{e}_j \mathbf{e}_k^* + \mathbf{e}_k \mathbf{e}_j^*$$

for any given $j, k \in [m]$ with $j \neq k$. If multiplied against $A \in \mathbb{C}^{m \times n}$ from the left, $P(j, k) \in \mathbb{C}^{m \times m}$ will swap the j^{th} and k^{th} rows of A . If multiplied against $A \in \mathbb{C}^{m \times n}$ on the right, $P(j, k) \in \mathbb{C}^{n \times n}$ will swap the j^{th} and k^{th} columns of A .

Example 2.5.3. Let $P(0, 1) = \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix} \in \mathbb{C}^{2 \times 2}$. Then, we can see that both

$$P(0, 1) \begin{pmatrix} a & b \\ c & d \end{pmatrix} = \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix} \begin{pmatrix} a & b \\ c & d \end{pmatrix} = \begin{pmatrix} c & d \\ a & b \end{pmatrix}, \text{ and}$$

$$\begin{pmatrix} a & b \\ c & d \end{pmatrix} P(0, 1) = \begin{pmatrix} a & b \\ c & d \end{pmatrix} \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix} = \begin{pmatrix} b & a \\ d & c \end{pmatrix}.$$

Exercise 2.5.6. Show that $(P(j, k))^{-1} = P(j, k) \in \mathbb{C}^{m \times m}$ for all $j, k \in [m]$ with $j \neq k$.

Exercise 2.5.7. Let $P = \prod_{\ell=0}^{q-1} P(j_\ell, k_\ell) \in \mathbb{C}^{n \times n}$, where $j_\ell, k_\ell \in [n]$ with $j_\ell \neq k_\ell$ for all $\ell \in [q]$, be a product of q atomic permutation matrices. Show that $P^{-1} = P^*$.

Having briefly reviewed Gaussian elimination, we will now return to our attempt to use a QR decomposition of a low rank matrix to try to compress it as much as possible. We will now show how Gaussian elimination can be used to help us improve on what we have

already achieved above.

Back to Near-Optimal Compression of Low Rank Matrices: Consider a QR decomposition of $A \in \mathbb{C}^{m \times n}$, $A = QR$. Recalling that $R \in \mathbb{C}^{r \times n}$ will be upper triangular, we further note that there will be a permutation matrix $P \in \mathbb{C}^{n \times n}$ so that RP will be both upper triangular *and* have $(RP)_{j,j} \neq 0$ for all $j \in [r]$.⁶ In particular, one can see that P can always be represented by a product of at most $r - 1$ atomic permutation matrices which encode the process of swapping column 1 of R with the first column, j_1 , of R that has $R_{1,j_1} \neq 0$, then swapping column 2 with the first column, j_2 , that has $R_{2,j_2} \neq 0$, etc.. As a result, we can see that remembering (i.e., storing) P requires us to remember at most $r - 1$ values in $[n]$ (i.e., the columns $j_1, j_2, \dots, j_{r-1} \in [n]$ of R satisfying $j_\ell = \min\{k \in [n] \mid R_{\ell,k} \neq 0\}$).

Using that RP will be both upper triangular and have $(RP)_{j,j} \neq 0$ for all $j \in [r]$, we can now further see that there will exist an invertible matrix $T \in \mathbb{C}^{r \times r}$ consisting of a product of at most $\frac{r^2+r}{2}$ elementary summing and rescaling matrices such that

$$TRP = \left(I_r | \tilde{R}\right) \in \mathbb{C}^{r \times n}. \quad (2.13)$$

That is, we can carry out Gaussian elimination to transform the first r columns of RP into the $r \times r$ identity matrix. Note that $\tilde{R} \in \mathbb{C}^{r \times (n-r)}$ in (2.13). Recalling that our goal is to compactly represent $A \in \mathbb{C}^{m \times n}$, we can now see that

$$A = QR = QT^{-1}TRPP^{-1} = (QT^{-1}) \left(I_r | \tilde{R}\right) P^*,$$

where we have used both (2.13) and that $P^{-1} = P^*$ in the final equality.

Letting $\tilde{Q} = QT^{-1} \in \mathbb{C}^{m \times r}$ we can finally see that $A = \tilde{Q} \left(I_r | \tilde{R}\right) P^*$. Thus, to represent $A \in \mathbb{C}^{m \times n}$ we need to store $\tilde{Q} \in \mathbb{C}^{m \times r}$, $\tilde{R} \in \mathbb{C}^{r \times (n-r)}$, and $P \in \mathbb{C}^{n \times n}$. Recalling from above that we can store the permutation matrix P by remembering at most $r - 1$ values in $[n]$, we finally see that we can always represent any rank r matrix $A \in \mathbb{C}^{m \times n}$ by storing just $mr + nr - r^2$ complex values (the optimal number!), plus at most $r - 1$ additional integers in $[n]$. This is a clear improvement over (2.12).

To conclude, we briefly mention that there are other factors we might want to consider when storing A in a factorized form beyond the total number of entries the factorization requires us to store. For example, we might also want to ensure that both $\tilde{Q} \in \mathbb{C}^{m \times r}$ and $\tilde{R} \in \mathbb{C}^{r \times (n-r)}$ are “well behaved”. We will describe in some more detail what “well behaved” might mean in Section 3.1, as well as how one might come up with a good low rank matrix to store in the first place. For a journal article that uses related ideas to those discussed in this section to produce a similar compressed representation of a low rank matrix we refer the interested reader to [7]. After finishing Chapter 2 and 3.1 the attentive reader will know everything they need to know in order to begin digesting its contents.

⁶One can revisit Example 2.4.7 see that we do generally need a permutation matrix for this to be true.

2.6 Set Addition, Orthogonal Projections, and Perpendicular Subspaces

We will now discuss even more of the useful properties possessed by orthonormal bases. The first of these are related to set addition.

Definition 2.6.1 (Set Sums, Subtractions, and Rescalings). *Let S and T be subsets of $\mathbb{C}^n$. We define the **(Minkowski) sum** of S and T , denoted by $S + T$, to be the set*

$$S + T := \{\mathbf{x} + \mathbf{y} \mid \mathbf{x} \in S, \mathbf{y} \in T\}.$$

*Similarly, for $\alpha \in \mathbb{C}$, we define the **set rescaling** $\alpha S \subset \mathbb{C}^n$ to be $\{\alpha \mathbf{x} \mid \mathbf{x} \in S\}$. We also define the subtraction of two sets to be*

$$S - T := S + (-1)T = \{\mathbf{x} - \mathbf{y} \mid \mathbf{x} \in S, \mathbf{y} \in T\} \subset \mathbb{C}^n.$$

Note that if $\mathbf{0} \in S$, then $T \subset S + T$. Similarly, if $\mathbf{0} \in T$, then $S \subset S + T$. As a result, if $\mathbf{0} \in S \cap T$, then $S \cup T \subset S + T$ (check this!). For similar reasons, the sum of two linear subspaces U and V of $\mathbb{C}^n$ will also always be a larger linear subspace of $\mathbb{C}^n$ containing both U and V (i.e., U and V will be subspaces of $U + V$).

Lemma 2.6.2. *Let $U, V \subset \mathbb{C}^n$ both be linear subspaces of $\mathbb{C}^n$. Then $U + V$ is also a linear subspace of $\mathbb{C}^n$.*

Proof. It suffices to show that $\text{span}(U + V) \subset U + V$ and we'll be finished (why?). We can see that

$$\begin{aligned} \mathbf{x} \in \text{span}(U + V) &\implies \exists p \in \mathbb{N} \text{ s.t. } \mathbf{x} = \sum_{\ell \in [p]} \beta_\ell \mathbf{x}_\ell \text{ with } \{\beta_\ell\}_{\ell \in [p]} \subset \mathbb{C} \text{ \& } \{\mathbf{x}_\ell\}_{\ell \in [p]} \subset U + V \\ &\implies \mathbf{x} = \sum_{\ell \in [p]} \beta_\ell (\mathbf{u}_\ell + \mathbf{v}_\ell) \text{ for } \{\mathbf{u}_\ell\}_{\ell \in [p]} \subset U \text{ \& } \{\mathbf{v}_\ell\}_{\ell \in [p]} \subset V \\ &\implies \mathbf{x} = \underbrace{\left(\sum_{\ell \in [p]} \beta_\ell \mathbf{u}_\ell \right)}_{=:\mathbf{u}} + \underbrace{\left(\sum_{\ell \in [p]} \beta_\ell \mathbf{v}_\ell \right)}_{=:\mathbf{v}}. \end{aligned}$$

We are now finished since, above, $\mathbf{u} \in \text{span}(U) = U$ and $\mathbf{v} \in \text{span}(V) = V$. Hence, $\mathbf{x} \in U + V$. □

Exercise 2.6.1. *Let $U, V \subset \mathbb{C}^n$ both be linear subspaces of $\mathbb{C}^n$. Show that*

$$\max\{\dim(U), \dim(V)\} \leq \dim(U + V) \leq \dim(U) + \dim(V),$$

where $\dim(U) \in [n+1]$ denotes the dimension of U , etc.. When will $\max\{\dim(U), \dim(V)\} = \dim(U + V)$? When will $\dim(U + V) = \dim(U) + \dim(V)$?

Exercise 2.6.2. Let $A, B \in \mathbb{C}^{m \times n}$. Show that if A has rank r and B has rank s , then $A + B$ has rank at most $r + s$.

As we shall soon see, the sum of two “orthogonal” linear subspaces of $\mathbb{C}^n$, U and V , will behave much more predictably than the sum of two arbitrary linear subspaces of $\mathbb{C}^n$. In particular, orthonormal bases of each summed subspace U and V can be directly combined to create a new orthonormal basis of $U + V$.

Definition 2.6.3 (Perpendicular Subspaces). Let U and V be linear subspaces of $\mathbb{C}^n$. We say that U and V are **perpendicular**, or **orthogonal**, if $\langle \mathbf{u}, \mathbf{v} \rangle = 0$ for all $\mathbf{u} \in U$ and $\mathbf{v} \in V$. We will also denote this by writing $U \perp V$.

Lemma 2.6.4. Suppose that B_U is an orthonormal basis of a linear subspace $U \subset \mathbb{C}^n$, B_V is an orthonormal basis of a linear subspace $V \subset \mathbb{C}^n$, and $U \perp V$. Then $B_U \cup B_V$ is an orthonormal basis of $U + V$.

Proof. Since B_U and B_V are orthonormal, every element in $B_U \cup B_V$ has norm 1. Thus, we only need to show that $B_U \cup B_V$ is orthogonal. Let $\mathbf{x}, \mathbf{y} \in B_U \cup B_V$. Since B_U is orthogonal, if $\mathbf{x} \in B_U$ and $\mathbf{y} \in B_U$, then $\langle \mathbf{x}, \mathbf{y} \rangle = 0$. Since B_V is orthogonal, if $\mathbf{x} \in B_V$ and $\mathbf{y} \in B_V$, then $\langle \mathbf{x}, \mathbf{y} \rangle = 0$. Since $U \perp V$, if $\mathbf{x} \in B_U$ and $\mathbf{y} \in B_V$, or $\mathbf{x} \in B_V$ and $\mathbf{y} \in B_U$, then $\langle \mathbf{x}, \mathbf{y} \rangle = 0$. Hence, $B_U \cup B_V$ is orthogonal.

Now we will show that $B_U \cup B_V$ is a basis of $U + V$. Since $B_U \cup B_V$ is orthonormal, its entries are linearly independent, so it remains to show that $\text{span}(B_U \cup B_V) = U + V$. Let $B_U = \{\mathbf{b}_0, \dots, \mathbf{b}_{r-1}\}$ and $B_V = \{\mathbf{d}_0, \dots, \mathbf{d}_{s-1}\}$. Then, for any vector $\mathbf{x} \in \mathbb{C}^n$, it holds that

$$\begin{aligned} \mathbf{x} \in U + V &\iff \exists \mathbf{u} \in U, \mathbf{v} \in V \text{ such that } \mathbf{x} = \mathbf{u} + \mathbf{v} \\ &\iff \mathbf{x} = \left(\sum_{j \in [r]} \alpha_j \mathbf{b}_j \right) + \left(\sum_{j \in [s]} \beta_j \mathbf{d}_j \right) \text{ for some } \{\alpha_j\}_{j \in [r]} \cup \{\beta_j\}_{j \in [s]} \subset \mathbb{C} \\ &\iff \mathbf{x} \in \text{span}(B_U \cup B_V). \end{aligned}$$

Therefore, $U + V = \text{span}(B_U \cup B_V)$. □

Corollary 2.6.5. If $U, V \subset \mathbb{C}^n$ are linear subspaces of $\mathbb{C}^n$, and $U \perp V$, then $\dim(U + V) = \dim(U) + \dim(V)$.

Given any linear subspace $U \subset \mathbb{C}^n$, we define

$$U^\perp := \{\mathbf{x} \in \mathbb{C}^n \mid \langle \mathbf{x}, \mathbf{y} \rangle = 0 \ \forall \mathbf{y} \in U\}.$$

In other words, $U^\perp$ is the set of all vectors orthogonal to everything in U . We will next show that $U^\perp$ is also a linear subspace of $\mathbb{C}^n$.

Lemma 2.6.6. Let $U \subset \mathbb{C}^n$ be a linear subspace of $\mathbb{C}^n$. Then, $U^\perp$ is also a linear subspace of $\mathbb{C}^n$.

Proof. It suffices to show that $\text{span}(U^\perp) \subset U^\perp$ (why?). Let $\mathbf{x} \in \text{span}(U^\perp)$. Then, $\mathbf{x}$ is a linear combination of elements in $U^\perp$ so that $\exists p \in \mathbb{N}$, $\{\mathbf{x}_\ell\}_{\ell \in [p]} \subset U^\perp$, and $\{\alpha_\ell\}_{\ell \in [p]} \subset \mathbb{C}$ with $\mathbf{x} = \sum_{\ell \in [p]} \alpha_\ell \mathbf{x}_\ell$. Now we can see that for every $\mathbf{y} \in U$ we have

$$\langle \mathbf{x}, \mathbf{y} \rangle = \left\langle \sum_{\ell \in [p]} \alpha_\ell \mathbf{x}_\ell, \mathbf{y} \right\rangle = \sum_{\ell \in [p]} \overline{\alpha_\ell} \langle \mathbf{x}_\ell, \mathbf{y} \rangle = 0$$

since $\{\mathbf{x}_\ell\}_{\ell \in [p]} \subset U^\perp$. Hence, $\mathbf{x} \in U^\perp$. $\square$

We are now prepared to define orthogonal projections with respect to a given orthonormal set. Let $U = \{\mathbf{u}_j\}_{j \in [r]}$ be an orthonormal subset of $\mathbb{C}^n$. We define the **orthogonal projection of $\mathbf{x}$ onto $\text{span}(U)$ in terms of U** to be the function $P_U : \mathbb{C}^n \rightarrow \text{span}(U)$ defined by

$$P_U(\mathbf{x}) = \sum_{j \in [r]} \langle \mathbf{u}_j, \mathbf{x} \rangle \mathbf{u}_j$$

for all $\mathbf{x} \in \mathbb{C}^n$. Note that this definition explicitly depends on the orthonormal basis U of $\text{span}(U)$ that we started with. The idea behind projecting onto a linear subspace $\text{span}(U)$, however, is that the projection should return the portion of $\mathbf{x}$ “living inside” the linear subspace $\text{span}(U)$. That is, it’s the *span* of U that matters to us, not the set U itself. If, e.g., we pick a new orthonormal set V with the same exact span as U , then it really shouldn’t matter whether we project onto $\text{span}(U) = \text{span}(V)$ using U or V . We should get the same answer either way. The next result will help us show that this is indeed the case.

Lemma 2.6.7. *Let $U = \{\mathbf{u}_\ell\}_{\ell \in [r]}$ and $V = \{\mathbf{v}_\ell\}_{\ell \in [r]}$ be two orthonormal bases of the same linear subspace $\mathcal{L} = \text{span}(U) = \text{span}(V) \subset \mathbb{C}^n$. Then,*

$$\langle \mathbf{u}_j, \mathbf{x} \rangle = \sum_{\ell \in [r]} \langle \mathbf{v}_\ell, \mathbf{x} \rangle \langle \mathbf{u}_j, \mathbf{v}_\ell \rangle$$

for all $\mathbf{x} \in \mathbb{C}^n$ and $j \in [r]$.

Proof. Let $\mathbf{x} \in \mathbb{C}^n$. Extend V to an orthonormal basis $\tilde{V}$ of all of $\mathbb{C}^n$ by appealing to Theorem 2.4.5. The orthonormal set $\tilde{V}$ will take the form $\tilde{V} = \{\mathbf{v}_0, \dots, \mathbf{v}_{r-1}, \mathbf{w}_r, \dots, \mathbf{w}_{n-1}\}$ for some $\mathbf{w}_r, \dots, \mathbf{w}_{n-1} \in \mathbb{C}^n$. Since $\tilde{V}$ is an orthonormal basis of $\mathbb{C}^n$ we can write $\mathbf{x}$ as

$$\mathbf{x} = \sum_{\ell \in [r]} \langle \mathbf{v}_\ell, \mathbf{x} \rangle \mathbf{v}_\ell + \sum_{\ell=r}^{n-1} \langle \mathbf{w}_\ell, \mathbf{x} \rangle \mathbf{w}_\ell.$$

Additionally, since each $\mathbf{u}_j \in U$ is in the span of V , we have for all $r \leq \ell \leq n-1$ that

$$\langle \mathbf{u}_j, \mathbf{w}_\ell \rangle = \left\langle \sum_{k \in [r]} \alpha_{j,k} \mathbf{v}_k, \mathbf{w}_\ell \right\rangle = \sum_{k \in [r]} \overline{\alpha_{j,k}} \langle \mathbf{v}_k, \mathbf{w}_\ell \rangle = 0$$

since $\tilde{V}$ is orthogonal. Hence,

$$\begin{aligned}\langle \mathbf{u}_j, \mathbf{x} \rangle &= \left\langle \mathbf{u}_j, \sum_{\ell \in [r]} \langle \mathbf{v}_\ell, \mathbf{x} \rangle \mathbf{v}_\ell + \sum_{\ell=r}^{n-1} \langle \mathbf{w}_\ell, \mathbf{x} \rangle \mathbf{w}_\ell \right\rangle \\ &= \sum_{\ell \in [r]} \langle \mathbf{v}_\ell, \mathbf{x} \rangle \langle \mathbf{u}_j, \mathbf{v}_\ell \rangle + \sum_{\ell=r}^{n-1} \langle \mathbf{w}_\ell, \mathbf{x} \rangle \langle \mathbf{u}_j, \mathbf{w}_\ell \rangle \\ &= \sum_{\ell \in [r]} \langle \mathbf{v}_\ell, \mathbf{x} \rangle \langle \mathbf{u}_j, \mathbf{v}_\ell \rangle.\end{aligned}$$

□

Using this lemma allows us to show that the orthogonal projection $P_U : \mathbb{C}^n \rightarrow \text{span}(U)$ only depends on $\text{span}(U)$, and not on the orthonormal set U itself.

Theorem 2.6.8. *Let $U = \{\mathbf{u}_j\}_{j \in [m]}$ and $V = \{\mathbf{v}_\ell\}_{\ell \in [m]}$ be two orthonormal bases of the same linear subspace $\mathcal{L} \subset \mathbb{C}^n$. Then, $P_U = P_V$.*

Proof. Let $\mathbf{x} \in \mathbb{C}^n$. Appealing to Lemma 2.6.7, we have that

$$\begin{aligned}P_U(\mathbf{x}) &= \sum_{j \in [m]} \langle \mathbf{u}_j, \mathbf{x} \rangle \mathbf{u}_j = \sum_{j \in [m]} \left(\sum_{\ell \in [m]} \langle \mathbf{v}_\ell, \mathbf{x} \rangle \langle \mathbf{u}_j, \mathbf{v}_\ell \rangle \right) \mathbf{u}_j \\ &= \sum_{\ell \in [m]} \langle \mathbf{v}_\ell, \mathbf{x} \rangle \left(\sum_{j \in [m]} \langle \mathbf{u}_j, \mathbf{v}_\ell \rangle \mathbf{u}_j \right) = \sum_{\ell \in [m]} \langle \mathbf{v}_\ell, \mathbf{x} \rangle \mathbf{v}_\ell = P_V(\mathbf{x}),\end{aligned}$$

where we have also used that each $\mathbf{v}_\ell \in V$ is in the span of U (recall (2.11)). □

We now know that an orthogonal projection only depends on the linear subspace of $\mathbb{C}^n$ onto which one projects. Thus, for any linear subspace $\mathcal{L} \subset \mathbb{C}^n$ we can define **the orthogonal projection onto $\mathcal{L}$** , denoted by $P_{\mathcal{L}} : \mathbb{C}^n \rightarrow \mathcal{L}$, to be $P_{\mathcal{L}} := P_U$ where U is any orthonormal basis of $\mathcal{L}$ you like.

Example 2.6.9. *The orthogonal projection onto the x -axis of $\mathbb{C}^2$ is the function $P_{\mathbf{x}\text{-axis}}$ from $\mathbb{C}^2$ to the x -axis which sends the vector $\mathbf{x} = \begin{pmatrix} x_0 \\ x_1 \end{pmatrix} \in \mathbb{C}^2$ to the vector*

$$P_{\mathbf{x}\text{-axis}}(\mathbf{x}) = \left\langle \begin{pmatrix} 1 \\ 0 \end{pmatrix}, \begin{pmatrix} x_0 \\ x_1 \end{pmatrix} \right\rangle \begin{pmatrix} 1 \\ 0 \end{pmatrix} = x_0 \begin{pmatrix} 1 \\ 0 \end{pmatrix} = \begin{pmatrix} x_0 \\ 0 \end{pmatrix}.$$

Exercise 2.6.3. *Let $\mathcal{L}$ be a linear subspace of $\mathbb{C}^n$. Verify that $P_{\mathcal{L}} : \mathbb{C}^n \rightarrow \mathcal{L}$ is a linear function (i.e., that $P_{\mathcal{L}}(\alpha \mathbf{x} + \beta \mathbf{y}) = \alpha P_{\mathcal{L}}(\mathbf{x}) + \beta P_{\mathcal{L}}(\mathbf{y})$ holds for all $\mathbf{x}, \mathbf{y} \in \mathbb{C}^n$ and $\alpha, \beta \in \mathbb{C}$).*

Exercise 2.6.4. Let $B = \{\mathbf{b}_j\}_{j \in [r]} \subset \mathbb{C}^n$ be an orthonormal basis of $\mathcal{L} = \text{span}(B)$. Complete B to be an orthonormal basis $\tilde{B} = \{\mathbf{b}_j\}_{j \in [r]} \cup \{\mathbf{u}_\ell\}_{\ell=r}^{n-1} \subset \mathbb{C}^n$ of all of $\mathbb{C}^n$ using Theorem 2.4.5. Prove that $\{\mathbf{u}_\ell\}_{\ell=r}^{n-1}$ is an orthonormal basis of $\mathcal{L}^\perp$.

The following theorem characterizes many of the most important properties of orthogonal projections.

Theorem 2.6.10. Let $\mathcal{L}$ be a linear subspace of $\mathbb{C}^n$, and let $\mathbf{x} \in \mathbb{C}^n$.

1. If $\mathbf{x} \in \mathcal{L}$ then $P_{\mathcal{L}}(\mathbf{x}) = \mathbf{x}$. As a consequence, $P_{\mathcal{L}}(P_{\mathcal{L}}(\mathbf{x})) = P_{\mathcal{L}}(\mathbf{x})$ always holds.
2. $\mathbf{x} - P_{\mathcal{L}}(\mathbf{x}) \in \mathcal{L}^\perp$.
3. $\|\mathbf{x}\|_2^2 = \|P_{\mathcal{L}}(\mathbf{x})\|_2^2 + \|\mathbf{x} - P_{\mathcal{L}}(\mathbf{x})\|_2^2$.

Proof. We prove each part below.

1. See Exercise 2.6.5.
2. Let $U = \{\mathbf{u}_j\}_{j \in [m]} \subset \mathbb{C}^n$ be an orthonormal basis of $\mathcal{L}$, and $\mathbf{y} \in \mathcal{L}$. Then $\mathbf{y} = \sum_{j \in [m]} \alpha_j \mathbf{u}_j$, so that

$$\begin{aligned}
 \langle \mathbf{x} - P_{\mathcal{L}}(\mathbf{x}), \mathbf{y} \rangle &= \sum_{j \in [m]} \alpha_j \langle \mathbf{x} - P_{\mathcal{L}}(\mathbf{x}), \mathbf{u}_j \rangle \\
 &= \sum_{j \in [m]} \alpha_j (\langle \mathbf{x}, \mathbf{u}_j \rangle - \langle P_{\mathcal{L}}(\mathbf{x}), \mathbf{u}_j \rangle) \\
 &= \sum_{j \in [m]} \alpha_j \left(\langle \mathbf{x}, \mathbf{u}_j \rangle - \left\langle \sum_{\ell \in [m]} \langle \mathbf{u}_\ell, \mathbf{x} \rangle \mathbf{u}_\ell, \mathbf{u}_j \right\rangle \right) \\
 &= \sum_{j \in [m]} \alpha_j (\langle \mathbf{x}, \mathbf{u}_j \rangle - \overline{\langle \mathbf{u}_j, \mathbf{x} \rangle}) = 0,
 \end{aligned}$$

since $\langle \mathbf{x}, \mathbf{u}_j \rangle = \overline{\langle \mathbf{u}_j, \mathbf{x} \rangle}$.

3. From (1) and (2) above we know that $P_{\mathcal{L}}(\mathbf{x})$ and $\mathbf{x} - P_{\mathcal{L}}(\mathbf{x})$ are orthogonal. Hence, normalizing them will produce an orthonormal set whose span contains $\mathbf{x}$. This part now follows from the Pythagorean theorem (see Theorem 2.3.9 and Figure 2.2).

□

Theorem 2.6.10 tells us that we can write $\mathbb{C}^n = \mathcal{L} + \mathcal{L}^\perp$ for any linear subspace $\mathcal{L} \subset \mathbb{C}^n$. In some sense we already know this though – recall, e.g., Exercise 2.6.4! The main contribution of Theorem 2.6.10 is that it expresses this fact in a much simple way using orthogonal projections. This more simply expressed property then also allows for a simpler application of the Pythagorean theorem (see Figure 2.2).

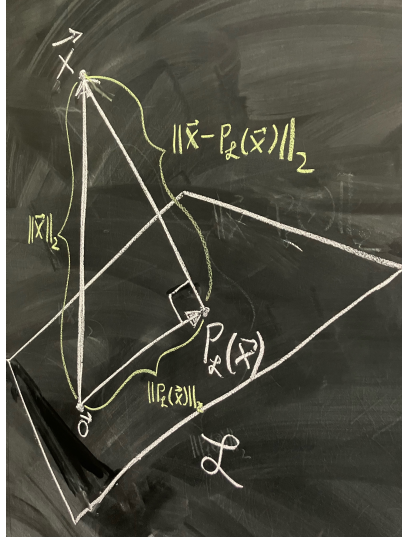


Figure 2.2: A pictorial representation of the projection of $\mathbf{x} \in \mathbb{C}^n$ onto a linear subspace $\mathcal{L} \subset \mathbb{C}^n$. Note that the right triangle whose hypotenuse is of length $\|\mathbf{x}\|_2$ will in fact be entirely contained in the two-dimensional linear subspace spanned by $\{P_{\mathcal{L}}(\mathbf{x}), \mathbf{x} - P_{\mathcal{L}}(\mathbf{x})\}$.

Exercise 2.6.5. Prove part (1) of Theorem 2.6.10.

The next lemma demonstrates yet another incredibly useful way of characterizing what the orthogonal projection onto a linear subspace actually does.

Lemma 2.6.11. Let $\mathbf{x} \in \mathbb{C}^n$. Then $\|\mathbf{x} - P_{\mathcal{L}}(\mathbf{x})\|_2 < \|\mathbf{x} - \mathbf{y}\|_2$ for all $\mathbf{y} \in \mathcal{L} \setminus \{P_{\mathcal{L}}(\mathbf{x})\}$ (i.e., for all $\mathbf{y} \in \mathcal{L}$ with $\mathbf{y} \neq P_{\mathcal{L}}(\mathbf{x})$).

Proof. We have from Theorem 2.6.10 that

$$\begin{aligned} \|\mathbf{x} - \mathbf{y}\|_2^2 &= \|P_{\mathcal{L}}(\mathbf{x} - \mathbf{y})\|_2^2 + \|(\mathbf{x} - \mathbf{y}) - P_{\mathcal{L}}(\mathbf{x} - \mathbf{y})\|_2^2 \\ &= \|P_{\mathcal{L}}(\mathbf{x}) - P_{\mathcal{L}}(\mathbf{y})\|_2^2 + \|\mathbf{x} - P_{\mathcal{L}}(\mathbf{x}) - \mathbf{y} + P_{\mathcal{L}}(\mathbf{y})\|_2^2 \\ &= \|P_{\mathcal{L}}(\mathbf{x}) - \mathbf{y}\|_2^2 + \|\mathbf{x} - P_{\mathcal{L}}(\mathbf{x})\|_2^2 \\ &> \|\mathbf{x} - P_{\mathcal{L}}(\mathbf{x})\|_2^2. \end{aligned}$$

Here we have used that $P_{\mathcal{L}}(\mathbf{y}) = \mathbf{y}$ for all $\mathbf{y} \in \mathcal{L}$ and that $\|P_{\mathcal{L}}(\mathbf{x}) - \mathbf{y}\|_2 > 0$ must hold since $P_{\mathcal{L}}(\mathbf{x}) - \mathbf{y} \neq \mathbf{0}$. $\square$

Looking at Lemma 2.6.11 we can see that $P_{\mathcal{L}}(\mathbf{x}) \in \mathcal{L}$ is the *unique* closest point to $\mathbf{x}$ in $\mathcal{L}$ with respect to ℓ^2 -distances (recall Figure 2.2 as well). As a consequence, we can see that in fact

$$P_{\mathcal{L}}(\mathbf{x}) = \arg \min_{\mathbf{y} \in \mathcal{L}} \|\mathbf{x} - \mathbf{y}\|_2$$

also holds. That is, we could have defined $P_{\mathcal{L}}(\mathbf{x})$ to be the closest point in $\mathcal{L}$ to $\mathbf{x}$ in the first place if we had wanted. We will next discuss how to represent $P_{\mathcal{L}}$ as a matrix.

2.6.1 Representing Orthogonal Projections with Matrices

The following fundamental matrices are used to represent all orthogonal projections.

Definition 2.6.12 (Orthonormal and Unitary Matrices). *A matrix $Q \in \mathbb{C}^{m \times n}$ with orthonormal columns will be called an **orthonormal matrix**. If $Q \in \mathbb{C}^{n \times n}$ is both orthonormal and square we will call it a **unitary matrix**.*

In fact the attentive reader will recognize that we have already been introduced to orthonormal matrices. In particular, the “ Q ” in a QR decomposition of a given matrix is always an orthonormal matrix. The next few highly recommended exercises will introduce you to some of the very useful properties of orthonormal matrices.

Exercise 2.6.6. Let $Q \in \mathbb{C}^{m \times n}$ be an orthonormal matrix. Prove that $n \leq m$.

Exercise 2.6.7. Let $Q \in \mathbb{C}^{m \times n}$. Show that $Q^*Q = I_n$ if and only if Q is orthonormal.

Exercise 2.6.8. Let $Q \in \mathbb{C}^{m \times n}$ be an orthonormal matrix and $\mathbf{x}, \mathbf{y} \in \mathbb{C}^m$. Show that $(I_m - QQ^*)\mathbf{x} = \mathbf{x} - QQ^*\mathbf{x}$ is orthogonal to $QQ^*\mathbf{y}$.

Let $B = \{\mathbf{q}_\ell\}_{\ell \in [r]} \subset \mathbb{C}^n$ be an orthonormal basis of a linear subspace $\mathcal{L}$. We can form an orthonormal matrix $Q \in \mathbb{C}^{n \times r}$ by letting the columns of Q be the elements of B so that $Q_{:, \ell} = \mathbf{q}_\ell$ for all $\ell \in [r]$, so that

$$Q = \begin{pmatrix} | & | & \cdots & | \\ \mathbf{q}_0 & \mathbf{q}_1 & \cdots & \mathbf{q}_{r-1} \\ | & | & \cdots & | \end{pmatrix}.$$

We can represent the orthogonal projection onto $\mathcal{L} = \mathcal{C}(Q) = \text{span}(B)$ using an **orthogonal projection matrix** $QQ^* \in \mathbb{C}^{n \times n}$ by

$$P_{\mathcal{L}} = QQ^* = \sum_{j \in [r]} \mathbf{q}_j \mathbf{q}_j^*. \quad (2.14)$$

To see that (2.14) holds it suffices to check that $P_{\mathcal{L}}(\mathbf{x}) = QQ^*\mathbf{x} = \left(\sum_{j \in [r]} \mathbf{q}_j \mathbf{q}_j^*\right) \mathbf{x}$ for all $\mathbf{x} \in \mathbb{C}^n$ (see Exercise 2.6.9). Let $\mathbf{x} \in \mathbb{C}^n$. We have that

$$\begin{aligned}
QQ^*\mathbf{x} &= Q \begin{pmatrix} -\overline{\mathbf{q}_0} & - \\ \vdots & \\ -\overline{\mathbf{q}_{r-1}} & - \end{pmatrix} \mathbf{x} = Q \begin{pmatrix} \langle \mathbf{q}_0, \mathbf{x} \rangle \\ \vdots \\ \langle \mathbf{q}_{r-1}, \mathbf{x} \rangle \end{pmatrix} = \begin{pmatrix} | & | & & | \\ \mathbf{q}_0 & \mathbf{q}_1 & \cdots & \mathbf{q}_{r-1} \\ | & | & & | \end{pmatrix} \begin{pmatrix} \langle \mathbf{q}_0, \mathbf{x} \rangle \\ \vdots \\ \langle \mathbf{q}_{r-1}, \mathbf{x} \rangle \end{pmatrix} \\
&= \sum_{j \in [r]} \mathbf{q}_j \langle \mathbf{q}_j, \mathbf{x} \rangle = P_{\mathcal{L}}(\mathbf{x}) \\
&= \sum_{j \in [r]} \mathbf{q}_j \mathbf{q}_j^* \mathbf{x} = \left(\sum_{j \in [r]} \mathbf{q}_j \mathbf{q}_j^* \right) \mathbf{x}.
\end{aligned}$$

Using (2.14) we can also establish the equivalence of orthogonal projection matrices built from orthonormal matrices with the same column spaces.

Lemma 2.6.13. *Let $Q, V \in \mathbb{C}^{m \times n}$ be two orthonormal matrices with the same column span (i.e., with $\mathcal{C}(Q) = \mathcal{C}(V)$). Then $QQ^* = VV^*$.*

Proof. It suffices to show that $QQ^*\mathbf{x} = VV^*\mathbf{x}$ for all $\mathbf{x} \in \mathbb{C}^n$ (see Exercise 2.6.9). Using Theorem 2.6.8 and (2.14) we have that

$$QQ^*\mathbf{x} = P_{\mathcal{C}(Q)}(\mathbf{x}) = P_{\mathcal{C}(V)}(\mathbf{x}) = VV^*\mathbf{x}$$

for all $\mathbf{x} \in \mathbb{C}^n$. We are now finished by Exercise 2.6.9. $\square$

Exercise 2.6.9. *Let $A, B \in \mathbb{C}^{m \times n}$. Suppose that $A\mathbf{x} = B\mathbf{x}$ for all $\mathbf{x} \in \mathbb{C}^n$. Show that $A = B$.*

The next theorem shows that orthonormal projection matrices built from unitary matrices are always equivalent to the identity matrix.

Theorem 2.6.14. *The following are equivalent:*

1. $U \in \mathbb{C}^{n \times n}$ is unitary,
2. $U^*U = I_n$,
3. U^* is unitary, and
4. $UU^* = I_n$.

Proof.

(2) $\iff$ (1): Let $U \in \mathbb{C}^{n \times n}$ and set $\mathbf{u}_j := U_{:,j} \in \mathbb{C}^n$ for all $j \in [n]$. Note that $(U^*U)_{\ell,k} = \langle \mathbf{u}_\ell, \mathbf{u}_k \rangle$ for all $\ell, k \in [n]$. As a result we can see that

$$\begin{aligned} U^*U = I_n &\iff (U^*U)_{\ell,k} = (I_n)_{\ell,k} \quad \forall \ell, k \in [n] \\ &\iff \langle \mathbf{u}_\ell, \mathbf{u}_k \rangle = \begin{cases} 1 & \text{if } \ell = k \\ 0 & \text{else} \end{cases} \\ &\iff \{\mathbf{u}_\ell\}_{\ell \in [n]} \subset \mathbb{C}^n \text{ is an orthonormal set} \\ &\iff U \text{ is unitary.} \end{aligned}$$

Hence, U is unitary if and only if $U^*U = I_n$.

(1) $\implies$ (4): Let $U \in \mathbb{C}^{n \times n}$ be unitary. Then $\mathcal{C}(U) = \mathbb{C}^n$ (see Exercise 2.4.4). Since $\mathbb{C}^n = \text{span}\{\mathbf{e}_j\}_{j \in [n]}$ (i.e., $\mathcal{C}(I_n) = \mathbb{C}^n$) Lemma 2.6.13 tells us that $UU^* = I_n I_n^* = I_n$.

(4) $\iff$ (3): This is the same as (1) $\iff$ (2) with U replaced by U^* .

(3) $\implies$ (2): This is the same as (1) $\implies$ (4) with U replaced by U^* . $\square$

An additional consequence of Theorem 2.6.14 is that a matrix $U \in \mathbb{C}^{n \times n}$ is unitary if and only if $U^* = U^{-1}$. Given this, we can see that we have already met an important family of unitary matrices – the permutation matrices (recall Exercise 2.5.7).

Exercise 2.6.10. Let $U, V \in \mathbb{C}^{n \times n}$ both be unitary. Show that both $UV \in \mathbb{C}^{n \times n}$ and $VU \in \mathbb{C}^{n \times n}$ are then also unitary.

Exercise 2.6.11. Let $U \in \mathbb{C}^{n \times n}$ be unitary. Prove that $\|U\mathbf{x}\|_2 = \|\mathbf{x}\|_2$ for all $\mathbf{x} \in \mathbb{C}^n$.

Exercise 2.6.12. Let $B = \{\mathbf{b}_j\}_{j \in [r]} \subset \mathbb{C}^n$ be an orthonormal basis of $\mathcal{L} = \text{span}(B)$. Complete B to be an orthonormal basis $\tilde{B} = \{\mathbf{b}_j\}_{j \in [r]} \cup \{\mathbf{u}_\ell\}_{\ell=r}^{n-1} \subset \mathbb{C}^n$ of all of $\mathbb{C}^n$ using Theorem 2.4.5. Let $Q \in \mathbb{C}^{n \times r}$ be the orthonormal matrix with $Q_{:,j} = \mathbf{b}_j$ for all $j \in [r]$ and $U \in \mathbb{C}^{n \times (n-r)}$ be the orthonormal matrix with $U_{:,k} = \mathbf{u}_{r+k}$ for all $k \in [n-r]$. Prove that the orthogonal projection onto $\mathcal{L}^\perp$, $P_{\mathcal{L}^\perp} : \mathbb{C}^n \rightarrow \mathcal{L}^\perp$, has the following properties.

1. $P_{\mathcal{L}^\perp} = UU^*$ (Hint: Recall Exercise 2.6.4.).
2. Show that $P_{\mathcal{L}}(\mathbf{x}) + P_{\mathcal{L}^\perp}(\mathbf{x}) = \mathbf{x} = I_n \mathbf{x}$ holds for all $\mathbf{x} \in \mathbb{C}^n$. Conclude that $P_{\mathcal{L}^\perp} = I_n - P_{\mathcal{L}}$.
3. Show that $UU^* = I_n - QQ^* \in \mathbb{C}^{n \times n}$.

Lemma 2.6.15. Let $\mathcal{L}$ and $\mathcal{T}$ be linear subspaces of $\mathbb{C}^n$ such that $\mathcal{L}^\perp = \mathcal{T}$. Then, $\mathcal{T}^\perp = \mathcal{L}$ also holds (i.e., $(\mathcal{L}^\perp)^\perp = \mathcal{L}$).

Proof. We must show that both $\mathcal{L} \subset \mathcal{T}^\perp$ and that $\mathcal{T}^\perp \subset \mathcal{L}$ hold.

$\mathcal{L} \subset \mathcal{T}^\perp$: Let $\mathbf{x} \in \mathcal{L}$. Then $\langle \mathbf{x}, \mathbf{y} \rangle = 0$ for all $\mathbf{y} \in \mathcal{L}^\perp = \mathcal{T}$ by definition of $\mathcal{L}^\perp$. Hence, $\mathbf{x} \in \mathcal{T}^\perp$.

$\mathcal{T}^\perp \subset \mathcal{L}$: Let $\mathbf{x} \in \mathcal{T}^\perp$. By the definition of $\mathcal{T}^\perp$, $\langle \mathbf{x}, \mathbf{y} \rangle = 0$ for all $\mathbf{y} \in \mathcal{T} = \mathcal{L}^\perp$. Hence, $P_{\mathcal{L}^\perp}(\mathbf{x}) = \mathbf{0}$. Now we can see that $\mathbf{x} = P_{\mathcal{L}}(\mathbf{x}) + P_{\mathcal{L}^\perp}(\mathbf{x}) = P_{\mathcal{L}}(\mathbf{x})$ (using, e.g., part (2) of Exercise 2.6.12). Thus, $\mathbf{x} \in \mathcal{L}$. $\square$

Now that we have achieved a good understanding of orthogonal projections and orthonormal matrices we are prepared to discuss the least-squares approach to solving systems of linear equations.

2.6.2 Least-Squares Theory for (Approximately) Solving Systems of Linear Equations

Let $A \in \mathbb{C}^{m \times n}$, $\mathbf{b} \in \mathbb{C}^m$, and suppose that we want to solve the equation $A\mathbf{x} = \mathbf{b}$ for $\mathbf{x} \in \mathbb{C}^n$. The least-squares approach aims to do this by minimizing $f(\mathbf{x}) := \|\mathbf{b} - A\mathbf{x}\|_2^2$ as a function of $\mathbf{x} \in \mathbb{C}^n$. To see why this makes sense, observe that $\mathbf{b} \in \mathcal{C}(A) \iff \exists \mathbf{y} \in \mathbb{C}^n$ such that $\mathbf{b} = A\mathbf{y}$ in which case $f(\mathbf{x}) = \|\mathbf{b} - A\mathbf{x}\|_2^2$ will attain its absolute minimum at $f(\mathbf{y}) = 0$. Furthermore, anytime $f(\mathbf{x}) = 0$ it must in fact be the case that $A\mathbf{x} = \mathbf{b}$. Hence, if $A\mathbf{x} = \mathbf{b}$ has solutions we can indeed find one by minimizing f down to 0.

If, on the other hand, $\mathbf{b} \notin \mathcal{C}(A)$ then $A\mathbf{x} = \mathbf{b}$ won't have any solutions and $\inf_{\mathbf{x} \in \mathbb{C}^n} f(\mathbf{x}) = \inf_{\mathbf{x} \in \mathbb{C}^n} \|\mathbf{b} - A\mathbf{x}\|_2^2 > 0$. Nonetheless, there is absolutely nothing stopping us from still minimizing f in hopes of getting "close" to a solution anyways. Observe that by Theorem 2.6.10

$$\begin{aligned} \|\mathbf{b} - A\mathbf{x}\|_2^2 &= \|P_{\mathcal{C}(A)}(\mathbf{b} - A\mathbf{x})\|_2^2 + \|\mathbf{b} - A\mathbf{x} - P_{\mathcal{C}(A)}(\mathbf{b} - A\mathbf{x})\|_2^2 \\ &= \|P_{\mathcal{C}(A)}(\mathbf{b}) - A\mathbf{x}\|_2^2 + \|\mathbf{b} - A\mathbf{x} - P_{\mathcal{C}(A)}(\mathbf{b}) + A\mathbf{x}\|_2^2 \\ &= \|P_{\mathcal{C}(A)}(\mathbf{b}) - A\mathbf{x}\|_2^2 + \|\mathbf{b} - P_{\mathcal{C}(A)}(\mathbf{b})\|_2^2. \end{aligned}$$

Above we can see that the first term $\|P_{\mathcal{C}(A)}(\mathbf{b}) - A\mathbf{x}\|_2^2$ can be minimized to 0 since $P_{\mathcal{C}(A)}(\mathbf{b}) \in \mathcal{C}(A)$, and also that $\|\mathbf{b} - P_{\mathcal{C}(A)}(\mathbf{b})\|_2^2$ does not depend on $\mathbf{x}$ at all. Hence,

$$\inf_{\mathbf{x} \in \mathbb{C}^n} f(\mathbf{x}) = \inf_{\mathbf{x} \in \mathbb{C}^n} \|\mathbf{b} - A\mathbf{x}\|_2^2 = \|\mathbf{b} - P_{\mathcal{C}(A)}(\mathbf{b})\|_2^2$$

with the minimum attained when $\mathbf{x}$ satisfies $A\mathbf{x} = P_{\mathcal{C}(A)}(\mathbf{b})$.

The end result of this analysis is that instead of solving $A\mathbf{x} = \mathbf{b}$ we might as well, whenever possible, instead solve $A\mathbf{x} = P_{\mathcal{C}(A)}(\mathbf{b})$ which we know always has a solution. Furthermore, we can use a QR decomposition of A to solve $A\mathbf{x} = P_{\mathcal{C}(A)}(\mathbf{b})$ efficiently. Let $A = QR$ be a QR decomposition of A . We have that

$$\begin{aligned} A\mathbf{x} = P_{\mathcal{C}(A)}(\mathbf{b}) &\iff QR\mathbf{x} = P_{\mathcal{C}(Q)}(\mathbf{b}) \iff QR\mathbf{x} = QQ^*\mathbf{b} \\ &\iff R\mathbf{x} = Q^*\mathbf{b}. \end{aligned}$$

Algorithm 4 ALGORITHM FOR (APPROXIMATELY) SOLVING $A\mathbf{x} = \mathbf{b}$

- 1: **Input:** $A \in \mathbb{C}^{m \times n}$, $\mathbf{b} \in \mathbb{C}^m$.
 - 2: **Output:** $\mathbf{x} \in \mathbb{C}^n$ minimizing $f(\mathbf{x}) = \|\mathbf{b} - A\mathbf{x}\|_2^2$.
 - 3: Compute a QR decomposition of A , so that $A = QR$.
 - 4: Solve $R\mathbf{x} = Q^*\mathbf{b}$ using back substitution.
 - 5: Return $\mathbf{x}$.
-

Furthermore, $R\mathbf{x} = Q^*\mathbf{b}$ can be solved efficiently by back substitution since R is upper triangular. Algorithm 4 outlines how to find the least-squares solution of $A\mathbf{x} = \mathbf{b}$ using a QR decomposition of A .

If $A \in \mathbb{C}^{m \times n}$ is small enough to fit into computer memory and/or accuracy is of principal concern, then one can safely default to directly computing a minimizer of $f(\mathbf{x}) = \|\mathbf{b} - A\mathbf{x}\|_2^2$ using Algorithm 4. If, on the other hand, an approximate least-squares solution suffices and/or A is too large or inaccessible to allow for easy use of Algorithm 4, then one can instead use optimization methods to minimize $f(\mathbf{x}) = \frac{1}{2}\|\mathbf{b} - A\mathbf{x}\|_2^2$ iteratively. In fact, this least-squares problem is important enough that we will discuss it several more times.

Finally, we note that when the rank of $A \in \mathbb{C}^{m \times n}$ is less than n there will be an entire $n - \text{rank}(A)$ dimensional affine subspace of equally good (approximate) solutions to $A\mathbf{x} = \mathbf{b}$. That is, $A(\mathbf{x} + \mathbf{n}) = A\mathbf{x} = \mathbf{b}$ will hold for all $\mathbf{n}$ in the “null space” of A . We will take this as initial motivation to review facts about the null space of a matrix next.

2.7 The Four Fundamental Linear Subspaces of a Matrix

Let $A \in \mathbb{C}^{m \times n}$. The four fundamental linear subspaces of A are:

1. **the column space of A** , $\mathcal{C}(A) = \text{span}\{A_{:,j} \mid j \in [n]\} \subset \mathbb{C}^m$,
2. **the null space of A , or kernel of A** , $\mathcal{N}(A) = \{\mathbf{x} \in \mathbb{C}^n \mid A\mathbf{x} = \mathbf{0}\} \subset \mathbb{C}^n$,
3. **the column space of A^* , or row space of $\overline{A}$** , $\mathcal{C}(A^*) = \text{span}\{A_{:,j}^* \mid j \in [m]\} \subset \mathbb{C}^n$,
and
4. **the null space of A^* , or kernel of A^*** , $\mathcal{N}(A^*) = \{\mathbf{y} \in \mathbb{C}^m \mid A^*\mathbf{y} = \mathbf{0}\} \subset \mathbb{C}^m$.

Exercise 2.7.1. Let $A \in \mathbb{C}^{m \times n}$. Show that the null space of A is a linear subspace of $\mathbb{C}^n$.

Reviewing facts about each of these linear subspaces, we recall that $r := \text{rank}(A)$ will always equal the dimension of $\mathcal{C}(A)$ by definition. In fact, it also turns out that $A^* \in \mathbb{C}^{n \times m}$ will also always have the same rank as $A \in \mathbb{C}^{m \times n}$.

Theorem 2.7.1. Let $A \in \mathbb{C}^{m \times n}$. It's always the case that $r = \text{rank}(A) = \text{rank}(A^*)$.

Proof. We will use a QR decomposition of A , $A = QR$, with $Q \in \mathbb{C}^{m \times r}$ and $R \in \mathbb{C}^{r \times n}$. Recall that $\text{rank}(Q) = \text{rank}(A) = r$. Additionally, $\text{rank}(A^*) = \dim(\mathcal{C}(A^*)) = \dim(\mathcal{C}(R^*Q^*))$. Since the columns of Q are orthonormal, so we can extend them to an orthonormal basis B of all of $\mathbb{C}^m$ which takes the form

$$B = \{Q_{:,0}, \dots, Q_{:,r-1}, \mathbf{q}_r, \dots, \mathbf{q}_{m-1}\} \subset \mathbb{C}^m.$$

Now observe that

$$\begin{aligned} \mathcal{C}(A^*) &= \{A^* \mathbf{y} \mid \mathbf{y} \in \mathbb{C}^m\} = \{R^* Q^* \mathbf{y} \mid \mathbf{y} \in \mathbb{C}^m\} \\ &= \left\{ R^* Q^* \left(\sum_{j \in [r]} \alpha_j Q_{:,j} + \sum_{\ell=r}^{m-1} \beta_\ell \mathbf{q}_\ell \right) \mid \{\alpha_j\}_{j \in [r]} \cup \{\beta_\ell\}_{\ell=r}^{m-1} \subset \mathbb{C} \right\} \\ &= \left\{ R^* \left(\sum_{j \in [r]} \alpha_j \mathbf{e}_j \right) \mid \{\alpha_j\}_{j \in [r]} \subset \mathbb{C} \right\} \\ &= \mathcal{C}(R^*). \end{aligned}$$

By the Exchange Lemma (Lemma 2.3.4) it follows that the rank of A^* , which is the size of any basis of $\mathcal{C}(A^*)$, must be less than the number of columns of R^* , which is $r = \text{rank}(A)$. Thus, $\text{rank}(A^*) \leq \text{rank}(A)$. Repeating the argument above with A replaced by A^* similarly shows that $\text{rank}(A) \leq \text{rank}(A^*)$. Combining these two results we learn that $\text{rank}(A) = \text{rank}(A^*)$ must hold. $\square$

Note that the proof of Theorem 2.7.1 above also shows that $\dim(\mathcal{C}(R^*)) = \dim(\mathcal{C}(A^*)) = \text{rank}(A) = r$. Thus, $\text{rank}(R^*)$ equals the number of columns of R^* . Similarly, R is also rank r which equals the number of rows of R . Generally, we will say that any $m \times n$ matrix whose rank matches $\min\{m, n\}$ is **full rank**. Hence, we can see from the argument above that the matrices Q and R resulting from the QR decomposition will always be full rank.

Lemma 2.7.2. *Let $A \in \mathbb{C}^{m \times n}$. Then $\mathcal{N}(A) = \mathcal{C}(A^*)^\perp \subset \mathbb{C}^n$ and $\mathcal{C}(A^*) = \mathcal{N}(A)^\perp \subset \mathbb{C}^m$.*

Proof. By Lemma 2.6.15 it suffices to show that $\mathcal{N}(A) = \mathcal{C}(A^*)^\perp$. Let $\mathbf{x} \in \mathcal{N}(A)$ and consider any given $\mathbf{z} \in \mathcal{C}(A^*)$. By definition, $A\mathbf{x} = \mathbf{0}$ and $\mathbf{z} = A^*\mathbf{y}$ for some $\mathbf{y} \in \mathbb{C}^m$. Hence, we can see that

$$\langle \mathbf{z}, \mathbf{x} \rangle = \langle A^* \mathbf{y}, \mathbf{x} \rangle = \langle \mathbf{y}, A\mathbf{x} \rangle = \langle \mathbf{y}, \mathbf{0} \rangle = 0.$$

Thus, $\mathcal{N}(A) \subset \mathcal{C}(A^*)^\perp$. To see that $\mathcal{C}(A^*)^\perp \subset \mathcal{N}(A)$ also holds, we note that if $\langle \mathbf{z}, \mathbf{x} \rangle = 0$ for all $\mathbf{z} \in \mathcal{C}(A^*)$, then $\langle A^* \mathbf{y}, \mathbf{x} \rangle = 0$ for all $\mathbf{y} \in \mathbb{C}^m$. This in turn implies that $\langle \mathbf{y}, A\mathbf{x} \rangle = 0$ for all $\mathbf{y} \in \mathbb{C}^m$ which means that $\langle A\mathbf{x}, A\mathbf{x} \rangle = \|A\mathbf{x}\|_2^2 = 0$. $\square$

Algorithm 5 ALGORITHM FOR COMPUTING AN ORTHONORMAL BASIS OF $\mathcal{N}(A)$

- 1: **Input:** A rank r matrix $A \in \mathbb{C}^{m \times n}$.
 - 2: **Output:** An orthonormal basis of A 's null space $\mathcal{N}(A) \subset \mathbb{C}^n$.
 - 3: Compute a QR decomposition of A^* , so that $A^* = QR$. Note that $\mathcal{C}(A^*) = \mathcal{C}(Q)$.
 - 4: Complete $B = \{Q_{:,j}\}_{j \in [r]}$ to be an orthonormal basis $\tilde{B} = B \cup S$ of all of $\mathbb{C}^n$. The set S will be an orthonormal basis of $\mathcal{C}(Q)^\perp = \mathcal{C}(A^*)^\perp = \mathcal{N}(A)$.
 - 5: Return S .
-

Using Lemma 2.7.2 we can further see that the dimension of $\mathcal{N}(A)$ is $n - r$ since $\mathbb{C}^n = \mathcal{C}(A^*) + \mathcal{C}(A^*)^\perp = \mathcal{C}(A^*) + \mathcal{N}(A)$. Hence, an orthonormal basis of $\mathcal{C}(A^*)$, which will consist of r vectors, can be completed into a larger orthonormal basis of all of $\mathbb{C}^n$ by adding $n - r$ new orthonormal vectors that span $\mathcal{N}(A)$. By encoding this argument as an algorithm we can also create a method for computing an orthonormal basis of the null space of any matrix $A \in \mathbb{C}^{m \times n}$. We can begin by computing an orthonormal basis B of $\mathcal{C}(A^*)$ by, e.g., running Algorithm 1 on the columns of A^* . We can then complete B to an orthonormal basis $\tilde{B} = B \cup S$ of all of $\mathbb{C}^n$ using Algorithm 3. The set S will be an orthonormal basis of $\mathcal{N}(A)$ of size $n - \text{rank}(A)$. See Algorithm 5 for pseudocode.

Exercise 2.7.2. Let $A \in \mathbb{C}^{m \times n}$ have rank r . Show that $\mathcal{N}(A^*) = \mathcal{C}(A)^\perp \subset \mathbb{C}^m$ and $\mathcal{C}(A) = \mathcal{N}(A^*)^\perp \subset \mathbb{C}^m$. Then, argue that $\dim(\mathcal{N}(A^*)) = m - r$.

Exercise 2.7.3. Show that $A \in \mathbb{C}^{n \times n}$ is full rank if and only if $\mathcal{N}(A) = \{\mathbf{0}\}$. Such square matrices are also said to be **invertible**.

The next lemma will be important soon in Section 3.1. We will prove it here since it depends crucially on our recent revelations regarding null spaces.

Lemma 2.7.3. Let $A \in \mathbb{C}^{m \times n}$. Then $\mathcal{C}(A^*A) = \mathcal{C}(A^*)$.

Proof. First we note that

$$\mathcal{C}(A^*A) = \{A^*A\mathbf{y} \mid \mathbf{y} \in \mathbb{C}^n\} = \{A^*\mathbf{z} \mid \mathbf{z} \in \mathcal{C}(A)\} = \{A^*P_{\mathcal{C}(A)}\mathbf{x} \mid \mathbf{x} \in \mathbb{C}^m\}.$$

Now we can re-express $\mathcal{C}(A^*)$ using that $\mathcal{C}(A)^\perp = \mathcal{N}(A^*) \subset \mathbb{C}^m$ to see that

$$\begin{aligned} \mathcal{C}(A^*) &= \left\{ A^* \left(P_{\mathcal{C}(A)}\mathbf{x} + P_{\mathcal{C}(A)^\perp}\mathbf{x} \right) \mid \mathbf{x} \in \mathbb{C}^m \right\} = \left\{ A^*P_{\mathcal{C}(A)}\mathbf{x} + A^*P_{\mathcal{N}(A^*)}\mathbf{x} \mid \mathbf{x} \in \mathbb{C}^m \right\} \\ &= \{A^*P_{\mathcal{C}(A)}\mathbf{x} \mid \mathbf{x} \in \mathbb{C}^m\} = \mathcal{C}(A^*A). \end{aligned}$$

□

Exercise 2.7.4. Let $A \in \mathbb{C}^{m \times n}$. Prove that $\mathcal{C}(AA^*) = \mathcal{C}(A)$.

As a consequence of the above, we can see that

$$\text{rank}(A^*A) = \text{rank}(A^*) = \text{rank}(A) = \text{rank}(AA^*).$$

2.8 The Spectral Theorem for Hermitian Matrices

We will now briefly concentrate on a very special type of square matrix which will serve as our doorway to the almighty singular value decomposition in Section 3.1.

Definition 2.8.1. A matrix $A \in \mathbb{C}^{n \times n}$ is called **Hermitian** if $A = A^*$.

Exercise 2.8.1. Let $A \in \mathbb{C}^{m \times n}$. Show that both $AA^* \in \mathbb{C}^{m \times m}$ and $A^*A \in \mathbb{C}^{n \times n}$ are Hermitian.

Exercise 2.8.2. Let $A \in \mathbb{C}^{n \times n}$ be Hermitian. Show that all entries on A 's diagonal are real numbers.

Exercise 2.8.3. Let $A \in \mathbb{C}^{n \times n}$ be Hermitian. Show that $\mathcal{N}(A) = \mathcal{C}(A)^\perp \subset \mathbb{C}^n$ and $\mathcal{C}(A) = \mathcal{N}(A)^\perp \subset \mathbb{C}^n$.

The eigenvalues and eigenvectors of Hermitian matrices have a lot of special properties that we will need later. We will discuss these properties next.

Definition 2.8.2. An **eigenvalue-eigenvector pair**, or **eigenpair**, of a matrix $A \in \mathbb{C}^{n \times n}$ is a pair $(\lambda, \mathbf{v}) \in \mathbb{C} \times \mathbb{C}^n \setminus \{\mathbf{0}\}$ such that $\mathbf{v} \neq \mathbf{0}$ satisfies $A\mathbf{v} = \lambda\mathbf{v}$.

Lemma 2.8.3. Let $A \in \mathbb{C}^{n \times n}$ be Hermitian. Then all eigenvalues of A are real numbers.

Proof. Let $(\lambda, \mathbf{v})$ be an eigenpair of A . If $\lambda = 0 \in \mathbb{R}$ we are done. Thus, suppose that $\lambda \neq 0$. Then we have that

$$\begin{aligned} \|\mathbf{v}\|_2^2 &= \langle \mathbf{v}, \mathbf{v} \rangle = \left\langle \frac{1}{\lambda} A\mathbf{v}, \mathbf{v} \right\rangle = (1/\bar{\lambda}) \langle \mathbf{v}, A^*\mathbf{v} \rangle = (1/\bar{\lambda}) \langle \mathbf{v}, A\mathbf{v} \rangle = (1/\bar{\lambda}) \langle \mathbf{v}, \lambda\mathbf{v} \rangle \\ &= (\lambda/\bar{\lambda}) \|\mathbf{v}\|_2^2. \end{aligned}$$

Since $\mathbf{v}$ is nonzero we know $\|\mathbf{v}\|_2 \neq 0$ so that $\lambda = \bar{\lambda}$ must hold. Hence, $\lambda \in \mathbb{R}$. $\square$

Note that every fixed eigenvalue $\lambda \in \mathbb{C}$ of $A \in \mathbb{C}^{n \times n}$ has an infinite number of associated eigenvectors. In fact, one can see that the set of all eigenvectors corresponding to λ (after adding in the zero vector) is closed under both addition and scalar multiplication so that it forms a linear subspace of $\mathbb{C}^n$. And, this subspace of $\mathbb{C}^n$ is exactly equal to the nullspace of $A - \lambda I_n \in \mathbb{C}^{n \times n}$,

$$\mathcal{N}(A - \lambda I_n) = \{\mathbf{v} \in \mathbb{C}^n \mid (\lambda, \mathbf{v}) \text{ is an eigenpair of } A\} \cup \{\mathbf{0}\}.$$

For this reason we will refer to $\mathcal{N}(A - \lambda I_n)$ as the **eigenspace associated with λ** . Furthermore, we will let an orthonormal basis of this linear subspace be denoted by $B_\lambda \subset \mathbb{C}^n$ for each eigenvalue λ .

Example 2.8.4. Let $U \in \mathbb{C}^{n \times n}$ be unitary. Then $UU^* = I_n$ so that UU^* has only one nontrivial eigenspace $\mathcal{N}(UU^* - I_n) = \mathbb{C}^n$ associated with its single eigenvalue $\lambda = 1$. Furthermore, its orthonormal basis B_1 will be an orthonormal basis of all of $\mathbb{C}^n$.

Exercise 2.8.4. Prove that every matrix $A \in \mathbb{C}^{n \times n}$ with $\text{rank} < n$ has at least one nontrivial eigenspace. What is it?

Exercise 2.8.5. Prove that $A \in \mathbb{C}^{n \times n}$ has exactly one eigenvalue if and only if it's a scalar multiple of the identity matrix I_n .

Another important property of Hermitian matrices is that all of their distinct eigenspaces must be orthogonal to one another. This fact is proven in the next lemma.

Lemma 2.8.5. Let $(\lambda, \mathbf{v})$ and $(\mu, \mathbf{u})$ be two eigenpairs of a Hermitian matrix $A \in \mathbb{C}^{n \times n}$ with $\lambda \neq \mu$. Then $\langle \mathbf{v}, \mathbf{u} \rangle = 0$. As a consequence, $\mathcal{N}(A - \lambda I_n) \perp \mathcal{N}(A - \mu I_n)$.

Proof. Since $\lambda, \mu \in \mathbb{R}$ are distinct, at least one is nonzero. Without loss of generality let λ be nonzero. Then, $\mu \neq \lambda \implies \frac{\mu}{\lambda} \neq 1 \implies 1 - \frac{\mu}{\lambda} \neq 0$. Since $\lambda \in \mathbb{R} \setminus \{0\}$ we can also see that

$$\langle \mathbf{v}, \mathbf{u} \rangle = \frac{1}{\lambda} \langle A\mathbf{v}, \mathbf{u} \rangle = \frac{1}{\lambda} \langle \mathbf{v}, A^*\mathbf{u} \rangle = \frac{1}{\lambda} \langle \mathbf{v}, A\mathbf{u} \rangle = \frac{1}{\lambda} \langle \mathbf{v}, \mu\mathbf{u} \rangle = \frac{\mu}{\lambda} \langle \mathbf{v}, \mathbf{u} \rangle.$$

Thus,

$$\left(1 - \frac{\mu}{\lambda}\right) \langle \mathbf{v}, \mathbf{u} \rangle = 0.$$

Hence, it must be the case that $\langle \mathbf{v}, \mathbf{u} \rangle = 0$ since $1 - \frac{\mu}{\lambda} \neq 0$. $\square$

Let $A \in \mathbb{C}^{n \times n}$ be a Hermitian matrix whose eigenvalues are $\lambda_0, \dots, \lambda_{m-1} \in \mathbb{R}$. Lemma 2.8.5 implies that the eigenspaces of A will all be orthogonal to one another. As a result, if we let B_{λ_j} be an orthonormal basis for each eigenspace $\mathcal{N}(A - \lambda_j I_n)$ of A , then

$$B := \bigcup_{j \in [m]} B_{\lambda_j} \subset \mathbb{C}^n \tag{2.15}$$

will be an orthonormal set. In fact, it will also always be the case that B is an orthonormal basis for all of $\mathbb{C}^n$ (we will not prove this here – see, e.g., [19, Chapter 2] or [15, Chapter 14] for corroborating evidence).

Fact 2.8.6. If $A \in \mathbb{C}^{n \times n}$ is Hermitian then there exists an orthonormal basis of all of $\mathbb{C}^n$ consisting of eigenvectors of A . In particular, the set B in (2.15) will be an orthonormal basis of $\mathbb{C}^n$.

Let $A \in \mathbb{C}^{n \times n}$ be Hermitian and $B = \{\mathbf{b}_j\}_{j \in [n]} \subset \mathbb{C}^n$ be an orthonormal basis of $\mathbb{C}^n$ consisting of eigenvectors of A as defined in (2.15). Form a unitary matrix $U \in \mathbb{C}^{n \times n}$ that contains the elements of B as its columns (i.e., so that $U_{:,j} = \mathbf{b}_j$ for all $j \in [n]$). By the definition of eigenpairs we can see that

$$\begin{aligned} AU &= A \begin{pmatrix} | & | & & | \\ \mathbf{b}_0 & \mathbf{b}_1 & \cdots & \mathbf{b}_{n-1} \\ | & | & & | \end{pmatrix} = \begin{pmatrix} | & | & & | \\ A\mathbf{b}_0 & A\mathbf{b}_1 & \cdots & A\mathbf{b}_{n-1} \\ | & | & & | \end{pmatrix} \\ &= \begin{pmatrix} | & | & & | \\ \lambda_0 \mathbf{b}_0 & \lambda_1 \mathbf{b}_1 & \cdots & \lambda_{n-1} \mathbf{b}_{n-1} \\ | & | & & | \end{pmatrix} = \begin{pmatrix} | & | & & | \\ \mathbf{b}_0 & \mathbf{b}_1 & \cdots & \mathbf{b}_{n-1} \\ | & | & & | \end{pmatrix} \text{diag}(\lambda_0, \dots, \lambda_{n-1}) \\ &= U \text{diag}(\lambda_0, \dots, \lambda_{n-1}). \end{aligned}$$

where $\lambda_j \in \mathbb{R}$ refers to the eigenvalue corresponding to $\mathbf{b}_j \in B$. Finally, recalling that U is unitary we can see that multiplying both sides of the equation just above on the right by U^* yields

$$A = AUU^* = U \text{diag}(\lambda_0, \dots, \lambda_{n-1})U^*.$$

This computation together with Lemma 2.8.3, Lemma 2.8.5, and Fact 2.8.6 prove the following theorem (see also, e.g., Theorem 2.5.6 in [19]).

Theorem 2.8.7 (The Full Spectral Decomposition of a Hermitian Matrix). *Let $A \in \mathbb{C}^{n \times n}$ be Hermitian. Then there exist $\lambda_0, \dots, \lambda_{n-1} \in \mathbb{R}$ and a unitary matrix $U \in \mathbb{C}^{n \times n}$ such that*

$$A = U \text{diag}(\lambda_0, \dots, \lambda_{n-1})U^*.$$

Exercise 2.8.6. *Let $A \in \mathbb{C}^{m \times n}$. Show that all the eigenvalues of the Hermitian matrices $A^*A \in \mathbb{C}^{n \times n}$ and $AA^* \in \mathbb{C}^{m \times m}$ are nonnegative real numbers.*

Theorem 2.8.7 is great, but we'd also like a version that allows us to store low-rank matrices in a compressed form. Let's think about how to develop such a variant – it'll also be good practice for Section 3.1. Recall from our definition of atomic permutation matrices $P(j, k) \in \mathbb{C}^{n \times n}$ (see Example 2.5.3 and the surrounding text) that $P(j, k)$ swaps the j^{th} and k^{th} rows of $A \in \mathbb{C}^{n \times n}$ when multiplied against it on the left, and swaps the j^{th} and k^{th} columns of $A \in \mathbb{C}^{n \times n}$ when multiplied against it on the right. Furthermore, every atomic permutation matrix $P(j, k) \in \mathbb{C}^{n \times n}$ is unitary, as are all products of atomic permutation matrices (see Exercise 2.5.7 and Theorem 2.6.14). Having refamiliarised ourselves with atomic permutation matrices, note that if $P(j, k)$ is applied to both sides of a diagonal matrix simultaneously it will swap its j^{th} and k^{th} diagonal entries. That is,

$$\begin{aligned} P(j, k) \text{diag}(\lambda_0, \dots, \lambda_{j-1}, \lambda_j, \lambda_{j+1}, \dots, \lambda_{k-1}, \lambda_k, \lambda_{k+1}, \dots, \lambda_{n-1}) P(j, k) \\ = \text{diag}(\lambda_0, \dots, \lambda_{j-1}, \lambda_k, \lambda_{j+1}, \dots, \lambda_{k-1}, \lambda_j, \lambda_{k+1}, \dots, \lambda_{n-1}). \end{aligned}$$

Example 2.8.8. Let $P(0, 2) = \begin{pmatrix} 0 & 0 & 1 \\ 0 & 1 & 0 \\ 1 & 0 & 0 \end{pmatrix} \in \mathbb{C}^{3 \times 3}$. We can see that

$$\begin{aligned} P(0, 2) \operatorname{diag}(a, b, c) P(0, 2) &= \begin{pmatrix} 0 & 0 & 1 \\ 0 & 1 & 0 \\ 1 & 0 & 0 \end{pmatrix} \begin{pmatrix} a & 0 & 0 \\ 0 & b & 0 \\ 0 & 0 & c \end{pmatrix} \begin{pmatrix} 0 & 0 & 1 \\ 0 & 1 & 0 \\ 1 & 0 & 0 \end{pmatrix} \\ &= \begin{pmatrix} 0 & 0 & c \\ 0 & b & 0 \\ a & 0 & 0 \end{pmatrix} \begin{pmatrix} 0 & 0 & 1 \\ 0 & 1 & 0 \\ 1 & 0 & 0 \end{pmatrix} = \begin{pmatrix} c & 0 & 0 \\ 0 & b & 0 \\ 0 & 0 & a \end{pmatrix} \\ &= \operatorname{diag}(c, b, a). \end{aligned}$$

Using these facts about atomic permutation matrices together with Theorem 2.8.7, we can see that there exists a permutation matrix $P = \prod_{\ell \in [q]} P(j_\ell, k_\ell)$ consisting of a product of $q \in \mathbb{N}$ atomic permutation matrices such that

$$\begin{aligned} A &= U \operatorname{diag}(\lambda_0, \dots, \lambda_{n-1}) U^* = U(PP^*) \operatorname{diag}(\lambda_0, \dots, \lambda_{n-1}) (PP^*)U^* \\ &= (UP)(P \operatorname{diag}(\lambda_0, \dots, \lambda_{n-1}) P)(P^*U^*) \\ &= (UP) \operatorname{diag}(\tilde{\lambda}_0, \dots, \tilde{\lambda}_{n-1}) (UP)^*, \end{aligned}$$

where $\tilde{\lambda}_0, \dots, \tilde{\lambda}_{n-1}$ is a permutation of $\lambda_0, \dots, \lambda_{n-1} \in \mathbb{R}$ satisfying

$$|\tilde{\lambda}_0| \geq |\tilde{\lambda}_1| \geq \dots \geq |\tilde{\lambda}_{n-1}|.$$

Let $\tilde{U} = UP$, and note that $\tilde{U}$ is still a unitary matrix (see, e.g., Exercise 2.6.10).

Continuing, now consider the case where A is not full rank so that $|\tilde{\lambda}_{n-1}| = 0$. In this case we can further compress our spectral decomposition of A using block matrix representations. To begin, let's re-express $\tilde{U}$ in block form by

$$\tilde{U} = \begin{pmatrix} V & \tilde{U}_{:,n-1} \end{pmatrix} \in \mathbb{C}^{n \times n}$$

where $V \in \mathbb{C}^{n \times (n-1)}$ is the orthonormal matrix formed by the first $n-1$ columns of $\tilde{U}$. Further, let's represent $\operatorname{diag}(\tilde{\lambda}_0, \dots, \tilde{\lambda}_{n-1})$ in block form as well by

$$\operatorname{diag}(\tilde{\lambda}_0, \dots, \tilde{\lambda}_{n-1}) = \begin{pmatrix} D & \mathbf{0} \\ \mathbf{0}^* & 0 \end{pmatrix} \in \mathbb{R}^{n \times n}$$

where $D = \operatorname{diag}(\tilde{\lambda}_0, \dots, \tilde{\lambda}_{n-2}) \in \mathbb{R}^{(n-1) \times (n-1)}$ and $\mathbf{0}$ is a suitably tall vector of zeroes. Then, we have that

$$A = \begin{pmatrix} V & \tilde{U}_{:,n-1} \end{pmatrix} \begin{pmatrix} D & \mathbf{0} \\ \mathbf{0}^* & 0 \end{pmatrix} \begin{pmatrix} V^* \\ (\tilde{U}_{:,n-1})^* \end{pmatrix} = \begin{pmatrix} V & \tilde{U}_{:,n-1} \end{pmatrix} \begin{pmatrix} DV^* \\ \mathbf{0}^* \end{pmatrix} = VDV^*.$$

Note that $V \in \mathbb{C}^{n \times (n-1)}$ above is no longer unitary since it isn't square, but it is still an orthonormal matrix, and D is still a diagonal matrix of real numbers. And, of course, we can repeat this process again if $\tilde{\lambda}_{n-2} = 0$ too, and so on, until we run out of 0 eigenvalues. When will that happen? Well, denote the rank of our Hermitian $A \in \mathbb{C}^{n \times n}$ by $r < n$. The eigenspace associated with the 0 eigenvalue of A is exactly the null space of A so that the orthonormal set B_0 in (2.15) will have $|B_0| = \dim(\mathcal{N}(A)) = n - r$. Hence, we carry out this process $n - r$ total times for all of $\tilde{\lambda}_{n-1} = \tilde{\lambda}_{n-2} = \dots = \tilde{\lambda}_r = 0$. Formalizing this discussion gives us the following result.

Corollary 2.8.9 (The Compact Spectral Decomposition of a Hermitian Matrix). *Let $A \in \mathbb{C}^{n \times n}$ be Hermitian with rank $r < n$. Then, there exists an orthonormal matrix $U \in \mathbb{C}^{n \times r}$, and $\lambda_0, \dots, \lambda_{r-1} \in \mathbb{R}$ satisfying*

$$|\lambda_0| \geq |\lambda_1| \geq \dots |\lambda_{r-1}| > 0,$$

such that

$$A = U \operatorname{diag}(\lambda_0, \dots, \lambda_{r-1}) U^*.$$

We end our discussion of the spectral theorem here by noting that Theorem 2.8.7 and Corollary 2.8.9 are really fantastic! They decompose every Hermitian matrix into a product of extremely well behaved (e.g., easily invertible in the full rank case) matrices. Given how much we have used the QR decomposition in this chapter, we hope that the reader can now instinctively anticipate the potential utility of yet another decomposition that in many ways is even nicer (let's be honest – the R in the QR decomposition is just not as nice as the diagonal/unitary combination Theorem 2.8.7 effectively replaces it with). Theorem 2.8.7 and Corollary 2.8.9 do have one major flaw, however. They only apply to one very special type of square matrix! In the next chapter we will remove this flaw by developing a generalization of these Hermitian matrix decompositions that applies to all (even rectangular) matrices.

2.9 Positive (Semi)Definite Matrices

We will now briefly discuss a special class of Hermitian (or symmetric, if real-valued) matrices that are important in optimization, statistics, and applied mathematics more generally.

Definition 2.9.1 (Positive and Negative Definite Matrices). *A Hermitian matrix $A \in \mathbb{C}^{n \times n}$ is **positive definite** if $\mathbf{x}^* A \mathbf{x} > 0$ holds $\forall \mathbf{x} \in \mathbb{C}^n \setminus \{\mathbf{0}\}$. A Hermitian matrix $A \in \mathbb{C}^{n \times n}$ is **negative definite** if $\mathbf{x}^* A \mathbf{x} < 0$ holds $\forall \mathbf{x} \in \mathbb{C}^n \setminus \{\mathbf{0}\}$.*

Definition 2.9.2 (Positive SemiDefinite (PSD) and Negative Semidefinite Matrices). *A Hermitian matrix $A \in \mathbb{C}^{n \times n}$ is **positive semidefinite** if $\mathbf{x}^* A \mathbf{x} \geq 0$ holds $\forall \mathbf{x} \in \mathbb{C}^n \setminus \{\mathbf{0}\}$. A Hermitian matrix $A \in \mathbb{C}^{n \times n}$ is **negative semidefinite** if $\mathbf{x}^* A \mathbf{x} \leq 0$ holds $\forall \mathbf{x} \in \mathbb{C}^n \setminus \{\mathbf{0}\}$.*

We encourage the serious reader to do the following exercises in order to help themselves absorb the definitions above.

Exercise 2.9.1. Let $A \in \mathbb{C}^{n \times n}$. Show that A is negative definite if and only if $-A$ is positive definite.

Exercise 2.9.2. Let $A \in \mathbb{C}^{n \times n}$. Prove that both AA^* and A^*A are PSD.

Exercise 2.9.3. Let $r \in \mathbb{R}^+$ and $A \in \mathbb{C}^{n \times n}$ be PSD. Show that rA is also PSD.

Exercise 2.9.4. Let $A, B \in \mathbb{C}^{n \times n}$ both be PSD. Show that $A + B$ is also PSD.

The following lemma will be useful below.

Lemma 2.9.3. Suppose that a Hermitian matrix $A \in \mathbb{C}^{n \times n}$ only has nonzero eigenvalues. Then, A is full rank.

Proof. We will prove the contrapositive of the desired result after recalling that A is full rank if and only if $\dim(\mathcal{C}(A)) = n$, which in turn will hold if and only if all n columns of A are linearly independent. Now suppose that A is not full rank. Then $\exists \alpha \in \mathbb{C}^n \setminus \{\mathbf{0}\}$ such that $A\alpha = \sum_{j \in [n]} \alpha_j A_{:,j} = \mathbf{0}$. Hence, 0 is an eigenvalue of A with eigenvector α . $\square$

The discerning reader will see that Lemma 2.9.3 immediately implies that all positive (and negative) definite matrices are full rank and, therefore, invertible (why?). In fact we can say a bit more.

Theorem 2.9.4. Let $A \in \mathbb{C}^{n \times n}$ be Hermitian. The following are equivalent:

1. A is positive definite.
2. All the eigenvalues of A are positive real numbers.
3. There exists an invertible matrix $B \in \mathbb{C}^{n \times n}$ such that $A = BB^*$.

Proof. (1) $\implies$ (2): Let λ be an eigenvalue of A with an associated eigenvector $\mathbf{v} \in \mathbb{C}^n \setminus \{\mathbf{0}\}$. We know that λ is a real number by Lemma 2.8.3. Furthermore, since A is positive definite we have that

$$0 < \mathbf{v}^* A \mathbf{v} = \mathbf{v}^* (\lambda \mathbf{v}) = \lambda \underbrace{\|\mathbf{v}\|_2^2}_{> 0} \implies \lambda > 0.$$

(2) $\implies$ (3): By Theorem 2.8.7 $A = U \text{diag}(\lambda_0, \dots, \lambda_{n-1}) U^*$ where $U \in \mathbb{C}^{n \times n}$ is unitary and the λ_j are the positive real eigenvalues (by assumption) of A . Because the λ_j are nonnegative we may define $B := U \text{diag}(\sqrt{\lambda_0}, \dots, \sqrt{\lambda_{n-1}}) \in \mathbb{C}^{n \times n}$. Note then that

$$A = U \text{diag}(\sqrt{\lambda_0}, \dots, \sqrt{\lambda_{n-1}}) \text{diag}(\sqrt{\lambda_0}, \dots, \sqrt{\lambda_{n-1}}) U^* = BB^*.$$

The matrix B will be full rank (and therefore invertible) by Exercise 2.7.4 and Lemma 2.9.3.

(3) $\implies$ (1): Let $\mathbf{x} \in \mathbb{C}^n \setminus \{\mathbf{0}\}$. Since $B \in \mathbb{C}^{n \times n}$ is invertible/full rank, $B^* \mathbf{x} \neq \mathbf{0}$. Hence,

$$\mathbf{x}^* A \mathbf{x} = \mathbf{x}^* B B^* \mathbf{x} = (B^* \mathbf{x})^* (B^* \mathbf{x}) = \|B^* \mathbf{x}\|_2^2 > 0.$$

□

Similar characterizations exist of PSD matrices, negative definite matrices, and negative semidefinite matrices. We leave them as exercises.

Exercise 2.9.5. Let $A \in \mathbb{C}^{n \times n}$ be Hermitian. Show that the following are equivalent:

1. A is PSD.
2. All the eigenvalues of A are nonnegative real numbers.
3. There exists a matrix $B \in \mathbb{C}^{n \times n}$ with $\text{rank}(B) = \text{rank}(A)$ such that $A = B B^*$.

Exercise 2.9.6. Prove a variant of Theorem 2.9.4 for negative definite matrices.

Exercise 2.9.7. Prove a variant of Theorem 2.9.4 for negative semidefinite matrices.

2.9.1 The Cholesky Decomposition

Let $B^* \in \mathbb{C}^{n \times n}$ be as in the third part of Theorem 2.9.4 and consider its QR decomposition $B^* = QR$, where $Q \in \mathbb{C}^{n \times n}$ is unitary and $R \in \mathbb{C}^{n \times n}$ is upper triangular. Note that the entries on the diagonal of R must be nonzero since B^* is full rank, but they can be complex. To help make them real we can let $D \in \mathbb{C}^{n \times n}$ be the unitary diagonal matrix

$$D = \text{diag} \left(\left\{ e^{-i \arg(R_{jj})} \right\}_{j \in [n]} \right)$$

and then write

$$B^* = QR = (QD^*)(DR) = \tilde{Q}\tilde{R},$$

where $\tilde{Q} := QD^*$ is unitary by Exercise 2.6.10, and $\tilde{R} := DR$ is upper triangular with positive real entries on its diagonal ($\tilde{R}_{j,j} = |R_{j,j}| > 0$ for all $j \in [n]$).

Returning to the third part of Theorem 2.9.4 we can see that any positive definite matrix A will have

$$A = B B^* = (B^*)^* B^* = (\tilde{Q}\tilde{R})^* \tilde{Q}\tilde{R} = \tilde{R}^* \tilde{Q}^* \tilde{Q} \tilde{R} = \tilde{R}^* \tilde{R},$$

since $\tilde{Q}$ is unitary. Note that since $\tilde{R}$ is upper triangular, $\tilde{R}^*$ is lower triangular. We have just proven the following theorem.

Theorem 2.9.5 (Cholesky Decomposition). Let $A \in \mathbb{C}^{n \times n}$ be positive definite. Then $\exists$ a lower triangular matrix $L \in \mathbb{C}^{n \times n}$ with positive real entries on its diagonal so that $A = LL^*$.

Exercise 2.9.8. Let $A \in \mathbb{C}^{n \times n}$. Prove that $\exists$ a lower triangular matrix $L \in \mathbb{C}^{n \times n}$ with positive real entries on its diagonal so that $AA^* = LL^*$

2.10 A Review of the Trace and Determinant Functions

In this section we will rapidly review some important properties of both the trace and the determinant of a matrix. Unlike more introductory texts (see, e.g., [13, 29]), we will give both of these functions a rather cursory treatment focused primarily on their computational aspects. As a result, we encourage any reader with more than a casual interest in either function to consult other sources if the presentation here is found to be lacking.

2.10.1 The Trace of a Matrix

We begin with the more computationally useful of the two functions.

Definition 2.10.1 (Trace). *The trace function, $\text{Trace} : \mathbb{C}^{m \times n} \rightarrow \mathbb{C}$, is defined by*

$$\text{Trace}(A) := \sum_{j \in [\min\{m,n\}]} A_{j,j} = \sum_{j=0}^{\min\{m,n\}-1} A_{j,j}$$

for all $A \in \mathbb{C}^{m \times n}$.

Exercise 2.10.1. *Prove that the trace function is linear. That is, for all $A, B \in \mathbb{C}^{m \times n}$ and $\lambda \in \mathbb{C}$ show that both*

1. $\text{Trace}(A + B) = \text{Trace}(A) + \text{Trace}(B)$, and
2. $\text{Trace}(\lambda A) = \lambda \text{Trace}(A)$

hold.

Exercise 2.10.2. *Let $A \in \mathbb{C}^{m \times n}$. Prove that $\text{Trace}(A^*) = \overline{\text{Trace}(A)}$.*

The following property of trace functions is particularly useful.

Lemma 2.10.2. *Let $A \in \mathbb{C}^{m \times n}$ and $B \in \mathbb{C}^{n \times m}$. Then $\text{Trace}(AB) = \text{Trace}(BA)$.*

Proof. We compute

$$\begin{aligned} \text{Trace}(AB) &= \sum_{j \in [m]} (AB)_{j,j} = \sum_{j \in [m]} \sum_{k \in [n]} A_{j,k} B_{k,j} = \sum_{k \in [n]} \sum_{j \in [m]} B_{k,j} A_{j,k} \\ &= \sum_{k \in [n]} (BA)_{k,k} = \text{Trace}(BA). \end{aligned}$$

□

The trace function has an important relationship to the Frobenius norm of a given matrix $A \in \mathbb{C}^{m \times n}$, allowing one to create an associated inner product. The following lemma explicitly demonstrates the relationship between the Frobenius norm and the trace function.

Lemma 2.10.3. *Let $A \in \mathbb{C}^{m \times n}$. Then $\|A\|_F^2 = \text{Trace}(A^*A) = \text{Trace}(AA^*)$.*

Proof. We have that

$$\begin{aligned} \|A\|_F^2 &= \sum_{j \in [m]} \sum_{k \in [n]} |A_{j,k}|^2 = \sum_{k \in [n]} \sum_{j \in [m]} A_{j,k} \overline{A_{j,k}} = \sum_{k \in [n]} \sum_{j \in [m]} (A^*)_{k,j} A_{j,k} \\ &= \sum_{k \in [n]} (A^*A)_{k,k} = \text{Trace}(A^*A). \end{aligned}$$

The second equality now follows from Lemma 2.10.2. $\square$

We may now go ahead and define the Frobenius inner product of two matrices $A, B \in \mathbb{C}^{m \times n}$ to be

$$\langle A, B \rangle_F := \text{Trace}(A^*B).$$

To see that this is indeed an inner product, recall that one can always reshape an $m \times n$ matrix into a vector of length mn using, e.g., the vectorization operator $\text{vec} : \mathbb{C}^{m \times n} \rightarrow \mathbb{C}^{mn}$ defined for all $\ell \in [mn]$ by $(\text{vec}(A))_\ell = A_{\ell \bmod m, \lceil \ell/m \rceil}$, where “ $\ell \bmod m$ ” is defined for all $\ell \in \mathbb{Z}$ and $m \in \mathbb{N}$ to be the single element contained in the set $\{\ell + km \mid k \in \mathbb{Z}\} \cap [m]$ (or, equivalently, to be the unique value $r \in [m] = \{0, 1, \dots, m-1\}$ such that $\exists k \in \mathbb{Z}$ satisfying $\ell = r + km$). We then have that

$$\begin{aligned} \langle A, B \rangle_F &= \text{Trace}(A^*B) = \sum_{j \in [n]} \sum_{k \in [m]} (A^*)_{j,k} B_{k,j} = \sum_{j \in [n]} \sum_{k \in [m]} \overline{A_{k,j}} B_{k,j} \\ &= \langle \text{vec}(A), \text{vec}(B) \rangle. \end{aligned}$$

Hence, we can see that the Frobenius inner product is simply the usual Euclidean inner product in disguise.

One benefit of realizing that the Frobenius inner product is indeed an inner product is that we can now apply previously established facts about inner products to learn additional facts about trace functions. For example, we can now apply the Cauchy-Schwarz inequality to see that

$$|\text{Trace}(A^*B)| = |\langle \text{vec}(A), \text{vec}(B) \rangle| \leq \|\text{vec}(A)\|_2 \|\text{vec}(B)\|_2 = \|A\|_F \|B\|_F.$$

Lemma 2.10.4 (Cauchy-Schwarz Example for Trace). *Let $A, B \in \mathbb{C}^{m \times n}$. Then*

$$|\text{Trace}(A^*B)| \leq \|A\|_F \|B\|_F.$$

We hope that this useful inequality spurs the reader to more deeply appreciate the utility of both abstraction and re-expression. Old facts often return in disguise as new facts, and abstract theory is a channel through which many such elegant re-expressions travel.

2.10.2 The Determinant of a Matrix

The determinant is a function from square matrices into $\mathbb{C}$, denoted herein by $\det : \mathbb{C}^{n \times n} \rightarrow \mathbb{C}$. Determinants have many interesting properties and are fundamentally important to both multivariable calculus and differential geometry as well as their relations and descendants. That said, if you ever find yourself actually computing a determinant you're likely doing something wrong. It does, however, happen on occasion. . . .

For a review of the basic properties of determinants we refer the interested reader to, e.g., [13, Chapter 4] and/or [29, Chapter 5]. Herein we will focus on a small subset of their many properties which are of the most value for computing them efficiently.

Theorem 2.10.5 (See, e.g., Appendix C of [15]). *Let $A, B \in \mathbb{C}^{n \times n}$. The determinant function $\det : \mathbb{C}^{n \times n} \rightarrow \mathbb{C}$ satisfies the following properties:*

1. $\det(A^*) = \overline{\det(A)}$,
2. $\det(AB) = \det(A)\det(B)$, and
3. If A is lower (or upper) triangular, then $\det(A) = \prod_{j \in [n]} A_{j,j}$.

There are many other useful properties of determinants as well, some of which you may now derive from Theorem 2.10.5.

Exercise 2.10.3. *Use Theorem 2.10.5 to prove the following additional standard facts about determinants.*

1. $\det(I_n) = 1$.
2. If $A \in \mathbb{C}^{n \times n}$ is invertible, then $\det(A^{-1}) = \frac{1}{\det(A)}$.
3. If $U \in \mathbb{C}^{n \times n}$ is unitary, then $\det(U) \in \mathbb{C}$ has magnitude 1 so that $|\det(A)| = 1$.

Exercise 2.10.4. *Prove the following additional standard facts about determinants.*

1. If $A \in \mathbb{C}^{n \times n}$ is Hermitian, then $\det(A)$ is the product of A 's eigenvalues.
2. $\det(\lambda A) = \lambda^n \det(A)$ for all $A \in \mathbb{C}^{n \times n}$ and $\lambda \in \mathbb{C}$.

Exercise 2.10.5. *Let $A = QR$ be a QR-decomposition of $A \in \mathbb{C}^{n \times n}$. Show that $|A| := |\det(A)| = \prod_{j \in [n]} |R_{j,j}|$.*

The exercise just above implies a relatively efficient and stable numerical algorithm for computing the absolute determinant of a given matrix.

Chapter 3

Some More Advanced Topics in Linear Algebra

3.1 One Factorization to Rule Them All: The Singular Value Decomposition

The Singular Value Decomposition (SVD) is arguably the most useful fact of Linear Algebra, which is itself arguably the most useful and ubiquitous of mathematical subjects (with respect to computation in particular). The SVD’s utility in data analysis is underscored by the fact that it has been (re)discovered at least three times in different scientific communities [32]. In this section we will review the SVD of a given matrix $A \in \mathbb{C}^{m \times n}$. Many sections of the book hereafter will use the SVD repeatedly and often – it is well worth refreshing yourself here, and familiarizing yourself with our notation, before moving on.

Finally, to re-emphasize our statement about linear algebra over the real versus complex numbers from the beginning of Chapter 2, we remind the reader that **replacing the symbol “ $\mathbb{C}$ ” everywhere it appears in this section with an “ $\mathbb{R}$ ” will not affect the correctness of the results herein in any way whatsoever**. In fact, the only cosmetic (and frankly, totally unnecessary) changes that might result by restricting ourselves to $\mathbb{R} \subset \mathbb{C}$ below would be on the order of, e.g., renaming real-valued Hermitian matrices “symmetric matrices”, calling the conjugate-transpose of a real-valued matrix just its “transpose”, etc..

We will now begin studying the SVD by proving a relatively simple lemma that establishes some notation as well as a large number of potential matrix factorizations which include the SVD as a special case.

Lemma 3.1.1. *Let $A \in \mathbb{C}^{m \times n}$ and $\{\mathbf{w}_0, \dots, \mathbf{w}_{n-1}\} \subset \mathbb{C}^n$ be an orthonormal basis for $\mathbb{C}^n$.*

Define $s_j := \|A\mathbf{w}_j\|_2$ (reordering the $\mathbf{w}_j$'s as needed so that $s_0 \geq s_1 \geq \dots \geq s_{n-1}$), and let

$$\mathbf{h}_j := \begin{cases} \mathbf{0} & \text{if } s_j = 0 \\ \frac{1}{s_j} A\mathbf{w}_j \in \mathbb{C}^m & \text{if } s_j \neq 0 \end{cases} . \quad (3.1)$$

Finally, let $W \in \mathbb{C}^{n \times n}$ be the unitary matrix with $W_{:,j} = \mathbf{w}_j$ for all $j \in [n]$ and $H \in \mathbb{C}^{m \times n}$ be the matrix with $H_{:,j} = \mathbf{h}_j$ for all $j \in [n]$. Then, we have

$$A = H \operatorname{diag}(s_0, \dots, s_{n-1}) W^*$$

where $s_0 \geq s_1 \geq \dots \geq s_{n-1} \in [0, \infty)$.

Proof. We have that

$$\begin{aligned} AW &= A \begin{pmatrix} | & | & & | \\ \mathbf{w}_0 & \mathbf{w}_1 & \cdots & \mathbf{w}_{n-1} \\ | & | & & | \end{pmatrix} = \begin{pmatrix} | & | & & | \\ A\mathbf{w}_0 & A\mathbf{w}_1 & \cdots & A\mathbf{w}_{n-1} \\ | & | & & | \end{pmatrix} \\ &= \begin{pmatrix} | & | & & | \\ s_0 \mathbf{h}_0 & s_1 \mathbf{h}_1 & \cdots & s_{n-1} \mathbf{h}_{n-1} \\ | & | & & | \end{pmatrix} = \begin{pmatrix} | & | & & | \\ \mathbf{h}_0 & \mathbf{h}_1 & \cdots & \mathbf{h}_{n-1} \\ | & | & & | \end{pmatrix} \begin{pmatrix} s_0 & \cdots & 0 \\ \vdots & \ddots & \vdots \\ 0 & \cdots & s_{n-1} \end{pmatrix} \\ &= H \operatorname{diag}(s_0, \dots, s_{n-1}). \end{aligned}$$

Thus, $A = AWW^* = H \operatorname{diag}(s_0, \dots, s_{n-1}) W^*$. $\square$

Lemma 3.1.1 already yields a large family of decompositions for any given $A \in \mathbb{C}^{m \times n}$ with several of the structural properties that will ultimately be provided by the singular value decomposition. The next lemma tells us how to choose the orthonormal basis $\{\mathbf{w}_j\}_{j \in [n]}$ of $\mathbb{C}^n$ in order to ensure that the $\mathbf{h}_j$ vectors defined in (3.1) can be used to form a unitary matrix. As a happy coincidence, our choice of $\{\mathbf{w}_j\}_{j \in [n]} \subset \mathbb{C}^n$ will also contain an orthonormal basis for the null space of A as subset of its columns, and guarantee the uniqueness of the ordered s_j values from Lemma 3.1.1.

As we shall see, choosing $\{\mathbf{w}_j\}_{j \in [n]} \subset \mathbb{C}^n$ in Lemma 3.1.1 to be an orthonormal basis of $\mathbb{C}^n$ consisting of eigenvectors of $A^*A \in \mathbb{C}^{n \times n}$ is the “correct” choice (at least, if our goal is to try to orthogonalize H as much as possible). And, it’s important to note, this choice is always possible by Fact 2.8.6 since A^*A will always be Hermitian no matter what $A \in \mathbb{C}^{m \times n}$ itself looks like. Toward seeing how nicely this works out, let’s quickly recall some facts about the four fundamental subspaces of both A and A^*A from Section 2.7. First, if $\mathbf{w}_j \in \mathbb{C}^n$ is an eigenvector of A^*A then $A\mathbf{w}_j = \mathbf{0}$ can only hold if $\mathbf{w}_j \in \mathcal{N}(A) = \mathcal{C}(A^*)^\perp = \mathcal{C}(A^*A)^\perp = \mathcal{N}(A^*A)$ (see, e.g., Lemmas 2.7.2 and 2.7.3). Second, A is rank r if and only if A^*A is rank r (see Theorem 2.7.1 and Lemma 2.7.3). Thus, if A is rank r there will be exactly r orthonormal eigenvectors of A^*A associated with nonzero eigenvalues, and they will span $\mathcal{C}(A^*A) = \mathcal{C}(A^*)$.

Exercise 3.1.1. Let $A \in \mathbb{C}^{m \times n}$ be rank r and $\{\mathbf{w}_j\}_{j \in [n]} \subset \mathbb{C}^n$ be an orthonormal basis of $\mathbb{C}^n$ consisting of eigenvectors of $A^*A \in \mathbb{C}^{n \times n}$. Prove that exactly r of the orthonormal eigenvectors of A^*A in $\{\mathbf{w}_j\}_{j \in [n]}$ will be associated with nonzero eigenvalues. Suppose, w.l.g., that these r orthonormal eigenvectors of A^*A are $\{\mathbf{w}_j\}_{j \in [r]}$. Show that they are an orthonormal basis of $\mathcal{C}(A^*)$.

Exercise 3.1.2. Let $A \in \mathbb{C}^{m \times n}$ be rank r and $\{\mathbf{w}_j\}_{j \in [n]} \subset \mathbb{C}^n$ be an orthonormal basis of $\mathbb{C}^n$ consisting of eigenvectors of $A^*A \in \mathbb{C}^{n \times n}$. Prove that $A\mathbf{w}_j = \mathbf{0}$ will hold if and only if $\mathbf{w}_j$ has eigenvalue 0 as an eigenvector of A^*A . Conclude that $A\mathbf{w}_j = \mathbf{0}$ will hold for exactly $n - r$ of the orthonormal eigenvectors of A^*A in $\{\mathbf{w}_j\}_{j \in [n]}$. Suppose, w.l.g., that these $n - r$ orthonormal eigenvectors of A^*A are $\{\mathbf{w}_j\}_{j=r}^{n-1}$. Argue that they are an orthonormal basis of $\mathcal{N}(A)$.

Let $A \in \mathbb{C}^{m \times n}$ be rank r . The next lemma shows that choosing $\{\mathbf{w}_j\}_{j \in [n]} \subset \mathbb{C}^n$ in Lemma 3.1.1 to be an orthonormal basis of $\mathbb{C}^n$ consisting of eigenvectors of $A^*A \in \mathbb{C}^{n \times n}$ will result in exactly r nonzero and orthonormal $\mathbf{h}_j$ vectors in (3.1).

Lemma 3.1.2. Let $A \in \mathbb{C}^{m \times n}$ be rank r . Choose $\{\mathbf{w}_j\}_{j \in [n]} \subset \mathbb{C}^n$ in Lemma 3.1.1 to be an orthonormal basis of $\mathbb{C}^n$ consisting of eigenvectors of $A^*A \in \mathbb{C}^{n \times n}$. Then the $\mathbf{h}_j$ vectors defined in (3.1) will be such that $\{\mathbf{h}_j\}_{j \in [r]} \subset \mathbb{C}^m$ form an orthonormal basis of $\mathcal{C}(A)$, and $\mathbf{h}_j = \mathbf{0}$ for all $j = r, \dots, n - 1$.

Proof. Exactly r of the $\mathbf{h}_j$ vectors defined in (3.1) will be nonzero by Exercise 3.1.2. Furthermore, these nonzero $\mathbf{h}_j$ vectors will be $\{\mathbf{h}_j\}_{j \in [r]}$ due to the ordering imposed on the $s_j = \|A\mathbf{w}_j\|_2$ values. Finally, each $\mathbf{h}_j \in \mathcal{C}(A)$ will have $\|\mathbf{h}_j\|_2 = 1$ for all $j \in [r]$ by the definition of the $\mathbf{h}_j$ vectors in (3.1). Thus, to finish the proof it suffices by Exercise 2.4.5 to prove that $\{\mathbf{h}_j\}_{j \in [r]}$ is orthogonal.

Let λ_ℓ be the eigenvalue of A^*A associated with an eigenvector $\mathbf{w}_\ell$ for all $0 \leq \ell < r$. Considering the inner product of any two nonzero $\mathbf{h}_j$ vectors from (3.1) we have that

$$\langle \mathbf{h}_j, \mathbf{h}_\ell \rangle = \frac{1}{s_j s_\ell} \langle A\mathbf{w}_j, A\mathbf{w}_\ell \rangle = \frac{1}{s_j s_\ell} (A\mathbf{w}_j)^* A\mathbf{w}_\ell = \frac{1}{s_j s_\ell} \mathbf{w}_j^* (A^* A \mathbf{w}_\ell) = \frac{\lambda_\ell}{s_j s_\ell} \mathbf{w}_j^* \mathbf{w}_\ell = 0$$

whenever $j \neq \ell$ due to the orthonormality of $\{\mathbf{w}_j\}_{j \in [n]}$. Hence, $\{\mathbf{h}_j\}_{j \in [r]}$ is an orthonormal basis of $\mathcal{C}(A)$. $\square$

Exercise 3.1.3. Let $A \in \mathbb{C}^{m \times n}$ be rank r . Suppose that some choice of the orthonormal basis $\{\mathbf{w}_j\}_{j \in [n]}$ of $\mathbb{C}^n$ in Lemma 3.1.1 results in exactly r orthonormal $\mathbf{h}_j$ vectors in (3.1). Prove that every $\mathbf{w}_j$ must then be an eigenvector of $A^*A \in \mathbb{C}^{n \times n}$.

Lemma 3.1.2 combined with Exercise 3.1.3 imply that there is essentially only one way to apply Lemma 3.1.1 so that its H matrix ends up having exactly $r = \text{rank}(A)$ nonzero and orthonormal columns $\{\mathbf{h}_j\}_{j \in [r]}$. We simply must choose $\{\mathbf{w}_j\}_{j \in [n]} \subset \mathbb{C}^n$ to be an orthonormal basis of $\mathbb{C}^n$ consisting of eigenvectors of $A^*A \in \mathbb{C}^{n \times n}$. Making that choice, we

then have that $\{\mathbf{h}_j\}_{j \in [r]}$ will be an orthonormal basis of $\mathcal{C}(A) \subset \mathbb{C}^m$. We can, therefore, complete $\{\mathbf{h}_j\}_{j \in [r]}$ to be larger orthonormal basis $B = \{\mathbf{h}_j\}_{j \in [r]} \cup \{\mathbf{u}_\ell\}_{\ell=r}^{m-1}$ of all of $\mathbb{C}^m$, where $\{\mathbf{u}_\ell\}_{\ell=r}^{m-1}$ will then be an orthonormal basis of $\mathcal{C}(A)^\perp = \mathcal{N}(A^*)$ by construction.

Let $U \in \mathbb{C}^{m \times m}$ be the unitary matrix with its columns given by

$$U_{:,j} = \begin{cases} \mathbf{h}_j & \text{if } j \in [r] \\ \mathbf{u}_j & \text{otherwise} \end{cases}.$$

In addition, let $V \in \mathbb{C}^{n \times n}$ be the unitary matrix whose columns are our well-chosen $\{\mathbf{w}_j\}_{j \in [n]}$ basis so that $V_{:,j} = \mathbf{w}_j$ for all $j \in [n]$. For our $A \in \mathbb{C}^{m \times n}$ we will then have that

$$\begin{aligned} AV &= \begin{pmatrix} | & | & & | \\ A\mathbf{w}_0 & A\mathbf{w}_1 & \cdots & A\mathbf{w}_{n-1} \\ | & | & & | \end{pmatrix} \\ &= \begin{pmatrix} | & | & & | & | & | & | \\ s_0\mathbf{h}_0 & s_1\mathbf{h}_1 & \cdots & s_{r-1}\mathbf{h}_{r-1} & \mathbf{0} & \cdots & \mathbf{0} \\ | & | & & | & | & | & | \end{pmatrix} \in \mathbb{C}^{m \times n} \\ &= \underbrace{\begin{pmatrix} | & | & & | & | & | & | \\ \mathbf{h}_0 & \mathbf{h}_1 & \cdots & \mathbf{h}_{r-1} & \mathbf{u}_r & \cdots & \mathbf{u}_{m-1} \\ | & | & & | & | & | & | \end{pmatrix}}_{\in \mathbb{C}^{m \times m}} \underbrace{\begin{pmatrix} \text{diag}(s_0, \dots, s_{r-1}) & \mathbf{0} & \cdots & \mathbf{0} \\ \mathbf{0} & \mathbf{0} & \cdots & \mathbf{0} & \mathbf{0} & \mathbf{0} & \cdots & \mathbf{0} \end{pmatrix}}_{\in \mathbb{C}^{m \times n}} \\ &= U\Sigma, \end{aligned} \tag{3.2}$$

where $\Sigma \in [0, \infty)^{m \times n}$ is a real-valued diagonal matrix whose entries are given by

$$\Sigma_{i,j} = \begin{cases} s_j & \text{if } i = j < r \\ 0 & \text{otherwise} \end{cases}.$$

Multiplying (3.2) through on the right by V^* we finally see that

$$A = AVV^* = U\Sigma V^*.$$

Example 3.1.3. To help the reader digest the abstract computation in (3.2) we will perform a specific example of it here. Let $A = \begin{pmatrix} 1 & -1 & 1 \\ 0 & 2 & 2 \end{pmatrix}$ so that $A^*A = \begin{pmatrix} 1 & -1 & 1 \\ -1 & 5 & 3 \\ 1 & 3 & 5 \end{pmatrix}$. One can then check that

$$\left\{ \begin{pmatrix} 0 \\ \frac{1}{\sqrt{2}} \\ \frac{1}{\sqrt{2}} \end{pmatrix}, \begin{pmatrix} \frac{1}{\sqrt{3}} \\ \frac{-1}{\sqrt{3}} \\ \frac{1}{\sqrt{3}} \end{pmatrix}, \begin{pmatrix} \frac{-2}{\sqrt{6}} \\ \frac{-1}{\sqrt{6}} \\ \frac{1}{\sqrt{6}} \end{pmatrix} \right\} \subset \mathbb{R}^3$$

is an orthonormal set of eigenvectors of A^*A (do check this!). Applying A to each of these vectors we obtain

$$A \begin{pmatrix} 0 \\ \frac{1}{\sqrt{2}} \\ \frac{1}{\sqrt{2}} \end{pmatrix} = 2\sqrt{2} \begin{pmatrix} 0 \\ 1 \\ 1 \end{pmatrix}, \quad A \begin{pmatrix} \frac{1}{\sqrt{3}} \\ \frac{-1}{\sqrt{3}} \\ \frac{1}{\sqrt{3}} \end{pmatrix} = \sqrt{3} \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix}, \quad \text{and} \quad A \begin{pmatrix} \frac{-2}{\sqrt{6}} \\ \frac{1}{\sqrt{6}} \\ \frac{1}{\sqrt{6}} \end{pmatrix} = \begin{pmatrix} 0 \\ 0 \\ 0 \end{pmatrix}.$$

Thus, in the terminology of Lemma 3.1.1, we have $s_0 = 2\sqrt{2}$, $s_1 = \sqrt{3}$, $s_2 = 0$, and

$$\mathbf{h}_0 = \begin{pmatrix} 0 \\ 1 \\ 1 \end{pmatrix}, \quad \mathbf{h}_1 = \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix}, \quad \mathbf{h}_2 = \begin{pmatrix} 0 \\ 0 \\ 0 \end{pmatrix}.$$

Forming the unitary matrices $U \in \mathbb{R}^{2 \times 2}$ and $V \in \mathbb{R}^{3 \times 3}$ used in (3.2) in this case and carrying out the computation to its conclusion we learn that

$$A = \begin{pmatrix} 1 & -1 & 1 \\ 0 & 2 & 2 \end{pmatrix} = \underbrace{\begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix}}_U \underbrace{\begin{pmatrix} 2\sqrt{2} & 0 & 0 \\ 0 & \sqrt{3} & 0 \end{pmatrix}}_\Sigma \underbrace{\begin{pmatrix} 0 & \frac{1}{\sqrt{2}} & \frac{1}{\sqrt{2}} \\ \frac{1}{\sqrt{3}} & \frac{\sqrt{3}}{\sqrt{6}} & \frac{\sqrt{3}}{\sqrt{6}} \\ \frac{-2}{\sqrt{6}} & \frac{-1}{\sqrt{6}} & \frac{1}{\sqrt{6}} \end{pmatrix}}_{V^*}.$$

Exercise 3.1.4. Repeat the calculation in Example 3.1.3 for the matrix $A = \begin{pmatrix} 1 & 1 \\ 1 & 1 \\ -1 & 1 \end{pmatrix}$.

Formalizing the discussion above allows us to prove the following theorem establishing the existence of the SVD for any matrix $A \in \mathbb{C}^{m \times n}$.

Theorem 3.1.4 (The Full Singular Value Decomposition). *Every rank r matrix $A \in \mathbb{C}^{m \times n}$ can be decomposed into $A = U\Sigma V^*$ where*

1. $U \in \mathbb{C}^{m \times m}$ and $V \in \mathbb{C}^{n \times n}$ are both unitary, and
2. $\Sigma \in [0, \infty)^{m \times n}$ is a unique diagonal matrix with entries

$$\Sigma_{i,j} = \begin{cases} \sigma_j(A) & \text{if } i = j \\ 0 & \text{otherwise} \end{cases}$$

satisfying $\sigma_0(A) \geq \sigma_1(A) \geq \cdots \geq \sigma_{r-1}(A) > 0 = \sigma_r(A) = \cdots = \sigma_{\min\{m,n\}-1}(A)$.

Here the j^{th} -largest diagonal entry of the diagonal matrix Σ , $\sigma_j(A) \in [0, \infty)$, is called the j^{th} **singular value** of A . Similarly, given a valid SVD of A , $A = U\Sigma V^*$, the vectors $\mathbf{u}_j = U_{:,j} \in \mathbb{C}^m$ and $\mathbf{v}_j = V_{:,j} \in \mathbb{C}^n$ are called the j^{th} **left and right (respectively) singular vectors of (the SVD of) A** .¹

¹These slightly awkward names for $\mathbf{u}_j = U_{:,j} \in \mathbb{C}^m$ and $\mathbf{v}_j = V_{:,j} \in \mathbb{C}^n$ are due to the fact that these vectors are not generally unique for a given matrix A . Note that there will be many unitary U and V matrix pairs that work as part of a valid SVD of A , especially when there are repeated singular values.

Exercise 3.1.5. Let $A \in \mathbb{C}^{m \times n}$ have the full SVD $A = U\Sigma V^*$. Set $r = \text{rank}(A)$. Show that

$$A = \sum_{j \in [r]} \sigma_j(A) \mathbf{u}_j \mathbf{v}_j^* \quad (3.3)$$

where $\sigma_j(A)$ is the j^{th} singular value of A , and $\mathbf{u}_j = U_{:,j} \in \mathbb{C}^m$, $\mathbf{v}_j = V_{:,j} \in \mathbb{C}^n$ are the j^{th} left/right singular vectors of the SVD of A . (*Hint: Consider using Exercise 2.6.9.*)

One can now prove the following corollary from Theorem 3.1.4 via an argument analogous to the one used to derive Corollary 2.8.9 from Theorem 2.8.7 (or, alternatively, by using (3.3) from Exercise 3.1.5 to build the new factorization more directly).

Corollary 3.1.5 (The Compact Singular Value Decomposition). *Every rank r matrix $A \in \mathbb{C}^{m \times n}$ can be decomposed into $A = U\Sigma V^*$ where*

1. $U \in \mathbb{C}^{m \times r}$ and $V \in \mathbb{C}^{n \times r}$ are both orthonormal matrices, and
2. $\Sigma = \text{diag}(\sigma_0(A), \dots, \sigma_{r-1}(A)) \in [0, \infty)^{r \times r}$ is a unique diagonal matrix containing the r nonzero singular values of A ordered so that $\sigma_0(A) \geq \sigma_1(A) \geq \dots \geq \sigma_{r-1}(A) > 0$.

Exercise 3.1.6. Prove Corollary 3.1.5.

However one proves Theorem 3.1.4 and Corollary 3.1.5, the *uniqueness* of the singular values of a matrix $A \in \mathbb{C}^{m \times n}$ ultimately follows from the fact that they must always be the square roots of the eigenvalues of $A^*A \in \mathbb{C}^{n \times n}$ (and $AA^* \in \mathbb{C}^{m \times m}$). For this reason (in addition to several others), we will now briefly review the properties that any valid SVD of a matrix A must share with the spectral decompositions of both A^*A and AA^* .

3.1.1 The Relationship Between any Valid SVD of A and the Spectral Decompositions of A^*A and AA^*

Let $A = U\Sigma V^*$ be a valid full SVD of a rank r matrix $A \in \mathbb{C}^{m \times n}$ (i.e., so that $U \in \mathbb{C}^{m \times m}$ and $V \in \mathbb{C}^{n \times n}$ are both unitary, and $\Sigma \in [0, \infty)^{m \times n}$ is a diagonal matrix satisfying $\Sigma_{0,0} \geq \dots \geq \Sigma_{r-1,r-1} > \Sigma_{r,r} = \dots = \Sigma_{q-1,q-1} = 0$, where $q = \min\{m, n\}$). Notice that then

$$A^*A = (U\Sigma V^*)^*(U\Sigma V^*) = V\Sigma^*U^*U\Sigma V^* = V(\Sigma^*\Sigma)V^*,$$

where $D = \Sigma^*\Sigma \in [0, \infty)^{n \times n}$ is a diagonal matrix with $D_{0,0} = \Sigma_{0,0}^2 \geq \dots \geq D_{r-1,r-1} = \Sigma_{r-1,r-1}^2 > D_{r,r} = \dots = D_{n-1,n-1} = 0$. As a consequence, we can see that every column $\mathbf{v}_j = V_{:,j}$ of V will be an eigenvector of A^*A with eigenvalue $D_{j,j}$ since

$$A^*A\mathbf{v}_j = V(\Sigma^*\Sigma)V^*\mathbf{v}_j = V(\Sigma^*\Sigma)\mathbf{e}_j = VD_{j,j}\mathbf{e}_j = D_{j,j}\mathbf{v}_j.$$

Thus, $D_{j,j}$ must be the j^{th} largest eigenvalue of $A^*A \in \mathbb{C}^{n \times n}$. Given that the eigenvalues of A^*A are both unique (with potential repetitions since they are the zeros of the characteristic polynomial of A^*A – see, e.g., [15, Chapter 10]), and always nonnegative real numbers (see Exercise 2.8.6), this further implies that each $\Sigma_{j,j} = \sqrt{D_{j,j}}$ is also uniquely determined by A . Hence, we'll call the value that $\Sigma_{j,j}$ must always take in any valid full SVD of A “ $\sigma_j(A)$ ”, and will later discuss it even in the absence of a particular SVD of A .

Exercise 3.1.7. Let $A = U\Sigma V^*$ be a valid full SVD of a rank r matrix $A \in \mathbb{C}^{m \times n}$. Show that $\Sigma_{j,j}$ must always equal the square-root of the j^{th} largest eigenvalue of $AA^* \in \mathbb{C}^{m \times m}$. Conclude that the nonzero eigenvalues of $AA^* \in \mathbb{C}^{m \times m}$ must always match the nonzero eigenvalues of $A^*A \in \mathbb{C}^{n \times n}$.

The following result can be proven by carefully considering the discussion so far.

Theorem 3.1.6. Let $A = U\Sigma V^*$ be a valid full SVD of a rank r matrix $A \in \mathbb{C}^{m \times n}$. The following statements must hold:

1. The r nonzero singular values of A are exactly the square roots of the positive eigenvalues of $A^*A \in \mathbb{C}^{n \times n}$ and $AA^* \in \mathbb{C}^{m \times m}$.
2. The first r columns of $U \in \mathbb{C}^{m \times m}$ are an orthonormal basis for the column space of A , $\mathcal{C}(A) \subset \mathbb{C}^m$.
3. The last $m - r$ columns of $U \in \mathbb{C}^{m \times m}$ form an orthonormal basis for the null space of A^* , $\mathcal{N}(A^*) \subset \mathbb{C}^m$.
4. The first r columns of $V \in \mathbb{C}^{n \times n}$ form an orthonormal basis for the column space of A^* , $\mathcal{C}(A^*) \subset \mathbb{C}^n$.
5. The last $n - r$ columns of $V \in \mathbb{C}^{n \times n}$ form an orthonormal basis for the null space of A , $\mathcal{N}(A) \subset \mathbb{C}^n$.
6. If $m = n$ and A is Hermitian, then A will have λ as an eigenvalue if and only if there exists a $j \in [n]$ such that
 - $|\lambda|$ is the j^{th} singular value of A (i.e., $\sigma_j = |\lambda|$),
 - the j^{th} column of V , $\mathbf{v}_j \in \mathbb{C}^n$, is an eigenvector of A associated with λ , and
 - the j^{th} column of $U = \text{sign}(\lambda)\mathbf{v}_j$.

Exercise 3.1.8. Prove Theorem 3.1.6.

Exercise 3.1.9. Let $U \in \mathbb{C}^{m \times m}$ and $V \in \mathbb{C}^{n \times n}$ both be unitary, $A \in \mathbb{C}^{m \times n}$, and $q := \min\{m, n\}$. Show that $\sigma_j(UA) = \sigma_j(A) = \sigma_j(AV)$ holds for all $j \in [q]$.

Exercise 3.1.10. Let $\alpha, \beta \in \mathbb{Z} \setminus \{0\}$. The $\frac{\alpha}{\beta}$ -power of a full rank matrix $A \in \mathbb{C}^{n \times n}$ is a matrix $B \in \mathbb{C}^{n \times n}$ with the property that $B^\beta = A^\alpha$ (e.g., when $\beta = 2$ and $\alpha = 1$ then B is called the matrix square root of A). Prove that there always exists a unitary matrix $W \in \mathbb{C}^{n \times n}$ such that any desired $\frac{\alpha}{\beta}$ -power of AW exists. When can W simply be the identity? How can one compute such a B and W for any given $A \in \mathbb{C}^{n \times n}$?

As Theorem 3.1.6 hopefully makes clear, a SVD of A conveniently encodes just about any standard information you might want to know about A . It is a commonly computed decomposition as a result. Numerically, a SVD of a small to moderately sized matrix $A \in \mathbb{C}^{m \times n}$ can be efficiently computed using a variety of standard methods (depending on how, e.g., m compares in size to n). We refer the interested reader to numerical linear algebra texts such as [34, Lecture 31] or [11, Chapter 5] for details. For an extremely large matrix $A \in \mathbb{C}^{m \times n}$ that might not be (able to be) stored on a single machine, however, one might have to utilize a distributed/incremental SVD algorithm instead (see, e.g., [4, 5, 20]).

3.1.2 The SVD and the Moore–Penrose Inverse of a Matrix

Note that every matrix $A \in \mathbb{C}^{m \times n}$ is a linear bijection from $\mathcal{C}(A^*)$ onto $\mathcal{C}(A)$. Hence, $A : \mathcal{C}(A^*) \rightarrow \mathcal{C}(A)$ always has a linear inverse with the same rank as A called the **Moore–Penrose (or, pseudo)inverse** of A , denoted by $A^\dagger : \mathcal{C}(A) \rightarrow \mathcal{C}(A^*)$. Furthermore, a compact SVD of $A^\dagger \in \mathbb{C}^{n \times m}$ can always be obtained from a compact SVD of A .

Let $A = U\Sigma V^*$ be a compact SVD of a rank r matrix $A \in \mathbb{C}^{m \times n}$ so that $U \in \mathbb{C}^{m \times r}$ and $V \in \mathbb{C}^{n \times r}$ are orthonormal matrices, and $\Sigma = \text{diag}(\sigma_0(A), \dots, \sigma_{r-1}(A)) \in [0, \infty)^{r \times r}$ is invertible (due to $\sigma_0(A) \geq \dots \geq \sigma_{r-1}(A) > 0$). One can now see that

$$A^\dagger = V\Sigma^{-1}U^* \quad (3.4)$$

must hold. To understand why, recall that the orthogonal projections $P_{\mathcal{C}(A)}$ and $P_{\mathcal{C}(A^*)}$ act as the identities on $\mathcal{C}(A)$ and $\mathcal{C}(A^*)$, respectively (see Theorem 2.6.10). And, e.g.,

$$A^\dagger A = (V\Sigma^{-1}U^*)(U\Sigma V^*) = V\Sigma^{-1}I_r\Sigma V^* = VV^* = P_{\mathcal{C}(A^*)}$$

by (2.14) and part (4) of Theorem 3.1.6. Hence, $A^\dagger : \mathcal{C}(A) \rightarrow \mathcal{C}(A^*)$ from (3.4) is indeed the left inverse of $A : \mathcal{C}(A^*) \rightarrow \mathcal{C}(A)$. A similar calculation shows that $AA^\dagger = P_{\mathcal{C}(A)}$ also holds.

Exercise 3.1.11. Let $A = U\Sigma V^*$ be a compact SVD of a rank r matrix $A \in \mathbb{C}^{m \times n}$. Show that $A^\dagger$ from (3.4) satisfies $AA^\dagger = P_{\mathcal{C}(A)}$.

Exercise 3.1.12. Suppose that $A \in \mathbb{C}^{n \times n}$ is full rank (so that $\text{rank}(A) = n$). Show that $A^\dagger = A^{-1} \in \mathbb{C}^{n \times n}$ in this case.

The exercise directly above demonstrates that $A^\dagger$ is a *strict generalization* of the “usual” matrix inverse A^{-1} . As a result, in some sense we always should (and really always effectively do) work with $A^{-1} := A^\dagger$ when thinking about inverting a matrix of any size.

Exercise 3.1.13. Suppose that $A \in \mathbb{C}^{n \times n}$ is full rank (so that $\text{rank}(A) = n$). Show that $\sigma_0(A^{-1}) = \frac{1}{\sigma_{n-1}(A)}$. More generally, show that $\sigma_j(A^{-1}) = \frac{1}{\sigma_{n-1-j}(A)}$ for all $j \in [n]$.

Exercise 3.1.14. Let $A \in \mathbb{C}^{m \times n}$ with $m \geq n$ have rank n . Show that $A^\dagger = (A^*A)^{-1}A^*$.

3.1.3 Singular Values, Matrix Norms, and Some Related Inequalities

If we have not yet convinced you that the SVD is potentially interesting and useful, we will try again here by showing that two of the most commonly used matrix norms from Section 2.2.3 are closely related to the singular values of a given matrix.

The Frobenius Norm

Given $A \in \mathbb{C}^{m \times n}$ recall that $\|A\|_F = \sqrt{\sum_{\ell,j} |A_{\ell,j}|^2}$. Let $A = U\Sigma V^*$ be a full SVD of A , and set $q := \min\{m, n\}$. Computing the squared Frobenius norm of A via its SVD we can see that

$$\begin{aligned} \|A\|_F^2 &= \|U\Sigma V^*\|_F^2 = \sum_{j \in [n]} \|(U\Sigma V^*)_{:,j}\|_2^2 = \sum_{j \in [n]} \|U(\Sigma V^*)_{:,j}\|_2^2 \\ &= \sum_{j \in [n]} \|(\Sigma V^*)_{:,j}\|_2^2 = \|\Sigma V^*\|_F^2 \end{aligned}$$

by Exercise 2.6.11 since U is unitary. Continuing, we can further see that since $\|A\|_F = \|A^*\|_F$ holds for all $A \in \mathbb{C}^{m \times n}$ we also have that

$$\|A\|_F^2 = \|V\Sigma^*\|_F^2 = \sum_{j \in [m]} \|V(\Sigma^*)_{:,j}\|_2^2 = \sum_{j \in [m]} \|\Sigma^*_{:,j}\|_2^2 = \sum_{j \in [q]} (\sigma_j(A))^2. \quad (3.5)$$

We will see that (3.5) has several important implications in later sections.

Exercise 3.1.15. Let $U \in \mathbb{C}^{m \times m}$ and $V \in \mathbb{C}^{n \times n}$ both be unitary. Show that $\|UA\|_F = \|A\|_F = \|AV\|_F$ holds for all $A \in \mathbb{C}^{m \times n}$.

The (ℓ^2, ℓ^2) -Operator Norm

Given $A \in \mathbb{C}^{m \times n}$ recall that $\|A\|_{2 \rightarrow 2} = \max_{\mathbf{x} \in \mathbb{C}^n \text{ s.t. } \|\mathbf{x}\|_2=1} \|A\mathbf{x}\|_2$. Let $A = U\Sigma V^*$ be a full SVD of A , and set $q := \min\{m, n\}$. Computing the (ℓ^2, ℓ^2) -operator norm of A via its SVD we can see that

$$\|A\|_{2 \rightarrow 2} = \max_{\mathbf{x} \in \mathbb{C}^n \text{ s.t. } \|\mathbf{x}\|_2=1} \|U\Sigma V^*\mathbf{x}\|_2 = \max_{\mathbf{x} \in \mathbb{C}^n \text{ s.t. } \|\mathbf{x}\|_2=1} \|\Sigma V^*\mathbf{x}\|_2$$

by Exercise 2.6.11 since U is unitary. Furthermore, since V is also unitary its columns form an orthonormal basis of $\mathbb{C}^n$ so that every $\mathbf{x} \in \mathbb{C}^n$ with $\|\mathbf{x}\|_2 = 1$ can be written as $\mathbf{x} = \sum_{j=1}^n \alpha_j V_{:,j}$ where $\|\boldsymbol{\alpha}\|_2 = \|\mathbf{x}\|_2 = 1$ (see Theorem 2.3.9). Thus, continuing we can see that

$$\begin{aligned} \|A\|_{2 \rightarrow 2} &= \max_{\boldsymbol{\alpha} \in \mathbb{C}^n \text{ s.t. } \|\boldsymbol{\alpha}\|_2=1} \left\| \Sigma V^* \left(\sum_{j=1}^n \alpha_j V_{:,j} \right) \right\|_2 = \max_{\boldsymbol{\alpha} \in \mathbb{C}^n \text{ s.t. } \|\boldsymbol{\alpha}\|_2=1} \left\| \Sigma \left(\sum_{j=1}^n \alpha_j \mathbf{e}_j \right) \right\|_2 \\ &= \max_{\boldsymbol{\alpha} \in \mathbb{C}^n \text{ s.t. } \|\boldsymbol{\alpha}\|_2=1} \|\Sigma \boldsymbol{\alpha}\|_2 = \max_{\boldsymbol{\alpha} \in \mathbb{C}^n \text{ s.t. } \|\boldsymbol{\alpha}\|_2=1} \sqrt{\sum_{j \in [q]} |\alpha_j|^2 (\sigma_j(A))^2}. \end{aligned}$$

Recalling that $\sigma_0(A) \geq \sigma_1(A) \geq \dots \geq \sigma_{q-1}(A)$ we can now see that this last expression is always maximized when $|\alpha_0| = 1$. Hence,

$$\|A\|_{2 \rightarrow 2} = \sigma_0(A). \quad (3.6)$$

We will see that (3.6) also has several important implications in later sections.

Exercise 3.1.16. Let $U \in \mathbb{C}^{m \times m}$ and $V \in \mathbb{C}^{n \times n}$ both be unitary. Show that $\|UA\|_{2 \rightarrow 2} = \|A\|_{2 \rightarrow 2} = \|AV\|_{2 \rightarrow 2}$ holds for all $A \in \mathbb{C}^{m \times n}$.

Exercise 3.1.17. Let $A \in \mathbb{C}^{m \times n}$ and set $q := \min\{m, n\}$. Prove that $\|A\|_{2 \rightarrow 2} \leq \|A\|_F \leq \sqrt{q} \|A\|_{2 \rightarrow 2}$ always holds. For what type of matrices will $\|A\|_{2 \rightarrow 2} = \|A\|_F$ hold? For what type of matrices will $\|A\|_F = \sqrt{q} \|A\|_{2 \rightarrow 2}$ hold?

Some Singular Value Inequalities

Now that we have seen a few reasons why we might want to compute a singular value decomposition of a matrix (e.g., to compute its Moore–Penrose inverse, or its (ℓ^2, ℓ^2) -operator norm), it's worth considering how robust a matrix SVD actually is to small errors. Imagine, for example, that we want to compute the singular values of a matrix $A \in \mathbb{C}^{m \times n}$ on a digital computer. We will encounter potential problems immediately since, unfortunately, we probably can't even store A exactly on our computer! Instead, we will actually store $A + E$, where $E \in \mathbb{C}^{m \times n}$ contains all the round-off errors that result from representing each entry of A with a finite number of binary digits (i.e., bits). Given that we can (at best) then compute the singular values of $A + E$ instead of A , it'd be good to know how close the singular values of $A + E$ are to the true singular values of A we actually want. If, e.g., E has a small Frobenius norm (and, therefore, small singular values by (3.5)) we want to make sure that $\sigma_j(A + E) \approx \sigma_j(A)$ holds for all relevant j . We will now state some very useful singular value inequalities which effectively show that singular values are indeed robust to small perturbations in this way (both additive and multiplicative).

Theorem 3.1.7 (See Theorem 3.3.16 in [18]). Let $A, B \in \mathbb{C}^{m \times n}$ and $q = \min\{m, n\}$. Then

$$(a) \sigma_{j+k}(A+B) \leq \sigma_j(A) + \sigma_k(B), \text{ and}$$

$$(b) \sigma_{j+k}(AB^*) \leq \sigma_j(A) \sigma_k(B)$$

for all $j, k \in [q]$ such that $j+k \in [q]$. In particular,

$$(c) |\sigma_j(A+B) - \sigma_j(A)| \leq \sigma_0(B) \quad \forall j \in [q], \text{ and}$$

$$(d) \sigma_j(AB^*) \leq \sigma_j(A) \sigma_0(B) \quad \forall j \in [q].$$

Exercise 3.1.18. Let $B \in \mathbb{C}^{m \times n}$ and $q = \min\{m, n\}$. Prove that $\sigma_j(-B) = \sigma_j(B) = \sigma_j(B^*)$ holds for all $j \in [q]$.

Exercise 3.1.19. Use parts (a) and (b) of Theorem 3.1.7 to prove parts (c) and (d).

Looking at Theorem 3.1.7 (c) one can see that if $B = E$ has a small largest singular value, $\sigma_0(E)$, then we will indeed have $\sigma_j(A+E) \approx \sigma_j(A)$ for all $j \in [q]$. Furthermore, one can also use these inequalities to see, e.g., that small perturbations to the entries of A won't influence how it behaves as a linear function too much either. This means that matrices can be applied as linear functions on digital computers without distorting their outputs too extremely.

Example 3.1.8. Suppose that $A \in \mathbb{C}^{m \times n}$ is stored on a digital computer as $\tilde{A} = A + E$, where $E \in \mathbb{C}^{m \times n}$ is, e.g., a round-off error matrix with $|E_{i,j}| \leq \epsilon$ for all i, j . How much can $\tilde{A}\mathbf{x}$ differ from $A\mathbf{x}$ on a worst-case input vector $\mathbf{x} \in \mathbb{C}^n$?

To answer this question we will upper bound $\|A\mathbf{x} - \tilde{A}\mathbf{x}\|_2$. Considering this error we can see that

$$\begin{aligned} \|A\mathbf{x} - \tilde{A}\mathbf{x}\|_2 &= \|\mathbf{x}\|_2 \left\| (A - \tilde{A}) \frac{\mathbf{x}}{\|\mathbf{x}\|_2} \right\|_2 = \|\mathbf{x}\|_2 \left\| E \left(\frac{\mathbf{x}}{\|\mathbf{x}\|_2} \right) \right\|_2 \\ &\leq \|\mathbf{x}\|_2 \|E\|_{2 \rightarrow 2} = \|\mathbf{x}\|_2 \sigma_0(E). \end{aligned}$$

If we want an upper bound in terms of ϵ we can now use the fact that $\|E\|_{2 \rightarrow 2} \leq \|E\|_F$ always holds (see Exercise 3.1.17) to get that

$$\|A\mathbf{x} - \tilde{A}\mathbf{x}\|_2 \leq \|\mathbf{x}\|_2 \|E\|_F \leq \epsilon \|\mathbf{x}\|_2 \sqrt{mn}.$$

Thus, the error is will always be small in ℓ^2 -norm as long as ϵ is small compared to $\|\mathbf{x}\|_2 \sqrt{mn}$.

The following two singular value inequalities will be useful in later chapters.

Lemma 3.1.9 (A Slight Generalization of Theorem 3.1.7 Part (d)). Let $A \in \mathbb{C}^{m \times n}$ and $B \in \mathbb{C}^{n \times p}$. Then,

$$\sigma_j(AB) \leq \min\{\sigma_j(A) \sigma_0(B), \sigma_j(B) \sigma_0(A)\} \quad \forall j \in [\min\{n, m, p\}].$$

Proof. Suppose, without loss of generality, that $m \leq p$ (else, we may instead apply the argument below to $\sigma_j((AB)^*) = \sigma_j(B^*A^*)$ using that $\sigma_j(AB) = \sigma_j((AB)^*)$). Since $m \leq p$ we can project all of $\mathbb{C}^m$ into $\mathbb{C}^p$ with $Q = \begin{pmatrix} I_m \\ \dots \mathbf{0} \dots \end{pmatrix} \in \mathbb{C}^{p \times m}$. Further, we may note that

$$\sigma_j(QA) = \sigma_j(A) \text{ and } \sigma_j(QAB) = \sigma_j(AB) \quad \forall j \in [\min\{m, n\}] = [\min\{n, m, p\}].$$

Applying part (d) of Theorem 3.1.7 we can now see that both

$$\sigma_j(AB) = \sigma_j(QAB) \leq \sigma_j(QA) \sigma_0(B^*) = \sigma_j(A) \sigma_0(B)$$

and

$$\sigma_j(AB) = \sigma_j(QAB) = \sigma_j(B^*(QA)^*) \leq \sigma_j(B^*) \sigma_0(QA) = \sigma_j(B) \sigma_0(A)$$

hold. The result follows. $\square$

Lemma 3.1.10. *Let $A \in \mathbb{C}^{m \times n}$ and $B \in \mathbb{C}^{n \times p}$ be such that*

- 1.) *B has a full SVD $B = U\Sigma V^*$, and*
- 2.) *AU has rank $r = n$ with a compact SVD $AU = Q\tilde{\Sigma}P^*$.*

Then,

$$\sigma_j(AB) \geq \sigma_r(AU) \sigma_j(B) \quad \forall j \in [r].$$

Proof. Let $j \in [r]$. Noting that P is unitary since $r = n$, we can see that

$$\sigma_j(B) = \sigma_j(\Sigma V^*) = \sigma_j\left(P\tilde{\Sigma}^{-1}\tilde{\Sigma}P^*\Sigma V^*\right) \leq \sigma_0\left(P\tilde{\Sigma}^{-1}\right) \cdot \sigma_j\left(\tilde{\Sigma}P^*\Sigma V^*\right)$$

by Lemma 3.1.9. Furthermore, $\sigma_0\left(P\tilde{\Sigma}^{-1}\right) = \sigma_0\left(\tilde{\Sigma}^{-1}\right) = \frac{1}{\sigma_r(\tilde{\Sigma})} = \frac{1}{\sigma_r(AU)}$. Hence,

$$\sigma_j(B) \leq \frac{\sigma_j\left(\tilde{\Sigma}P^*\Sigma V^*\right)}{\sigma_r(AU)} \implies \sigma_j\left(\tilde{\Sigma}P^*\Sigma V^*\right) \geq \sigma_r(AU) \sigma_j(B). \quad (3.7)$$

Finally, since $m \geq r = n$ we can see that $Q^*Q = I_n$ so that

$$\begin{aligned} \sigma_j\left(\tilde{\Sigma}P^*\Sigma V^*\right) &= \sigma_j\left(Q^*Q\tilde{\Sigma}P^*\Sigma V^*\right) = \sigma_j(Q^*AU\Sigma V^*) = \sigma_j(Q^*AB) \\ &\leq \sigma_j(AB) \sigma_0(Q^*) = \sigma_j(AB) \end{aligned} \quad (3.8)$$

by Lemma 3.1.9. Combining (3.7) and (3.8) now finishes the proof. $\square$

Though perturbation bounds for singular values such as those in Theorem 3.1.7 are both more commonly used and far more robust, it's also worth knowing about the existence of similarly useful perturbation theory for singular vectors/subspaces as well. We urge the interested reader to peruse, e.g., [30, 31] to get a good overview of these results.

3.1.4 Optimal Low-Rank Approximation

Recalling Section 1.2.3, suppose that we have trained a deep FNN resulting in a large number of huge weight matrices, $W_j \in \mathbb{R}^{d_j \times d_{j-1}}$, where both d_j and d_{j-1} are “big” for most $j \in [L]$. Our goal is to compress these huge weight matrices as much as possible so that our FNN is easier to store. Simultaneously, we want to accurately preserve each weight matrix as a linear function so that our overall FNN still does what we need it to do after compression. Motivated by, e.g., Section 2.5 we can aim to accomplish our goal by approximating each huge weight matrix W_j by a new low-rank matrix $\tilde{W}_j$ that we can then store in an optimally compressed form. At the same time, Example 3.1.8 implies that it would also be helpful to, e.g., produce $\tilde{W}_j$ in a way that reduces the value of $\|W_j - \tilde{W}_j\|_{2 \rightarrow 2} = \sigma_0(W_j - \tilde{W}_j)$ as much as possible since doing so will help to keep $W_j \mathbf{x} \approx \tilde{W}_j \mathbf{x}$ for all $\mathbf{x} \in \mathbb{R}^{d_{j-1}}$.

These considerations collectively suggest the following two step low-rank compression approach for our FNN weight matrices:

1. Approximate each of $W_0, \dots, W_L$ using low-rank matrices $\tilde{W}_0, \dots, \tilde{W}_L$ so that, e.g., $\|W_j - \tilde{W}_j\|_{2 \rightarrow 2}$ is small for all $j \in [L]$, and then
2. store $\tilde{W}_0, \dots, \tilde{W}_L$ in a compressed format.

We have already discussed step 2 above in Section 2.5, so we will focus on step 1 here. As we shall see, the SVD is once again extremely useful in this setting, and ultimately allows us to accomplish step 1 in an optimal way.

Let $A \in \mathbb{C}^{m \times n}$ be an arbitrary (e.g., full rank) matrix, and suppose that we want to approximate A with a rank s matrix $A_s \in \mathbb{C}^{m \times n}$ that, e.g., minimizes $\|A - A_s\|_{2 \rightarrow 2}$ over all possible choices of rank s matrices in $\mathbb{C}^{m \times n}$ so that

$$\|A - A_s\|_{2 \rightarrow 2} = \inf_{\text{rank } s \ B \in \mathbb{C}^{m \times n}} \|A - B\|_{2 \rightarrow 2}.$$

To find $A_s \in \mathbb{C}^{m \times n}$, let $A = U\Sigma V^*$ be a full SVD of A and recall that we can then always write

$$A = \sum_{j \in [q]} \sigma_j(A) \mathbf{u}_j \mathbf{v}_j^*,$$

where $q = \min\{m, n\}$, $\mathbf{u}_j = U_{:,j}$, and $\mathbf{v}_j = V_{:,j}$ (see Exercise 3.1.5). We claim that

$$A_s := \sum_{j \in [s]} \sigma_j(A) \mathbf{u}_j \mathbf{v}_j^* \tag{3.9}$$

is then an optimal rank s approximation to A with respect to both the Frobenius and the (ℓ^2, ℓ^2) -operator norms.

Exercise 3.1.20. Let $A \in \mathbb{C}^{m \times n}$, $q = \min\{m, n\}$, and $A_s \in \mathbb{C}^{m \times n}$ be as in (3.9). Show that $\sigma_j(A - A_s) = \sigma_{j+s}(A)$ for all $j \in [q - s]$, and that $\sigma_j(A - A_s) = 0$ for all $q - s \leq j < q$.

Optimality of A_s in the Frobenius Norm

Observe that for the Frobenius norm we have

$$\|A - A_s\|_F^2 = \left\| \sum_{j=s}^q \sigma_j(A) \mathbf{u}_j \mathbf{v}_j^* \right\|_F^2 = \sum_{j=s}^{q-1} \sigma_j^2(A) \quad (3.10)$$

by (3.5) and Exercise 3.1.20. The next theorem shows that this approximation error is minimal.

Theorem 3.1.11. *Let $A, B \in \mathbb{C}^{m \times n}$, $q = \min\{m, n\}$, and $A_s \in \mathbb{C}^{m \times n}$ be as in (3.9). Furthermore, suppose that B is rank s . Then*

$$\|A - B\|_F \geq \|A - A_s\|_F.$$

That is, A_s is a best rank s approximation to A with respect to Frobenius norm error.

Proof. Note that $\sigma_s(B) = 0$ by Theorem 3.1.4 since B is rank s . Thus, Theorem 3.1.7 implies that

$$\sigma_{j+s}(A) = \sigma_{j+s}((A - B) + B) \leq \sigma_j(A - B) + \sigma_s(B) = \sigma_j(A - B)$$

for all $j \in [q - s]$. As a result, (3.5) and (3.10) now reveal that

$$\begin{aligned} \|A - B\|_F^2 &= \sum_{j \in [q]} \sigma_j^2(A - B) = \sum_{j \in [q-s]} \sigma_j^2(A - B) + \sum_{j \geq q-s} \sigma_j^2(A - B) \\ &\geq \sum_{j \in [q-s]} \sigma_{j+s}^2(A) = \|A - A_s\|_F^2. \end{aligned}$$

Hence, A_s achieves the smallest possible Frobenius norm approximation error achievable by any rank s matrix. $\square$

Optimality of A_s in the (ℓ^2, ℓ^2) -Operator Norm

Observe that for the (ℓ^2, ℓ^2) -operator norm we have

$$\|A - A_s\|_{2 \rightarrow 2} = \left\| \sum_{j=s}^q \sigma_j(A) \mathbf{u}_j \mathbf{v}_j^* \right\|_{2 \rightarrow 2} = \sigma_s(A) \quad (3.11)$$

by (3.6) and Exercise 3.1.20. The next theorem shows that this approximation error is also minimal.

Theorem 3.1.12. *Let $A, B \in \mathbb{C}^{m \times n}$, $q = \min\{m, n\}$, and $A_s \in \mathbb{C}^{m \times n}$ be as in (3.9). Furthermore, suppose that B is rank s . Then*

$$\|A - B\|_{2 \rightarrow 2} \geq \|A - A_s\|_{2 \rightarrow 2}.$$

That is, A_s is a best rank s approximation to A with respect to (ℓ^2, ℓ^2) -operator norm error.

Proof. Since B is rank s we can write it in terms of a QR decomposition $B = QR$, where $Q \in \mathbb{C}^{m \times s}$ and $R \in \mathbb{C}^{s \times n}$. Similarly, let $A = U\Sigma V^*$ be a full SVD of A . Since V is unitary, $\mathcal{L} = \text{span}\{V_{:,0}, \dots, V_{:,s}\} \subset \mathbb{C}^n$ has dimension $s + 1$. Also, we know that $\mathcal{C}(R^*)^\perp = \mathcal{N}(R)$ has dimension $n - s$ by (the discussion around) Lemma 2.7.2. Hence, it must be the case that

$$\text{span}\{V_{:,0}, \dots, V_{:,s}\} \cap \mathcal{N}(R)$$

is a linear subspace of $\mathbb{C}^n$ of dimension at least 1 by Exercise 2.3.12. Thus, there exists $\mathbf{n} \in \text{span}\{V_{:,0}, \dots, V_{:,s}\} \cap \mathcal{N}(R)$ with $\|\mathbf{n}\|_2 = 1$.

Using the fact that $\mathbf{n} \in \mathcal{N}(R) \subset \mathcal{N}(B)$, and writing $\mathbf{n}$ as $\sum_{j \in [s+1]} \alpha_j V_{:,j}$ for some $\alpha_0, \dots, \alpha_s \in \mathbb{C}$ with $\|\boldsymbol{\alpha}\|_2 = 1$ (recall Theorem 2.3.9), we can now see that

$$\begin{aligned} \|A - B\|_{2 \rightarrow 2} &\geq \|(A - B)\mathbf{n}\|_2 = \|A\mathbf{n}\|_2 = \left\| U\Sigma V^* \left(\sum_{j \in [s+1]} \alpha_j V_{:,j} \right) \right\|_2 \\ &= \sqrt{\sum_{j \in [s+1]} |\alpha_j|^2 \sigma_j^2(A)}. \end{aligned}$$

Recalling that $\|\boldsymbol{\alpha}\|_2 = 1$, we can now see that the expression above is minimized when $\alpha_j = 0$ for all $j < s$ so that $\alpha_s = 1$. Therefore,

$$\|A - B\|_{2 \rightarrow 2} \geq \sigma_s(A).$$

We are now finished by (3.11). □

3.2 Solving Ill-Conditioned and Noisy Linear Systems

Noise, like death and taxes, is a sad and inescapable fact of life. Every single scientific, engineering, and computational problem you ever encounter will be littered with noise. In this section we will discuss the effect that noise has on the most fundamental linear algebra problem there is – solving a system of linear equations. In the process, we will present two classical strategies for mitigating the effects of noise.

To make our discussion more concrete, suppose that we are given access to a nonzero vector $\mathbf{b} \in \mathbb{C}^n$ and an invertible matrix $A \in \mathbb{C}^{n \times n}$. Our mission is to find a vector

$\mathbf{x} \in \mathbb{C}^n$ such that $A\mathbf{x} = \mathbf{b}$. Not wanting to complete this task by hand (suppose $n > 100$) we instead opt to solve this simple linear system on a computer.² In the process of loading/transferring/typing the matrix A and the vector $\mathbf{b}$ into the computer, though, something unpleasant happens. Though we somehow miraculously get A into computer memory with no errors, we do end up introducing small errors into $\mathbf{b}$ in the process, instead ending up with $\mathbf{b}' := \mathbf{b} + \boldsymbol{\epsilon}$ in its place. Note that this could happen in a thousand different ways. One obvious source of errors in our very simple example would be due to round-off error (i.e., due to the fact that digital computers can only represent a few of the most significant digits of any given number). Our only consolation in that case is that we can assume the “noise” vector $\boldsymbol{\epsilon} \in \mathbb{C}^n$ has a small norm compared to $\mathbf{b}$ (e.g., suppose we can tell that $\|\boldsymbol{\epsilon}\|_2 \leq 10^{-8}\|\mathbf{b}\|_2$).

Taking solace in the fact that $\boldsymbol{\epsilon} \in \mathbb{C}^n$ has a small norm, we go ahead and solve the linear system we currently have in computer memory,

$$A\mathbf{y} = \mathbf{b}' = \mathbf{b} + \boldsymbol{\epsilon}, \quad (3.12)$$

for $\mathbf{y} \in \mathbb{C}^n$ *instead* of solving the original equation $A\mathbf{x} = \mathbf{b}$ that we actually care about. After solving for $\mathbf{y}$ we then ask ourselves how close it is to the true solution $\mathbf{x}$ we want. Doing a rough calculation of our relative error we see that

$$\begin{aligned} \frac{\|\mathbf{x} - \mathbf{y}\|_2}{\|\mathbf{x}\|_2} &= \frac{\|A^{-1}\mathbf{b} - A^{-1}(\mathbf{b} + \boldsymbol{\epsilon})\|_2}{\|\mathbf{x}\|_2} = \frac{\|A^{-1}\boldsymbol{\epsilon}\|_2}{\|\mathbf{x}\|_2} \leq \frac{\|A^{-1}\|_{2 \rightarrow 2} \|\boldsymbol{\epsilon}\|_2}{\|\mathbf{x}\|_2} \\ &\leq \frac{\|A^{-1}\|_{2 \rightarrow 2} (10^{-8}\|\mathbf{b}\|_2)}{\|\mathbf{x}\|_2} = \frac{\|A^{-1}\|_{2 \rightarrow 2} (10^{-8}\|A\mathbf{x}\|_2)}{\|\mathbf{x}\|_2} \\ &\leq \frac{\|A^{-1}\|_{2 \rightarrow 2} (10^{-8}\|A\|_{2 \rightarrow 2}\|\mathbf{x}\|_2)}{\|\mathbf{x}\|_2} = 10^{-8} (\|A^{-1}\|_{2 \rightarrow 2} \|A\|_{2 \rightarrow 2}) \\ &= 10^{-8} \left(\frac{\sigma_0(A)}{\sigma_{n-1}(A)} \right). \end{aligned} \quad (3.13)$$

Hence, our relative error will be small if $\frac{\sigma_0(A)}{\sigma_{n-1}(A)}$ is not too large. If this quantity is larger than 10^8 , however, all bets are off – the relative error could be huge, and our approximate solution $\mathbf{y}$ effectively worthless.

Exercise 3.2.1. Suppose that $\|\boldsymbol{\epsilon}\|_2 = 10^{-8}\|\mathbf{b}\|_2$ in (3.13). Show that in fact $\frac{\|\mathbf{x} - \mathbf{y}\|_2}{\|\mathbf{x}\|_2} = 10^{-8} \left(\frac{\sigma_0(A)}{\sigma_{n-1}(A)} \right)$ is achieved for a particular worst-case $\mathbf{x}$ and $\boldsymbol{\epsilon}$. What are they?

²In fact, digital computers were invented to solve large systems of linear equations in the first place [1]. Somewhat interestingly, the inventor was an American mathematician of Bulgarian descent working at Iowa State University (then College), not at some ivy league university as one might mistakenly assume. Anyone working anywhere can have a huge impact!*

*To anyone with connections to Iowa State who might be reading this footnote some day: oh yes, I am very aware that my motivational message is implicitly casting shade. Michigan State is obviously superior in every conceivable way....).

Given its importance in applications, the quantity $\frac{\sigma_0(A)}{\sigma_{n-1}(A)}$ has a special name.

Definition 3.2.1 (Condition Number). *Let $A \in \mathbb{C}^{m \times n}$ have rank r . The condition number of A is $\kappa(A) := \|A^\dagger\|_{2 \rightarrow 2} \|A\|_{2 \rightarrow 2} = \frac{\sigma_0(A)}{\sigma_{r-1}(A)} \geq 1$.*

Exercise 3.2.2. *Let $U \in \mathbb{C}^{n \times n}$ be a unitary matrix. Show that $\kappa(U) = 1$.*

Exercise 3.2.3. *Let $A, B \in \mathbb{C}^{n \times n}$ both be full rank. Show that $\kappa(AB) \leq \kappa(A)\kappa(B)$.*

Exercise 3.2.4 (Scale Invariance). *Let $A \in \mathbb{C}^{m \times n}$ and $\alpha \in \mathbb{C}$. Show that $\kappa(\alpha A) = \kappa(A)$.*

Exercise 3.2.5. *Let $A \in \mathbb{C}^{m \times n}$. Show that $\kappa(A^\dagger) = \kappa(A)$.*

Exercise 3.2.6. *Let $A_s \in \mathbb{C}^{m \times n}$ be the best rank- s approximation to $A \in \mathbb{C}^{m \times n}$ with respect the Frobenius and (ℓ^2, ℓ^2) -operator norms as per (3.9). Prove that $\kappa(A_s) \leq \kappa(A)$.*

Generally it indeed turns out to be true that the larger the condition number of A is, the larger the relative error $\frac{\|\mathbf{x} - \mathbf{y}\|_2}{\|\mathbf{x}\|_2}$ in (3.13) will end up being. For this reason matrices A with large condition numbers are said to be **ill-conditioned**, and specialized techniques are often used to try to invert them more accurately. We will now briefly discuss two such techniques.

3.2.1 Improving Conditioning by SVD Truncation

Let $A \in \mathbb{C}^{n \times n}$ be the full rank matrix in (3.12), and let $A_s \in \mathbb{C}^{n \times n}$ be its best rank- s approximation as per (3.9). As shown in Exercise 3.2.6, the condition number of A_s will always be less than or equal to the condition number of A . This suggests a relatively obvious strategy to combat poor conditioning. Instead of solving (3.12) directly by inverting A , we can instead solve the related least-squares problem

$$\mathbf{y} = \arg \min_{\mathbf{z} \in \mathbb{C}^n} \|A_s \mathbf{z} - \mathbf{b}'\|_2^2 \quad (3.14)$$

using the better-conditioned and lower-rank $A_s \in \mathbb{C}^{n \times n}$ in place of A . Recalling Sections 2.6.2 and 3.1.2, we can see that $\mathbf{y} = (A_s)^\dagger \mathbf{b}' = (A_s)^\dagger (\mathbf{b} + \boldsymbol{\epsilon})$ will solve (3.14). We can now ask ourselves what the relative error between this new $\mathbf{y}$ and the true solution $\mathbf{x} = A^{-1} \mathbf{b}$ will be in this case.

To help bound the relative error between $\mathbf{x} = A^{-1} \mathbf{b}$ and $\mathbf{y} = (A_s)^\dagger (\mathbf{b} + \boldsymbol{\epsilon})$ it will be useful to have a compact SVD of both A_s and $(A_s)^\dagger$. Toward that end, suppose that $A = U \Sigma V^*$ is a full SVD of the matrix A . From (3.9) we can then see that

$$A_s = U_s \begin{pmatrix} \sigma_0(A) & & \\ & \ddots & \\ & & \sigma_{s-1}(A) \end{pmatrix} (V_s)^*,$$

where $U_s, V_s \in \mathbb{C}^{n \times s}$ contain the first s orthonormal columns of $U, V \in \mathbb{C}^{n \times n}$, respectively. The pseudoinverse of A_s is therefore

$$(A_s)^\dagger = V_s \begin{pmatrix} 1/\sigma_0(A) & & \\ & \ddots & \\ & & 1/\sigma_{s-1}(A) \end{pmatrix} (U_s)^*. \quad (3.15)$$

Finally, using (3.15) we can further see, e.g., that

$$\begin{aligned} (A_s)^\dagger A &= V_s \begin{pmatrix} 1/\sigma_0(A) & & \\ & \ddots & \\ & & 1/\sigma_{s-1}(A) \end{pmatrix} (U_s)^* U \begin{pmatrix} \sigma_0(A) & & \\ & \ddots & \\ & & \sigma_{n-1}(A) \end{pmatrix} V^* \\ &= V_s \begin{pmatrix} 1/\sigma_0(A) & & \\ & \ddots & \\ & & 1/\sigma_{s-1}(A) \end{pmatrix} (I_s \mid 0_{s \times n-s}) \begin{pmatrix} \sigma_0(A) & & \\ & \ddots & \\ & & \sigma_{n-1}(A) \end{pmatrix} V^* \\ &= V_s \begin{pmatrix} 1/\sigma_0(A) & & \\ & \ddots & \\ & & 1/\sigma_{s-1}(A) \end{pmatrix} \begin{pmatrix} \sigma_0(A) & & 0 & \cdots & 0 \\ & \ddots & 0 & \cdots & 0 \\ & & \sigma_{s-1}(A) & 0 & \cdots & 0 \end{pmatrix} V^* \\ &= V_s (I_s \mid 0_{s \times n-s}) V^* = V_s (V_s)^* = P_{\mathcal{C}((A_s)^*)}. \end{aligned} \quad (3.16)$$

Exercise 3.2.7. Let $A \in \mathbb{C}^{n \times n}$ be full rank. Show that $A(A_s)^\dagger = P_{\mathcal{C}(A_s)}$.

Using (3.16) to help bound the relative error between $\mathbf{x} = A^{-1}\mathbf{b}$ and $\mathbf{y} = (A_s)^\dagger(\mathbf{b} + \boldsymbol{\epsilon})$ we have that

$$\begin{aligned} \frac{\|\mathbf{x} - \mathbf{y}\|_2}{\|\mathbf{x}\|_2} &= \frac{\|A^{-1}\mathbf{b} - (A_s)^\dagger(\mathbf{b} + \boldsymbol{\epsilon})\|_2}{\|\mathbf{x}\|_2} \leq \frac{\|(A_s)^\dagger \boldsymbol{\epsilon}\|_2}{\|\mathbf{x}\|_2} + \frac{\|((A_s)^\dagger - A^{-1})\mathbf{b}\|_2}{\|\mathbf{x}\|_2} \\ &\leq \frac{\|(A_s)^\dagger\|_{2 \rightarrow 2} \|\boldsymbol{\epsilon}\|_2}{\|\mathbf{x}\|_2} + \frac{\|((A_s)^\dagger A - I_n)\mathbf{x}\|_2}{\|\mathbf{x}\|_2} \\ &\leq \frac{\|(A_s)^\dagger\|_{2 \rightarrow 2} \cdot 10^{-8} \cdot \|\mathbf{b}\|_2}{\|\mathbf{x}\|_2} + \frac{\|(P_{\mathcal{C}((A_s)^*)} - I_n)\mathbf{x}\|_2}{\|\mathbf{x}\|_2} \\ &\leq 10^{-8} \left(\frac{\sigma_0(A)}{\sigma_{s-1}(A)} \right) + \left\| P_{\mathcal{C}((A_s)^*)^\perp} \left(\frac{\mathbf{x}}{\|\mathbf{x}\|_2} \right) \right\|_2. \end{aligned} \quad (3.17)$$

Suppose that $\sigma_{s-1}(A) > \sigma_{n-1}(A)$. Comparing (3.17) with (3.13) in this case we can see that (3.17) provides a smaller upper bound on the relative error $\|\mathbf{x} - \mathbf{y}\|_2/\|\mathbf{x}\|_2$ than (3.13) does whenever, e.g., $\mathbf{x} \in \mathcal{C}((A_s)^*)$. In particular, Exercise 3.2.1 gives an example where the relative error provided by (3.17) is a strict improvement over the relative error provided by (3.13) (check this!). More generally, since $\kappa(A_s) = \left(\frac{\sigma_0(A)}{\sigma_{s-1}(A)} \right) \leq \left(\frac{\sigma_0(A)}{\sigma_{n-1}(A)} \right) = \kappa(A)$ holds for

all $s \in [n+1] \setminus \{0\}$, using $\mathbf{y} = (A_s)^\dagger(\mathbf{b} + \boldsymbol{\epsilon})$ will often lead to a better relative error than using $\mathbf{y} = A^{-1}(\mathbf{b} + \boldsymbol{\epsilon})$ when s is chosen so that

$$\left\| P_{\mathcal{C}((A_s)^*)^\perp} \left(\frac{\mathbf{x}}{\|\mathbf{x}\|_2} \right) \right\|_2 < (\kappa(A) - \kappa(A_s)) \frac{\|\boldsymbol{\epsilon}\|_2}{\|\mathbf{b}\|_2}.$$

We will next briefly discuss a different strategy for inverting ill-conditioned systems which doesn't require us to entirely discard all of $P_{\mathcal{C}((A_s)^*)^\perp} \left(\frac{\mathbf{x}}{\|\mathbf{x}\|_2} \right)$.

3.2.2 Tikhonov Regularization

Let $A \in \mathbb{C}^{n \times n}$ again be the full rank matrix in (3.12). This time, though, suppose that we pick a positive real $\alpha \in (0, \infty)$ and then solve the modified least-squares problem

$$\mathbf{y} = \arg \min_{\mathbf{z} \in \mathbb{C}^n} \|A\mathbf{z} - \mathbf{b}'\|_2^2 + \alpha \|\mathbf{z}\|_2^2 \quad (3.18)$$

instead of solving (3.12). As $\alpha \rightarrow 0$ we can see that the solution of (3.18) will converge to the solution of (3.12), so we expect that using a small α should still give us about the same answer. As we shall see below, though, there are benefits to solving (3.18) instead of (3.12) when $\kappa(A)$ is very large.

Solving (3.18) in order to estimate $\mathbf{x} = A^{-1}\mathbf{b}$ is known as *Tikhonov regularization*, and the number α above is referred to as the *regularization parameter*. Note that (3.18) has less overhead to set up and start solving than (3.14) does because (3.18) doesn't require us to compute A_s . For this reason solving (3.18) is often preferable to solving (3.14) when A is large and we don't expect to reuse it many times.

HERE!

We can now ask ourselves what the relative error between this new $\mathbf{y}$ and the true solution $\mathbf{x} = A^{-1}\mathbf{b}$ will be in this case.

[24, Theorem 5.9] – . One can see that the solution is $\mathbf{y} = (\alpha I_n + A^*A)^{-1} A^*\mathbf{b}'$.

Let $A \in \mathbb{C}^{m \times n}$ with $(A) = r$ and $\in \mathbb{C}^m$. Again, we want to solve $A =$ for $\in \mathbb{C}^n$. To do so we will instead consider $\alpha + A^*A = A^*$ for some $\alpha \in \mathbb{R}_{>0}$. We hope to find that $\approx$.

Some questions:

1. Can we solve $\alpha + A^*A = A^*$? This is equivalent to asking if $(\alpha I + A^*A) = A^*$ is solvable.

Answer: Suppose $A = U\Sigma V^*$ is the full SVD of A . Then,

$$\begin{aligned} A^*A &= V\Sigma U^*U\Sigma V^* \\ &= V\Sigma^*\Sigma V^* \\ &= V \begin{pmatrix} \sigma_1(A)^2 & & & & \\ & \ddots & & & \\ & & \sigma_r(A)^2 & 0 & \\ & & 0 & 0 & \\ & & & & \ddots & \\ & & & & & 0 \end{pmatrix} V^* \end{aligned}$$

Thus,

$$\alpha I + A^*A = V \begin{pmatrix} \sigma_1(A)^2 + \alpha & & & & \\ & \ddots & & & \\ & & \sigma_r(A)^2 + \alpha & 0 & \\ & & 0 & \alpha & \\ & & & & \ddots & \\ & & & & & \alpha \end{pmatrix} V^*.$$

Thus, $\alpha I + A^*A$ is positive definite since the above decomposition represents its eigenvalue decomposition with eigenvalues $\lambda_j = \alpha + \sigma_j(A)^2 > 0$ for $j = 1, \dots, r$ and $\lambda_j = \alpha > 0$ for $j = r + 1, \dots, n$. Therefore, $\alpha I + A^*A$ always has full rank and so the desired equation is always solvable.

2. What does solving $\alpha I + A^*A = A^*$ have to do with solving $A = ?$

Suppose we have α such that $(\alpha I + A^*A) = A^*$, then $x = (\alpha I + A^*A)^{-1}A^*$. So, if $\alpha \approx 0$, then $x \approx (A^*A)^{-1}A^*$, giving an approximation to the least squares solution. Recall that, for $(A) = n$, if $A = U\Sigma V^*$ is the thin SVD of A , then $A^*A = V\Sigma^2V^*$ since $\Sigma \in \mathbb{C}^{n \times n}$. Further, $(A^*A)^{-1} = V\Sigma^{-2}V^*$ since A having full rank implies $V \in \mathbb{C}^{n \times n}$, so V is unitary. Therefore, the pseudoinverse $A^+ = V\Sigma^{-1}U^* = (V\Sigma^{-2}V^*)(U\Sigma V^*)^* = (A^*A)^{-1}A^*$. Therefore, if $\alpha \approx 0$, then $x \approx A^+$, meaning x approximates the least squares solution.

(The case where A does not have full rank is left as an exercise for the reader.)

3. Why $\alpha > 0$? Why not solve least squares directly?

We take $\alpha > 0$ because $(\alpha I + A^*A)^{-1}A^*$ is better conditioned than $(A^*A)^{-1}A^*$. From

the eigenvalue decomposition

$$(\alpha I + A^* A)^{-1} = V \begin{pmatrix} \frac{1}{\alpha + \sigma_1^2(A)} & & & & \\ & \ddots & & & \\ & & \frac{1}{\alpha + \sigma_r^2(A)} & & \\ & & & 0 & \\ & & & & \ddots \\ & & & & & 0 \end{pmatrix} V^*,$$

we get the SVD

$$(\alpha I + A^* A)^{-1} A^* = V \begin{pmatrix} \frac{\sigma_1(A)}{\alpha + \sigma_1^2(A)} & & & \\ & \ddots & & \\ & & \frac{\sigma_r(A)}{\alpha + \sigma_r^2(A)} & \\ & & & 0 \end{pmatrix} U^*.$$

So $(A) = \frac{\sigma_1(A)(\alpha + \sigma_r^2(A))}{\sigma_r(A)(\alpha + \sigma_1^2(A))}$, which shows $\alpha \rightarrow \infty$ sends $(A) \rightarrow 1$. Note that if we choose α to be too large, we will lose a significant amount of information when solving the modified equation.

Theorem 3.2.2 (Theorem 5.9 in Rainer Kress). *Suppose that $(\alpha I + A^* A) = A^*$. Then is the unique vector minimizing*

$$\|A - \|\|_2^2 + \alpha\|\|_2^2 = \inf_{\in \mathbb{C}^n} \|A - \|\|_2^2 + \alpha\|\|_2^2.$$

Proof. In Progress

□

Why is this useful?

Suppose you are solving an optimization problem. The minimizing vector above could serve as a good initialization step, so start relatively close.

3.3 Linear Least Squares Regression

TODO...

3.4 Discrete Convolution and Fourier Transform Matrices

We begin this section by defining a general class of matrices which are important in many applications including, e.g., as the weight matrices used in a special type of neural network layer known as a “convolutional” neural network layer (recall Definition 1.2.4). As will be clear soon, one advantage of this type of matrix is that it’s defined with many fewer parameters than a generic matrix requires.

Definition 3.4.1 (Toeplitz Matrix). The **Toeplitz matrix** $A \in \mathbb{C}^{m \times n}$ generated by the vector $\mathbf{a} \in \mathbb{C}^{m+n-1}$ is the $\mathbb{C}^{m \times n}$ matrix with entries given by

$$A_{j,k} := a_{(m-1)+k-j}$$

for all $j \in [m], k \in [n]$. We will also denote this matrix by $A = \text{Toep}_{m,n}(\mathbf{a})$. More generally, we will say that a matrix $A \in \mathbb{C}^{m \times n}$ is **Toeplitz** if there exists a vector $\mathbf{a} \in \mathbb{C}^{m+n-1}$ such that $A = \text{Toep}_{m,n}(\mathbf{a})$. We will also define $\text{Toep}_n(\mathbf{a}) := \text{Toep}_{n,n}(\mathbf{a})$ in the case of square matrices.

Example 3.4.2. The Toeplitz matrix $A \in \mathbb{C}^{3 \times 4}$ generated by $\mathbf{a} \in \mathbb{C}^6$ is

$$\text{Toep}_{3,4}(\mathbf{a}) = \begin{pmatrix} a_2 & a_3 & a_4 & a_5 \\ a_1 & a_2 & a_3 & a_4 \\ a_0 & a_1 & a_2 & a_3 \end{pmatrix}.$$

The Toeplitz matrix $A \in \mathbb{C}^{4 \times 3}$ generated by $\mathbf{a} \in \mathbb{C}^6$ is

$$\text{Toep}_{4,3}(\mathbf{a}) = \begin{pmatrix} a_3 & a_4 & a_5 \\ a_2 & a_3 & a_4 \\ a_1 & a_2 & a_3 \\ a_0 & a_1 & a_2 \end{pmatrix}.$$

Note that the entries of $\mathbf{a}$ appear along the bottom row, and then up the rightmost row, of the Toeplitz matrix it generates in a “backwards-L” shape (displayed in blue above). The rest of the Toeplitz matrix is then determined by its being constant along all of its diagonals.

Exercise 3.4.1. Show that $A \in \mathbb{C}^{m \times n}$ is Toeplitz if and only if $A_{j,k} = A_{j+1,k+1}$ holds for all $j \in [m-1]$ and $k \in [n-1]$.

Exercise 3.4.2. Show that A is Toeplitz if and only if A^* is Toeplitz.

Exercise 3.4.3. Given $\mathbf{a} \in \mathbb{C}^n$ let $\text{Reverse}(\mathbf{a}) \in \mathbb{C}^n$ be the vector with entries given by

$$(\text{Reverse}(\mathbf{a}))_j = \overline{a_{n-1-j}}.$$

Show that if $A \in \mathbb{C}^{m \times n}$ is the Toeplitz matrix generated by $\mathbf{a} \in \mathbb{C}^{m+n-1}$, then A^* is the Toeplitz matrix generated by $\text{Reverse}(\mathbf{a})$.

Definition 3.4.3 (Convolutional Layer of Neurons). A **Convolutional Layer of Neurons** $\ell : \mathbb{R}^N \rightarrow \mathbb{R}^d$ is a layer of neurons (recall Definition 1.2.4), $\sigma(W\mathbf{x} + \mathbf{b})$, where the weight matrix $W \in \mathbb{R}^{d \times N}$ is Toeplitz.

We can now see that an $m \times n$ Toeplitz matrix is entirely defined using only $m+n-1 < mn$ parameters. This can have potential benefits during, e.g., NN training. Of more immediate interest in this section, however, is that these matrices can also have runtime advantages as linear functions when coupled with Discrete Fourier Transform techniques. This will be discussed in Sections 3.4.2 and 3.4.3. Before we can understand how these computational advantages appear, however, we first have to discuss a special type of square Toeplitz matrices known as “circulant matrices”.

3.4.1 Circulant and Toeplitz Matrices

As we will see later in Section 3.4.2, the following special class of square Toeplitz matrices is crucial to realizing fast matrix-vector multiplication algorithms for more arbitrary Toeplitz matrices.

Definition 3.4.4 (Circulant Matrix). *The **circulant matrix** generated by a vector $\mathbf{v} \in \mathbb{C}^n$ is the matrix $\text{circ}(\mathbf{v}) \in \mathbb{C}^{n \times n}$ defined by*

$$(\text{circ}(\mathbf{v}))_{j,k} = \mathbf{v}_{(j-k) \bmod n}.$$

We will say that a matrix $A \in \mathbb{C}^{n \times n}$ is **circulant** if there exists a vector $\mathbf{v} \in \mathbb{C}^n$ such that $A = \text{circ}(\mathbf{v})$. Recall that “ $j \bmod n$ ” is defined for all $j \in \mathbb{Z}$ and $n \in \mathbb{N}$ to be the single element contained in the set $\{j + kn \mid k \in \mathbb{Z}\} \cap [n]$ (or, equivalently, to be the unique value $r \in [n] = \{0, 1, \dots, n-1\}$ such that $\exists k \in \mathbb{Z}$ satisfying $j = r + kn$).

Example 3.4.5. The circulant matrix $A \in \mathbb{C}^{4 \times 4}$ generated by $\mathbf{v} \in \mathbb{C}^4$ is

$$\text{circ}(\mathbf{v}) = \begin{pmatrix} v_0 & v_3 & v_2 & v_1 \\ v_1 & v_0 & v_3 & v_2 \\ v_2 & v_1 & v_0 & v_3 \\ v_3 & v_2 & v_1 & v_0 \end{pmatrix}.$$

Note that this matrix is also Toeplitz due to the fact that it's constant along its diagonals (recall Exercise 3.4.1).

Exercise 3.4.4. Show that every circulant matrix is also Toeplitz.

Exercise 3.4.5. Let $A \in \mathbb{C}^{2n \times 2n}$ be circulant. Show that $A_{j,k} = A_{j+n,k+n}$ for all $j, k \in [n]$. More generally, show that $A_{j,k} = A_{(j \pm n) \bmod 2n, (k \pm n) \bmod 2n}$ holds for all $j, k \in [2n]$.

Not only is every circulant matrix a Toeplitz matrix, but any square Toeplitz matrix can be embedded into a larger circulant matrix. Hence, e.g., any algorithm which efficiently multiplies circulant matrices against vectors can also be used to efficiently multiply square Toeplitz matrices against vectors.

Let $\mathbf{a} \in \mathbb{C}^{2n-1}$ and consider the square Toeplitz matrix generated by $\mathbf{a}$, $\text{Toep}_n(\mathbf{a}) \in \mathbb{C}^{n \times n}$, with entries given by

$$(\text{Toep}_n(\mathbf{a}))_{j,k} = a_{(n-1)-(j-k)}. \quad (3.19)$$

Now let $\mathbf{c} \in \mathbb{C}^{2n}$ be defined by $\mathbf{c}^T = (a_{n-1}, a_{n-2}, \dots, a_0, 0, a_{2n-2}, \dots, a_n)$ so that

$$c_\ell = \begin{cases} a_{n-1-\ell} & 0 \leq \ell \leq n-1 \\ 0 & \ell = n \\ a_{3n-1-\ell} & n+1 \leq \ell \leq 2n-1 \end{cases} \quad (3.20)$$

for all $\ell \in [2n]$. Then, the circulant matrix $\text{circ}(\mathbf{c}) \in \mathbb{C}^{2n \times 2n}$ will always take the block form

$$\text{circ}(\mathbf{c}) = \begin{pmatrix} \text{Toep}_n(\mathbf{a}) & A \\ A & \text{Toep}_n(\mathbf{a}) \end{pmatrix} \in \mathbb{C}^{2n \times 2n}, \quad (3.21)$$

where $A \in \mathbb{C}^{n \times n}$.

Example 3.4.6. Let $\mathbf{a} \in \mathbb{C}^3$. The 2×2 Toeplitz matrix generated by $\mathbf{a}$ is

$$\text{Toep}_2(\mathbf{a}) = \begin{pmatrix} a_1 & a_2 \\ a_0 & a_1 \end{pmatrix}.$$

If we form the vector $\mathbf{c} \in \mathbb{C}^4$ defined by $\mathbf{c}^T = (a_1, a_0, 0, a_2)$ then

$$\text{circ}(\mathbf{c}) = \begin{pmatrix} a_1 & a_2 & 0 & a_0 \\ a_0 & a_1 & a_2 & 0 \\ 0 & a_0 & a_1 & a_2 \\ a_2 & 0 & a_0 & a_1 \end{pmatrix} = \begin{pmatrix} \text{Toep}_2(\mathbf{a}) & A \\ A & \text{Toep}_2(\mathbf{a}) \end{pmatrix} \in \mathbb{C}^{4 \times 4}.$$

The following lemma guarantees the upper-left $n \times n$ block of $\text{circ}(\mathbf{c}) \in \mathbb{C}^{2n \times 2n}$ is indeed always $\text{Toep}_n(\mathbf{a}) \in \mathbb{C}^{n \times n}$ as claimed above in (3.21).

Lemma 3.4.7. Let $\mathbf{a} \in \mathbb{C}^{2n-1}$. Build $\mathbf{c} \in \mathbb{C}^{2n}$ from $\mathbf{a}$ entry-wise via (3.20). Then, $(\text{circ}(\mathbf{c}))_{j,k} = (\text{Toep}_n(\mathbf{a}))_{j,k}$ for all $j, k \in [n]$.

Proof. We can see that $-(n-1) \leq j-k \leq (n-1)$ since $j, k \in [n]$. Furthermore,

$$(j-k) \bmod 2n = \begin{cases} j-k & 0 \leq j-k \leq n-1 \\ 2n+(j-k) & -(n-1) \leq j-k < 0 \end{cases}.$$

Hence, if $0 \leq j-k \leq n-1$ we have that

$$(\text{circ}(\mathbf{c}))_{j,k} = c_{(j-k) \bmod n} = c_{j-k} = a_{n-1-(j-k)},$$

and if $-(n-1) \leq j-k < 0$ we have that

$$(\text{circ}(\mathbf{c}))_{j,k} = c_{(j-k) \bmod n} = c_{2n+(j-k)} = a_{3n-1-(2n+(j-k))} = a_{n-1-(j-k)}.$$

Thus, $(\text{circ}(\mathbf{c}))_{j,k} = (\text{Toep}_N(\mathbf{a}))_{j,k} = a_{(n-1)-(j-k)}$ for all $j, k \in [n]$ by (3.19). $\square$

Exercise 3.4.6. Use Lemma 3.4.7 together with Exercise 3.4.5 to show that (3.21) holds for all $n \in \mathbb{N}$.

Computationally, observe that via (3.21) we have for any $\mathbf{v} \in \mathbb{C}^n$ that

$$\text{circ}(\mathbf{c}) \begin{pmatrix} \mathbf{v} \\ \mathbf{0} \end{pmatrix} = \begin{pmatrix} \text{Toep}_n(\mathbf{a}) & A \\ A & \text{Toep}_n(\mathbf{a}) \end{pmatrix} \begin{pmatrix} \mathbf{v} \\ \mathbf{0} \end{pmatrix} = \begin{pmatrix} \text{Toep}_n(\mathbf{a})\mathbf{v} \\ A\mathbf{v} \end{pmatrix}. \quad (3.22)$$

Thus, we can always recover $\text{Toep}_n(\mathbf{a})\mathbf{v}$ from $\text{circ}(\mathbf{c}) \begin{pmatrix} \mathbf{v} \\ \mathbf{0} \end{pmatrix}$ by taking its first n entries. Hence, as previously mentioned, any algorithm which efficiently multiplies circulant matrices against arbitrary vectors can also be used to efficiently multiply square Toeplitz matrices against arbitrary vectors. We will use this fact to our advantage later in Section 3.4.3.

3.4.2 Discrete Fourier Transforms and Circular Convolutions

In this section we will discuss a particular orthonormal basis of $\mathbb{C}^n$, known as the discrete Fourier basis, which is important for a large number of computational reasons involving convolutions. As we shall see, its many remarkable properties are in fact due to the periodic nature of the unit magnitude complex numbers $\{e^{i\theta} \mid \theta \in [0, 2\pi]\} \subset \mathbb{C}$. In particular, given $n \in \mathbb{N}$ the unit magnitude n^{th} root of unity

$$f_n := e^{\frac{-2\pi i}{n}} \in \mathbb{C}$$

will be the atomic building block of the basis, and its properties are therefore crucial.

Exercise 3.4.7. Show that $(f_n)^{kn} = f_n^{kn} = 1$ for all $k \in \mathbb{Z}$.

Exercise 3.4.8. Show that $(f_n)^k = f_n^k \neq 1$ for all nonzero $k \in [n]$.

Exercise 3.4.9. Show that $(f_n)^{\omega j} = f_n^{\omega j} = f_n^{(\omega \bmod n)(j \bmod n)} = f_n^{(\omega j \bmod n)}$ for all $j, \omega \in \mathbb{Z}$.³

Exercise 3.4.10. Suppose that $p, n \in \mathbb{N}$ are such that $\frac{n}{p} \in \mathbb{N}$ (so that p divides n). Show that $(f_n)^{pj} = f_n^{pj} = f_{\frac{n}{p}}^{(j \bmod \frac{n}{p})}$ holds for all $j \in \mathbb{Z}$.

Let $F \in \mathbb{C}^{n \times n}$ be the $n \times n$ matrix whose entries are given by

$$F_{\omega, j} := \frac{f_n^{\omega \cdot j}}{\sqrt{n}}$$

for all $\omega, j \in [n]$. The matrix F is called the **Discrete Fourier Transform (DFT) matrix of size n** . Importantly, the columns of F^* form an orthonormal basis of $\mathbb{C}^n$ (i.e., one can show that F is a unitary matrix – see Exercise 3.4.11). This basis is called the **discrete Fourier basis of $\mathbb{C}^n$** .

³Let $j \in \mathbb{Z}$. Recall that “ $j \bmod n$ ” denotes the unique integer $r \in [n]$ satisfying $j = r + k \cdot n$ for some $k \in \mathbb{Z}$.

Example 3.4.8. Recall that $\mathbf{1} \in \mathbb{C}^n$ denotes the vector of all ones. We have that

$$F\mathbf{1} = \frac{1}{\sqrt{n}} \begin{pmatrix} \sum_{j=0}^{n-1} f_n^{0 \cdot j} \\ \sum_{j=0}^{n-1} f_n^{1 \cdot j} \\ \vdots \\ \sum_{j=0}^{n-1} f_n^{(n-1) \cdot j} \end{pmatrix}.$$

Considering the k^{th} entry of $F\mathbf{1} \in \mathbb{C}^n$ for all $k \neq 0$ we can see that

$$(F\mathbf{1})_k = \frac{1}{\sqrt{n}} \sum_{j=0}^{n-1} f_n^{k \cdot j} = \frac{1}{\sqrt{n}} \left(\frac{1 - f_n^{kn}}{1 - f_n^k} \right) = \frac{1}{\sqrt{n}} \left(\frac{1 - 1}{1 - f_n^k} \right) = 0$$

by Exercises 3.4.7 and 3.4.8. On the other hand, for $k = 0$ we have that

$$(F\mathbf{1})_0 = \frac{1}{\sqrt{n}} \sum_{j=0}^{n-1} f_n^{0 \cdot j} = \frac{1}{\sqrt{n}} \sum_{j=0}^{n-1} 1 = \frac{n}{\sqrt{n}} = \sqrt{n}.$$

Hence, $F\mathbf{1} = \sqrt{n} \mathbf{e}_0$.

Exercise 3.4.11. Prove that the DFT matrix, F , is unitary. (*HINT:* Recall Theorem 2.6.14.)

Exercise 3.4.12. Prove that $\|F\mathbf{v}\|_2^2 = \|\mathbf{v}\|_2^2$ holds for all $\mathbf{v} \in \mathbb{C}^n$. This equality is sometimes referred to as “Parseval’s identity” in the context of the discrete Fourier basis.

The **Discrete Fourier Transform (DFT)** of a vector $\mathbf{v} \in \mathbb{C}^n$ is simply

$$\widehat{\mathbf{v}} := F\mathbf{v} \tag{3.23}$$

with entries given by $\widehat{v}_\omega = \frac{1}{\sqrt{n}} \sum_{j=0}^{n-1} v_j f_n^{\omega \cdot j}$ for all $\omega \in [n] = \{0, \dots, n-1\} \subset \mathbb{N}$. Similarly, the **Inverse Discrete Fourier Transform (IDFT)** of a vector $\mathbf{v} \in \mathbb{C}^n$ is

$$\widehat{\mathbf{v}}^{-1} := F^{-1}\mathbf{v} = F^*\mathbf{v}.$$

As we shall see, the DFT walks hand in hand with our next definition.

Exercise 3.4.13. Suppose $p, n \in \mathbb{N}$ are such that $n/p \in \mathbb{N}$ (i.e., p divides n). Given $\mathbf{u} \in \mathbb{C}^p$, let $\mathbf{v} \in \mathbb{C}^n$ be a longer vector with entries given by

$$v_j = \begin{cases} u_{pj/n} & \text{if } j \equiv 0 \pmod{(n/p)} \\ 0 & \text{else} \end{cases},$$

and let $\mathbf{w} \in \mathbb{C}^n$ be another longer vector with entries given by $w_j = u_{j \bmod p}$. Compute the n -length DFTs $\widehat{\mathbf{v}}, \widehat{\mathbf{w}} \in \mathbb{C}^n$ in terms of the p -length DFT of $\mathbf{u}$.

Exercise 3.4.14. Let $a, b, c \in [n]$ be such that a is invertible modulo n .⁴ Furthermore, suppose that $\mathbf{u}, \mathbf{v} \in \mathbb{C}^n$ satisfy

$$v_j = \mathbb{E}^{\frac{2\pi i c j}{n}} u_{aj+b \bmod n} = f_n^{-cj} u_{aj+b \bmod n}$$

for all $j \in [n]$. Write $\widehat{v}_\omega$ in terms of one or more entries of $\widehat{\mathbf{u}}$ for a given $\omega \in [n]$. How does a affect the entries of $\widehat{\mathbf{v}}$ when $c = b = 0$? How does b affect the entries of $\widehat{\mathbf{v}}$ when $a = 1$ and $c = 0$? How does c affect the entries of $\widehat{\mathbf{v}}$ when $a = 1$ and $b = 0$?

The **discrete (circular) convolution** of two vectors $\mathbf{u}, \mathbf{v} \in \mathbb{C}^n$, denoted by $\mathbf{u} \star \mathbf{v} \in \mathbb{C}^n$, is defined entrywise via

$$(\mathbf{u} \star \mathbf{v})_k := \sum_{j=0}^{n-1} u_j \cdot v_{(k-j) \bmod n} = \sum_{j=0}^{n-1} u_{j \bmod n} \cdot v_{(k-j) \bmod n}$$

for all $k \in [n]$. Note that, in fact, $\mathbf{u} \star \mathbf{v} = \text{circ}(\mathbf{v})\mathbf{u}$ for all $\mathbf{u}, \mathbf{v} \in \mathbb{C}^n$.

Example 3.4.9. Let $\mathbf{u}, \mathbf{v} \in \mathbb{C}^4$. Then,

$$\mathbf{u} \star \mathbf{v} = \begin{pmatrix} \sum_{j=0}^3 u_j v_{-j \bmod 4} \\ \sum_{j=0}^3 u_j v_{1-j \bmod 4} \\ \sum_{j=0}^3 u_j v_{2-j \bmod 4} \\ \sum_{j=0}^3 u_j v_{3-j \bmod 4} \end{pmatrix} = \begin{pmatrix} v_0 & v_3 & v_2 & v_1 \\ v_1 & v_0 & v_3 & v_2 \\ v_2 & v_1 & v_0 & v_3 \\ v_3 & v_2 & v_1 & v_0 \end{pmatrix} \begin{pmatrix} u_0 \\ u_1 \\ u_2 \\ u_3 \end{pmatrix} = \text{circ}(\mathbf{v})\mathbf{u}.$$

The discrete convolution has the following useful relationship with the discrete Fourier transform.

Theorem 3.4.10. Let $\mathbf{u}, \mathbf{v} \in \mathbb{C}^n$. Then

$$(\widehat{\mathbf{u} \star \mathbf{v}})_\omega = \sqrt{n} \widehat{u}_\omega \widehat{v}_\omega \quad (3.24)$$

holds for all $\omega \in [n]$.

Proof: To obtain (3.24) we compute

$$(\widehat{\mathbf{u} \star \mathbf{v}})_\omega = \frac{1}{\sqrt{n}} \sum_{k=0}^{n-1} (\mathbf{u} \star \mathbf{v})_k f_n^{\omega \cdot k} = \frac{1}{\sqrt{n}} \sum_{k=0}^{n-1} \left(\sum_{j=0}^{n-1} u_j \cdot v_{(k-j) \bmod n} \right) f_n^{\omega \cdot k}.$$

Exchanging the final double sum we obtain that

$$(\widehat{\mathbf{u} \star \mathbf{v}})_\omega = \frac{1}{\sqrt{n}} \sum_{j=0}^{n-1} u_j f_n^{\omega \cdot j} \left(\sum_{k=0}^{n-1} v_{(k-j) \bmod n} f_n^{\omega \cdot (k-j)} \right) = \sqrt{n} \widehat{u}_\omega \widehat{v}_\omega.$$

⁴A value $a \in [n]$ is invertible modulo n if there exists an $h \in [n]$ such that $a h \equiv 1 \pmod{n}$. Any $a \in [n]$ that is relatively prime to n will be invertible modulo n by the Fermat-Euler Theorem (see, e.g., [25, Theorem 2.8]).

Here we have used the fact that $f_n^{\ell \cdot n} = 1$ for all $\ell \in \mathbb{Z}$ so that $f_n^{\omega \cdot (k-j)} = f_n^{\omega \cdot ((k-j) \bmod n)}$ always holds (see, e.g, Exercise 3.4.9). $\square$

Exercise 3.4.15. Show that $\text{circ}(\mathbf{u})\mathbf{v} = \mathbf{v} \star \mathbf{u} = \mathbf{u} \star \mathbf{v} = \text{circ}(\mathbf{v})\mathbf{u}$ holds for all $\mathbf{u}, \mathbf{v} \in \mathbb{C}^n$.

Theorem 3.4.10 tells us that the DFT of the convolution of two vectors is, up to rescaling by $\sqrt{n}$, equal to the entrywise product of the DFTs of the two vectors. Using this relationship we can compute the discrete convolution of $\mathbf{u}$ and $\mathbf{v}$ using their DFTs. Let $\mathbf{u} \odot \mathbf{v} \in \mathbb{C}^n$ denote the entrywise (or Hadamard) product of the two vectors $\mathbf{u}, \mathbf{v} \in \mathbb{C}^n$. That is, let

$$(\mathbf{u} \odot \mathbf{v})_j := u_j v_j$$

for all $j \in [n]$. Theorem 3.4.10 now directly implies that

$$\mathbf{u} \star \mathbf{v} = \sqrt{n} \widehat{\mathbf{u} \odot \mathbf{v}}^{-1} = \sqrt{n} F^* (F\mathbf{u} \odot F\mathbf{v}). \quad (3.25)$$

Note that the last expression of (3.25) could be computed quickly if we could find a way to quickly calculate both $F\mathbf{u}$ and $F^*\mathbf{u}$ for any given $\mathbf{u}$. This is in fact possible as we shall see in Section 3.4.3.

The following additional fact relating IDFT matrices to circulant matrices is closely related to Theorem 3.4.10: Every column of the IDFT matrix F^* is an eigenvector of every circulant matrix. As a result, the $n \times n$ IDFT matrix F^* simultaneously diagonalizes this entire class of $n \times n$ matrices.

Theorem 3.4.11. Let $\mathbf{v} \in \mathbb{C}^n$. Every column of $F^* \in \mathbb{C}^{n \times n}$ is an eigenvector of $\text{circ}(\mathbf{v})$.

Proof. Let $\mathbf{u} = F^* \mathbf{e}_j \in \mathbb{C}^n$ be the j^{th} column of F^* . By (3.25),

$$\begin{aligned} \text{circ}(\mathbf{v})\mathbf{u} &= \mathbf{u} \star \mathbf{v} = \sqrt{n} F^* (F\mathbf{u} \odot F\mathbf{v}) = \sqrt{n} F^* ((FF^* \mathbf{e}_j) \odot F\mathbf{v}) \\ &= \sqrt{n} F^* (\mathbf{e}_j \odot \widehat{\mathbf{v}}) = \sqrt{n} F^* (\widehat{v}_j \mathbf{e}_j) = \sqrt{n} \widehat{v}_j \mathbf{u} \end{aligned}$$

Thus, the j^{th} column of F^* is an eigenvector of $\text{circ}(\mathbf{v})$ with eigenvalue $\sqrt{n} \widehat{v}_j$. $\square$

Exercise 3.4.16. Let $\mathbf{v} \in \mathbb{C}^n$. Show that $\text{circ}(\mathbf{v}) \in \mathbb{C}^{n \times n}$ is invertible if and only if $\widehat{v}_\omega \neq 0$ for all $\omega \in [n]$.

Exercise 3.4.17. Order the Fourier coefficients of $\mathbf{v} \in \mathbb{C}^n$ by magnitude so that

$$|\widehat{v}_{\omega_0}| \geq |\widehat{v}_{\omega_2}| \geq \cdots \geq |\widehat{v}_{\omega_{n-1}}|.$$

Prove that the j^{th} singular value of $\text{circ}(\mathbf{v}) \in \mathbb{C}^{n \times n}$ satisfies $\sigma_j(\text{circ}(\mathbf{v})) = \sqrt{n} |\widehat{v}_{\omega_j}|$.

One important consequence of the proof above is that the DFT gives us an easy way to compute the eigenvalues of all circulant matrices. Of even more consequence, though, is that (3.25) can also be used in many other applications where convolutions naturally appear. We have already seen, e.g., that the Toeplitz weight matrices of convolutional layers of neurons can be embedded into circulant matrices (recall definition 3.4.3 and (3.21)). Hence, (3.25) can potentially help evaluate convolutional layers of neurons more quickly via (3.22). In addition, convolutions also appear in numerous other important applications, two of which we will briefly discuss next.

Example 3.4.12 (Deblurring). *Consider the following “deblurring” problem: given $\mathbf{u} \star \mathbf{v} \in \mathbb{C}^n$ (the blurry signal) and knowledge of the blur kernel $\mathbf{v} \in \mathbb{C}^n$ (e.g., a Gaussian blur kernel), recover the unblurred signal $\mathbf{u} \in \mathbb{C}^n$. Such problems are common in imaging applications where a blurred image can indeed be thought of as a crisp/unblurred image convolved with a blur kernel. The question then becomes how one can try to “undo the blur” in order to get $\mathbf{u} \in \mathbb{C}^n$ back from its blurry version $\mathbf{u} \star \mathbf{v} \in \mathbb{C}^n$.*

Somewhat amazingly, this is easy to do efficiently if we have both the blurry signal $\mathbf{u} \star \mathbf{v} \in \mathbb{C}^n$ and knowledge of how the original image was likely blurred (i.e., we also know $\mathbf{v} \in \mathbb{C}^n$). In that case one can compute

$$\hat{u}_\omega = \frac{\widehat{\mathbf{u} \star \mathbf{v}}_\omega}{\hat{v}_\omega}$$

for all $\omega \in [n]$, and then set $\mathbf{u} = F^ \hat{\mathbf{u}}$. This of course assumes that the Fourier coefficients $\hat{v}_\omega \neq 0$ for all $\omega \in [n]$. If there are zero Fourier coefficients, then one can instead note that we are equivalently simply trying to solve the linear system $\text{circ}(\mathbf{v})\mathbf{u} = \mathbf{u} \star \mathbf{v}$ for $\mathbf{u} \in \mathbb{C}^n$. In such a case we can instead always find an approximate solution by returning, e.g., the least-squares estimate $\tilde{\mathbf{u}} = \text{circ}(\mathbf{v})^\dagger (\mathbf{u} \star \mathbf{v}) = P_{C(\text{circ}(\mathbf{v}))} \mathbf{u}$ (recall Section 3.1.2). Furthermore, an SVD of $\text{circ}(\mathbf{v})$, and therefore $\text{circ}(\mathbf{v})^\dagger$, can be constructed efficiently using Theorem 3.4.11.*

Example 3.4.13 (Polynomial Multiplication). *Convolutions also appear naturally as part of polynomial multiplication. Let $q(x) = \sum_{j=0}^{n-1} q_j x^j$ and $r(x) = \sum_{j=0}^{n-1} r_j x^j$ be two polynomials. Then $t(x) = q(x) \cdot r(x)$ is a polynomial of degree $\leq 2n-2$ that can be expressed as $t(x) = \sum_{j=0}^{2n-2} t_j x^j$. Writing the coefficients of q and r as vectors $\mathbf{q}, \mathbf{r} \in \mathbb{C}^n$, respectively, and the coefficients of t as a vector $\mathbf{t} \in \mathbb{C}^{2n-1}$, we have that*

$$\mathbf{t} = \begin{pmatrix} \mathbf{q} \\ \mathbf{0} \end{pmatrix} \star \begin{pmatrix} \mathbf{r} \\ \mathbf{0} \end{pmatrix}.$$

For example, when $n = 3$ we have that

$$\begin{aligned} t(x) &= (q_2 x^2 + q_1 x + q_0)(r_2 x^2 + r_1 x + r_0) \\ &= \underbrace{q_2 r_2}_{t_4} x^4 + \underbrace{(q_2 r_1 + q_1 r_2)}_{t_3} x^3 + \underbrace{(q_2 r_0 + q_1 r_1 + q_0 r_2)}_{t_2} x^2 + \underbrace{(q_1 r_0 + q_0 r_1)}_{t_1} x + \underbrace{q_0 r_0}_{t_0}. \end{aligned}$$

In vector form this corresponds to

$$\begin{pmatrix} t_0 \\ t_1 \\ t_2 \\ t_3 \\ t_4 \end{pmatrix} = \begin{pmatrix} q_0 & 0 & 0 & q_2 & q_1 \\ q_1 & q_0 & 0 & 0 & q_2 \\ q_2 & q_1 & q_0 & 0 & 0 \\ 0 & q_2 & q_1 & q_0 & 0 \\ 0 & 0 & q_2 & q_1 & q_0 \end{pmatrix} \begin{pmatrix} r_0 \\ r_1 \\ r_2 \\ 0 \\ 0 \end{pmatrix} = \text{circ} \left(\begin{pmatrix} \mathbf{q} \\ 0 \end{pmatrix} \right) \begin{pmatrix} r_0 \\ r_1 \\ r_2 \\ 0 \\ 0 \end{pmatrix} = \begin{pmatrix} q_0 \\ q_1 \\ q_2 \\ 0 \\ 0 \end{pmatrix} \star \begin{pmatrix} r_0 \\ r_1 \\ r_2 \\ 0 \\ 0 \end{pmatrix}.$$

Hence, if we can compute (I)DFTs quickly then we can also multiply polynomials quickly via (3.25).

Exercise 3.4.18. Consider the “finite difference” matrix $D_2 \in \mathbb{R}^{n \times n}$ whose entries are given by

$$(D_2)_{i,j} = \begin{cases} -2 & \text{if } i = j \\ 1 & \text{if } (i - j) \equiv 1 \pmod{n} \\ 1 & \text{if } (i - j) \equiv n - 1 \pmod{n} \\ 0 & \text{otherwise} \end{cases}. \quad (3.26)$$

This is an example of a circulant matrix. Show that $FD_2 = EF$, where $E \in \mathbb{R}^{n \times n}$ is a diagonal matrix with entries given by

$$(E)_{i,j} = \begin{cases} 2 \cos(2\pi j/n) - 2 & \text{if } i = j \\ 0 & \text{if } i \neq j \end{cases}. \quad (3.27)$$

Exercise 3.4.19. Let $D_{2r} \in \mathbb{R}^{n \times n}$ be defined by $D_{2r} := D_2^r$. Use the previous exercise to show that $FD_{2r} = E^r F$ for all $r \in \mathbb{Z}^+$.

As we will discuss in the next section, there is indeed a fast algorithm for computing both $F\mathbf{u}$ and $F^*\mathbf{u}$ for all $\mathbf{u} \in \mathbb{C}^n$. As a result, there are fast (i.e., computationally efficient) algorithms based on (3.25) for rapidly computing the convolutions involved in all of the applications mentioned in this section. Before explaining how any of these fast algorithms work, however, let's first briefly discuss what we actually mean when we say an algorithm is “fast”.

Big- $\mathcal{O}$ Notation and the Basic Art of Runtime Analysis

Throughout this text we will approach runtime discussions/analysis by counting six general types of atomic computational operations which we will assume any reasonable computer can do in a constant amount of time. These six types of constant-cost operations are:

1. Assigning a complex value to/reading a complex value from a variable or vector entry (e.g., setting $x_j = y \in \mathbb{C}$).
2. Adding/subtracting two machine numbers (e.g., adding any two real or complex numbers to a fixed precision).

3. Multiplying/dividing two machine numbers (e.g., multiplying any two real or complex numbers to a fixed precision).
4. Comparing two machine numbers (e.g., deciding whether one real number is larger/smaller/equal to another real number).
5. Evaluating basic logical expressions and conditional statements (e.g., deciding if “(boolean value A) AND/OR (boolean value B)” is True or False).
6. Evaluating simple functions $f : \mathbb{R} \rightarrow \mathbb{R}$ to a fixed precision. Herein, this class of “simple functions” includes (i) functions with rapidly convergent Maclaurin (i.e., 0-centered Taylor) series expansions such as the exponential, sine, and cosine functions, (ii) related complex-valued functions like $e^{i\theta} = \cos(\theta) + i \sin(\theta)$, and (iii) other rapidly approximable functions such as $f(t) = t^\alpha$ for a given (e.g., non-integer) $\alpha \in \mathbb{R}$.

Looking at the “constant-cost” operations above the invested reader’s eyebrows should be at least slightly raised. The sixth type of operation (evaluating simple functions) seems particularly fishy, doesn’t it?⁵ Even the second type of operation (i.e., simple addition) being “constant-cost” should inculcate suspicion in anyone who was expected to add 6 digit numbers to one another by hand in elementary school. I urge anyone who is not skeptical to grab a piece of chalk and investigate the claim that adding two 300 digit numbers together is the same “constant-cost” operation as adding 9 to 8.⁶ That said, let me urge you to allow the escape clause “to a fixed precision”, as well as the related term “machine numbers”, to save you from your skepticism, at least enough to believe that there is indeed some value to such simple types of operation counts.

Generally speaking, a digital computer can only guarantee the calculation of a fixed number of the leading digits of any real number one aims to compute/store. This is simply a fact of life. All of our algorithms here (or any others you see that are analyzed in a similar way) only guarantee you numerical answers up to some precision, or number of digits of accuracy – if that number of digits is not enough to be meaningful, then the algorithms are computing garbage. If, however, the answer you are after can be expressed accurately enough to satisfy you by its most significant, e.g., ~ 16 decimal digits, then the type of accounting we do here will be completely adequate for you. Even if you want many more digits of accuracy, though, a computer algorithm that needs to use only a few higher-precision operations will still be much faster to execute than one that uses many more higher-precision operations. As a result, even if our operation counts don’t truly

⁵We urge the interested reader to consult, e.g., [23, Chapters 1 and 3] to learn why this sixth type of constant-cost operation is indeed not *too* fishy after all...

⁶Really even adding two numbers should not be considered “constant-time” if you are doing serious computations involving, e.g., large number (and, therefore, extended precision) arithmetic. More precisely, the complexity of addition should depend on the the number of digits in each sum that you want to be able to correctly compute. Of course, this type of more complicated accounting then only gets more involved as you consider the other types of operations above.

represent an algorithm's runtime complexity with 100% accuracy in all cases (they don't), they do at least correlate well enough to be informative.⁷

Example 3.4.14 (Matrix-vector Multiplication). *As an illustrative example, let's consider the runtime complexity of computing the matrix-vector product $A\mathbf{x} \in \mathbb{C}^m$ for an arbitrary matrix $A \in \mathbb{C}^{m \times n}$ and vector $\mathbf{x} \in \mathbb{C}^n$. Noting that each entry of $\mathbf{y} = A\mathbf{x} \in \mathbb{C}^m$ is computed by*

$$y_j = (A\mathbf{x})_j = \sum_{k \in [n]} A_{j,k} x_k,$$

we can see that calculating $y_j \in \mathbb{C}$ requires $4n$ operations (we must read the $2n$ $A_{j,k}/x_k$ values into memory and then perform n multiplications, $n - 1$ additions, and finally one assignment of the correct value to y_j). Given that we must compute y_j for all $j \in [m]$ in order to calculate $\mathbf{y} = A\mathbf{x}$ we can now conclude that computing $\mathbf{y}$ will require at most $4nm$ constant-cost operations.⁸

In the example above the constant 4 we ended up with matters much less in general than the parameters m and n which will be significantly larger than 4 for big matrices A . As a result, it's standard practice to simplify operation counts by ignoring all such constants via big- $\mathcal{O}$ notation.

Definition 3.4.15 (Big- $\mathcal{O}$ Notation). *Let $f, g : (0, 1)^n \times (1, \infty)^m \rightarrow [0, \infty)$ be two functions of $n + m \geq 1$ variables for nonnegative $n, m \in \mathbb{Z}$. We say that f is $\mathcal{O}(g)$ if there exists a constant $C \in [1, \infty)$ and values $(\delta_0, \dots, \delta_{n-1}) \times (y_0, \dots, y_{m-1}) \in (0, 1)^n \times (1, \infty)^m$ such that*

$$f(\epsilon_0, \dots, \epsilon_{n-1}, x_0, \dots, x_{m-1}) \leq Cg(\epsilon_0, \dots, \epsilon_{n-1}, x_0, \dots, x_{m-1})$$

whenever $\epsilon_j < \delta_j$ and $x_k > y_k$ hold for all $j \in [n]$ and $k \in [m]$.

We can now see from Example 3.4.14 that computing $A\mathbf{x}$ can always be done using $\mathcal{O}(mn)$ operations (check!). The following example will also help illustrate the proper use of big- $\mathcal{O}$ notation.

Example 3.4.16. *Consider $f, g : (0, 1)^2 \times (1, \infty)^2 \rightarrow \mathbb{R}^+$ given by $f(\epsilon, \delta, m, n) = n/\epsilon + m/\delta + 100n + 200$ and $g(\epsilon, \delta, m, n) = \frac{\max\{m, n\}}{\min\{\epsilon, \delta\}}$. We will show that f is $\mathcal{O}(g)$.*

⁷We urge the interested reader to consult, e.g. [23, Chapter 2] and [10, Chapters 2 and 3] to learn more about numerical precision, machine numbers, and algorithmic runtime analysis. To begin understanding how one might make complexity analysis more rigorous one can also consult, e.g., [26].

⁸Here we say that computing $\mathbf{y}$ will require *at most* $4nm$ constant cost operations because we have demonstrated a way to compute $\mathbf{y}$ using this number of operations. However, there might be better ways to do it that use fewer operations by, e.g., avoiding rereading x_k multiple times for every different y_j calculation.

To start we note that $1 < \frac{1}{\min\{\epsilon, \delta\}}$ holds for all $\epsilon, \delta \in (0, 1)$. Similarly, $1 < \max\{m, n\}$ holds for all $m, n \in (1, \infty)$. Hence,

$$\begin{aligned} f(\epsilon, \delta, m, n) &= n/\epsilon + m/\delta + 100n + 200 \leq \max\{m, n\} \left(\frac{1}{\epsilon} + \frac{1}{\delta} + 100 \right) + 200 \\ &\leq \max\{m, n\} \left(\frac{2}{\min\{\epsilon, \delta\}} + 100 \right) + 200 \leq \max\{m, n\} \left(\frac{102}{\min\{\epsilon, \delta\}} \right) + 200 \\ &\leq 302 \left(\frac{\max\{m, n\}}{\min\{\epsilon, \delta\}} \right) = 302 \cdot g(\epsilon, \delta, m, n) \end{aligned}$$

holds for all $(\epsilon, \delta) \times (m, n) \in (0, 1)^2 \times (1, \infty)^2$. As a consequence, we can see that f is indeed $\mathcal{O}(g)$ in accordance with Definition 3.4.15.

Exercise 3.4.20. Let $g : (0, 1) \times [1, \infty)^2 \rightarrow [0, \infty)$ be given by $g(\epsilon, x, y) = \frac{x \log y}{\epsilon}$. Which of these functions $f : (0, 1) \times [1, \infty)^2 \rightarrow [0, \infty)$ are $\mathcal{O}(g)$? Explain your answers.

(a) $f(\epsilon, x, y) = \frac{300x \log y}{\epsilon} + 50$

(b) $f(\epsilon, x, y) = 500x^{0.34} + 600/\epsilon + 10^6 \log y$

(c) $f(\epsilon, x, y) = \frac{0.2x}{\epsilon^2} + \log y$

(d) $f(\epsilon, x, y) = \frac{20\sqrt{x} \log(100 + \ln(y))}{\sqrt{\epsilon}}$

Exercise 3.4.21. Let $A \in \mathbb{C}^{m \times n}$ and $B \in \mathbb{C}^{n \times p}$. Show that computing $AB \in \mathbb{C}^{m \times p}$ can be done using $\mathcal{O}(mnp)$ operations.⁹

3.4.3 The Fast Fourier Transform (FFT)

As seen above, computing the DFT of a vector $\mathbf{v} \in \mathbb{C}^n$ requires the computation of $F\mathbf{v}$. Computing $F\mathbf{v}$ directly via a generic matrix-vector multiply as per Example 3.4.14 uses $\mathcal{O}(n^2)$ operations. In this section we will discuss the Fast Fourier Transform (FFT) algorithm which can compute the DFT of a vector using only $\mathcal{O}(n \log n)$ operations. Though this reduction in computational complexity might seem slight at first glance, this speedup has had such far reaching impacts that the FFT has been lauded as one of the ten most important algorithmic developments of the twentieth century as a result [8].¹⁰

The FFT was first published and analyzed as a computer algorithm by Cooley and Tukey in 1965 [9], despite similar techniques being utilized much earlier (e.g., by Gauss and many others [17]). Cooley and Tukey's algorithm is particularly efficient for vector

⁹There are in fact faster (though not terribly practical) matrix multiplication algorithms out there for arbitrary matrices! We direct the interested reader to, e.g., [10, Chapter 28], [33, 21, 35].

¹⁰The QR decomposition discussed in Section 2.4 also made the list of the top 10 most important algorithm developments by the way!

dimensions, n , whose prime factorizations contain only small prime factors. Later variants of the FFT [2, 28] allowed the FFT to also be utilized effectively for vector sizes whose prime factorizations contain larger primes. This section has primarily followed [9, 2, 28]. For more information on Fourier methods and algorithms we recommend that the interested reader consult the relevant chapters of [27], [23], [10], or [3]. For a fast FFT implementation we recommend FFTW [14]. In what follows we will outline the recursive construction of the FFT algorithm via sum splitting techniques.

Let $\mathbf{u} \in \mathbb{C}^n$, and suppose that its dimension, n , has the prime factorization

$$n = p_1 \cdot p_2 \cdots p_m, \text{ where } p_1 \leq p_2 \leq \cdots \leq p_m \text{ are } n\text{'s prime factors.}$$

Choose $\omega \in [n]$. Recalling the definition of the DFT we have that

$$\hat{u}_\omega = \frac{1}{\sqrt{n}} \sum_{j=0}^{n-1} u_j f_n^{\omega \cdot j}. \quad (3.28)$$

By splitting the sum (3.28) for $\hat{u}_\omega$ into p_1 smaller subsums, one for each possible residue modulo p_1 , we can see that

$$\hat{u}_\omega = \frac{1}{\sqrt{n}} \sum_{k=0}^{p_1-1} f_n^{\omega \cdot k} \left(\sum_{j=0}^{\frac{n}{p_1}-1} u_{k+p_1 \cdot j} f_n^{\omega \cdot p_1 \cdot j} \right). \quad (3.29)$$

Let's now rewrite the internal sum of (3.29) in order to realize some progress.

Given $k \in [p_1]$, define $\mathbf{u}^{(k,p_1)} \in \mathbb{C}^{n/p_1}$ to be the vector whose entries are the entries of $\mathbf{u}$ having indexes that are congruent to k modulo p_1 ,

$$\left(\mathbf{u}^{(k,p_1)} \right)_j := u_{k+j \cdot p_1} \quad (3.30)$$

for all $j \in [n/p_1]$.¹¹ Our equation (3.29) for $\hat{u}_\omega$ now becomes

$$\hat{u}_\omega = \frac{1}{\sqrt{p_1}} \sum_{k=0}^{p_1-1} f_n^{\omega \cdot k} \left(\frac{1}{\sqrt{n/p_1}} \sum_{j=0}^{\frac{n}{p_1}-1} \left(\mathbf{u}^{(k,p_1)} \right)_j f_n^{p_1 \cdot \omega \cdot j} \right) \quad (3.31)$$

$$\begin{aligned} &= \frac{1}{\sqrt{p_1}} \sum_{k=0}^{p_1-1} f_n^{\omega \cdot k} \left(\frac{1}{\sqrt{n/p_1}} \sum_{j=0}^{\frac{n}{p_1}-1} \left(\mathbf{u}^{(k,p_1)} \right)_j f_{\frac{n}{p_1}}^{\left(\omega \bmod \frac{n}{p_1} \right) \cdot j} \right) \\ &= \frac{1}{\sqrt{p_1}} \sum_{k=0}^{p_1-1} f_n^{\omega \cdot k} \left(\widehat{\mathbf{u}^{(k,p_1)}} \right)_{\omega \bmod \frac{n}{p_1}}. \end{aligned} \quad (3.32)$$

¹¹Note that we used an integer divisor of n , i.e. p_1 , exactly to ensure that $\frac{n}{p_1} \in \mathbb{N}$.

Algorithm 6 FAST FOURIER TRANSFORM (FFT)

```

1: Input:  $\mathbf{u} \in \mathbb{C}^n$ , Dimension  $n$ , and  $n$ 's prime factorization  $p_1 \leq \dots \leq p_m$ 
2: Output:  $\hat{\mathbf{u}} \in \mathbb{C}^n$ 
3: if  $n = 1$  then
4:   Return  $\mathbf{u}$ 
5: end if
6: for  $k$  from 0 to  $p_1 - 1$  do
7:    $\widehat{\mathbf{u}^{(k,p_1)}} \leftarrow \text{FFT} \left( \mathbf{u}^{(k,p_1)}, \frac{n}{p_1}, p_2 \leq p_3 \leq \dots \leq p_m \right)$ 
8: end for
9: for  $\omega$  from 0 to  $n - 1$  do
10:   $\hat{u}_\omega \leftarrow \frac{1}{\sqrt{p_1}} \left( \sum_{k=0}^{p_1-1} f_n^{k\omega} \left( \widehat{\mathbf{u}^{(k,p_1)}} \right)_{\omega \bmod \frac{n}{p_1}} \right)$ 
11: end for
12: Return  $\hat{\mathbf{u}}$ 

```

For the sake of clarity we emphasize that the vector $\widehat{\mathbf{u}^{(k,p_1)}} \in \mathbb{C}^{n/p_1}$ in (3.32) is exactly $F\mathbf{u}^{(k,p_1)}$, where $F \in \mathbb{C}^{\frac{n}{p_1} \times \frac{n}{p_1}}$ is now the DFT matrix of size n/p_1 . We strongly recommend that you verify the equality of (3.31) – (3.32) for yourself before reading further.

At this point it's useful to ask ourselves what we've managed to accomplish by reformulating (3.28) into (3.32). Mainly, we can now compute $\hat{\mathbf{u}} \in \mathbb{C}^n$ with fewer operations than before by computing it in two steps. First, we compute $\widehat{\mathbf{u}^{(k,p_1)}} \in \mathbb{C}^{\frac{n}{p_1}}$ for all $k \in [p_1]$. Next, we use the vectors $\widehat{\mathbf{u}^{(0,p_1)}}, \dots, \widehat{\mathbf{u}^{(p_1-1,p_1)}}$ computed in the first step in order to compute each entry of $\hat{\mathbf{u}}$ via (3.32). The first step can be accomplished with p_1 matrix-vector multiplications, each of which can be computed using $\mathcal{O}(n^2/p_1^2)$ operations (recall that $\widehat{\mathbf{u}^{(k,p_1)}} = F\mathbf{u}^{(k,p_1)}$, where F is the DFT matrix of size n/p_1). Hence, the first step can be completed using $\mathcal{O}(n^2/p_1)$ total operations. Step two only requires $\mathcal{O}(p_1 n)$ total operations in order to finish calculating $\hat{\mathbf{u}}$, $\mathcal{O}(p_1)$ -operations for each entry $\hat{u}_\omega$. Putting it all together, we can see that (3.32) allows us compute $\hat{\mathbf{u}} \in \mathbb{C}^n$ using a grand total of $\mathcal{O}(p_1 n + n^2/p_1)$ operations, as opposed to computing it directly via (3.23) using $\sim n^2$ operations.

Although the computational gain obtained from (3.32) is modest when $p_1 \ll n$, it is important to note that the sum-splitting technique used to obtain it can now be employed again in order to compute each $\widehat{\mathbf{u}^{(k,p_1)}}$, $k \in [0, p_1)$, more quickly. That is, we may split up the sum for $(\widehat{\mathbf{u}^{(k,p_1)}})_\omega$ into p_2 additional sums, etc.. Repeatedly sum-splitting in this fashion leads to the recursive **Fast Fourier Transform (FFT)** shown in Algorithm 6. Analogous sum-splitting leads to the **Inverse Fast Fourier Transform (IFFT)** which can be obtained from Algorithm 6 by replacing line 10's $f_n^{k\omega}$ by $f_n^{-k\omega}$ and replacing each $\hat{\mathbf{u}}$ by a $\hat{\mathbf{u}}^{-1}$.

We are now ready to analyze the computational complexity of the FFT. Let T_n be the

total number of operations used by Algorithm 6 to compute $\hat{\mathbf{u}} \in \mathbb{C}^n$. In order to determine an equation for T_n we note that lines 6 – 8 of Algorithm 6 require $p_1 \cdot T_{\frac{n}{p_1}}$ operations while lines 9 – 11 use $\mathcal{O}(p_1 n)$ operations. Therefore we have

$$T_n = \mathcal{O}(p_1 n) + p_1 \cdot T_{\frac{n}{p_1}}.$$

However, Algorithm 6 is recursively invoked again to compute $\widehat{\mathbf{u}^{(0, p_1)}}, \dots, \widehat{\mathbf{u}^{(p_1-1, p_1)}}$ by sum-splitting in line 7. Taking this into account we can see that

$$T_{\frac{n}{p_1}} = \mathcal{O}\left(p_2 \frac{n}{p_1}\right) + p_2 \cdot T_{\frac{n}{p_1 p_2}}.$$

We now have that

$$T_n = \mathcal{O}(p_1 n) + p_1 \cdot \left(\mathcal{O}\left(\frac{p_2 n}{p_1}\right) + p_2 \cdot T_{\frac{n}{p_1 p_2}} \right) = \mathcal{O}(n(p_1 + p_2)) + p_1 p_2 \cdot T_{\frac{n}{p_1 p_2}}.$$

Repeating this recursive sum-splitting $j \leq m$ times shows us that

$$T_n = \mathcal{O}\left(n \cdot \sum_{\ell=1}^j p_\ell\right) + \prod_{\ell=1}^j p_\ell \cdot T_{\frac{n}{p_1 \cdots p_j}}.$$

Using that $T_1 = \mathcal{O}(1)$ (see Algorithm 6's lines 3 – 5) we have

$$T_n = \mathcal{O}\left(n \cdot \sum_{\ell=1}^m p_\ell\right) + \mathcal{O}(n) = \mathcal{O}(m \cdot p_m \cdot n). \quad (3.33)$$

Note that $m \leq \log_2 n$ while p_m is n 's largest prime factor. We have proven the following theorem in the course of the prior discussion.

Theorem 3.4.17. *Let $\mathbf{u} \in \mathbb{C}^n$ and suppose that $n \in \mathbb{N}$ has the prime factorization $n = p_1 \cdots p_m$, where $p_1 \leq p_2 \leq \cdots \leq p_m$ are the prime factors of n ordered from smallest to largest. Then, we may compute $\hat{\mathbf{u}} = F\mathbf{u}$ using $\mathcal{O}(n \cdot \sum_{\ell=1}^m p_\ell)$ operations.*

Theorem 3.4.17 tells us that the FFT can significantly speed up computation of the DFT. For example, if n is a power of 2 we'll have $m = \log_2 n$ and $p_m = 2$ leaving Algorithm 6 with an $\mathcal{O}(n \ln n)$ operation count. This is a clear improvement over the $\sim n^2$ operations required to in order to compute (3.23) directly. However, if n has large prime factors the improvement is less impressive. In the worst case, when n is itself a prime number, we have $m = 1$ and $p_1 = n$. This leaves Algorithm 6 with a $\mathcal{O}(n^2)$ runtime which, in practice, is even slower than the direct method (3.23).

The inability of Algorithm 6 to efficiently handle vectors with sizes containing large prime factors isn't a setback when one may dictate, with little or no repercussions, the

dimension of the vectors they work with. A popular choice is to simply force n to be a power of 2. However, sometimes one simply needs to compute the DFT of a vector whose size contains (or may contain) large prime factors. In the next subsection we discuss how to do this efficiently.

Exercise 3.4.22 (Computational Exercise). *Implement both the FFT and the IFFT for vectors of size 2^n , $n \in \mathbb{N}$. Produce a plot showing that each is indeed faster than the corresponding naive method for directly computing the (I)DFT of an arbitrary vector.*

3.4.4 The FFT for Vectors of Arbitrary Size

As discussed in the previous subsection, Algorithm 6 may not be a very efficient means of computing $\hat{\mathbf{u}} \in \mathbb{C}^n$ when n contains large prime factors. One way of addressing this issue is to rewrite $\hat{\mathbf{u}}$ as a discrete convolution of two vectors of a slightly larger dimension, $\tilde{n} > n$, that does contain only small prime factors. This discrete convolution can then be computed quickly using Algorithm 6 which will be efficient for vectors of dimension $\tilde{n}$.

Let $\omega \in [n]$. We may rewrite $\hat{u}_\omega$ as

$$\hat{u}_\omega = f_n^{\frac{\omega^2}{2}} f_n^{-\frac{\omega^2}{2}} \cdot \hat{u}_\omega = \frac{f_n^{\frac{\omega^2}{2}}}{\sqrt{n}} \cdot \sum_{j=0}^{n-1} u_j f_n^{\omega \cdot j - \frac{\omega^2}{2}} = \frac{f_n^{\frac{\omega^2}{2}}}{\sqrt{n}} \cdot \sum_{j=0}^{n-1} u_j f_n^{\frac{j^2}{2}} f_n^{\frac{-(\omega-j)^2}{2}} \quad (3.34)$$

Note that the last sum in (3.34) resembles a convolution. In order to make the resemblance more concrete we will define two new vectors.

Let $\tilde{n} = 2^{\lceil \log_2 n \rceil + 1} \geq 2n$ and define $\tilde{\mathbf{u}} \in \mathbb{C}^{\tilde{n}}$ by

$$\tilde{u}_j := \begin{cases} u_j \cdot f_n^{\frac{j^2}{2}} & \text{if } 0 \leq j < n \\ 0 & \text{if } n \leq j < \tilde{n} \end{cases},$$

and $\mathbf{v} \in \mathbb{C}^{\tilde{n}}$ by

$$v_h := \begin{cases} f_n^{\frac{-h^2}{2}} & \text{if } 0 \leq h < n \\ 0 & \text{if } n \leq h \leq \tilde{n} - n \\ f_n^{\frac{-(h-\tilde{n})^2}{2}} & \text{if } \tilde{n} - n < h < \tilde{n} \end{cases}.$$

Computing a weighted convolution of $\tilde{\mathbf{u}}, \mathbf{v} \in \mathbb{C}^{\tilde{n}}$ we can see that

$$\begin{aligned} \frac{f_n^{\frac{\omega^2}{2}}}{\sqrt{n}} \cdot (\tilde{\mathbf{u}} \star \mathbf{v})_\omega &= \frac{f_n^{\frac{\omega^2}{2}}}{\sqrt{n}} \sum_{j=0}^{\tilde{n}-1} \tilde{u}_j \cdot v_{(\omega-j) \bmod \tilde{n}} \\ &= \frac{f_n^{\frac{\omega^2}{2}}}{\sqrt{n}} \left(\sum_{j=0}^{\omega} \tilde{u}_j \cdot v_{\omega-j} + \sum_{j=\omega+1}^{n-1} \tilde{u}_j \cdot v_{(\omega-j)+\tilde{n}} \right) \\ &= \frac{f_n^{\frac{\omega^2}{2}}}{\sqrt{n}} \sum_{j=0}^{n-1} u_j \cdot f_n^{\frac{j^2}{2}} f_n^{\frac{-(\omega-j)^2}{2}}. \end{aligned}$$

Comparing to (3.34) now reveals that

$$\hat{u}_\omega = \frac{f_n^{\frac{\omega^2}{2}}}{\sqrt{n}} \cdot (\tilde{\mathbf{u}} \star \mathbf{v})_\omega \quad \forall \omega \in [n]. \quad (3.35)$$

Note that the convolution in (3.35) can be computed efficiently by the FFT and IFFT using (3.25) since $\tilde{n}$ is a power of two. Furthermore, $\tilde{n} \leq 4n$ by definition. Hence, we have established the following theorem.

Theorem 3.4.18. *Let $\mathbf{u} \in \mathbb{C}^n$. Then, both $\hat{\mathbf{u}}, \hat{\mathbf{u}}^{-1} \in \mathbb{C}^n$ can be calculated using $\mathcal{O}(n \ln n)$ operations.*

Exercise 3.4.23. *Finish the proof of Theorem 3.4.18 by arguing that $\hat{\mathbf{u}}^{-1} \in \mathbb{C}^n$, like $\hat{\mathbf{u}} \in \mathbb{C}^n$, can also always be calculated using $\mathcal{O}(n \ln n)$ operations. What changes need to be made to (3.34) – (3.35)?*

Theorem 3.4.18 generalizes Theorem 3.4.17 to handle all values of n efficiently. We are now in the position to declare that the DFT of any vector in $\mathbb{C}^n$ can be calculated using only $\mathcal{O}(n \ln n)$ operations! We are now prepared to prove that any (square) Toeplitz matrix has a fast matrix-vector multiplication algorithm.

3.4.5 Fast Matrix Multiplication for Toeplitz Matrices

Let $\mathbf{a} \in \mathbb{C}^{2n-1}$ and consider the $n \times n$ Toeplitz matrix generated by $\mathbf{a}$, $\text{Toep}_n(\mathbf{a}) \in \mathbb{C}^{n \times n}$. Given $\mathbf{v} \in \mathbb{C}^n$, we want to compute $\text{Toep}_n(\mathbf{a})\mathbf{v} \in \mathbb{C}^n$ using as few operations as possible. Recalling Lemma 3.4.7 we can begin by embedding $\text{Toep}_n(\mathbf{a})$ into a $2n \times 2n$ circulant matrix. Specifically, we have seen that the vector $\mathbf{c} = (a_{n-1}, a_{n-2}, \dots, a_0, 0, a_{2n-2}, \dots, a_n)^T \in \mathbb{C}^{2n}$ satisfies $(\text{circ}(\mathbf{c}))_{j,k} = (\text{Toep}_n(\mathbf{a}))_{j,k}$ for all $j, k \in [n]$. This then further implies that $(\text{Toep}_n(\mathbf{a})\mathbf{v})_j = \left(\text{circ}(\mathbf{c}) \begin{pmatrix} \mathbf{v} \\ \mathbf{0} \end{pmatrix} \right)_j$ for all $j \in [n]$ by (3.22). Hence, we can compute $\text{Toep}_n(\mathbf{a})\mathbf{v}$

Algorithm 7 FAST TOEPLITZ MATRIX MULTIPLICATION

- 1: **Input:** $\mathbf{a} \in \mathbb{C}^{2n-1}$, $\mathbf{v} \in \mathbb{C}^n$
 - 2: **Output:** $\text{Toep}_n(\mathbf{a})\mathbf{v} \in \mathbb{C}^n$
 - 3: $\mathbf{c} \leftarrow (a_{n-1}, a_{n-2}, \dots, a_0, 0, a_{2n-2}, \dots, a_n)^T \in \mathbb{C}^{2n}$
 - 4: Compute $\widehat{\mathbf{c}} \in \mathbb{C}^{2n}$ using the FFT
 - 5: Compute $\widehat{\begin{pmatrix} \mathbf{v} \\ \mathbf{0} \end{pmatrix}} \in \mathbb{C}^{2n}$ using the FFT
 - 6: $\widehat{\mathbf{b}} \leftarrow \sqrt{2n} \widehat{\mathbf{c}} \odot \widehat{\begin{pmatrix} \mathbf{v} \\ \mathbf{0} \end{pmatrix}} \in \mathbb{C}^{2n}$
 - 7: Compute $\mathbf{b} \in \mathbb{C}^{2n}$ using the IFFT
 - 8: Return $(b_0, b_1, \dots, b_{n-1})^T \in \mathbb{C}^n$
-

by computing the convolution $\mathbf{c} \star \begin{pmatrix} \mathbf{v} \\ \mathbf{0} \end{pmatrix}$. Finally, we further seen in (3.25) that this convolution can be computed efficiently via

$$\mathbf{c} \star \begin{pmatrix} \mathbf{v} \\ \mathbf{0} \end{pmatrix} = \sqrt{2n} F^* \left(\widehat{\mathbf{c}} \odot \widehat{\begin{pmatrix} \mathbf{v} \\ \mathbf{0} \end{pmatrix}} \right).$$

See Algorithm 7 for streamlined pseudocode.

Considering the runtime of Algorithm 7, we can see that forming both $\begin{pmatrix} \mathbf{v} \\ \mathbf{0} \end{pmatrix}$, $\mathbf{c} \in \mathbb{C}^{2n}$ can be accomplished in $\mathcal{O}(n)$ time. In addition, the (entrywise) Hadamard product of any two vectors in $\mathbb{C}^{2n}$, as well as selecting the first n entries of any vector $\mathbf{b} \in \mathbb{C}^{2n}$, can also always be accomplished in $\mathcal{O}(n)$ time. Finally, each (I)FFT can be computed using $\mathcal{O}(n \log n)$ operations by Theorem 3.4.18. Hence, Algorithm 7 will always utilize a total of $\mathcal{O}(n \log n)$ operations in order to compute $\text{Toep}_n(\mathbf{a})\mathbf{v} \in \mathbb{C}^n$. This is significantly faster than direct $\mathcal{O}(n^2)$ -time matrix-vector multiplication when n is large.

Fast Matrix Multiplication for Rectangular Toeplitz Matrices

Note that Algorithm 7 only applies to *square* Toeplitz matrices. A very natural next question then becomes what we can do if we instead need to quickly multiply a large non-square (rectangular) Toeplitz matrix, $\text{Toep}_{m,n}(\mathbf{a}) \in \mathbb{C}^{m \times n}$, against a vector $\mathbf{v} \in \mathbb{C}^n$? One simple strategy for handling such problems involves re-expressing the rectangular Toeplitz matrix in a block-matrix form, where each resulting block is a smaller square Toeplitz submatrix. The large rectangular matrix $\text{Toep}_{m,n}(\mathbf{a}) \in \mathbb{C}^{m \times n}$ can then be multiplied against a given $\mathbf{v} \in \mathbb{C}^n$ by combining the results of its smaller square Toeplitz submatrices multiplied against (appropriate pieces of) $\mathbf{v}$, each of which can now be computed efficiently using, e.g.,

Algorithm 7.¹²

Example 3.4.19. Let $\mathbf{a} \in \mathbb{C}^6$ and $\mathbf{v} \in \mathbb{C}^2$. Suppose that we want to compute $\text{Toep}_{5,2}(\mathbf{a})\mathbf{v} \in \mathbb{C}^5$. Instead of computing the result directly we can instead, e.g., decompose $\text{Toep}_{5,2}(\mathbf{a})$ into two 2×2 and two 1×1 Toeplitz submatrices, and then compute $\text{Toep}_{5,2}(\mathbf{a})\mathbf{v}$ using the resulting block-matrix form via

$$\text{Toep}_{5,2}(\mathbf{a})\mathbf{v} = \begin{pmatrix} a_4 & a_5 \\ a_3 & a_4 \\ a_2 & a_3 \\ a_1 & a_2 \\ a_0 & a_1 \end{pmatrix} \mathbf{v} = \begin{pmatrix} \begin{pmatrix} a_4 & a_5 \\ a_3 & a_4 \end{pmatrix} \\ \begin{pmatrix} a_2 & a_3 \\ a_1 & a_2 \end{pmatrix} \\ (a_0) \quad (a_1) \end{pmatrix} \mathbf{v} = \begin{pmatrix} \begin{pmatrix} a_4 & a_5 \\ a_3 & a_4 \end{pmatrix} \mathbf{v} \\ \begin{pmatrix} a_2 & a_3 \\ a_1 & a_2 \end{pmatrix} \mathbf{v} \\ (a_0) v_0 + (a_1) v_1 \end{pmatrix}.$$

Note that the fact that $\text{Toep}_{5,2}(\mathbf{a})$ has constant diagonals ensures that all of its square submatrices above are also Toeplitz.

Exercise 3.4.24. Suppose that we are given $\text{Toep}_{m,n}(\mathbf{a}) \in \mathbb{C}^{m \times n}$ and integers $1 \leq p \leq \min\{m, n\}$, $h \in [m - p + 1]$, and $\ell \in [n - p + 1]$. Let $A \in \mathbb{C}^{p \times p}$ be such that $A_{j,k} := (\text{Toep}_{m,n}(\mathbf{a}))_{h+j, \ell+k}$ for all $j, k \in [p]$. Show that A is Toeplitz.

Exercise 3.4.25. Let $p, n \in \mathbb{N}$, $\text{Toep}_{pn,n}(\mathbf{a}) \in \mathbb{C}^{pn \times n}$, and $\mathbf{v} \in \mathbb{C}^n$. Show that $\text{Toep}_{pn,n}(\mathbf{a})\mathbf{v} \in \mathbb{C}^{pn}$ can be computed in $\mathcal{O}(pn \log n)$ operations.

Exercise 3.4.26. Let $q(x) = \sum_{j=0}^{n-1} q_j x^j$ and $r(x) = \sum_{j=0}^{n-1} r_j x^j$ be two polynomials of degree at most $n - 1$. Let $t(x) = q(x) \cdot r(x)$ be their product. We know that $t(x)$ is a polynomial of degree at most $2n - 2$ which can be written as $t(x) = \sum_{j=0}^{2n-2} t_j x^j$. Write pseudocode for an algorithm that will compute the coefficients t_j of $t(x)$ in $\mathcal{O}(n \ln n)$ total operations.

Exercise 3.4.27. Let $g : [0, 1] \rightarrow \mathbb{R}$ be a twice continuously differentiable and periodic function. Any such g will have a Fourier series expansion of the form

$$g(x) = \sum_{\omega \in \mathbb{Z}} c_\omega e^{2\pi i \omega x} \quad \forall x \in [0, 1],$$

where the Fourier series coefficients $c_\omega \in \mathbb{C}$ satisfy (i) $c_\omega = \overline{c_{-\omega}}$ for all $\omega \in \mathbb{Z}$, and (ii) $\sum_{\omega \in \mathbb{Z}} |c_\omega| < \infty$. Let $\mathbf{u} \in \mathbb{R}^n$ be a vector whose entries are given by $u_j = g(j/n)$ for all $j \in [n]$. Show that the vector $F\mathbf{u} \in \mathbb{C}^n$ has entries

$$(F\mathbf{u})_j = \sqrt{n} \sum_{\omega \equiv j \pmod n} c_\omega.$$

¹²Of course, the best way to decompose a given $m \times n$ Toeplitz matrix into square Toeplitz submatrices depends on how m and n compare with one another.

Rapidly Evaluating Convolutional Layers of Neurons

In addition to having fewer parameters than general layers of neurons, we can now see that convolutional layers of neurons (recall Definition 3.4.3) also have other computational advantages. Consider, e.g., a convolutional layer of neurons $\ell : \mathbb{R}^N \rightarrow \mathbb{R}^N$ defined by $\ell(\mathbf{x}) := \sigma(W\mathbf{x} + \mathbf{b})$ with Toeplitz weight matrix $W = \text{Toep}_N(\mathbf{w}) \in \mathbb{R}^{N \times N}$. Evaluating $\ell(\mathbf{x})$ as part of a neural network forward-evaluation requires us to: (i) Compute $W\mathbf{x}$, (ii) add $\mathbf{b}$ to $W\mathbf{x}$ to compute $W\mathbf{x} + \mathbf{b}$, and then (iii) compute $\sigma(W\mathbf{x} + \mathbf{b})$ by applying the activation function $\sigma : \mathbb{R} \rightarrow \mathbb{R}$ to each entry of $W\mathbf{x} + \mathbf{b}$. Assuming that σ is a simple function, both steps (ii) and (iii) can always be accomplished in $\mathcal{O}(N)$ operations. The first step (i) therefore always dominates the layer's evaluation cost.

Focussing on step (i) above, we can see that it can be accomplished for $W = \text{Toep}_N(\mathbf{w})$ in $\mathcal{O}(N \log N)$ -operations via Algorithm 7. Hence, evaluating $\ell(\mathbf{x})$ can also be accomplished in $\mathcal{O}(N \log N)$ total operations in this case. In contrast, a general layer of neurons must generally rely on direct $\mathcal{O}(N^2)$ -time matrix-vector multiplication to complete step (i). Thus, convolutional layers of neurons require fewer operations to evaluate than general layers of neurons – yet another advantage of their Toeplitz structure!

3.5 ℓ^p -norms and Associated Inequalities

Proof of Minkowski Inequality follows [22]

IGNORE – STUFF FOR POSSIBLE(?) FUTURE INCLUSION:

– ℓ^p -norms and Holder inequality

- If A is then $\det(A) \in \mathbb{R}$ is the product of A' 's eigenvalues,

Of: A function $f : \mathbb{R} \rightarrow \mathbb{R}$ is convex on (asb) if $f(tx + (1-t)y) \leq tf(x) + (1-t)f(y)$ holds $\forall t \in [0, 1]$ and $x, y \in (a, b)$.

$\leadsto$ Note that the inequiling automatically holds for $t = 0$ and so we can always $f_0(u)$ on $t \in (0, 1)$ instead.

Sef: A function $f : \mathbb{R} \rightarrow \mathbb{R}$ is concave on (a, b) if

$(-f)(x) := -f(x)$ is convex on (a, b)

Lemma 1: f is convex on (a, b) iff $\forall x_1, x_2, x_3 \in (a, b)$ with $x_1 < x_2 < x_3$ we have

$$\frac{f(x_2) - f(x_1)}{x_2 - x_1} \leq \frac{f(x_3) - f(x_2)}{x_3 - x_2}$$

pf: Suppose welge that $x < y$. Set $x_1 = x, x_2 = tx + (1-t)y$ and $x_3 = y$. in this case are have

$$x_2 - x_1 = (t-1)x_1 + (1-t)x_3 = (1-t)(x_3 - x_1)$$

so that $t = \frac{x_3 - x_2}{x_3 - x_1}, (1-t) = \frac{x_2 - x_1}{x_3 - x_1}$, pase 2

and $\frac{1-t}{t} = \frac{x_2 - x_1}{x_3 - x_2}$.

$\leadsto$ Now we can see that f is convex on (a, b) iff

$$\begin{aligned} f(x_2) &\leq tf(x_1) + (1-t)f(x_3) \\ t(f(x_2) - f(x_1)) &\leq (1-t)(f(x_3) - f(x_2)) \\ f(x_2) - f(x_1) &\leq \frac{1-t}{t}(f(x_3) - f(x_2)) \end{aligned}$$

π

$$f(x_2) - f(x_1) \leq \left(\frac{x_2 - x_1}{x_3 - x_2} \right) (f(x_3) - f(x_2))$$

$\sim$) Dividing both sides of the last inequality on $x_2 - x$, proves the serult.

$\rightarrow$ Using this last lerman we can now prove that functions

We will need the Mean Value Theorem from calcular page 3 (see, eg.) Thun 7,20 in "Elementang Beel Analys is" by Thomson, Bruchuer, Brace Mean value Then: Suppose $f : \mathbb{R} \rightarrow \mathbb{R}$ is continous on $[a, b]$ and differentiable on (a, b) . Then $\exists c \in (a, b)$ s.t.

$$f'(c) = \frac{f(b) - f(a)}{b - a}$$

Thm: Suppose that $f : \mathbb{R} \rightarrow \mathbb{R}$ is continuous on $[a, b]$ and diff. on (a, b) with a non-decreasing derivative. Then, f is convex on (a, b) .

pt: Choose any $x_1 < x_2 < x_3$ you like in (a, b) . By the MVT $\exists c_1 \in (x_1, x_2)$ and $c_2 \in (x_2, x_3)$ s.t.

$$\frac{f(x_2) - f(x_1)}{x_2 - x_1} = f'(c_1) \leq f'(c_2) = \frac{f(x_3) - f(x_2)}{x_3 - x_2}.$$

The result now follows via lemma 1.

HWS Suppose the $f^{(1)}$ has 2 derivatives on (a, b) , s.t. on $[a, b]$ a $f''(x) \geq 0 \quad \forall x \in [a, b]$. Show that f is convex on (a, b) . (Hint: recall the FTC)

HW: Let $p \geq 1$. Show that given by $f(x) = X^p$ is convex on $(0, \infty)$.

Hw: Show that the natural logarithm, $\log : \mathbb{R} \rightarrow \mathbb{R}$ is concave.

Back to $\mathbb{C}$:

any incitant result for the function

Will now show $f(z) =$

Lemma(t): Let $p \geq 1$. Then

$$|(1 - \lambda)z + \lambda\omega|^p \leq (1 - \lambda)|z|^p + \lambda|\omega|^p$$

holds $\forall \lambda \in [0, 1]$ and $z, \omega \in \mathbb{R}$.

The result clearly holds if either z, ω is 0, so we assume both are non-zero. Let $g : \mathbb{R} \rightarrow \mathbb{R}$ be $g(x) = X^p$ and note that g is both non-decreasing and convex on $(0, \infty)$ [see Ha]

Using the properties of the magnitude function $|z| = \sqrt{(\operatorname{Re}(z))^2 + (\operatorname{Im}(z))^2}$ together with the fact that g is non-decreasing we have that

$$|(1 - \lambda)z + \lambda\omega| \leq |(1 - \lambda)z| + |\lambda\omega| = (1 - \lambda)|z| + \lambda|\omega| \Rightarrow$$

$$g(|(1 - \lambda)z + \lambda\omega|) \leq g((1 - \lambda)|z| + \lambda|\omega|) \leq g((1 - \lambda)|z|) + \lambda g(|\omega|) = (1 - \lambda)|z|^p + \lambda|\omega|^p$$

Since g is convex on $(0, \infty)$ we have that

$$|(1 - \lambda)z + \lambda\omega|^p \leq ((1 - \lambda)|z| + \lambda|\omega|)^p \leq (1 - \lambda)|z|^p + \lambda|\omega|^p \text{ as we wished to show.}$$

Def: Let $p \geq 1$. The p -norm of a vector $\vec{z} \in \mathbb{C}^n$ is

$$\|\vec{z}\|_p := \left(\sum_{j \in [n]} |z_j|^p \right)^{1/p}$$

- Note that we have seen l_1, l_2 , and l_∞ norms already. The ideal (after a little thought in the case of how it fits) special case of this more general definition. give it this

HW: Let $p \geq 1$. Begin showing that $\|\cdot\|_p : \mathbb{C}^n \rightarrow \mathbb{R}^+$ is inds a vector norm by establishing that

- (i) $\|\vec{z}\|_p = 0 \Leftrightarrow \vec{z} = \vec{0} \quad \forall z \in \mathbb{C}^n$
- (ii) $\|q\vec{z}\|_p = |q|\|\vec{z}\|_p \quad \forall q \in \mathbb{C} \text{ and } z \in \mathbb{C}^n$.

As usual, establishing the triangle inegulity (also known as the Minkowshi Ineguality) is the hardest Dont.

Thus: Let $p \geq 1$. Then $\|\vec{x} + \vec{y}\|_p \leq \|\vec{x}\|_p + \|\vec{y}\|_p$ $\left[\frac{pase6}{\forall \vec{x}, \vec{y} \in \mathbb{C}} \right]$

pf: This trivially holds if either $\vec{x}$ or $\vec{y}$ is $\vec{0}$, so we assum bothere unzers s.t. $c = \|\vec{x}\|_p + \|\vec{y}\|_p > 0$. Let $\tilde{x} = \vec{x}/c$ and $\tilde{y} = \vec{y}/c$. Nete that by our chice of C we have

$\|\tilde{x}\|_p = \lambda$ and $\|\tilde{y}\|_p = 1 - \lambda$ for some $\lambda \in (0, 1)$

Lett

that

$$\begin{aligned} \|\tilde{x} + \tilde{y}\|_p &\leq \|\tilde{x}\|_p + \|\tilde{y}\|_p \Leftrightarrow \|\lambda\tilde{x} + (1 - \lambda)\tilde{y}\|_p \leq 1 \\ &\Leftrightarrow \|\lambda\tilde{x} + (1 - \lambda)\tilde{y}\|_p^p \leq 1 \end{aligned}$$

We may row use Lemma (T) to verify this finel ineguliz by seein that

$$\begin{aligned} \|\lambda\tilde{x} + (1 - \lambda)\tilde{y}\|_p^p &= \sum_{j \in [n]} |\lambda\tilde{x}_j + (1 - \lambda)\tilde{y}_j|^p \\ &\leq \sum_{j \in [n]} \lambda |\tilde{x}_j|^p + (1 - \lambda) |\tilde{y}_j|^p \\ &= \lambda \|\tilde{x}\|_p^p + (1 - \lambda) \|\tilde{y}\|_p^p = 1. \end{aligned}$$

Hew: Show that $\|\vec{x}\|_\infty \leq \|\vec{x}\|_p \quad \forall p \in [1, \infty)$

Hw: Let $q \geq p \geq 1$. Show that $\|\vec{x}\|_q \leq \|\vec{x}\|_p$,

(Hint: Let $\vec{y} = \vec{x}/\|\vec{x}\|_c$ and prove that $\|\vec{y}\|_q \leq \|\vec{y}\|_p$)

Young's Inegulity for products & the Hölder Ineperlis

Lemma (Yonng's Faop.): Let $p, q \in (1, \infty)$ satisfy $\frac{1}{p} + \frac{1}{q} = 1$.

Then, $ab \leq \frac{ap}{p} + \frac{bq}{q} \quad \forall a, b \in [0, \infty)$.

The equalis is Wiviny drue if aor $b = 0$, 50 assame bothare

pf: Sine $\ln : (0, \infty) \rightarrow \mathbb{R}$ is concave on $(0, \infty)$ we havitive. that

$$-\ln \left(\frac{1}{p} q^p + \frac{b^q}{q} \right) = -\ln \left(\frac{1}{p} q^p + \left(1 - \frac{1}{p} \right) b^q \right)$$

$$\begin{aligned} &\leq -\frac{1}{p} \ln(a^p) - \left(1 - \frac{1}{p}\right) \ln(b^q) \\ &= -\ln(a) - \ln(b) = \ln(ab). \end{aligned}$$

Hence, $\ln(ab) \leq \ln\left(\frac{1}{p}a^p + \frac{b^q}{q}\right)$. Expocientiation both sides of this list inepr. proves the result.
 $\leadsto$ We can row prove probalh the inst usefil incurr related to general lp-norms. It, ergo gearelies the Carchy-Selwarz ineputility.

Theorem 3.5.1. *Thm (Hölder's Inegraliz)*

s.t. $\frac{1}{p} + \frac{1}{q} = 1$. Then, $|\langle \vec{x}, \vec{y} \rangle| \leq \sum_{j \in [n]} |\vec{x}_j| |\vec{y}_j| \leq$

$$\|\vec{x}\|_p \|\vec{y}\|_q \text{ holds } \forall \vec{x}, \vec{y} \in \mathbb{C}^n,$$

pf: The first inequality alwings holds by properties of the magnitude function $|\cdot| : \mathbb{R} \rightarrow \mathbb{R}^+$ (Chech this!). Thus, it suffices to check that

$$\sum_{j \in [n]} |x_j| |y_j| \leq \|\vec{x}\|_p \|\vec{y}\|_q.$$

Let $\tilde{x} = \frac{\vec{x}}{\|\vec{x}\|_p}$ and $\tilde{y} = \frac{\vec{y}}{\|\vec{y}\|_q}$. Then, by Yamis' In

$$\begin{aligned} \sum_{j \in [n]} \frac{|x_j| |y_j| \zeta}{w_j |w_j| \in [n]} |\tilde{x}_j| |\tilde{y}_j| &\leq \sum_{j \in [n]} \frac{|\tilde{x}_j|^p}{p} + \frac{|\tilde{y}_j|^q}{q} \\ &= \frac{\|\tilde{x}\|_p^p}{p} + \frac{\|\tilde{y}\|_q^q}{q} = \frac{1}{p} + \frac{1}{q} = 1 \end{aligned}$$

Multiolysio both sides les $\|\hat{x}\|_0 \|\hat{y}\|$. Suichor the broot.

Hw: Prove that $|\langle \vec{x}, \vec{y} \rangle| \leq \|\vec{x}\|_\infty \|\vec{y}\|_2$ also page 9 bold, $\forall z, z \in \mathbb{C}^n$.

HW: Frove that $\|\vec{x}\|_p \leq \|\vec{x}\|_\infty n^{1/p}$ for all $\vec{x} \in C^n, p \geq$

Hw! Plan that $\|\vec{x}\|_2 \leq \sqrt{\|\vec{x}\|_b \|\vec{x}\|_1} \quad \forall \vec{x} \in \mathbb{C}^n$

HW: Prove that $\|\vec{x}\|_{2p} \leq \sqrt{\|\vec{x}\|_{pq} \|\vec{y}\|_{pr}} \quad \forall p \geq 1, \vec{x} \in C^n$, and $r, g \in (0, 1)$ satisfing $\frac{1}{r} + \frac{1}{q} = 1$.

Maxix Opensor Norms

We can now semulite the $\|A\|_{2 \rightarrow 2}$ -norm from Chp.2 ?

$$\vec{x} \in \mathbb{C}^n, A \in \mathbb{C}^{m+1},$$

Lemma: Let $A \in \mathbb{C}^m$

If $\vec{x} = \vec{0}$ we're done, so assure $\vec{x} \neq \vec{0}$, Then

$$\|A\vec{x}\|_q = \frac{\|A\vec{x}\|_q}{\|\vec{x}\|_p} \|\vec{x}\|_p \leq \left(\sup_{\vec{x} \neq \vec{0}} \frac{\|A\vec{x}\|_q}{\|\vec{x}\|_p} \right) \|\vec{x}\|_p = \|A\|_{p \rightarrow q}$$

HV: Prove that $\|AB\|_{p \rightarrow p} \leq \|A\|_{p \rightarrow p} \|B\|_{p \rightarrow p} \quad \forall p \geq 1,$

$$A \in \mathbb{C}^{m+n}, B \in \mathbb{C}^{n \times P}$$

Hw: Prove that $\|AB\|_{p \rightarrow q} \leq \|A\|_{r \rightarrow q} \|B\|_{p \rightarrow r} \forall r \geq 1$. Lif $A \in C^{\text{mon}} \|A\|_{1+1} = \max_{k \in [n]} \sum_{j \in [m]} |A_{jk}|$.
Lemma

$$K \in [n], j \in [m]$$

HW: Prove that $\|A\|_{\infty \rightarrow \infty} = \max_{j \in [M]} \sum_{k \in [M]} |A_{jk}|$

Note that

$$\text{pf: } \|A\|_{1 \rightarrow 1} = \max_{\|A\vec{x}\|_1 = \max} \sum_{i \in C^m} | \sum_{L \in C(n)} A_{j,k} x_k \| \vec{x} \|_1 = 1 \quad \| \vec{x} \|_1 = 1 \quad \vec{j} \in [m] \quad \text{Lot}(n)$$

Furthumale, this upen bound is ackeved since $\|A_{e_k}^{-3}\|_j = \|A, DM \setminus K$

- Cite Bergh, Lotstron, "Interpolation Spaces", Springer 1976
- Application in "Fart Dimension Reduction Using Retemacher Series on Duel BCH Codes" by Ailar & Libuty

Thm (Discrete Riesz-Thorin) Let $A \in \mathbb{C}^{m \times n}$, $\theta \in [0, 1]$, a $\left\{ \begin{array}{l} q_1 p_1 q_1) p_2, q_2 \in [1, \theta] \text{ be s.t } \frac{1}{p} = \frac{\theta}{p_1} + \frac{(1-\theta)}{p_2} \\ \frac{1}{q} = \frac{\theta}{q_1} + \frac{(1-\theta)}{q_2} \end{array} \right.$ Then, $\|A\|_{p_1 q} \leq (\|A\|_{p_1 + q_1})^\theta (\|A\|_{p_2 + q_2})^{1-\theta}$

HW: Let $F \in \mathbb{C}^{n \times n}$ be the DFT madix of size n , shor that $\|F\|_{2 \rightarrow 2} = 1$ and

$$\|F\|_{1 \rightarrow \infty} = \frac{1}{\sqrt{n}}$$

$\rightsquigarrow$ Using Riesz-Tborin we can prove some addition. useful borm bounds for F .

$\rightsquigarrow$ Chorse $\theta \in [0, 1]$ and Sappose that

$$\begin{aligned} \frac{1}{p} &= \frac{\theta}{1} + \frac{(1-\theta)}{2} = \frac{1+\theta}{2} \\ \frac{1}{q} &= \frac{\theta}{\infty} + \frac{(1-\theta)}{2} = \frac{1-\theta}{2} \end{aligned}$$

Summing these expression, we can see that $\frac{1}{p} + \frac{1}{q} = 1$
and $\theta = 1 - \frac{2}{q}$

Applying Riesz-Thorin we have that

$$\begin{aligned} \|F\|_{p \rightarrow y} &\leq (\|A\|_{1-\infty})^{1-\frac{2}{q}} (\|A\|_{2-2})^{\frac{2}{q}} = n^{\frac{1}{q}-\frac{1}{2}} \\ &= n^{1/2-\frac{1}{p}} \end{aligned}$$

$\rightarrow$ Noting that since $\theta \in [0, 1]$ we have $p \in [1, 2]$ we obtain the following resalt

Thun (Discrete Hawsdoff Y Yang) Let $\rho \in [1, 2]$ and $q = \frac{p}{p-1}$ ($\frac{1}{0} := \infty$). Then $\|F\|_{p \rightarrow q} \leq n^{\frac{1}{2}-\frac{1}{p}}$;

HW: Let $\vec{V} \in \mathbb{C}^n$, Prove that and $\|\text{circ}(\vec{V})\|_{1+1} = \|\text{circ}(\vec{V})\|_{\infty+\infty} = \|\vec{V}\|_1$

HW: Let $\vec{V} \in \mathbb{C}^U$. Prove that

$$\|\text{cif}(\vec{v})\|_{p \rightarrow \infty} \leq \begin{cases} \|\vec{v}\|_1^{\frac{2}{p}-1} (n^{1/2} \|F - \vec{v}\|_\infty)^{2-\frac{2}{p}} & \forall p \in [1, 2] \\ \|V\|_1^{1-\frac{2}{p}} (n^{1/2} \|F - \tau\|_1)^2 / p & p \gg 1 \end{cases}$$

Hw: Prove the $\|F\vec{V}\|_\infty \leq \frac{1}{\sqrt{n}} \|\vec{V}\|_1 \forall \vec{V} \in \mathbb{C}^n$

Conclude that $\|\text{Circ}(\vec{v})\|_{p \rightarrow p} \leq \|\vec{v}\|_1 \forall \vec{v} \in \mathbb{C}^n$ and

Hw: (H) Prove that $\|\vec{x} * \vec{y}\|_p \leq \min\{\|\vec{x}\|_q \|\vec{y}\|_p, \|\vec{y}\|_1 \|\vec{x}\|_p\} = 1, \forall p \geq 1$

Discrete Konny's Convolution

H. Let $p, q \in (\phi, \infty)$ be s.t. $\frac{1}{p} + \frac{1}{q} = 1$. Show that $\|\vec{x} \times \vec{y}\|_\infty \leq \|\vec{x}\|_p \|\vec{y}\|_q \rightarrow \vec{x}, \vec{y} \in \mathbb{C}^n$.

$$\|\vec{x}\|_q \|\vec{y}\|_p\}$$

Let $\vec{V} \in \mathbb{C}^n$ and

Let $p, s, r, s \geq 1$ be s.t. $\frac{1}{p} + \frac{1}{s} = 1$. Then $\|\text{Circ}(\vec{V})\|_{1 \rightarrow p} \leq \|\vec{V}\|_p$ by $H\omega(t)$, and

$\|\text{circ}(\vec{v})\|_{g \rightarrow \infty} \leq \|\vec{V}\|_p$ by $H\omega(*)$. towingt

$\Rightarrow$ Applying Riesz - That is we have $\forall \theta \in [a, 1]$ that if $\frac{1}{n} = \frac{\theta}{1} + \frac{(1-\theta)}{n}, \frac{1}{s} = \frac{\theta}{p} + \frac{1-\theta}{\infty} = \frac{\theta}{p}$

Riesz-Tharin tells us that

$$\begin{aligned} \|\text{circ}(\vec{v})\|_{pr \rightarrow s} &\leq (\|\text{circ}(\vec{v})\|_{1 \rightarrow q_0}) (\|\text{cinc}(\vec{v})\|_{q=-\varepsilon}^\infty \\ &= \|\vec{V}\|_p. \end{aligned}$$

$\rightarrow$ Considering the relationships between Pg, r, s above we have that $\frac{1}{r} \in [1, \frac{1}{\varepsilon}] \Rightarrow r \in [1, q]$, and

$$\begin{aligned} \frac{1}{f} &= \theta + (1-\theta) \frac{1}{q} = \theta + \frac{1}{q} - \frac{\theta}{q} \\ &= \theta + \frac{1}{q} - \theta \left(1 - \frac{1}{p}\right) \\ &= \frac{1}{q} + \frac{\theta}{p} = \frac{1}{q} + \frac{1}{s}. \end{aligned}$$

$\rightarrow$ Putting This all together we obtain the following Yorrerete Incepa.): Let $\vec{V} \in \mathbb{C}^n$ and $p, q \geq 1$ be s.t. $\frac{1}{p} + \frac{1}{q} = 1$. Pick $r \in [1, q]$ and set $\frac{1}{s} = \frac{1}{r} - \frac{\theta}{q}$. Then $\|\text{circ}(\vec{v})\|_{r \rightarrow s} \leq \|\vec{v}\|_p$. HW: Let $\vec{u}, \vec{v} \in \mathbb{C}^n$ and $p, q = 1$ be s.t. $\frac{1}{p} + \frac{1}{1-1} = 1$, Then $\forall r \in [1, q]$.

ADD BELOW INTO SVD SECTION ?

DATA FITTING:

- Given $P = \{\vec{x}_1, \dots, \vec{x}_N\} \subseteq \mathbb{R}^D$.
- Our fitness measure for an affine subspace H is $R_\tau(H, P) = (\sum_{\vec{x}_j \in P} d(\vec{x}_j, H)^\tau)^{1/\tau}$ for some $\tau \in \mathbb{R}^+$. Here $d(\cdot, \cdot)$ is Hausdorff distance.
- Assume that P has mean $\frac{1}{N} \sum_{j=1}^N \vec{x}_j = \vec{0}$
- For $\tau = 2$ we get a least squares approximation to P .
- Review $\tau = 2$: This can be solved exactly in $O(ND \min\{N, D\})$ -time

Goal : Minimize $(R_2(H, P))^2 = \sum_{\vec{x} \in P} d(\vec{x}_j, H)^2$ over all $n < D$ dimensional subspace H .

- Let $X_P = (\vec{x}_1, \dots, \vec{x}_N) \in \mathbb{R}^{D \times N}$
- Represent an n -dimensional H (subspace) by a projection matrix $\Pi_H \in \mathbb{R}^{D \times D}$ (rank n) that projects onto H .

$$\begin{aligned}
 (R_2(H, P))^2 &= \sum_{\vec{x}_j \in P} d(\vec{x}_j, H)^2 \\
 &= \sum_{\vec{x}_j \in P} \|\vec{x}_j - \Pi_H \vec{x}_j\|_2^2 \\
 &= \|X_P - \Pi_H X_P\|_F^2 \text{ (Recall } \|A\|_F^2 = \sum \sum |a_{ij}|^2) \\
 &= \|(I - \Pi_H)X_P\|_F^2
 \end{aligned}$$

- We want to minimize this $\|\cdot\|_F$ over all H . Recall that

$$\begin{aligned}
\|A\|_F &= \sqrt{\text{trace}(A^T A)} \\
&= \sqrt{\text{trace}(V \Sigma^2 V^T)} \\
&= \sqrt{\text{trace}(\Sigma^2)} \text{ (when } A = U \Sigma V^T, \text{ the SVD of } A) \\
&= \sqrt{\sum_{j=1}^N \sigma_j(A)^2}
\end{aligned}$$

So we want to minimize $\sum_{j=1}^{\min(N,D)} \sigma_j((I - \Pi_H)X_P)^2$ over all H .

- If H is n -dimensional, $(I - \Pi_H)$ is $(D - n)$ -dimensional projection.
- Let $X_P = U \Sigma V^T$ (SVD of X_P).
- We should let $I - \Pi_H$ project onto the subspace spanned by $D - n$ columns of U associated with $\sigma_D, \dots, \sigma_{n+1}$.

To minimize $R_2(P, H)$ over H , we want to

1. calculate SVD of X_P , $X_P = U \Sigma V^T$.
2. set $\Pi_H = U_n U_n^T$ where $U_n = (\vec{u}_1 \ \dots \ \vec{u}_n \ \vec{0} \ \dots \vec{0})$; $U = (\vec{u}_1 \ \dots \ \vec{u}_D)$.

Bibliography

- [1] J. V. Atanasoff. Advent of electronic digital computing. *Annals of the History of Computing*, 6(3):229–282, 2008.
- [2] L. I. Bluestein. A Linear Filtering Approach to the Computation of Discrete Fourier Transform. *IEEE Transactions on Audio and Electroacoustics*, 18:451–455, 1970.
- [3] J. P. Boyd. *Chebyshev and Fourier Spectral Methods*. Dover Publications, Inc., 2001.
- [4] M. Brand. Incremental singular value decomposition of uncertain data with missing values. In *Computer Vision—ECCV 2002: 7th European Conference on Computer Vision Copenhagen, Denmark, May 28–31, 2002 Proceedings, Part I* 7, pages 707–720. Springer, 2002.
- [5] M. Brand. Fast low-rank modifications of the thin singular value decomposition. *Linear algebra and its applications*, 415(1):20–30, 2006.
- [6] J. W. Brown and R. V. Churchill. *Complex variables and applications*. McGraw-Hill,, 2009.
- [7] H. Cheng, Z. Gimbutas, P.-G. Martinsson, and V. Rokhlin. On the compression of low rank matrices. *SIAM Journal on Scientific Computing*, 26(4):1389–1404, 2005.
- [8] B. A. Cipra. The best of the 20th century: Editors name top 10 algorithms. *SIAM news*, 33(4):1–2, 2000.
- [9] J. Cooley and J. Tukey. An algorithm for the machine calculation of complex Fourier series. *Math. Comput.*, 19:297–301, 1965.
- [10] T. H. Cormen, C. E. Leiserson, R. L. Rivest, and C. Stein. Introduction to algorithms. *2nd Edition*, 2001.
- [11] J. W. Demmel. *Applied numerical linear algebra*. SIAM, 1997.
- [12] S. Foucart. *Mathematical pictures at a data science exhibition*. Cambridge University Press, 2022.

- [13] S. H. Friedberg, A. J. Insel, and L. E. Spence. *Linear Algebra (Fourth Edition)*. Pearson Higher Ed, 2003.
- [14] M. Frigo and S. Johnson. The design and implementation of fftw3. *Proceedings of IEEE 93 (2)*, pages 216–231, 2005.
- [15] S. R. Garcia and R. A. Horn. *Matrix Mathematics: A Second Course in Linear Algebra*. Cambridge University Press, 2023.
- [16] V. Guillemin and A. Pollack. *Differential topology*, volume 370. American Mathematical Soc., 2010.
- [17] M. Heideman, D. Johnson, and C. Burrus. Gauss and the history of the fast fourier transform. *IEEE ASSP Magazine*, 1(4):14–21, 1984.
- [18] R. Horn and C. Johnson. *Topics in Matrix Analysis*. Cambridge University Press, 1991.
- [19] R. A. Horn and C. R. Johnson. *Matrix analysis*. Cambridge university press, 2013.
- [20] M. A. Iwen and B. Ong. A distributed and incremental svd algorithm for agglomerative data analysis on large networks. *SIAM Journal on Matrix Analysis and Applications*, 37(4):1699–1718, 2016.
- [21] M. A. Iwen and C. V. Spencer. A note on compressed sensing and the complexity of matrix multiplication. *Information Processing Letters*, 109(10):468–471, 2009.
- [22] R. Kantrowitz and M. M. Neumann. Yet another proof of minkowski’s inequality. *The American Mathematical Monthly*, 115(5):445–447, 2008.
- [23] D. R. Kincaid and E. W. Cheney. *Numerical analysis: mathematics of scientific computing*, volume 2. Brooks/Cole Pacific Grove, CA, 2002.
- [24] R. Kress. *Numerical Analysis*, volume 181. Springer-Verlag, 1998.
- [25] I. Niven, H. S. Zuckerman, and H. L. Montgomery. *An introduction to the theory of numbers*. John Wiley & Sons, 1991.
- [26] C. H. Papadimitriou. *Computational Complexity*. Addison-Wesley Publishing Company, Inc., 1994.
- [27] G. Plonka, D. Potts, G. Steidl, and M. Tasche. *Numerical Fourier Analysis*. Springer, 2018.
- [28] L. Rabiner, R. Schafer, and C. Rader. The Chirp z-Transform Algorithm. *IEEE Transactions on Audio and Electroacoustics*, AU-17(2):86–92, June 1969.

- [29] T. Shifrin and M. Adams. *Linear algebra: A geometric approach (Second Edition)*. Macmillan, 2011.
- [30] G. Stewart. Perturbation theory for the singular value decomposition. *Digital Repository at the University of Maryland, UMIACS-TR-90-124 CS-TR 2539*, September 1990.
- [31] G. Stewart and J.-g. Sun. *Matrix Perturbation Theory*. Computer Science and Scientific Computing/Academic Press, Inc, 1990.
- [32] G. W. Stewart. On the early history of the singular value decomposition. *SIAM review*, 35(4):551–566, 1993.
- [33] V. Strassen. Gaussian elimination is not optimal. *Numerische mathematik*, 13(4):354–356, 1969.
- [34] L. N. Trefethen and D. Bau. *Numerical linear algebra*. SIAM, 2022.
- [35] V. V. Williams. Multiplying matrices faster than coppersmith-winograd. In *Proceedings of the forty-fourth annual ACM symposium on Theory of computing*, pages 887–898, 2012.