

NOTES FOR SECOND YEAR LINEAR ALGEBRA

JESSE RATZKIN

1. INTRODUCTION AND BACKGROUND

We begin with some motivational remarks and then outline a list (not necessarily complete!) of things you should already know.

1.1. Motivations. In these notes we will explore the algebra and geometry of a certain family of very nice transformations: the linear transformations. In previous classes, you've seen that one can add together vectors in space, and multiply a vector by a number (*i.e.* scalar). This gives three-dimensional space (in fact, any finite dimensional Euclidean space) a certain degree of algebraic structure. The 2LA module is all about the transformations of Euclidean space which preserve this structure. Along the way we will also explore some abstract properties of linear transformations and the natural spaces they act on, which are vector spaces. The material in this course will seem very abstract at times, but it has a wealth of applications to mathematical modeling, dynamical systems, differential equations, and many many other areas. For instance, it turns out that linear algebra greatly simplifies the classification of isometries of the plane and of three-dimensional Euclidean space (and of any dimensional Euclidean space). Some of you will take the 2DE module next semester, and there you will see many applications of eigenvalues and eigenvectors (which we will learn about in this module).

1.2. Background. Here we list some topics you should already be familiar with from MAM1000. This is by no means a complete list!

- vectors in two and three dimensions
- the dot product
- the cross product
- matrices, and operations on them (*e.g.* matrix addition and multiplication)
- systems of linear equations; in particular, you should know how to solve a system of linear equations by converting it to a matrix equation and then row-reducing the matrix.
- the complex numbers
- some elementary linear mappings of the plane (particularly rotations and reflections)

If you don't already know an item listed above, please review it as soon as possible.

1.3. Other resources. I have more or less cribbed these notes from the book *Linear Algebra* by S. Friedberg, A. Insel, and L. Spence. I have placed this book on short loan, and you can borrow it for 3 hours at a time.

There are many other great books on linear algebra, and you can find them under the call number 512.5 in the main library. I'd encourage you to page through several of them until you find one you like. Here are some others I like:

- *Elementary Linear Algebra* by H. Anton
- *Linear Algebra with Applications* by O. Bretscher
- *Introduction to Linear Algebra* by G. Strang

1.4. Notation. For future reference, we collect here a table of some common notation.

N	the set of all natural numbers $1, 2, 3, \dots$
Z	the set of all integers
Q	the set of all rational numbers
R	the set of all real numbers
C	the set of all complex numbers
$M_{m \times n}(\mathbf{R})$	the set of all matrices with m rows and n columns, whose entries are real numbers
$M_{m \times n}$	the set of all matrices with m rows and n columns, whose entries are real numbers
$M_{m \times n}(\mathbf{C})$	the set of all matrices with m rows and n column, whose entries are complex numbers
$\mathbf{R}_n[x]$	the set of all polynomials of degree at most n and real coefficients in the variable x
$\mathbf{R}[x]$	$\bigcup_{n=1}^{\infty} \mathbf{R}_n[x]$
$\det(A)$	the determinant of the matrix A
A^t	the transpose of a matrix A

2. REVISION: MATRICES AND TRANSFORMATIONS OF EUCLIDEAN SPACE

We will discuss transformations $T : \mathbf{R}^n \rightarrow \mathbf{R}^m$ which are linear. That is, they have the property

$$T(av + bu) = aT(v) + bT(u)$$

for all vectors $v, u \in \mathbf{R}^n$ and real numbers $a, b \in \mathbf{R}$.

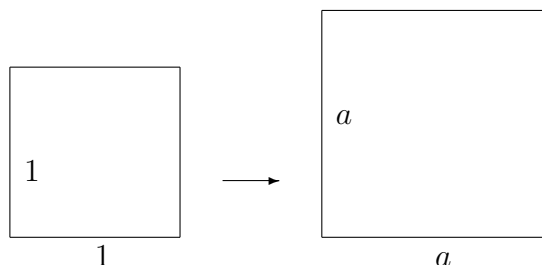
Remark 1. • *Most of this section is revision of material from MAM1000, so we will not cover it in lecture. However, it is left in the notes for your convenience.*

- *Throughout the discussion below, we will always use the standard basis $\{e_1, \dots, e_n\}$ for $\mathbf{R}^n$, where $e_j \in \mathbf{R}^n$ is the vector with 1 in the j th component and 0 everywhere else. Strictly speaking, we should denote this somehow in writing down the matrix associated with a linear transformation, for instance writing $[T]_{\mathcal{B}}$ instead of $[T]$. However, this would introduce a great amount of clutter into our formulas, so we suppress the basis dependence.*

See section 4.5 below for a discussion of how changing the basis in the domain and/or target effects the matrix involved.

We'll begin by considering linear transformations $T : \mathbf{R}^2 \rightarrow \mathbf{R}^2$, and (with a tiny bit more generality) linear transformations $T : \mathbf{R}^n \rightarrow \mathbf{R}^m$, where n, m are positive integers.

Scaling: The simplest sort of linear transformation of the plane we can write down is a rescaling (which is also called a dilation). If a is a positive number, we can send (x, y) to $a(x, y) = (ax, ay)$. Geometrically, we can imagine this transformation as taking the unit square $\{0 \leq x \leq 1, 0 \leq y \leq 1\}$ to a similiar square $\{0 \leq x \leq a, 0 \leq y \leq a\}$, which we represent in the following picture.



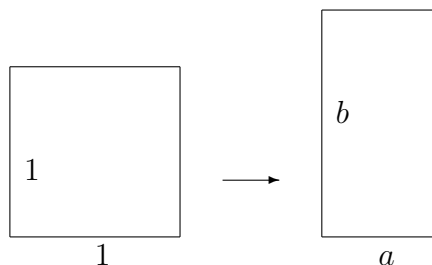
Actually, if we adopt the convention that we always start with the unit square $\{0 \leq x \leq 1, 0 \leq y \leq 1\}$, we really only need to draw the square on the right to have a geometric picture of our transformation. If we want to write this transformation in terms of matrices, we can write

$$\begin{bmatrix} x \\ y \end{bmatrix} \mapsto \begin{bmatrix} a & 0 \\ 0 & a \end{bmatrix} \begin{bmatrix} x \\ y \end{bmatrix} = \begin{bmatrix} ax \\ ay \end{bmatrix}.$$

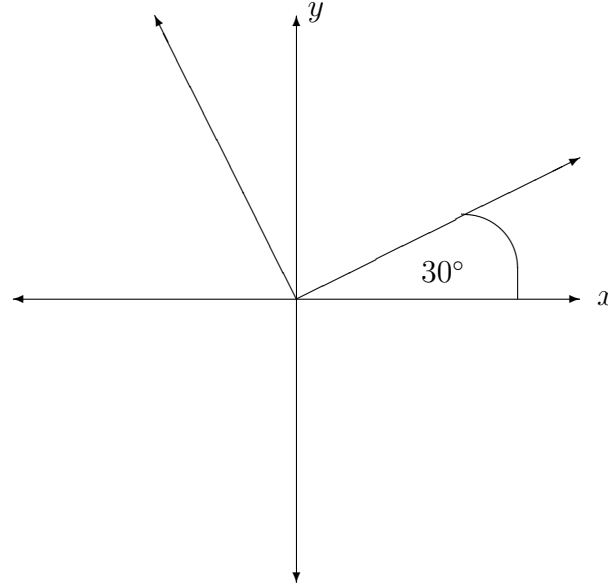
In the previous example, we scaled the horizontal and vertical axes by the same factor, but there's no reason we have to do this. More generally, we might scale the horizontal axis by $a > 0$ and the vertical axis by $b > 0$. This time, we can write the transformation as

$$\begin{bmatrix} x \\ y \end{bmatrix} \mapsto \begin{bmatrix} a & 0 \\ 0 & b \end{bmatrix} \begin{bmatrix} x \\ y \end{bmatrix} = \begin{bmatrix} ax \\ by \end{bmatrix}.$$

As before, we can represent this transformation with a picture.



Rotations and reflections: Let's suppose we want to write down a formula for a 30° counterclockwise rotation in the plane; call this rotation R_{30} . For instance, we may have a collection of datapoints we'd like to put into a database, but the coordinates of these datapoints are all rotated by 30° in the clockwise direction, and so we want to undo this rotation, by rotating through the same angle in the opposite direction. So let's find out how to write down a formula for the rotation. We start with a picture of the vectors $(1, 0)$ and $(0, 1)$ rotated by 30° .



We see that $(1, 0)$ gets mapped to the point

$$R_{30} \left(\begin{bmatrix} 1 \\ 0 \end{bmatrix} \right) = \begin{bmatrix} \cos(30^\circ) \\ \sin(30^\circ) \end{bmatrix} = \begin{bmatrix} \sqrt{3}/2 \\ 1/2 \end{bmatrix},$$

and that $(0, 1)$ gets mapped to the point

$$R_{30} \left(\begin{bmatrix} 0 \\ 1 \end{bmatrix} \right) = \begin{bmatrix} -\sin(30^\circ) \\ \cos(30^\circ) \end{bmatrix} = \begin{bmatrix} -1/2 \\ \sqrt{3}/2 \end{bmatrix}.$$

Next observe that we can rescale the vectors $(1, 0)$ and $(0, 1)$, and, because rotations don't change lengths, the rotation will carry these scalings along:

$$R_{30} \left(\begin{bmatrix} x \\ 0 \end{bmatrix} \right) = \begin{bmatrix} \sqrt{3}x/2 \\ x/2 \end{bmatrix}, \quad R_{30} \left(\begin{bmatrix} 0 \\ y \end{bmatrix} \right) = \begin{bmatrix} -y/2 \\ \sqrt{3}y/2 \end{bmatrix}.$$

Finally, we can put this all together, because R_{30} acts independently on the vectors $(1, 0)$ and $(0, 1)$. This means R_{30} rotates $\begin{bmatrix} 1 \\ 0 \end{bmatrix}$ and $\begin{bmatrix} 0 \\ 1 \end{bmatrix}$ independently, which we can write as

$$R_\theta \left(\begin{bmatrix} 1 \\ 0 \end{bmatrix} + \begin{bmatrix} 0 \\ 1 \end{bmatrix} \right) = R_\theta \left(\begin{bmatrix} 1 \\ 0 \end{bmatrix} \right) + R_\theta \left(\begin{bmatrix} 0 \\ 1 \end{bmatrix} \right).$$

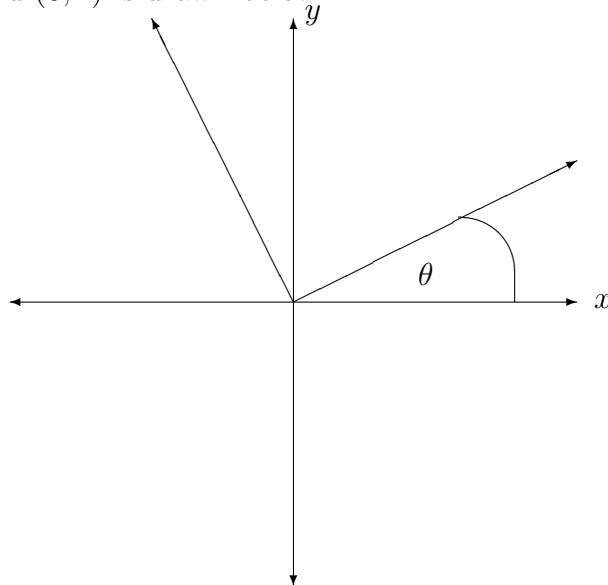
4

(You can verify this formula geometrically, by seeing where the rotation carries the top right corner of the unit square.) Adding these two vectors together, we have

$$R_{30} \left(\begin{bmatrix} x \\ y \end{bmatrix} \right) = \begin{bmatrix} \sqrt{3}x/2 - y/2 \\ x/2 + \sqrt{3}y/2 \end{bmatrix} = \begin{bmatrix} \sqrt{3}/2 & -1/2 \\ 1/2 & \sqrt{3}/2 \end{bmatrix} \begin{bmatrix} x \\ y \end{bmatrix}.$$

In this last step we used the rule for multiplying matrices we stated previously in the notes, which starts to explain why we defined matrix multiplication the way we did.

We can redo this whole discussion with a rotation through any angle. Let R_θ be the rotation through angle θ in the counterclockwise direction, whose action on the vectors $(1, 0)$ and $(0, 1)$ is drawn below.



Then, just as before, we have

$$R_\theta \left(\begin{bmatrix} 1 \\ 0 \end{bmatrix} \right) = \begin{bmatrix} \cos \theta \\ \sin \theta \end{bmatrix}, \quad R_\theta \left(\begin{bmatrix} 0 \\ 1 \end{bmatrix} \right) = \begin{bmatrix} -\sin \theta \\ \cos \theta \end{bmatrix},$$

and, by the same argument we have above,

$$R_\theta \left(\begin{bmatrix} x \\ y \end{bmatrix} \right) = \begin{bmatrix} x \cos \theta - y \sin \theta \\ x \sin \theta + y \cos \theta \end{bmatrix} = \begin{bmatrix} \cos \theta & -\sin \theta \\ \sin \theta & \cos \theta \end{bmatrix} \begin{bmatrix} x \\ y \end{bmatrix}.$$

In this way, we can say that the rotation R_θ is given by multiplication (on the left) by the matrix

$$[R_\theta] = \begin{bmatrix} \cos \theta & -\sin \theta \\ \sin \theta & \cos \theta \end{bmatrix}.$$

Exercise: Verify that

$$[R_{-\theta}] = \begin{bmatrix} \cos \theta & \sin \theta \\ -\sin \theta & \cos \theta \end{bmatrix}.$$

Exercise: Verify that $[R_\theta]^2 = [R_{2\theta}]$. (You'll need to remember double angle formulas from trigonometry.)

Exercise: Verify that $[R_\theta][R_\phi] = [R_{\theta+\phi}]$. (You'll need to remember the angle addition formulas from trigonometry.) Notice that rotation matrices commute! That is, $[R_\theta][R_\phi] = [R_\phi][R_\theta]$.

Now that we've figured out how to write any rotation as multiplication by a matrix, let's be a little ambitious and see what else we can write. The next natural thing to consider is a reflection. The reflection through the x axis sends (x, y) to $(x, -y)$. We can write this as a matrix product by

$$\begin{bmatrix} x \\ y \end{bmatrix} \mapsto \begin{bmatrix} x \\ -y \end{bmatrix} = \begin{bmatrix} 1 & 0 \\ 0 & -1 \end{bmatrix} \begin{bmatrix} x \\ y \end{bmatrix}.$$

Now that we know how to write reflection through the y axis and any rotation, we can write down any reflection through a line that intersects the origin $(0, 0)$. Indeed, let l be a line passing through the origin making an angle θ with the positive x axis, and let r_l be reflection through the line l . We build the matrix for r_l by performing three transformations in succession. We first rotate our coordinates by $-\theta$, then reflect through the x axis, and then rotate back by the angle θ . The result is a reflection fixing the line l , so it must be r_l , and it has the matrix representation

$$\begin{aligned} [r_l] &= \begin{bmatrix} \cos \theta & -\sin \theta \\ \sin \theta & \cos \theta \end{bmatrix} \begin{bmatrix} 1 & 0 \\ 0 & -1 \end{bmatrix} \begin{bmatrix} \cos \theta & \sin \theta \\ -\sin \theta & \cos \theta \end{bmatrix} \\ &= \begin{bmatrix} \cos^2 \theta - \sin^2 \theta & 2 \cos \theta \sin \theta \\ 2 \cos \theta \sin \theta & \sin^2 \theta - \cos^2 \theta \end{bmatrix} \\ &= \begin{bmatrix} \cos(2\theta) & \sin(2\theta) \\ \sin(2\theta) & -\cos(2\theta) \end{bmatrix}. \end{aligned}$$

Let's check quickly that we have the right matrix for the reflection through the line l . This line l is uniquely determined by the two points $(0, 0)$ and $(\cos \theta, \sin \theta)$, so we only need to check that $[r_l]$ fixes these two vectors. It is quite easy to check that

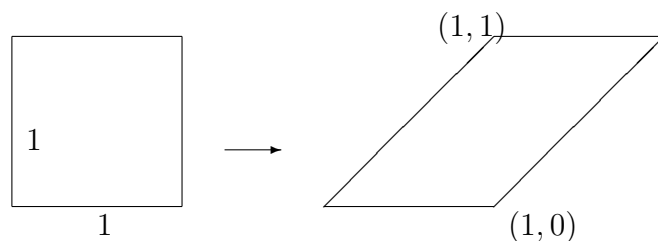
$$[r_l] \begin{bmatrix} 0 \\ 0 \end{bmatrix} = \begin{bmatrix} \cos(2\theta) & \sin(2\theta) \\ \sin(2\theta) & -\cos(2\theta) \end{bmatrix} \begin{bmatrix} 0 \\ 0 \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \end{bmatrix}.$$

Now we check that $[r_l]$ fixes $(\cos \theta, \sin \theta)$:

$$\begin{aligned} [r_l] \begin{bmatrix} \cos \theta \\ \sin \theta \end{bmatrix} &= \begin{bmatrix} \cos^2 \theta - \sin^2 \theta & 2 \cos \theta \sin \theta \\ 2 \cos \theta \sin \theta & \sin^2 \theta - \cos^2 \theta \end{bmatrix} \begin{bmatrix} \cos \theta \\ \sin \theta \end{bmatrix} \\ &= \begin{bmatrix} \cos^3 \theta - \sin^2 \theta \cos \theta + 2 \sin^2 \theta \cos \theta \\ 2 \cos^2 \theta \sin \theta - \cos^2 \theta \sin \theta + \sin^3 \theta \end{bmatrix} \\ &= \begin{bmatrix} \cos \theta (\cos^2 \theta + \sin^2 \theta) \\ \sin \theta (\cos^2 \theta + \sin^2 \theta) \end{bmatrix} = \begin{bmatrix} \cos \theta \\ \sin \theta \end{bmatrix}. \end{aligned}$$

Sheers: The next sort of transformation we'll talk about is a *sheer*, which you can imagine as what happens to a deck of cards (as viewed from the side) when you push the top card to the side and hold the bottom card still. This means a sheer will fix one direction, say the direction of $\begin{bmatrix} 1 \\ 0 \end{bmatrix}$, but it will move the other

directions. We draw a picture of what this transformation does the unit square $\{0 \leq x \leq 1, 0 \leq y \leq 1\}$ below.



We'll call this sheer S .

We'll construct the matrix for this sheer mapping S by seeing what it does to the two coordinate vectors $\begin{bmatrix} 1 \\ 0 \end{bmatrix}$ and $\begin{bmatrix} 0 \\ 1 \end{bmatrix}$, which, in a way, is how we constructed the rotation matrix. We can see from the picture that $\begin{bmatrix} 1 \\ 0 \end{bmatrix}$ keeps the same direction, so we can rescale in the horizontal direction to make

$$S\left(\begin{bmatrix} 1 \\ 0 \end{bmatrix}\right) = \begin{bmatrix} 1 \\ 0 \end{bmatrix},$$

which tells us

$$[S] = \begin{bmatrix} 1 & * \\ 0 & * \end{bmatrix}.$$

Here the $*$'s can stand for any number, because we haven't figured out yet what these parts of $[S]$ are.

On the other hand, the vector $\begin{bmatrix} 0 \\ 1 \end{bmatrix}$ gets tilted in the clockwise direction, and it looks like it gets stretched as well. If we look a little more closely, we see

$$S\left(\begin{bmatrix} 0 \\ 1 \end{bmatrix}\right) = \begin{bmatrix} 1 \\ 1 \end{bmatrix},$$

which tells us

$$[S] = \begin{bmatrix} 1 & 1 \\ 0 & 1 \end{bmatrix}.$$

We can check this last formula by separating out the components:

$$S\left(\begin{bmatrix} x \\ 0 \end{bmatrix} + \begin{bmatrix} 0 \\ y \end{bmatrix}\right) = S\left(\begin{bmatrix} x \\ y \end{bmatrix}\right) = \begin{bmatrix} 1 & 1 \\ 0 & 1 \end{bmatrix} \begin{bmatrix} x \\ y \end{bmatrix}.$$

We've just constructed the matrix of a particular sheer which fixes the horizontal direction. In general, the vertical direction will go to some other direction, so that

$$\begin{bmatrix} 0 \\ 1 \end{bmatrix} \mapsto \begin{bmatrix} a \\ 1 \end{bmatrix},$$

where $a \neq 0$ is a number. Notice that we have the second component equal to 1, which we can arrange by rescaling if necessary. We always have the second component nonzero, because otherwise the shear would collapse the unit square down to a (horizontal) line segment. Now, following the same reasoning as we did above, we find the matrix of this shear is

$$[S] = \begin{bmatrix} 1 & a \\ 0 & 1 \end{bmatrix}.$$

Exercise: Notice that a can be negative in the formula just above. What does the parallelogram which is the image of the unit square $\{0 \leq x \leq 1, 0 \leq y \leq 1\}$ look like in this case? In particular, what can you say about the angle at the origin $(0, 0)$?

Exercise: Show that the general shear which fixes the y axis is given by a matrix of the form

$$[S] = \begin{bmatrix} 1 & 0 \\ a & 1 \end{bmatrix},$$

where $a \neq 0$ is a number.

Exercise: Construct the general shear which fixes the $\begin{bmatrix} 1 \\ 1 \end{bmatrix}$ direction. Hint: you might want to apply a rotation.

General matrices as mappings: We just saw how to construct the matrix associated to a shear by tracking where the shear transformation sends the basis vectors $\begin{bmatrix} 1 \\ 0 \end{bmatrix}$ and $\begin{bmatrix} 0 \\ 1 \end{bmatrix}$. In fact, this technique is exactly how we can produce the matrix associated to any linear transformation. Let $T : \mathbf{R}^2 \rightarrow \mathbf{R}^2$ be any linear transformation, which means $T(v + w) = T(v) + T(w)$ for all vectors $v, w \in \mathbf{R}^2$ and $T(av) = aT(v)$ for all scalars a .

Exercise: Prove that $T\left(\begin{bmatrix} 0 \\ 0 \end{bmatrix}\right) = \begin{bmatrix} 0 \\ 0 \end{bmatrix}$ for any linear mapping. Hint: suppose otherwise; then what is $T\left(2\begin{bmatrix} 0 \\ 0 \end{bmatrix}\right)$?

We can construct a matrix associated to T , which we call $[T]$, as follows. The first column of $[T]$ is $T\left(\begin{bmatrix} 1 \\ 0 \end{bmatrix}\right)$, and the second column of $[T]$ is $T\left(\begin{bmatrix} 0 \\ 1 \end{bmatrix}\right)$.

Let's check this is actually the right matrix. Suppose we have a linear mapping $T : \mathbf{R}^2 \rightarrow \mathbf{R}^2$ with

$$T\left(\begin{bmatrix} 1 \\ 0 \end{bmatrix}\right) = \begin{bmatrix} a \\ c \end{bmatrix}, \quad T\left(\begin{bmatrix} 0 \\ 1 \end{bmatrix}\right) = \begin{bmatrix} b \\ d \end{bmatrix}.$$

In this case we'd like to check that the matrix associated with T is

$$[T] = \begin{bmatrix} a & b \\ c & d \end{bmatrix}.$$

Indeed,

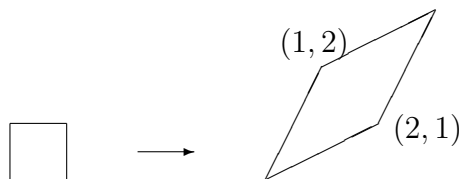
$$\begin{aligned}
 T\left(\begin{bmatrix} x \\ y \end{bmatrix}\right) &= T\left(\begin{bmatrix} x \\ 0 \end{bmatrix} + \begin{bmatrix} 0 \\ 1 \end{bmatrix}\right) \\
 &= xT\left(\begin{bmatrix} 1 \\ 0 \end{bmatrix}\right) + yT\left(\begin{bmatrix} 0 \\ 1 \end{bmatrix}\right) \\
 &= x\begin{bmatrix} a \\ c \end{bmatrix} + y\begin{bmatrix} b \\ d \end{bmatrix} \\
 &= \begin{bmatrix} ax + by \\ cx + dy \end{bmatrix} \\
 [T]\begin{bmatrix} x \\ y \end{bmatrix} &= \begin{bmatrix} a & b \\ c & d \end{bmatrix}\begin{bmatrix} x \\ y \end{bmatrix} \\
 &= \begin{bmatrix} ax + by \\ cx + dy \end{bmatrix}.
 \end{aligned}$$

In both computations we end up with the same answer, regardless of which x and y we choose, so this matrix must be the correct choice.

Let's look at an example. Suppose we want to find the linear map which takes the unit square $\{0 \leq x \leq 1, 0 \leq y \leq 1\}$ to the parallelogram with the vertices

$$(0, 0), \quad (2, 1), \quad (1, 2), \quad (3, 3).$$

Here's a picture.



In fact, we have two choices for this linear mapping; we can either have

$$T\left(\begin{bmatrix} 1 \\ 0 \end{bmatrix}\right) = \begin{bmatrix} 2 \\ 1 \end{bmatrix}, \quad T\left(\begin{bmatrix} 0 \\ 1 \end{bmatrix}\right) = \begin{bmatrix} 1 \\ 2 \end{bmatrix},$$

or we can have

$$T\left(\begin{bmatrix} 1 \\ 0 \end{bmatrix}\right) = \begin{bmatrix} 1 \\ 2 \end{bmatrix}, \quad T\left(\begin{bmatrix} 0 \\ 1 \end{bmatrix}\right) = \begin{bmatrix} 2 \\ 1 \end{bmatrix}.$$

In the first case we have

$$[T] = \begin{bmatrix} 2 & 1 \\ 1 & 2 \end{bmatrix},$$

and in the second case we have

$$[T] = \begin{bmatrix} 1 & 2 \\ 2 & 1 \end{bmatrix}.$$

Notice that we can get from one of these matrices to the other by swapping the columns, which geometrically corresponds to the swapping $\begin{bmatrix} 1 \\ 0 \end{bmatrix}$ and $\begin{bmatrix} 0 \\ 1 \end{bmatrix}$. We can write this in terms of matrix multiplication as

$$\begin{bmatrix} 2 & 1 \\ 1 & 2 \end{bmatrix} = \begin{bmatrix} 1 & 2 \\ 2 & 1 \end{bmatrix} \begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix}.$$

(You should check the matrix product.) This should not surprise you. The matrix $\begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix}$ corresponds to the reflection through the line $y = x$, which maps our parallelogram to itself and interchanges the vectors $\begin{bmatrix} 1 \\ 0 \end{bmatrix}$ and $\begin{bmatrix} 0 \\ 1 \end{bmatrix}$. Thus we see that **we represent the composition of linear mappings as matrix multiplication**. We will return to this important idea later on in these notes.

Exercise: Why can't we have $T\left(\begin{bmatrix} 1 \\ 0 \end{bmatrix}\right) = \begin{bmatrix} 3 \\ 3 \end{bmatrix}$? Hint: what is $\begin{bmatrix} 1 \\ 0 \end{bmatrix} + \begin{bmatrix} 0 \\ 1 \end{bmatrix}$?

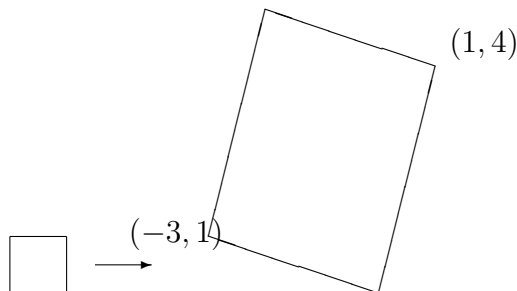
In fact, we can reverse this process. Suppose we have a matrix, let's say

$$[T] = \begin{bmatrix} 1 & -3 \\ 4 & 1 \end{bmatrix},$$

and we want to understand the linear transformation associated to this matrix. We can draw the parallelogram that T sends the unit square $\{0 \leq x \leq 1, 0 \leq y \leq 1\}$ onto, which gives us all the geometric information about T . We see from the matrix that

$$T\left(\begin{bmatrix} 1 \\ 0 \end{bmatrix}\right) = \begin{bmatrix} 1 \\ 4 \end{bmatrix}, \quad T\left(\begin{bmatrix} 0 \\ 1 \end{bmatrix}\right) = \begin{bmatrix} -3 \\ 1 \end{bmatrix}.$$

To draw the parallelogram, all we need to do is draw in these two edges starting at $\begin{bmatrix} 0 \\ 0 \end{bmatrix}$ and connect them. We end up with the following picture.



Exercise: Can you explain why the image of a square is always a parallelogram? (Or a line segment, which is really a degenerate parallelogram, with one pair of opposite angles collapsed to 0 ...)

Beyond two dimensions: So far we've seen how to write down the matrix of a linear transformation taking the unit square to an arbitrary parallelogram, and how to draw the parallelogram which is the image of the unit square under an arbitrary

linear mapping. However, nothing we've done so far is special to two dimensions, and everything works in higher dimensions. Let's suppose $T : \mathbf{R}^n \rightarrow \mathbf{R}^m$ is a linear mapping. Also let $\{e_1, e_2, e_3, \dots, e_n\}$ be the vectors in $\mathbf{R}^n$ where e_i has a 1 in the i th component and 0 elsewhere. (This is known as the *standard basis* of $\mathbf{R}^n$.) Then we can write down a matrix $[T]$, where the i th column of T is $T(e_i)$. We have

$$[T] = \begin{bmatrix} \vdots & & & \vdots \\ T(e_1) & T(e_2) & \cdots & T(e_n) \\ \vdots & & & \vdots \end{bmatrix}.$$

We can do a quick example, and write down the linear mapping $T : \mathbf{R}^3 \rightarrow \mathbf{R}^3$ taking the unit cube $\{0 \leq x \leq 1, 0 \leq y \leq 1, 0 \leq z \leq 1\}$ to the parallelepiped which has the vertices

$$\begin{pmatrix} 0, 0, 0 & 3, 1, 1 & 1, 3, 1 & 1, 1, 3 \\ 4, 4, 2 & 4, 2, 4 & 2, 4, 4 & 5, 5, 5 \end{pmatrix}.$$

This time there are actually six such examples; we'll write one of them down, and leave the other five to you. If we want

$$T\left(\begin{bmatrix} 1 \\ 0 \\ 0 \end{bmatrix}\right) = \begin{bmatrix} 3 \\ 1 \\ 1 \end{bmatrix}, \quad T\left(\begin{bmatrix} 0 \\ 1 \\ 0 \end{bmatrix}\right) = \begin{bmatrix} 1 \\ 3 \\ 1 \end{bmatrix}, \quad T\left(\begin{bmatrix} 0 \\ 0 \\ 1 \end{bmatrix}\right) = \begin{bmatrix} 1 \\ 1 \\ 3 \end{bmatrix},$$

then we must have

$$[T] = \begin{bmatrix} 3 & 1 & 1 \\ 1 & 3 & 1 \\ 1 & 1 & 3 \end{bmatrix}.$$

Exercise: Write down the matrices of the other five possible linear mappings which carry the unit cube onto this parallelepiped.

Composition of mappings and matrix multiplication: There is one very important thing we mentioned above, which we emphasize here. Namely, the composition of linear transformations is given by matrix multiplication. More precisely, if $T_1 : \mathbf{R}^n \rightarrow \mathbf{R}^k$ and $T_2 : \mathbf{R}^k \rightarrow \mathbf{R}^m$ are linear mappings, then the composition

$$T_2 \circ T_1 : \mathbf{R}^n \rightarrow \mathbf{R}^m, \quad T_2 \circ T_1(v) = T_2(T_1(v))$$

is also linear, and the matrix associated to the composition is the product of the matrices:

$$[T_2 \circ T_1] = [T_2][T_1].$$

This explains why the definition of matrix multiplication is the way it is. As a quick check, it's good to see that the matrix product is well-defined. We have $T_1 : \mathbf{R}^n \rightarrow \mathbf{R}^k$, so that $[T_1]$ is a $k \times n$ matrix, and (similarly) $[T_2]$ is a $m \times k$ matrix. Then the product $[T_2][T_1]$ is well-defined, and it is an $m \times n$ matrix. Also, the composition $T_2 \circ T_1 : \mathbf{R}^n \rightarrow \mathbf{R}^m$ corresponds to an $m \times n$ matrix. And so everything fits together nicely.

3. VECTOR SPACES

This section is going to seem for a while as if it has nothing to do with the previous (revision) section, but in fact the two are very much related. We would like to generalize much of what we wrote out for matrices with m rows and n columns corresponding to nice transformation from $\mathbf{R}^n$ to $\mathbf{R}^m$. In our generalization, we'd like to replace our Euclidean spaces $\mathbf{R}^n$ and $\mathbf{R}^m$ with other objects, and it turns out that the correct objects are vector spaces. In order to formulate the theory correctly, we must first define important properties of vector spaces, so please be patient while we do this.

We begin with the definition of a vector space.

3.1. Definitions and examples.

Definition 1. *A vector space V over the real numbers $\mathbf{R}$ is a set equipped with two operations: vector addition and scalar multiplication. Moreover, addition and multiplication obey the following rules, where $u, v, w \in V$ are vectors and $a, b \in \mathbf{R}$ are scalars (numbers).*

- $av + bw$ is a well-defined vector in V
- $v + w = w + v$ and $(u + v) + w = u + (v + w)$
- $(a + b)v = av + bv$ and $a(v + w) = av + aw$
- $(ab)v = a(bv)$
- There is a vector $0 \in V$, the additive identity, such that $0 + v = v$ for all $v \in V$.
- For any $v \in V$ we have $v + (-1) \cdot v = 0$, where the 0 on the right hand side of this equation is the vector 0.
- For any $v \in V$ we have $1 \cdot v = v$, where $1 \in \mathbf{R}$.

Remark 2. *One can replace the real numbers $\mathbf{R}$ with the complex numbers $\mathbf{C}$, and then one gets a vector space over the complex numbers. We will see that every vector space over $\mathbf{C}$ is also a vector space over $\mathbf{R}$, but the reverse is not generally true. In fact, vector spaces over $\mathbf{C}$ inherit additional structure.*

Example: We use the rules of a vector space to show that for all $v \in V$ we have $0 \cdot v = 0$, where the 0 on the left hand side of the equation is the scalar 0 and the 0 on the right hand side is the vector 0. Indeed,

$$0 \cdot v = (1 + (-1)) \cdot v = 1 \cdot v + (-1) \cdot v = v + (-1) \cdot v = 0.$$

It is useful to keep in mind some examples.

- the real line, with the usual addition and multiplication
- the complex numbers, with complex addition and multiplication
- The set of vectors in the plane, with usual vector addition and scalar multiplication is a real vector space.
- The usual set of vectors in three-dimensional space, with the usual vector addition and scalar multiplication, is a real vector space.
- In fact, there is nothing special about two or three dimensions. The set of vectors in any dimension, with the usual vector addition and scalar multiplication, is a real vector space.

- Let $n = 0, 1, 2, 3, \dots$ and let $\mathbf{R}_n[x]$ be the set of polynomials of degree at most n . We can write any such polynomial as

$$p(x) = a_0 + a_1x + a_2x^2 + \cdots + a_nx^n.$$

This is a vector space, where the vector addition is the sum of two polynomials and scalar multiplication is the multiplication of a polynomial by a number.

- Let $V = \{v \in \mathbf{R}^3 : v \perp (1, 1, 1)\} = \{v = (x, y, z) : x + y + z = 0\}$. This is a vector space sitting inside $\mathbf{R}^3$, which you might recognize as a plane passing through the origin.

This last example is a particular kind of vector space.

Definition 2. Let V be a vector space. A (vector) subspace W of V is a subset of V which is itself a vector space, with vector addition and scalar multiplication in W being the restriction of those operations in V .

Proposition 1. Let V be a vector space and let $W \subset V$ be a subset. Then W is a subspace if and only if it is closed under addition and scalar multiplication and $0 \in W$.

Proof. First suppose that W is a subspace of V . Then by definition, $0 \in W$ and, for any $w_1, w_2 \in W$ and $a, b \in \mathbf{R}$ we have $aw_1 + bw_2 \in W$.

Now suppose that $W \subset V$ contains 0 and is closed under addition and scalar multiplication. Then W inherits all the algebraic properties of V , such as the associative and distributive laws. Thus it is a straightforward exercise to check that W satisfies all the conditions of being a vector space. $\square$

Proposition 2. Let W_1 and W_2 be subspaces of a vector space V . Then $W_1 \cap W_2$ is also a subspace.

Proof. We only need to show that $W_1 \cap W_2$ contains 0 and is closed under linear combinations. Indeed, $0 \in W_1$ and $0 \in W_2$, so $0 \in W_1 \cap W_2$. Let $u, v \in W_1 \cap W_2$ and let $a, b \in \mathbf{R}$. Then $au + bv \in W_1$ and $au + bv \in W_2$, so $au + bv \in W_1 \cap W_2$. $\square$

Example: It is not always true that the union of two subspaces is a subspace. Let $V = \mathbf{R}^2$, let $W_1 = \{(x, 0) : x \in \mathbf{R}\}$, and let $W_2 = \{(0, y) : y \in \mathbf{R}\}$. Then W_1 and W_2 are subspaces of V , but $W_1 \cup W_2$ is not. Indeed, $(1, 0) \in W_1$ and $(0, 1) \in W_2$, but $(1, 1) = (1, 0) + (0, 1) \notin W_1 \cup W_2$. Thus $W_1 \cup W_2$ is not closed under vector addition, so it cannot be a subspace.

Exercise: Prove that $W_1 \cup W_2$ is a subspace if and only if either $W_1 \subset W_2$ or $W_2 \subset W_1$.

In the next chapter we will see many very important examples of subspaces in vector spaces, but there are two large families of examples you already know.

- Any line through the origin in either the plane or three-space is a vector subspace.
- Any plane through the origin in three-space is a vector subspace.

3.2. Linear dependence and independence. We return to vector spaces over $\mathbf{R}$.

Definition 3. Let V be a vector space over $\mathbf{R}$ and let $v_1, v_2, \dots, v_k \in V$. We say $\{v_1, \dots, v_k\}$ are linearly dependent if there is a choice of scalars $a_1, a_2, \dots, a_k$, not all of which are zero, such that

$$a_1v_1 + a_2v_2 + \dots + a_kv_k = 0.$$

Otherwise, we say $\{v_1, \dots, v_k\}$ are linearly independent.

The left hand side of this equation above,

$$a_1v_1 + a_2v_2 + \dots + a_kv_k,$$

is called a (finite) linear combination of the vectors $v_1, \dots, v_k$, and the scalars $a_1, \dots, a_k$ are called the coefficients of the linear combination. Thus we can say a set $A = \{v_1, \dots, v_k\}$ is linearly dependent if and only if there is some linear combination, not all of whose coefficients are zero, where the linear combination itself is zero.

Example: We prove that $\{v_1, \dots, v_k\}$ are linearly independent if and only if

$$a_1v_1 + a_2v_2 + \dots + a_kv_k = 0$$

implies $a_1 = a_2 = \dots = a_k = 0$. Suppose there exist coefficients $a_1, a_2, \dots, a_k$, not all of which are zero, such that

$$a_1v_1 + a_2v_2 + \dots + a_kv_k = 0.$$

Then by definition $\{v_1, v_2, \dots, v_k\}$ is linearly dependent, so it is not a linearly independent set. Now suppose the only set of coefficients $a_1, a_2, \dots, a_k$ such that

$$a_1v_1 + a_2v_2 + \dots + a_kv_k = 0$$

is $a_1 = a_2 = \dots = a_k = 0$. Then by definition $\{v_1, v_2, \dots, v_k\}$ cannot be linearly dependent, so it must be linearly independent.

Example: Consider the sets $A = \{v_1, v_2, v_3\}$ and $B = \{v_1, v_2, v_3, v_4\}$. We prove that if A is linearly dependent then so is B , and that if B is linearly independent then so is A . First suppose that A is linearly dependent. Then there must exist coefficients a_1, a_2, a_3 , not all of which are zero, so that

$$a_1v_1 + a_2v_2 + a_3v_3 = 0.$$

Setting $b_1 = a_1, b_2 = a_2, b_3 = a_3, b_4 = 0$, we now have coefficients b_1, b_2, b_3, b_4 , not all of which are zero, such that

$$b_1v_1 + b_2v_2 + b_3v_3 + b_4v_4 = 0,$$

which proves that B is linearly dependent. Now suppose that B is linearly independent, and suppose that

$$a_1v_1 + a_2v_2 + a_3v_3 = 0$$

for some coefficients a_1, a_2, a_3 . Again, choose $b_1 = a_1, b_2 = a_2, b_3 = a_3, b_4 = 0$, and we have coefficients such that

$$0 = b_1v_1 + b_2v_2 + b_3v_3 + b_4v_4.$$

However, B is linearly independent, so we must have $b_1 = b_2 = b_3 = b_4 = 0$, so in particular we must have $a_1 = a_2 = a_3 = 0$. Thus A is also linearly independent.

One can use exactly the same reasoning to prove the following result.

Proposition 3. *Let V be a vector space and let $A \subset B \subset V$ be subsets. If A is linearly dependent then so is B , and if B is linearly independent then so is A .*

3.3. Span and basis.

Definition 4. *If $A = \{v_1, v_2, \dots, v_k\}$ is a subset of a vector space V , we define the span of A as*

$$\text{span}(A) = \{a_1v_1 + a_2v_2 + \dots + a_kv_k : a_1, a_2, \dots, a_k \in \mathbf{R}\}.$$

By convention, the span of the empty set $\emptyset$ is the singleton set $\{0\}$. We also say that the span of A is the linear combination of all vectors in A .

Observe, that, by definition, the span of any set A contains sums of only **finitely many** elements of A , and not infinite sums. This is not an important distinction for finite-dimensional vector spaces, but it is important when the vector space is infinite-dimensional (*e.g.* the space of all polynomials).

We have already seen some examples. For instance,

$$\mathbf{R}^2 = \text{span}\{e_1, e_2\} = \text{span}\left\{\begin{pmatrix} 1 \\ 0 \end{pmatrix}, \begin{pmatrix} 0 \\ 1 \end{pmatrix}\right\}$$

and

$$\mathbf{R}_2[x] = \text{span}\{1, x, x^2\}.$$

Definition 5. *Let V be a vector space. A basis for V is a set of vectors $\{v_1, v_2, \dots\}$ which is linearly independent and whose span is all of V . If V has a basis $\mathcal{B} = \{v_1, v_2, \dots, v_n\}$ with finitely many elements, then we say V is finite dimensional, and has dimension $n = \#(\mathcal{B})$. Otherwise, we say V is infinite dimensional.*

We will see below that, if V is finite dimensional, then the number of elements in a basis for is it fixed, and so the dimension of V is well-defined.

Observe that the choice of basis is definitely not unique. Any vector space will have many different bases. For instance, we can observe that

$$\mathcal{B}_1 = \left\{\begin{bmatrix} 1 \\ 0 \end{bmatrix}, \begin{bmatrix} 0 \\ 1 \end{bmatrix}\right\}, \quad \mathcal{B}_2 = \left\{\begin{bmatrix} 1 \\ -1 \end{bmatrix}, \begin{bmatrix} 1 \\ 1 \end{bmatrix}\right\}$$

are both perfectly fine bases for $\mathbf{R}^2$. Can you find other (distinct) bases for $\mathbf{R}^2$?

Example: We show that a set $\mathcal{B} = \{v_1, v_2, \dots\}$ is a basis for the vector space V if and only if it satisfies the following two conditions:

- the vectors $v_1, v_2, \dots$ in $\mathcal{B}$ are linearly independent
- one can write any $w \in V$ as a linear combination

$$w = \sum_{v_j \in \mathcal{B}} a_j v_j, \text{ where only finitely many } a_j \text{ are nonzero.}$$

If $\mathcal{B}$ is a basis then by definition the vectors $v_1, v_2, \dots$ form a linearly independent set which spans V . The fact that $\mathcal{B}$ spans V says exactly that, given any $w \in V$, there are coefficients $a_1, a_2, \dots$ such that

$$w = \sum_{v_j \in \mathcal{B}} a_j v_j.$$

Now suppose we can write any w as a finite linearly combination of elements of $\mathcal{B}$. This means $\mathcal{B}$ spans V . If, in addition, $\mathcal{B}$ is also linearly independent, then it must be a basis. Notice that in this example we should not take an infinite sum. A good example to keep in mind here is that of $V = \mathbf{R}[x]$, the space of polynomials of (finite) degree, and $\mathcal{B} = \{1, x, x^2, \dots\}$. In this case, any finite sum of elements of $\mathcal{B}$ is a polynomial in $\mathbf{R}[x]$, but infinite sums are not.

Above we defined the dimension of a vector space as the number of elements in a basis. To show this is a consistent definition, we need to show that any two bases have the same number of elements, which we will prove in two steps.

Theorem 4. (*The Replacement Theorem*) Let V be a vector space and let $\mathcal{B} = \{v_1, \dots, v_n\}$ be a basis for V with n elements. Choose an integer $m \leq n$ and let $S = \{w_1, \dots, w_m\}$ be a finite set of linearly independent vectors. Then there is a subset $\hat{S} \subset \mathcal{B}$ containing exactly $n - m$ elements such that $\text{span}(\{S \cup \hat{S}\}) = V$.

Remark 3. Some people like to call this theorem the "Exchange Theorem."

We will prove theorem by induction. As you might recall from MAM1000, the general idea behind induction is that you first prove a base case (in this case, $m = 0$). Next you prove that if the theorem is true for some integer m it is also true for $m + 1$. This completes the proof in the following way. If you want to conclude the statement of the theorem when $m = 2$, you only need to know it's true for $m = 1$. However, if you want to conclude the statement of the theorem when $m = 1$, you only need to know it's true when $m = 0$, which is the base case we proved in the first place. The statement of the theorem is true for $m = 0$, which implies it's true for $m = 1$, which in turn implies it's true for $m = 2$.

Proof. First consider the case of $m = 0$. Then $S = \emptyset$, and we choose $\hat{S} = \mathcal{B}$.

Now we suppose the statement of the theorem holds for some value of $m < n$, and we wish to prove it for $m + 1$. Let $S = \{w_1, \dots, w_{m+1}\}$ be a linearly independent set with $m + 1$ elements, and let $S_1 = \{w_1, \dots, w_m\}$. By the induction hypothesis (*i.e.* our assumption that the statement of the theorem is true for m), there is a subset $\hat{S}_1 = \{v_1, \dots, v_{n-m}\} \subset \mathcal{B}$ such that $\text{span}(\{S_1 \cup \hat{S}_1\}) = V$. In particular, we can write

$$w_{m+1} = a_1 w_1 + a_2 w_2 + \dots + a_m w_m + b_1 v_1 + b_2 v_2 + \dots + b_{n-m} v_{n-m},$$

for some coefficients $a_1, \dots, a_m, b_1, \dots, b_{n-m}$. Observe that, because S is linearly independent, w_{m+1} is not in the space of $S_1 = \{w_1, \dots, w_m\}$. This means at least one of the coefficients $b_1, \dots, b_{n-m}$ is non-zero; without loss of generality, we take

that coefficient to be b_1 . This means we can now rewrite the equation above as

$$v_1 = -\frac{1}{b_1} [a_1 w_1 + a_2 w_2 + \cdots a_m w_m - w_{m+1} + b_2 v_2 + b_3 v_3 + \cdots + b_{n-m} v_{n-m}].$$

We conclude that

$$v_1 \in \text{span}\{w_1, \dots, w_m, w_{m+1}, v_2, \dots, v_{n-m}\},$$

and so

$$V = \text{span}\{w_1, \dots, w_m, v_1, \dots, v_{n-m}\} \subset \text{span}\{w_1, \dots, w_{m+1}, v_2, \dots, v_{n-m}\}.$$

However, this latter space $\text{span}\{w_1, \dots, w_{m+1}, v_2, \dots, v_{n-m}\}$ is formed by linear combination of vectors in V , so it must be contained in V . We conclude that

$$\text{span}\{w_1, \dots, w_{m+1}, v_2, \dots, v_{n-m}\} = V,$$

as we claimed. $\square$

Theorem 5. *Let V be a vector space and let $\mathcal{B} = \{v_1, \dots, v_n\}$ be a basis for V with n elements. If S is any linearly independent set with exactly n elements, then S is also a basis.*

Proof. Let $S = \{w_1, \dots, w_n\}$ be a linearly independent set with exactly n elements. Applying the replacement theorem, we see there is a set $\hat{S}$ with $n - n = 0$ elements such that $\text{span}(S \cup \hat{S}) = V$. However, since $\hat{S}$ has zero elements, we conclude that $\text{span}(S) = V$, and so (by definition) S is a basis. $\square$

Theorem 6. *If V is a vector space with a basis $\mathcal{B} = \{v_1, \dots, v_n\}$ then any set $S = \{w_1, \dots, w_{n+1}\}$ of $n + 1$ elements is linearly dependent. As a result, any linearly independent set in V has at most n elements.*

Proof. Suppose that $S = \{w_1, \dots, w_{n+1}\}$ is linearly independent, and let $S_1 = \{w_1, \dots, w_n\}$. By the previous theorem, S_1 must be a basis, and so (in particular) we can write

$$w_{n+1} = a_1 w_1 + \cdots + a_n w_n,$$

which contradicts the assumption that S is linearly independent. $\square$

Theorem 7. *If V is a vector space with a basis $\mathcal{B} = \{v_1, \dots, v_n\}$ then any basis of V will have n elements.*

Proof. Let $\mathcal{C} = \{w_1, \dots, w_m\}$ be another basis. Then the previous theorem implies that both $m \leq n$ and $n \leq m$, which is only possible if $m = n$. $\square$

Example: Recall that $\mathbf{R}_n[x]$, the set of polynomials of degree at most n , is a vector space. We form the basis $\mathcal{B} = \{1, x, x^2, \dots, x^n\}$. First observe that each element of $\mathcal{B}$ is in $\mathbf{R}_n[x]$. Next, each is a polynomial of different degree, so if

$$0 = a_0 + a_1 x + a_2 x^2 + \cdots + a_n x^n$$

for all x then we must have $a_0 = a_1 = a_2 = \cdots = a_n = 0$, so $\mathcal{B}$ is linearly independent. Finally, we can write any polynomial of degree at most n as

$$p = a_0 + a_1 x + \cdots + a_n x^n$$

for some coefficients $a_0, a_1, \dots, a_n$, which means $\mathcal{B}$ spans $\mathbf{R}_n[x]$. Therefore $\mathcal{B}$ is a basis for $\mathbf{R}_n[x]$, and so $\dim(\mathbf{R}_n[x]) = \#(\mathcal{B}) = n + 1$.

Exercise: Recall that the space $\mathbf{R}_2[x] = \{a_0 + a_1x + a_2x^2\}$ of quadratic polynomials is a vector space. Show that

$$\{p_1(x) = 1 - x^2, p_2(x) = 1 + x^2, p_3(x) = x + x^2\}$$

is a basis for $\mathbf{R}_2[x]$. Find the three numbers a_1, a_2, a_3 such that

$$q(x) = x^2 + 4x + 4 = a_1p_1(x) + a_2p_2(x) + a_3p_3(x).$$

Exercise: Let $v_1 = \begin{bmatrix} 1 \\ 0 \\ 1 \end{bmatrix}$, $v_2 = \begin{bmatrix} 0 \\ 1 \\ 1 \end{bmatrix}$, and $v_3 = \begin{bmatrix} t \\ t \\ 1 \end{bmatrix}$. Find necessary and

sufficient conditions on t so that $\{v_1, v_2, v_3\}$ are linearly independent, and prove that in this case $\text{span}(v_1, v_2, v_3) = \mathbf{R}^3$. (Hint: it might be easier to find conditions that v_1, v_2, v_3 are linearly dependent.)

Exercise: Let v_1, v_2, v_3 be vectors in $\mathbf{R}^2$. Is it possible that they are all linearly independent? Find a necessary and sufficient condition that $\text{span}(v_1, v_2, v_3) = \mathbf{R}^2$.

Example: Let $V = M_{m \times n}$ be the space of matrices with real entries with m rows and n columns. We show that V is a vector space, and $\dim(V) = mn$. Indeed, we construct a basis

$$\mathcal{B} = \{A^{ij} : 1 \leq i \leq m, 1 \leq j \leq n\}, \quad [A^{ij}]_{kl} = \begin{cases} 1 & i = k \text{ and } l = j \\ 0 & i \neq k \text{ or } l \neq j \end{cases}$$

In other words, A^{ij} is the matrix which has a 1 in the i th row, j th column, and 0 elsewhere. It is now a straightforward exercise (which you should do!) to verify that $\mathcal{B}$ is a basis for $M_{m \times n}$.

The previous exercises highlight the connection between abstract vector spaces and matrices. Indeed, we just saw that in order to express the polynomial $q(x)$ as a linear combination of the three basis vectors p_1, p_2, p_3 , we had to solve a system of three linear equations in three unknowns. You probably solved that system of equations by setting up a matrix equation, where the matrix was 3×3 .

We will return to this point later, but here are some questions to keep in the back of your mind.

- If you're given a basis $\mathcal{B} = \{v_1, \dots, v_n\}$ and another set $S = \{w_1, \dots, w_m\}$, can you always solve the system of equations to find the coefficients $b_{i1}, \dots, b_{in}$, so that

$$w_i = b_{i1}v_1 + b_{i2}v_2 + \dots + b_{in}v_n?$$

- If you can solve the system of equations above, is the solution unique?
- Suppose I have some set $\mathcal{B} = \{v_1, \dots, v_n\}$ and another vector w . When can I solve the system of equations to find the coefficients c_j so that

$$w = c_1v_1 + c_2v_2 + \dots + c_nv_n?$$

- In the above question, when is the solution unique?

The answer to all the questions immediately above is encoded in the following theorem.

Theorem 8. Let V be a vector space and let $\mathcal{B} = \{v_1, \dots, v_n\}$ be a basis for V . Given any other vector $w \in V$, there is a unique choice of coefficients $a_1, a_2, \dots, a_n$ such that

$$w = \sum_{i=1}^n a_i v_i = a_1 v_1 + a_2 v_2 + \dots + a_n v_n.$$

Proof. The fact that $\mathcal{B}$ spans V implies that there is some choice of coefficients $a_1, \dots, a_n$ such that

$$w = a_1 v_1 + \dots + a_n v_n.$$

It remains to show uniqueness. Suppose that there are some other coefficients $b_1, \dots, b_n$ such that

$$w = b_1 v_1 + \dots + b_n v_n.$$

Subtracting, we then have

$$0 = (a_1 v_1 + \dots + a_n v_n) - (b_1 v_1 + \dots + b_n v_n) = (a_1 - b_1) v_1 + \dots + (a_n - b_n) v_n.$$

We have now produced a linear combination of $v_1, \dots, v_n$ which sums to zero. However, $\mathcal{B}$ is linearly independent, so all the coefficients $a_i - b_i = 0$, for $i = 1, 2, 3, \dots, n$, which in turn implies $a_i = b_i$ for $i = 1, 2, 3, \dots, n$. $\square$

If we're given a vector space and a basis, we can now define an element in $\mathbf{R}^n$ associated to any $v \in V$ as follows.

Definition 6. Let V be a vector space with basis $\mathcal{B} = \{v_1, \dots, v_n\}$. Given $v \in V$ we define $[v]_{\mathcal{B}} = (a_1, a_2, \dots, a_n) \in \mathbf{R}^n$, where $v = a_1 v_1 + a_2 v_2 + \dots + a_n v_n$. We call the coefficients $a_1, a_2, \dots, a_n$ the coordinates of V with respect to the basis $\mathcal{B}$, and call $[v]_{\mathcal{B}}$ the coordinate representation to v .

The previous theorem tells us this choice of coefficients is unique, so $[v]_{\mathcal{B}}$ is well-defined.

We close this section with two examples. First, we reconsider the complex numbers as 2×2 matrices. Recall that $z \in \mathbf{C}$ has the form $z = x + iy$, where $x, y \in \mathbf{R}$ and $i^2 = -1$. Also, complex multiplication is defined as

$$z \cdot w = (x + iy)(u + iv) = (xu - yv) + i(yu + xv).$$

If we consider the complex numbers as two copies of the real numbers, we can write (using vector notation)

$$z = \begin{bmatrix} x \\ y \end{bmatrix}, \quad w = \begin{bmatrix} u \\ v \end{bmatrix}, \quad z \cdot w = \begin{bmatrix} xu - yv \\ yu + xv \end{bmatrix}.$$

Observe that, if we fix $z = x + iy$, this gives us a set of two linear equations for u and v , and we can rewrite this as

$$z \cdot w = \begin{bmatrix} xu - yv \\ yu + xv \end{bmatrix} = \begin{bmatrix} x & -y \\ y & x \end{bmatrix} \cdot \begin{bmatrix} u \\ v \end{bmatrix}.$$

(You might want to check the computation is correct.) This means we can identify the complex number $z = x + iy$ with the 2×2 matrix

$$\begin{bmatrix} x & -y \\ y & x \end{bmatrix}.$$

We can rewrite this matrix further if we let $r = \sqrt{x^2 + y^2} = |z|$, $\cos \theta = \frac{x}{r}$, and $\sin \theta = \frac{y}{r}$. Now we

$$z \simeq \begin{bmatrix} x & -y \\ y & x \end{bmatrix} = r \begin{bmatrix} \cos \theta & -\sin \theta \\ \sin \theta & \cos \theta \end{bmatrix},$$

and we can recognize complex multiplication as nothing more than a (counterclockwise) rotation followed by a dilation. This is a that the complex numbers carry some special algebraic structure.

For our second example, we discuss the Lagrange interpolation formula. You know that two points determine a line (in the plane), and if I hand you two points you can find the line passing through both of them. You'd probably imagine that three points determine a parabola, and that if I hand you three points then (with a little more work) you can find the parabola passing through all three of them. What if I hand you 17 points and ask you to find the degree 16 polynomial passing through them. Can you do that?

It turns out this is not too hard, if you use the following clever basis for $\mathbf{R}_{16}[x]$, the space of degree 16 polynomials. We will outline the process for a polynomial of some arbitrary (fixed) degree n . Start by choosing $n + 1$ distinct real numbers $c_0, c_1, \dots, c_n$; these are the points where you will evaluate the polynomial (*i.e.* the x -coordinates of the $n + 1$ points I hand to you). Now define the polynomials

$$p_i(x) = \frac{(x - c_0) \cdots (x - c_{i-1})(x - c_{i+1}) \cdots (x - c_n)}{(c_i - c_0) \cdots (c_i - c_{i-1})(c_i - c_{i+1}) \cdots (c_i - c_n)} = \prod_{j=0, j \neq i}^{j=n} \frac{x - c_j}{c_i - c_j}.$$

(Here the Π means take a product of all those terms.) For instance,

$$p_0(x) = \frac{(x - c_1) \cdots (x - c_n)}{(c_0 - c_1) \cdots (c_0 - c_n)}, \quad p_n(x) = \frac{(x - c_0) \cdots (x - c_{n-1})}{(c_n - c_0) \cdots (c_n - c_{n-1})}.$$

Observe that

$$p_i(c_j) = \begin{cases} 0 & i \neq j \\ 1 & i = j. \end{cases}$$

We will use this property to show that $\mathcal{B} = \{p_0, \dots, p_n\}$ is a basis for $\mathbf{R}_n[x]$, the degree n polynomials in the variable x . First suppose that there are coefficients $a_0, a_1, \dots, a_n$ such that

$$a_0 p_0(x) + a_1 p_1(x) + \cdots + a_n p_n(x) = 0$$

for all x . Evaluate this expression at each c_j , so we get

$$0 = a_0 p_0(c_j) + \cdots + a_n p_n(c_j) = a_j$$

for each j , which implies (by definition) that $\{p_0, \dots, p_n\}$ are linearly independent. We already know that $\dim(\mathbf{R}_n[x]) = n + 1$, and $\mathcal{B} = \{p_0, \dots, p_n\}$ contains exactly $n + 1$ elements, so $\mathcal{B}$ must be a basis.

Now we return to the original problem. We have $n + 1$ points of the form $(c_0, y_0), (c_1, y_1), \dots, (c_n, y_n)$, and we want to find the polynomial $q(x)$ of degree n such that

$$q(c_0) = y_0, \quad q(c_1) = y_1, \quad \dots, \quad q(c_n) = y_n.$$

This is simple now, we must have

$$q(x) = y_0 p_0(x) + y_1 p_1(x) + \dots + y_n p_n(x).$$

Exercise: Show that the graph of the polynomial $q(x)$ given above passes through all the points $(c_0, y_0), \dots, (c_n, y_n)$.

Exercise: Show that the polynomial $q(x)$ given above is the **only** degree n polynomial passing through all the points $(c_0, y_0), \dots, (c_n, y_n)$.

This is the Lagrange interpolation formula, and it is very useful for finding nice functions passing through a given set of points.

Exercise: Show the following: if $f \in \mathbf{R}_n[x]$ is a degree n polynomial and $f(c_0) = f(c_1) = \dots = f(c_n) = 0$ for $n + 1$ distinct numbers $c_0, \dots, c_n$ then f is the zero function.

3.4. Subspaces revisited.

Proposition 9. *Let V be a finite dimensional vector space and let $W \subset V$ be a subspace. Then W is also finite dimensional, and $\dim(W) \leq \dim(V)$. Moreover, if $\dim(W) = \dim(V)$ then $W = V$.*

One can find a quick proof of this result using Theorem 6, but we offer another proof for the readers enlightenment.

Proof. Let $\dim(V) = n$. If $W = \{0\}$, then by inspection W is finite-dimensional, and $\dim(W) = 0$. If $W \neq \{0\}$ then it contains a nonzero vector w_1 , and so $\{w_1\}$ is a linearly independent set. If $W = \text{span}\{w_1\}$, then we stop, and conclude $\dim(W) = 1$. Otherwise, we continue to add elements $w_2, w_3, \dots, w_k$ to our list, until $W = \text{span}\{w_1, w_2, \dots, w_k\}$, all the while maintaining that $\{w_1, \dots, w_k\}$ is a linearly independent set. Moreover, by the replacement theorem, we cannot have $k \geq n$, because otherwise we'd have a linearly independent set in V with more than n elements. We have just shown that $\dim(W) = k \leq n$, and in particular that W is finite dimensional.

Finally, if $\dim(W) = n$ then we have a basis $\{w_1, \dots, w_n\}$ for W with n elements. By the replacement theorem, this is also a basis for V , and so $W = V$. $\square$

Example: Let $W \subset \mathbf{R}^5$ be defined by

$$W = \{(x_1, \dots, x_5) : x_1 + x_3 + x_5 = 0, x_2 = x_4\}.$$

First we show that W is a subspace. Indeed, by inspection $0 \in W$. If $x, y \in W$ then we have

$$x_1 + x_3 + x_5 = 0 = y_1 + y_3 + y_5, \quad x_2 = x_4, \quad y_2 = y_4.$$

This implies

$$(x_1 + y_1) + (x_3 + y_3) + (x_5 + y_5) = (x_1 + x_3 + x_5) + (y_1 + y_3 + y_5) = 0, \quad x_2 + y_2 = x_4 + y_4,$$

and so $x + y \in W$. Similarly, $ax + by \in W$ for any $a, b \in \mathbf{R}$.

Next we find a basis for W . We can find 3 linearly independent vectors in W :

$$w_1 = (0, 1, 0, 1, 0), \quad w_2 = (1, 0, 0, 0, -1), \quad w_3 = (0, 0, 1, 0, -1).$$

Suppose $\mathcal{B} = \{w_1, w_2, w_3\}$ were not a linearly independent set. Then we could find three numbers a_1, a_2, a_3 , not all of which are zero, such that

$$(0, 0, 0, 0, 0) = a_1 w_1 + a_2 w_2 + a_3 w_3 = (a_2, a_1, a_3, a_1, -a_2 - a_3).$$

However, we can read off from the first three coefficients that $a_1 = a_2 = a_3 = 0$. We conclude that $\mathcal{B}$ is linearly independent.

Does $\mathcal{B} = \{w_1, w_2, w_3\}$ span W ? Write some element of W as $w = (a_1, a_2, a_3, a_4, a_5)$; we claim that

$$w = a_2 w_1 + a_1 w_2 + a_3 w_3 = (a_1, a_2, a_3, a_2, -a_1 - a_3).$$

Indeed, the first three components match, so we only need to verify the fourth and fifth components. We check that, because $w \in W$ we must have $a_2 = a_4$, and so the fourth component matches. Similarly, we must have $a_5 = -a_1 - a_3$, and so the fifth component matches.

In particular, we see that $\dim(W) = 3$.

Exercise: Is it an accident that $\dim(W) = 3 = 5 - 2$ and W is a subspace of $\mathbf{R}^5$ determined by 2 linear equations?

Example: We already know that the set of $n \times n$ matrices is a vector space of dimension n^2 . Let V be the set of symmetric $n \times n$ matrices. That is, $A = [a_{ij}] \in V$ if and only if $a_{ij} = a_{ji}$, *i.e.* $A^t = A$ (here A^t is the transpose of A , the matrix you get by swapping rows and columns). First we show V is a subspace. First of all, if $Z = [z_{ij}]$ is the zero matrix, that is $z_{ij} = 0$ for all i, j , then we have $z_{ij} = 0 = z_{ji}$. Thus the zero vector (in the vector space of $n \times n$ matrices) is also in V . Next let $A = [a_{ij}] \in V$ and $B = [b_{ij}] \in V$, and observe $a_{ij} = a_{ji}, b_{ij} = b_{ji}$. Then, for any $\alpha, \beta \in \mathbf{R}$ we have

$$(\alpha A + \beta B)_{ij} = \alpha a_{ij} + \beta b_{ij} = \alpha a_{ji} + \beta b_{ji} = (\alpha A + \beta B)_{ji},$$

and so $\alpha A + \beta B \in V$. Thus V is a vector subspace of the space of $n \times n$ matrices.

Now we find a basis for V . For $i \geq j$ we define the $n \times n$ matrix

$$A^{ij} = [(a^{ij})_{kl}], \quad (a^{ij})_{kl} = \begin{cases} 1 & i = k, j = l \\ 1 & i = l, j = k \\ 0 & \text{otherwise.} \end{cases}$$

If $i \neq i'$ or $j \neq j'$ then A^{ij} and $A^{i'j'}$ have 1's in different entries, and so the set

$$\mathcal{B} = \{A^{ij} : 1 \leq i \leq n, 1 \leq j \leq n, i \geq j\}$$

is linearly independent. Now let $A = [a_{ij}] \in V$. It is easy to check that

$$A = \sum_{i=1}^n \sum_{j=1}^i a_{ij} A^{ij},$$

and so $V = \text{span}(\mathcal{B})$. Thus $\mathcal{B}$ is a basis for V , and, in particular,

$$\dim(V) = \#\mathcal{B} = n + (n-1) + \cdots + 2 + 1 = \frac{n(n+1)}{2}.$$

Exercise: Repeat the exercise directly above with W , the set of skew-symmetric matrices. That is, $A = [a_{ij}] \in W$ if and only if $a_{ij} = -a_{ji}$.

4. LINEAR TRANSFORMATIONS

By themselves, vector spaces are boring. This subject only becomes alive when we talk about linear transformations, which are precisely those transformations between vector spaces that preserve the vector space structure.

4.1. Definitions and examples.

Definition 7. Let V and W be vector spaces. A transformation (i.e. mapping, i.e. function) $T : V \rightarrow W$ is linear if for all $v_1, v_2 \in V$ and $a_1, a_2 \in \mathbf{R}$ we have

$$T(a_1v_1 + a_2v_2) = a_1T(v_1) + a_2T(v_2).$$

Example: We use induction to show that if $T : V \rightarrow W$ is linear then

$$T(a_1v_1 + a_2v_2 + \cdots + a_kv_k) = a_1T(v_1) + a_2T(v_2) + \cdots + a_kT(v_k).$$

As our base case, we have

$$T(a_1v_1 + a_2v_2) = a_1T(v_1) + a_2T(v_2)$$

by definition. Now suppose

$$T(a_1v_1 + \cdots + a_kv_k) = a_1T(v_1) + \cdots + a_kT(v_k)$$

for some $k \geq 2$. Then

$$\begin{aligned} T(a_1v_1 + \cdots + a_kv_k + a_{k+1}v_{k+1}) &= T((a_1v_1 + \cdots + a_kv_k) + a_{k+1}v_{k+1}) \\ &= T(a_1v_1 + \cdots + a_kv_k) + a_{k+1}T(v_{k+1}) \\ &= a_1T(v_1) + \cdots + a_kT(v_k) + a_{k+1}T(v_{k+1}). \end{aligned}$$

This completes the induction step, so we've proved the result we wanted to prove.

Example: Let V and W be vector spaces. There are two very simple, yet important, examples of linear transforms which we should mention to begin. First,

$$I : V \rightarrow V, \quad I(v) = v$$

is the identity transformation. This mapping always sends a vector v to itself. The second example is

$$Z : V \rightarrow W, \quad Z(v) = 0,$$

is the zero mapping. This sends every vector v to the zero vector 0 (in W). We verify this quickly:

$$I(a_1v_1 + a_2v_2) = a_1v_1 + a_2v_2 = a_1I(v_1) + a_2I(v_2)$$

and

$$Z(a_1v_1 + a_2v_2) = 0 = a_1 \cdot 0 + a_2 \cdot 0 = a_1Z(v_1) + a_2Z(v_2).$$

Lemma 10. Let $T : V \rightarrow W$ be a linear transformation. Then $T(0) = 0$.

Proof.

$$T(0) = T(0 + 0) = T(0) + T(0) \Rightarrow 0 = T(0)$$

□

Example: Let $\mathbf{R}_n[x]$ be the set of polynomials in the variable x of degree at most n . The transformation

$$T : \mathbf{R}_n[x] \rightarrow \mathbf{R}_{n-1}[x], \quad T(p) = p'$$

is linear. Indeed, for any two polynomials p and q and real numbers a and b we have

$$T(ap + bq) = (ap + bq)' = ap' + bq' = aT(p) + bT(q).$$

Example: Let $\mathcal{C}[0, 1]$ be the space of continuous function on the interval $[0, 1]$; that is, any continuous function $f(x)$, defined for $0 \leq x \leq 1$, is in the space $\mathcal{C}[0, 1]$. we saw in class that this is a vector space. The map

$$T : \mathcal{C}[0, 1] \rightarrow \mathbf{R}, \quad T(f) = \int_0^1 f(x)dx$$

is linear. Indeed, for any two continuous functions f and g and real numbers a and b we have

$$T(af + bg) = \int_0^1 (af(x) + bg(x))dx = a \int_0^1 f(x)dx + b \int_0^1 g(x)dx = aT(f) + bT(g).$$

Example: Let $\mathcal{C}^2[0, 1]$ be the space of functions $f(x)$ defined for $0 \leq x \leq 1$ such that f'' exists and is continuous. The map

$$T : \mathcal{C}^2[0, 1] \rightarrow \mathcal{C}[0, 1], \quad T(f) = f'' - 3f' + 2f$$

is a linear transformation

4.2. Kernel and range. There are two important vector subspaces associated to any linear transformation. We define them now.

Definition 8. Let $T : V \rightarrow W$ be a linear transformation. Define the sets

$$\ker(T) = \{v \in V : T(v) = 0\}, \quad \mathcal{R}(T) = \{w \in W : w = T(v) \text{ for some } v \in V\}.$$

We call $\ker(T)$ the "kernel" of T and $\mathcal{R}(T)$ the "range" of T . (The range is not to be confused with the target W , which in general is a larger vector space.)

Theorem 11. If $T : V \rightarrow W$ is linear then $\ker(T)$ is a subspace of V and $\mathcal{R}(T)$ is a subspace of W .

Proof. First, we have $T(0) = 0$, which implies $0 \in \ker(T)$ and $0 \in \mathcal{R}(T)$. Next, suppose $v_1, v_2 \in \ker(T)$. Then, for any $a_1, a_2 \in \mathbf{R}$, we have

$$T(a_1v_1 + a_2v_2) = a_1T(v_1) + a_2T(v_2) = a_1 \cdot 0 + a_2 \cdot 0 = 0,$$

and so $a_1v_1 + a_2v_2 \in \ker(T)$. We conclude that $\ker(T)$ is a subspace of V . Finally, suppose that $w_1, w_2 \in \mathcal{R}(T)$. Then there are $v_1, v_2 \in V$ such that $T(v_1) = w_1$ and $T(v_2) = w_2$. Now, for any $a_1, a_2 \in \mathbf{R}$, we have

$$a_1w_1 + a_2w_2 = a_1T(v_1) + a_2T(v_2) = T(a_1v_1 + a_2v_2) \in \mathcal{R}(T),$$

because $a_1v_1 + a_2v_2 \in V$. □

Remark 4. In general, this is the most useful tool to prove something is a vector space.

Example: We know one can write the equation of any plane passing through the origin in $\mathbf{R}^3$ as

$$n_1x_1 + n_2x_2 + n_3x_3 = 0,$$

where $n = (n_1, n_2, n_3)$ is the normal vector to the plane. This formulation tells us immediately that any plane through the origin is a subspace of $\mathbf{R}^3$, because it is the kernel of the linear transformation

$$T : \mathbf{R}^3 \rightarrow \mathbf{R}, \quad T(x_1, x_2, x_3) = n_1x_1 + n_2x_2 + n_3x_3.$$

Exercise: Mimic the example above to show that any line passing through $(0, 0)$ in the plane is a vector space. Also, prove that any line passing through the origin in three-space is a vector space.

Example: Above we considered the linear transformation

$$T : \mathcal{C}[0, 1] \rightarrow \mathbf{R}, \quad T(f) = \int_0^1 f(x)dx.$$

It is easy to check that $\mathcal{R}(T) = \mathbf{R}$. Indeed, the target $\mathbf{R}$ is a one-dimensional vector space, so we must either have $\mathcal{R}(T) = \mathbf{R}$ or $\mathcal{R}(T) = \{0\}$. Also, $T(2x) = 1 \neq 0$, so $\mathcal{R}(T) \neq \{0\}$. Thus we must have $\mathcal{R}(T) = \mathbf{R}$. The kernel $\ker(T)$ is the space of functions on $[0, 1]$ with average value 0.

Example: Recall that $\mathcal{C}^2[0, 1]$, the space of functions $f(x)$, for $0 \leq x \leq 1$, with continuous second derivatives, is a vector space. We define the linear transformation

$$T : \mathcal{C}^2[0, 1] \rightarrow \mathcal{C}[0, 1], \quad T(f) = f'' + (1 + x^2)f = 0.$$

The set of functions

$$\ker(T) = \{f \in \mathcal{C}^2[0, 1] : T(f) = f'' + (1 + x^2)f = 0\}$$

is the kernel of a linear transformation, so it is a vector space.

Exercise: Use the same technique to prove that the set of solutions to any homogeneous, second order ordinary differential equation,

$$K = \{f \in \mathcal{C}^2[0, 1] : f'' + p(x)f' + q(x)f = 0\}$$

is a vector space. Can you generalize this result at all?

Theorem 12. Let $T : V \rightarrow W$ be linear and let $\mathcal{B} = \{v_1, \dots, v_n\}$ be a basis for V . Then

$$\mathcal{R}(T) = \text{span}\{T(v_1), T(v_2), \dots, T(v_n)\}.$$

Proof. Each $T(v_i) \in \mathcal{R}(T)$, and $\mathcal{R}(T)$ is a subspace of W , so we must have

$$\text{span}\{T(v_1), \dots, T(v_n)\} \subset \mathcal{R}(T).$$

Conversely, suppose $w \in \mathcal{R}(T)$. Then $w = T(v)$ for some $v \in V$, but

$$v = a_1v_1 + \dots + a_nv_n.$$

Therefore,

$$w = T(a_1v_1 + \cdots + a_nv_n) = a_1T(v_1) + \cdots + a_nT(v_n) \in \text{span}\{T(v_1), \dots, T(v_n)\},$$

and so

$$\mathcal{R}(T) \subset \text{span}\{T(v_1), \dots, T(v_n)\}.$$

□

Example: Let $M_{2 \times 2}$ be the vector space of 2×2 matrices and let $\mathbf{R}_2[x]$ be the space of quadratic polynomials. Consider

$$T : \mathbf{R}_2[x] \rightarrow M_{2 \times 2}, \quad T(p) = \begin{bmatrix} p(1) - p(2) & 0 \\ 0 & p(0) \end{bmatrix}.$$

We wish to find $\mathcal{R}(T)$. Recall that $\mathcal{B} = \{1, x, x^2\}$ is a basis for $\mathbf{R}_2[x]$, so we can find $\mathcal{R}(T)$ by finding taking the span of the image of our basis elements under T . Thus we have

$$\begin{aligned} \mathcal{R}(T) &= \text{span}\{T(1), T(x), T(x^2)\} \\ &= \text{span}\left\{\begin{bmatrix} 0 & 0 \\ 0 & 1 \end{bmatrix}, \begin{bmatrix} -1 & 0 \\ 0 & 0 \end{bmatrix} \begin{bmatrix} -3 & 0 \\ 0 & 0 \end{bmatrix}\right\} \\ &= \text{span}\left\{\begin{bmatrix} 0 & 0 \\ 0 & 1 \end{bmatrix}, \begin{bmatrix} -1 & 0 \\ 0 & 0 \end{bmatrix}\right\}. \end{aligned}$$

In fact, the computation above finds a basis of two elements for $\mathcal{R}(T)$, and so $\dim(\mathcal{R}(T)) = 2$.

Theorem 13. *Let $T : V \rightarrow W$ be a linear transformation. Then T is one-to-one if and only if $\ker(T) = \{0\}$.*

Proof. First suppose $\ker(T) = \{0\}$ and let $T(v_1) = T(v_2)$. Then

$$0 = T(v_1) - T(v_2) = T(v_1 - v_2) \Rightarrow v_1 - v_2 \in \ker(T) = \{0\}.$$

We conclude $v_1 - v_2 = 0$, i.e. that $v_1 = v_2$, and so T is one-to-one.

Conversely, suppose that T is one-to-one, and let $v \in \ker(T)$. We already know that $T(0) = 0$, and so (because T is one-to-one) we must have $v = 0$. □

The following theorem is the most important theorem you will learn in this linear algebra class; it is called the **Rank-Nullity Theorem**.

Theorem 14. (*Rank-Nullity Theorem*) *Let $T : V \rightarrow W$ be linear and suppose V is finite dimensional. Then*

$$\dim(\ker(T)) + \dim(\mathcal{R}(T)) = \dim(V).$$

Remark 5. • Notice that we do not require W to be finite dimensional.

- This is called the rank-nullity theorem because one can call $\dim(\mathcal{R}(T))$ the **rank** of T and $\dim(\ker(T))$ the **nullity** of T . Then the rank-nullity theorem reads

$$\text{rank}(T) + \text{nullity}(T) = n.$$

- Some people like to call this result "the dimension theorem."

Proof. Let $\dim(V) = n$, and choose a basis $\{v_1, \dots, v_k\}$ for $\ker(T)$. Notice that $\ker(T)$ is a subspace of V , so we necessarily have $k \leq n$, and $k = n$ if and only if $\ker(T) = V$. If $\ker(T) = V$, we must have $\mathcal{R}(T) = \{0\}$, which is 0-dimensional.

Otherwise, we have $k < n$, and (by the replacement theorem) we can choose $n - k$ linearly independent vectors $\{v_{k+1}, \dots, v_n\}$ such that $\{v_1, \dots, v_n\}$ forms a basis of V . Once we show that

$$\mathcal{B} = \{T(v_{k+1}), T(v_{k+2}), \dots, T(v_n)\}$$

is a basis for $\mathcal{R}(T)$, we're done, because in this case we show that

$$\dim(\ker(T)) + \dim(\mathcal{R}(T)) = k + n - k = n = \dim(V).$$

Let $w \in \mathcal{R}(T)$, and write $w = T(v)$. Then there are coefficients $a_1, \dots, a_n$ such that

$$v = a_1 v_1 + a_2 v_2 + \dots + a_n v_n.$$

Use the fact that $T(v_i) = 0$ for $i = 1, 2, \dots, k$ to see

$$w = T(v) = T\left(\sum_{i=1}^n a_i v_i\right) = \sum_{i=1}^n a_i T(v_i) = \sum_{i=k+1}^n a_i T(v_i),$$

which implies (together with the previous theorem) that

$$\mathcal{R}(T) = \text{span}\{T(v_{k+1}), \dots, T(v_n)\} = \text{span}(\mathcal{B}).$$

Finally, we show that $\mathcal{B}$ is linearly independent. Suppose there are coefficients $b_{k+1}, \dots, b_n$ such that

$$0 = \sum_{i=k+1}^n b_i T(v_i) = T\left(\sum_{i=k+1}^n b_i v_i\right).$$

This means $\sum_{i=k+1}^n b_i v_i \in \ker(T)$, and so we can find coefficients $a_1, a_2, \dots, a_k$ such that

$$\sum_{i=k+1}^n b_i v_i = \sum_{j=1}^k a_j v_j \Leftrightarrow 0 = -\sum_{j=1}^k a_j v_j + \sum_{i=k+1}^n b_i v_i.$$

However, $\{v_1, \dots, v_k, v_{k+1}, \dots, v_n\}$ is a basis for V , and in particular it is a linearly independent set, so we must have $b_{k+1} = 0, b_{k+2} = 0, \dots, b_n = 0$. We have just shown that $\mathcal{B}$ is basis for $\mathcal{R}(T)$, completing the proof. $\square$

Corollary 15. *Let V and W be finite dimensional vector spaces of equal dimension, and let $T : V \rightarrow W$ be linear. Then T is one-to-one if and only if T is onto.*

Compare this to the case of nonlinear functions $f : \mathbf{R} \rightarrow \mathbf{R}$. It is quite easy to find such functions that are one-to-one but not onto (such as $f(x) = e^x$) or onto but not one-to-one (such as $f(x) = x^3 - x$).

Proof. Let $n = \dim(V) = \dim(W)$. If T is one-to-one, then $\dim(\ker(T)) = 0$, and so by the rank-nullity theorem we have $\dim(\mathcal{R}(T)) = n = \dim(W)$. However, the only way $\mathcal{R}(T)$ can be a subspace of W of the same dimension is if $\mathcal{R}(T) = W$, i.e. if T is onto. Conversely, suppose T is onto. Then $\dim(\mathcal{R}(T)) = n$, and so

by the rank-nullity theorem $\dim(\ker(T)) = 0$, which means $\ker(T) = \{0\}$. This is equivalent to T being one-to-one. $\square$

Example: In the example above we considered

$$T : \mathbf{R}_2[x] \rightarrow M_{2 \times 2}, \quad T(p) = \begin{bmatrix} p(1) - p(2) & 0 \\ 0 & p(0) \end{bmatrix}.$$

We found that $\dim(\mathcal{R}(T)) = 2$, and we know that $\dim(\mathbf{R}_2[x]) = 3$. Thus the rank-nullity theorem tells us $\dim(\ker(T)) = 1$.

Exercise: Find a nonzero vector in the kernel of the map T above. (Hint: $T(x^2) = 3T(x)$.)

Example: Let $\mathbf{R}_2[x]$ be the space of quadratic polynomials and $\mathbf{R}_3[x]$ the space of cubic polynomials. Define

$$T : \mathbf{R}_2[x] \rightarrow \mathbf{R}_3[x], \quad T(p) = 2p' + \int_0^x 3p(t)dt.$$

We have

$$\mathcal{R}(T) = \text{span}\{T(1), T(x), T(x^2)\} = \text{span}\left\{3x, 2 + \frac{3}{2}x^2, 4x + x^3\right\},$$

and so $\dim(\mathcal{R}(T)) = 3$. By the rank-nullity theorem

$$\dim(\ker(T)) + 3 = \dim(\mathbf{R}_2[x]) = 3 \Rightarrow \dim(\ker(T)) = 0,$$

which implies T is one-to-one. On the other hand, $\dim(\mathbf{R}_3[x]) = 4 > \dim(\mathcal{R}(T))$, and so T is not onto.

The following theorem and its corollary describe how a linear transformation is uniquely determined by its action on a basis. In practical terms, this means you only need to compute what a linear transformation does to basis vectors; once you've done that you know everything about it.

Theorem 16. *Let V and W be vector spaces, and suppose V is finite dimensional with a basis $\mathcal{B} = \{v_1, v_2, \dots, v_n\}$. Then, for any choice of vectors $\{w_1, w_2, \dots, w_n\} \subset W$ there is a unique linear transformation $T : V \rightarrow W$ such that*

$$T(v_1) = w_1, \quad T(v_2) = w_2, \quad \dots, T(v_n) = w_n.$$

Proof. If $v \in V$ there is a unique choice of coefficients $a_1, \dots, a_n$ such that $v = \sum_{i=1}^n a_i v_i$, and then define

$$T(v) = T\left(\sum_{i=1}^n a_i v_i\right) = \sum_{i=1}^n a_i w_i.$$

It is clear that with this definition $T(v_i) = w_i$, as required. It remains to check that T is linear and that T is the only linear transformation such that maps v_i to w_i for all $i = 1, 2, \dots, n$.

We first check that T is linear. Let

$$v = \sum_{i=1}^n a_i v_i, \quad u = \sum_{i=1}^n b_i v_i.$$

Then

$$T(\alpha v + \beta u) = T\left(\sum_{i=1}^n (\alpha a_i + \beta b_i) v_i\right) = \sum_{i=1}^n (\alpha a_i + \beta b_i) w_i = \alpha T(v) + \beta T(u).$$

Now suppose there is another linear transformation $U : V \rightarrow W$ such that $U(v_i) = w_i$ for all $i = 1, 2, \dots, n$. Writing

$$v = \sum_{i=1}^n a_i v_i$$

again, we have (because U is linear)

$$U(v) = \sum_{i=1}^n a_i U(v_i) = \sum_{i=1}^n a_i w_i = T(v),$$

and so $T = U$. □

Corollary 17. *Let V and W be vector spaces, and suppose V is finite dimensional with a basis $\{v_1, \dots, v_n\}$. If $T, U : V \rightarrow W$ are two linear transformations such that $T(v_i) = U(v_i)$ for all $i = 1, 2, \dots, n$ then $T = U$.*

Example: We consider two bases in $\mathbf{R}^3$:

$$\begin{aligned}\mathcal{B} &= \{e_1 = (1, 0, 0), e_2 = (0, 1, 0), e_3 = (0, 0, 1)\} \\ \tilde{\mathcal{B}} &= \{v_1 = (1, 0, -1), v_2 = (1, 0, 1), v_3 = (0, 1, 0)\}.\end{aligned}$$

We wish to find a linear transformation such that $T(e_1) = v_1, T(e_2) = v_2, T(e_3) = v_3$. By the theorem above, we must have

$$T\left(\begin{bmatrix} a \\ b \\ c \end{bmatrix}\right) = T(a_1 e_1 + a_2 e_2 + a_3 e_3) = a_1 v_1 + a_2 v_2 + a_3 v_3 = \begin{bmatrix} a_1 + a_2 \\ a_3 \\ -a_1 + a_2 \end{bmatrix}.$$

In fact, if you multiply the matrices out, you'll find that

$$T\left(\begin{bmatrix} a_1 \\ a_2 \\ a_3 \end{bmatrix}\right) = \begin{bmatrix} 1 & 1 & 0 \\ 0 & 0 & 1 \\ -1 & 1 & 0 \end{bmatrix} \begin{bmatrix} a_1 \\ a_2 \\ a_3 \end{bmatrix}.$$

We have already seen this relationship between linear transformations and matrices in the case of linear transformations mapping $\mathbf{R}^n$ to $\mathbf{R}^m$. We will see in the next section that, after choosing bases for the domain and target of a linear transformation, one obtains the same sort of matrix representation of a linear mapping.

However, before we do that, we need several more terms.

Definition 9. *Let V, W be vector spaces. A linear transformation $T : V \rightarrow W$ which is both one-to-one and onto is called a (linear) isomorphism. Two vector spaces V, W are isomorphic if there exists a (linear) isomorphism $T : V \rightarrow W$. We write this condition as $V \simeq W$.*

Isomorphic is a Greek word; it means "same structure."

Corollary 18. *Let V, W be vector spaces and let $T : V \rightarrow W$ be linear. Then the restriction $T : V \rightarrow \mathcal{R}(T)$ is an isomorphism if and only if $\dim(\ker(T)) = 0$.*

This follows from the fact that T is one-to-one if and only if $\dim(\ker(T)) = 0$.

Corollary 19. *Let V, W be finite dimensional vector spaces of the same dimension, and let $T : V \rightarrow W$ be linear. Then T is an isomorphism if and only if T is one-to-one if and only if T is onto.*

Example: Let V and W be finite dimensional vector spaces, and let $\dim(V) \neq \dim(W)$. Then V and W cannot be isomorphic. Indeed, suppose that $T : V \rightarrow W$ is an isomorphism. Then, because T would have to be both one-to-one and onto, we must have

$$\dim(V) = \dim(\mathcal{R}(T)) + \dim(\ker(T)) = \dim(W) + 0,$$

which is impossible.

Corollary 20. *Let V be a finite dimensional vector space with $\dim(V) = n$. Then V is isomorphic to $\mathbf{R}^n$.*

Proof. Choose a basis $\mathcal{B} = \{v_1, v_2, \dots, v_n\}$ for V and let $\{e_1, e_2, \dots, e_n\}$ be the standard basis for $\mathbf{R}^n$; that is, e_i is the vector with a 1 in the i th component and a 0 everywhere else. By Theorem 16 there is a linear map $T : \mathbf{R}^n \rightarrow V$ with $T(e_i) = v_i$ for $i = 1, 2, \dots, n$. The map T is onto because $\mathcal{B}$ spans V , and it is one-to-one because $\mathcal{B}$ is linearly independent. $\square$

Remark 6. *If all finite dimensional are the same (i.e. isomorphic) to some Euclidean space, as we have just shown, then why do we study abstract vector spaces at all? The answer has to do with the way we constructed the isomorphism in the proof above. To prove our corollary, we had to choose a basis. Indeed, if we choose a different basis for V then we will get a different isomorphism $\mathbf{R}^n \rightarrow V$. This freedom to choose the representation of our vector space is very powerful, and changing the representation can turn a difficult problem into an easy one. Geometrically and physically, you can think of the following analogy. Some problems are easier in the usual Euclidean coordinates, and some are easier in rotated coordinates. For instance, if you're tracking the motion of a planet around a star, you'd certainly want to choose coordinates so that the motion of the planet lies in a coordinate plane. This choice of a particular coordinate system is (at least locally) the same as choosing a basis for the vector space $\mathbf{R}^3$.*

4.3. Matrix representations of general linear transformations. In fact, one can represent a linear transformation between any two finite dimensional vector space as a matrix, after one chooses bases.

Let V and W be finite dimensional vector spaces, and choose a basis $\mathcal{B} = \{v_1, v_2, \dots, v_n\}$ for V and a basis $\mathcal{A} = \{w_1, w_2, \dots, w_m\}$ for W . By Theorem 16 we can represent any linear transformation

$$T : V \rightarrow W$$

as

$$T(v_j) = \sum_{i=1}^m a_{ij}w_i = a_{1j}w_1 + \cdots + a_{mj}w_m.$$

We have now represented T by the $m \times n$ matrix $[T]_{\mathcal{B}}^{\mathcal{A}} = [a_{ij}]$, which has a_{ij} as the entry in its i th row and j th column. This is a straight-forward generalization of the previous subsection, using the fact that $V \simeq \mathbf{R}^n$ and $W \simeq \mathbf{R}^m$. In fact, choosing the bases $\{v_1, \dots, v_n\}$ and $\{w_1, \dots, w_m\}$ realizes these isomorphisms, which in turn determines the matrix $[T]_{\mathcal{B}}^{\mathcal{A}}$. In the next section we will consider the effect changing basis has on $[T]_{\mathcal{B}}^{\mathcal{A}}$, but first we will write out some examples. We will also say something brief about the set of linear transformations between V and W .

Example: Above we considered

$$T : \mathbf{R}_2[x] \rightarrow M_{2 \times 2}, \quad T(p) = \begin{bmatrix} p(1) - p(2) & 0 \\ 0 & p(0) \end{bmatrix}.$$

In order to find the matrix $[T]$ we need to choose bases for $\mathbf{R}_2[x]$ and $M_{2 \times 2}$. We choose

$$\mathcal{B} = \{p_0(x) = 1, p_1(x) = x, p_2(x) = x^2\}$$

as a basis for $\mathbf{R}_2[x]$ and

$$\mathcal{A} = \left\{ e_{11} = \begin{bmatrix} 1 & 0 \\ 0 & 0 \end{bmatrix}, e_{12} = \begin{bmatrix} 0 & 1 \\ 0 & 0 \end{bmatrix}, e_{21} = \begin{bmatrix} 0 & 0 \\ 1 & 0 \end{bmatrix}, e_{22} = \begin{bmatrix} 0 & 0 \\ 0 & 1 \end{bmatrix} \right\}$$

as a basis for $M_{2 \times 2}$. Notice we have

$$T(1) = \begin{bmatrix} 0 & 0 \\ 0 & 1 \end{bmatrix} = e_{22}, \quad T(x) = \begin{bmatrix} -1 & 0 \\ 0 & 0 \end{bmatrix} = -e_{11}, \quad T(x^2) = \begin{bmatrix} -3 & 0 \\ 0 & 0 \end{bmatrix} = -3e_{11},$$

which implies

$$[T]_{\mathcal{B}}^{\mathcal{A}} = \begin{bmatrix} 0 & -1 & -3 \\ 0 & 0 & 0 \\ 0 & 0 & 0 \\ 1 & 0 & 0 \end{bmatrix}.$$

Example: We also considered the linear map

$$T : \mathbf{R}_2[x] \rightarrow \mathbf{R}_3[x], \quad T(p) = 2p' + \int_0^x 3p(t)dt.$$

We choose bases

$$\mathcal{B} = \{p_0(x) = 1, p_1(x) = x, p_2(x) = x^2\}$$

for $\mathbf{R}_2[x]$ and

$$\mathcal{A} = \{q_0(x) = 1, q_1(x) = x, q_2(x) = x^2, q_3(x) = x^3\}$$

for $\mathbf{R}_3[x]$. We computed above that

$$T(1) = 3x = 3q_1, \quad T(x) = 2 + \frac{3}{2}x^2 = 2q_0 + \frac{3}{2}q_2, \quad T(x^2) = 4x + x^3 = 4q_1 + q_3,$$

so that

$$[T]_{\mathcal{B}}^{\mathcal{A}} = \begin{bmatrix} 0 & 2 & 0 \\ 3 & 0 & 4 \\ 0 & \frac{3}{2} & 0 \\ 0 & 0 & 1 \end{bmatrix}.$$

There is a nice correspondence between the range of T and the span of the column vectors of $[T]_{\mathcal{B}}^{\mathcal{A}}$.

Proposition 21. *Let V and W be vector spaces of dimension n and m , respectively, and let $T : V \rightarrow W$ be linear. Choose bases $\mathcal{B} = \{v_1, \dots, v_n\}$ for V and $\mathcal{A} = \{w_1, \dots, w_m\}$ for W , and let $[T]_{\mathcal{B}}^{\mathcal{A}}$ be the matrix corresponding to T under these choices of bases. Then we can identify the range of T with the span of the columns of $[T]_{\mathcal{B}}^{\mathcal{A}}$.*

Proof. We know that $\mathcal{R}(T) = \text{span}\{T(v_1), \dots, T(v_n)\}$, and that the j th column of $[T]_{\mathcal{B}}^{\mathcal{A}}$ is $T(v_j)$, written in terms of $\{w_1, \dots, w_m\}$. The result follows. $\square$

4.4. Systems of linear equations and linear transformations. We have just seen how one can represent a linear transformation $T : V \rightarrow W$ as an $m \times n$ matrix, where $\dim(V) = n$ and $\dim(W) = m$. Before, you also saw $m \times n$ matrices as tools to solve m linear equations in n unknowns. In this section we describe the connection between these two phenomena.

For the rest of the section, we take $A \in M_{m \times n}$ to be a matrix with m rows and n columns, whose entries are real numbers. We will also regard A as a linear transformation $A : \mathbf{R}^n \rightarrow \mathbf{R}^m$. In the context of linear equations, we're used to solving the matrix equation $Ax = b$ for the unknown x , given the matrix A and the right hand side b .

It will be convenient to isolate the homogeneous case, where the right hand side b is zero, first.

Lemma 22. *Let $A \in M_{m \times n}$. Then x solves the equations $Ax = 0$ if and only if $x \in \ker(A)$.*

Proof. By definition, $\ker(A) = \{x \in \mathbf{R}^n : Ax = 0\}$, so the lemma follows. $\square$

Lemma 23. *Let $A \in M_{m \times n}$ with $m < n$. Then the set of solutions to the homogeneous equation, which is $\{x : Ax = 0\}$ has positive dimension.*

Proof. By the rank-nullity theorem,

$$\dim(\ker(T)) = n - \dim(\mathcal{R}(T)) \geq n - m > 0.$$

$\square$

Next we proceed to the general case, where the right hand side b of the equation $Ax = b$ can be any vector in $\mathbf{R}^m$.

Lemma 24. *Let $A \in M_{m \times n}$. One can solve the matrix equation $Ax = b$ for x if and only if $b \in \mathcal{R}(T)$.*

Proof. By definition, $b \in \mathcal{R}(T)$ if and only if there exists a vector $x \in \mathbf{R}^n$ such that $Ax = b$, that is, if and only if there is a solution to the equation $Ax = b$. $\square$

Lemma 25. *Let $A \in M_{m \times n}$ and let $b \in \mathcal{R}(T)$. The solution x to the matrix equation $Ax = b$ is unique if and only if $\ker(T) = \{0\}$.*

Proof. Suppose $\ker(T) = \{0\}$, and let x_1 and x_2 both solve $Ax = b$. Then

$$Ax_1 = b = Ax_2 \Leftrightarrow 0 = Ax_1 - Ax_2 = A(x_1 - x_2) \Leftrightarrow x_1 - x_2 = 0,$$

and so the solution to the equation $Ax = b$ is unique. Conversely, suppose there are two distinct solutions $x_1 \neq x_2$ to $Ax = b$. Then

$$Ax_1 = b = Ax_2 \Leftrightarrow 0 = A(x_1 - x_2) \Leftrightarrow x_1 - x_2 \in \ker(T),$$

and so $\ker(T) \neq \{0\}$. □

Theorem 26. *Let $A \in M_{m \times n}$.*

- (1) *$A : \mathbf{R}^n \rightarrow \mathbf{R}^m$ is onto if and only if the matrix equation $Ax = b$ admits at least one solution x for every choice of b on the right hand side.*
- (2) *$A : \mathbf{R}^n \rightarrow \mathbf{R}^m$ is one-to-one if and only if the matrix equation $Ax = b$ admits at most one solution x for every choice of b on the right hand side.*

Proof. The first statement follows from the first lemma above, and the second statement follows from the second lemma above. □

Corollary 27. *Let $A \in M_{m \times n}$.*

- (1) *If $m > n$ then there are $b \in \mathbf{R}^m$ such that one cannot find a solution x for $Ax = b$.*
- (2) *if $m < n$ then there are $b \in \mathbf{R}^m$ such that one can find many solutions x to the equation $Ax = b$.*

Proof. By the rank-nullity theorem, A cannot be onto if $m > n$, and so the first statement follows from the previous theorem. Similarly, A cannot be one-to-one if $m < n$, and so the second statement follows. □

The case of $n = m$ bears specific mention, and we can summarize it here. Recall (from MAM1000) that $A \in M_{n \times n}$ is invertible if there exists a matrix $B = A^{-1}$ such that $AB = BA = I$. (See also Definition 10.)

Theorem 28. *Let $A \in M_{n \times n}$. Then the following statements are equivalent.*

- *$A : \mathbf{R}^n \rightarrow \mathbf{R}^n$ is invertible.*
- *A is an invertible matrix.*
- *One can find a unique solution x to the equation $Ax = b$ for any b .*
- *One can find at least one solution x to the equation $Ax = b$ for any b .*
- *One can find at most one solution x to the equation $Ax = b$ for any b .*

Proof. This theorem follows from the previous theorem and the fact that A is invertible if and only if A is onto if and only if A is one-to-one. (This latter part is particular to $n \times n$ matrices.) □

4.5. Isomorphisms, invertibility, and the effect of changing basis. In this section we relate isomorphisms and the invertibility of a matrix, and discuss how the matrix representing a linear transformation changes if one changes the basis of the vector space. We begin with some general properties of linear transformations and use this to discuss inversion of matrices, leading to the effect of a change of basis.

Theorem 29. *Let V and W be vector spaces and let $T : V \rightarrow W$ be linear and invertible. Then $T^{-1} : W \rightarrow V$ is also linear.*

Proof. Let $w_1, w_2 \in W$ and let $a_1, a_2 \in \mathbf{R}$. Since T is both one-to-one and onto, there are unique vectors $v_1, v_2 \in V$ such that $T(v_1) = w_1$ and $T(v_2) = w_2$, which implies $v_1 = T^{-1}(w_1)$ and $T^{-1}(w_2) = v_2$. Thus

$$\begin{aligned} T^{-1}(a_1 w_1 + a_2 w_2) &= T^{-1}(a_1 T(v_1) + a_2 T(v_2)) = T^{-1}(T(a_1 v_1 + a_2 v_2)) \\ &= a_1 v_1 + a_2 v_2 = a_1 T^{-1}(w_1) + a_2 T^{-1}(w_2). \end{aligned}$$

□

Using our identification of a matrix as a linear transformation, we can now define the inverse of an $n \times n$ matrix.

Definition 10. *Let $A \in M_{n \times n}$. Then A is invertible if there exists $B \in M_{n \times n}$ such that*

$$AB = BA = I,$$

where I is the $n \times n$ identity matrix (with 1's on the leading diagonal, and 0's in every other entry).

Example: If

$$A = \begin{bmatrix} 5 & 7 \\ 2 & 3 \end{bmatrix}, \quad \begin{bmatrix} 3 & -7 \\ -2 & 5 \end{bmatrix}$$

it is easy to check that $AB = BA = I$, so $B = A^{-1}$.

Theorem 30. *Let V and W be finite dimensional vector spaces, and choose bases $\mathcal{B} = \{v_1, \dots, v_n\}$ for V and $\mathcal{A} = \{w_1, \dots, w_n\}$ for W . Let $T : V \rightarrow W$ be linear. Then T is an isomorphism if and only if $[T]_{\mathcal{B}}^{\mathcal{A}}$ is an invertible matrix. In this case, $[T^{-1}]_{\mathcal{A}}^{\mathcal{B}} = ([T]_{\mathcal{B}}^{\mathcal{A}})^{-1}$.*

Proof. We already know (by the rank-nullity theorem) that T can be an isomorphism if and only if $\dim(V) = \dim(W) = n$, so $[T]_{\mathcal{B}}^{\mathcal{A}}$ must be an $n \times n$ matrix. We have also seen that we can represent the composition of linear transformations as the product of matrices. Thus, if T is an isomorphism then we have $T \circ T^{-1} = I_W : W \rightarrow W$ and $T^{-1} \circ T = I_V : V \rightarrow V$. Thus

$$I = [I_V]_{\mathcal{B}}^{\mathcal{B}} = [T^{-1} \circ T]_{\mathcal{B}}^{\mathcal{B}} = [T^{-1}]_{\mathcal{A}}^{\mathcal{B}} [T]_{\mathcal{B}}^{\mathcal{A}}, \quad I = [I_W]_{\mathcal{A}}^{\mathcal{A}} = [T \circ T^{-1}]_{\mathcal{A}}^{\mathcal{A}} = [T]_{\mathcal{B}}^{\mathcal{A}} [T^{-1}]_{\mathcal{A}}^{\mathcal{B}},$$

so $([T]_{\mathcal{B}}^{\mathcal{A}})^{-1} = [T^{-1}]_{\mathcal{A}}^{\mathcal{B}}$. In particular, if T is an isomorphism then $[T]$ must be an invertible matrix.

Now we suppose $A = [T]_{\mathcal{B}}^{\mathcal{A}}$ be an invertible matrix and prove that T is an isomorphism. There must be $B \in M_{n \times n}$ such that $AB = BA = I$, and write the

components of A as a_{ij} and the components of B as b_{ij} . There is a unique linear transformation $U : W \rightarrow V$ such that

$$U(w_j) = \sum_{i=1}^n b_{ij} v_i$$

for each $j = 1, 2, \dots, n$. By construction, $[U]_{\mathcal{A}}^{\mathcal{B}} = B$. Moreover,

$$[UT]_{\mathcal{B}}^{\mathcal{B}} = [U]_{\mathcal{A}}^{\mathcal{B}}[T]_{\mathcal{B}}^{\mathcal{A}} = BA = I = [I_V]_{\mathcal{B}}^{\mathcal{B}}, \quad [TU]_{\mathcal{A}}^{\mathcal{A}} = [T]_{\mathcal{B}}^{\mathcal{A}}[U]_{\mathcal{A}}^{\mathcal{B}} = AB = I = [I_W]_{\mathcal{A}}^{\mathcal{A}},$$

and so U is the inverse transformation of T . In particular, T must be an isomorphism. $\square$

Example: Consider $T : \mathbf{R}_1[x] \rightarrow \mathbf{R}^2$ defined by $T(a + bx) = (a, a + b)$. We can check directly that T is an isomorphism, and the inverse transformation is given by $T^{-1}(c, d) = c + (d - c)x$. Indeed,

$$T^{-1}(T(a + bx)) = T^{-1}(a, a + b) = a + (a + b - a)x = a + bx$$

and

$$T(T^{-1}(c, d)) = T(c + (d - c)x) = (c, d - c + c) = (c, d).$$

Now choose the bases $\mathcal{B} = \{1, x\}$ for $\mathbf{R}_1[x]$ and $\mathcal{A} = \{e_1 = (1, 0), e_2 = (0, 1)\}$ for $\mathbf{R}^2$. With respect to these bases, it is easy to check that

$$[T]_{\mathcal{B}}^{\mathcal{A}} = \begin{bmatrix} 1 & 0 \\ 1 & 1 \end{bmatrix}, \quad [T^{-1}]_{\mathcal{A}}^{\mathcal{B}} = \begin{bmatrix} 1 & 0 \\ -1 & 1 \end{bmatrix}.$$

A straightforward computation gives

$$[T]_{\mathcal{B}}^{\mathcal{A}}[T^{-1}]_{\mathcal{A}}^{\mathcal{B}} = [T^{-1}]_{\mathcal{A}}^{\mathcal{B}}[T]_{\mathcal{B}}^{\mathcal{A}} = \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix}.$$

We have seen before that an n -dimensional vector space V is isomorphic to $\mathbf{R}^n$, where the particular isomorphism depends on the choice of basis $\mathcal{B} = \{v_1, \dots, v_n\}$ for V . We have written this isomorphism as

$$\Phi_{\mathcal{B}} : V \rightarrow \mathbf{R}^n, \quad v \mapsto \Phi_{\mathcal{B}}(v) = [v]_{\mathcal{B}}.$$

We can call $[v]_{\mathcal{B}}$ the **coordinate representation** of the vector v with respect to the basis $\mathcal{B}$. However, we do need to be a bit careful, because if we change our basis the coordinate representation will certainly change! Below we will track exactly how $[v]_{\mathcal{B}}$ changes if we change $\mathcal{B}$.

Example: We define $T : \mathbf{R}_3[x] \rightarrow \mathbf{R}_2[x]$ by $T(p)(x) = p'(x)$. Now choose bases $\mathcal{B} = \{1, x, x^2, x^3\}$ for $\mathbf{R}_3[x]$ and $\mathcal{A} = \{1, x, x^2\}$ for $\mathbf{R}_2[x]$. With respect to these bases, we have

$$[T]_{\mathcal{B}}^{\mathcal{A}} = \begin{bmatrix} 0 & 1 & 0 & 0 \\ 0 & 0 & 2 & 0 \\ 0 & 0 & 0 & 3 \end{bmatrix}.$$

If we let $p(x) = 2 + x - 3x^2 + 5x^3$ then

$$\Phi_{\mathcal{B}}(p) = [p]_{\mathcal{B}} = \begin{bmatrix} 2 \\ 1 \\ -3 \\ 5 \end{bmatrix},$$

and so

$$[T(p)]_{\mathcal{A}} = [T]_{\mathcal{B}}^{\mathcal{A}} \Phi_{\mathcal{B}}(p) = \begin{bmatrix} 0 & 1 & 0 & 0 \\ 0 & 0 & 2 & 0 \\ 0 & 0 & 0 & 3 \end{bmatrix} \begin{bmatrix} 2 \\ 1 \\ -3 \\ 5 \end{bmatrix} = \begin{bmatrix} 1 \\ -6 \\ 15 \end{bmatrix},$$

which we can verify by noticing

$$T(p) = 1 - 6x + 15x^2.$$

At this point, we turn our attention to the effect of changing basis. As a way of motivating this discussion, we begin by considering the plane curve in $\mathbf{R}^2$ given by $\{(x, y) \in \mathbf{R}^2 : 2x^2 - 4xy + 5y^2 = 1\}$. It might be difficult to recognize this plane curve as an ellipse in the standard coordinates, but if we make the substitution

$$\begin{aligned} x &= \frac{2}{\sqrt{5}}x' - \frac{1}{\sqrt{5}}y' \\ y &= \frac{1}{\sqrt{5}}x' + \frac{2}{\sqrt{5}}y' \end{aligned},$$

then the equation of the curve becomes $\{(x', y') \in \mathbf{R}^2 : (x')^2 + 6(y')^2 = 1\}$. This is rather obviously an ellipse. In fact, all we are doing in changing from the (x, y) coordinates to the (x', y') coordinates is rotating by the angle $\theta = -\arcsin(1/\sqrt{5})$.

It is a useful exercise to write out the change of variables matrix for the coordinate transformation we have just done. Implicitly, we have two bases for $\mathbf{R}^2$ in the computation we have just done. The first is the standard basis

$$\mathcal{B} = \{e_1 = (1, 0), e_2 = (0, 1)\},$$

and the second is

$$\mathcal{A} = \left\{ v_1 = \left(\frac{2}{\sqrt{5}}, \frac{1}{\sqrt{5}} \right) = \frac{2}{\sqrt{5}}e_1 + \frac{1}{\sqrt{5}}e_2, v_2 = \left(\frac{-1}{\sqrt{5}}, \frac{2}{\sqrt{5}} \right) = \frac{-1}{\sqrt{5}}e_1 + \frac{2}{\sqrt{5}}e_2 \right\}.$$

The matrix representing the change of basis between these two bases is the matrix representing the **identity** linear transformation, written with respect to two different bases. We have

$$Q = [I]_{\mathcal{A}}^{\mathcal{B}} = \frac{1}{\sqrt{5}} \begin{bmatrix} 2 & -1 \\ 1 & 2 \end{bmatrix},$$

and so

$$\begin{bmatrix} x' \\ y' \end{bmatrix} = [v]_{\mathcal{A}} = [I]_{\mathcal{A}}^{\mathcal{B}}[v]_{\mathcal{B}} = Q \begin{bmatrix} x \\ y \end{bmatrix} = \frac{1}{\sqrt{5}} \begin{bmatrix} 2 & -1 \\ 1 & 2 \end{bmatrix} \begin{bmatrix} x \\ y \end{bmatrix}.$$

We call $Q = [I]_{\mathcal{A}}^{\mathcal{B}}$ the **change of basis matrix**, because it changes the coordinates of a vector from the coordinates from the $\mathcal{A}$ coordinates to the $\mathcal{B}$ coordinates. There is in fact nothing particular about the change of coordinates we chose above,

and the entire process carries through with any change of basis. We summarize this with the following theorem.

Theorem 31. *Let $\mathcal{B}$ and $\mathcal{A}$ be two bases for a finite dimensional vector space V , and let $Q = [I_V]_{\mathcal{A}}^{\mathcal{B}}$, the matrix representing the identity map from V to itself, written with respect to the basis $\mathcal{A}$ in the domain and the basis $\mathcal{B}$ in the target. Then*

- (1) *Q is invertible.*
- (2) *For any $v \in V$ we have $[v]_{\mathcal{B}} = Q[v]_{\mathcal{A}}$.*

Proof. We have seen that the matrix representing a linear transformation is invertible if and only if the linear transformation is an isomorphism, and the identity map is certainly an isomorphism. Thus the first statement holds. We verify the second statement as follows:

$$[v]_{\mathcal{B}} = [I_V(v)]_{\mathcal{B}} = [I_V(v)]_{\mathcal{A}}^{\mathcal{B}}[v]_{\mathcal{A}} = Q[v]_{\mathcal{A}}.$$

□

Example: Take $V = \mathbf{R}^2$, and choose the two bases

$$\mathcal{B} = \{(1, 1); (1, -1)\}, \quad \mathcal{A} = \{(2, 4), (3, 1)\}.$$

Observe that

$$(2, 4) = 3(1, 1) - 1(1, -1), \quad (3, 1) = 2(1, 1) + 1(1, -1),$$

so in this case the change of basis matrix is

$$Q = \begin{bmatrix} 3 & 2 \\ -1 & 1 \end{bmatrix}.$$

If $T : V \rightarrow V$ is linear, we can write a matrix $[T]_{\mathcal{B}}^{\mathcal{B}}$ representing T with respect to the basis $\mathcal{B}$ or a matrix $[T]_{\mathcal{A}}^{\mathcal{A}}$ with respect to $\mathcal{A}$. How are these two matrices related?

Theorem 32. *Let $T : V \rightarrow V$ be linear, let $\mathcal{B}$ and $\mathcal{A}$ be two bases for V , let $[T]_{\mathcal{B}}^{\mathcal{B}}$ be the matrix of T with respect to $\mathcal{B}$, and let $[T]_{\mathcal{A}}^{\mathcal{A}}$ be the matrix of T with respect to $\mathcal{A}$. Then*

$$[T]_{\mathcal{A}}^{\mathcal{A}} = Q^{-1}[T]_{\mathcal{B}}^{\mathcal{B}}Q.$$

Proof. Let I be the identity transformation on V , so that $T = I \circ T = T \circ I$. Then

$$\begin{aligned} Q[T]_{\mathcal{A}}^{\mathcal{A}} &= [I]_{\mathcal{A}}^{\mathcal{B}}[T]_{\mathcal{A}}^{\mathcal{A}} = [IT]_{\mathcal{A}}^{\mathcal{B}} \\ &= [TI]_{\mathcal{A}}^{\mathcal{B}} = [T]_{\mathcal{B}}^{\mathcal{B}}[I]_{\mathcal{A}}^{\mathcal{B}} = [T]_{\mathcal{B}}^{\mathcal{B}}Q. \end{aligned}$$

The result now follows if we multiply this equation on the left by Q^{-1} . □

Example: Let $V = \mathbf{R}^3$ and define

$$T : \mathbf{R}^3 \rightarrow \mathbf{R}^3, \quad T(a_1, a_2, a_3) = (2a_1 + a_2, a_1 + a_2 + a_3, -a_2).$$

We choose $\mathcal{B} = \{e_1, e_2, e_3\}$ to be the standard basis of $\mathbf{R}^3$ and choose a second basis

$$\mathcal{A} = \{v_1 = (-1, 0, 0), v_2 = (2, 1, 0), v_3 = (1, 1, 1)\}.$$

We have

$$[T]_B^B = \begin{bmatrix} 2 & 1 & 0 \\ 1 & 1 & 3 \\ 0 & -1 & 0 \end{bmatrix}, \quad Q = \begin{bmatrix} -1 & 2 & 1 \\ 0 & 1 & 1 \\ 0 & 0 & 1 \end{bmatrix}, \quad Q^{-1} = \begin{bmatrix} -1 & 2 & -1 \\ 0 & 1 & -1 \\ 0 & 0 & 1 \end{bmatrix},$$

and if we multiply these matrices out we see

$$[T]_A^A = Q^{-1}[T]_B^B Q = \begin{bmatrix} 0 & 2 & 8 \\ -1 & 4 & 6 \\ 0 & -1 & -1 \end{bmatrix}.$$

We can verify that the columns of $[T]_A^A$ given the image of $\{v_1, v_2, v_3\}$, written as linear combinations of themselves. For instance,

$$T(v_2) = T(2, 1, 0) = (5, 3, -1) = 2v_1 + 4v_2 + (-1)v_3.$$

The coefficients 2, 4, and -1 are indeed the second column of $[T]_A^A$.

Observe, that we now have many different ways of writing down a matrix representing the same linear transformation, but they are all related by the change of basis matrix. Thus it makes sense to have the following definition.

Definition 11. Let $A, B \in M_{n \times n}$. We say that B is similar to A if there is an invertible matrix $Q \in M_{n \times n}$ such that $B = Q^{-1}AQ$.

This definition says that two matrices are similar precisely when they represent the same linear transformation written with respect to two different bases.

5. DETERMINANTS

Area and mappings from the plane to itself: Recall that in Section 2 we found a linear mapping to take the unit square $S = \{0 \leq x \leq 1, 0 \leq y \leq 1\}$ to any parallelogram P with one corner at the origin. We can write the parallelogram P as

$$P = \{xv + yw : 0 \leq x \leq 1, 0 \leq y \leq 1\},$$

where v and w are the two vectors which form the edges of P starting at the origin $(0, 0)$. Then we can write the linear transformation as

$$T\left(\begin{bmatrix} x \\ y \end{bmatrix}\right) = xv + yw, \quad [T] = \begin{bmatrix} v_1 & w_1 \\ v_2 & w_2 \end{bmatrix},$$

where $v = (v_1, v_2)$ and $w = (w_1, w_2)$ in components. Notice that the mapping T is invertible precisely when it does not collapse S down to a line segment (or a point), which happens precisely when the area of the parallelogram P is non-zero.

You might recall that in MAM1000 you defined an object called the determinant, written $\det([T]) = v_1w_2 - w_1v_2$, and were told

$$\det([T]) \neq 0 \Leftrightarrow T \text{ invertible} \Leftrightarrow \text{Area}(P) \neq 0.$$

We'll see next that $\det([T])$ is the area of P , up to a sign.

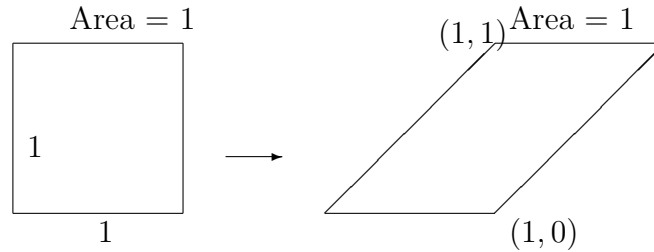
This is easiest to see with the shear map we examined in the last set of notes. Start with the shear map T whose matrix representation is

$$[T] = \begin{bmatrix} 1 & b \\ 0 & 1 \end{bmatrix}.$$

(In the earlier set of notes we wrote the entry in the upper right corner of $[T]$ as a , but it will turn out to be convenient to call it b for our later discussion.) In this case, T maps the unit square $S = \{0 \leq x \leq 1, 0 \leq y \leq 1\}$ to the parallelogram P spanned by the two vectors $v = \begin{bmatrix} 1 \\ 0 \end{bmatrix}$ and $w = \begin{bmatrix} b \\ 1 \end{bmatrix}$; in other words,

$$T(S) = P = \{xv + yw : 0 \leq x \leq 1, 0 \leq y \leq 1\} = \left\{ x \begin{bmatrix} 1 \\ 0 \end{bmatrix} + y \begin{bmatrix} b \\ 1 \end{bmatrix} : 0 \leq x \leq 1, 0 \leq y \leq 1 \right\}.$$

We reproduce a picture here:



We already know that the unit square S has area 1, but let's see that P also has area 1. The area of a parallelogram is equal to its base times its height, and the height and base of P are both 1, so the area of P is $1 \cdot 1 = 1$. On the other hand,

$$\det([T]) = \det \left(\begin{bmatrix} 1 & b \\ 0 & 1 \end{bmatrix} \right) = 1 \cdot 1 - 0 \cdot b = 1 = \text{Area}(P).$$

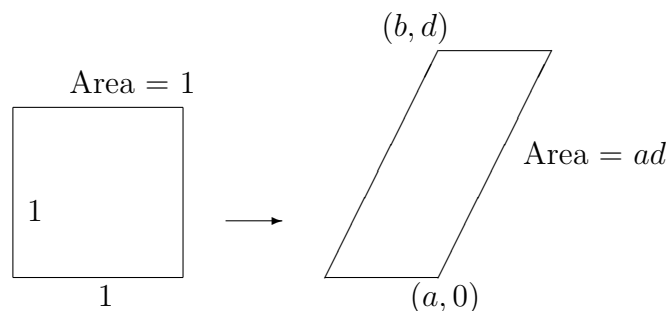
Now we can rescale the shear T by a in the horizontal direction and by d and vertical direction, to have something more general. This time we have

$$[T] = \begin{bmatrix} a & b \\ 0 & d \end{bmatrix},$$

and

$$T(S) = P = \{xv + yw : 0 \leq x \leq 1, 0 \leq y \leq 1\} = \left\{ x \begin{bmatrix} a \\ 0 \end{bmatrix} + y \begin{bmatrix} b \\ d \end{bmatrix} : 0 \leq x \leq 1, 0 \leq y \leq 1 \right\},$$

and the picture looks like



(In this particular picture $a = 1/2$ and $d = 2$, but this choice of scaling factors is not important.) This time the height of the parallelogram P is d while its base is a , so $\text{Area}(P) = \text{base} \cdot \text{height} = ad$. Again, we have

$$|\det([T])| = \left| \det \left(\begin{bmatrix} a & b \\ 0 & d \end{bmatrix} \right) \right| = |a \cdot d - b \cdot 0| = |ad| = \text{Area}(P).$$

Notice that the absolute value here is necessary, because a and d could have opposite signs.

Now that we know $|\det([T])|$ gives the area of the image of the unit square if $[T] = \begin{bmatrix} a & b \\ 0 & d \end{bmatrix}$, it's not too hard to see this is true for any linear map. We'll first need a technical fact.

Example: Choose any angle θ . Then

$$\begin{aligned} \det([R_\theta]) &= \det \left(\begin{bmatrix} \cos \theta & -\sin \theta \\ \sin \theta & \cos \theta \end{bmatrix} \right) \\ &= \cos^2 \theta - (-\sin^2 \theta) = \cos^2 \theta + \sin^2 \theta = 1. \end{aligned}$$

Geometrically, this computation says that a rotation leaves area unchanged.

Example: We prove $\det(AB) = \det(A)\det(B)$ for 2×2 matrices A and B . We have

$$\begin{aligned} \det(AB) &= \det \left(\begin{bmatrix} a & b \\ c & d \end{bmatrix} \begin{bmatrix} e & f \\ g & h \end{bmatrix} \right) = \det \begin{bmatrix} ae + bg & af + bh \\ ce + dg & cf + gh \end{bmatrix} \\ &= (ae + bg)(cf + dh) - (ce + dg)(af + bh) = adeh + bcfg - bceh - adfg \end{aligned}$$

and

$$\begin{aligned} \det(A)\det(B) &= \det \begin{bmatrix} a & b \\ c & d \end{bmatrix} \det \begin{bmatrix} e & f \\ g & h \end{bmatrix} = (ad - bc)(eh - fg) \\ &= adeh - bceh - adfg + bcfg. \end{aligned}$$

Notice that this means $\det(AB) = \det(A)\det(B) = \det(B)\det(A) = \det(BA)$ for any pair of 2×2 matrices. In fact, this is true for $A, B \in M_{n \times n}$, though the proof is a little bit messier.

Now let $T : \mathbf{R}^2 \rightarrow \mathbf{R}^2$ be a linear mapping of the plane to itself, and suppose

$$[T] = \begin{bmatrix} a & b \\ c & d \end{bmatrix}.$$

This means $T(e_1) = \begin{bmatrix} a \\ c \end{bmatrix}$ and $T(e_2) = \begin{bmatrix} b \\ d \end{bmatrix}$, where $e_1 = \begin{bmatrix} 1 \\ 0 \end{bmatrix}$ and $e_2 = \begin{bmatrix} 0 \\ 1 \end{bmatrix}$ as before. Now, the vector $T(e_1) = \begin{bmatrix} a \\ c \end{bmatrix}$ makes some angle θ with the positive x axis, so we apply the rotation $R_{-\theta}$ to T to get a new mapping

$$\tilde{T} = R_{-\theta} \circ T, \quad [\tilde{T}] = [R_{-\theta}][T] = \begin{bmatrix} \cos \theta & \sin \theta \\ -\sin \theta & \cos \theta \end{bmatrix} \begin{bmatrix} a & b \\ c & d \end{bmatrix} = \begin{bmatrix} \tilde{a} & \tilde{b} \\ 0 & \tilde{d} \end{bmatrix},$$

and

$$\det([\tilde{T}]) = \det([R_{-\theta}][T]) = \det([R_{-\theta}]) \det([T]) = \det([T]).$$

By the computation we did above, $\text{Area}(\tilde{P}) = |\det([\tilde{T}])|$. We also have that T sends the unit square S to a parallelogram P , and $\tilde{T}$ sends S to a parallelogram $\tilde{P}$. These two parallelograms P and $\tilde{P}$ differ by a rotation, so they have the same area. Thus we see

$$\text{Area}(P) = \text{Area}(\tilde{P}) = |\det([\tilde{T}])| = |\det([T])|.$$

In particular, we have just proven that $\det([T]) \neq 0$ precisely when T is invertible, because this is precisely when the image parallelogram P has nonzero area.

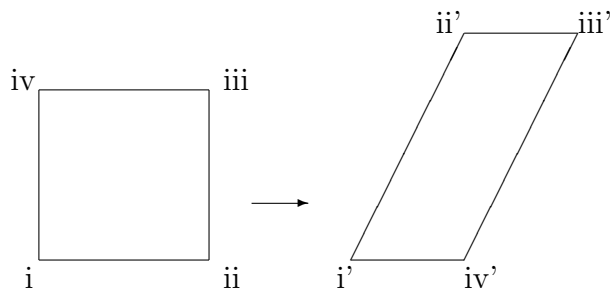
Orientation and the sign of the determinant: As we saw in the previous notes, there are actually two linear transformations which map the unit square S onto this parallelogram P , we can also have

$$T\left(\begin{bmatrix} 1 \\ 0 \end{bmatrix}\right) = \begin{bmatrix} b \\ d \end{bmatrix}, \quad T\left(\begin{bmatrix} 0 \\ 1 \end{bmatrix}\right) = \begin{bmatrix} a \\ 0 \end{bmatrix}, \quad [T] = \begin{bmatrix} b & a \\ d & 0 \end{bmatrix}.$$

In this case we see that

$$\det([T]) = -ad = -\text{Area}(P).$$

Why do we have the minus sign? To understand what's going on, it will help to label the corners of the unit square S and the parallelogram P as in the picture below.



What does this labeling mean? The mapping T sends the vector e_1 , which goes from i to ii in the square on the left to the vector $\begin{bmatrix} b \\ d \end{bmatrix}$, which also goes from i' to ii' in the parallelogram on the right. Similarly, the mapping T sends the vector e_2 , which goes from i to iv in the square on the left to the vector $\begin{bmatrix} a \\ 0 \end{bmatrix}$, which also goes

from i' to iv' in the parallelogram on the right. Now, if we follow the labeling of the corners of the square in order, as in i to ii to iii to iv, then we traverse along the boundary of the square counter-clockwise. However, if we follow the labeling of the corners of the parallelogram in order, as in i' to ii' to iii' to iv', we traverse along the boundary of the parallelogram clockwise. This means the mapping T reversed the direction we traversed along the boundary of the shape. In other words, T reversed the orientation. We have discovered the following general principle:

$$\det([T]) < 0 \Leftrightarrow T \text{ reverses orientation.}$$

This principle is exactly why we wrote $|\det([T])| = \text{Area}(P)$ before. In general, if $T : \mathbf{R}^2 \rightarrow \mathbf{R}^2$ preserves orientation then $\det([T]) = \text{Area}(P)$, but if T reverses orientation then $\det([T]) = -\text{Area}(P)$.

Higher dimensions: So far we've seen that the determinant of a 2×2 matrix is the area (up to a sign) of the parallelogram which is the image of the unit square. In fact, a similar thing is true in higher dimensions. Let $[T]$ be an $n \times n$ matrix, which we've seen corresponds to a linear map $T : \mathbf{R}^n \rightarrow \mathbf{R}^n$. Then T sends the unit cube $S = \{0 \leq x_i \leq 1 : i = 1, 2, \dots, n\}$ to a parallelepiped P , which is spanned by the columns of $[T]$. Then $|\det([T])| = \text{Vol}(P)$, where Vol gives the n -dimensional volume.

We begin with a quick illustrative example. Consider

$$[T] = \begin{bmatrix} a & b & 0 \\ c & d & 0 \\ 0 & 0 & e \end{bmatrix}, \quad e > 0.$$

Then the image of the unit cube S under T is

$$P = \{(x, y, z) : (x, y) \in \bar{P}, 0 \leq z \leq e\},$$

where

$$[\bar{T}] = \begin{bmatrix} a & b \\ c & d \end{bmatrix}, \quad \bar{S} = \{0 \leq x \leq 1, 0 \leq y \leq 1\}, \quad \bar{P} = T(\bar{S}).$$

By slicing P with horizontal slices, we see

$$\text{Vol}(P) = e \text{Area}(\bar{P}) = e |\det([\bar{T}])|.$$

So, by any reasonable definition of the determinant for 3×3 matrices which fits with our definition for 2×2 matrices, we must have

$$\det([T]) = e \det([\bar{T}]) = e(ad - bc).$$

Exercise: Let $[T]$ be a 3×3 matrix. Show that you can always perform a rotation to make the last row of $[T]$ into $[0 \ 0 \ e]$. (Hint: geometrically, you want to rotate the parallelepiped so that one of its faces lies in a coordinate plane. What are the columns of $[T]$?)

At this point, we can write down a reasonable formula for the determinant of a 3×3 matrix. Let

$$[T] = \begin{bmatrix} a & b & c \\ d & e & f \\ g & h & i \end{bmatrix},$$

then

$$\det[T] = g \det \left(\begin{bmatrix} b & c \\ e & f \end{bmatrix} \right) - h \det \left(\begin{bmatrix} a & c \\ d & f \end{bmatrix} \right) + i \det \left(\begin{bmatrix} a & b \\ d & e \end{bmatrix} \right).$$

Here we've singled out the last row, but we can do the same thing by picking out any row or column. To do this properly, we need some notation. Let $[T] = [A_{ij}]$, so that the entry of $[T]$ in the i th row, j th column is A_{ij} . Also, let $[\bar{T}_{ij}]$ be the 2×2 matrix you get from $[T]$ by crossing out the i th row and j th column. Then for any choice of $j = 1, 2, 3$ we can write

$$\det([T]) = (-1)^{1+j} A_{1j} \det([\bar{T}_{1j}]) + (-1)^{2+j} A_{2j} \det([\bar{T}_{2j}]) + (-1)^{3+j} A_{3j} \det([\bar{T}_{3j}]),$$

which computes $\det([T])$ by expanding along the j th column. Alternatively, for any choice of $i = 1, 2, 3$ we can write

$$\det([T]) = (-1)^{i+1} A_{i1} \det([\bar{T}_{i1}]) + (-1)^{i+2} A_{i2} \det([\bar{T}_{i2}]) + (-1)^{i+3} A_{i3} \det([\bar{T}_{i3}]),$$

which computes $\det([T])$ by expanding along the i th row.

The same idea will compute the determinant of any square matrix inductively. That is, you write the determinant of an $n \times n$ matrix as a sum of determinants of $(n-1) \times (n-1)$ matrices. We write the general formula as follows. Again, we let A_{ij} be the entry of $[T]$ in the i th row, j th column, and we let $[\bar{T}_{ij}]$ be the $(n-1) \times (n-1)$ matrix you get from $[T]$ by crossing out the i th row and the j th column. Then for any choice of $j = 1, 2, \dots, n$ we compute $\det([T])$ by expanding along the j th column using the formula

$$\det([T]) = (-1)^{1+j} A_{1j} \det([\bar{T}_{1j}]) + (-1)^{2+j} A_{2j} \det([\bar{T}_{2j}]) + \dots + (-1)^{n+j} A_{nj} \det([\bar{T}_{nj}]).$$

Alternatively, for any choice of $i = 1, 2, \dots, n$ we compute $\det([T])$ by expanding along the i th row using the formula

$$\det([T]) = (-1)^{i+1} A_{i1} \det([\bar{T}_{i1}]) + (-1)^{i+2} A_{i2} \det([\bar{T}_{i2}]) + \dots + (-1)^{i+n} A_{in} \det([\bar{T}_{in}]).$$

We summarize some important properties of the determinant here.

- (1) The determinant is linear in each row and column. That is, if A is an $n \times n$ matrix and $\tilde{A}$ is the same as A except that you multiply the i th row by c , then $\det(\tilde{A}) = c \det(A)$. Also, A_1 and A_2 are the same except at the i th row and A is what you get by adding together the i th row of A_1 and A_2 then $\det(A) = \det(A_1) + \det(A_2)$. The same goes for columns.
- (2) Consequently, if A is an $n \times n$ matrix and c is a number then $\det(cA) = c^n \det(A)$.
- (3) An $n \times n$ matrix A is invertible if and only if $\det(A) \neq 0$.
- (4) In fact, $|\det(A)|$ is the n -dimensional volume of the parallelepiped P which is the image of the unit cube $S = \{0 \leq x_1 \leq 1, \dots, 0 \leq x_n \leq 1\}$ under the linear transformation associated to A . (You can prove this by induction, in a very similar way we got the geometric interpretation for three dimensions from the two-dimensional version.)
- (5) Let A and B be $n \times n$ matrices, then $\det(AB) = \det(A) \det(B)$.
- (6) Let A be an $n \times n$ matrix and let $\tilde{A}$ be the matrix you get by swapping two rows of A (or by swapping two columns). Then $\det(\tilde{A}) = -\det(A)$.

6. EIGENVALUES AND EIGENVECTORS

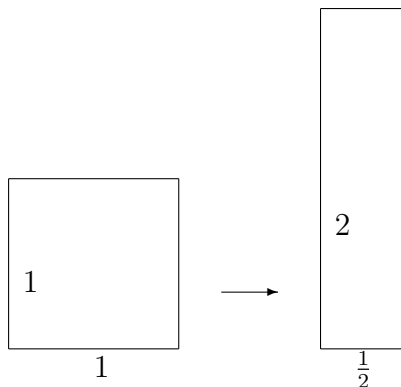
In this section we discuss eigenvalues and eigenvectors of a linear transformation. The bulk of this section is concerned with linear transformations $T : \mathbf{R}^n \rightarrow \mathbf{R}^n$, and so we can work entirely in terms of $n \times n$ matrices. However, the situation is really no more complicated for a general linear transformation $T : V \rightarrow V$, where V is an arbitrary finite-dimensional vector space. So we will close with some more general examples.

6.1. Eigenvalues of linear transformations from $\mathbf{R}^n$ to itself.

6.1.1. *Some motivation.* We saw in the last section that the determinant of a 2×2 matrix tells us the effect the associated linear map has on area. In other words, if $\det([T]) = 2$ then $T : \mathbf{R}^2 \rightarrow \mathbf{R}^2$ will scale the areas of squares by a factor of 2. It's not too hard to show that T scales the areas of all shapes by the same factor. (Hint: cut whatever shape you're interested in into a bunch of little tiny squares. You won't be able to do this exactly, but what you have left over has a negligible area.) However, it's easy to find a linear map which preserves area but distorts lengths by a lot. For instance, consider the linear

$$T : \mathbf{R}^2 \rightarrow \mathbf{R}^2, \quad [T] = \begin{bmatrix} \frac{1}{2} & 0 \\ 0 & 2 \end{bmatrix}.$$

We draw a picture of what T does to the unit square below.



This map preserves area, but it changes lengths by a lot. It shrinks length in some directions by a factor of $1/2$ and it stretches lengths in other directions by a factor of 2. We can make this picture much worse by choosing, for instance, a horizontal scale factor of $1/100$ and a vertical scale factor of 100. This example tells us we need at least two numbers to keep track of how a linear map $T : \mathbf{R}^2 \rightarrow \mathbf{R}^2$ deforms lengths. We'll see in a bit that, at least in some special cases, we only need two numbers, and that these numbers are (essentially) the eigenvalues.

6.1.2. *Definitions.* If you know a little about the German language, you might be able to guess what an eigenvector is. The German word *eigen* means *own*, and an

eigenvector of a linear transformation keeps its own direction. It can get rescaled, but the direction remains the same.

Definition 12. Let $T : \mathbf{R}^n \rightarrow \mathbf{R}^n$ be a linear transformation. Then a nonzero vector $v \in \mathbf{R}^n$ is an eigenvector with eigenvalue λ if $T(v) = \lambda v$. Notice that, even though v is not allowed to be zero, it's possible that $\lambda = 0$.

Exercise: Why is it necessary to have $v \neq 0$ in the definition of an eigenvector v ?

This definition is a little awkward for doing computations, so the first thing we'll do is reformulate it a little. Let v be an eigenvector of T with eigenvalue λ . Then

$$[T][v] = \lambda[v] = \lambda[I][v] \Leftrightarrow ([T] - \lambda[I])[v] = 0.$$

Now, $v \neq 0$, so the linear transformation $T - \lambda I$ sends a nonzero vector to 0, which means it can't be one-to-one. This means $T - \lambda I$ isn't invertible, and so

$$\det([T] - \lambda[I]) = 0.$$

This last equation is an n -th degree polynomial equation for the unknown λ . We know that any n -th degree polynomial has exactly n roots in the complex numbers $\mathbf{C}$ (so long as we remember to count repeated roots), which means we've just proved the following

Theorem 33. Let $T : \mathbf{R}^n \rightarrow \mathbf{R}^n$ be linear. Then a complex number $\lambda \in \mathbf{C}$ is an eigenvalue of T if and only if

$$\det([T] - \lambda[I]) = 0.$$

Moreover, every $n \times n$ matrix has precisely n complex numbers $\lambda_1, \dots, \lambda_n$ (counted with multiplicity) which are eigenvalues.

The polynomial $\det([T] - \lambda[I])$ is called the **characteristic polynomial** of T ; it is a polynomial of degree n if $T \in M_{n \times n}$, and carries much important information about T . The Cayley-Hamilton theorem states that T is a root of its own characteristic polynomial.

This theorem tells us how to compute eigenvalues of a square matrix: we write down the polynomial $\det([T] - \lambda[I])$ and find its roots. In practice this can be a little sticky, for instance, if we want to find the eigenvalues of a 5×5 matrix. However, for the case of 2×2 matrices, which is most of what we'll discuss in this class, the eigenvalues are the roots of a second order polynomial, which we can always find using the quadratic formula. So, for the time being at least, let's say we can find eigenvalues and continue, to see how to find the eigenvectors.

Let $[T]$ be an $n \times n$ matrix, and let λ be an eigenvalue of $[T]$. We want to find the associated eigenvector(s), that is the nonzero vectors v such that $T(v) = \lambda v$. We write this equation as a matrix equation

$$[T][v] = \lambda[v]$$

and try to solve it using our favorite method (like row reduction).

Exercise: Show that if v is an eigenvector of the matrix A with eigenvalue λ , then $2v$ is also an eigenvector of A , with the same eigenvalue λ . Is there anything special about the scale factor of 2?

Exercise: Show that the linear system $[T][v] = \lambda[v]$ for finding an eigenvector will always have many many solutions. Usually, this system will have one free variable, so it might be convenient to set one of the components of v to 1. However, it is possible that this linear system has more than one free variables.

6.1.3. *Some properties of eigenvalues and eigenvectors.* Here we list some properties of the eigenvalues and eigenvectors.

Recall that v is an eigenvector of A with eigenvalue λ if $Av = \lambda v$. If λ is a real number as well, this means $A(v)$ is colinear with v , *i.e.* either $A(v)$ points in the same direction or the opposite direction as v . In other words, if λ is a real eigenvalue of A then, considered as a linear map, A preserves the direction of the associated eigenvector v .

Exercise: Recall that we constructed the 2×2 rotation matrices

$$[R_\theta] = \begin{bmatrix} \cos \theta & -\sin \theta \\ \sin \theta & \cos \theta \end{bmatrix}.$$

Show that $[R_\theta]$ has a real eigenvalue if and only if the angle θ is an integer multiple of π (when measured in radians).

Exercise: We also constructed reflection matrices. Show that 1 is an eigenvalue of any reflection matrix.

Exercise: Let A be the 3×3 matrix associated to a rotation of 3-dimensional space. Show that 1 is an eigenvalue of A , and describe the relation between this associated eigenvector and the rotation.

Exercise: Show that 0 is an eigenvalue of an $n \times n$ matrix A if and only if A is not invertible. (This is completely general.)

Now we consider an $n \times n$ matrix A with real entries A_{ij} in the i th row, j th column. We have that λ is an eigenvalue of A precisely when

$$\det(A - \lambda I) = 0.$$

This is an n th degree polynomial, and the coefficients of this polynomial are sums of products of the entries of A . This means λ is a root of a polynomial with real coefficients. Now, it can happen that λ is not a real number, but it is a complex number, but these complex roots occur in conjugate pairs. We have the following

Proposition 34. *Let A be an $n \times n$ matrix with real entries. Then a non-real complex number $\lambda = a + ib$ is an eigenvalue of A if and only if its complex conjugate $\bar{\lambda} = a - ib$ is also an eigenvalue. In fact, in this case the eigenvectors are also complex conjugates. That is, if v is an eigenvector associated to the eigenvalue λ then $\bar{v}$ is an eigenvector associated to the eigenvalue $\bar{\lambda}$.*

The last sentence of the proposition follows immediately from taking the complex conjugate of the equation $Av = \lambda v$ to get $A\bar{v} = \bar{\lambda}\bar{v}$.

Exercise: Let A be an $n \times n$ matrix with real entries, and let $\lambda = a + ib$ be a non-real eigenvalue. Show that the components of the associated eigenvector v are also non-real.

Some times we can find n independent eigenvectors $v_1, \dots, v_n$ for an $n \times n$ matrix A . This means we can find n linearly independent vectors $v_1, \dots, v_n$ such that $A(v_j) = \lambda_j v_j$, and that we can't write v_j as the weighted sum of the other v_i 's. In this case, we say that A is **diagonalizable**, for the following reason. We can write any vector w as a sum $w = c_1 v_1 + c_2 v_2 + \dots + c_n v_n$, and then

$$(1) \quad A(w) = A(c_1 v_1 + \dots + c_n v_n) = c_1 A(v_1) + \dots + c_n A(v_n) = c_1 \lambda_1 v_1 + \dots + c_n \lambda_n v_n.$$

In other words, if we let $\mathcal{B} = \{v_1, \dots, v_n\}$ be a basis of $\mathbf{R}^n$ consisting of eigenvectors of A then we have

$$[A]_{\mathcal{B}}^{\mathcal{B}} = \begin{bmatrix} \lambda_1 & 0 & \dots & 0 \\ 0 & \lambda_2 & \dots & 0 \\ \vdots & & & \vdots \\ 0 & \dots & 0 & \lambda_n \end{bmatrix};$$

in other words, $[A]$ is a diagonal matrix in the right coordinates. We see immediately from equation (1) that, at least if A is diagonalizable, that the eigenvalues $\lambda_1, \dots, \lambda_n$ encode the stretch factors we were looking for at the beginning of this section.

We need to know the facts that the determinant and the trace of a matrix do not depend on the basis; that is, if you change coordinates as we just did the determinant and the trace remain the same.

Exercise: Show that, for a diagonalizable, $n \times n$ matrix, the determinant is the product of the eigenvalues and the trace is the sum of the eigenvalues.

Exercise: Not all matrices are diagonalizable. In fact, show that $\begin{bmatrix} 1 & 1 \\ 0 & 1 \end{bmatrix}$ is not diagonalizable.

We saw above that if A is diagonalizable then $\det(A)$ is the product of all the eigenvalues and $\text{tr}(A)$ is their sum. In fact, this is true for any $n \times n$ matrix, as you'll see in a second year linear algebra course when you discuss the Jordan canonical form of a matrix.

Proposition 35. *For any $n \times n$ matrix A , it holds that $\det(A)$ is the product of the eigenvalues of A , and $\text{tr}(A)$ is their sum.*

Exercise: Let A be a 2×2 matrix with complex eigenvalues $\lambda_{\pm} = a \pm ib$. Show that $\text{tr}(A) = 2a$ and $\det(A) = a^2 + b^2$. In particular, $\det(A) \geq 0$.

Exercise: Let A be a 2×2 matrix with real eigenvalues λ_1 and λ_2 . Show that $\det(A) > 0$ if and only if λ_1 and λ_2 have the same sign. Then show that λ_1 and λ_2 are both positive if and only if both $\det(A) > 0$ and $\text{tr}(A) > 0$.

Finally, we mention **symmetric** matrices, that is matrices such that $A_{ij} = A_{ji}$ where A_{ij} is the entry of A in the i th row, j th column. These are particularly nice, as we see from the following theorem (which we will not prove here).

Theorem 36. *A symmetric $n \times n$ matrix is diagonalizable and has n real eigenvalues $\lambda_1 \leq \lambda_2 \leq \dots \leq \lambda_n \in \mathbf{R}$.*

6.1.4. *Examples.* We'll compute the eigenvalues and eigenvectors of some 2×2 matrices here, just so we have some examples written down.

First let

$$A = \begin{bmatrix} 4 & -2 \\ 3 & -3 \end{bmatrix}.$$

We want to find the eigenvalues of A , so we set

$$\begin{aligned} 0 &= \det(A - \lambda I) = \det \left(\begin{bmatrix} 4 - \lambda & -2 \\ 3 & -3 - \lambda \end{bmatrix} \right) \\ &= \lambda^2 - \lambda - 6 = (\lambda - 3)(\lambda + 2), \end{aligned}$$

and we see that the eigenvalues of A are $\lambda_1 = -2$ and $\lambda_2 = 3$.

Now we find the eigenvector associated to the eigenvalue $\lambda_1 = -2$. We want to solve the linear equation

$$Av = -2v \Leftrightarrow \begin{bmatrix} 4 & -2 \\ 3 & -3 \end{bmatrix} \begin{bmatrix} v_1 \\ v_2 \end{bmatrix} = \begin{bmatrix} -2v_1 \\ -2v_2 \end{bmatrix}.$$

Of course, you can solve this using row reduction, but I find that for a small system like this, it's easier to just write out the equations. We have

$$4v_1 - 2v_2 = -2v_1, \quad 3v_1 - 3v_2 = -2v_2,$$

and both these equations reduce to $v_2 = 3v_1$. (You might want to think about why you'll always reduce from two equations to one when you're finding the eigenvectors of a 2×2 matrix.) So, up to a scale factor, the eigenvector of A associated to $\lambda_1 = -2$ is

$$v = \begin{bmatrix} 1 \\ 3 \end{bmatrix}.$$

Finally we find the eigenvector associated to $\lambda_2 = 3$. This time the linear equation is

$$Aw = 3w \Leftrightarrow \begin{bmatrix} 4 & -2 \\ 3 & -3 \end{bmatrix} \begin{bmatrix} w_1 \\ w_2 \end{bmatrix} = \begin{bmatrix} 3w_1 \\ 3w_2 \end{bmatrix},$$

which we rewrite as

$$4w_1 - 2w_2 = 3w_1, \quad 3w_1 - 3w_2 = 3w_2.$$

This reduces to $w_1 = 2w_2$, and so the eigenvector is (again, up to scale)

$$w = \begin{bmatrix} 2 \\ 1 \end{bmatrix}.$$

For our next example, we take the matrix

$$A = \begin{bmatrix} 1 & -1 \\ 1 & 1 \end{bmatrix}.$$

Again, we find eigenvalues of A by setting

$$\begin{aligned} 0 &= \det(A - \lambda I) = \det \left(\begin{bmatrix} 1 - \lambda & -1 \\ 1 & 1 - \lambda \end{bmatrix} \right) \\ &= \lambda^2 - 2\lambda + 2. \end{aligned}$$

Using the quadratic formula we see that the eigenvalues are $\lambda_+ = 1 + i$ and $\lambda_- = 1 - i$. Notice that, just as we said earlier, the eigenvalues occur in conjugate pairs.

We set up the equation for the eigenvector associated to $\lambda_+ = 1 + i$ as before, and get

$$Av = (1 + i)v \Leftrightarrow \begin{bmatrix} 1 & -1 \\ 1 & 1 \end{bmatrix} \begin{bmatrix} v_1 \\ v_2 \end{bmatrix} = \begin{bmatrix} (1 + i)v_1 \\ (1 + i)v_2 \end{bmatrix},$$

which reduces to

$$v_1 = iv_2 \Leftrightarrow -v_2 = iv_1.$$

(You might want to recall here that $\frac{1}{i} = -i$.) So we see that the eigenvector associated to $\lambda_+ = 1 + i$ is

$$v = \begin{bmatrix} i \\ 1 \end{bmatrix}.$$

There's a short cut to finding the other eigenvector w : since

$$Aw = \lambda_- w = \bar{\lambda}_+ w$$

and $\bar{A} = A$ we must have

$$w = \bar{v} = \begin{bmatrix} -i \\ 1 \end{bmatrix}.$$

We can also do this computation directly:

$$\begin{bmatrix} 1 & -1 \\ 1 & 1 \end{bmatrix} \begin{bmatrix} w_1 \\ w_2 \end{bmatrix} = \begin{bmatrix} (1 - i)w_1 \\ (1 - i)w_2 \end{bmatrix} \Leftrightarrow w_2 = iw_1 \Leftrightarrow w_1 = -iw_2$$

and we recover

$$w = \begin{bmatrix} -i \\ 1 \end{bmatrix}.$$

6.2. Eigenvalues of linear transforms from V to itself. In fact, there is absolutely nothing special about the case $V = \mathbf{R}^n$ in everything we have done regarding eigenvalues and eigenvectors.

Definition 13. Let $T : V \rightarrow V$ be a linear transformation. A nonzero vector $v \in V$ is an eigenvector of T , with eigenvalue λ , if $T(v) = \lambda v$.

As before, $T(v) = \lambda v$ if and only if $v \in \ker(T - \lambda I)$, which in particular implies $T - \lambda I$ is not invertible. If V is finite dimensional, we can choose a basis $\{v_1, \dots, v_n\}$ for it and construct the associated matrix $[T]$, and then compute the eigenvalues by setting $\det([T] - \lambda[I]) = 0$ as before. Below we carry this exercise out for two examples involving $\mathbf{R}_2[x]$, the space of quadratic polynomials.

Example: Let $T : \mathbf{R}_2[x] \rightarrow \mathbf{R}_2[x]$ be given by $T(p)(x) = p(x) - 2xp'(x)$. Observe that

$$T(1) = 1, \quad T(x) = -x, \quad T(x^2) = -3x^2,$$

so our three eigenvectors are $p_0 = 1$ with eigenvalue $\lambda_0 = 1$, $p_1 = x$ with eigenvalue $\lambda_1 = -1$, and $p_2 = x^2$ with eigenvalue $\lambda_2 = -3$. We can verify all this by choosing

the usual basis $\mathcal{B} = \{1, x, x^2\}$; with respect to this basis, we have

$$[T] = \begin{bmatrix} 1 & 0 & 0 \\ 0 & -1 & 0 \\ 0 & 0 & 3 \end{bmatrix}.$$

Example: This time we take $T : \mathbf{R}_2[x] \rightarrow \mathbf{R}_2[x]$ given by $T(p)(x) = 2p(x) + 3p'(x) - 2x^2p''(x)$. We can find one eigenvalue-eigenvector pair immediately by observing $T(1) = 2$; so $p_0 = 1$ is an eigenvector with eigenvalue $\lambda_0 = 2$. To find the others, we choose the usual basis $\mathcal{B} = \{1, x, x^2\}$ for $\mathbf{R}_2[x]$; with respect to this basis, the matrix of T is

$$[T] = \begin{bmatrix} 2 & 5 & 0 \\ 0 & 0 & 6 \\ 0 & 0 & -2 \end{bmatrix},$$

so eigenvalues are roots of the characteristic polynomial

$$0 = \det \begin{bmatrix} 2 - \lambda & 5 & 0 \\ 0 & -\lambda & 6 \\ 0 & 0 & -2 - \lambda \end{bmatrix} = -\lambda(\lambda - 2)(\lambda + 2).$$

We can now read off that the three eigenvalues are $\lambda_0 = 2$ (which we already knew), $\lambda_1 = -2$, and $\lambda_2 = 0$. Finally, we find the eigenvectors. To find the eigenvector p_1 associated to $\lambda_1 = -2$, we solve the system of equations

$$\begin{bmatrix} 2 & 5 & 0 \\ 0 & 0 & 6 \\ 0 & 0 & -2 \end{bmatrix} \begin{bmatrix} a \\ b \\ c \end{bmatrix} = \begin{bmatrix} 2a + 5b \\ 6c \\ -2c \end{bmatrix} = -2 \begin{bmatrix} a \\ b \\ c \end{bmatrix}.$$

The general solution is $a = -\frac{5}{4}b$, $c = -\frac{1}{3}b$, with b being a free variable. We can set $b = 1$, to get the eigenvector $p_1 = -\frac{5}{4} + x - \frac{1}{3}x^2$. Similarly, we find the eigenvector p_2 associated to $\lambda_2 = 0$ by solving the system of equations

$$\begin{bmatrix} 2 & 5 & 0 \\ 0 & 0 & 6 \\ 0 & 0 & -2 \end{bmatrix} \begin{bmatrix} a \\ b \\ c \end{bmatrix} = \begin{bmatrix} 2a + 5b \\ 6c \\ -2c \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \\ 0 \end{bmatrix}.$$

This time the general solution is $c = 0$, $a = -\frac{5}{2}b$. We can take $b = 1$ to get the eigenvector $p_2 = -\frac{5}{2} + x$.

7. SPACES OF LINEAR TRANSFORMATIONS

7.1. Generalities. We have seen now that we can associate a matrix to any linear transformation. We also learned (last year) that one can perform some operations on matrices, particularly addition, multiplication by scalars, and matrix multiplication. We've already seen that matrix multiplication corresponds to the composition of linear transformation, but what does the linear combination of matrices correspond to?

Definition 14. Let V, W be vector spaces, and let $T, U : V \rightarrow W$ both be linear transformations. Then, for any $a, b \in \mathbf{R}$, we can define the linear transformation

$$(aT + bU) : V \rightarrow W, \quad (aT + bU)(v) = aT(v) + bU(v).$$

The following proposition is easy to prove, so we will sketch part of the proof and leave the rest to the reader.

Proposition 37. Let V, W be vector spaces, and let $T, U : V \rightarrow W$ be linear transformations. Then, for any $a, b \in \mathbf{R}$ the map $aT + bU$ defined above is linear. In fact, the set $\mathcal{L}(V, W)$ of linear transformations from V to W is a vector space, with this notion of addition and scalar multiplication.

Remark 7. Be careful that you do not confuse $\mathcal{L}(V, W)$ with $\mathcal{L}(W, V)$. They are the same space if and only if $V = W$.

Proof. Let $v_1, v_2 \in V$ and $\alpha_1, \alpha_2 \in \mathbf{R}$. Then

$$(aT + bU)(\alpha_1 v_1 + \alpha_2 v_2) = aT(\alpha_1 v_1 + \alpha_2 v_2) + bU(\alpha_1 v_1 + \alpha_2 v_2) = \alpha_1(aT + bU)(v_1) + \alpha_2(aT + bU)(v_2),$$

and so $aT + bU$ is indeed a linear transformation.

The zero transformation $Z : V \rightarrow W$ which sends everything to zero, plays the role of the zero vector. It is now easy to check all the vector space axioms. For instance, for any $T : V \rightarrow W$ and $v \in V$ we have

$$(T + Z)(v) = T(v) + Z(v) = T(v) = Z(v) + T(v) = (Z + T)(v).$$

□

We have now shown the following theorem.

Theorem 38. Let V, W be finite dimension vector spaces, with $\dim(V) = n$, $\dim(W) = m$. Choosing bases $\{v_1, \dots, v_n\}$ for V and $\{w_1, \dots, w_m\}$ for W identifies $\mathcal{L}(V, W)$ with $M_{m \times n}$, the set of matrices with m rows and n columns. This identification is a linear isomorphism.

Remark 8. • In fact, in your next algebra course, you will learn that this identification is an isomorphism of rings precisely when $m = n$ and one uses the same basis for both the domain and the target.

• The particular isomorphism $\mathcal{L}(V, W) \simeq M_{m \times n}$ depends very much on our choice of the bases $\{v_1, \dots, v_n\}$ for V and $\{w_1, \dots, w_m\}$ for W .

We check one thing quickly. Let $T, U : V \rightarrow W$ be linear, choose bases $\mathcal{B} = \{v_1, \dots, v_n\}$ and $\mathcal{A} = \{w_1, \dots, w_m\}$ for V and W respectively, and let $\alpha, \beta \in \mathbf{R}$. Then we claim

$$[\alpha T + \beta U]_{\mathcal{B}}^{\mathcal{A}} = \alpha [T]_{\mathcal{B}}^{\mathcal{A}} + \beta [U]_{\mathcal{B}}^{\mathcal{A}}.$$

Indeed, we only need to check this formula on the basis vectors $v_1, \dots, v_n$ in $\mathcal{B}$. We must have numbers $a_{ij}, b_{ij} \in \mathbf{R}$ such that

$$T(v_j) = \sum_{i=1}^m a_{ij} w_i, \quad U(v_j) = \sum_{i=1}^m b_{ij} w_i.$$

Now,

$$(\alpha T + \beta U)(v_j) = \sum_{i=1}^m (\alpha a_{ij} + \beta b_{ij}) w_i,$$

which shows

$$[\alpha T + \beta U]_{\mathcal{B}}^A = \alpha[a_{ij}] + \beta[b_{ij}] = \alpha[T]_{\mathcal{B}}^A + \beta[U]_{\mathcal{B}}^A.$$

7.2. Dual spaces. Given two vector spaces V and W , we have just seen the vector space $\mathcal{L}(V, W)$, the vector space of linear transformations from V to W . In the special case that the target $W = \mathbf{R}$, we obtain something called the dual space to V .

Definition 15. Let V be a vector space. The dual space V^* to V is the vector space $V^* = \mathcal{L}(V, \mathbf{R})$.

An element $\phi \in V^*$ is a linear function $\phi : V \rightarrow \mathbf{R}$; such a linear map is often called a **linear functional**. In the special case that $\mathcal{B} = \{v_1, \dots, v_n\}$ is a basis for V , we have a particular set of n linear functionals $\mathcal{B}^* = \{\phi_1, \dots, \phi_n\}$, which are defined by the relation

$$\phi_i(v_j) = \begin{cases} 1 & i = j \\ 0 & i \neq j. \end{cases}$$

The main content of the proof of the theorem below is to show that $\mathcal{B}^*$ is a basis of V^* ; it is called the **dual basis** to $\mathcal{B}$.

Theorem 39. Let V be a finite dimensional vector space. Then V^* is isomorphic to V .

Proof. We pick a basis $\mathcal{B} = \{v_1, \dots, v_n\}$ and prove the theorem. In fact, we already know that V is isomorphic to $\mathbf{R}^n$, so we only need to prove that V^* is also isomorphic to $\mathbf{R}^n$ as well. To show this, we only need to show that V^* is an n -dimensional vector space, *i.e.* that V^* has a basis with n elements. We already have an excellent candidate for this basis, namely the dual basis $\mathcal{B}^* = \{\phi_1, \dots, \phi_n\}$.

Let $\phi \in V^*$. We show that there is a unique choice of coefficients $a_1, \dots, a_n$ such that

$$\phi = a_1\phi_1 + a_2\phi_2 + \dots + a_n\phi_n;$$

as we have seen before, this shows both that $\mathcal{B}^*$ spans V^* (by the existence of a solution) and that it is linearly independent (by the uniqueness of the solution). Given our choice of $\phi \in V^*$, we can pick

$$a_1 = \phi(v_1), \quad a_2 = \phi(v_2), \quad \dots, \quad a_n = \phi(v_n).$$

Now we compare ϕ to the map

$$\tilde{\phi} = a_1\phi_1 + a_2\phi_2 + \dots + a_n\phi_n.$$

We have, by definition, $\phi(v_j) = a_j = \tilde{\phi}(v_j)$, which implies $\phi = \tilde{\phi}$. Thus we have represented our arbitrary $\phi \in V^*$ as a linear combination of elements of $\mathcal{B}^*$. Furthermore, if there is some other choice of coefficients $b_1, \dots, b_n$ such that

$$\phi = b_1\phi_1 + \dots + b_n\phi_n,$$

then for any $j = 1, \dots, n$ we have

$$b_j = (b_1\phi_1 + \dots + b_n\phi_n)(v_j) = \phi(v_j) = (a_1\phi_1 + \dots + a_n\phi_n)(v_j) = a_j.$$

Thus the choice of coefficients is unique, completing the proof. $\square$

Remark 9. • Notice that isomorphism $V \rightarrow V^*$ depends quite heavily on the choice of basis $\mathcal{B}$ (or, equivalently, on the choice of dual basis $\mathcal{B}^*$). If we change basis in V , we will obtain a very different isomorphism $V \rightarrow V^*$.
 • Also notice that in order to construct the dual basis $\mathcal{B}^*$, we need to first pick the basis $\mathcal{B}$. We cannot pick them simultaneously, say choosing the first two elements of $\mathcal{B}$, then the elements of $\mathcal{B}^*$, then the rest of $\mathcal{B}$. This is because in order to describe even the first linear functional $\phi_1 \in \mathcal{B}^*$, we need to know what it does to each and every basis vector in $\mathcal{B}$.

Example: As our first example, we take $V = \mathbf{R}^n$ with the standard basis $\mathcal{B} = \{e_1, \dots, e_n\}$, and we denote the elements of the dual basis by $\{e_1^*, \dots, e_n^*\}$. By definition we must have

$$e_i^*(x_1, x_2, \dots, x_n) = x_i,$$

and we can write the matrix of e_i^* (with respect to the standard basis) as the row with n element, which has a 1 in the i th column and 0's everywhere else.

Example: Let $V = \mathbf{R}^3$ and choose the basis $\mathcal{B} = \{e_1 - e_2, e_1 + e_2, e_3\}$. Then the dual basis is

$$\mathcal{B}^* = \left\{ \frac{1}{2}(e_1^* - e_2^*), \frac{1}{2}(e_1^* + e_2^*), e_3^* \right\}.$$

Exercise: Let $V = \mathbf{R}^3$, and choose the basis

$$\mathcal{B} = \{e_1 - e_2, e_1 + e_2, e_1 + e_2 + e_3\}.$$

Find the dual basis $\mathcal{B}^*$.

We have seen that the isomorphism $V \rightarrow V^*$ depends very much on our choice of basis. Thus it is remarkable that the isomorphism $V \rightarrow (V^*)^* = V^{**}$ is canonical; that is, this isomorphism doesn't depend on anything at all.

Theorem 40. *Let V be a finite dimensional vector space. Then V is canonically isomorphic to its double-dual V^{**} . In other words, the isomorphism $V \rightarrow V^{**}$ does not depend on anything at all.*

Proof. Given a vector $v \in V$, we define $\hat{v} \in V^{**}$ by $\hat{v}(\phi) = \phi(v)$ for all $\phi \in V^*$. This now defines a mapping

$$\Psi : V \rightarrow V^{**}, \quad \Psi(v) = \hat{v},$$

which, it turns out, will be our isomorphism. Observe that we have now defined $\hat{v}$ and Ψ without choosing anything at all, and so Ψ is canonically defined.

We first check that Ψ is linear. Let $v_1, v_2 \in V$ and choose two real numbers $a_1, a_2 \in \mathbf{R}$. Then, for any $\phi \in V^*$ we have

$$\Psi(a_1v_1 + a_2v_2)(\phi) = \phi(a_1v_1 + a_2v_2) = a_1\phi(v_1) + a_2\phi(v_2) = a_1\Psi(v_1)(\phi) + a_2\Psi(v_2)(\phi).$$

Since this formula holds for all $\phi \in V^*$, we conclude that

$$\Psi(a_1v_1 + a_2v_2) = a_1\Psi(v_1) + a_2\Psi(v_2).$$

Next we check that $\ker(\Psi) = \{0\}$. Suppose that $\Psi(v) = 0 \in V^{**}$. This means that for all $\phi \in V^*$ we must have

$$0 = \Psi(v)(\phi) = \phi(v).$$

If $v \neq 0$ then we could choose an ordered basis $\{v_1 = v, v_2, \dots, v_n\}$, and choosing $\phi = \phi_1$ to be the first element of the dual basis would force $\phi(v) = 1 \neq 0$. Clearly this is impossible in light of $\phi(v) = 0$ for all $\phi \in V^*$.

In fact, we're done with the proof now. We know that $\dim(V) = \dim(V^*) = \dim(V^{**})$, and we've just shown that $\Psi : V \rightarrow V^{**}$ is a one-to-one linear transformation. By our previous theorems we know this is only possible if Ψ is an isomorphism. $\square$

Remark 10. *Some people call the isomorphism Ψ the **tautological isomorphism** between V and V^{**} .*

Remark 11. *We have done everything here for finite dimensional vector spaces. Some of these ideas extend to infinite dimensions (e.g. the idea of constructing a dual vector space), but many of the proofs don't carry through in higher dimensions. In particular, the dual space to an infinite dimensional vector space is bigger than the original.*

8. INNER PRODUCTS AND OTHER QUADRATIC FORMS

In this section we discuss inner products, which generalize the usual dot product in Euclidean space.

8.1. Definitions, examples, and basic properties. We begin with some definitions, examples, and basic properties.

Definition 16. *Let V be a vector space over $\mathbf{R}$. An inner product is a function*

$$\langle \cdot, \cdot \rangle : V \times V \rightarrow \mathbf{R}.$$

which satisfies the following rules for any choice $u, v, w \in V$ and $a, b \in \mathbf{R}$.

- (1) $\langle au + bv, w \rangle = a\langle u, w \rangle + b\langle v, w \rangle$
- (2) $\langle v, w \rangle = \langle w, v \rangle$
- (3) $\langle v, v \rangle \geq 0$, with equality if and only if $v = 0$.

As stated, this definition applies only to vector spaces over the real numbers. There is a similar notion for an inner product on a vector space over the complex numbers $\mathbf{C}$. For the sake of completeness, we mention this definition now, even though we will mostly only work with vector spaces over $\mathbf{R}$ and real-valued inner products.

Definition 17. *Let V be a vector space over the complex numbers $\mathbf{C}$. A hermitian inner product is a function*

$$\langle \cdot, \cdot \rangle : V \times V \rightarrow \mathbf{C}$$

which satisfies the following rules for any choice of $u, v, w \in V$ and $a, b \in \mathbf{C}$.

- (1) $\langle au + bv, w \rangle = a\langle u, w \rangle + b\langle v, w \rangle$

- (2) $\langle v, w \rangle = \overline{\langle w, v \rangle}$, where the over-bar denotes the complex conjugate
(3) $\langle v, v \rangle \geq 0$, with equality if and only if $v = 0$

Notice that, in particular, $\langle v, v \rangle$ is a non-negative real number for any $v \in V$.

At this point, we describe some examples.

Example: We take $V = \mathbf{R}^n$ and $\langle \cdot, \cdot \rangle$ to be the usual dot product. In the usual coordinates, we have

$$\langle (x_1, x_2, \dots, x_n), (y_1, y_2, \dots, y_n) \rangle = x_1 y_1 + x_2 y_2 + \dots + x_n y_n = \sum_{j=1}^n x_j y_j.$$

With a little more generality, we can take $V = \mathbf{C}^n$ and define the standard hermitian inner product on $\mathbf{C}^n$ by

$$\langle (z_1, z_2, \dots, z_n), (w_1, w_2, \dots, w_n) \rangle = z_1 \bar{w}_1 + z_2 \bar{w}_2 + \dots + z_n \bar{w}_n = \sum_{j=1}^n z_j \bar{w}_j,$$

where $z = (z_1, \dots, z_n), w = (w_1, \dots, w_n) \in \mathbf{C}^n$. For instance, if we take $z = (1 + i, 4) \in \mathbf{C}^2$ and $w = (2 - 3i, 4 + 5i) \in \mathbf{C}^2$ we have

$$\langle z, w \rangle = (1 + i)(2 + 3i) + (4)(4 - 5i) = 15 - 15i.$$

Example: If $A \in M_{n \times n}$ we have already seen that we can define the transpose of A as A^t , the matrix we obtain by swapping the rows and the columns. In components, if a_{ij} is the entry in the i th row, j th column of A , then A^t has a_{ij} in the i th column, j th row. Now we can define an inner product

$$\langle \cdot, \cdot \rangle : M_{n \times n} \times M_{n \times n} \rightarrow \mathbf{R}, \quad \langle A, B \rangle = \text{tr}(B^t A).$$

We check this is actually an inner product; we need to recall that $\text{tr}(c_1 T_1 + c_2 T_2) = c_1 \text{tr}(T_1) + c_2 \text{tr}(T_2)$, that $\text{tr}(T^t) = \text{tr}(T)$ and the formula for transposes and matrix products. If $A_1, A_2 \in M_{n \times n}$ and $c_1, c_2 \in \mathbf{R}$ then

$$\langle c_1 A_1 + c_2 A_2, B \rangle = \text{tr}(B^t(c_1 A_1 + c_2 A_2)) = c_1 \text{tr}(B^t A_1) + c_2 \text{tr}(B^t A_2) = c_1 \langle A_1, B \rangle + c_2 \langle A_2, B \rangle.$$

Also,

$$\langle B, A \rangle = \text{tr}(A^t B) = \text{tr}((B^t A)^t) = \text{tr}(B^t A) = \langle A, B \rangle.$$

Finally, if $A = [a_{ij}]$ then

$$\langle A, A \rangle = \sum_{i=1}^n (A^t A)_{ii} = \sum_{i=1}^n \sum_{j=1}^n a_{ji} a_{ji}.$$

This is a sum of non-negative numbers, so it must be non-negative. Moreover, it is only zero if all the entries of A are zero, namely only if A is the zero matrix.

Example: Let $V = \mathcal{C}[0, 2\pi]$, the space of continuous, real-valued functions on $[0, 2\pi]$. We can define the inner product

$$\langle \cdot, \cdot \rangle : \mathcal{C}[0, 2\pi] \times \mathcal{C}[0, 2\pi] \rightarrow \mathbf{R}, \quad \langle f, g \rangle = \frac{1}{2\pi} \int_0^{2\pi} f(x)g(x)dx.$$

This is an important inner product for defining Fourier series, and the factor of $\frac{1}{2\pi}$ in front of the integral makes many of the formulas come out nicely.

It will be useful to have two more notions for the rest of this chapter: the notion of a length (or a norm), and the notion of an orthogonal set.

Definition 18. Let V be a vector space over $\mathbf{R}$ with an inner product $\langle \cdot, \cdot \rangle$. Then we can define the length of any vector $v \in V$ to be $\|v\| = \sqrt{\langle v, v \rangle}$

It is immediate from the definitions that $\|v\| \geq 0$ for any vector $v \in V$, with equality if and only if $v = 0$. By convention, we say that $v \in V$ is a **unit vector** if $\|v\| = 1$.

Definition 19. Let V be a vector space over $\mathbf{R}$ with an inner product $\langle \cdot, \cdot \rangle$. A set $S = \{v_1, v_2, \dots, v_k\}$ is orthogonal if $\langle v_i, v_j \rangle = 0$ whenever $i \neq j$. If in addition $\|v_i\| = 1$ for each $i = 1, 2, \dots, k$ then we say S is an orthonormal set.

We will revisit orthonormal sets in the next subsection.

We close this section by listing some basic properties of inner products.

Theorem 41. Let V be a vector space over $\mathbf{R}$ with an inner product $\langle \cdot, \cdot \rangle$. Then for any $u, v, w \in V$ and $a, b \in \mathbf{R}$ we have

- (1) $\langle u, av + bw \rangle = a\langle u, v \rangle + b\langle u, w \rangle$
- (2) If $\langle u, v \rangle = \langle u, w \rangle$ for all $u \in V$ then $v = w$.

Proof. The first statement follows from the properties of the inner product and the fact that $\langle u, v \rangle = \langle v, u \rangle$. Suppose that $\langle u, v \rangle = \langle u, w \rangle$ for all $u \in V$. We can rearrange this equation to read

$$0 = \langle u, v \rangle - \langle u, w \rangle = \langle u, v - w \rangle$$

for all u . Now substitute $u = v - w$ into this equation to get $\langle v - w, v - w \rangle = 0$, which implies $v - w = 0$, i.e. $v = w$. $\square$

Theorem 42. Let V be a vector space over $\mathbf{R}$ with an inner product $\langle \cdot, \cdot \rangle$. Then for any $u, v \in V$ and $c \in \mathbf{R}$ we have

- (1) $\|cv\| = |c|\|v\|$
- (2) $|\langle u, v \rangle| \leq \|u\|\|v\|$
- (3) $\|u + v\| \leq \|u\| + \|v\|$

The second statement is called the Cauchy-Schwarz inequality, and the third statement is the familiar triangle inequality.

Proof. The first equation follows from the fact that $\langle cv, cv \rangle = c^2\langle v, v \rangle$.

Next we prove the Cauchy Schwarz inequality. If $v = 0$ then both sides of the equation are zero, so the inequality holds. If $v \neq 0$ then we have, for any c ,

$$0 \leq \|u - cv\|^2 = \langle u - cv, u - cv \rangle = \|u\|^2 + c^2\|v\|^2 - 2c\langle u, v \rangle$$

Now choose

$$c = \frac{\langle u, v \rangle}{\|v\|^2}$$

and plug this into the inequality we have just derived to get

$$0 \leq \|u\|^2 - \frac{|\langle u, v \rangle|^2}{\|v\|^2},$$

which we can rearrange to give the Cauchy-Schwarz inequality.

Finally we use Cauchy-Schwarz to prove the triangle inequality.

$$\begin{aligned}\|u + v\|^2 &= \langle u + v, u + v \rangle = \|u\|^2 + 2\langle u, v \rangle + \|v\|^2 \\ &\leq \|u\|^2 + 2|\langle u, v \rangle| + \|v\|^2 \leq \|u\|^2 + 2\|u\|\|v\| + \|v\|^2 \\ &\leq (\|u\| + \|v\|)^2\end{aligned}$$

Taking square roots of both sides of this inequality gives the triangle inequality. $\square$

8.2. Gram-Schmidt orthogonalization. For many applications it is convenient to have an orthonormal basis for a vector space with an inner product. However, when one originally picks a basis for a vector space it is usually not orthogonal, so it is useful to have a process to turn an arbitrary basis into an orthonormal basis. This subsection will outline an algorithm for doing so, and present some examples.

Recall that an orthogonal set $S = \{v_1, \dots, v_k\} \subset V$ is a set such that $\langle v_i, v_j \rangle = 0$ if $i \neq j$. We say that S is orthonormal if we also have $\|v_i\|^2 = \langle v_i, v_i \rangle = 1$ for $i = 1, 2, \dots, k$. Thus an orthonormal basis for V is a basis $\{v_1, v_2, \dots, v_n\}$ such that

$$\langle v_i, v_j \rangle = \begin{cases} 1 & i = j \\ 0 & i \neq j. \end{cases}$$

Example: We list some orthonormal bases for $\mathbf{R}^2$ equipped with the usual dot product. The first such basis is the familiar standard basis

$$\{e_1 = (1, 0), e_2 = (0, 1)\}.$$

We can rotate this basis through an angle $\pi/4$ to get

$$\{v_1 = \frac{1}{\sqrt{2}}(1, 1), v_2 = \frac{1}{\sqrt{2}}(1, -1)\},$$

or through an angle $\pi/6$ to get

$$\{v_1 = \frac{1}{2}(\sqrt{3}, 1), v_2 = \frac{1}{2}(-1, \sqrt{3})\},$$

or (for a general angle θ)

$$\{v_1 = (\cos \theta, \sin \theta), v_2 = (-\sin \theta, \cos \theta)\}.$$

Before outlining the Gram-Schmidt process we should see why an orthonormal basis is useful.

Theorem 43. *Let V be a vector space over $\mathbf{R}$ with an inner product $\langle \cdot, \cdot \rangle$, and let $S = \{v_1, \dots, v_k\}$ be an orthogonal set of nonzero vectors. If $w = \sum_{j=1}^k a_j v_j$ then $a_j = \frac{\langle w, v_j \rangle}{\|v_j\|^2}$ for all $j = 1, 2, \dots, k$.*

Proof. We take some inner products. For $i = 1, \dots, k$ we have

$$\langle w, v_i \rangle = \langle \sum_{j=1}^k a_j v_j, v_i \rangle = \sum_{j=1}^k a_j \langle v_j, v_i \rangle = a_i \|v_i\|^2.$$

The theorem follows. $\square$

Corollary 44. *Under the hypotheses of the last theorem, if S is also orthonormal then*

$$w = \sum_{j=1}^k \langle w, v_j \rangle v_j.$$

Corollary 45. *Let V be a vector space over $\mathbf{R}$ with an inner product $\langle \cdot, \cdot \rangle$, and let $S = \{v_1, v_2, \dots, v_k\}$ be an orthogonal set of nonzero vectors. Then S is also linearly independent.*

Proof. Suppose there exist coefficients $a_1, \dots, a_k$ such that

$$a_1 v_1 + a_2 v_2 + \dots + a_k v_k = 0.$$

Then, by the theorem above, for each $i = 1, 2, \dots, k$ we have $a_i = \frac{\langle 0, v_i \rangle}{\|v_i\|^2} = 0$. We conclude that S must be linearly independent. $\square$

The following theorem is the main tool we will use to construct an orthonormal basis from an arbitrary basis.

Theorem 46. *Let V be a vector space over $\mathbf{R}$ with an inner product $\langle \cdot, \cdot \rangle$ and let $S = \{w_1, w_2, \dots, w_n\} \subset V$ be a linearly independent subset. If we define*

$$v_1 = w_1, \quad v_k = w_k - \sum_{j=1}^{k-1} \frac{\langle w_k, v_j \rangle}{\|v_j\|^2} v_j \text{ for } k = 2, 3, \dots, n$$

then $S' = \{v_1, v_2, \dots, v_n\}$ is an orthogonal set with $\text{span}(S') = \text{span}(S)$.

If we want to complete the process of finding an orthonormal basis, we apply the above theorem to a basis to get an orthogonal basis $\{v_1, \dots, v_n\}$, and then replace v_j with $\frac{1}{\|v_j\|} v_j$. Thus the main difficulty lies in achieving $\langle v_i, v_j \rangle = 0$ for $i \neq j$.

Proof. We prove this by induction on n . If $n = 1$ then there is nothing to check, and so the theorem holds. Assume now that the statement of the theorem holds for $S_n = \{w_1, \dots, w_n\}$, namely that we can construct $S'_n = \{v_1, \dots, v_n\}$ using the formulas above to find an orthogonal set such that $\text{span}(S_n) = \text{span}(S'_n)$. Now suppose we have $S_{n+1} = \{w_1, \dots, w_n, w_{n+1}\}$ which is linearly independent; we wish to construct $S'_{n+1} = \{v_1, \dots, v_n, v_{n+1}\}$ according to our formula, and see that we get an orthogonal set with $\text{span}(S_{n+1}) = \text{span}(S'_{n+1})$.

By the induction hypothesis, we already have $\{v_1, \dots, v_n\}$ which are orthogonal (namely, $\langle v_i, v_j \rangle = 0$ for $i \neq j$ so long as $1 \leq i \leq n$ and $1 \leq j \leq n$), so we only need to check that

$$v_{n+1} = w_{n+1} - \sum_{j=1}^n \frac{\langle w_{n+1}, v_j \rangle}{\|v_j\|^2} v_j$$

is nonzero and orthogonal to $v_1, v_2, \dots, v_n$. If $v_{n+1} = 0$ then

$$w_{n+1} = \sum_{j=1}^n \frac{\langle w_{n+1}, v_j \rangle}{\|v_j\|^2} v_j \in \text{span}\{v_1, \dots, v_n\} = \text{span}\{w_1, \dots, w_n\}.$$

However, this contradicts the fact that S_{n+1} is linearly independent, so it is impossible that $v_{n+1} = 0$. Next we check the inner products. Using the fact that $\{v_1, \dots, v_n\}$ are already orthogonal, for $1 \leq k \leq n$ we have

$$\langle v_{n+1}, v_k \rangle = \langle w_{n+1}, v_k \rangle - \sum_{j=1}^n \frac{\langle w_{n+1}, v_j \rangle}{\|v_j\|^2} \langle v_j, v_k \rangle = \langle w_{n+1}, v_k \rangle - \frac{\langle w_{n+1}, v_k \rangle}{\|v_k\|^2} \langle v_k, v_k \rangle = 0$$

We conclude that S'_{n+1} is indeed orthogonal. By our formula for $\{v_1, v_2, \dots, v_{n+1}\}$ we have $\text{span}(S'_{n+1}) \subset \text{span}(S_{n+1})$. However, we also know that S'_{n+1} is linearly independent, so

$$\dim(\text{span}(S'_{n+1})) = n + 1 = \dim(\text{span}(S_{n+1})),$$

which implies $\text{span}(S'_{n+1}) = \text{span}(S_{n+1})$. $\square$

We can combine the previous theorems to see the following.

Theorem 47. *Let V be a finite dimensional vector space over $\mathbf{R}$ with an inner product. Then V has an orthonormal basis $\mathcal{B} = \{v_1, \dots, v_n\}$. Moreover, if $v \in V$ then we can represent*

$$v = \sum_{i=1}^n a_i v_i, \quad a_i = \langle v, v_i \rangle.$$

Proof. Let $\mathcal{A}_1 = \{w_1, \dots, w_n\}$ be any basis for V and apply Gram-Schmidt to $\mathcal{A}_1$ to obtain $\mathcal{A}_2 = \{u_1, \dots, u_n\}$. This is now an orthogonal basis, so we let

$$v_i = \frac{u_i}{\|u_i\|}, \quad \mathcal{B} = \{v_1, \dots, v_n\}.$$

This completes the proof. $\square$

Corollary 48. *Let V be an n -dimensional vector product space over $\mathbf{R}$ with an inner product, and let $\mathcal{B} = \{v_1, \dots, v_n\}$ be an orthonormal basis for V . Also, let $T : V \rightarrow V$ be a linear transformation. Then*

$$[T]_{\mathcal{B}}^{\mathcal{B}} = [a_{ij}], \quad a_{ij} = \langle T(v_i), v_j \rangle.$$

9. NEW VECTOR SPACES FROM OLD

In this section we describe some ways to create new vector spaces from old ones, namely the direct sum construction and the quotient construction. The key tool in each construction is a general notion of a projection.

Definition 20. *Let V be a vector space. A linear map $P : V \rightarrow V$ is a projection if $P^2 = P$.*

It is worthwhile to check that the usual coordinate projections

$$P_1 : \mathbf{R}^2 \rightarrow \mathbf{R}^2, \quad P_1(x_1, x_2) = (x_1, 0)$$

and

$$P_2 : \mathbf{R}^2 \rightarrow \mathbf{R}^2, \quad P_2(x_1, x_2) = (0, x_2)$$

satisfy $P_1^2 = P_1$ and $P_2^2 = P_2$.

9.1. Direct sums. We begin with some definitions and basic properties.

Definition 21. Let V be a vector space, and let $W_1, W_2 \subset V$ be subspace. Then

$$W_1 + W_2 = \{w_1 + w_2 : w_1 \in W_1, w_2 \in W_2\}.$$

Proposition 49. The sum $W_1 + W_2$ is a subspace of V . In fact, it is the smallest subspace containing both W_1 and W_2 .

Proof. We need first to verify that if $w, \tilde{w} \in W_1 + W_2$ and $a, \tilde{a} \in \mathbf{R}$ then $aw + \tilde{a}\tilde{w} \in W_1 + W_2$. We have $w = w_1 + w_2$, where $w_1 \in W_1$ and $w_2 \in W_2$. Similarly, $\tilde{w} = \tilde{w}_1 + \tilde{w}_2$, where $\tilde{w}_1 \in W_1$ and $\tilde{w}_2 \in W_2$. Then

$$aw + \tilde{a}\tilde{w} = a(w_1 + w_2) + \tilde{a}(\tilde{w}_1 + \tilde{w}_2) = (aw_1 + \tilde{a}\tilde{w}_1) + (aw_2 + \tilde{a}\tilde{w}_2),$$

and we have now written $aw + \tilde{a}\tilde{w}$ as the sum of a vector in W_1 and a vector in W_2 . Thus $W_1 + W_2$ is closed under linear combinations, so it must be a subspace.

Now suppose U is another subspace of V containing both W_1 and W_2 , and choose $w = w_1 + w_2 \in W_1 + W_2$. However, w_1 and w_2 must both be elements of U , and U is closed under addition, so $w_1 + w_2 \in U$. We have just shown that $W_1 + W_2 \subset U$. $\square$

Definition 22. A vector space V is said to be the direct sum of the subspaces W_1 and W_2 if $W_1 \cap W_2 = \{0\}$ and $V = W_1 + W_2$. In this case we write $V = W_1 \oplus W_2$.

Example: We can write $\mathbf{R}^n = W_1 \oplus W_2$ where $W_1 = \{(x_1, x_2, \dots, x_n) : x_n = 0\}$ and $W_2 = \{(x_1, x_2, \dots, x_n) : x_1 = x_2 = \dots = x_{n-1} = 0\}$. We can write any vector in $\mathbf{R}^n$ as

$$(x_1, \dots, x_n) = (x_1, \dots, x_{n-1}, 0) + (0, \dots, 0, x_n),$$

and so $\mathbf{R}^n = W_1 + W_2$. Moreover, the only vector in common with W_1 and W_2 must have all its components equal to zero, so $W_1 \cap W_2 = \{0\}$.

Example: Let $V = M_{n \times n}$, let $W_1 = \{A \in M_{n \times n} : A^t = A\}$ be the set of symmetric matrices, and let $W_2 = \{A \in M_{n \times n} : A^t = -A\}$ be the set of skew-symmetric matrices. Then $V = W_1 \oplus W_2$. Indeed, for any matrix $A \in M_{n \times n}$ we have

$$\frac{1}{2}(A + A^t) \in W_1, \quad \frac{1}{2}(A - A^t) \in W_2, \quad A = \frac{1}{2}(A + A^t) + \frac{1}{2}(A - A^t),$$

so $V = W_1 + W_2$. Now suppose $A \in W_1 \cap W_2$. Then

$$A \in W_1 \Rightarrow A = A^t, \quad A \in W_2 \Rightarrow A = -A^t.$$

Putting these last two equations together we get $A = -A$, which is only possible if $A = 0$.

Exercise: Let $V = \mathbf{R}_n[x]$, the space of polynomials of degree at most n . Define

$$W_1 = \{p = a_0 + a_1x + \dots + a_nx^n : a_k = 0 \text{ if } k \text{ is even}\},$$

and

$$W_2 = \{p = a_0 + a_1x + \dots + a_nx^n : a_k = 0 \text{ if } k \text{ is odd}\}.$$

Prove that $V = W_1 \oplus W_2$.

Theorem 50. Let W_1 and W_2 be finite-dimensional subspaces of the vector space V . Then $W_1 + W_2$ is also finite dimensional, and

$$\dim(W_1 + W_2) = \dim(W_1) + \dim(W_2) - \dim(W_1 \cap W_2).$$

In particular, $V = W_1 \oplus W_2$ if and only if

$$\dim(V) = \dim(W_1) + \dim(W_2).$$

Proof. Let $\mathcal{B}_1 = \{v_1, \dots, v_k\}$ be a basis for W_1 and let $\mathcal{B}_2 = \{w_1, \dots, w_m\}$ be a basis for W_2 . Then

$$W_1 + W_2 \subset \text{span}(\mathcal{B}_1 \cup \mathcal{B}_2) = \text{span}\{v_1, \dots, v_k, w_1, \dots, w_m\},$$

so in particular

$$\dim(W_1 + W_2) \leq k + m = \dim(W_1) + \dim(W_2),$$

which implies $W_1 + W_2$ is finite dimensional.

If $W_1 \cap W_2 = \{0\}$ then $\text{span}(\mathcal{B}_1)$ and $\text{span}(\mathcal{B}_2)$ cannot contain any vectors in common, and so $\mathcal{B}_1 \cup \mathcal{B}_2$ must be linearly independent. In this case, we have shown $\dim(W_1 + W_2) = \dim(W_1) + \dim(W_2) - \dim\{0\} = \dim(W_1) + \dim(W_2) - \dim(W_1 \cap W_2)$ as desired.

Otherwise, we still have that $W_1 \cap W_2$ is a subspace of W_2 . (It is also a subspace of W_1 and of V , but for the purposes of the proof we will work with $W_1 \cap W_2$ as a subspace of W_2 .) By the replacement theorem, we can re-order $\{w_1, \dots, w_m\}$ so that $\{w_1, w_2, \dots, w_{m'}\}$ span $W_1 \cap W_2$, for some $m' \leq m$. As each of $w_1, w_2, \dots, w_{m'}$ must also be in $W_1 = \text{span}(\mathcal{B}_1)$, each must be a linear combination of $v_1, v_2, \dots, v_k$. Removing these vectors $w_1, \dots, w_{m'}$ from $\mathcal{B}_2$ we still have a set

$$\mathcal{B}' = \{v_1, \dots, v_k, w_{m'+1}, \dots, w_m\}$$

such that $W_1 + W_2 = \text{span}(\mathcal{B}')$. If we show $\mathcal{B}'$ is linearly independent then we're done, because in this case

$$\dim(W_1 + W_2) = k + (m - m') = k + m - m' = \dim(W_1) + \dim(W_2) - \dim(W_1 \cap W_2).$$

So, suppose there is some choice of coefficients $a_1, \dots, a_k, b_{m'+1}, \dots, b_m \in \mathbf{R}$, not all of which are zero, such that

$$0 = a_1 v_1 + a_2 v_2 + \dots + a_k v_k + b_{m'+1} w_{m'+1} + b_{m'+2} w_{m'+2} + \dots + b_m w_m.$$

If one of the a_j 's is non-zero, we can re-order the vectors $v_1, \dots, v_k$ so that $a_1 \neq 0$. Then we rearrange our equation to read

$$v_1 = -\frac{1}{a_1}(a_2 v_2 + \dots + a_k v_k + b_{m'+1} w_{m'+1} + \dots + b_m w_m).$$

However, we know that $\{v_1, \dots, v_k\}$ are linearly independent, and $w_{m'+1}, \dots, w_m$ are not in W_1 at all, so this is impossible. Similarly, assuming that one of the b_j 's is nonzero will also give a contradiction. We have just proven that $\mathcal{B}'$ is linearly independent, completing the proof of our theorem. $\square$

If we have $V = W_1 \oplus W_2$ then we can examine two natural projection operators:

$$\Pi_1 : V \rightarrow W_1, \quad \Pi_2 : V \rightarrow W_2.$$

To understand these operators, we need some preliminaries. First observe that we can write any $v \in V$ as $v = \tilde{w}_1 + \tilde{w}_2$, where $\tilde{w}_1 \in W_1$ and $\tilde{w}_2 \in W_2$. In fact, this choice of $\tilde{w}_1$ and $\tilde{w}_2$ is unique. To see this, choose bases $\mathcal{B}_1 = \{v_1, \dots, v_k\}$ for W_1 and $\mathcal{B}_2 = \{w_1, \dots, w_m\}$ for W_2 . By the proof of the theorem immediately above, we know that $\mathcal{B}_1 \cup \mathcal{B}_2$ is a basis for $V = W_1 \oplus W_2$, which means that for any $v \in V$ there is a unique choice of coefficients $a_1, \dots, a_k, b_1, \dots, b_m$ such that

$$v = a_1 v_1 + \dots + a_k v_k + b_1 w_1 + \dots + b_m w_m,$$

and so

$$\tilde{w}_1 = a_1 v_1 + \dots + a_k v_k, \quad \tilde{w}_2 = b_1 w_1 + \dots + b_m w_m.$$

Because $\mathcal{B}_1 = \{v_1, \dots, v_k\}$ is a basis for W_1 and $\mathcal{B}_2 = \{w_1, \dots, w_m\}$ is a basis for W_2 , these representations of $\tilde{w}_1$ and $\tilde{w}_2$ uniquely determine them. Then our projection formulas read

$$\Pi_1(v) = \Pi_1(\tilde{w}_1 + \tilde{w}_2) = \tilde{w}_1, \quad \Pi_2(v) = \Pi_2(\tilde{w}_1 + \tilde{w}_2) = \tilde{w}_2.$$

Remark 12. *It is a straight-forward exercise to extend everything in this section to finite sums (or finite direct sums) of subspaces $W_1, W_2, \dots, W_k \subset V$ using induction.*

9.2. Quotients. To define quotients we will need to introduce equivalence relations.

Definition 23. *If X is any set, a relation R on X is a subset $R \subset X \times X$. If $x_1, x_2 \in X$ and $(x_1, x_2) \in R$ we write $x \sim y$. An equivalence relation is a relation such that*

- (1) *For all $x \in X$ we have $x \sim x$ (reflexivity).*
- (2) *If $x \sim y$ then $y \sim x$ (symmetry).*
- (3) *If $x \sim y$ and $y \sim z$ then $x \sim z$ (transitivity).*

Also, if X is a set with an equivalence relation R then the equivalence class of an element $x \in X$, written $[x]$, is the set of everything equivalent to x . That is, $[x] = \{\tilde{x} \in X : \tilde{x} \sim x\}$.

The first example of an equivalence relation you should think of is the following. Let $X = \mathbf{Z}$, the set of integers, and fix a positive integer n . We say $x \sim y$ if $x - y$ is an integer multiple of n . For instance, if $n = 2$, then $x \sim y$ if $x - y$ is an even integer, and $x \not\sim y$ if $x - y$ is odd.

We'll verify transitivity here, the other two properties are easier to check. If $x \sim y$ and $y \sim z$ then there are integers a and b such that

$$x - y = an, \quad y - z = bn,$$

and so

$$x - z = (an + y) - (-bn + y) = (a + b)n$$

is indeed a multiple of n .

We write

$$n\mathbf{Z} = 0, \pm n, \pm 2n, \pm 3n, \dots$$

for the set of all integer multiples of n , and we see that we can identify the set of equivalence classes

$$\mathbf{Z}_n = \mathbf{Z}/\sim = \mathbf{Z}/n\mathbf{Z} = \{0, 1, 2, \dots, n-1\}.$$

This means that for any integer x we can find another integer $y \in \{0, 1, 2, \dots, n-1\}$ such that $x \sim y$, and that this choice of y is unique. Indeed, we can write any integer x as $x = an + y$, where y is the remainder of x divided by n . By definition, we have $x \sim y$, and that $y \in \{0, 1, 2, \dots, n-1\}$. Also, if $y \sim z$ and $y, z \in \{0, 1, 2, \dots, n-1\}$ then $y - z = an$ for some integer a . However, this is only possible if $a = 0$, which implies $y = z$.

Now we are ready to define the quotient of a vector space V by a subspace W .

Definition 24. Let V be a vector space and let $W \subset V$ be a subspace. For any $v \in V$ we define the coset $v + W = \{v + w : w \in W\}$, and define an equivalence relation on the set of cosets by saying

$$v_1 + W \sim v_2 + W \Leftrightarrow v_1 - v_2 \in W.$$

Then the quotient space V/W is the vector space formed by the set of equivalence classes.

In order for this to be a sensible definition, we must check some things.

Theorem 51. The equivalence relation

$$v_1 + W \sim v_2 + W \Leftrightarrow v_1 - v_2 \in W$$

is indeed an equivalence relation. Moreover, the set of equivalence relations V/W inherits a vector space structure from V .

Proof. We first check that $\sim$ is an equivalence relation. First of all, $v + W \sim v + W$ because $v - v = 0 \in W$, as W is a subspace. Second, if $v_1 + W \sim v_2 + W$ then $v_1 - v_2 \in W$, which implies $v_2 - v_1 = -(v_1 - v_2) \in W$. Third, if $v_1 + W \sim v_2 + W$ and $v_2 + W \sim v_3 + W$, then $v_1 - v_2 = w_1 \in W$ and $v_2 - v_3 = w_2 \in W$. Thus

$$v_1 - v_3 = (v_1 - v_2) - (v_2 - v_3) = w_1 - w_2 \in W,$$

and so $v_1 + W \sim v_3 + W$.

At this point, it will be convenient to notice that $v + W = 0 + W = W$ (with equality as equivalence classes) if and only if $v \in W$. To see this, observe that

$$v + W = 0 + W = W \Leftrightarrow v = v - 0 \in W.$$

Next we define the vector addition and scalar multiplication on the cosets. Let $v_1 + W, v_2 + W \in V/W$ and let $a_1, a_2 \in \mathbf{R}$. Then we define

$$a_1(v_1 + W) + a_2(v_2 + W) = (a_1v_1 + a_2v_2) + W.$$

We need to check that this definition of vector addition and scalar multiplication does not depend on the choice we have made, namely, the representatives v_1 and v_2 we chose for the cosets. Let

$$v_1 + W \sim v'_1 + W \Leftrightarrow v_1 - v'_1 = w_1 \in W, \quad v_2 + W = v'_2 + W \Leftrightarrow v_2 - v'_2 = w_2 \in W$$

Then

$$(a_1v_1 + a_2v_2) - (a_1v'_1 + a_2v'_2) = a_1(v_1 - v'_1) + a_2(v_2 - v'_2) = a_1w_1 + a_2w_2 \in W,$$

so the coset we get by choosing a representative to do the vector addition and scalar multiplication does not in fact depend on our choice of representative.

It is now quite easy to see that V/W inherits all its vector space structure from that of V . $\square$

Just as was the case with direct sums, there is a natural projection operator associated to any quotient, which we define with the following theorem.

Theorem 52. *Let V be a vector space and let $W \subset V$ be a subspace. Define the map*

$$\eta : V \rightarrow V/W, \quad \eta(v) = v + W.$$

Then η is a linear transformation, $\ker(\eta) = W$, and $\dim(V) = \dim(W) + \dim(V/W)$.

Proof. Let $v_1, v_2 \in V$ and $a_1, a_2 \in \mathbf{R}$. Then

$$\eta(a_1v_1 + a_2v_2) = (a_1v_1 + a_2v_2) + W = a_1(v_1 + W) + a_2(v_2 + W) = a_1\eta(v_1) + a_2\eta(v_2),$$

and so η is linear. We have already noticed in the middle of the previous proof that $v + W = 0 + W = W$ if and only if $v \in W$, which (by the definition of η) proves that $\ker(\eta) = W$. The last statement of the theorem now follows immediately from the rank-nullity theorem. $\square$