



Lecture Notes for Linear Algebra

James S. Cook

Liberty University

Department of Mathematics

Fall 2026

preface

Before we begin, I should warn you that I assume a few things from the reader. These notes are intended for someone who has already grappled with the problem of constructing proofs. I assume you know the difference between $\Rightarrow$ and $\Leftrightarrow$. I assume the phrase "iff" is known to you. I assume you are ready and willing to do a proof by induction, strong or weak. I assume you know what $\mathbb{R}$, $\mathbb{C}$, $\mathbb{Q}$, $\mathbb{N}$ and $\mathbb{Z}$ denote. I assume you know what a subset of a set is. I assume you know how to prove two sets are equal. I assume you are familiar with basic set operations such as union and intersection. More importantly, I assume you have started to appreciate that mathematics is more than just calculations. Calculations without context, without theory, are doomed to failure. At a minimum theory and proper mathematics allows you to communicate analytical concepts to other like-educated individuals.

Some of the most seemingly basic objects in mathematics are insidiously complex. We've been taught they're simple since our childhood, but as adults, mathematical adults, we find the actual definitions of such objects as $\mathbb{R}$ or $\mathbb{C}$ are rather involved. I will not attempt to provide foundational arguments to build numbers from basic set theory. I believe it is possible, I think it's well-thought-out mathematics, but we take the existence of the real numbers as a given truth for these notes. We assume that $\mathbb{R}$ exists and that the real numbers possess all their usual properties. In fact, I assume $\mathbb{R}$, $\mathbb{C}$, $\mathbb{Q}$, $\mathbb{N}$ and $\mathbb{Z}$ all exist complete with their standard properties. In short, I assume we have numbers to work with. We leave the rigorization of numbers to a different course.

Just a bit more advice before I get to the good part. How to study? I have a few points:

- spend several days on the homework. Try it by yourself to begin. Later, compare with your study group. Leave yourself time to ask questions.
- come to class, take notes, think about what you need to know to solve problems.
- assemble a list of definitions, try to gain an intuitive picture of each concept, be able to give examples and counter-examples
- learn the notation, a significant part of this course is learning to deal with new notation.
- methods of proof, how do we prove things in linear algebra? There are a few standard proofs, know them.
- method of computation, I show you tools, learn to use them.
- it's not impossible. You can do it. Moreover, doing it the right way will make the courses which follow this easier. Mathematical thinking is something that takes time for most of us to master. You began the process in Math 200, now we continue that process.

style guide

I use a few standard conventions throughout these notes. They were prepared with $\text{\LaTeX}$ which automatically numbers sections and the `hyperref` package provides links within the pdf copy from the Table of Contents as well as other references made within the body of the text.

I use color and some boxes to set apart some points for convenient reference. In particular,

1. definitions are in **green**.
2. remarks are in **red**.
3. theorems, propositions, lemmas and corollaries are in **blue**.
4. proofs start with a **Proof:** and are concluded with a $\square$.

However, I do make some definitions within the body of the text. As a rule, I try to put what I am defining in **bold**. Doubtless, I have failed to live up to my legalism somewhere. If you keep a list of these transgressions to give me at the end of the course it would be worthwhile for all involved.

The symbol $\square$ indicates that a proof is complete. The symbol ∇ indicates part of a proof is done, but it continues.

ERRORS: there may be errors. I make every effort to eliminate them, but, I do make mistakes. When you find one, simply send me an email about your concern anytime, day, night, even the weekend. I usually see it fairly soon and I can confirm or deny the error. Most of the examples are not new, so previous generations of students have combed through them for mistakes. If in doubt email. Please. Thanks!

Note on the 2025 version: I made a significant edit relative to the 2024 version. In particular, I've moved quotient space and dual space to a later chapter so more attention is given to basic theory early in the course. Chapters 4, 5 and 6 are newly formatted. I've also added some silly things here and there including an "about the author" section as well as the final Appendix on history, be sure to read the final remark if you read that Appendix.

Note on the 2026 version: I made a significant edit relative to the 2025 version. In particular, I've moved quotient space and dual space to an earlier chapter so more attention is given to basic theory early in the course. I also removed much of the AI slop which I generously slathered over the previous version. Ultimately, it was just out of place¹. I've added a few verses at the start of chapters to remind us of deeper truths.

¹As a ring of gold in a swine's snout So is a beautiful woman who turns away from discretion. Proverbs 11:22, LSB

reading guide

A number of excellent texts have helped me gain deeper insight into linear algebra. Let me discuss a few of them here.

1. Charles W. Curtis' *Linear Algebra: An Introductory Approach (Undergraduate Texts in Mathematics)*, 4-th Edition was the required text for the Spring 2016 Semester. This is a very good book. Apparently too good. Oh, the complaining. So much complaining. Hence, we use it no more. Sorry for those of you who wanted something deeper. Still, certain aspects of Curtis' excellent text remain to influence these notes.
2. Damiano and Little's *A Course in Linear Algebra* published by Dover. I chose this as the required text in Spring 2015 as it is a well-written book, inexpensive and has solutions in the back to many exercises. The notation is fairly close to the notation used in these notes. I also liked the appearance of some diagrammatics for understanding Jordan forms. The section on minimal and characteristic polynomials is lucid.
3. Berberian's *Linear Algebra* published by Dover. This book is a joy. The exercises are challenging for this level and there were no solutions in the back of the text. This book is full of things I would like to cover, but, don't quite have time to do.
4. Takahashi and Inoue's *The Manga Guide to Linear Algebra*. Hillarious. Fun. Probably a better algorithm for Gaussian elimination than is given in my notes.
5. Hefferon's *Linear Algebra*: this text has nice gentle introductions to many topics as well as an appendix on proof techniques. The emphasis is linear algebra and the matrix topics are delayed to a later part of the text. Furthermore, the term linear transformation as supplanted by homomorphism and there are a few other, in my view, non-standard terminologies. All in all, very strong, but we treat matrix topics much earlier in these notes. Many theorems in this set of notes were inspired from Hefferon's excellent text. Also, it should be noted the solution manual to Hefferon, like the text, is freely available as a pdf.
6. Anton and Rorres' *Linear Algebra: Applications Version* or Lay's *Linear Algebra*, or Larson and Edwards *Linear Algebra*, or... standard linear algebra text. Written with non-math majors in mind. Many theorems in my notes borrowed from these texts.
7. Insel, Spence and Friedberg's *Elementary Linear Algebra*. This text is a little light on applications in comparison to similar texts, however, the theory of Gaussian elimination and other basic algorithms are extremely clear. This text focus on column vectors for the most part.
8. Insel, Spence and Friedberg's *Linear Algebra*. It begins with the definition of a vector space essentially. Then all the basic and important theorems are given. Theory is well presented in this text and it has been invaluable to me as I've studied the theory of adjoints, the problem of simultaneous diagonalization and of course the Jordan and rational canonical forms.
9. Strang's *Linear Algebra*. If geometric intuition is what you seek and/or are energized by then you should read this in parallel to these notes. This text introduces the dot product early on and gives geometric proofs where most others use an algebraic approach. We'll take the algebraic approach whenever possible in this course. We relegate geometry to the place of motivational side comments. This is due to the lack of prerequisite geometry on the part of a significant portion of the students who use these notes.

10. my advanced calculus notes. I review linear algebra and discuss multilinear algebra in some depth. I've heard from some students that they understood linear in much greater depth after the experience of my notes. Ask if interested, I'm always editing these.
11. Olver and Shakiban *Applied Linear Algebra*. For serious applications and an introduction to modeling this text is excellent for an engineering, science or applied math student. This book is somewhat advanced, but not as sophisticated as those further down this list.
12. Sadun's *Applied Linear Algebra: The Decoupling Principle* this is a second book in linear algebra. It presents much of the theory in terms of a unifying theme; decoupling. Probably this book is very useful to the student who wishes deeper understanding of linear system theory. Includes some Fourier analysis as well as a Chapter on Green's functions.
13. Curtis' *Abstract Linear Algebra*. Great supplement for a clean presentation of theorems. Written for math students without apology. His treatment of the wedge product as an abstract algebraic system is fun to read.
14. Roman's *Advanced Linear Algebra*. Treats all the usual topics as well as the generalization to modules. Some infinite dimensional topics are discussed. This has excellent insight into topics beyond this course.
15. Dummit and Foote *Abstract Algebra*. Part III contains a good introduction to the theory of modules. A module is roughly speaking a vector space over a ring. I believe many graduate programs include this material in their core algebra sequence. If you are interested in going to math graduate school, studying this book puts you ahead of the game a bit. Understanding Dummit and Foote by graduation is a nontrivial, but worthwhile, goal.
16. Golub and Van Loan's *Matrix Computations* 3rd edition. One of my students went to graduate school and reported back about this masterpiece of advanced matrix manipulations. Wonder what is beyond what I cover? Here's a good place to start if you prefer matrices to transformations.
17. Brown's *Matrices over Commutative Rings* takes a long look at what happens when we fill matrices with entries from a commutative ring. Interesting modifications of the theorems we covered this semester are seen here. Probably need Math 421 before reading this.
18. My Applied Linear Algebra notes from the Spring 2024 Semester. Those notes do have much in common with my old Math 321 notes as well. At the moment some of my lectures in Math 221 are not from these notes, but rather from the text by Lay itself. There are a few topics like the QR-decomposition, the Singular Value Decomposition which I have not yet edited. Eventually I want to have a set of notes which has part I for Math 221 then part II for Math 321.
19. My Linear Algebra notes from say 2019 which were for this course before we had the Math 221 prerequisite. The current version marks a major revision for the Fall 2024 Semester. If you want to read more background material and additional examples of Math 221 calculations then these notes are a good option.

notation in these notes

- Definitions given for an arbitrary field $\mathbb{F}$. Typical examples include complex numbers $\mathbb{C}$, real numbers $\mathbb{R}$, rational numbers $\mathbb{Q}$ or a finite field such as $\mathbb{Z}_2$ or $\mathbb{Z}_3$.
- If S is a set then $|S|$ denotes the **cardinality** of S . For example, $|\mathbb{Z}_3| = 3$ whereas $|\mathbb{Z}_2| = 2$. We denote $|\mathbb{N}| = \aleph_0$ whereas $|\mathbb{R}| = \mathfrak{c}$.
- For a function $f : A \rightarrow B$ and $U \subseteq A$ and $V \subseteq B$ then the set $f(U) = \{f(x) \mid x \in U\}$ is the **image of f under U** and the set $f^{-1}(V) = \{x \in A \mid f(x) \in V\}$ is the **inverse image of V under f** . We say f is a **surjection** if $f(A) = B$ and we say f is an **injection** if every nonempty inverse image of a singleton is a singleton. Equivalently, f is **injective** if $x, y \in A$ with $f(x) = f(y)$ implies $x = y$. Notice, the notation f^{-1} does not require that f has an inverse as a function.
- $R^{m \times n}$ is the set of $m \times n$ matrices with entries in R .
- If $x \in R^n = R^{n \times 1}$ then $x = (x_1, x_2, \dots, x_n) = \begin{bmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{bmatrix}$ is a **column vector**.
- If A is an $m \times n$ matrix and $1 \leq i \leq m$ and $1 \leq j \leq n$ then A_{ij} is the **component** of A in the i -th row and j -th column. Likewise, $col_j(A)$ is the j -th **column** of A and $row_i(A)$ is the i -th **row**. We write $A = [col_1(A) \mid \dots \mid col_n(A)]$ to denote the fact that A is the concatenation of all its columns.
- Given a matrix A the **reduced row echelon form** or **rref** of A is the matrix produced by applying the Gauss-Jordan row reduction to the matrix A . We denote this by $rref(A)$. The Column Correspondence Property or **CCP** is an important theorem connected to interpreting the $rref(A)$ and its various applications, see Proposition 1.7.6.
- $\delta_{ij} = 1$ if $i = j$ and $\delta_{ij} = 0$ if $i \neq j$. We call δ_{ij} the **Kronecker delta**.
- I or $I_n \in \mathbb{F}^{n \times n}$ is the **identity matrix** which has $I_{ij} = \delta_{ij}$.
- The column vector $e_i \in \mathbb{F}^n$ defined by $(e_i)_j = \delta_{ij}$ is the i -th **standard basis** element.
- The matrix $E_{ij} \in \mathbb{F}^{m \times n}$ defined by $(E_{ij})_{kl} = \delta_{ik}\delta_{jl}$ is the (i, j) -th **matrix unit**.
- If $\mathbb{F}$ is a field and V is a set then $V(\mathbb{F})$ denotes the vector space V with field of scalars $\mathbb{F}$.
- V is a vector space and writing $U \subseteq V$ means U is a **subset** of V whereas $U \leq V$ means U is a **subspace** of V . These are not the same.
- If S is a subset of a vector space V then $span(S) \subseteq V$ is the set of all finite linear combinations of vectors in S . We sometimes use $span_{\mathbb{F}}(S)$ to emphasize that the coefficients for the linear combinations are taken from $\mathbb{F}$.
- We use **LI** as an abbreviation for **Linear Independence**.
- If β is basis for $V(\mathbb{F})$ and $x \in V$ then $[x]_{\beta}$ is the **coordinate vector** of x with respect to β . Likewise, $\Phi_{\beta} : V \rightarrow \mathbb{F}^n$ is the **coordinate map** defined by $\Phi_{\beta}(x) = [x]_{\beta}$.

- We use $[T]_{\alpha,\beta}$ for the matrix of $T : V \rightarrow W$ where $V = \text{span}(\alpha)$ and $W = \text{span}(\beta)$.
- $P_n(\mathbb{F})$ is the set of up to n -th order polynomials with coefficients in the field $\mathbb{F}$. Likewise, $\mathbb{F}[t]$ is the set of polynomials with coefficients in $\mathbb{F}$ where we are using the variable t to express the polynomials. Likewise, $\mathbb{F}[x, y] = (\mathbb{F}[x])[y]$ is the set of bivariate polynomials with coefficients in $\mathbb{F}$ and $\mathbb{F}[x_1, \dots, x_n]$ is the set of n -multivariate polynomials with coefficients in $\mathbb{F}$.
- The **null space** of a matrix A is given by $\text{Null}(A) = \{x \in \mathbb{F}^n \mid Ax = 0\}$ whereas the **column space** of A is denoted $\text{Col}(A)$ which is the span of the columns of A . Likewise, the **nullity** of A is the dimension of $\text{Null}(A)$ whereas the **rank** of A is the dimension of $\text{Col}(A)$.
- Given vector spaces V, W the set of all **linear transformations** is denoted $\mathcal{L}(V, W)$. If $T \in \mathcal{L}(V, V) = \mathcal{L}(V)$ then we say T is a **linear transformation** on V or an **endomorphism** of V . The set of all functions from a set A to a set B is denoted $\mathcal{F}(A, B)$.
- Given $A \in \mathbb{F}^{m \times n}$ the **left multiplication** map $L_A : \mathbb{F}^n \rightarrow \mathbb{F}^m$ is defined by $L_A(x) = Ax$ for each $x \in \mathbb{F}^n$.
- For linear map $T : V \rightarrow W$, the **kernel** of T is $\text{Ker}(T) = \{x \in V \mid T(x) = 0\}$ and the **range** or **image** of T is $\text{Range}(T) = \{T(x) \mid x \in V\}$.
- The map $\text{Id}_V : V \rightarrow V$ is known as the **identity map** on V ; $\text{Id}_V(x) = x$ for all $x \in V$.
- Given appropriate linear maps T, S we call $T \circ S$ the **composition** of T with S . If T is a linear map on V then $T^0 = \text{Id}_V$ and $T^1 = T$ whereas $T^k = T \circ T^{k-1}$ for all $k \in \mathbb{N}$ with $k \geq 2$.
- If $W_1, W_2 \leq V$ and $W_1 \cap W_2 = 0$ then $W_1 \oplus W_2$ is the **direct sum** or **internal direct sum** of W_1 and W_2 . Given two vector spaces V_1, V_2 we call the Cartesian product $V_1 \times V_2$ the **external direct sum** or **direct sum** of V_1 and V_2 . Similarly, $W_1 \oplus \dots \oplus W_k$ is the **direct sum** of the **independent subspaces** $W_1, \dots, W_k$. See Theorem 3.6.4 for the five ways to characterize independent subspaces. Notice, §6.4 is an example of an external direct sum.
- For subspace W of V the set $x + W = \{x + w \mid w \in W\}$ is the **coset** of W with representative x . We also say $x + W$ is a **linear manifold** of V .
- The **quotient space** V/W is the set of all cosets of W .
- If vector spaces V and W are **isomorphic** then we write $V \cong W$.
- Let $T : V \rightarrow V$ be a linear map and $W \leq V$. We say $T|_W : W \rightarrow V$ is the **restriction** of T to W given by $T|_W(x) = T(x)$ for each $x \in W$. If $T(W) \leq W$ then W is an **invariant subspace** of T and we define $T_W : W \rightarrow W$ by $T_W(x) = T(x)$ for each $x \in W$.
- The **dual space** of a vector space $V(\mathbb{F})$ is given by $V^* = \mathcal{L}(V, \mathbb{F})$. Likewise, the **double dual** of V is V^{**} is the set of all linear functions from V^* to $\mathbb{F}$.
- The **annihilator** of $W \leq V$ is $\text{ann}(W) = \{\alpha \in V^* \mid \alpha(W) = 0\}$.
- The **transpose** of $A \in \mathbb{F}^{m \times n}$ is $A^T \in \mathbb{F}^{n \times m}$ where $(A^T)_{ij} = A_{ji}$ for all i, j . The **transpose** of $T \in \mathcal{L}(V, W)$ is $L^t \in \mathcal{W}^*, \mathcal{V}^*$.
- Matrices A, B are **congruent** if there exist invertible matrices P, Q for which $B = PAQ$.
- Matrices A, B are **similar** if there exists an invertible matrix P for which $B = P^{-1}AP$.

- $\Lambda_0 V = \mathbb{F}$, $\Lambda_1 V = V$ and generally $\Lambda_k V$ consists of sums of k -fold wedge products of vectors. We denote the **wedge product** by $\wedge$.
- The **determinant** of A is the scalar for which $Ae_1 \wedge \cdots \wedge Ae_n = \det(A)e_1 \wedge \cdots \wedge e_n$.
- Given a matrix A or an endomorphism T of a finite dimensional vector space $V(\mathbb{F})$ we define the **characteristic polynomial** $P(t) \in \mathbb{F}[t]$ by $P(t) = \det(A - tI_n)$ or $P(t) = \det(T - tId_V)$. The **characteristic equation** is $P(t) = 0$.
- Given T a linear map on $V(\mathbb{F})$, if there exists $x \neq 0$ and $\lambda \in \mathbb{F}$ for which $T(x) = \lambda x$ then x is an **eigenvector** with **eigenvalue** λ .
- If λ is an eigenvalue of $T \in \mathcal{L}(V(\mathbb{F}))$ then $\mathcal{E}_\lambda = \text{Ker}(T - \lambda Id_V)$ is the λ -**eigenspace** of T . If λ is an eigenvalue of $A \in \mathbb{F}^{n \times n}$ then $\mathcal{E}_\lambda = \text{Null}(A - \lambda I)$ is the λ -**eigenspace** of A .
- A **generalized eigenvector** of order k with eigenvalue λ with respect to a linear transformation $T : V \rightarrow V$ is a nonzero vector v such that

$$(T - \lambda Id)^k v = 0 \quad \& \quad (T - \lambda Id)^{k-1} v \neq 0.$$

- A **k -chain with eigenvalue** λ of a linear transformation $T : V \rightarrow V$ is set of k nonzero vectors $v_1, v_2, \dots, v_k$ such that $(T - \lambda Id)(v_j) = v_{j-1}$ for $j = 2, \dots, k$ and v_1 is an eigenvector with eigenvalue λ ; $(T - \lambda Id)(v_1) = 0$.
- Let $N = E_{12} + E_{23} + \cdots + E_{d-1,d} \in \mathbb{F}^{d \times d}$ be the matrix which is everywhere zero except where it is one on its superdiagonal. We define the $d \times d$ -Jordan block by $J_d(\lambda) = \lambda_d I + N$. A matrix $J \in \mathbb{F}^{n \times n}$ is said to be in **Jordan Form** if it is a block-diagonal matrix with Jordan blocks on the diagonal; $J = J_{d_1}(\lambda_1) \oplus \cdots \oplus J_{d_k}(\lambda_k)$.
- Let $T \in \mathcal{L}(V)$ and $x \in V$ then the **T -cyclic subspace** generated from x is $\langle x \rangle = \text{span}\{T^k(x) \mid k \in \mathbb{N} \cup \{0\}\}$
- A **generalized eigenspace** of eigenvalue λ for a linear transformation $T : V \rightarrow V$ is denoted K_λ . We define $x \in K_\lambda$ if there exists a positive integer k such that

$$(T - \lambda)^k x = 0.$$

- Suppose V is a vector space over $\mathbb{R}$, then $V_{\mathbb{C}}$ is the **complexification** of V . Likewise, if $T : V \rightarrow V$ is a linear then its **complexification** is $T_{\mathbb{C}} : V_{\mathbb{C}} \rightarrow V_{\mathbb{C}}$ which is defined by $T_{\mathbb{C}}(x + iy) = T(x) + iT(y)$ for all $x + iy \in V_{\mathbb{C}}$. A **complex eigenvector** for T is an eigenvector for $T_{\mathbb{C}}$ whose eigenvalue $\lambda \in \mathbb{C} - \mathbb{R}$.
- The matrices below are known as **real Jordan blocks**. Suppose $\alpha + i\beta \in \mathbb{C}$ with $\beta \neq 0$ then

$$R_2(\alpha + i\beta) = \begin{bmatrix} \alpha & \beta \\ -\beta & \alpha \end{bmatrix} \quad \& \quad R_4(\alpha + i\beta) = \begin{bmatrix} \alpha & \beta & 1 & 0 \\ -\beta & \alpha & 0 & 1 \\ 0 & 0 & \alpha & \beta \\ 0 & 0 & -\beta & \alpha \end{bmatrix}$$

Generally,

$$R_{2k}(\alpha + i\beta) = R_2(\alpha + i\beta) \otimes I_k + I_k \otimes N_k.$$

If A is the block-diagonal with each block either of the form $R_{2k}(\lambda)$ or $J_k(\lambda)$ then A is said to be in **real Jordan form**.

- e^A is the **matrix exponential** of A given by $e^A = I + A + \frac{1}{2}A^2 + \frac{1}{3!}A^3 + \dots$.
- If $x, y \in V$ then $\langle x, y \rangle$ is the **inner-product** of x and y .
- If $x \in V$ then $\|x\|$ is the **norm** or **length** of x .
- A set of vectors in an inner product space is **orthogonal** if $\langle x, y \rangle = 0$ for any distinct pair x, y in the set. A **unit-vector** is a vector of length one. A set of orthogonal unit-vectors is an **orthonormal** set.
- Given an inner product space V with inner product $\langle \cdot, \cdot \rangle : V \times V \rightarrow \mathbb{F}$ and $S \subseteq V$ we define the **orthogonal complement** or **perp** of S by $S^\perp = \{x \in V \mid \langle x, y \rangle = 0 \text{ for all } y \in S\}$

This image inspired by my time with Benjamin. We have been reviewing the adventures of KITT lately and that prompted this art:



Contents

1	Matrices and linear systems	1
1.1	matrices	1
1.2	matrix addition and scalar multiplication	5
1.2.1	standard column and matrix bases	7
1.3	matrix multiplication	10
1.3.1	multiplication of row or column concatenations	13
1.3.2	all your base are belong to us (e_i and E_{ij} that is)	16
1.4	matrix algebra	18
1.4.1	identity and inverse matrices	18
1.4.2	matrix powers	22
1.4.3	symmetric and antisymmetric matrices	23
1.4.4	triangular and diagonal matrices	24
1.4.5	nilpotent matrices	25
1.4.6	block matrices	26
1.5	systems of linear equations	29
1.5.1	superposition of solutions	36
1.6	elementary matrices	37
1.6.1	definition of elementary matrices	37
1.6.2	superaugmented matrices and inverse matrix calculation	39
1.7	spans of column vectors and the CCP	41
1.7.1	solving several spanning questions simultaneously	43
1.7.2	CCP and linear dependence of column vectors	43
2	Vector Spaces	47
2.1	definition and examples	48
2.2	subspaces	54
2.3	spanning sets and subspaces	59
2.4	linear independence	62
2.5	theory of dimension	65
2.6	coordinate mappings	69
2.6.1	how to calculate a basis for a span of row or column vectors	72
2.6.2	calculating the basis of the null space of a matrix	75
2.6.3	coordinates for an infinite dimensional vector space	77
2.6.4	trace based calculation of dimension	78
2.7	theory of subspaces	80

3	linear transformations	83
3.1	definition, examples and basic theory	84
3.1.1	examples of linear transformations	84
3.1.2	basic theory of linear transformations	86
3.1.3	standard matrices and their properties	91
3.2	restriction, extension, isomorphism	93
3.2.1	examples of isomorphisms	96
3.3	matrix of linear transformation	99
3.4	coordinate change	106
3.4.1	coordinate change of abstract vectors	106
3.4.2	coordinate change for column vectors	108
3.4.3	coordinate change of abstract linear transformations	110
3.4.4	coordinate change of linear transformations of column vectors	112
3.5	theory of dimensions for maps	113
3.5.1	a detour into matrix theory	116
3.6	structure of subspaces	118
3.6.1	independent subspaces	118
3.6.2	invariant subspaces and internal direct products	123
3.7	quotient vector space	125
3.7.1	the first isomorphism theorem	129
3.7.2	on direct sums and quotients	132
3.8	dual space	134
3.8.1	transpose of linear map and the up-down notation	137
4	Jordan Form	139
4.1	eigenvectors and diagonalization	140
4.2	eigenbases and eigenspaces	146
4.3	Invariant Subspaces and the Cayley Hamilton Theorem	149
4.4	Theory of the Jordan Form	152
4.5	generalized eigenvectors and Jordan blocks	153
4.6	complexification	160
4.6.1	concerning matrices and vectors with complex entries	160
4.6.2	the complexification	161
4.6.3	complexification for 2nd order constant coefficient problem	162
4.7	complex eigenvalues and vectors for real linear maps	163
4.8	examples of real and complex eigenvectors	166
4.8.1	characteristic equations	166
4.8.2	real eigenvector examples	167
4.8.3	complex eigenvector examples	171
4.9	real Jordan form	173
4.10	systems of differential equations	177
4.10.1	calculus of matrices	178
4.10.2	the normal form and theory for systems	179
4.10.3	solutions by eigenvector	184
4.10.4	solutions by matrix exponential	193

5	Linear Algebra with Geometry	205
5.1	analytic geometry for vector spaces	205
5.2	orthogonality	214
5.3	theory of adjoints	221
5.4	diagonalization in inner product spaces	225
6	Multilinear Algebra	231
6.1	multilinearity and the tensor product	232
6.1.1	bilinear maps	232
6.1.2	trilinear maps	236
6.1.3	multilinear maps	238
6.2	wedge product	241
6.2.1	wedge product of dual basis generates basis for ΛV	241
6.2.2	the exterior algebra	243
6.2.3	connecting vectors and forms in $\mathbb{R}^3$	249
6.3	an introduction to algebra	249
6.4	wedge product and determinants	251
6.4.1	exterior algebra over $\mathbb{F}^2$	252
6.4.2	exterior algebra over $\mathbb{F}^3$	252
6.4.3	the exterior algebra over a vector space	253
6.4.4	definition of determinant	253
6.4.5	wedge product proofs for determinants	255
6.5	linear dependence and the wedge product	257
6.6	bilinear forms and geometry, metric duality	259
6.6.1	metric geometry	259
7	Appendix on Modular Arithmetic	261
7.1	$\mathbb{Z}$ -Basics	261
7.1.1	division algorithm	262
7.1.2	divisibility in $\mathbb{Z}$	264
7.2	On modular arithmetic and groups	269
7.3	matrices of modular integers	275
8	Appendix on Sets and Functions	277
8.1	set theory	277
8.2	functions	279

Chapter 1

Matrices and linear systems

The fear of the LORD is the beginning of knowledge:
but fools despise wisdom and instruction.

PROVERBS 1:7, KJV.

The calculations in this Chapter should be largely a review from your first course in Linear Algebra. On the other hand, I do intend to test on the definitions of matrix addition, multiplication, scalar multiplication and transposition. We give precise definitions which allow us to make arguments for matrices of arbitrary size without using explicit listing of elements. I do hope you will learn the definitions for the standard basis, matrix units as well as all the basic matrix operations. These make for nice entry-level proof questions.

In addition, we seek to review the technique of row-reduction for the solution of linear systems of equations. I introduce explicit definitions of elementary matrices and we work out their structure via the matrix algebra introduced earlier in this Chapter. The general structure of solution sets is described both for infinite and finite fields. We also study the question of spanning for column vectors. We show how the Column Correspondence Property (CCP) gives an efficient and elegant method to understand the meaning of row-reductions. This CCP is special to the context of column vectors. We will soon see that corresponding questions of spanning in an abstract vector space require a bit more calculation.

It is unlikely I cover all of this material in lecture. I do hope you will review it. Certainly I do assume students of Math 321 already have a complete understanding of how to find solution sets over $\mathbb{R}$ or $\mathbb{C}$. Systems of equations and matrix algebra over finite fields is new and I probably should offer some homework to help us all understand the quirks of finite field arithmetic.

1.1 matrices

An array of objects is a collection of objects where we can keep track of which row and column each object resides. A finite sequence has the form $\{a_1, a_2, \dots, a_n\}$. There is a bijective correspondence between finite sequences¹ in a set S and functions from $\mathbb{N}_n$ to S . In particular, given $\{a_1, a_2, \dots, a_n\}$ we define $a(j) = a_j$ for each $j \in \mathbb{N}_n$. Likewise, if $a : \mathbb{N}_n \rightarrow S$ is a function then $\{a(1), a(2), \dots, a(n)\}$ is a finite, ordered list in S ; that is, a finite sequence in S . An $m \times n$ **array** of objects in S is likewise in bijective correspondence between functions from $\mathbb{N}_m \times \mathbb{N}_n$ to S . In particular, given

¹you studied the more subtle topic of infinite sequences in second semester calculus, there the sequences are functions from the positive integers to the real numbers typically

$a : \mathbb{N}_m \times \mathbb{N}_n \rightarrow S$ we may construct an array as follows:

$$\begin{bmatrix} a(1,1) & a(1,2) & \cdots & a(1,n) \\ a(2,1) & a(2,2) & \cdots & a(2,n) \\ \vdots & \vdots & \cdots & \vdots \\ a(m,1) & a(m,2) & \cdots & a(m,n) \end{bmatrix}$$

For the foundationalist, you might wonder how we construct an array. There are various answers to this question. One, we could adopt the viewpoint that an $m \times n$ array is simply a notation for a function from $\mathbb{N}_m \times \mathbb{N}_n$. Alternatively, you could view an array as a vector of vectors. The second view is sometimes used as a basis for the syntax used to manipulate matrices in Computer Algebra Systems² (CASs). Anyway, the construction of arrays from basic principles is really just something we assume in this course so I've probably already said too much about this substructure. Definition 1.1.3 captures the essential feature of an array we wish to exploit.

Definition 1.1.1.

An $m \times n$ matrix is an array of objects with m rows and n columns. The elements in the array are called **entries** or **components**. If A is an $m \times n$ matrix then A_{ij} denotes the object in the i -th row and the j -th column. We denote:

$$A = [A_{ij}] = \begin{bmatrix} A_{11} & A_{12} & \cdots & A_{1n} \\ A_{21} & A_{22} & \cdots & A_{2n} \\ \vdots & \vdots & \cdots & \vdots \\ A_{m1} & A_{m2} & \cdots & A_{mn} \end{bmatrix}$$

The label i is a **row index** and $1 \leq i \leq m$. The index j is a **column index** and $1 \leq j \leq n$.

$$\text{row}_i(A) = [A_{i1}, \dots, A_{in}] \quad \& \quad \text{col}_j(A) = \begin{bmatrix} A_{1j} \\ \vdots \\ A_{mj} \end{bmatrix}$$

Generally, if S is a set then $S^{m \times n}$ is the set of $m \times n$ arrays of objects from S . If a matrix has the same number of rows and columns then it is called a **square matrix**.

The set $m \times n$ of matrices with real number entries is denoted $\mathbb{R}^{m \times n}$. The set of $m \times n$ matrices with complex entries is $\mathbb{C}^{m \times n}$. An $m \times n$ matrix can be seen as a **concatenation** of rows or columns:

$$A = \begin{bmatrix} A_{11} & \cdots & A_{1n} \\ \vdots & \cdots & \vdots \\ A_{m1} & \cdots & A_{mn} \end{bmatrix} = [\text{col}_1(A) | \cdots | \text{col}_n(A)] = \begin{bmatrix} \text{row}_1(A) \\ \vdots \\ \text{row}_m(A) \end{bmatrix}$$

To concatenate matrices is to join them together to make a larger matrix. The horizontal and vertical lines simply point to where the matrices have been glued together.

It is important to distinguish between the matrix A and the³ i, j -th component A_{ij} . It is simply **not true** that $A = A_{ij}$. However, $A = [A_{ij}]$ as the brackets denote the array of all components.

²not to be confused with the ever more interesting ROUSs

³It is understood in this course that $i, j, k, l, m, n, p, q, r, s$ are in $\mathbb{N}$. I will not belabor this point. Please ask if in doubt.

Example 1.1.2. Let $A = \begin{bmatrix} 1 & 2 & 3 \\ 4 & 5 & 6 \\ 7 & 8 & 9 \end{bmatrix}$ and $B = \begin{bmatrix} i & 10 \\ 0 & 3+i \\ 11 & 12 \end{bmatrix}$ then $A \in \mathbb{R}^{3 \times 3}$ and $B \in \mathbb{C}^{3 \times 2}$. If $M = [A|B]$ and $N = [B|A]$ then $M, N \in \mathbb{C}^{3 \times 5}$. Notice, $M_{25} = 3 + i$ whereas $N_{25} = 6$.

In the example above we observed that the 2, 5 components of M and N differ. It follows $M \neq N$. Let us pause to remind what is required for two arrays to be identical:

Definition 1.1.3.

If $A, B \in S^{m \times n}$ then $A = B$ if and only if $A_{ij} = B_{ij}$ for all $1 \leq i \leq m$ and $1 \leq j \leq n$.

It is simple to see that $A = B$ is equivalent⁴ to $\text{row}_i(A) = \text{row}_i(B)$ for all $i \in \mathbb{N}_m$. Likewise, $A = B$ is equivalent to $\text{col}_j(A) = \text{col}_j(B)$ for all $j \in \mathbb{N}_n$. We can measure the verity of a matrix equation at the level of components, rows, or columns. We use all three notations throughout this study.

Example 1.1.4. Matrices may contain things other than numbers. For instance, if $f, g, h : \mathbb{R} \rightarrow \mathbb{R}$ are functions then $A = \begin{bmatrix} f & g \\ h & f \end{bmatrix}$ is a matrix of functions.

Our typical examples involve matrices with numbers as components. What is a *number*? That question is more philosophical than mathematical. I generally think of a number as an object which I can add, subtract and multiply.

Remark 1.1.5. *an overview of abstract algebraic terminology*

A **group** is a set paired with an operation which is associative, unital and is closed under inverses. If the operation of the group is commutative then the group is said to be **abelian**. A set R is called a **ring** if it has a pair of operations known as addition and multiplication. In particular, it is assumed that R paired with addition forms an abelian group. If R has a **unity** it is also assumed there exists $1 \in R$ for which $1x = x$ for each $x \in R$. Finally, multiplication of the ring must satisfy the following: for all $a, b, c, x, y \in R$,

$$a(x + y) = ax + ay \quad \& \quad (a + b)x = ax + bx.$$

In short, a ring is a place where you can do arithmetic as we usually practice. If $ab = ba$ for all $a, b \in R$ then R is a **commutative ring**. The study of commutative rings occupies a large part of the abstract algebra sequence at many universities. Even commutative rings are a bit more perilous than you might expect. For example, there are rings for which $ab = ac$ with $a \neq 0$ does not imply $b = c$. For example, $\mathbb{Z}/4\mathbb{Z} = \{\bar{0}, \bar{1}, \bar{2}, \bar{3}\}$ has $\bar{2}\bar{2} = \bar{0}$ and $\bar{2}\bar{0} = \bar{0}$. The number $\bar{2}$ is a **zero-divisor** in $\mathbb{Z}/4\mathbb{Z}$. If $r \in R$ has $s \in R$ for which $rs = 1$ then r is said to be a **unit** with multiplicative inverse s usually denoted $r^{-1} = s$. **Zero-divisors** are nonzero elements $a, b \in R$ for which $ab = 0$. A commutative ring with no zero divisors is called an **integral domain**. The quintessential example of an integral domain is $\mathbb{Z}$. If every nonzero element of a commutative ring is a unit then we say that ring is a **field**. It is a fun exercise to prove that no unit is a zero divisor (use proof by contradiction). It follows that every field is an integral domain. In the finite case, the converse is also true. Every finite integral domain is a field. This is the sort of claim we will prove in the study of abstract algebra course sequence. I mention it here for your informational edification. I'll leave you with a few claims whose proof I leave to another course; $\mathbb{Q}, \mathbb{R}$ and $\mathbb{C}$ are fields. If p is a prime then $\mathbb{Z}/p\mathbb{Z}$ is also a field. The distinction between $\mathbb{Q}, \mathbb{R}, \mathbb{C}$ and the finite field $\mathbb{Z}/p\mathbb{Z}$ is quite clearly seen by the concept of **characteristic**. If $1 + 1 + \cdots + 1 \neq 0$ then the field has **characteristic zero**. In contrast, if we add p -fold copies of 1 in $\mathbb{Z}/p\mathbb{Z}$ then $1 + 1 + \cdots + 1 = p1 = 0$. We say the characteristic of $\mathbb{Z}/p\mathbb{Z}$ is p . However, if n is composite then $\mathbb{Z}/n\mathbb{Z}$ is not a field as all divisors of n produce zero divisors.

⁴two rows are equal iff the given pair of rows have the same components in the same order. Likewise, equality of columns is defined by equality of matching components.

Row and column matrices deserve further discussion. First, we need to define transposition⁵ of a matrix: once more let S be a set,

Definition 1.1.6.

If $A \in S^{m \times n}$ then $A^T \in S^{n \times m}$ is defined by $(A^T)_{ji} = A_{ij}$ for all $(i, j) \in \mathbb{N}_m \times \mathbb{N}_n$. We say A^T is the **transpose** of A .

If we think through this definition in terms of rows and columns we can identify that transposition converts columns to rows and rows to columns.

Example 1.1.7. Let $A = \begin{bmatrix} A_{11} & A_{12} \\ A_{21} & A_{22} \\ A_{31} & A_{32} \end{bmatrix}$ then $A^T = \begin{bmatrix} A_{11} & A_{21} & A_{31} \\ A_{12} & A_{22} & A_{32} \end{bmatrix}$. Observe

$$\text{row}_1(A^T) = [A_{11}, A_{21}, A_{31}] = \begin{bmatrix} A_{11} \\ A_{21} \\ A_{31} \end{bmatrix}^T = (\text{col}_1(A))^T.$$

Likewise, $\text{row}_2(A^T) = (\text{col}_2(A))^T$.

Proposition 1.1.8.

Let $A \in S^{m \times n}$ then (i.) $(A^T)^T = A$. Furthermore, for each $i \in \mathbb{N}_m$ and $j \in \mathbb{N}_n$,

$$(ii.) \text{row}_j(A^T) = (\text{col}_j(A))^T \quad \& \quad (iii.) \text{col}_i(A^T) = (\text{row}_i(A))^T.$$

Proof: To prove (i.) simply note that $((A^T)^T)_{ij} = (A^T)_{ji} = A_{ij}$ for all $(i, j) \in \mathbb{N}_m \times \mathbb{N}_n$. Notice, $(\text{col}_j(A))_i = A_{ij}$ whereas $(\text{row}_i(A))_j = A_{ij}$. Thus consider,

$$(\text{row}_j(A^T))_i = (A^T)_{ji} = A_{ij} = (\text{col}_j(A))_i = ((\text{col}_j(A))^T)_i$$

the last step is simply that the i -th component of a row vector is the i -th component of the transpose of the row vector. I leave the proof of (iii.) to the reader. $\square$

In principle one can use column vectors for everything or row vectors for everything. I choose a subtle **convention** that allows us to use both⁶

Definition 1.1.9. *hidden column notation.*

We denote $(v_1, v_2, \dots, v_n) = \begin{bmatrix} v_1 \\ v_2 \\ \vdots \\ v_n \end{bmatrix}$ and we write $S^n = S^{n \times 1}$.

If I want to denote a real row vector then we will just write $[v_1, v_2, \dots, v_n]$. This convention means we view points as column vectors. This is just a notational choice.

⁵some authors prefer the notation ${}^t A$ in the place of A^T

⁶On the one hand it is nice to write vectors as rows since the typesetting is easier. However, once you start talking about matrix multiplication then it is natural to write the vector to the right of the matrix and we will soon see that the vector should be written as a column vector for that to be reasonable.

1.2 matrix addition and scalar multiplication

Often we only need to consider numbers in fields, but, I'll prove a few things just assuming a ring structure for the elements. This generality costs us nothing and helps the reader get in the habit of thinking abstractly. Let us collect the essential algebraic features of a commutative ring R with identity. There is an additive identity $0 \in R$ such that $x + 0 = x$ for each $x \in R$. There is also a multiplicative identity $1 \in R$ such that $1x = x$ for each $x \in R$. For each $x \in R$ there exists $-x \in R$ such that $x + (-x) = 0$. We also have $(-1)(x) = -x$ and $0(x) = 0$ for each $x \in R$. Furthermore, for each for $a, b, c \in R$,

- (1.) associativity of addition: $a + (b + c) = (a + b) + c$,
- (2.) commutativity of addition: $a + b = b + a$,
- (3.) left and right distributivity: $a(b + c) = ab + ac$ and $(a + b)c = ac + bc$,
- (4.) associativity of multiplication; $a(bc) = (ab)c$

Some mathematicians include the existence of $1 \in R$ as part of the definition of a ring, but, I do not assume that here or in the abstract algebra courses. That said, in the remainder of this section **please assume R is a commutative ring with identity 1**. Pragmatically, this means I allow $R = \mathbb{R}, \mathbb{Q}, \mathbb{C}, \mathbb{Z}, \mathbb{Z}/n\mathbb{Z}$ for $n \geq 2$. However, we could also have R be the set of continuous functions on $\mathbb{R}$. There are many sets of things which have the structure of a commutative ring with identity. In all those many cases we can construct a matrix of such objects. The arguments in this Chapter demonstrate that a common algebra is shared for this multitude of examples. Abstraction allows us to carry many loads at once. We cut away all the irrelevant features of R and focus just on the arithmetic properties above. Those suffice to develop the matrix algebra.

Definition 1.2.1. *Let R be a commutative ring with $1 \in R$.*

If $A, B \in R^{m \times n}$ then the **sum** of A, B is $A + B$ and the **scalar multiple** of A by c is cA . these are defined as follows:

$$(A + B)_{ij} = A_{ij} + B_{ij} \quad \& \quad (cA)_{ij} = cA_{ij}$$

for all $(i, j) \in \mathbb{N}_m \times \mathbb{N}_n$.

In the special case of row or column vectors we understand the Definition above to reduce to:

$$(x + y)_i = x_i + y_i \quad \& \quad (cx)_i = cx_i$$

for $c \in R$ and $x, y \in R^n = R^{n \times 1}$ or for $x, y \in R^{1 \times n}$.

Let us pause to consider a few computational examples⁷

Example 1.2.2. $\begin{bmatrix} 1 & 2 \\ 3 & 4 \end{bmatrix} + \begin{bmatrix} 5 & 6 \\ 7 & 8 \end{bmatrix} = \begin{bmatrix} 6 & 8 \\ 10 & 12 \end{bmatrix}$.

Example 1.2.3. Let $A = \begin{bmatrix} -2 & 4 \\ 10 & 4 \end{bmatrix}$ and $B = \begin{bmatrix} x & x^2 \\ 7+y & 2z^2 \end{bmatrix}$. Can we solve $A = B$? Notice, the equality $A = B$ gives four equations we must solve concurrently:

$$-2 = x, \quad 4 = x^2, \quad 10 = 7 + y, \quad 4 = 2z^2.$$

We find two solutions $x = -2, y = 3$ and $z = \pm\sqrt{2}$

⁷you might take a moment to notice my examples tend to fit into one of the following three types: question and answer, discussion-discovery or show-case a theorem or definition. If you're looking to identify the problem type so you can solve it when I ask you it again, you need to adjust your thinking here...

Example 1.2.4. Let $A, B \in \mathbb{R}^{m \times n}$ be defined by $A_{ij} = 3i + 5j$ and $B_{ij} = i^2$ for all i, j . Then we can calculate $(A + 7B)_{ij} = 3i + 5j + 7i^2$ for all i, j .

Example 1.2.5. Over $R = \mathbb{Z}/6\mathbb{Z}$ we calculate

$$\bar{3} \begin{bmatrix} \bar{1} & \bar{2} \\ \bar{3} & \bar{4} \end{bmatrix} = \begin{bmatrix} \bar{3} \cdot \bar{1} & \bar{3} \cdot \bar{2} \\ \bar{3} \cdot \bar{3} & \bar{3} \cdot \bar{4} \end{bmatrix} = \begin{bmatrix} \bar{3} & \bar{6} \\ \bar{9} & \bar{12} \end{bmatrix} = \begin{bmatrix} \bar{3} & \bar{0} \\ \bar{3} & \bar{0} \end{bmatrix}$$

Definition 1.2.1 says we define matrix addition and scalar multiplication **component-wise**. We show next how index notation provides us an elegant formalism to easily prove facts about the algebra of matrices. Notice how each arithmetic property rings induces a similar matrix property:

Proposition 1.2.6. *linearity of matrix addition and scalar multiplication*

Let R be a commutative ring with unity. If $A, B, C \in R^{m \times n}$ and $c_1, c_2 \in R$ then

(1.) $(A + B) + C = A + (B + C),$

(2.) $A + B = B + A,$

(3.) $c_1(A + B) = c_1A + c_2B,$

(4.) $(c_1 + c_2)A = c_1A + c_2A,$

(5.) $(c_1c_2)A = c_1(c_2A),$

(6.) $1A = A,$

Proof: Nearly all of these properties are proved by breaking the statement down to components then appealing to a ring property. I supply proofs of (1.) and (5.) and leave (2.), (3.), (4.) and (6.) to the reader.

Proof of (1.): assume A, B, C are given as in the statement of the Theorem. Observe that

$$\begin{aligned} ((A + B) + C)_{ij} &= (A + B)_{ij} + C_{ij} && \text{defn. of matrix add.} \\ &= (A_{ij} + B_{ij}) + C_{ij} && \text{defn. of matrix add.} \\ &= A_{ij} + (B_{ij} + C_{ij}) && \text{assoc. of ring addition} \\ &= A_{ij} + (B + C)_{ij} && \text{defn. of matrix add.} \\ &= (A + (B + C))_{ij} && \text{defn. of matrix add.} \end{aligned}$$

for all i, j . Therefore $(A + B) + C = A + (B + C)$. $\square$

Proof of (5.): assume c_1, c_2, A are given as in the statement of the Theorem. Observe that

$$\begin{aligned} ((c_1c_2)A)_{ij} &= (c_1c_2)A_{ij} && \text{defn. scalar multiplication.} \\ &= c_1(c_2A_{ij}) && \text{assoc. of ring multiplication} \\ &= (c_1(c_2A))_{ij} && \text{defn. scalar multiplication.} \end{aligned}$$

for all i, j . Therefore $(c_1c_2)A = c_1(c_2A)$. $\square$

The proofs of the other items are similar, we consider the i, j -th component of the identity and then apply the definition of the appropriate matrix operation's definition. This reduces the problem to

a statement about arithmetic in a ring so we can use the ring properties at the level of components. After applying the crucial fact about rings, we then reverse the steps. Since the calculation works for arbitrary i, j it follows that the matrix equation holds true. This Proposition provides a foundation for later work where we may find it convenient to prove a statement without resorting to a proof by components. Which method of proof is best depends on the question. However, I can't see another way of proving most of 1.2.6.

Definition 1.2.7.

The **zero matrix** in $R^{m \times n}$ is denoted 0 and defined by $0_{ij} = 0$ for all $(i, j) \in \mathbb{N}_m \times \mathbb{N}_n$. The additive inverse of $A \in R^{m \times n}$ is the matrix $-A$ such that $A + (-A) = 0$. The components of the additive inverse matrix are given by $(-A)_{ij} = -A_{ij}$ for all $(i, j) \in \mathbb{N}_m \times \mathbb{N}_n$.

The zero matrix joins a long list of other objects which are all denoted by 0 . Usually the meaning of 0 is clear from the context, the size of the zero matrix is chosen as to be consistent with the equation in which it is found.

Example 1.2.8. Solve the following matrix equation,

$$0 = \begin{bmatrix} x & y \\ z & w \end{bmatrix} + \begin{bmatrix} -1 & -2 \\ -3 & -4 \end{bmatrix} \Rightarrow \begin{bmatrix} 0 & 0 \\ 0 & 0 \end{bmatrix} = \begin{bmatrix} x-1 & y-2 \\ z-3 & w-4 \end{bmatrix}$$

The definition of matrix equality means this single matrix equation reduces to 4 scalar equations: $0 = x - 1, 0 = y - 2, 0 = z - 3, 0 = w - 4$. The solution is $x = 1, y = 2, z = 3, w = 4$.

Theorem 1.2.9.

If $A \in R^{m \times n}$ then

- (1.) $0 \cdot A = 0$, (scalar multiplication by 0 produces the zero matrix)
- (2.) $A + 0 = 0 + A = A$.

Proof: To prove (1.). Let $A \in R^{m \times n}$ and consider by definition of scalar multiplication of a matrix:

$$(0 \cdot A)_{ij} = 0(A_{ij}) = 0$$

for all i, j . Thus $0 \cdot A = 0$. To see (2.), observe by the definition of matrix addition and the zero-matrix:

$$(0 + A)_{ij} = 0_{ij} + A_{ij} = 0 + A_{ij} = A_{ij}$$

and as the above holds for all i, j we find $0 + A = A$. $\square$

1.2.1 standard column and matrix bases

The notation introduced in this subsection is near to my heart. It frees you to calculate a multitude of stupidly general claims with a minimum of writing.

Definition 1.2.10.

The symbol $\delta_{ij} = \begin{cases} 1 & , i = j \\ 0 & , i \neq j \end{cases}$ is called the **Kronecker delta**.

For example, $\delta_{22} = 1$ while $\delta_{12} = 0$.

Definition 1.2.11.

Let $e_i \in R^{n \times 1}$ be defined by $(e_i)_j = \delta_{ij}$. The size of the vector e_i is determined by context. We call e_i the i -th standard basis vector.

Example 1.2.12. Let me expand on what I mean by "context" in the definition above:

In R we have $e_1 = (1) = 1$ (by convention we drop the brackets in this case)

In R^2 we have $e_1 = (1, 0)$ and $e_2 = (0, 1)$.

In R^3 we have $e_1 = (1, 0, 0)$ and $e_2 = (0, 1, 0)$ and $e_3 = (0, 0, 1)$.

In R^4 we have $e_1 = (1, 0, 0, 0)$ and $e_2 = (0, 1, 0, 0)$ and $e_3 = (0, 0, 1, 0)$ and $e_4 = (0, 0, 0, 1)$.

Example 1.2.13. Any vector in R^n can be written as a sum of these basic vectors. For example,

$$\begin{aligned} v = (1, 2, 3) &= (1, 0, 0) + (0, 2, 0) + (0, 0, 3) \\ &= 1(1, 0, 0) + 2(0, 1, 0) + 3(0, 0, 1) \\ &= e_1 + 2e_2 + 3e_3. \end{aligned}$$

We say that v is a **finite linear combination** of e_1, e_2 and e_3 .

The concept of a finite linear combination is very important⁸.

Definition 1.2.14.

A **finite R -linear combination** of objects $A_1, A_2, \dots, A_k$ is a finite sum

$$c_1 A_1 + c_2 A_2 + \dots + c_k A_k = \sum_{i=1}^k c_i A_i$$

where the **coefficients** $c_i \in R$ for each i . If $c_1 = 0, c_2 = 0, \dots, c_k = 0$ then we say the linear combination is **trivial**. We also say that $\{0\}$ is formed by an **empty sum**. That is, a linear combination of $\emptyset$ is just $\{0\}$. We denote

$$\text{Span}_R \{A_1, \dots, A_k\} = \{c_1 A_1 + \dots + c_k A_k \mid c_1, \dots, c_k \in R\}.$$

The statement about the empty set $\emptyset$ helps theorems we state in future sections to be generally true. We will look at linear combinations of vectors, matrices and even functions in this course. The proposition below generalizes the calculation from Example 1.2.13.

Proposition 1.2.15.

Every vector in R^n is a linear combination of $e_1, e_2, \dots, e_n$.

Proof: Let $v = (v_1, v_2, \dots, v_n) \in R^n$. By the definition of vector addition and zero in R :

$$\begin{aligned} v &= (v_1 + 0, 0 + v_2, \dots, 0 + v_n) \\ &= (v_1, 0, \dots, 0) + (0, v_2, \dots, v_n) \\ &= (v_1, 0, \dots, 0) + (0, v_2, \dots, 0) + \dots + (0, 0, \dots, v_n) \\ &= (v_1 \cdot 1, v_1 \cdot 0, \dots, v_1 \cdot 0) + (v_2 \cdot 0, v_2 \cdot 1, \dots, v_2 \cdot 0) + \dots + (v_n \cdot 0, \dots, v_n \cdot 1) \end{aligned}$$

⁸since we only consider finite linear combinations in this course we generally omit the term *finite*

In the last step I rewrote each zero to emphasize that each entry of the k -th summand has a v_k factor. Continue by applying the definition of scalar multiplication to each vector in the sum above we find,

$$\begin{aligned} v &= v_1(1, 0, \dots, 0) + v_2(0, 1, \dots, 0) + \dots + v_n(0, 0, \dots, 1) \\ &= v_1 e_1 + v_2 e_2 + \dots + v_n e_n. \end{aligned}$$

Therefore, every vector in R^n is a linear combination of $e_1, e_2, \dots, e_n$. For each $v \in R^n$ we have $v = \sum_{i=1}^n v_i e_i$. $\square$

We can define a standard basis for matrices of arbitrary size in much the same manner.

Definition 1.2.16.

The ij -th **standard basis matrix** for $R^{m \times n}$ is denoted E_{ij} for $1 \leq i \leq m$ and $1 \leq j \leq n$. The matrix E_{ij} is zero in all entries except for the (i, j) -th slot where it has a 1. In other words, we define $(E_{ij})_{kl} = \delta_{ik} \delta_{jl}$.

Proposition 1.2.17.

Every matrix in $R^{m \times n}$ is a linear combination of the E_{ij} where $1 \leq i \leq m$ and $1 \leq j \leq n$.

Proof: Let $A \in R^{m \times n}$ then

$$\begin{aligned} A &= \begin{bmatrix} A_{11} & A_{12} & \cdots & A_{1n} \\ A_{21} & A_{22} & \cdots & A_{2n} \\ \vdots & \vdots & \cdots & \vdots \\ A_{m1} & A_{m2} & \cdots & A_{mn} \end{bmatrix} \\ &= A_{11} \begin{bmatrix} 1 & 0 & \cdots & 0 \\ 0 & 0 & \cdots & 0 \\ \vdots & \vdots & \cdots & 0 \\ 0 & 0 & \cdots & 0 \end{bmatrix} + A_{12} \begin{bmatrix} 0 & 1 & \cdots & 0 \\ 0 & 0 & \cdots & 0 \\ \vdots & \vdots & \cdots & 0 \\ 0 & 0 & \cdots & 0 \end{bmatrix} + \cdots + A_{mn} \begin{bmatrix} 0 & 0 & \cdots & 0 \\ 0 & 0 & \cdots & 0 \\ \vdots & \vdots & \cdots & 0 \\ 0 & 0 & \cdots & 1 \end{bmatrix} \\ &= A_{11} E_{11} + A_{12} E_{12} + \cdots + A_{mn} E_{mn}. \end{aligned}$$

The calculation above follows from repeated mn -applications of the definition of matrix addition and another mn -applications of the definition of scalar multiplication of a matrix. We can restate the final result in a more precise language,

$$A = \sum_{i=1}^m \sum_{j=1}^n A_{ij} E_{ij}.$$

As we claimed, any matrix can be written as a linear combination of the E_{ij} . $\square$

Alternate Proof: Let $A \in \mathbb{R}^{m \times n}$ then let $B = \sum_{i=1}^m \sum_{j=1}^n A_{ij} E_{ij}$ and calculate⁹,

$$B_{kl} = \left(\sum_{i=1}^m \sum_{j=1}^n A_{ij} E_{ij} \right)_{kl} = \sum_{i=1}^m \sum_{j=1}^n A_{ij} (E_{ij})_{kl} = \sum_{i=1}^m \sum_{j=1}^n A_{ij} \delta_{ik} \delta_{jl} = A_{kl}.$$

⁹here I have to repeatedly apply the definition of matrix addition and scalar multiplication. I will probably add a homework problem where you get to prove this follows from the definition by a simple induction argument

Thus $B_{kl} = A_{kl}$ for all $(k, l) \in \mathbb{N}_m \times \mathbb{N}_n$ hence $A = B$. $\square$

The term "basis" has a technical meaning which we will discuss at length in due time. For now, just think of it as part of the names of e_i and E_{ij} . These are the basic building blocks for matrix theory.

1.3 matrix multiplication

I don't seek to motivate the definition below. Well, not yet anyway. Rest assured there are many good reasons to multiply matrices in this way. We shall discover them as the course unfolds.

Definition 1.3.1. *Let R be a commutative ring.*

Let $A \in R^{m \times n}$ and $B \in R^{n \times p}$ then we say A and B are **multipliable** or **compatible** or **conformable**. The product of A and B is denoted by AB and $AB \in R^{m \times p}$ is defined by:

$$(AB)_{ij} = \sum_{k=1}^n A_{ik} B_{kj}$$

for each $1 \leq i \leq m$ and $1 \leq j \leq p$. In the case $m = p = 1$ the indices i, j are omitted in the equation since the matrix product is simply a number which needs no index. In particular, for $y \in R^{1 \times n}$ and $B \in R^{n \times p}$ then $yB \in R^{1 \times p}$ is a row-vector and

$$(yB)_j = \sum_{k=1}^n y_k B_{kj}.$$

Likewise, for $A \in R^{m \times n}$ and $x \in R^{n \times 1}$ then $Ax \in R^{m \times 1}$ is a column-vector and

$$(Ax)_i = \sum_{k=1}^n A_{ik} x_k.$$

If $y \in R^{1 \times n}$, $x \in R^{n \times 1}$ then $yx = \sum_{k=1}^n y_k x_k$. It is also possible for $n = 1$ in which case the summation is not needed, $(AB)_{ij} = A_{i1} B_{1j}$.

Notice, if $m = n = 1$ matrix multiplication reduces to multiplication by a scalar on the left and if $n = p = 1$ matrix multiplication reduces to multiplication by a scalar on the right.

Example 1.3.2. Suppose $A = (1, 2, 3, 4) \in R^{4 \times 1}$ and $B = [5, 6] \in R^{1 \times 2}$ then

$$AB = \begin{bmatrix} 1 \\ 2 \\ 3 \\ 4 \end{bmatrix} \begin{bmatrix} 5 & 6 \end{bmatrix} = \begin{bmatrix} 5 & 6 \\ 10 & 12 \\ 15 & 18 \\ 20 & 24 \end{bmatrix}.$$

The product studied in the example above is a column-row product. It is one of many ways to create a matrix from row and column vectors. What follows next is more common in applications. This definition is very nice for general proofs and we will need to know it for proofs. However, for explicit numerical examples, it is also useful to define *dot-products* (notice, I allow dot-products of row and column vectors for convenience of exposition)

Definition 1.3.3.

Let $v, w \in R^n \cup R^{1 \times n}$ then the **dot-product** of v and w is the number defined below:

$$v \bullet w = v_1 w_1 + v_2 w_2 + \cdots + v_n w_n = \sum_{k=1}^n v_k w_k.$$

In a later chapter we study the geometric¹⁰ content of the dot-product¹¹ when R is the field $\mathbb{R}$.

Proposition 1.3.4. *dot-product as a row-column multiplication:*

Let $v, w \in R^n$ then $v \bullet w = v^T w$.

Proof: Since v^T is an $1 \times n$ matrix and w is an $n \times 1$ matrix the definition of matrix multiplication indicates $v^T w$ should be a 1×1 matrix which is a number. Note in this case the outside indices ij are absent in the boxed equation so the equation reduces to

$$v^T w = v^T_1 w_1 + v^T_2 w_2 + \cdots + v^T_n w_n = v_1 w_1 + v_2 w_2 + \cdots + v_n w_n = v \bullet w. \quad \square$$

Proposition 1.3.5.

Let $A \in R^{m \times n}$ and $B \in R^{n \times p}$ then

$$AB = \begin{bmatrix} \text{row}_1(A) \cdot \text{col}_1(B) & \text{row}_1(A) \cdot \text{col}_2(B) & \cdots & \text{row}_1(A) \cdot \text{col}_p(B) \\ \text{row}_2(A) \cdot \text{col}_1(B) & \text{row}_2(A) \cdot \text{col}_2(B) & \cdots & \text{row}_2(A) \cdot \text{col}_p(B) \\ \vdots & \vdots & \cdots & \vdots \\ \text{row}_m(A) \cdot \text{col}_1(B) & \text{row}_m(A) \cdot \text{col}_2(B) & \cdots & \text{row}_m(A) \cdot \text{col}_p(B) \end{bmatrix}$$

Proof: The formula above claims $(AB)_{ij} = \text{row}_i(A) \cdot \text{col}_j(B)$ for all i, j . Recall that $(\text{row}_i(A))_k = A_{ik}$ and $(\text{col}_j(B))_k = B_{kj}$ thus

$$(AB)_{ij} = \sum_{k=1}^n A_{ik} B_{kj} = \sum_{k=1}^n (\text{row}_i(A))_k (\text{col}_j(B))_k$$

Hence, using definition of the dot-product, $(AB)_{ij} = \text{row}_i(A) \cdot \text{col}_j(B)$. This argument holds for all i, j therefore the Proposition is true. $\square$

Example 1.3.6. Let $A = \begin{bmatrix} 1 & 2 \\ 3 & 4 \end{bmatrix}$ and $v = \begin{bmatrix} x \\ y \end{bmatrix}$ then we may calculate the product Av as follows:

$$Av = \begin{bmatrix} 1 & 2 \\ 3 & 4 \end{bmatrix} \begin{bmatrix} x \\ y \end{bmatrix} = x \begin{bmatrix} 1 \\ 3 \end{bmatrix} + y \begin{bmatrix} 2 \\ 4 \end{bmatrix} = \begin{bmatrix} x + 2y \\ 3x + 4y \end{bmatrix}.$$

Notice that the product of an $n \times k$ matrix with a $k \times 1$ vector yields another vector of size $k \times 1$. In the example above we observed the pattern $(2 \times 2)(2 \times 1) \rightarrow (2 \times 1)$.

¹⁰In fact, there is some analog of the dot-product for complex numbers and quaternions. Many interesting *matrix groups* arise as *isometries* for these *inner products*.

¹¹The definition I give above is a bit unusual as it allows us to take the dot-product of row and column vectors. This is mostly a convenience of notation as to avoid writing a multitude of transposes in the Proposition below.

Example 1.3.7. The product of a 3×2 and 2×3 is a 3×3

$$\begin{bmatrix} 1 & 0 \\ 0 & 1 \\ 0 & 0 \end{bmatrix} \begin{bmatrix} 4 & 5 & 6 \\ 7 & 8 & 9 \end{bmatrix} = \begin{bmatrix} [1,0][4,7]^T & [1,0][5,8]^T & [1,0][6,9]^T \\ [0,1][4,7]^T & [0,1][5,8]^T & [0,1][6,9]^T \\ [0,0][4,7]^T & [0,0][5,8]^T & [0,0][6,9]^T \end{bmatrix} = \begin{bmatrix} 4 & 5 & 6 \\ 7 & 8 & 9 \\ 0 & 0 & 0 \end{bmatrix}$$

Example 1.3.8. The product of a 3×1 and 1×3 is a 3×3

$$\begin{bmatrix} 1 \\ 2 \\ 3 \end{bmatrix} \begin{bmatrix} 4 & 5 & 6 \end{bmatrix} = \begin{bmatrix} 4 \cdot 1 & 5 \cdot 1 & 6 \cdot 1 \\ 4 \cdot 2 & 5 \cdot 2 & 6 \cdot 2 \\ 4 \cdot 3 & 5 \cdot 3 & 6 \cdot 3 \end{bmatrix} = \begin{bmatrix} 4 & 5 & 6 \\ 8 & 10 & 12 \\ 12 & 15 & 18 \end{bmatrix}$$

Example 1.3.9. Let $A = \begin{bmatrix} 1 & 2 & 3 \\ 4 & 5 & 6 \\ 7 & 8 & 9 \end{bmatrix}$ and $v = \begin{bmatrix} 1 \\ 0 \\ -3 \end{bmatrix}$ calculate Av .

$$Av = \begin{bmatrix} 1 & 2 & 3 \\ 4 & 5 & 6 \\ 7 & 8 & 9 \end{bmatrix} \begin{bmatrix} 1 \\ 0 \\ -3 \end{bmatrix} = \begin{bmatrix} (1,2,3) \cdot (1,0,-3) \\ (4,5,6) \cdot (1,0,-3) \\ (7,8,9) \cdot (1,0,-3) \end{bmatrix} = \begin{bmatrix} -8 \\ -14 \\ -20 \end{bmatrix}.$$

Example 1.3.10. Let $A = \begin{bmatrix} 1 & 2 \\ 3 & 4 \end{bmatrix}$ and $B = \begin{bmatrix} 5 & 6 \\ 7 & 8 \end{bmatrix}$. We calculate

$$AB = \begin{bmatrix} 1 & 2 \\ 3 & 4 \end{bmatrix} \begin{bmatrix} 5 & 6 \\ 7 & 8 \end{bmatrix} = \begin{bmatrix} [1,2][5,7]^T & [1,2][6,8]^T \\ [3,4][5,7]^T & [3,4][6,8]^T \end{bmatrix} = \begin{bmatrix} 5+14 & 6+16 \\ 15+28 & 18+32 \end{bmatrix} = \begin{bmatrix} 19 & 22 \\ 43 & 50 \end{bmatrix}$$

Notice the product of square matrices is square. For numbers $a, b \in \mathbb{R}$ it we know the product of a and b is commutative ($ab = ba$). Let's calculate the product of A and B in the opposite order,

$$BA = \begin{bmatrix} 5 & 6 \\ 7 & 8 \end{bmatrix} \begin{bmatrix} 1 & 2 \\ 3 & 4 \end{bmatrix} = \begin{bmatrix} [5,6][1,3]^T & [5,6][2,4]^T \\ [7,8][1,3]^T & [7,8][2,4]^T \end{bmatrix} = \begin{bmatrix} 5+18 & 10+24 \\ 7+24 & 14+32 \end{bmatrix} = \begin{bmatrix} 23 & 34 \\ 31 & 46 \end{bmatrix}$$

Clearly $AB \neq BA$ thus matrix multiplication is **noncommutative** or **not commutative**.

Remark 1.3.11. *commutators*

The **commutator** of two square matrices A, B is given by $[A, B] = AB - BA$. If $[A, B] \neq 0$ then clearly $AB \neq BA$. There are many interesting properties of the commutator. It has deep physical significance in quantum mechanics. It is also the quintessential example of a **Lie Bracket**. It turns out that if the commutator of two observables is zero then they can be measured simultaneously to arbitrary precision. However, if the commutator of two observables is nonzero (such as is the case with position and momentum) then they cannot be simultaneously measured with arbitrary precision. The more precisely you know position, the less you know momentum and vice-versa. This is Heisenberg's **uncertainty principle** of quantum mechanics.

Properties of matrix multiplication are given in the theorem below. To summarize, matrix math works as you would expect with the exception that matrix multiplication is not commutative. We must be careful about the order of letters in matrix expressions.

Theorem 1.3.12.

If $A, B, C \in R^{m \times n}$, $X, Y \in R^{n \times p}$, $Z \in R^{p \times q}$ and $c_1, c_2 \in R$ then

- (1.) $(AX)Z = A(XZ)$,
- (2.) $(c_1A)X = c_1(AX) = A(c_1X) = (AX)c_1$,
- (3.) $A(X + Y) = AX + AY$,
- (4.) $A(c_1X + c_2Y) = c_1AX + c_2AY$,
- (5.) $(A + B)X = AX + BX$,

Proof: I leave the proofs of (1.), (2.), (4.) and (5.) to the reader. Proof of (3.): assume A, X, Y are given as in the statement of the Theorem. Observe that

$$\begin{aligned}
 ((A(X + Y))_{ij} &= \sum_k A_{ik}(X + Y)_{kj} && \text{defn. matrix multiplication,} \\
 &= \sum_k A_{ik}(X_{kj} + Y_{kj}) && \text{defn. matrix addition,} \\
 &= \sum_k (A_{ik}X_{kj} + A_{ik}Y_{kj}) && \text{dist. prop. of rings,} \\
 &= \sum_k A_{ik}X_{kj} + \sum_k A_{ik}Y_{kj} && \text{prop. of finite sum,} \\
 &= (AX)_{ij} + (AY)_{ij} && \text{defn. matrix multiplication}(\times 2), \\
 &= (AX + AY)_{ij} && \text{defn. matrix addition,}
 \end{aligned}$$

for all i, j . Therefore $A(X + Y) = AX + AY$. $\square$

The proofs of the other items are similar, I invite the reader to try to prove them in a style much like the proof I offer above.

We began our study of transpose in Proposition 1.1.8. Let us continue it:

Proposition 1.3.13. *Let R be a commutative ring with identity.*

- (1.) $(A^T)^T = A$ for all $A \in R^{m \times n}$,
- (2.) $(AB)^T = B^T A^T$ for all $A \in R^{m \times n}$ and $B \in R^{n \times p}$ (socks-shoes),
- (3.) $(cA)^T = cA^T$ for all $A \in R^{m \times n}$ and $c \in R$,
- (4.) $(A + B)^T = A^T + B^T$ for all $A, B \in R^{m \times n}$.

Proof: We proved (1.) for Proposition 1.1.8. Proof of (2.) is left to the reader. Proof of (3.) and (4.) is simple enough,

$$((A + cB)^T)_{ij} = (A + cB)_{ji} = A_{ji} + cB_{ji} = (A^T)_{ij} + ((cB)^T)_{ij}$$

for all i, j . Set $A = 0$ to obtain (3.) and set $c = 1$ to obtain (4.). $\square$

1.3.1 multiplication of row or column concatenations

Proposition 1.3.5 is not the only way to calculate the matrix product. In this subsection we find several new ways to decompose a product which are ideal to reveal such row or column patterns. In

some sense, this section is just a special case of the later section on block-multiplication. However, you could probably just as well say block multiplication is a simple outgrowth of what we study here. In any event, we need this material to properly understand the method to calculate A^{-1} and the final proposition of this section is absolutely critical to properly understand the structure of the solution set for $Ax = b$.

Example 1.3.14. *The product of a 2×2 and 2×1 is a 2×1 . Let $A = \begin{bmatrix} 1 & 2 \\ 3 & 4 \end{bmatrix}$ and let $v = \begin{bmatrix} 5 \\ 7 \end{bmatrix}$,*

$$Av = \begin{bmatrix} 1 & 2 \\ 3 & 4 \end{bmatrix} \begin{bmatrix} 5 \\ 7 \end{bmatrix} = \begin{bmatrix} [1, 2][5, 7]^T \\ [3, 4][5, 7]^T \end{bmatrix} = \begin{bmatrix} 19 \\ 43 \end{bmatrix}$$

Likewise, define $w = \begin{bmatrix} 6 \\ 8 \end{bmatrix}$ and calculate

$$Aw = \begin{bmatrix} 1 & 2 \\ 3 & 4 \end{bmatrix} \begin{bmatrix} 6 \\ 8 \end{bmatrix} = \begin{bmatrix} [1, 2][6, 8]^T \\ [3, 4][6, 8]^T \end{bmatrix} = \begin{bmatrix} 22 \\ 50 \end{bmatrix}$$

Something interesting to observe here, recall that in Example 1.3.10 we calculated

$AB = \begin{bmatrix} 1 & 2 \\ 3 & 4 \end{bmatrix} \begin{bmatrix} 5 & 6 \\ 7 & 8 \end{bmatrix} = \begin{bmatrix} 19 & 22 \\ 43 & 50 \end{bmatrix}$. But these are the same numbers we just found from the two matrix-vector products calculated above. We identify that B is just the **concatenation** of the vectors v and w ; $B = [v|w] = \begin{bmatrix} 5 & 6 \\ 7 & 8 \end{bmatrix}$. Observe that:

$$AB = A[v|w] = [Av|Aw].$$

The term **concatenate** is sometimes replaced with the word **adjoin**. I think of the process as gluing matrices together. This is an important operation since it allows us to lump together many solutions into a single matrix of solutions. (I will elaborate on that in detail in a future section)

Proposition 1.3.15. *the concatenation proposition for columns*

Let $A \in R^{m \times n}$ and $B \in R^{n \times p}$ then we can understand the matrix multiplication of A and B as the concatenation of several matrix-vector products,

$$AB = A[col_1(B)|col_2(B)|\cdots|col_p(B)] = [Acol_1(B)|Acol_2(B)|\cdots|Acol_p(B)]$$

Proof: see the Problem Set. You should be able to follow the same general strategy as the Proof of Proposition 1.3.5. Show that the i, j -th entry of the L.H.S. is equal to the matching entry on the R.H.S. Good hunting. $\square$

There are actually many many different ways to perform the calculation of matrix multiplication. Proposition 1.3.15 essentially parses the problem into a bunch of (matrix)(column vector) calculations. You could go the other direction and view AB as a bunch of (row vector)(matrix) products glued together. In particular,

Proposition 1.3.16. *the concatenation proposition for rows*

Let $A \in R^{m \times n}$ and $B \in R^{n \times p}$ then we can understand the matrix multiplication of A and B as the concatenation of several matrix-vector products,

$$AB = \begin{bmatrix} \text{row}_1(A) \\ \text{row}_2(A) \\ \vdots \\ \text{row}_m(A) \end{bmatrix} B = \begin{bmatrix} \text{row}_1(A)B \\ \text{row}_2(A)B \\ \vdots \\ \text{row}_m(A)B \end{bmatrix}.$$

Proof: let $R_i = \text{row}_i(A)$ hence $A^T = [R_1^T | R_2^T | \cdots | R_m^T]$. Use Proposition 1.3.15 to calculate:

$$B^T A^T = B^T [R_1^T | R_2^T | \cdots | R_m^T] = [B^T R_1^T | B^T R_2^T | \cdots | B^T R_m^T] \star$$

But, $(B^T A^T)^T = (A^T)^T (B^T)^T = AB$. Thus, taking the transpose of $\star$ yields

$$AB = [B^T R_1^T | B^T R_2^T | \cdots | B^T R_m^T]^T = \begin{bmatrix} \frac{(B^T R_1^T)^T}{(B^T R_2^T)^T} \\ \vdots \\ \frac{(B^T R_m^T)^T}{(B^T R_m^T)^T} \end{bmatrix} = \begin{bmatrix} \frac{R_1 B}{R_2 B} \\ \vdots \\ \frac{R_m B}{R_m B} \end{bmatrix}.$$

where we used $(B^T R_i^T)^T = (R_i^T)^T (B^T)^T = R_i B$ for each i in the last step. $\square$

There are stranger ways to calculate the product. You can also assemble the product by adding together a bunch of outer-products of the rows of A with the columns of B . The dot-product of two vectors is an example of an inner product and we saw $v \cdot w = v^T w$. The outer-product of two vectors goes the other direction: given $v \in R^n$ and $w \in R^m$ we find $vw^T \in R^{n \times m}$.

Proposition 1.3.17. *matrix multiplication as sum of outer products.*

Let $A \in R^{m \times n}$ and $B \in R^{n \times p}$ then

$$AB = \text{col}_1(A)\text{row}_1(B) + \text{col}_2(A)\text{row}_2(B) + \cdots + \text{col}_n(A)\text{row}_n(B).$$

Proof: consider the i, j -th component of AB , by definition we have

$$(AB)_{ij} = \sum_{k=1}^n A_{ik} B_{kj} = A_{i1} B_{1j} + A_{i2} B_{2j} + \cdots + A_{in} B_{nj}$$

but note that $(\text{col}_k(A)\text{row}_k(B))_{ij} = \text{col}_k(A)_i \text{row}_k(B)_j = A_{ik} B_{kj}$ for each $k = 1, 2, \dots, n$ and the proposition follows. $\square$

A corollary is a result which falls immediately from a given result. Take the case $B = v \in R^{n \times 1}$ to prove the following:

Corollary 1.3.18. *matrix-column product is linear combination of columns.*

Let $A \in R^{m \times n}$ and $v \in R^n$ then

$$Av = v_1 \text{col}_1(A) + v_2 \text{col}_2(A) + \cdots + v_n \text{col}_n(A).$$

Some texts use the result above as the foundational definition for matrix multiplication. We took a different approach in these notes, largely because I wish for students to gain better grasp of index calculation. If you'd like to know more about the other approach, I can recommend some reading.

1.3.2 all your base are belong to us (e_i and E_{ij} that is)

Example 1.3.19. Suppose $A \in R^{m \times n}$ and $e_i \in R^n$ is a standard basis vector,

$$(Ae_i)_j = \sum_{k=1}^n A_{jk}(e_i)_k = \sum_{k=1}^n A_{jk}\delta_{ik} = A_{ji}$$

Thus, $\boxed{Ae_i = \text{col}_i(A)}$. We find that multiplication of a matrix A by the standard basis e_i yields the i -th column of A .

Example 1.3.20. Suppose $A \in R^{m \times n}$ and $e_i \in R^{m \times 1}$ is a standard basis vector,

$$(e_i^T A)_j = \sum_{k=1}^n (e_i)_k A_{kj} = \sum_{k=1}^n \delta_{ik} A_{kj} = A_{ij}$$

Thus, $\boxed{e_i^T A = \text{row}_i(A)}$. We find multiplication of a matrix A by the transpose of standard basis e_i yields the i -th row of A .

Example 1.3.21. Again, suppose $e_i, e_j \in R^n$ are standard basis vectors. The product $e_i^T e_j$ of the $1 \times n$ and $n \times 1$ matrices is just a 1×1 matrix which is just a number. In particular consider,

$$e_i^T e_j = \sum_{k=1}^n (e_i^T)_k (e_j)_k = \sum_{k=1}^n \delta_{ik} \delta_{jk} = \delta_{ij}$$

The product is zero unless the vectors are identical.

Example 1.3.22. Suppose $e_i \in R^{m \times 1}$ and $e_j \in R^n$. The product of the $m \times 1$ matrix e_i and the $1 \times n$ matrix e_j^T is an $m \times n$ matrix. In particular,

$$(e_i e_j^T)_{kl} = (e_i)_k (e_j^T)_l = \delta_{ik} \delta_{jl} = (E_{ij})_{kl}.$$

Thus the standard basis matrices are constructed from the standard basis vectors; $E_{ij} = e_i e_j^T$.

Example 1.3.23. What about the matrix E_{ij} ? What can we say about multiplication by E_{ij} on the right of an arbitrary matrix? Let $A \in R^{m \times n}$ and consider,

$$(AE_{ij})_{kl} = \sum_{p=1}^n A_{kp} (E_{ij})_{pl} = \sum_{p=1}^n A_{kp} \delta_{ip} \delta_{jl} = A_{ki} \delta_{jl}$$

Notice the matrix above has zero entries unless $j = l$ which means that the matrix is mostly zero except for the j -th column. We can select the j -th column by multiplying the above by e_j , using Examples 1.3.21 and 1.3.19,

$$(AE_{ij} e_j)_k = (Ae_i e_j^T e_j)_k = (Ae_i \delta_{jj})_k = (Ae_i)_k = (\text{col}_i(A))_k$$

This means,

$$AE_{ij} = \begin{bmatrix} & \text{column } j \\ 0 & 0 & \cdots & A_{1i} & \cdots & 0 \\ 0 & 0 & \cdots & A_{2i} & \cdots & 0 \\ \vdots & \vdots & \cdots & \vdots & \cdots & \vdots \\ 0 & 0 & \cdots & A_{mi} & \cdots & 0 \end{bmatrix}$$

Right multiplication of matrix A by E_{ij} moves the i -th column of A to the j -th column of AE_{ij} and all other entries are zero. It turns out that left multiplication by E_{ij} moves the j -th row of A to the i -th row and sets all other entries to zero.

Example 1.3.24. Let $A = \begin{bmatrix} 1 & 2 \\ 3 & 4 \end{bmatrix}$ consider multiplication by E_{12} ,

$$AE_{12} = \begin{bmatrix} 1 & 2 \\ 3 & 4 \end{bmatrix} \begin{bmatrix} 0 & 1 \\ 0 & 0 \end{bmatrix} = \begin{bmatrix} 0 & 1 \\ 0 & 3 \end{bmatrix} = [0 \mid \text{col}_1(A)]$$

Which agrees with our general abstract calculation in the previous example. Next consider,

$$E_{12}A = \begin{bmatrix} 0 & 1 \\ 0 & 0 \end{bmatrix} \begin{bmatrix} 1 & 2 \\ 3 & 4 \end{bmatrix} = \begin{bmatrix} 3 & 4 \\ 0 & 0 \end{bmatrix} = \begin{bmatrix} \text{row}_2(A) \\ 0 \end{bmatrix}.$$

Example 1.3.25. Calculate the product of E_{ij} and E_{kl} .

$$(E_{ij}E_{kl})_{mn} = \sum_p (E_{ij})_{mp}(E_{kl})_{pn} = \sum_p \delta_{im}\delta_{jp}\delta_{kp}\delta_{ln} = \delta_{im}\delta_{jk}\delta_{ln}$$

For example,

$$(E_{12}E_{34})_{mn} = \delta_{1m}\delta_{23}\delta_{4n} = 0.$$

In order for the product to be nontrivial we must have $j = k$,

$$(E_{12}E_{24})_{mn} = \delta_{1m}\delta_{22}\delta_{4n} = \delta_{1m}\delta_{4n} = (E_{14})_{mn}.$$

We can make the same identification in the general calculation,

$$(E_{ij}E_{kl})_{mn} = \delta_{jk}(E_{il})_{mn}.$$

Since the above holds for all m, n ,

$$\boxed{E_{ij}E_{kl} = \delta_{jk}E_{il}}$$

this is at times a very nice formula to know about.

Remark 1.3.26.

The proofs in these examples are much longer if written without the benefit of index notation. It usually takes most students a little time to master the idea of index notation. There are a few homeworks assigned which require this sort of thinking, I do expect all students of Math 321 to gain proficiency in index calculation.

Example 1.3.27. Let $A \in R^{m \times n}$ and suppose $e_i \in R^{m \times 1}$ and $e_j \in R^n$. Consider,

$$(e_i)^T A e_j = \sum_{k=1}^m ((e_i)^T)_k (A e_j)_k = \sum_{k=1}^m \delta_{ik} (A e_j)_k = (A e_j)_i = A_{ij}$$

This is a useful observation. If we wish to select the (i, j) -entry of the matrix A then we can use the following simple formula,

$$A_{ij} = (e_i)^T A e_j$$

This is analogous to the idea of using dot-products to select particular components of vectors in analytic geometry; (reverting to calculus III notation for a moment) recall that to find v_1 of $\vec{v}$ we learned that the dot product by $\hat{i} = \langle 1, 0, 0 \rangle$ selects the first components $v_1 = \vec{v} \cdot \hat{i}$. The following theorem is simply a summary of our results for this subsection.

Theorem 1.3.28.

Let $A \in R^{m \times n}$ and $v \in R^n$ if $(E_{ij})_{kl} = \delta_{ik} \delta_{jl}$ and $(e_i)_j = \delta_{ij}$ then,

$v = \sum_{i=1}^n v_i e_i$	$A = \sum_{i=1}^m \sum_{j=1}^n A_{ij} E_{ij}$	$e_i^T A = \text{row}_i(A)$	$A e_i = \text{col}_i(A)$
$A_{ij} = (e_i)^T A e_j$	$E_{ij} E_{kl} = \delta_{jk} E_{il}$	$E_{ij} = e_i e_j^T$	$e_i^T e_j = \delta_{ij}$

The reader should understand I am abusing notation in the case $m \neq n$. For example, to build the $m \times n$ matrix units for rectangular matrices we probably should use a notation like $E_{ij} = \bar{e}_i e_j^T$ where $\bar{e}_i \in R^m$ and $e_j \in R^n$. Likewise, $A_{ij} = \bar{e}_i^T A e_j$ for $1 \leq i \leq m$ and $1 \leq j \leq n$.

1.4 matrix algebra

In this subsection we discover the matrix analog of the number 1, the formulation of the multiplicative inverse and raising a matrix to a power.

1.4.1 identity and inverse matrices

We begin by studying the 2×2 case.

Example 1.4.1. Let $I = \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix}$ and $A = \begin{bmatrix} a & b \\ c & d \end{bmatrix}$. We calculate

$$IA = \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix} \begin{bmatrix} a & b \\ c & d \end{bmatrix} = \begin{bmatrix} a & b \\ c & d \end{bmatrix}$$

Likewise calculate,

$$AI = \begin{bmatrix} a & b \\ c & d \end{bmatrix} \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix} = \begin{bmatrix} a & b \\ c & d \end{bmatrix}$$

Since the matrix A was arbitrary we conclude that $IA = AI$ for all $A \in R^{2 \times 2}$.

Definition 1.4.2.

The identity matrix in $R^{n \times n}$ is the $n \times n$ square matrix I which has components $I_{ij} = \delta_{ij} = \begin{cases} 1 & i = j \\ 0 & i \neq j \end{cases}$. The notation I_n is sometimes used if the size of the identity matrix needs emphasis, otherwise the size of the matrix I is to be understood from the context.

$$I_2 = \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix} \quad I_3 = \begin{bmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{bmatrix} \quad I_4 = \begin{bmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \end{bmatrix}$$

You might wonder, do other square matrices D satisfy $AD = DA$ for each square matrix A ? That sounds like an excellent homework problem¹². For now, let's see how Example 1.4.1 generalizes:

Proposition 1.4.3.

If $X \in R^{n \times p}$ then $XI_p = X$ and $I_n X = X$.

Proof: I omit the p in I_p to reduce clutter below. Consider the i, j component of XI ,

$$\begin{aligned} (XI)_{ij} &= \sum_{k=1}^p X_{ik} I_{kj} && \text{defn. matrix multiplication} \\ &= \sum_{k=1}^p X_{ik} \delta_{kj} && \text{defn. of } I \\ &= X_{ij} \end{aligned}$$

The last step follows from the fact that all other terms in the sum are made zero by the Kronecker delta. Finally, observe the calculation above holds for all i, j hence $XI = X$. The proof of $IX = X$ is left to the reader. $\square$

Before we define the inverse of a matrix it is wise to prove the following:

Proposition 1.4.4. *Let R be a commutative ring.*

Suppose $A \in R^{n \times n}$. If $B, C \in R^{n \times n}$ satisfy $AB = BA = I$ and $AC = CA = I$ then $B = C$.

Proof: suppose $A, B, C \in R^{n \times n}$ and $AB = BA = I$ and $AC = CA = I$ thus $AB = AC$. Multiply B on the left of $AB = AC$ to obtain $BAB = BAC$ hence $IB = IC \Rightarrow B = C$. $\square$

The identity matrix plays the role of the multiplicative identity for matrix multiplication. If $AB = I$ then we **do not write** $B = I/A$, instead, the following notation is customary:

Definition 1.4.5.

Let $A \in R^{n \times n}$. If there exists $B \in R^{n \times n}$ such that $AB = I$ and $BA = I$ then we say that A is **invertible** and $A^{-1} = B$. Invertible matrices are also called **nonsingular**. If a matrix has no inverse then it is called a **noninvertible** or **singular** matrix.

¹²see § 1.3.2 for notation which is helpful to characterize D for which $AD = DA$ for all A

We'll discuss how and when it is possible to calculate A^{-1} for a given square matrix A , however, we need to develop a few tools before we're ready for that problem. This much I can share now:

Example 1.4.6. Consider the problem of inverting a 2×2 matrix $A = \begin{bmatrix} a & b \\ c & d \end{bmatrix}$. We seek $B = \begin{bmatrix} x & y \\ z & w \end{bmatrix}$ such that $AB = I$ and $BA = I$. The resulting algebra would lead you to conclude $x = d/t, y = -b/t, z = -c/t, w = a/t$ where $t = ad - bc$.

$$\begin{bmatrix} a & b \\ c & d \end{bmatrix}^{-1} = \frac{1}{ad - bc} \begin{bmatrix} d & -b \\ -c & a \end{bmatrix}$$

It's not hard to show this formula works,

$$\begin{aligned} \frac{1}{ad - bc} \begin{bmatrix} a & b \\ c & d \end{bmatrix} \begin{bmatrix} d & -b \\ -c & a \end{bmatrix} &= \frac{1}{ad - bc} \begin{bmatrix} ad - bc & -ab + ab \\ cd - dc & -bc + da \end{bmatrix} \\ &= \frac{1}{ad - bc} \begin{bmatrix} ad - bc & 0 \\ 0 & ad - bc \end{bmatrix} = \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix} \end{aligned}$$

Proof that $BA = I$ is similar.

the quantity $ad - bc$ for $A = \begin{bmatrix} a & b \\ c & d \end{bmatrix}$ is known as the **determinant** of A . In particular, we denote $\det(A) = ad - bc$. If $R = \mathbb{R}$ then the significance of the determinant is that it provides the signed-area of the parallelogram with sides $\langle a, b \rangle, \langle c, d \rangle$. The sign tells us if $\langle a, b \rangle$ is rotated clockwise (CW) or counterclockwise (CCW) to reach $\langle c, d \rangle$. If you study it carefully, you'll find positive determinant indicates the second row is obtained from the first by a CCW rotation. Determinants are discussed in more detail later in this course. In fact, this formula is generalized to n -th order matrices. However, the formula is so complicated that only a truly silly student¹³ would try to implement it for anything beyond the 2×2 case. Computationally, we find an efficient algorithm for finding inverses larger matrices in § 1.6.2.

Example 1.4.7. One interesting application of 2×2 matrices is that they can be used to generate rotations in the plane. In particular, a counterclockwise **rotation** by angle θ in the plane can be represented by a matrix $R(\theta) = \begin{bmatrix} \cos \theta & \sin \theta \\ -\sin \theta & \cos \theta \end{bmatrix}$. Calculate via the 2×2 inverse formula with $a = d = \cos \theta$ and $b = -c = \sin \theta$

$$(R(\theta))^{-1} = \frac{1}{\cos^2 \theta + \sin^2 \theta} \begin{bmatrix} \cos \theta & -\sin \theta \\ \sin \theta & \cos \theta \end{bmatrix} = \begin{bmatrix} \cos(-\theta) & \sin(-\theta) \\ -\sin(-\theta) & \cos(-\theta) \end{bmatrix} = R(-\theta)$$

We observe the inverse matrix corresponds to a rotation by angle $-\theta$; $R(\theta)^{-1} = R(-\theta)$. Notice that $R(0) = \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix}$ thus $R(\theta)R(-\theta) = R(0) = I$. Rotations are very special invertible matrices, we shall see them again.

Noninvertible matrices challenge our intuition. For example, if A^{-1} does not exist for A it is possible to have $Av = Aw$ and yet $v \neq w$. For example, $A = \begin{bmatrix} 1 & 1 \\ 0 & 0 \end{bmatrix}$ is not invertible and we observe $Ae_1 = e_1 = Ae_2$ and obviously $e_1 \neq e_2$. Invertible matrices allow some of our usual thinking:

¹³his name is Minh

Theorem 1.4.8.

If $A, B \in R^{n \times n}$ are invertible, $X, Y \in R^{m \times n}$, $Z, W \in R^{n \times m}$ and nonzero $c \in R$ then

- (1.) $(AB)^{-1} = B^{-1}A^{-1}$,
- (2.) $(cA)^{-1} = \frac{1}{c}A^{-1}$,
- (3.) $XA = YA$ implies $X = Y$,
- (4.) $AZ = AW$ implies $Z = W$,
- (5.) $(A^T)^{-1} = (A^{-1})^T$.
- (6.) $(A^{-1})^{-1} = A$.

Proof: To prove (1.) simply notice that

$$(AB)B^{-1}A^{-1} = A(BB^{-1})A^{-1} = A(I)A^{-1} = AA^{-1} = I.$$

also¹⁴ note $B^{-1}A^{-1}(AB) = I$. The proof of (2.) follows from the calculation below,

$$(\frac{1}{c}A^{-1})cA = \frac{1}{c}cA^{-1}A = A^{-1}A = I.$$

and note $cA(\frac{1}{c}A^{-1}) = I$ by nearly the same calculation. To prove (3.) assume that $XA = YA$ and multiply both sides by A^{-1} on the right to obtain $XAA^{-1} = YAA^{-1}$ which reveals $XI = YI$ or simply $X = Y$. To prove (4.) multiply by A^{-1} on the left. Finally, consider $AA^{-1} = I$ and $A^{-1}A = I$ implies by the socks-shoes identity for the transpose that $(A^{-1})^TA^T = I^T = I$ and $A^T(A^{-1})^T = I^T = I$ therefore $(A^T)^{-1} = (A^{-1})^T$. Finally, (6.) is immediate from the definition. $\square$

Remark 1.4.9.

The proofs just given were all matrix arguments. These contrast the component level proofs needed for 1.2.6. We could give component level proofs for the Theorem above but that is not necessary and those arguments would only obscure the point. I hope you gain your own sense of which type of argument is most appropriate as the course progresses.

The importance of inductive arguments in linear algebra ought not be overlooked.

Proposition 1.4.10.

If $A_1, A_2, \dots, A_k \in R^{n \times n}$ are invertible then $(A_1A_2 \cdots A_k)^{-1} = A_k^{-1}A_{k-1}^{-1} \cdots A_1^{-1}$.

Proof: follows from induction on k . In particular, $k = 1$ is trivial. Assume inductively the proposition is true for some k with $k \geq 2$,

$$(\underbrace{A_1A_2 \cdots A_k}_B A_{k+1})^{-1} = (BA_{k+1})^{-1} = A_{k+1}^{-1}B^{-1}$$

by Theorem 1.4.8 part (1.). Applying the induction hypothesis to B yields

$$(A_1A_2 \cdots A_{k+1})^{-1} = A_{k+1}^{-1}A_k^{-1} \cdots A_1^{-1} \quad \square.$$

¹⁴My apologies to the reader who already knows that $AB = I$ implies $BA = I$ for square matrices A, B . We have yet to learn that. We shall soon, but, for now these proofs have a bit extra.

1.4.2 matrix powers

The power of a matrix is defined in the natural way. Notice we need for A to be square in order for the product AA to be defined.

Definition 1.4.11.

Let $A \in R^{n \times n}$. We define $A^0 = I$, $A^1 = A$ and $A^m = AA^{m-1}$ for all $m \geq 1$. If A is invertible then $A^{-p} = (A^{-1})^p$.

As you would expect, $A^3 = AA^2 = AAA$.

Proposition 1.4.12. laws of exponents

Consider nonzero $A, B \in R^{n \times n}$ and

- (1.) $A^p A^q = A^{p+q}$ for all $p, q \in \mathbb{N} \cup \{0\}$,
- (2.) $(A^p)^q = A^{pq}$ for all $p, q \in \mathbb{N} \cup \{0\}$,
- (3.) if A^{-1} exists then the $A^p A^q = A^{p+q}$ and $(A^p)^q = A^{pq}$ for $p, q \in \mathbb{Z}$.

Proof: we prove (1.) by induction on q . Fix $p \in \mathbb{N} \cup \{0\}$. If $q = 0$ then $A^q = A^0 = I$ thus $A^p A^q = A^p I = A^p = A^{p+0} = A^{p+q}$ thus (1.) is true for $q = 0$. Suppose inductively that (1.) is true for some $q \in \mathbb{N}$. Consider,

$$\begin{aligned} A^p A^{q+1} &= A^p A^q A && \text{by definition of matrix power} \\ &= A^{p+q} A && \text{by induction hypothesis} \\ &= A^{p+q+1} && \text{by definition of matrix power} \end{aligned}$$

thus (1.) holds for $q+1$ and we conclude (1.) is true for all $q \in \mathbb{N} \cup \{0\}$ for arbitrary $p \in \mathbb{N} \cup \{0\}$. I leave (2.) for the reader, it can be shown by a similar inductive argument. To prove $A^p A^q = A^{p+q}$ fix $p \in \mathbb{Z}$ and notice the claim is true for $q = 0$. Our argument for $q \in \mathbb{N}$ still is valid when we take $p \in \mathbb{Z}$ (it was non-negative in our argument for (1.)). Hence, consider $q \in \mathbb{Z}$ with $q \leq 0$. Let $q = -r$ and observe $r \geq 0$. We intend to prove $A^p A^{-r} = A^{p+(-r)}$ for all $r \in \mathbb{N} \cup \{0\}$ by induction on r . Note $A^p A^{-r} = A^{p+(-r)}$ is true for $r = 0$. Suppose inductively $A^p A^{-r} = A^{p+(-r)}$ for some $r \in \mathbb{N}$. Consider,

$$\begin{aligned} A^p A^{-(r+1)} &= A^p (A^{-1})^{r+1} && \text{by definition of matrix power} \\ &= A^p (A^{-1})^r A^{-1} && \text{by definition of matrix power} \\ &= A^p A^{-r} A^{-1} && \text{by definition of matrix power} \\ &= A^{p+(-r)} A^{-1} && \text{by induction hypothesis} \\ &= A^{-(r-p)} A^{-1} && \text{arithmetic} \\ &= (A^{-1})^{r-p} A^{-1} && \text{definition of matrix power} \\ &= (A^{-1})^{r-p+1} && \text{definition of matrix power} \\ &= A^{p-(r+1)} && \text{definition of matrix power} \end{aligned}$$

hence the claim is true for $r+1$ and it follows by induction it is true for $r \in \mathbb{N} \cup \{0\}$. Hence, $A^p A^q = A^{p+q}$ is true for all $p, q \in \mathbb{Z}$. I leave proof of the other half of (3.) to the reader, the

argument should be similar. $\square$

You should notice that $(AB)^p \neq A^p B^p$ for matrices. Instead,

$$(AB)^2 = ABAB, \quad (AB)^3 = ABABAB, \text{ etc...}$$

This means the binomial theorem will not hold for matrices. For example,

$$(A + B)^2 = (A + B)(A + B) = A(A + B) + B(A + B) = AA + AB + BA + BB$$

hence $(A + B)^2 \neq A^2 + 2AB + B^2$ as the matrix product is not generally commutative. However, in the special case that $AB = BA$ and we can prove that $(AB)^p = A^p B^p$ and the binomial theorem holds true as well. I may have assigned you the proof of the binomial theorem in homework.

1.4.3 symmetric and antisymmetric matrices

Definition 1.4.13.

Let $A \in \mathbb{R}^{n \times n}$. We say A is **symmetric** iff $A^T = A$. We say A is **antisymmetric** iff $A^T = -A$.

At the level of components, $A^T = A$ gives $A_{ij} = A_{ji}$ for all i, j . Whereas, $A^T = -A$ gives $A_{ij} = -A_{ji}$ for all i, j . Both symmetric and antisymmetric matrices appear in common physical applications. For example, the inertia tensor which describes the rotational motion of a body is represented by a symmetric 3×3 matrix. The Faraday tensor is represented by an antisymmetric 4×4 matrix. The Faraday tensor includes both the electric and magnetic fields. Physics aside, we'll see later that symmetric matrices play an important role in multivariate Taylor series and the Spectral Theorem makes symmetric matrices especially simple to analyze in general. We might study the Spectral Theorem towards the end of this course.

Example 1.4.14. *Examples and non-examples of symmetric and antisymmetric matrices:*

$$\underbrace{I, O, E_{ii}, \begin{bmatrix} 1 & 2 \\ 2 & 0 \end{bmatrix}}_{\text{symmetric}} \quad \underbrace{O, \begin{bmatrix} 0 & 2 \\ -2 & 0 \end{bmatrix}}_{\text{antisymmetric}} \quad \underbrace{[1, 2], E_{i,i+1}, \begin{bmatrix} 1 & 2 \\ 3 & 4 \end{bmatrix}}_{\text{neither}}$$

Proposition 1.4.15.

Let $A \in \mathbb{R}^{m \times n}$ then $A^T A$ is symmetric.

Proof: Proposition 1.3.13 yields $(A^T A)^T = A^T (A^T)^T = A^T A$. Thus $A^T A$ is symmetric. $\square$

Proposition 1.4.16.

If A is symmetric then A^k is symmetric for all $k \in \mathbb{N}$.

Proof: Suppose $A^T = A$. Proceed inductively. Clearly $k = 1$ holds true since $A^1 = A$. Assume inductively that A^k is symmetric.

$$\begin{aligned} (A^{k+1})^T &= (AA^k)^T && \text{defn. of matrix exponents,} \\ &= (A^k)^T A^T && \text{socks-shoes prop. of transpose,} \\ &= A^k A && \text{using induction hypothesis.} \\ &= A^{k+1} && \text{defn. of matrix exponents,} \end{aligned}$$

thus by proof by mathematical induction A^k is symmetric for all $k \in \mathbb{N}$. $\square$

1.4.4 triangular and diagonal matrices

Definition 1.4.17.

Let $A \in R^{m \times n}$. If $A_{ij} = 0$ for all i, j such that $i \neq j$ then A is called a **diagonal matrix**. If A has components $A_{ij} = 0$ for all i, j such that $i \leq j$ then we call A a **upper triangular matrix**. If A has components $A_{ij} = 0$ for all i, j such that $i \geq j$ then we call A a **lower triangular matrix**. If the diagonal of a matrix is zero then the matrix is **hollow**.

Example 1.4.18. Let me illustrate a generic example of each case for 3×3 matrices:

$$\begin{bmatrix} A_{11} & 0 & 0 \\ 0 & A_{22} & 0 \\ 0 & 0 & A_{33} \end{bmatrix} \quad \begin{bmatrix} A_{11} & A_{12} & A_{13} \\ 0 & A_{22} & A_{23} \\ 0 & 0 & A_{33} \end{bmatrix} \quad \begin{bmatrix} A_{11} & 0 & 0 \\ A_{21} & A_{22} & 0 \\ A_{31} & A_{32} & A_{33} \end{bmatrix}$$

As you can see the diagonal matrix only has nontrivial entries on the diagonal, and the names lower triangular and upper triangular are likewise natural.

If an upper triangular matrix has zeros on the diagonal then it is said to be **strictly upper triangular**. Likewise, if a lower triangular matrix has zeros on the diagonal then it is said to be **strictly lower triangular**. Obviously any matrix can be written as a sum of a diagonal and strictly upper and strictly lower matrix,

$$\begin{aligned} A &= \sum_{i,j} A_{ij} E_{ij} \\ &= \sum_i A_{ii} E_{ii} + \sum_{i < j} A_{ij} E_{ij} + \sum_{i > j} A_{ij} E_{ij} \end{aligned}$$

There is an algorithm called *LU-factorization* which for many matrices¹⁵ finds a lower triangular matrix L and an upper triangular matrix U such that $A = LU$. It is one of several factorization schemes which is computationally advantageous for large systems. There are many many ways to solve a system, but some are faster methods. Algorithmics is the study of which method is optimal.

Example 1.4.19. In the 2×2 case it is simple to verify the product of upper(lower) triangular matrices is once more (upper)lower triangular:

$$\begin{bmatrix} a & 0 \\ b & c \end{bmatrix} \begin{bmatrix} x & 0 \\ y & z \end{bmatrix} = \begin{bmatrix} ax & 0 \\ bx + cy & cz \end{bmatrix} \quad \& \quad \begin{bmatrix} a & b \\ 0 & c \end{bmatrix} \begin{bmatrix} x & y \\ 0 & z \end{bmatrix} = \begin{bmatrix} ax & ay + bz \\ 0 & cz \end{bmatrix}$$

Generally, the 2×2 case is surprisingly insightful. This is such a case:

Proposition 1.4.20.

Let $A, B \in R^{n \times n}$.

- (1.) If A, B are diagonal then AB is diagonal.
- (2.) If A, B are upper triangular then AB is upper triangular.
- (3.) If A, B are lower triangular then AB is lower triangular.

¹⁵An LU decomposition exists iff the principal minors are all positive. However, a PLU (permutation, lower, upper) factorization always exists. I will discuss this in Math 221.

Proof of (1.): Suppose A and B are diagonal. It follows there exist a_i, b_j such that $A = \sum_i a_i E_{ii}$ and $B = \sum_j b_j E_{jj}$. Calculate,

$$AB = \sum_i a_i E_{ii} \sum_j b_j E_{jj} = \sum_i \sum_j a_i b_j E_{ii} E_{jj} = \sum_i \sum_j a_i b_j \delta_{ij} E_{ij} = \sum_i a_i b_i E_{ii}$$

thus the product matrix AB is also diagonal and we find that the diagonal of the product AB is just the product of the corresponding diagonals of A and B .

Proof of (2.): Suppose A and B are upper triangular. It follows there exist A_{ij}, B_{ij} such that¹⁶ $A = \sum_{i \leq j} A_{ij} E_{ij}$ and $B = \sum_{k \leq l} B_{kl} E_{kl}$. Calculate,

$$AB = \sum_{i \leq j} A_{ij} E_{ij} \sum_{k \leq l} B_{kl} E_{kl} = \sum_{i \leq j} \sum_{k \leq l} A_{ij} B_{kl} E_{ij} E_{kl} = \sum_{i \leq j} \sum_{k \leq l} A_{ij} B_{kl} \delta_{jk} E_{il} = \sum_{i \leq j} \sum_{j \leq l} A_{ij} B_{jl} E_{il}.$$

Notice that every term in the sum above has $i \leq j$ and $j \leq l$ hence $i \leq l$. It follows the product is upper triangular since it is a sum of upper triangular matrices. The proof of (3.) is similar. $\square$.

I hope you can appreciate these arguments are superior to component level calculations with explicit listing of components and The notations e_i and E_{ij} are extremely helpful on many such questions. Furthermore, a proof captured in the notation of this section will more clearly show the root cause for the truth of the identity in question. What is easily lost in several pages of brute-force can be elegantly seen in a couple lines of carefully crafted index calculation.

1.4.5 nilpotent matrices

Definition 1.4.21.

Let $N \in R^{n \times n}$ be nonzero then N is **nilpotent** of degree k if k is the first positive integer for which $N^k = 0$.

Nilpotent matrices are easy to find. Here is an important example:

Example 1.4.22. Let $N = \begin{bmatrix} 0 & 1 & 0 \\ 0 & 0 & 1 \\ 0 & 0 & 0 \end{bmatrix}$ hence $N^2 = \begin{bmatrix} 0 & 0 & 1 \\ 0 & 0 & 0 \\ 0 & 0 & 0 \end{bmatrix}$ and $N^3 = 0$ thus N is nilpotent of degree 3. If $N \in R^{n \times n}$ and $N_{i,i+1} = 1$ for $i = 1, \dots, n-1$ and $N_{ij} = 0$ otherwise then we can show through similar calculation that $N^{n-1} = E_{1n}$ and $N^n = 0$.

A fun question to ponder: which of the matrix units are nilpotent? Moving on, another interesting aspect of a nilpotent matrix is that if we modify the identity matrix by N then it is still invertible. For example:

Example 1.4.23. Suppose N is nilpotent of degree 2 then $I + N$ has inverse $I - N$ as is easily seen by $(I + N)(I - N) = I + N - N - N^2 = I$ and $(I - N)(I + N) = I - N + N - N^2 = I$ hence $(I + N)^{-1} = I - N$.

The inverse in the example above is not hard to guess. Try out the next case, can you find $(I + N)^{-1}$ for N nilpotent of degree 3. As a formal intuition, you might think about the geometric series.

¹⁶the notation $\sum_{i \leq j}$ indicates we sum over all pairs i, j for which $i \leq j$. For example, if $n = 3$ then we sum over $(i, j) = (1, 1), (1, 2), (1, 3), (2, 2), (2, 3), (3, 3)$.

1.4.6 block matrices

If you look at most undergraduate linear algebra texts they will not bother to even attempt much of a proof that block-multiplication holds in general. I will foolishly attempt it here. However, I'm going to cheat a little and employ uber-sneaky physics notation.

The Einstein summation convention states that if an index is repeated then it is assumed to be summed over its values. This means that the letters used for particular indices are reserved. If i, j, k are used to denote components of a spatial vector then you cannot use them for a spacetime vector at the same time. A typical notation in physics would be that v^j is a vector in xyz -space whereas v^μ is a vector in $txyz$ -spacetime. A spacetime vector could be written as a sum of space components and a time component; $v = v^\mu e_\mu = v^0 e_0 + v^1 e_1 + v^2 e_2 + v^3 e_3 = v^0 e_0 + v^j e_j$. This is not the sort of language we tend to use in mathematics. For us notation is usually not reserved. Anyway, cultural commentary aside, if we were to use Einstein-type notation in linear algebra then we would likely omit sums as follows:

$$v = \sum_i v_i e_i \longrightarrow v = v_i e_i$$

$$A = \sum_{ij} A_{ij} E_{ij} \longrightarrow A = A_{ij} E_{ij}$$

We wish to partition a matrices A and B into 4 parts, use indices M, N which split into subindices m, μ and n, ν respectively. In this notation there are 4 different types of pairs possible:

$$A = [A_{MN}] = \left[\begin{array}{c|c} A_{mn} & A_{m\nu} \\ \hline A_{\mu n} & A_{\mu\nu} \end{array} \right] \quad B = [B_{NJ}] = \left[\begin{array}{c|c} B_{nj} & B_{n\gamma} \\ \hline B_{\nu j} & B_{\nu\gamma} \end{array} \right]$$

Then the sum over M, N breaks into 2 cases,

$$A_{MN} B_{NJ} = A_{Mn} B_{nJ} + A_{M\nu} B_{\nu J}$$

But, then there are 4 different types of M, J pairs,

$$[AB]_{mj} = A_{mN} B_{NJ} = A_{mn} B_{nj} + A_{m\nu} B_{\nu j}$$

$$[AB]_{m\gamma} = A_{mN} B_{N\gamma} = A_{mn} B_{n\gamma} + A_{m\nu} B_{\nu\gamma}$$

$$[AB]_{\mu j} = A_{\mu N} B_{NJ} = A_{\mu n} B_{nj} + A_{\mu\nu} B_{\nu j}$$

$$[AB]_{\mu\gamma} = A_{\mu N} B_{N\gamma} = A_{\mu n} B_{n\gamma} + A_{\mu\nu} B_{\nu\gamma}$$

Let me summarize,

$$\left[\begin{array}{c|c} A_{mn} & A_{m\nu} \\ \hline A_{\mu n} & A_{\mu\nu} \end{array} \right] \left[\begin{array}{c|c} B_{nj} & B_{n\gamma} \\ \hline B_{\nu j} & B_{\nu\gamma} \end{array} \right] = \left[\begin{array}{c|c} [A_{mn}][B_{nj}] + [A_{m\nu}][B_{\nu j}] & [A_{mn}][B_{n\gamma}] + [A_{m\nu}][B_{\nu\gamma}] \\ \hline [A_{\mu n}][B_{nj}] + [A_{\mu\nu}][B_{\nu j}] & [A_{\mu n}][B_{n\gamma}] + [A_{\mu\nu}][B_{\nu\gamma}] \end{array} \right]$$

Let me again summarize, but this time I'll drop the annoying indices:

Theorem 1.4.24. *block multiplication.*

Suppose $A \in \mathbb{R}^{m \times n}$ and $B \in \mathbb{R}^{n \times p}$ such that both A and B are partitioned as follows:

$$A = \left[\begin{array}{c|c} A_{11} & A_{12} \\ \hline A_{21} & A_{22} \end{array} \right] \quad \text{and} \quad B = \left[\begin{array}{c|c} B_{11} & B_{12} \\ \hline B_{21} & B_{22} \end{array} \right]$$

where A_{11} is an $m_1 \times n_1$ block, A_{12} is an $m_1 \times n_2$ block, A_{21} is an $m_2 \times n_1$ block and A_{22} is an $m_2 \times n_2$ block. Likewise, $B_{n_k p_k}$ is an $n_k \times p_k$ block for $k = 1, 2$. We insist that $m_1 + m_2 = m$ and $n_1 + n_2 = n$. If the partitions are compatible as described above then we may multiply A and B by multiplying the blocks as if they were scalars and we were computing the product of 2×2 matrices:

$$\left[\begin{array}{c|c} A_{11} & A_{12} \\ \hline A_{21} & A_{22} \end{array} \right] \left[\begin{array}{c|c} B_{11} & B_{12} \\ \hline B_{21} & B_{22} \end{array} \right] = \left[\begin{array}{c|c} A_{11}B_{11} + A_{12}B_{21} & A_{11}B_{12} + A_{12}B_{22} \\ \hline A_{21}B_{11} + A_{22}B_{21} & A_{21}B_{12} + A_{22}B_{22} \end{array} \right].$$

To give a careful proof we'd just need to write out many sums and define the partition with care from the outset of the proof. In any event, notice that once you have this partition you can apply it twice to build block-multiplication rules for matrices with more blocks. The basic idea remains the same: you can parse two matrices into matching partitions then the matrix multiplication follows a pattern which is as if the blocks were scalars. However, the blocks are not scalars so the multiplication of the blocks is nonabelian. For example,

$$AB = \left[\begin{array}{c|c} A_{11} & A_{12} \\ \hline A_{21} & A_{22} \\ \hline A_{31} & A_{32} \end{array} \right] \left[\begin{array}{c|c} B_{11} & B_{12} \\ \hline B_{21} & B_{22} \end{array} \right] = \left[\begin{array}{c|c} A_{11}B_{11} + A_{12}B_{21} & A_{11}B_{12} + A_{12}B_{22} \\ \hline A_{21}B_{11} + A_{22}B_{21} & A_{21}B_{12} + A_{22}B_{22} \\ \hline A_{31}B_{11} + A_{32}B_{21} & A_{31}B_{12} + A_{32}B_{22} \end{array} \right].$$

where if the partitions of A and B are compatible it follows that the block-multiplications on the RHS are all well-defined.

Example 1.4.25. Let $R(\theta) = \begin{bmatrix} \cos(\theta) & \sin(\theta) \\ -\sin(\theta) & \cos(\theta) \end{bmatrix}$ and $B(\gamma) = \begin{bmatrix} \cosh(\gamma) & \sinh(\gamma) \\ \sinh(\gamma) & \cosh(\gamma) \end{bmatrix}$. Further-
more construct 4×4 matrices Λ_1 and Λ_2 as follows:

$$\Lambda_1 = \left[\begin{array}{c|c} B(\gamma_1) & 0 \\ \hline 0 & R(\theta_1) \end{array} \right] \quad \Lambda_2 = \left[\begin{array}{c|c} B(\gamma_2) & 0 \\ \hline 0 & R(\theta_2) \end{array} \right]$$

Multiply Λ_1 and Λ_2 via block multiplication:

$$\begin{aligned} \Lambda_1 \Lambda_2 &= \left[\begin{array}{c|c} B(\gamma_1) & 0 \\ \hline 0 & R(\theta_1) \end{array} \right] \left[\begin{array}{c|c} B(\gamma_2) & 0 \\ \hline 0 & R(\theta_2) \end{array} \right] \\ &= \left[\begin{array}{c|c} B(\gamma_1)B(\gamma_2) + 0 & 0 + 0 \\ \hline 0 + 0 & 0 + R(\theta_1)R(\theta_2) \end{array} \right] \\ &= \left[\begin{array}{c|c} B(\gamma_1 + \gamma_2) & 0 \\ \hline 0 & R(\theta_1 + \theta_2) \end{array} \right]. \end{aligned}$$

The last calculation is actually a few lines in detail, if you know the adding angles formulas for cosine, sine, cosh and sinh it's easy. If $\theta = 0$ and $\gamma \neq 0$ then Λ would represent a **velocity boost** on spacetime. Since it mixes time and the first coordinate the velocity is along the x -coordinate. On the other hand, if $\theta \neq 0$ and $\gamma = 0$ then Λ gives a **rotation** in the yz spatial coordinates in space

time. If both parameters are nonzero then we can say that Λ is a **Lorentz transformation** on spacetime. Of course there is more to say here, perhaps we could offer a course in special relativity if enough students were interested in concert.

Example 1.4.26. Problem: Suppose M is a square matrix with submatrices $A, B, C, 0$ where A, C are square. What conditions should we insist on for $M = \left[\begin{array}{c|c} A & B \\ \hline 0 & C \end{array} \right]$ to be invertible.

Solution: I propose we partition the potential inverse matrix $M^{-1} = \left[\begin{array}{c|c} D & E \\ \hline F & G \end{array} \right]$. We seek to find conditions on A, B, C such that there exist D, E, F, G and $MM^{-1} = I$. Each block of the equation $MM^{-1} = I$ gives us a separate submatrix equation:

$$MM^{-1} = \left[\begin{array}{c|c} A & B \\ \hline 0 & C \end{array} \right] \left[\begin{array}{c|c} D & E \\ \hline F & G \end{array} \right] = \left[\begin{array}{c|c} AD + BF & AE + BG \\ \hline 0D + CF & 0E + CG \end{array} \right] = \left[\begin{array}{c|c} I & 0 \\ \hline 0 & I \end{array} \right]$$

We must solve simultaneously the following:

$$(1.) AD + BF = I, \quad (2.) AE + BG = 0, \quad (3.) CF = 0, \quad (4.) CG = I$$

If C^{-1} exists then $G = C^{-1}$ from (4.). Moreover, (3.) then yields $F = C^{-1}0 = 0$. Our problem thus reduces to (1.) and (2.) which after substituting $F = 0$ and $G = C^{-1}$ yield

$$(1.) AD = I, \quad (2.) AE + BC^{-1} = 0.$$

Equation (1.) says $D = A^{-1}$. Finally, let's solve (2.) for E ,

$$E = -A^{-1}BC^{-1}.$$

Let's summarize the calculation we just worked through. If A, C are invertible then the matrix $M = \left[\begin{array}{c|c} A & B \\ \hline 0 & C \end{array} \right]$ is invertible with inverse

$$M^{-1} = \left[\begin{array}{c|c} A^{-1} & -A^{-1}BC^{-1} \\ \hline 0 & C^{-1} \end{array} \right].$$

Consider the case that M is a 2×2 matrix and $A, B, C \in \mathbb{R}$. Then the condition of invertibility reduces to the simple conditions $A, C \neq 0$ and $-A^{-1}BC^{-1} = \frac{-B}{AC}$ we find the formula:

$$M^{-1} = \left[\begin{array}{c|c} \frac{1}{A} & \frac{-B}{AC} \\ \hline 0 & \frac{1}{C} \end{array} \right] = \frac{1}{AC} \left[\begin{array}{c|c} C & -B \\ \hline 0 & A \end{array} \right].$$

This is of course the formula for the 2×2 matrix in this special case where $M_{21} = 0$.

Of course the real utility of formulas like those in the last example is that they work for partitions of arbitrary size. If we can find a block of zeros somewhere in the matrix then we may reduce the size of the problem. The time for a computer calculation is largely based on some power of the size of the matrix. For example, if the calculation in question takes n^2 steps then parsing the matrix into 3 nonzero blocks which are $n/2 \times n/2$ would result in something like $[n/2]^2 + [n/2]^2 + [n/2]^2 = \frac{3}{4}n^2$ steps. If the calculation took on order n^3 computer operations (flops) then my toy example of 3 blocks would reduce to something like $[n/2]^3 + [n/2]^3 + [n/2]^3 = \frac{3}{8}n^3$ flops. A savings of more than 60% of computer time. If the calculation was typically order n^4 for an $n \times n$ matrix then the saving

is even more dramatic. If the calculation is a determinant then the cofactor formula depends on the factorial of the size of the matrix. Try to compare $10!+10!$ verses say $20!$. Hope your calculator has a big display:

$$10! = 3628800 \Rightarrow 10! + 10! = 7257600 \quad \text{or} \quad 20! = 2432902008176640000.$$

Perhaps you can start to appreciate why numerical linear algebra software packages often use algorithms which make use of block matrices to streamline large matrix calculations.

In quantum mechanics, it is good to find a basis of state vectors which makes the Hamiltonian into a block-diagonal matrix. Each block corresponds to a certain set of statevectors sharing a common energy. The goal of representation theory in physics is basically to break down matrices into blocks with nice physical meanings. On the other hand, abstract algebraists also use blocks to rip apart a matrix into it's most basic form. For linear algebraists, the so-called Jordan form is full of blocks. Wherever reduction of a linear system into smaller subsystems is of interest there will be blocks.

1.5 systems of linear equations

We now introduce some notation which will help bring organization to our method of solving linear systems: we assume $\mathbb{F}$ is a field in what follows:

Definition 1.5.1. *augmented coefficient matrix for m -equations in n -unknowns*

The augmented coefficient matrix is an array of numbers which provides an abbreviated notation for a system of linear equations.

$$\left[\begin{array}{cccc|c} A_{11}x_1 + A_{12}x_2 + \cdots + A_{1n}x_n = b_1 \\ A_{21}x_1 + A_{22}x_2 + \cdots + A_{2n}x_n = b_2 \\ \vdots \quad \vdots \quad \vdots \quad \vdots \quad \vdots \\ A_{m1}x_1 + A_{m2}x_2 + \cdots + A_{mn}x_n = b_m \end{array} \right] \quad \text{replaced by} \quad \left[\begin{array}{cccc|c} A_{11} & A_{12} & \cdots & A_{1n} & b_1 \\ A_{21} & A_{22} & \cdots & A_{2n} & b_2 \\ \vdots & \vdots & \vdots & \vdots & \vdots \\ A_{m1} & A_{m2} & \cdots & A_{mn} & b_m \end{array} \right].$$

We say $A = [A_{ij}]$ is the **coefficient matrix** of the system and $b = [b_i]$ is the inhomogenous term. When $b = 0$ the system is **homogeneous**. The system of equations can also be expressed as a **matrix equation** $Ax = b$. The **solution set** of the system is the set of all $x \in \mathbb{F}^n$ for which $Ax = b$. If the solution set is empty then the system is said to be **inconsistent**. If there exists a solution to $Ax = b$ then the system is **consistent**.

The vertical bar is optional, I include it to draw attention to the distinction between the matrix of coefficients A_{ij} and the nonhomogeneous terms b_i . Let us briefly review the method of solving systems of equation via row-reduction.

Definition 1.5.2. Elementary Row operations: Let $A \in \mathbb{F}^{m \times n}$ we define

	Effect on the linear system:		Effect on the matrix:
Type I	Interchange equation i and equation j (List the equations in a different order.)	$\iff$	Swap Row i and Row j
Type II	Multiply both sides of equation i by a non-zero scalar c	$\iff$	Multiply Row i by c where $c \neq 0$
Type III	Multiply both sides of equation i by c and add to equation j	$\iff$	Add c times Row i to Row j where c is any scalar

If we can get matrix A from matrix B by performing a series of elementary row operations, then A and B are called **row equivalent matrices**.

Of course, there are also corresponding *elementary column operations*. If we can get matrix A from matrix B by performing a series of elementary column operations, we call A and B **column equivalent matrices**. Both of these *equivalences* are in fact equivalence relations¹⁷. While both row and column operations are important, we will (for now) focus on row operations since they correspond to steps used when solving linear systems.

The **Gauss-Jordan Elimination** is an “algorithm” which given a matrix returns a row equivalent matrix in reduced row echelon form (RREF). Let us give a precise account of this algorithm:

Definition 1.5.3. *A matrix is in **Row Echelon Form** (or **REF**) if...*

Given a matrix over a field $\mathbb{F}$. We first perform a **forward pass**:

- (1.) Determine the leftmost non-zero column. This is a **pivot column** and the topmost entry is a **pivot position**. If “0” is in this pivot position, swap (an unignored) row with the topmost row (use a Type I operation) so that there is a non-zero entry in the pivot position.
- (2.) Add appropriate multiples of the topmost (unignored) row to the rows beneath it so that only “0” appears below the pivot (use several Type III operations).
- (3.) Ignore the topmost (unignored) row. If any non-zero rows remain, go to step 1.

The forward pass is now complete, such matrices can be denoted $ref(A)$. Now let’s finish Gauss-Jordan Elimination by performing a **backward pass**:

- (1.) If necessary, scale the rightmost unfinished pivot to 1 (use a Type II operation).
- (2.) Add appropriate multiples of the current pivot’s row to rows above it so that only 0 appears above the current pivot (using several Type III operations).
- (3.) The current pivot is now “finished”. If any unfinished pivots remain, go to step 4.
- (4.) Let the matrix you have obtained is denoted $rref(A)$, this is short-hand for the **reduced row echelon form** of A .

Proof that $rref(A)$ is unique can be found in many texts, I hope the reader will forgive me for omitting such proof here.

One advice to always keep in mind, you should think of the Gauss-Jordan algorithm as a sort of road-map. It’s ok to take detours to avoid fractions and such but the end goal should remain in sight. Also, keep in mind, it is sometimes far simpler to simply add equations directly or make strategic substitutions instead of robotically following the algorithm. For homework, it is important to use technology to check your work on row-reductions. I offer a number of examples, but most of this should be a review from the previous course.

¹⁷recall an equivalence relation on a set is a relation which is reflexive, symmetric and transitive.

Example 1.5.4. The equations $x + 2y - 3z = 1$, $2x + 4y = 7$ and $-x + 3y + 2z = 0$ can be solved by row operations on the matrix $[A|b]$ below: Given $[A|b] = \left[\begin{array}{ccc|c} 1 & 2 & -3 & 1 \\ 2 & 4 & 0 & 7 \\ -1 & 3 & 2 & 0 \end{array} \right]$ calculate $\text{rref}(A|b)$.

$$[A|b] = \left[\begin{array}{ccc|c} 1 & 2 & -3 & 1 \\ 2 & 4 & 0 & 7 \\ -1 & 3 & 2 & 0 \end{array} \right] \xrightarrow{r_2 - 2r_1} \left[\begin{array}{ccc|c} 1 & 2 & -3 & 1 \\ 0 & 0 & 6 & 5 \\ -1 & 3 & 2 & 0 \end{array} \right] \xrightarrow{r_1 + r_3} \left[\begin{array}{ccc|c} 1 & 2 & -3 & 1 \\ 0 & 0 & 6 & 5 \\ 0 & 5 & -1 & 1 \end{array} \right] \xrightarrow{r_2 \leftrightarrow r_3} \left[\begin{array}{ccc|c} 1 & 2 & -3 & 1 \\ 0 & 5 & -1 & 1 \\ 0 & 0 & 6 & 5 \end{array} \right] = \text{ref}[A|b]$$

that completes the forward pass. We begin the backwards pass,

$$\begin{aligned} \text{ref}[A|b] &= \left[\begin{array}{ccc|c} 1 & 2 & -3 & 1 \\ 0 & 5 & -1 & 1 \\ 0 & 0 & 6 & 5 \end{array} \right] \xrightarrow{\frac{1}{6}r_3} \left[\begin{array}{ccc|c} 1 & 2 & -3 & 1 \\ 0 & 5 & -1 & 1 \\ 0 & 0 & 1 & 5/6 \end{array} \right] \xrightarrow{r_2 + r_3} \left[\begin{array}{ccc|c} 1 & 2 & -3 & 1 \\ 0 & 5 & 0 & 11/6 \\ 0 & 0 & 1 & 5/6 \end{array} \right] \xrightarrow{r_1 + 3r_3} \left[\begin{array}{ccc|c} 1 & 2 & 0 & 21/6 \\ 0 & 5 & 0 & 11/6 \\ 0 & 0 & 1 & 5/6 \end{array} \right] \xrightarrow{\frac{1}{5}r_2} \left[\begin{array}{ccc|c} 1 & 2 & 0 & 21/6 \\ 0 & 1 & 0 & 11/30 \\ 0 & 0 & 1 & 5/6 \end{array} \right] \xrightarrow{r_1 - 2r_2} \left[\begin{array}{ccc|c} 1 & 0 & 0 & 83/30 \\ 0 & 1 & 0 & 11/30 \\ 0 & 0 & 1 & 5/6 \end{array} \right] = \text{rref}(A) \end{aligned}$$

Thus, we've found the system of equations $x + 2y - 3z = 1$, $2x + 4y = 7$ and $-x + 3y + 2z = 0$ has solution $x = 83/30$, $y = 11/30$ and $z = 5/6$. This means the matrix equation $Ax = b$ where

$$Ax = \underbrace{\begin{bmatrix} 1 & 2 & -3 \\ 2 & 4 & 0 \\ -1 & 3 & 2 \end{bmatrix}}_A \underbrace{\begin{bmatrix} x_1 \\ x_2 \\ x_3 \end{bmatrix}}_x = \underbrace{\begin{bmatrix} 1 \\ 7 \\ 0 \end{bmatrix}}_b \quad \text{has vector solution} \quad x = \begin{bmatrix} 83/30 \\ 11/30 \\ 5/6 \end{bmatrix}.$$

Remark 1.5.5.

The geometric interpretation of the last example is interesting. The equation of a plane with normal vector $\langle a, b, c \rangle$ is $ax + by + cz = d$. Each of the equations in the system of Example 1.5.4 has a solution set which is in one-one correspondence with a particular plane in $\mathbb{R}^3$. The intersection of those three planes is the single point $(83/30, 11/30, 5/6)$.

Example 1.5.6. Solve the following system of real linear equations if possible,

$$\begin{aligned} x - y &= 1 \\ 3x - 3y &= 0 \\ 2x - 2y &= -3 \end{aligned}$$

Calculate,

$$\begin{aligned}
 [A|b] &= \left[\begin{array}{cc|c} 1 & -1 & 1 \\ 3 & -3 & 0 \\ 2 & -2 & -3 \end{array} \right] \xrightarrow{r_2 - 3r_1} \left[\begin{array}{cc|c} 1 & -1 & 1 \\ 0 & 0 & -3 \\ 2 & -2 & -3 \end{array} \right] \xrightarrow{r_3 - 2r_1} \\
 &\left[\begin{array}{cc|c} 1 & -1 & 1 \\ 0 & 0 & -3 \\ 0 & 0 & -5 \end{array} \right] \xrightarrow{\substack{3r_3 \\ 5r_2}} \left[\begin{array}{cc|c} 1 & -1 & 1 \\ 0 & 0 & -15 \\ 0 & 0 & -15 \end{array} \right] \xrightarrow{\substack{r_3 - r_2 \\ -\frac{1}{15}r_2}} \\
 &\left[\begin{array}{cc|c} 1 & -1 & 1 \\ 0 & 0 & 1 \\ 0 & 0 & 0 \end{array} \right] \xrightarrow{r_1 - r_2} \boxed{\left[\begin{array}{cc|c} 1 & -1 & 0 \\ 0 & 0 & 1 \\ 0 & 0 & 0 \end{array} \right]} = \text{rref}(A|b)
 \end{aligned}$$

which shows the system has no solutions since row two in the rref corresponds to the equation $0x + 0y = 1$. The given equations are inconsistent.

Example 1.5.7. Solve the following system of real linear equations if possible,

$$\begin{aligned}
 x - y + z &= 0 \\
 3x - 3y &= 0 \\
 2x - 2y - 3z &= 0
 \end{aligned}$$

Gaussian elimination on the augmented coefficient matrix reveals

$$\text{rref} \left[\begin{array}{ccc|c} 1 & -1 & 1 & 0 \\ 3 & -3 & 0 & 0 \\ 2 & -2 & -3 & 0 \end{array} \right] = \left[\begin{array}{ccc|c} 1 & -1 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 \end{array} \right] \Rightarrow \boxed{\begin{array}{l} x - y = 0 \\ z = 0 \end{array}}$$

The row of zeros indicates that we will not find a unique solution. We have a choice to make, either x or y can be stated as a function of the other. Typically in linear algebra we will solve for the variables that correspond to the pivot columns in terms of the non-pivot column variables. In this problem the pivot columns are the first column which corresponds to the variable x and the third column which corresponds the variable z . The variables x, z are called **dependent variables** while y is called a **free variable**¹⁸. The solution set is $\{(y, y, 0) \mid y \in \mathbb{R}\}$; in other words, $x = y, y = y$ and $z = 0$ for all $y \in \mathbb{R}$.

You might object to the last example. You might ask why is y the free variable and not x . This is roughly equivalent to asking the question why is y the dependent variable and x the independent variable in the usual calculus. However, the roles are reversed. In the preceding example the variable x depends on y . Physically there may be a reason to distinguish the roles of one variable over another. There may be a clear cause-effect relationship which the mathematics fails to capture. For example, velocity of a ball in flight depends on time, but does time depend on the ball's velocity? I'm guessing no. So time would seem to play the role of independent variable. However, when we write equations such as $v = v_o - gt$ we can just as well write $t = \frac{v-v_o}{-g}$; the algebra alone does not reveal which variable should be taken as "independent". Hence, a choice must be made. In the case of infinitely many solutions, we customarily **choose** the pivot variables as the "dependent" or

¹⁸the choice of free and dependent variables is suggested by the pivot positions, however, there may also be other reasonable choices

"basic" variables and the non-pivot variables as the "free" variables. Sometimes the word *parameter* is used instead of variable, it is synonymous.

We can calculate the rref of a matrix even when no equations are given.

Example 1.5.8. Find the rref of the matrix $A \in \mathbb{F}^{3 \times 5}$ given below:

$$\begin{aligned}
 A &= \begin{bmatrix} 1 & 1 & 1 & 1 & 1 \\ 1 & -1 & 1 & 0 & 1 \\ -1 & 0 & 1 & 1 & 1 \end{bmatrix} \xrightarrow{r_2 - r_1} \begin{bmatrix} 1 & 1 & 1 & 1 & 1 \\ 0 & -2 & 0 & -1 & 0 \\ -1 & 0 & 1 & 1 & 1 \end{bmatrix} \xrightarrow{r_3 + r_1} \\
 &\begin{bmatrix} 1 & 1 & 1 & 1 & 1 \\ 0 & -2 & 0 & -1 & 0 \\ 0 & 1 & 2 & 2 & 2 \end{bmatrix} \xrightarrow{r_2 / -2} \begin{bmatrix} 1 & 1 & 1 & 1 & 1 \\ 0 & 1 & 0 & 1/2 & 0 \\ 0 & 1 & 2 & 2 & 2 \end{bmatrix} \xrightarrow{r_3 - r_2} \\
 &\begin{bmatrix} 1 & 1 & 1 & 1 & 1 \\ 0 & 1 & 0 & 1/2 & 0 \\ 0 & 0 & 2 & 3/2 & 2 \end{bmatrix} \xrightarrow{r_3/2} \begin{bmatrix} 1 & 1 & 1 & 1 & 1 \\ 0 & 1 & 0 & 1/2 & 0 \\ 0 & 0 & 1 & 3/4 & 1 \end{bmatrix} \xrightarrow{r_1 - r_3} \\
 &\begin{bmatrix} 1 & 1 & 0 & 1/4 & 0 \\ 0 & 1 & 0 & 1/2 & 0 \\ 0 & 0 & 1 & 3/4 & 1 \end{bmatrix} \xrightarrow{r_1 - r_2} \begin{bmatrix} 1 & 0 & 0 & -1/4 & 0 \\ 0 & 1 & 0 & 1/2 & 0 \\ 0 & 0 & 1 & 3/4 & 1 \end{bmatrix} = \text{rref}(A)
 \end{aligned}$$

The equation $Ax = 0$ can be solved via the reduction above since row operations act column by column. With no further calculation, we note: $\text{rref}[A|0] = \left[\begin{array}{ccccc|c} 1 & 0 & 0 & -1/4 & 0 & 0 \\ 0 & 1 & 0 & 1/2 & 0 & 0 \\ 0 & 0 & 1 & 3/4 & 1 & 0 \end{array} \right]$. Therefore, we find solution set¹⁹: $\{(x_4/4, -x_4/2, -3x_4/4 - x_5, x_4, x_5) \mid x_4, x_5 \in \mathbb{F}\}$.

Example 1.5.9. We can rewrite the following system of linear equations

$$\begin{aligned}
 x_1 + x_4 &= 0 \\
 2x_1 + 2x_2 + x_5 &= 0 \\
 4x_1 + 4x_2 + 4x_3 &= 1
 \end{aligned}$$

in matrix form this system of equations is $Av = b$ where

$$Av = \underbrace{\begin{bmatrix} 1 & 0 & 0 & 1 & 0 \\ 2 & 2 & 0 & 0 & 1 \\ 4 & 4 & 4 & 0 & 0 \end{bmatrix}}_A \underbrace{\begin{bmatrix} x_1 \\ x_2 \\ x_3 \\ x_4 \\ x_5 \end{bmatrix}}_v = \underbrace{\begin{bmatrix} 0 \\ 0 \\ 1 \end{bmatrix}}_b.$$

Gaussian elimination on the augmented coefficient matrix reveals

$$\text{rref} \left[\begin{array}{ccccc|c} 1 & 0 & 0 & 1 & 0 & 0 \\ 2 & 2 & 0 & 0 & 1 & 0 \\ 4 & 4 & 4 & 0 & 0 & 1 \end{array} \right] = \left[\begin{array}{ccccc|c} 1 & 0 & 0 & 1 & 0 & 0 \\ 0 & 1 & 0 & -1 & 1/2 & 0 \\ 0 & 0 & 1 & 0 & -1/2 & 1/4 \end{array} \right].$$

¹⁹In the span notation, $\text{Span}_{\mathbb{F}}\{(1, -2, -3, 1, 0), (0, 0, -1, 0, 1)\}$ and you might recall this calculation as **finding the basis for the nullspace of A** . We discuss this further in a more abstract context a bit later in the course.

Consequently, x_4, x_5 are free and solutions are of the form

$$\begin{aligned} x_1 &= -x_4 \\ x_2 &= x_4 - \frac{1}{2}x_5 \\ x_3 &= \frac{1}{4} + \frac{1}{2}x_5 \end{aligned}$$

for all $x_4, x_5 \in \mathbb{R}$. The vector form of the solution is as follows:

$$v = \begin{bmatrix} -x_4 \\ x_4 - \frac{1}{2}x_5 \\ \frac{1}{4} + \frac{1}{2}x_5 \\ x_4 \\ x_5 \end{bmatrix} = x_4 \begin{bmatrix} -1 \\ 1 \\ 0 \\ 1 \\ 0 \end{bmatrix} + x_5 \begin{bmatrix} 0 \\ -\frac{1}{2} \\ \frac{1}{2} \\ 0 \\ 1 \end{bmatrix} + \begin{bmatrix} 0 \\ 0 \\ \frac{1}{4} \\ 0 \\ 0 \end{bmatrix}.$$

Remark 1.5.10.

You might ask the question: what is the geometry of the solution set above? Let $S = \text{Sol}_{[A|b]} \subset \mathbb{R}^5$, we see S is formed by tracing out all possible linear combinations of the vectors $v_1 = (-1, 1, 0, 1, 0)$ and $v_2 = (0, -\frac{1}{2}, \frac{1}{2}, 0, 1)$ based from the point $p_o = (0, 0, \frac{1}{4}, 0, 0)$. In other words, this is a two-dimensional plane containing the vectors v_1, v_2 and the point p_o . This plane is placed in a 5-dimensional space, this means that at any point on the plane you could go in three different directions away from the plane.

The examples which follow are probably not a review for the students of Math 321. Gaussian elimination over $\mathbb{F}$ when $\mathbb{F} \neq \mathbb{R}, \mathbb{Q}, \mathbb{C}$ follows the same rules. The main difference is we have to keep in mind the rules for arithmetic for the finite field in question. See Appendix Chapter 7 for a brief introduction.

Example 1.5.11. In $\mathbb{Z}/2\mathbb{Z}$ we calculate (See Chapter 7 if you seek background):

$$A = \begin{bmatrix} 1 & 1 & 1 & 0 \\ 1 & 1 & 0 & 1 \\ 1 & 0 & 1 & 0 \end{bmatrix} \xrightarrow[r_3 + r_1]{r_2 + r_1} \begin{bmatrix} 1 & 1 & 1 & 0 \\ 0 & 0 & 1 & 1 \\ 0 & 1 & 0 & 0 \end{bmatrix} \xrightarrow{r_1 + r_3} \begin{bmatrix} 1 & 0 & 1 & 0 \\ 0 & 0 & 1 & 1 \\ 0 & 1 & 0 & 0 \end{bmatrix} \xrightarrow{r_1 + r_2} \begin{bmatrix} 1 & 0 & 0 & 1 \\ 0 & 0 & 1 & 1 \\ 0 & 1 & 0 & 0 \end{bmatrix}$$

then swap rows 2 and 3 to obtain

$$\text{rref}(A) = \begin{bmatrix} 1 & 0 & 0 & 1 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 1 \end{bmatrix}.$$

Example 1.5.12. In this example we calculate in $\mathbb{Z}/5\mathbb{Z}$ (See Chapter 7 explains this in depth):

$$\begin{aligned} A = \begin{bmatrix} 2 & 1 & 1 & 0 \\ 1 & 4 & 0 & 3 \\ 3 & 1 & 1 & 1 \end{bmatrix} &\xrightarrow{r_1 \leftrightarrow r_2} \begin{bmatrix} 1 & 4 & 0 & 3 \\ 2 & 1 & 1 & 0 \\ 3 & 1 & 1 & 1 \end{bmatrix} \xrightarrow[r_3 - 3r_1]{r_2 - 2r_1} \begin{bmatrix} 1 & 4 & 0 & 3 \\ 0 & -7 & 1 & -6 \\ 0 & -11 & 1 & -8 \end{bmatrix} = \\ &\begin{bmatrix} 1 & 4 & 0 & 3 \\ 0 & 3 & 1 & -1 \\ 0 & -1 & 1 & 2 \end{bmatrix} \xrightarrow{2r_2} \begin{bmatrix} 1 & 4 & 0 & 3 \\ 0 & 1 & 2 & -2 \\ 0 & -1 & 1 & 2 \end{bmatrix} \xrightarrow[r_3 + r_2]{r_1 + r_2} \begin{bmatrix} 1 & 0 & 2 & 1 \\ 0 & 1 & 2 & -2 \\ 0 & 0 & 3 & 0 \end{bmatrix} \xrightarrow{2r_3} \\ &\begin{bmatrix} 1 & 0 & 2 & 1 \\ 0 & 1 & 2 & -2 \\ 0 & 0 & 1 & 0 \end{bmatrix} \xrightarrow[r_2 - 2r_3]{r_1 - 2r_3} \begin{bmatrix} 1 & 0 & 0 & 1 \\ 0 & 1 & 0 & -2 \\ 0 & 0 & 1 & 0 \end{bmatrix} = \boxed{\begin{bmatrix} 1 & 0 & 0 & 1 \\ 0 & 1 & 0 & 3 \\ 0 & 0 & 1 & 0 \end{bmatrix}} = \text{rref}(A). \end{aligned}$$

Theorem 1.5.13.

Given a system of m linear equations and n unknowns over an infinite field, the solution set falls into one of the following cases:

- (i.) the solution set is empty.
- (ii.) the solution set has only one element.
- (iii.) the solution set is infinite and is parametrized by $(n - k)$ -parameters where k is the number of pivot columns in the reduced row echelon form of the augmented coefficient matrix for the system.

Proof: Consider the augmented coefficient matrix $[A|b] \in \mathbb{F}^{m \times (n+1)}$ for the given system of m -linear equations in n -unknowns over the infinite field $\mathbb{F}$. Apply the Gauss-Jordan Algorithm to compute $rref[A|b]$ and consider the possible cases:

If $rref[A|b]$ contains a row of zeros with a 1 in the last column then the system is inconsistent and we find no solutions thus the solution set is empty. This brings us to case (i.).

Suppose $rref[A|b]$ does not contain a row of zeros with a 1 in the far right position. Then there are less than $n + 1$ pivot columns and we may break into two possible subcases:

- (a.) Suppose there are n pivot columns, let c_i for $i = 1, 2, \dots, m$ be the entries in the rightmost column. We find $x_1 = c_1, x_2 = c_2, \dots, x_n = c_m$. Consequently the solution set is $\{(c_1, c_2, \dots, c_m)\}$ which we identify as case (ii.).
- (b.) If $rref[A|b]$ has $k < n$ pivot columns then there are $(n + 1 - k)$ -non-pivot positions. Since the last column corresponds to b it follows there are $n - k \geq 1$ free variables. Examining $rref[A|b]$ we find the k -pivot variables can be written as affine linear combinations of the k -free variables. In short, the solution set is parametrized by the $(n - k)$ -free variables and since $n - k \geq 1$ and each free variable takes as many values as $\mathbb{F}$ we find the cardinality of the solution set is infinite.

Naturally, the last case considered provides case (iii.) and the proof is complete $\square$

In the case of a finite field we find a very similar theorem. The proof is nearly the same so we omit all but the most interesting detail.

Theorem 1.5.14.

Given a system of m linear equations and n unknowns over a finite field with P elements, the solution set falls into one of the following cases:

- (i.) the solution set is empty.
- (ii.) the solution set has P^{n-k} solutions which are parametrized by $n - k$ -parameters where k is the number of pivot columns in the reduced row echelon form of the augmented coefficient matrix for the system.

Proof: If there are $k = n$ pivot columns then we find a unique solution and this is consistent with the formula $P^{n-n} = P^0 = 1$. On the other hand, in the case the system is consistent and there are $k < n$ pivot columns there are $n - k$ free variables. Each free variable ranges over the P elements of $\mathbb{F}$ hence there are P^{n-k} possible solutions. $\square$

Example 1.5.15. To give an easy and possibly interesting example, consider $x - y - z = 1$ in $\mathbb{Z}/2\mathbb{Z}$ the finite field with 2 elements. Observe y, z serve as parameters of the solution set as $x = 1 + y + z$:

$$\text{solution set} = \{(1 + y + z, y, z) \mid y, z \in \mathbb{Z}/2\mathbb{Z}\}$$

To be explicit, the solution set has $2^2 = 4$ solutions:

$$\text{solution set} = \{(1, 0, 0), (0, 0, 1), (0, 1, 0), (1, 1, 1)\}.$$

If this system was given over $\mathbb{Z}/3\mathbb{Z}$ then we would find $3^2 = 9$ solutions.

1.5.1 superposition of solutions

Theorem 1.5.16. *Superposition of solutions:*

Let $A \in \mathbb{F}^{m \times n}$. Let $c_1, c_2 \in \mathbb{F}$ and suppose there exist $x_1, x_2 \in \mathbb{F}^n$ for which $Ax_1 = b_1$ and $Ax_2 = b_2$ then $x = c_1x_1 + c_2x_2$ is a solution of $Ax = c_1b_1 + c_2b_2$. In particular, if $Ax_1 = 0$ and $Ax_2 = 0$ then $c_1x_1 + c_2x_2$ is a solution to $Ax = 0$.

Proof: with x_1, x_2 as in the Theorem we note:

$$A(c_1x_1 + c_2x_2) = c_1Ax_1 + c_2Ax_2 = c_1b_1 + c_2b_2.$$

The homogeneous case follows from setting $b_1 = b_2 = 0$. $\square$

Example 1.5.17. If we have two nonhomogeneous solutions of the same linear system then it is easy to generate a homogeneous solution. To see this, suppose $Ax_1 = b$ and $Ax_2 = b$. Notice $A(x_2 - x_1) = Ax_2 - Ax_1 = b - b = 0$ thus $x_2 - x_1$ is a solution of $Ax = 0$.

The example above is interesting for physical systems. We can subject a given linear system to two known forces and from the difference in the response functions it is possible to determine the intrinsic character of the system in the absense of external force. In a linear system, the net-response is a superposition of the responses to each source driving the system.

Theorem 1.5.18. *General solution is sum of particular and homogeneous solutions.*

Let $A \in \mathbb{F}^{m \times n}$ and $b \neq 0$. Suppose $Ax_p = b$ for some $x_p \neq 0$ then any solution of $Ax = b$ has the form $x = x_h + x_p$ where $Ax_h = 0$.

Proof: if $Ax_p = b$ and $Ax = b$ then $A(x - x_p) = Ax - Ax_p = b - b = 0$ hence $x_h = x - x_p$ has $Ax_h = 0$ and we conclude $x = x_p + x_h$ as desired. $\square$

If $Ax_h = 0$ then $x_h \in \text{Null}(A)$. The dimension of the nullspace of a matrix is called its **nullity**. The nullity is the number of non-pivot columns in A .

1.6 elementary matrices

1.6.1 definition of elementary matrices

Since we already know about matrix multiplication, it is interesting to note that an elementary row operation performed on A can be accomplished by multiplying A on the *left* by a square matrix (called an elementary matrix)²⁰. Likewise, multiplying A on the *right* by an elementary matrix performs a column operation. We assume $\mathbb{F}$ is a field in what follows.

Recall, $e_i = (0, \dots, 1, \dots, 0) \in \mathbb{F}^n$ where the non-zero entry is located in the i -position²¹. For example, $e_1 = \begin{bmatrix} 1 \\ 0 \end{bmatrix}$ and $e_2 = \begin{bmatrix} 0 \\ 1 \end{bmatrix}$ in $\mathbb{F}^2$. We saw in Subsection 1.3.2 that we can construct the identity matrix and the standard basis in $\mathbb{F}^{n \times n}$ from the standard basis vectors $e_i \in \mathbb{F}^n$:

$$I = [e_1 | e_2 | \dots | e_n] \quad \& \quad E_{ij} = e_i e_j^T = [0, \dots, 0, e_i, 0, \dots, 0]$$

The matrix E_{ij} has a 1 in the (i, j) -position and 0's elsewhere. An illustration from the 2×2 case:

Example 1.6.1. In $\mathbb{F}^{2 \times 2}$, $E_{21} = e_2 e_1^T = \begin{bmatrix} 0 \\ 1 \end{bmatrix} \begin{bmatrix} 1 & 0 \end{bmatrix} = \begin{bmatrix} 0 & 0 \\ 1 & 0 \end{bmatrix}$

In Example 1.3.19 we learned multiplication by e_j on the right of A produces the j -th column of A . Likewise, in Example 1.3.20 we saw multiplication of e_i^T on the left of A allows us to select the i -th row of A :

$$A e_j = \text{col}_j(A) \quad \& \quad e_i^T A = \text{row}_i(A).$$

Again, we illustrate these identities in the 2×2 case:

Example 1.6.2. Let $A = \begin{bmatrix} 1 & 2 \\ 3 & 4 \end{bmatrix}$. Thus,

$$A e_2 = \begin{bmatrix} 1 & 2 \\ 3 & 4 \end{bmatrix} \begin{bmatrix} 0 \\ 1 \end{bmatrix} = \begin{bmatrix} 2 \\ 4 \end{bmatrix} \quad \& \quad e_2^T A = \begin{bmatrix} 0 & 1 \end{bmatrix} \begin{bmatrix} 1 & 2 \\ 3 & 4 \end{bmatrix} = \begin{bmatrix} 3 & 4 \end{bmatrix}.$$

Proposition 1.3.15 which allows us to view matrix multiplication as

$$A[\mathbf{v}_1 \ \mathbf{v}_2 \ \dots \ \mathbf{v}_\ell] = [A\mathbf{v}_1 \ A\mathbf{v}_2 \ \dots \ A\mathbf{v}_\ell]$$

(i.e. done column-by-column) and Proposition 1.3.16 allows us to view matrix multiplication as:

$$\begin{bmatrix} \mathbf{w}_1 \\ \vdots \\ \mathbf{w}_k \end{bmatrix} A = \begin{bmatrix} \mathbf{w}_1 A \\ \vdots \\ \mathbf{w}_k A \end{bmatrix}$$

(i.e. done row-by-row). Therefore,

$$A E_{ij} = A[0 \ \dots \ 0 \ e_i \ 0 \ \dots \ 0] = [0 \ \dots \ 0 \ A e_i \ 0 \ \dots \ 0]$$

is merely the i^{th} column of A slapped into the j^{th} column of the zero matrix. Likewise, $E_{ij} A$ is the j^{th} row of A slapped into the i^{th} row of the zero matrix. Illustrate via the 2×2 case once more:

²⁰in a compressed treatment of linear algebra I might avoid discussion of these, but, as we will see shortly, these elementary matrices allow for concrete proofs of many important aspects of row reduction, the structure of inverse matrices, even the product identity for determinants. Skipping these is not wise if we care about why things work

²¹in $\mathbb{R}^2$ we sometimes call $\mathbf{e}_1 = \mathbf{i}$ and $\mathbf{e}_2 = \mathbf{j}$ and in $\mathbb{R}^3$ sometimes we say $\mathbf{e}_1 = \mathbf{i}$, $\mathbf{e}_2 = \mathbf{j}$, $\mathbf{e}_3 = \mathbf{k}$. However, in my multivariate calculus notes I use the notation $e_i = \hat{x}_i$ to denote unit-vectors in the direction of the i -th Cartesian coordinate, or $\hat{x}, \hat{y}, \hat{z}$ for the usual three-dimensional applications.

Example 1.6.3. Let $A = \begin{bmatrix} 1 & 2 \\ 3 & 4 \end{bmatrix}$. Then,

$$AE_{21} = \begin{bmatrix} 1 & 2 \\ 3 & 4 \end{bmatrix} \begin{bmatrix} 0 & 0 \\ 1 & 0 \end{bmatrix} = \begin{bmatrix} 2 & 0 \\ 4 & 0 \end{bmatrix} \quad \& \quad E_{21}A = \begin{bmatrix} 0 & 0 \\ 1 & 0 \end{bmatrix} \begin{bmatrix} 1 & 2 \\ 3 & 4 \end{bmatrix} = \begin{bmatrix} 0 & 0 \\ 1 & 2 \end{bmatrix}.$$

Proposition 1.6.4. *Elementary Matrix of Type I:*

If $E = I_m - E_{ii} - E_{jj} + E_{ij} + E_{ji} \in \mathbb{F}^{m \times m}$ and $A \in \mathbb{F}^{m \times n}$ then EA is A with the i -th and j -th row of A swapped.

Proof: let $E = I_m - E_{ii} - E_{jj} + E_{ij} + E_{ji}$. Let's track through how the various parts of E act on A by left-multiplication:

- (1.) $E_{ii}A$ would be the i^{th} row of A left in place with all other rows zeroed out. Likewise $E_{jj}A$ provides be the j^{th} row of A left in place with all other rows zeroed out
- (2.) by (1.) we find $(I_m - E_{ii} - E_{jj})A = A - E_{ii}A - E_{jj}A$ wipes out rows i and j
- (3.) $E_{ij}A$ would be the j^{th} row of A moved to the i^{th} row with all other rows zeroed out. Thus, by adding in $(E_{ij} + E_{ji})A$, we put rows i and j back but interchanging their locations in $E = I_m - E_{ii} - E_{jj} + E_{ij} + E_{ji}$.

In summary, EA swaps rows i and j . In particular, $E = EI_m$ is the identity matrix with rows i and j swapped (this gives an explicit formula for E). $\square$

It is interesting to note that $E^T = I_m^T - E_{ii}^T - E_{jj}^T + E_{ij}^T + E_{ji}^T = I_m - E_{ii} - E_{jj} + E_{ji} + E_{ij} = E$. Also, since swapping twice undoes the swap, $E^{-1} = E$ (E is its own inverse).

Example 1.6.5. We obtain an elementary matrix E for Type I operation formed by swapping rows 1 and 3 (so $E = I_3 - E_{11} - E_{33} + E_{13} + E_{31}$).

$$\begin{bmatrix} 0 & 0 & 1 \\ 0 & 1 & 0 \\ 1 & 0 & 0 \end{bmatrix} \begin{bmatrix} 2 & -1 & 3 & : & -1 \\ 0 & 5 & -6 & : & 0 \\ -1 & 0 & 4 & : & 7 \end{bmatrix} = \begin{bmatrix} -1 & 0 & 4 & : & 7 \\ 0 & 5 & -6 & : & 0 \\ 2 & -1 & 3 & : & -1 \end{bmatrix}$$

Proposition 1.6.6. *Elementary Matrix of Type II: for $c \neq 0$ in $\mathbb{F}$:*

If $E = I_m - E_{ii} + cE_{ii} \in \mathbb{F}^{m \times m}$ and $A \in \mathbb{F}^{m \times n}$ then EA is A with the i -th row of A multiplied by c .

Proof: Let $E = I_m - E_{ii} + cE_{ii}$. Observe $(I_m - E_{ii})A$ gives A with the i -th row zeroed out. Note $cE_{ii}A$ provides a matrix with zero all rows except the i -th row. In the i -th row of $cE_{ii}A$ we find $\text{crow}_i(A)$. Hence, EA gives A with the i -th row multiplied by c and the remaining rows unaltered. $\square$

Notice that $E^{-1} = I_m - E_{ii} + c^{-1}E_{ii}$ (to undo scaling row i by c we should scale row i by $1/c$). So E^{-1} corresponds to a type II operation.

Example 1.6.7. We find E for a Type II operation of scaling row 3 by -2 (so we set $E = I_3 - E_{33} + (-2)E_{33}$).

$$\begin{bmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & -2 \end{bmatrix} \begin{bmatrix} 2 & -1 & 3 & : & -1 \\ 0 & 5 & -6 & : & 0 \\ -1 & 0 & 4 & : & 7 \end{bmatrix} = \begin{bmatrix} 2 & -1 & 3 & : & -1 \\ 0 & 5 & -6 & : & 0 \\ 2 & 0 & -8 & : & -14 \end{bmatrix}$$

Proposition 1.6.8. *Elementary Matrix of Type III: for $s \in \mathbb{F}$ and $i \neq j$:*

If $E = I_m + sE_{ji} \in \mathbb{F}^{m \times m}$ and $A \in \mathbb{F}^{m \times n}$ then EA has $\text{row}_k(EA) = \text{row}_k(A)$ for $k \neq j$ and $\text{row}_j(EA) = \text{row}_j(A) + s \text{row}_i(A)$.

Proof: Let $E = I_m + sE_{ji}$. Again, recall left multiplication by E_{ji} (note the subscripts) will copy row i into row j 's place. So $E = I_m + sE_{ji}$ will add s times row i to j . $\square$

To undo adding s times row j to row i we should subtract s times row j from row i . Therefore, $E^{-1} = I_m - sE_{ji}$, so yet again the inverse of an elementary operation is an elementary operation of the same type.

Example 1.6.9. *A Type III operation can be formed by adding 3 times row 3 to row 2 (so we construct $E = I_3 + 3E_{23}$).*

$$\begin{bmatrix} 1 & 0 & 0 \\ 0 & 1 & 3 \\ 0 & 0 & 1 \end{bmatrix} \begin{bmatrix} 2 & -1 & 3 & : & -1 \\ 0 & 5 & -6 & : & 0 \\ -1 & 0 & 4 & : & 7 \end{bmatrix} = \begin{bmatrix} 2 & -1 & 3 & : & -1 \\ -3 & 5 & 6 & : & 21 \\ -1 & 0 & 4 & : & 7 \end{bmatrix}$$

Notice, the discussion in this section shows:

Proposition 1.6.10.

Each elementary matrix is invertible with inverse matrix of the same type.

Proposition 1.6.11.

Let $A \in \mathbb{F}^{m \times n}$ then there exist elementary matrices $E_1, E_2, \dots, E_k$ such that $\text{rref}(A) = E_k \cdots E_2 E_1 A$.

Proof: Gauss Jordan elimination consists of a sequence of k elementary row operations. Each row operation can be implemented by multiplying the corresponding elementary matrix on the left. Hence successively left-multiplying A by each elementary matrix corresponding to a row operation in the sequence will produce $\text{rref}(A)$. $\square$

1.6.2 supraugmented matrices and inverse matrix calculation

Suppose we have a system of equations $Ax = b_1$ and another system $Ax = b_2$. Both systems share the same coefficient matrix A . Suppose E is a product of elementary matrices for which $\text{rref}(A) = EA$. Since row-reduction is done column-by-column,

$$\text{rref}[A|b_1|B_2] = E[A|b_1|b_2]$$

Apply Corollary 1.3.18,

$$\text{rref}[A|b_1|b_2] = [EA|Eb_1|Eb_2] = [\text{rref}(A)|Eb_1|Eb_2].$$

This means we can calculate the solution to multiple systems with the same row-reduction.

Example 1.6.12. Solve the systems given below,

$$\begin{array}{rcl} x + y + z = 1 \\ x - y + z = 0 \\ -x + z = 1 \end{array} \quad (\text{I.}) \quad \text{and} \quad \begin{array}{rcl} x + y + z = 1 \\ x - y + z = 1 \\ -x + z = 1 \end{array} \quad (\text{II.})$$

The systems above share the same coefficient matrix, however $b_1 = (1, 0, 1)$ whereas $b_2 = (1, 1, 1)$. We can solve both at once by making an extended augmented coefficient matrix $[A|b_1|b_2]$

$$[A|b_1|b_2] = \left[\begin{array}{ccc|c|c} 1 & 1 & 1 & 1 & 1 \\ 1 & -1 & 1 & 0 & 1 \\ -1 & 0 & 1 & 1 & 1 \end{array} \right] \quad \text{rref}[A|b_1|b_2] = \left[\begin{array}{ccc|c|c} 1 & 0 & 0 & -1/4 & 0 \\ 0 & 1 & 0 & 1/2 & 0 \\ 0 & 0 & 1 & 3/4 & 1 \end{array} \right]$$

Thus (I.) has solution $(-1/4, 1/2, 3/4)$ and (II.) has solution $(0, 0, 1)$.

If $A \in \mathbb{F}^{n \times n}$ is invertible then there exist $v_1, \dots, v_n \in \mathbb{F}^n$ such that $A^{-1} = [v_1 | \dots | v_n]$ and

$$AA^{-1} = I \Rightarrow A[v_1 | \dots | v_n] = [Av_1 | \dots | Av_n] = [e_1 | \dots | e_n] \Rightarrow Av_j = e_j$$

for $j = 1, \dots, n$. Therefore, the inverse of A exists if and only if the equations $Av_1 = e_1, \dots, Av_n = e_n$ have a solution. We can calculate these n -solutions by simply row reducing the super-augmented coefficient matrix $[A|e_1 | \dots | e_n]$.

Suppose A is square and $\text{rref}(A) = EA$ then

$$\text{rref}[A|I] = E[A|I] = [EA|EI] = [\text{rref}(A)|E]$$

If $\text{rref}(A) \neq I$ then at least one pivot column is found in the last n -columns of the $n \times (2n)$ matrix $[\text{rref}(A)|E]$ and so there exists some j for which $Ax = e_j$ has no solution. In other words, if $\text{rref}(A) \neq I$ then A^{-1} does not exist. On the other hand, if $\text{rref}(A) = I$ then $\text{rref}[A|I] = [I|E]$ and we see $EA = I$ thus $E = A^{-1}$. This is the logic behind the algorithm taught to find the inverse by row-reducing $[A|I]$.

Example 1.6.13. Recall that in Example 1.5.8 we worked out the details of

$$\text{rref} \left[\begin{array}{ccc|ccc} 1 & 0 & 0 & 1 & 0 & 0 \\ 2 & 2 & 0 & 0 & 1 & 0 \\ 4 & 4 & 4 & 0 & 0 & 1 \end{array} \right] = \left[\begin{array}{ccc|ccc} 1 & 0 & 0 & 1 & 0 & 0 \\ 0 & 1 & 0 & -1 & 1/2 & 0 \\ 0 & 0 & 1 & 0 & -1/2 & 1/4 \end{array} \right]$$

Thus,

$$\left[\begin{array}{ccc} 1 & 0 & 0 \\ 2 & 2 & 0 \\ 4 & 4 & 4 \end{array} \right]^{-1} = \left[\begin{array}{ccc} 1 & 0 & 0 \\ -1 & 1/2 & 0 \\ 0 & -1/2 & 1/4 \end{array} \right].$$

The theorem that follows here collects the many things we have learned about an $n \times n$ invertible matrices and corresponding linear systems:

Theorem 1.6.14.

Let A be an $n \times n$ matrix over a field $\mathbb{F}$ then the following are equivalent:

- (1.) A is invertible,
- (2.) $rref[A] = I$,
- (3.) $Ax = 0$ iff $x = 0$,
- (4.) A is the product of elementary matrices,
- (5.) there exists $B \in \mathbb{R}^{n \times n}$ such that $AB = I$,
- (6.) there exists $B \in \mathbb{R}^{n \times n}$ such that $BA = I$,
- (7.) for each $b \in \mathbb{F}^m$ there exists $x \in \mathbb{F}^n$ for which $rref[A|b] = [I|x]$,
- (8.) $Ax = b$ is consistent for every $b \in \mathbb{F}^m$,
- (9.) $Ax = b$ has exactly one solution for each $b \in \mathbb{F}^m$,
- (10.) A^T is invertible.

1.7 spans of column vectors and the CCP

Recall that the span of a set of vectors in $\mathbb{F}^n$ is simply the set of all finite $\mathbb{F}$ -linear combinations of vectors taken from the set. The problem of solving a linear system and the problem of spanning are naturally linked as is explained in the Proposition:

Proposition 1.7.1.

If $A = [v_1|v_2|\cdots|v_n] \in \mathbb{F}^{m \times n}$ and $b \in \mathbb{F}^m$ then the matrix equation $Ax = b$ has the same set of solutions as the vector equation

$$x_1v_1 + x_2v_2 + \cdots + x_nv_n = b.$$

Thus $b \in \text{span}\{v_1, v_2, \dots, v_n\}$ if and only if $[v_1|v_2|\cdots|v_n]x = b$ has a solution.

Example 1.7.2. Problem: Let $b_1 = (1, 1, 0)$, $b_2 = (0, 1, 1)$ and $b_3 = (0, 1, -1)$. Is²² $e_3 \in \text{span}\{b_1, b_2, b_3\}$?

Solution: Find the explicit linear combination of b_1, b_2, b_3 that produces e_3 . We seek to find $x, y, z \in \mathbb{R}$ such that $xb_1 + yb_2 + zb_3 = e_3$,

$$x \begin{bmatrix} 1 \\ 1 \\ 0 \end{bmatrix} + y \begin{bmatrix} 0 \\ 1 \\ 1 \end{bmatrix} + z \begin{bmatrix} 0 \\ 1 \\ -1 \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \\ 1 \end{bmatrix} \Rightarrow \begin{bmatrix} x \\ x+y+z \\ y-z \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \\ 1 \end{bmatrix}$$

Following the Proposition above, we answer the question by gluing the given vectors into a matrix and doing row reduction. In particular, we can solve the vector equation above by solving the

²²challenge: once you understand this example for e_3 try answering it for other vectors or for an arbitrary vector $v = (v_1, v_2, v_3)$. How would you calculate $x, y, z \in \mathbb{R}$ such that $v = xb_1 + yb_2 + zb_3$?

corresponding system below:

$$\left[\begin{array}{ccc|c} 1 & 0 & 0 & 0 \\ 1 & 1 & 1 & 0 \\ 0 & 1 & -1 & 1 \end{array} \right] \xrightarrow{r_2 - r_1} \left[\begin{array}{ccc|c} 1 & 0 & 0 & 0 \\ 0 & 1 & 1 & 0 \\ 0 & 1 & -1 & 1 \end{array} \right] \xrightarrow{r_3 - r_2} \left[\begin{array}{ccc|c} 1 & 0 & 0 & 0 \\ 0 & 1 & 1 & 0 \\ 0 & 0 & -2 & 1 \end{array} \right] \xrightarrow{\begin{array}{l} -r_3/2 \\ r_2 - r_3 \\ r_1 - r_3 \end{array}} \left[\begin{array}{ccc|c} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 1/2 \\ 0 & 0 & 1 & -1/2 \end{array} \right]$$

Therefore, $x = 0, y = \frac{1}{2}$ and $z = -\frac{1}{2}$. We find that $e_3 = \frac{1}{2}b_2 - \frac{1}{2}b_3$ thus $e_3 \in \text{span}\{b_1, b_2, b_3\}$.

Example 1.7.3. Problem: Let $b_1 = (1, 2, 3, 4)$, $b_2 = (0, 1, 0, 1)$ and $b_3 = (0, 0, 1, 1)$. Is $w = (1, 1, 4, 4) \in \text{span}\{b_1, b_2, b_3\}$?

Solution: Following the same method as the last example we seek to find x_1, x_2 and x_3 such that $x_1b_1 + x_2b_2 + x_3b_3 = w$ by solving the aug. coeff. matrix as is our custom:

$$[b_1|b_2|b_3|w] = \left[\begin{array}{cccc|c} 1 & 0 & 0 & 1 \\ 2 & 1 & 0 & 1 \\ 3 & 0 & 1 & 4 \\ 4 & 1 & 1 & 4 \end{array} \right] \xrightarrow{\begin{array}{l} r_2 - 2r_1 \\ r_3 - 3r_1 \\ r_4 - 4r_1 \end{array}} \left[\begin{array}{cccc|c} 1 & 0 & 0 & 1 \\ 0 & 1 & 0 & -1 \\ 0 & 0 & 1 & 1 \\ 0 & 1 & 1 & 0 \end{array} \right] \xrightarrow{r_4 - r_2} \left[\begin{array}{cccc|c} 1 & 0 & 0 & 1 \\ 0 & 1 & 0 & -1 \\ 0 & 0 & 1 & 1 \\ 0 & 0 & 1 & 1 \end{array} \right] \xrightarrow{r_4 - r_3} \left[\begin{array}{cccc|c} 1 & 0 & 0 & 1 \\ 0 & 1 & 0 & -1 \\ 0 & 0 & 1 & 1 \\ 0 & 0 & 0 & 0 \end{array} \right] = \text{rref}[b_1|b_2|b_3|w]$$

We find $x_1 = 1, x_2 = -1, x_3 = 1$ thus $w = b_1 - b_2 + b_3$. Therefore, $w \in \text{span}\{b_1, b_2, b_3\}$.

Pragmatically, if the question is sufficiently simple you may not need to use the augmented coefficient matrix to solve the question. I use them here to illustrate the method.

Example 1.7.4. Problem: Let $b_1 = (1, 1, 0)$ and $b_2 = (0, 1, 1)$. Is $e_2 \in \text{span}\{b_1, b_2\}$?

Solution: Attempt to find the explicit linear combination of b_1, b_2 that produces e_2 . We seek to find $x, y \in \mathbb{R}$ such that $xb_1 + yb_2 = e_3$,

$$x \begin{bmatrix} 1 \\ 1 \\ 0 \end{bmatrix} + y \begin{bmatrix} 0 \\ 1 \\ 1 \end{bmatrix} = \begin{bmatrix} 0 \\ 1 \\ 0 \end{bmatrix} \Rightarrow \begin{bmatrix} x \\ x+y \\ y \end{bmatrix} = \begin{bmatrix} 0 \\ 1 \\ 0 \end{bmatrix}$$

We don't really need to consult the augmented matrix to solve this problem. Clearly $x = 0$ and $y = 0$ is found from the first and third components of the vector equation above. But, the second component yields $x + y = 1$ thus $0 + 0 = 1$. It follows that this system is inconsistent and we may conclude that $w \notin \text{span}\{b_1, b_2\}$. For the sake of curiosity let's see how the augmented solution matrix looks in this case: omitting details of the row reduction,

$$\text{rref} \left[\begin{array}{cc|c} 1 & 0 & 0 \\ 1 & 1 & 1 \\ 0 & 1 & 0 \end{array} \right] = \left[\begin{array}{cc|c} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{array} \right]$$

note the last row again confirms that this is an inconsistent system.

1.7.1 solving several spanning questions simultaneously

If we are given $B = \{b_1, b_2, \dots, b_k\} \subset \mathbb{F}^n$ and $T = \{w_1, w_2, \dots, w_r\} \subset \mathbb{F}^n$ and we wish to determine if $T \subset \text{span}(B)$ then we can answer the question by examining if $[b_1|b_2|\dots|b_k]x = w_j$ has a solution for each $j = 1, 2, \dots, r$. Or we may solve it by calculating

$$\text{rref}[b_1|b_2|\dots|b_k||w_1|w_2|\dots|w_r].$$

The notation $||$ is used to distinguish between the generating vectors from B and the target vectors in T . The question is whether or not the vectors in T can be attained by some linear combination of the vectors in B . If there is a row with zeros in the first k -columns and a nonzero entry in the last r -columns then this means that at least one vector w_k is not in the span of B (moreover, the vector not in the span corresponds to the nonzero entrie(s)). Otherwise, each vector is in the span of B and we can read the precise linear combination from the matrix. I will illustrate this in the example that follows.

Example 1.7.5. Let $W = \text{span}\{e_1 + e_2, e_2 + e_3, e_1 - e_3\}$ and suppose $T = \{e_1, e_2, e_3 - e_1\}$. Is $T \subseteq W$? If not, which vectors in T are not in W ? Consider, $[e_1 + e_2|e_2 + e_3|e_1 - e_3||e_1|e_2|e_3 - e_1] =$

$$\begin{aligned} &= \left[\begin{array}{ccc|ccc} 1 & 0 & 1 & 1 & 0 & -1 \\ 1 & 1 & 0 & 0 & 1 & 0 \\ 0 & 1 & -1 & 0 & 0 & 1 \end{array} \right] \xrightarrow{r_2 - r_1} \left[\begin{array}{ccc|ccc} 1 & 0 & 1 & 1 & 0 & -1 \\ 0 & 1 & -1 & -1 & 1 & 1 \\ 0 & 1 & -1 & 0 & 0 & 1 \end{array} \right] \\ &\xrightarrow{r_3 - r_2} \left[\begin{array}{ccc|ccc} 1 & 0 & 1 & 1 & 0 & -1 \\ 0 & 1 & -1 & -1 & 1 & 1 \\ 0 & 0 & 0 & 1 & -1 & 0 \end{array} \right] \xrightarrow[r_1 - r_3]{r_2 + r_3} \left[\begin{array}{ccc|ccc} 1 & 0 & 1 & 0 & 1 & -1 \\ 0 & 1 & -1 & 0 & 0 & 1 \\ 0 & 0 & 0 & 1 & -1 & 0 \end{array} \right] \end{aligned}$$

Let me summarize the calculation:

$$\text{rref}[e_1 + e_2|e_2 + e_3|e_1 - e_3||e_1|e_2|e_3 - e_1] = \left[\begin{array}{ccc|ccc} 1 & 0 & 1 & 0 & 1 & -1 \\ 0 & 1 & -1 & 0 & 0 & 1 \\ 0 & 0 & 0 & 1 & -1 & 0 \end{array} \right]$$

We deduce that e_1 and e_2 are not in W . However, $e_3 - e_1 \in W$ and we can read from the matrix $-(e_1 + e_2) + (e_2 + e_3) = e_3 - e_1$. I added the double vertical bar for book-keeping purposes, as usual the vertical bars are just to aid the reader in parsing the matrix.

In short, if we wish to settle if several vectors are in the span of a given generating set then we can do one sweeping row reduction to determine what is in the span. This is a great reduction in labor from what you might naively expect. That said, the CCP is even faster. Let's get to it.

1.7.2 CCP and linear dependence of column vectors

Recall a set of column vectors in $\mathbb{F}^n$ is said to be **linearly dependent** if there exists a vector in the set which can be written as an $\mathbb{F}$ -linear combination of other vectors in the set. In particular, if $b \in \text{span}\{v_1, \dots, v_k\}$ then $\{b, v_1, \dots, v_k\}$ is a linearly dependent set. Our main goal here is to understand a particular technique which is special to the context of column vectors. I call this method the **Column Correspondence Property** or **CCP**. Others call it the **linear correspondence**. In short, the CCP allows us to decide questions of linear dependence simply by inspection of the rref and common sense.

Proposition 1.7.6. *Column Correspondence Property (CCP)*

If $A, R \in \mathbb{F}^{m \times n}$ are row-equivalent matrices then any linear dependence of the columns of R is shared by the columns of A . In particular, any linear dependence found amongst the columns of $rref(A)$ is likewise found in the columns of A .

Proof: if A and R are row-equivalent then there exist elementary matrices $E_1, \dots, E_k$ for which $R = E_k \cdots E_1 A$. Let $E = E_k \cdots E_1$ for brevity of notation. Let $A = [A_1 | \cdots | A_n]$ and $R = [R_1 | \cdots | R_n]$. Observe,

$$EA = R \Rightarrow E[A_1 | \cdots | A_n] = [EA_1 | \cdots | EA_n] = [R_1 | \cdots | R_n] \Rightarrow EA_j = R_j$$

for $j = 1, \dots, k$. Suppose $c_1, \dots, c_k \in \mathbb{F}$ such that $\sum_{j=1}^k c_j A_j = 0$. Multiply by E on the left,

$$E \left(\sum_{j=1}^k c_j A_j \right) = E(0) = 0 \Rightarrow \sum_{j=1}^k c_j EA_j = \sum_{j=1}^k c_j R_j = 0$$

Likewise, if we suppose there exist $b_1, \dots, b_k \in \mathbb{F}$ such that $\sum_{j=1}^k b_j R_j = 0$ then multiplication by E^{-1} on the left, paired with the observation $A_j = E^{-1} R_j$, will likewise yield $\sum_{j=1}^k b_j A_j = 0$. Finally, the claim about the $rref(A)$ follows immediately since $rref(A)$ and A are row-equivalent matrices. $\square$

Example 1.7.7. In Example 1.7.2 we studied $b_1 = (1, 1, 0)$, $b_2 = (0, 1, 1)$ and $b_3 = (0, 1, -1)$ and asked if $e_3 \in \text{span}(b_1, b_2, b_3)$? We calculated:

$$rref[b_1 | b_2 | b_3 | e_3] = \left[\begin{array}{ccc|c} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 1/2 \\ 0 & 0 & 1 & -1/2 \end{array} \right] \Rightarrow e_3 = \frac{1}{2}b_2 - \frac{1}{2}b_3$$

by the CCP and inspection of the row reduced matrix above.

Example 1.7.8. In Example 1.7.3 we studied $b_1 = (1, 2, 3, 4)$, $b_2 = (0, 1, 0, 1)$ and $b_3 = (0, 0, 1, 1)$ and asked if $w = (1, 1, 4, 4) \in \text{span}\{b_1, b_2, b_3\}$? Observe that the CCP yields the implication below:

$$rref[b_1 | b_2 | b_3 | w] = \left[\begin{array}{ccc|c} 1 & 0 & 0 & 1 \\ 0 & 1 & 0 & -1 \\ 0 & 0 & 1 & 1 \\ 0 & 0 & 0 & 0 \end{array} \right] \Rightarrow w = b_1 - b_2 + b_3.$$

Example 1.7.9. In Example 1.7.4 we set $b_1 = (1, 1, 0)$ and $b_2 = (0, 1, 1)$ and we asked if $e_2 \in \text{span}\{b_1, b_2\}$? Observe, by the CCP applied to $rref[b_1 | b_2 | e_2]$, the answer is clearly no,

$$rref[b_1 | b_2 | e_2] = \left[\begin{array}{cc|c} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{array} \right] \Rightarrow e_2 \notin \text{span}\{b_1, b_2\}$$

Example 1.7.10. Let us revisit Example 1.7.5. Let $v_1 = e_1 + e_2$, $v_2 = e_2 + e_3$, $v_3 = e_3 - e_1$ and set $W = \text{span}\{v_1, v_2, v_3\}$. Since

$$rref[v_1 | v_2 | v_3 | e_1 | e_2 | e_3 - e_1] = \left[\begin{array}{ccc|ccc} 1 & 0 & 1 & 0 & 1 & -1 \\ 0 & 1 & -1 & 0 & 0 & 1 \\ 0 & 0 & 0 & 1 & -1 & 0 \end{array} \right]$$

the CCP implies $e_1, e_2 \notin W$ whereas $e_3 - e_1 = v_2 - v_1 \in W$

You should notice that the CCP saves us the trouble of expressing how the constants c_i are related. If we are only interested in how the vectors are related the CCP gets straight to the point quicker. I invite the reader to look at all the row-reductions in this Chapter in light of the CCP.

Chapter 2

Vector Spaces

And we know that all things work together for good to them that love God,
to them who are the called according to his purpose. ROMANS 8:28, KJV

Up to this point the topics we have discussed loosely fit into the category of matrix theory. The concept of a matrix is millennia old. If I trust my source, and I think I do, the Chinese even had an analog of Gaussian elimination about 2000 years ago. The modern notation likely stems from the work of Cauchy in the 19-th century. Cauchy's prolific work colors much of the notation we still use. The concept of coordinate geometry as introduced by Descartes and Fermat around 1644 is what ultimately led to the concept of a vector space.¹ Grassmann, Hamilton, and many many others worked out volumous work detailing possible transformations on what we now call $\mathbb{R}^2, \mathbb{R}^3, \mathbb{R}^4, \dots$. Argand(complex numbers) and Hamilton(quaternions) had more than what we would call a vector space. They had a linear structure plus some rule for multiplication of vectors. A vector space with a multiplication is called an *algebra* in the modern terminology.

Honestly, I think once the concept of the Cartesian plane was discovered the concept of a vector space almost certainly must follow. That said, it took a while for the definition I state in the next section to appear. Giuseppe Peano gave the modern definition for a vector space in 1888². In addition he put forth some of the ideas concerning linear transformations. Peano is also responsible for the modern notations for intersection and unions of sets³. He made great contributions to proof by induction and the construction of the natural numbers from basic set theory.

I should mention the work of Hilbert, Lebesgue, Fourier, Banach and others were greatly influential in the formation of infinite dimensional vector spaces. Our focus is on the finite dimensional case.⁴

Let me summarize what a vector space is before we define it properly. In short, a vector space over a field $\mathbb{F}$ is simply a set which allows you to add its elements and multiply by the numbers in $\mathbb{F}$. A field is a set with addition and multiplication defined such that every nonzero element has a multiplicative inverse. Typical examples, $\mathbb{F} = \mathbb{R}, \mathbb{C}, \mathbb{Q}, \mathbb{Z}/p\mathbb{Z}$ where p is prime.

¹Bourbaki 1969, ch. "Algebre lineaire et algebre multilineaire", pp. 78-91.

²Peano, Giuseppe (1888), *Calcolo Geometrico secondo l'Ausdehnungslehre di H. Grassmann preceduto dalle Operazioni della Logica Deduttiva*, Turin

³see Pg 87 of A Transition to Advanced Mathematics: A Survey Course By William Johnston

⁴this history is flawed, one-sided and far too short. You should read a few more books if you're interested.

Vector spaces are found throughout modern mathematics. Moreover, the theory we cover in this chapter is applicable to a myriad of problems with real world content. This is the beauty of linear algebra: it simultaneously illustrates the power of application and abstraction in mathematics.

2.1 definition and examples

Axioms are not derived from a more basic logic. They are the starting point. Their validity is ultimately judged by their use. However, this definition is naturally motivated by the structure of vector addition and scalar multiplication in $\mathbb{R}^n$ (or $\mathbb{F}^n$ if that is where your intuition rests)

Definition 2.1.1.

A vector space V over a field $\mathbb{F}$ is a nonempty set V together with a function $+: V \times V \rightarrow V$ called **vector addition** and another function $\cdot: \mathbb{F} \times V \rightarrow V$ called **scalar multiplication**. We require that the operations of vector addition and scalar multiplication satisfy the following 10 axioms: for all $x, y, z \in V$ and $a, b \in \mathbb{F}$,

- (A1.) $x + y = y + x$ for all $x, y \in V$,
- (A2.) $(x + y) + z = x + (y + z)$ for all $x, y, z \in V$,
- (A3.) there exists $0 \in V$ such that $x + 0 = x = 0 + x$ for all $x \in V$,
- (A4.) for each $x \in V$ there exists $-x \in V$ such that $x + (-x) = 0 = (-x) + x$,
- (A5.) $1 \cdot x = x$ for all $x \in V$,
- (A6.) $(ab) \cdot x = a \cdot (b \cdot x)$ for all $x \in V$ and $a, b \in \mathbb{F}$,
- (A7.) $a \cdot (x + y) = a \cdot x + a \cdot y$ for all $x, y \in V$ and $a \in \mathbb{F}$,
- (A8.) $(a + b) \cdot x = a \cdot x + b \cdot x$ for all $x \in V$ and $a, b \in \mathbb{F}$,
- (A9.) If $x, y \in V$ then $x + y$ is a single element in V ,
(we say V is closed with respect to addition)
- (A10.) If $x \in V$ and $c \in \mathbb{F}$ then $c \cdot x$ is a single element in V .
(we say V is closed with respect to scalar multiplication)

We call 0 in axiom 3 the **zero vector** and the vector $-x$ is called the **additive inverse** of x . We will sometimes omit the $\cdot$ and instead denote scalar multiplication by juxtaposition; $a \cdot x = ax$. We write $V(\mathbb{F})$ to communicate that the pointset V is a vectorspace over $\mathbb{F}$.

Axioms (9.) and (10.) are admittedly redundant given that those automatically follow from the statements that $+: V \times V \rightarrow V$ and $\cdot: \mathbb{F} \times V \rightarrow V$ are functions. I've listed them so that you are less likely to forget they must be checked.

The terminology "vector" does not necessarily indicate an explicit geometric interpretation in this general context. Sometimes I'll insert the word "abstract" to emphasize this distinction. We'll see that matrices, polynomials and functions in general can be thought of as abstract vectors.

Example 2.1.2. Real Matrices form real vector spaces: $\mathbb{R}$ is a vector space if we identify addition of real numbers as the vector addition and multiplication of real numbers as the scalar

multiplication. Likewise, $\mathbb{R}^n$ forms a vector space over $\mathbb{R}$ with respect to the usual vector addition and scalar multiplication:

$$(x + y)_i = x_i + y_i \quad \& \quad (cx)_i = cx_i$$

for each $i \in \mathbb{N}_n$ and $x, y \in \mathbb{R}^n$ and $c \in \mathbb{R}$. In fact, even $\mathbb{R}^n$ is just the $n \times 1$ case of $\mathbb{R}^{m \times n}$. Indeed, $\mathbb{R}^{m \times n}$ forms a vector space over $\mathbb{R}$ with addition and scalar multiplication of matrices defined as we studied in previous chapters:

$$(A + B)_{ij} = A_{ij} + B_{ij} \quad \& \quad (cA)_{ij} = cA_{ij}$$

for all $(i, j) \in \mathbb{N}_m \times \mathbb{N}_n$ and $A, B \in \mathbb{R}^{m \times n}$ and $c \in \mathbb{R}$.

In the previous example, I introduced the standard interpretation of $\mathbb{R}$, $\mathbb{R}^n$ and $\mathbb{R}^{m \times n}$ as **real vector spaces**. To say V is a **real** vector space is just another way of saying V is a vector space with the field of scalars being the real numbers. Proof that Axioms 1-10 are met was already given in part in Proposition 1.2.6 and Theorem 1.2.9 (if we prove something for an arbitrary commutative ring then this naturally includes the case $R = \mathbb{R}$). I should mention, we can also view $\mathbb{R}$, $\mathbb{R}^n$ and $\mathbb{R}^{m \times n}$ as vector spaces over the rational numbers $\mathbb{Q}$. Our notation for such vector spaces would be:

$$\mathbb{R}(\mathbb{Q}), \quad \mathbb{R}^n(\mathbb{Q}), \quad \mathbb{R}^{m \times n}(\mathbb{Q})$$

which indicates vector spaces of real numbers, vectors and matrices such that the scalar multiplication by rational numbers. However, even $\mathbb{R}$ is infinite dimensional over the rational numbers. In contrast, $\mathbb{R}$ is one-dimensional over $\mathbb{R}$. For now, I use the term **dimensional** as an intuitive term, we shall soon give it a rigorous meaning. The next example should not be surprising in view of Example 2.1.2.

Example 2.1.3. Complex matrices form complex vector spaces: $\mathbb{C}$ is a vector space if we identify addition of complex numbers as the vector addition and multiplication of complex numbers as the scalar multiplication. Likewise, $\mathbb{C}^n$ forms a vector space over $\mathbb{C}$ with respect to the usual vector addition and scalar multiplication:

$$(x + y)_i = x_i + y_i \quad \& \quad (cx)_i = cx_i$$

for each $i \in \mathbb{N}_n$ and $x, y \in \mathbb{C}^n$ and $c \in \mathbb{C}$. In fact, even $\mathbb{C}^n$ is just the $n \times 1$ case of $\mathbb{C}^{m \times n}$. Indeed, $\mathbb{C}^{m \times n}$ forms a vector space over $\mathbb{C}$ with addition and scalar multiplication of matrices defined as we studied in previous chapters:

$$(A + B)_{ij} = A_{ij} + B_{ij} \quad \& \quad (cA)_{ij} = cA_{ij}$$

for all $(i, j) \in \mathbb{N}_m \times \mathbb{N}_n$ and $A, B \in \mathbb{C}^{m \times n}$ and $c \in \mathbb{C}$.

If a given point-set permits the assignment of a vector space structure (meaning we can define addition and scalar multiplication which adhere to Axioms 1-10) then it may be possible to assign a different vector space structure to the set as well. In Example 2.1.5 we discussed $\mathbb{C}^{m \times n}$ as a complex vector space. In contrast, in the example below we give $\mathbb{C}^{m \times n}$ a real vector space structure:

Example 2.1.4. Let $V = \mathbb{C}^{m \times n}(\mathbb{R})$ denote the vector space over $\mathbb{R}$ where addition and scalar multiplication are defined by:

$$(A + B)_{ij} = A_{ij} + B_{ij} \quad \& \quad (cA)_{ij} = cA_{ij}$$

for all $(i, j) \in \mathbb{N}_m \times \mathbb{N}_n$ and $A, B \in \mathbb{C}^{m \times n}$ and $c \in \mathbb{R}$.

Notice, $\mathbb{C}^{m \times n}(\mathbb{R})$ is a **real vector space** whereas $\mathbb{C}^{m \times n}(\mathbb{Q})$ is a **rational vector space**. The choice of scalars is what decides the nomenclature *complex* or *real* or *rational*. Our custom is to let $\mathbb{C}^{m \times n} = \mathbb{C}^{m \times n}(\mathbb{C})$ and $\mathbb{R}^{m \times n} = \mathbb{R}^{m \times n}(\mathbb{R})$ and $\mathbb{Q}^{m \times n} = \mathbb{Q}^{m \times n}(\mathbb{Q})$ by default. If in doubt about the intended choice of scalars for a given problem or example then please ask. I am here to help!

Generalizing a bit:

Example 2.1.5. *Matrices over a field $\mathbb{F}$ form vector spaces over $\mathbb{F}$: in particular, $\mathbb{F}^{m \times n}$ forms a vector space over $\mathbb{F}$ with addition and scalar multiplication of matrices defined as we studied in previous chapters:*

$$(A + B)_{ij} = A_{ij} + B_{ij} \quad \& \quad (cA)_{ij} = cA_{ij}$$

for all $(i, j) \in \mathbb{N}_m \times \mathbb{N}_n$ and $A, B \in \mathbb{F}^{m \times n}$ and $c \in \mathbb{F}$. We understand $\mathbb{F}$ and $\mathbb{F}^n$ as sub-cases.

I think the next example will seem a bit different.

Example 2.1.6. *Let S be a set and denote⁵ the set of all real-valued functions from S to $\mathbb{R}$ by $\mathcal{F}(S, \mathbb{R})$. Let $f, g \in \mathcal{F}(S, \mathbb{R})$ and suppose $c \in \mathbb{R}$, define addition and scalar multiplication of functions by*

$$(f + g)(x) \equiv f(x) + g(x) \quad \& \quad (cf)(x) = cf(x)$$

for all $x \in S$. In short, we define addition and scalar multiplication by the natural "point-wise" rules. This is an example of a function space. Notice that no particular structure is needed for the domain. The vector space structure is inherited from the codomain of the functions. I invite the reader to check Axioms 1-10 for this point-set. For example, define $z(x) = 0$ for all $x \in S$ then we can prove $z + f = f + z = f$ for each $f \in \mathcal{F}(S, \mathbb{R})$. This shows $z : S \rightarrow \mathbb{R}$ serves as the **zero-vector** for $\mathcal{F}(S, \mathbb{R})$.

In the interest of confusing the students, we often write $z = 0$ for the zero-function of the last example. The notation 0 really means just about nothing. Or, perhaps it means everything. For example, 0 is used to denote

$$0, [0, 0], [0, 0, 0], \begin{bmatrix} 0 \\ 0 \end{bmatrix}, \begin{bmatrix} 0 & 0 \\ 0 & 0 \end{bmatrix}, \begin{bmatrix} 0 & 0 & 0 \\ 0 & 0 & 0 \end{bmatrix}$$

in appropriate contexts. Let us move on to a less trivial discussion:

Example 2.1.7. *Let S be a set and let W be a vector space over $\mathbb{F}$ and let $V = \mathcal{F}(S, W)$ denotes the set of functions from S to W . If we define addition and scalar multiplication of functions in V in the same fashion as Example 2.1.6 then once more we have V as a vector space over $\mathbb{F}$. For example, functions from $\mathbb{R}$ to $\mathbb{C}^2$ are naturally viewed as a complex vector space. Or, functions from $\{a, b, c\}$ to $\mathbb{Z}/11\mathbb{Z}$ naturally form a vector space over $\mathbb{Z}/11\mathbb{Z}$.*

Example 2.1.8. *Let $P_2(\mathbb{R}) = \{ax^2 + bx + c \mid a, b, c \in \mathbb{R}\}$, the set of all real polynomials up to quadratic order. Define addition and scalar multiplication by the usual operations on polynomials. Notice that if $ax^2 + bx + c, dx^2 + ex + f \in P_2(\mathbb{R})$ then*

$$(ax^2 + bx + c) + (dx^2 + ex + f) = (a + d)x^2 + (b + e)x + (c + f) \in P_2(\mathbb{R})$$

⁵another popular notation for the set of functions from a set A to a set B is simple B^A . That is, $\mathcal{F}(A, B) = B^A$. In particular, this is in some sense consistent with the notation $\mathbb{R}^3$ as in we can view triples of real numbers as functions from $\{1, 2, 3\}$ to $\mathbb{R}$.

thus $+: P_2(\mathbb{R}) \times P_2(\mathbb{R}) \rightarrow P_2(\mathbb{R})$ (it is a binary operation on $P_2(\mathbb{R})$). Similarly,

$$d(ax^2 + bx + c) = dax^2 + dbx + dc \in P_2(\mathbb{R})$$

thus scalar multiplication maps $\mathbb{R} \times P_2(\mathbb{R}) \rightarrow P_2(\mathbb{R})$ as it ought. Verification of the other 8 axioms is straightforward. We denote the set of polynomials of order n or less via $P_n(\mathbb{R}) = \{a_n x^n + \cdots + a_1 x + a_0 | a_i \in \mathbb{R}\}$. Naturally, $P_n(\mathbb{R})$ also forms a vector space. Finally, if we take the set of all polynomials $\mathbb{R}[x]$ it forms a real vector space. Notice,

$$\mathbb{R} \subset P_1(\mathbb{R}) \subset P_2(\mathbb{R}) \subset P_3(\mathbb{R}) \subset P_4(\mathbb{R}) \subset \cdots \subset \mathbb{R}[x]$$

where $\mathbb{R}$ is naturally identified with the set of constant real polynomials $P_0(\mathbb{R})$

Real polynomials in a different variable are denoted in the natural fashion; for example, $\mathbb{R}[t]$ is the set of real polynomials in the variable t . Furthermore, once the context is clear we are free to drop the $\mathbb{R}$ notation from $P_n(\mathbb{R})$ and simply write P_n . We can generalize the last example by replacing $\mathbb{R}$ with an arbitrary field $\mathbb{F}$:

Example 2.1.9. We denote the set of polynomials with coefficients in $\mathbb{F}$ of order n or less via $P_n(\mathbb{F}) = \{a_n x^n + \cdots + a_1 x + a_0 | a_i \in \mathbb{F}\}$. Naturally, $P_n(\mathbb{F})$ also forms a vector space over $\mathbb{F}$. The set of all polynomials with coefficients in $\mathbb{F}$ is denoted $\mathbb{F}[x]$ and we can show it forms a vector space over $\mathbb{F}$ with respect to the usual addition and scalar multiplication of polynomials.

This list of examples is nowhere near comprehensive. We can also mix together examples we've covered thus far to create new vector spaces.

Example 2.1.10. Let $[P_2(\mathbb{R})]^{2 \times 2}$ denote the set of 2×2 matrices of possibly-degenerate quadratic polynomials. Addition and scalar multiplication are defined in the natural manner:

$$\begin{bmatrix} A_{11}(x) & A_{12}(x) \\ A_{21}(x) & A_{22}(x) \end{bmatrix} + \begin{bmatrix} B_{11}(x) & B_{12}(x) \\ B_{21}(x) & B_{22}(x) \end{bmatrix} = \begin{bmatrix} A_{11}(x) + B_{11}(x) & A_{12}(x) + B_{12}(x) \\ A_{21}(x) + B_{21}(x) & A_{22}(x) + B_{22}(x) \end{bmatrix}$$

and $c \begin{bmatrix} A_{11}(x) & A_{12}(x) \\ A_{21}(x) & A_{22}(x) \end{bmatrix} = \begin{bmatrix} cA_{11}(x) & cA_{12}(x) \\ cA_{21}(x) & cA_{22}(x) \end{bmatrix}$. We can similarly define $m \times n$ matrices of possibly degenerate n -th order real polynomials by $[P_n(\mathbb{R})]^{m \times n}$. Furthermore, we can replace $\mathbb{R}$ with $\mathbb{F}$ to obtain even more varied examples.

Given a pair of vector spaces over a particular field there is a natural way to combine them to make a larger vector space⁶:

Example 2.1.11. Let V, W be vector spaces over $\mathbb{F}$. The Cartesian product $V \times W$ has a natural vector space structure inherited from V and W : if $(v_1, w_1), (v_2, w_2) \in V \times W$ then we define

$$(v_1, w_1) + (v_2, w_2) = (v_1 + v_2, w_1 + w_2) \quad \& \quad c \cdot (v_1, w_1) = (c \cdot v_1, c \cdot w_1)$$

where the vector and scalar operations on the L.H.S. of the above equalities are given from the vector space structure of V and W . All the axioms of a vector space for $V \times W$ are easily verified from the corresponding axioms for V and W .

⁶this is roughly like adding the vector spaces, in contrast, much later we study $V \otimes W$ which is like multiplying the spaces. The tensor product $\otimes$ also gives us a larger new space from a given pair of vector spaces

Example 2.1.12. Let $V_1, V_2, \dots, V_k$ be vector spaces over $\mathbb{F}$. The Cartesian product $V_1 \times V_2 \times \dots \times V_k$ has a natural vector space structure inherited from $V_1, \dots, V_k$: if $x = (x_1, x_2, \dots, x_k), y = (y_1, y_2, \dots, y_k) \in V_1 \times V_2 \times \dots \times V_k$ then we define

$$x + y = (x_1 + y_1, x_2 + y_2, \dots, x_k + y_k) \quad \& \quad c \cdot x = (c \cdot x_1, c \cdot x_2, \dots, c \cdot x_k)$$

where the vector and scalar operations on the L.H.S. of the above equalities are given from the vector space structures of $V_1, \dots, V_k$ respective. All the axioms of a vector space for $V_1 \times V_2 \times \dots \times V_k$ are easily verified from the corresponding axioms for $V_1, \dots, V_k$ respective.

Example 2.1.13. Let $V = \mathbb{R}[x, y] = (\mathbb{R}[x])[y]$ then V contains elements such as

$$1, x, y, x^2, xy, y^2, x^3, x^2y, xy^2, y^3, \dots$$

more generally, suppose $c_{ij} \neq 0$ for only finitely many i, j then $f \in V$ has coefficients c_{ij} and

$$f = \sum_{i,j=0}^{\infty} c_{ij} x^i y^j$$

If f, g have coefficients c_{ij} and b_{ij} respectively then $f + g$ is defined to have coefficients $c_{ij} + b_{ij}$. Notice $f + g$ once more has only finitely many nonzero coefficients and is hence in V . Likewise, define $\alpha f = \alpha \left(\sum_{i,j=0}^{\infty} a_{ij} x^i y^j \right) = \sum_{i,j=0}^{\infty} \alpha a_{ij} x^i y^j$. It follows that V is a real vector space. There are many ways to extend this example.

Example 2.1.14. Let $V = \mathbb{F}[x_1, x_2, \dots, x_n]$ denote the set of multivariate polynomials in $x_1, \dots, x_n$ with coefficients taken from $\mathbb{F}$. If $v \in V$ then there exist finitely many nonzero $c_{i_1, i_2, \dots, i_k} \in \mathbb{F}$ such that

$$v = \sum_{k=0}^{\infty} c_{k, i_1, i_2, \dots, i_k} x_1^{i_1} x_2^{i_2} \dots x_k^{i_k} = \sum_I c_I x^I$$

where I have introduced multi-index notation $x^I = x_1^{i_1} x_2^{i_2} \dots x_k^{i_k}$ and $c_{k, i_1, i_2, \dots, i_k} = c_I$. To be less concise,

$$v = c_0 + \sum_{i=1}^n c_{1,i} x^i + \sum_{i,j=1}^n c_{2,i,j} x^i x^j + \sum_{i,j,k=1}^n c_{3,i,j,k} x^i x^j x^k + \dots$$

If $w = \sum_I b_I x^I$ and $\alpha \in \mathbb{F}$ then we define:

$$v + w = \sum_I (c_I + b_I) x^I \quad \& \quad \alpha v = \sum_I \alpha c_I x^I.$$

Convergence of these seemingly infinite sums is clear since only finitely many multi-indices have either $c_I \neq 0$ or $b_I \neq 0$. In summary, to add multivariate polynomials simply add coefficients of like terms. Verification of all the vector space axioms is a routine exercise. Finally, if you prefer, we could also have denoted $v \in V$ by

$$v = \sum_{i_1, i_2, \dots, i_n=0}^{\infty} b_{i_1, i_2, \dots, i_n} x_1^{i_1} x_2^{i_2} \dots x_n^{i_n} = b_{0,0,\dots,0} + b_{1,0,\dots,0} x_1 + b_{0,1,\dots,0} x_2 + \dots + b_{0,0,\dots,1} x_n + \dots$$

The theorem that follows is full of seemingly obvious facts. I show how each of these facts follow from the vector space axioms.

Theorem 2.1.15.

Let $\mathbb{F}$ be a field and V a vector space over $\mathbb{F}$ with zero vector 0 and let $c \in \mathbb{F}$,

- (1.) $0 \cdot x = 0$ for all $x \in V$,
- (2.) $c \cdot 0 = 0$ for all $c \in \mathbb{F}$,
- (3.) $(-1) \cdot x = -x$ for all $x \in V$,
- (4.) if $cx = 0$ then $c = 0$ or $x = 0$.

Lemma 2.1.16. Law of Cancellation:

Let a, x, y be vectors in a vector space V . If $x + a = y + a$ then $x = y$.

Proof of Lemma: Suppose $x + a = y + a$. By A4 there exists $-a$ such that $a + (-a) = 0$. Thus $x + a = y + a$ implies $(x + a) + (-a) = (y + a) + (-a)$. By A2 we find $x + (a + (-a)) = y + (a + (-a))$ which gives $x + 0 = y + 0$. Continuing we use A3 to obtain $x + 0 = 0$ and $y + 0 = y$ and consequently $x = y$. $\square$

We now seek to prove (1.). Consider:

$$\begin{aligned}
 0 \cdot x + 0 &= 0 \cdot x && \text{by A3} \\
 &= (0 + 0) \cdot x && \text{defn. of zero scalar} \\
 &= 0 \cdot x + 0 \cdot x && \text{by A8}
 \end{aligned}$$

Finally, apply the cancellation lemma to conclude $0 \cdot x = 0$. Note x was arbitrary thus (1.) has been shown true. $\square$

We now prove (2.). Suppose $c \in \mathbb{F}$.

$$\begin{aligned}
 c \cdot 0 + 0 &= c \cdot 0 && \text{by A3} \\
 &= c \cdot (0 + 0) && \text{by A3} \\
 &= c \cdot 0 + c \cdot 0 && \text{by A7}
 \end{aligned}$$

Consequently, by the cancellation lemma we find $c \cdot 0 = 0$ for all $c \in \mathbb{F}$. $\square$

The proof of (3.) is similar. Consider,

$$\begin{aligned}
 0 &= 0 \cdot x && \text{by (1.)} \\
 &= (1 + (-1)) \cdot x && \text{scalar arithmetic} \\
 &= 1 \cdot x + (-1) \cdot x && \text{by A8} \\
 &= x + (-1) \cdot x && \text{by A5}
 \end{aligned}$$

Thus, adding $-x$ to the equation above,

$$(-x) + 0 = (-x) + x + (-1) \cdot x = 0 + (-1) \cdot x$$

thus, using A3 once more, $(-1) \cdot x = -x$ for all $x \in V$. $\square$

To prove (4.), suppose $c \cdot x = 0$. If $c = 0$ then we have that the claim of (4.) is verified. If $c \neq 0$ then $1 = \frac{c}{c} = \frac{1}{c}c$ hence using A5 in the first equality and A6 in the third equality we find:

$$x = 1 \cdot x = \left(\frac{1}{c}c\right) \cdot x = \frac{1}{c} \cdot (c \cdot x) = \frac{1}{c} \cdot 0 = 0$$

where we used (2.) in the final equality. In summary, $c \cdot x = 0$ implies $c = 0$ or $x = 0$. $\square$

Perhaps we should pause to appreciate what was not in the last page or two of proofs. There were no components, no reference to the standard basis. The arguments offered depended only on the definition of the vector space itself. This means the truths we derived above are completely general; they hold for all vector spaces. In what follows past this point we sometimes use Theorem 2.1.15 without explicit reference. That said, I would like you to understand the results of the theorem do require proof and that is why we have taken some effort here to supply that proof.

2.2 subspaces

Definition 2.2.1.

Let V be a vector space over a field $\mathbb{F}$. If $W \subseteq V$ such that W is a vector space over $\mathbb{F}$ with respect to the operations of V restricted to W then we say W is a **subspace** of V and write $W \leq V$.

Example 2.2.2. Let V be a vector space. Notice that $V \subseteq V$ and obviously V is a vector space with respect to its operations. Therefore $V \leq V$. Likewise, the set containing the zero vector $\{0\} \leq V$. Notice that $0 + 0 = 0$ and $c \cdot 0 = 0$ so Axioms 9 and 10 are satisfied. I leave the other axioms to the reader. The subspace $\{0\}$ is called the **trivial subspace**.

Example 2.2.3. Let $L = \{(x, y) \in \mathbb{R}^2 \mid ax + by = 0\}$. Define addition and scalar multiplication by the natural rules in $\mathbb{R}^2$. Note if $(x, y), (z, w) \in L$ then $(x, y) + (z, w) = (x + z, y + w)$ and $a(x + z) + b(y + w) = ax + by + az + bw = 0 + 0 = 0$ hence $(x, y) + (z, w) \in L$. Likewise, if $c \in \mathbb{R}$ and $(x, y) \in L$ then $ax + by = 0$ implies $acx + bcy = 0$ thus $(cx, cy) = c(x, y) \in L$. We find that L is closed under vector addition and scalar multiplication. The other 8 axioms are naturally inherited from $\mathbb{R}^2$. This makes L a **subspace** of $\mathbb{R}^2$.

Example 2.2.4. If $V = \mathbb{R}^3$ then

1. $\{(0, 0, 0)\}$ is a subspace,
2. any line through the origin is a subspace,
3. any plane through the origin is a subspace.

Example 2.2.5. Let $V = \mathbb{R}$ and let $W = [a, b]$ where $a, b > 0$. Observe $a + b \notin [a, b]$ hence $W \not\leq \mathbb{R}$. Indeed, we can generalize this observation, if $W \neq \{0\}$ is a proper subset of $\mathbb{R}$ then it cannot form a subspace of $\mathbb{R}$. Vector addition and scalar multiplication will take us outside W . For example $W = \mathbb{Z}$ has $\sqrt{2}z \notin \mathbb{Z}$ for each $z \in \mathbb{Z}$.

Example 2.2.6. Let $W = \{(x, y, z) \mid x + y + z = 1\}$. Is this a subspace of $\mathbb{R}^3$ with the standard⁷ vector space structure? The answer is no. There are many reasons,

1. $(0, 0, 0) \notin W$ thus W has no zero vector, axiom 3 fails. Notice we cannot change the idea of "zero" for the subspace, if $(0, 0, 0)$ is zero for $\mathbb{R}^3$ then it is the only zero for potential subspaces. Why? Because subspaces inherit their structure from the vector space which contains them.
2. Observe $(1, 0, 0), (0, 1, 0) \in W$ yet $(1, 0, 0) + (0, 1, 0) = (1, 1, 0)$ is not in W since $1 + 1 + 0 = 2 \neq 1$. Thus W is not closed under vector addition (A9 fails).
3. Again $(1, 0, 0) \in W$ yet $2(1, 0, 0) = (2, 0, 0) \notin W$ since $2 + 0 + 0 = 2 \neq 1$. Thus W is not closed under scalar multiplication (A10 fails).

Of course, one reason is all it takes.

My focus on the last two axioms is not without reason. Let me explain this obsession⁸.

Theorem 2.2.7. Subspace Test:

Let V be a vector space over a field $\mathbb{F}$ and suppose $W \subseteq V$ with $W \neq \emptyset$ then $W \leq V$ if and only if the following two conditions hold true

- (1.) if $x, y \in W$ then $x + y \in W$ (W is closed under addition),
- (2.) if $x \in W$ and $c \in \mathbb{F}$ then $c \cdot x \in W$ (W is closed under scalar multiplication).

Proof: ($\Rightarrow$) If $W \leq V$ then W is a vector space with respect to the operations of addition and scalar multiplication thus (1.) and (2.) hold true.

($\Leftarrow$) Suppose W is a nonempty set which is closed under vector addition and scalar multiplication of V . We seek to prove W is a vector space with respect to the operations inherited from V . Let $x, y, z \in W$ then as $W \subseteq V$ we have $x, y, z \in V$. Use A1 and A2 for V (which were given to begin with) to find

$$x + y = y + x \quad \text{and} \quad (x + y) + z = x + (y + z).$$

Thus A1 and A2 hold for W . By (3.) of Theorem 2.1.15 we know that $(-1) \cdot x = -x$ and $-x \in W$ since we know W is closed under scalar multiplication. Consequently, $x + (-x) = 0 \in W$ since W is closed under addition. It follows A3 is true for W . Then by the arguments just given A4 is true for W . Let $a, b \in \mathbb{F}$ and notice that by A5, A6, A7, A8 for V we find

$$1 \cdot x = x, \quad (ab) \cdot x = a \cdot (b \cdot x), \quad a \cdot (x + y) = a \cdot x + a \cdot y, \quad (a + b) \cdot x = a \cdot x + b \cdot x.$$

Thus A5, A6, A7, A8 likewise hold for W . Finally, we assumed closure of addition and scalar multiplication on W so A9 and A10 are likewise satisfied and we conclude that W is a vector space over $\mathbb{F}$. Thus $W \leq V$. (if you're wondering where we needed W nonempty it was to argue that there exists at least one vector x and consequently the zero vector is in W .) $\square$

⁷yes, there is a non-standard addition which gives this space a vector space structure

⁸notice that Charles Curtis uses this Theorem as his definition for subspace. It is equivalent in view of the proof below

Remark 2.2.8.

The application of Theorem 2.2.7 is a four-step process

1. check that $W \subset V$
2. check that $0 \in W$ (this is a matter of convenience, and if it fails it's usually blatant)
3. take arbitrary $x, y \in W$ and show $x + y \in W$
4. take arbitrary $x \in W$ and $c \in \mathbb{F}$ and show $cx \in W$

We usually omit comment about (1.) since it is obviously true for examples we encounter.

Example 2.2.9. The function space $\mathcal{F}(\mathbb{R}) = \mathcal{F}(\mathbb{R}, \mathbb{R})$ has many subspaces.

1. continuous functions: $C(\mathbb{R})$
2. differentiable functions: $C^1(\mathbb{R})$
3. smooth functions: $C^\infty(\mathbb{R})$
4. polynomial functions (which are naturally identified with $\mathbb{R}[x]$)
5. analytic functions
6. solution set of a linear homogeneous ODE with no singular points

The proof that each of these follows from Theorem 2.2.7. For example, $f(x) = x$ is continuous therefore $C(\mathbb{R}) \neq \emptyset$. Moreover, the sum of continuous functions is continuous and a scalar multiple of a continuous function is continuous. Thus $C(\mathbb{R}) \leq \mathcal{F}(\mathbb{R})$. The arguments for (2.), (3.), (4.), (5.) and (6.) are identical. The solution set example is one of the most important examples for engineering and physics, linear ordinary differential equations. Also, we should note that $\mathbb{R}$ can be replaced with some subset I of real numbers. $\mathcal{F}(I)$ likewise has subspaces $C(I), C^1(I), C^\infty(I)$ etc.

Example 2.2.10. Let V, W be vector spaces over a field $\mathbb{F}$. If a function $T : V \rightarrow W$ satisfies

(1.) $T(x + y) = T(x) + T(y)$ for all $x, y \in V$; T is **additive** or T **preserves addition**

(2.) $T(cx) = cT(x)$ for all $x \in V$ and $c \in \mathbb{F}$; T is **preserves scalar multiplication**

then we say T is a **linear transformation** from V to W . The set of all linear transformations from V to W is denoted $\mathcal{L}(V, W)$. Also, $\mathcal{L}(V, V) = \mathcal{L}(V)$ and $T \in \mathcal{L}(V)$ is called a linear transformation **on** V . We know the set of all functions from V to W forms the vector space $\mathcal{F}(V, W)$ with respect to the addition of functions and scalar multiplication defined pointwise: if $f, g \in \mathcal{F}(V, W)$ and $c \in \mathbb{F}$ then

$$(f + g)(x) = f(x) + g(x) \quad \& \quad (cf)(x) = cf(x)$$

for all $x \in V$. Given $T, S \in \mathcal{L}(V, W)$ and $c \in \mathbb{F}$ we can prove $T + S, cT \in \mathcal{L}(V, W)$. Likewise, and $f(x) = 0$ for each $x \in V$ implies $f \in \mathcal{L}(V, W)$. Identify f serves as the additive identity in $\mathcal{L}(V, W) \neq \emptyset$. Thus $\mathcal{L}(V, W) \leq \mathcal{F}(V, W)$ by the subspace test. This is a very important example, much of future chapters are centered around its study.

Example 2.2.11. Let $Ax = 0$ denote a homogeneous system of m -equations in n -unknowns over $\mathbb{F}$. Let W be the solution set of this system; $W = \{x \in \mathbb{F}^n \mid Ax = 0\}$. Observe that $A0 = 0$ hence $0 \in W$ so the solution set is nonempty. Suppose $x, y \in W$ and $c \in \mathbb{F}$,

$$A(x + cy) = Ax + cAy = 0 + c(0) = 0$$

thus $x + cy \in W$. Closure of addition for W follows from $c = 1$ and closure of scalar multiplication follows from $x = 0$ in the just completed calculation. Thus $W \leq \mathbb{F}^n$ by the Subspace Test Theorem. We define the **null space** of A by:

$$\text{Null}(A) = \{x \in \mathbb{F}^n \mid Ax = 0\}$$

This example proves $\text{Null}(A) \leq \mathbb{F}^n$.

Sometimes it's easier to check both scalar multiplication and addition at once. It saves some writing. If you don't understand it then don't use the trick I just used, we should understand our work.

The example that follows here illustrates an important point in abstract math. Given a particular point set, there is often more than one way to define a structure on the set. Therefore, it is important to view things as more than mere sets. Instead, think about sets paired with a structure.

Example 2.2.12. Let⁹ V_p be the set of all vectors with base point $p \in \mathbb{R}^n$,

$$V_p = \{p + v \mid v \in \mathbb{R}^n\}$$

We define a nonstandard vector addition on V_p , if $p + v, p + w \in V_p$ and $c \in \mathbb{R}$ define:

$$(p + v) +_p (p + w) = p + v + w \quad \& \quad c \cdot_p (p + v) = p + cv.$$

Clearly $+_p : V_p \times V_p \rightarrow V_p$ and $\cdot_p : \mathbb{R} \times V_p \rightarrow V_p$ are closed and verification of the other axioms is straightforward. Observe $0_p = p$ as $(p + v) +_p (p + 0) = p + v + 0 = p + v$ hence $0_p = p + 0 = p$. Mainly, the vector space axioms for V_p follow from the corresponding axioms for $\mathbb{R}^n$. Geometrically, $+_p$ corresponds to the **tip-to-tail** rule we use in physics to add vectors. Consider S_p defined below:

$$S_p = \{p + v \mid v \in W \leq \mathbb{R}^n\}$$

Notice $0_p \in S_p$ as $0 \in W$ and $0_p = p + 0$. Furthermore, consider $p + v, p + w \in S_p$ and $c \in \mathbb{R}$

$$(p + v) +_p (p + w) = p + (v + w) \quad \& \quad c \cdot_p (p + v) = p + cv$$

note $v + w, cv \in W$ as $W \leq \mathbb{R}^n$ is closed under addition and scalar multiplication. We find $(p + v) +_p (p + w), c \cdot_p (p + v) \in S_p$ thus $S_p \leq V_p$ by the subspace test Theorem 2.2.7.

In the previous example, S_p need not be a subspace with respect to the standard vector addition of column vectors. However, with the modified addition based at p it is a subspace. We often say the solution set to $Ax = b$ with $b \neq 0$ is not a subspace. It should be understood that what is meant is that the solution set of $Ax = b$ is not a subspace with respect to the usual vector addition. It is possible to define a different vector addition which gives the solution set of $Ax = b$ a vector space structure.

⁹it may be better to use the notation (p, v) for $p + v$, this has the advantage of making the base-point p explicit whereas p can be obscured in the more geometrically direct $p + v$ notation. Another choice is to use v_p .

Example 2.2.13. Let $W = \{A \in \mathbb{R}^{n \times n} \mid A^T = A\}$. This is the set of **symmetric matrices**, it is nonempty since $I^T = I$ (of course there are many other examples, we only need one to show it's nonempty). Let $A, B \in W$ and suppose $c \in \mathbb{R}$ then

$$\begin{aligned} (A + B)^T &= A^T + B^T && \text{prop. of transpose} \\ &= A + B && \text{since } A, B \in W \end{aligned}$$

thus $A + B \in W$ and we find W is closed under addition. Likewise let $A \in W$ and $c \in \mathbb{R}$,

$$\begin{aligned} (cA)^T &= cA^T && \text{prop. of transpose} \\ &= cA && \text{since } A \in W \end{aligned}$$

thus $cA \in W$ and we find W is closed under scalar multiplication. Therefore, by the subspace test Theorem 2.2.7, $W \leq \mathbb{R}^{n \times n}$.

I invite the reader to modify the example above to show the set of antisymmetric matrices also forms a subspace of the vector space of square matrices.

Example 2.2.14. Let $W = \{f \in \mathcal{F}(\mathbb{R}) \mid \int_{-1}^1 f(x) dx = 0\}$. Notice the zero function $0(x) = 0$ is in W since $\int_{-1}^1 0 dx = 0$. Let $f, g \in W$, use linearity property of the definite integral to calculate

$$\int_{-1}^1 (f(x) + g(x)) dx = \int_{-1}^1 f(x) dx + \int_{-1}^1 g(x) dx = 0 + 0 = 0$$

thus $f + g \in W$. Likewise, if $c \in \mathbb{R}$ and $f \in W$ then

$$\int_{-1}^1 cf(x) dx = c \int_{-1}^1 f(x) dx = c(0) = 0$$

thus $cf \in W$ and by subspace test Theorem 2.2.7 $W \leq \mathcal{F}(\mathbb{R})$.

Example 2.2.15. Here we continue discussion of the product space introduced in Example 2.1.11. Suppose $V = \mathbb{C}$ and $W = P_2$ then $V \times W = \{(a + ib, cx^2 + dx + e) \mid a, b, c, d, e \in \mathbb{R}\}$. Let $U = \{(a, b) \mid a, b \in \mathbb{R}\}$. We can easily show $U \leq V \times W$ by the subspace test Theorem 2.2.7 $W \leq \mathcal{F}(\mathbb{R})$. Can you think of other subspaces? Is it possible to have a subspace of $V \times W$ which is not formed from a pair of subspaces from V and W respective?

Example 2.2.16. Let W be the set of real-valued functions on $\mathbb{R}$ for which $f(a) = 0$ for some fixed value $a \in \mathbb{R}$. If $f, g \in W$ and $c \in \mathbb{R}$ then $(f + cg)(a) = f(a) + cg(a) = 0 + c(0) = 0$ thus $f + cg \in W$. Observe W is closed under addition by the case $c = 1$ and W is closed under scalar multiplication by the case $f = 0$. Furthermore, $f(x) = 0$ for all $x \in \mathbb{R}$ defines the zero function which is in W . Hence $W \leq \mathcal{F}(\mathbb{R})$ by subspace test Theorem 2.2.7.

Example 2.2.17. Let W be the set of solutions of the differential equation $ay'' + by' + cy = 0$ where $a, b, c \in \mathbb{R}$ and $a \neq 0$. Notice that the zero function $y = 0$ solves $ay'' + by' + cy = 0$ thus $\emptyset \neq W \subseteq \mathcal{F}(\mathbb{R})$. Suppose $y_1, y_2 \in W$ and let $\alpha \in \mathbb{R}$. We are given that $ay_1'' + by_1' + cy_1 = 0$ and $ay_2'' + by_2' + cy_2 = 0$. Consider, by properties of the first and second derivative,

$$\begin{aligned} a(\alpha y_1 + y_2)'' + b(\alpha y_1 + y_2)' + c(\alpha y_1 + y_2) &= \alpha(ay_1'' + by_1' + cy_1) + ay_2'' + by_2' + cy_2 \\ &= \alpha(0) + 0 \\ &= 0 \end{aligned}$$

Thus $\alpha y_1 + y_2 \in W$. Therefore, by subspace test Theorem 2.2.7, $W \leq \mathcal{F}(\mathbb{R})$.

Example 2.2.18. Let V_1, V_2 be vector spaces over $\mathbb{F}$. We can show $\{0\} \times V_2 \leq V_1 \times V_2$. Since $0 \in V_2$ we find $(0, 0) \in \{0\} \times V_2 \neq \emptyset$. Suppose $x, y \in \{0\} \times V_2$ and $c \in \mathbb{F}$. Then there exists $x_2, y_2 \in V_2$ for which $x = (0, x_2)$ and $y = (0, y_2)$. Thus,

$$\alpha x + y = \alpha(0, x_2) + (0, y_2) = (\alpha 0, \alpha x_2) + (0, y_2) = (0, \alpha x_2 + y_2)$$

Since V_2 is a vector space we have $\alpha x_2 + y_2 \in V_2$ thus $\alpha x + y = (0, \alpha x_2 + y_2) \in \{0\} \times V_2$ and we conclude $\{0\} \times V_2 \leq V_1 \times V_2$ by the subspace test Theorem 2.2.7.

2.3 spanning sets and subspaces

In a vector space V over a field $\mathbb{F}$ we are free to form $\mathbb{F}$ -linear combinations¹⁰ of vectors; we say $v \in V$ is a linear combination of $v_1, \dots, v_k \in V$ if there exist $c_1, \dots, c_k \in \mathbb{F}$ such that $v = c_1 v_1 + \dots + c_k v_k$. In other words, $v = \sum_{j=1}^k c_j v_j$ is a **finite linear combination** of vectors v_j in V with coefficients c_j in $\mathbb{F}$. The Lemma below is quite useful¹¹:

Lemma 2.3.1. Let V be a vector space over the field $\mathbb{F}$ and let S be a nonempty subset of V ,

The finite $\mathbb{F}$ -linear combination of finite $\mathbb{F}$ -linear combinations of vectors from S is once more a finite $\mathbb{F}$ -linear combination of vectors from S .

Proof: Suppose V is a vector space over a field $\mathbb{F}$. Let $s_i = \sum_{j=1}^{n_i} c_{ij} t_{ij}$ where $c_{ij} \in \mathbb{F}$ and $t_{ij} \in S$ for $n_i, i \in \mathbb{N}$ with $i = 1, 2, \dots, k$. Let $b_1, \dots, b_k \in \mathbb{F}$ and consider by (2.) of Proposition ??

$$\sum_{i=1}^k b_i s_i = \sum_{i=1}^k b_i \left(\sum_{j=1}^{n_i} c_{ij} t_{ij} \right) = \sum_{i=1}^k \sum_{j=1}^{n_i} b_i c_{ij} t_{ij}.$$

Notice, this is a $\mathbb{F}$ -linear combination of vectors in S as $b_i c_{ij} \in \mathbb{F}$. $\square$

Definition 2.3.2.

Let V be a vector space over $\mathbb{F}$ and suppose $S \subseteq V$. Then $\text{span}(S)$ is defined to be the set of all **finite**- $\mathbb{F}$ -linear combinations of vectors taken from S ;

$$\text{span}(S) = \bigcup_{k \in \mathbb{N}} \left\{ \sum_{i=1}^k c_i v_i \mid c_i \in \mathbb{F}, v_i \in S \right\}.$$

If $W = \text{span}(S)$ then we say that S is a **generating set** for W . We also say S **spans** W in this case. To be clear, if $S = \emptyset$ then $\text{span}(\emptyset) = \{0\}$.

Note, in the case that $|S| < \infty$ we can write $S = \{v_1, \dots, v_n\}$ and the definition above simply reads:

$$\text{span}\{v_1, \dots, v_n\} = \{c_1 v_1 + \dots + c_n v_n \mid c_1, \dots, c_n \in \mathbb{F}\}.$$

However, the infinite case is important. A rather famous example is polynomials.

Example 2.3.3. If $f \in \mathbb{F}[x]$ then $f(x) = a_0 + a_1 x + \dots + a_n x^n$ for $a_i \in \mathbb{F}$ for $i = 0, \dots, n$. But, $n \in \mathbb{N}$ can be as large as we wish. In fact, $\mathbb{F}[x] = \text{span}(S)$ where $S = \{1, x^i \mid i \in \mathbb{N}\}$.

¹⁰following Definition 1.2.14 with $R = \mathbb{F}$

¹¹Charles Curtis' proof of (4.4) on page 27-28 has an induction-based refinement of proof I offer for the Lemma

Example 2.3.4. If we set $R = \mathbb{F}$ then Proposition 1.2.15 explains how $\mathbb{F}^n$ was spanned by the standard basis; $\mathbb{F}^n = \text{span}\{e_i\}_{i=1}^n$. Likewise, Proposition 1.2.17 showed the $m \times n$ matrix units E_{ij} spanned the set of all $m \times n$ matrices; $\mathbb{F}^{m \times n} = \text{span}\{E_{ij}\}_{i,j=1}^n$.

Spans are important because they are subspaces which are presented in a particularly lucid manner.

Theorem 2.3.5. $\text{span}(S)$ is a subspace.

Let V be a vector space over a field $\mathbb{F}$ and suppose $S \subseteq V$ then $\text{span}(S) \leq V$. Furthermore, if $W \leq V$ and $S \subseteq W$ then $S \subseteq \text{span}(S) \subseteq W$; that is, $\text{span}(S)$ is the smallest subspace of V which contains S .

Proof: If $S = \emptyset$ then $\text{span}(\emptyset) = \{0\} \leq V$. Otherwise, $S \neq \emptyset$ hence consider $x, y \in \text{span}(S)$ and $c \in \mathbb{F}$. Apply Lemma 2.3.1 to see the linear combination of linear combinations $x + y$ and cx is once more a linear combination of vectors in S . Thus $x + y, cx \in \text{span}(S)$ and we conclude by the Subspace Test Theorem $\text{span}(S) \leq V$.

Suppose W is any subspace of V which contains S . By definition W is closed under scalar multiplication and vector addition thus all linear combinations of vectors from S must be in W hence $\text{span}(S) \subseteq W$. Lastly, if $v \in S$ then $v = 1 \cdot v \in \text{span}(S)$ hence $S \subseteq \text{span}(S)$. $\square$

Example 2.3.6. Let $S = \{1, x, x^2, \dots, x^n\}$ then $\text{span}_{\mathbb{F}}(S) = P_n(\mathbb{F})$. For example,

$$\text{span}_{\mathbb{R}}\{1, x, x^2\} = \{ax^2 + bx + c \mid a, b, c \in \mathbb{R}\} = P_2(\mathbb{R})$$

Notice, $P_0(\mathbb{F}) \leq P_1(\mathbb{F}) \leq P_2(\mathbb{F}) \leq \dots \leq P_n(\mathbb{F}) \leq \dots \leq \mathbb{F}[x]$.

Example 2.3.7. Let $W = \{(s + t, 2s + t, 3s + t) \mid s, t \in \mathbb{R}\}$. Observe,

$$(s + t, 2s + t, 3s + t) = s(1, 2, 3) + t(1, 1, 1)$$

thus $W = \{s(1, 2, 3) + t(1, 1, 1) \mid s, t \in \mathbb{R}\} = \text{span}\{(1, 2, 3), (1, 1, 1)\}$. Therefore, Theorem 2.3.5 gives us $W \leq \mathbb{R}^3$.

The lesson of the last example is that we can show a particular space is a subspace by finding its generating set. Theorem 2.3.5 tells us that any set generated by a span is a subspace. This test is only convenient for subspaces which are defined as some sort of span. In that case we can immediately conclude the subset is in fact a subspace.

Example 2.3.8. Consider $y' = y$. Or, taking t as the independent variable, $\frac{dy}{dt} = y$. Separation of variables (that you are expected to know from calculus II) shows $\frac{dy}{y} = dt$ hence $\ln|y| = t + c$. It follows that $y = \pm e^c e^t$. Note $y = 0$ is also a solution of $y' = y$. In total, we find solutions of the form $y = c_1 e^t$. The solution set of this differential equation is a span; $S = \text{span}\{e^t\} \leq \mathcal{F}(\mathbb{R})$.

Example 2.3.9. Consider, $y''' = 0$. Integrate both sides to find $y'' = c_1$. Integrate again to find $y' = c_1 t + c_2$. Integrate once more, $y = c_1 \frac{1}{2} t^2 + c_2 t + c_3$. The general solution of $y''' = 0$ is a subspace S of function space:

$$S = \text{span}\left\{\frac{1}{2}t^2, t, 1\right\} \leq \mathcal{F}(\mathbb{R})$$

Physically, we often consider the situation $c_1 = -g$.

The analysis in Examples 2.3.8 and 2.3.9 simply derive from combining prerequisite calculus knowledge with linear algebra. In our Differential Equations course you learn how to construct the fundamental solution set for an n -th order homogeneous constant coefficient differential equation. It turns out that any solution can be written $y = c_1y_1 + \cdots + c_ny_n$ thus $\text{span}\{y_1, \dots, y_n\}$ is the solution set of the differential equation. If you've already taken differential equations then you can gain much intuition for linear algebra from your previous work on differential equations. On the other hand, if you haven't taken differential equations then the good news is that course becomes even easier when you apply the theory of linear algebra.

Example 2.3.10. Subspaces associated with a given matrix: Let $A \in \mathbb{F}^{m \times n}$. Define **column space** of A as the *span* of the columns of A :

$$\text{Col}(A) = \text{span}\{\text{col}_j(A) \mid j = 1, 2, \dots, n\}$$

this is clearly a subspace of $\mathbb{F}^m$ since each column has as many components as there are rows in A . We also define **row space** as the *span* of the rows:

$$\text{Row}(A) = \text{span}\{\text{row}_i(A) \mid i = 1, 2, \dots, m\}$$

this is clearly a subspace of $\mathbb{R}^{1 \times n}$ since it is formed as a *span* of vectors. Since the columns of A^T are the rows of A and the rows of A^T are the columns of A we can conclude that $\text{Col}(A^T) = \text{Row}(A)$ and $\text{Row}(A^T) = \text{Col}(A)$. Finally, we already defined the **null space** of A by:

$$\text{Null}(A) = \{x \in \mathbb{F}^n \mid Ax = 0\}$$

We studied this in Example 2.2.11 where we showed $\text{Null}(A) \leq \mathbb{F}^n$. With some effort and the insight of the $\text{rref}(A)$ we can write the null space as a *span* of an appropriate set of solutions to $Ax = 0$. See Example 2.2.11 for instance.

I would remind the reader we have the CCP and associated techniques to handle spanning questions for column vectors. In contrast, the following example requires a direct assault¹²:

Example 2.3.11. Is $E_{11} \in \text{span}\{E_{12} + 2E_{11}, E_{12} - E_{11}\}$? Assume $E_{ij} \in \mathbb{R}^{2 \times 2}$ for all i, j . We seek to find solutions of

$$E_{11} = a(E_{12} + 2E_{11}) + b(E_{12} - E_{11})$$

in explicit matrix form the equation above reads:

$$\begin{aligned} \begin{bmatrix} 1 & 0 \\ 0 & 0 \end{bmatrix} &= a \left(\begin{bmatrix} 0 & 1 \\ 0 & 0 \end{bmatrix} + \begin{bmatrix} 2 & 0 \\ 0 & 0 \end{bmatrix} \right) + b \left(\begin{bmatrix} 0 & 1 \\ 0 & 0 \end{bmatrix} + \begin{bmatrix} -1 & 0 \\ 0 & 0 \end{bmatrix} \right) \\ &= \begin{bmatrix} 2a & a \\ 0 & 0 \end{bmatrix} + \begin{bmatrix} -b & b \\ 0 & 0 \end{bmatrix} \\ &= \begin{bmatrix} 2a - b & a + b \\ 0 & 0 \end{bmatrix} \end{aligned}$$

¹²However, once we have the idea of *coordinates* ironed out then we can use the CCP tricks on the coordinate vectors then push back the results to the world of abstract vectors. For now we'll just confront each question by brute force. For an example such as this, the method used here is as good as our later methods.

thus $1 = 2a - b$ and $0 = a + b$. Substitute $a = -b$ to find $1 = 3a$ hence $a = \frac{1}{3}$ and $b = -\frac{1}{3}$. Indeed,

$$\frac{1}{3}(E_{12} + 2E_{11}) - \frac{1}{3}(E_{12} - E_{11}) = \frac{2}{3}E_{11} + \frac{1}{3}E_{11} = E_{11}.$$

Therefore, $E_{11} \in \text{span}\{E_{12} + 2E_{11}, E_{12} - E_{11}\}$.

Example 2.3.12. Find a generating set for the set of symmetric 2×2 matrices. That is find a set S of matrices such that $\text{span}(S) = \{A \in \mathbb{R}^{2 \times 2} \mid A^T = A\} = W$. There are many approaches, but I find it most natural to begin by studying the condition which defines W . Let $A = \begin{bmatrix} a & b \\ c & d \end{bmatrix} \in W$ and note

$$A^T = A \Leftrightarrow \begin{bmatrix} a & c \\ b & d \end{bmatrix} = \begin{bmatrix} a & b \\ c & d \end{bmatrix} \Leftrightarrow A = \begin{bmatrix} a & b \\ b & d \end{bmatrix}.$$

In summary, we find an equivalent way to express $A \in W$ is as:

$$A = \begin{bmatrix} a & b \\ b & d \end{bmatrix} = a \begin{bmatrix} 1 & 0 \\ 0 & 0 \end{bmatrix} + b \begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix} + d \begin{bmatrix} 0 & 0 \\ 0 & 1 \end{bmatrix} = aE_{11} + b(E_{12} + E_{21}) + dE_{22}.$$

Consequently $W = \text{span}\{E_{11}, E_{12} + E_{21}, E_{22}\}$ and the set $\{E_{11}, E_{12} + E_{21}, E_{22}\}$ generates W . This is not unique, there are many other sets which also generate W . For example, if we took $\bar{S} = \{E_{11}, E_{12} + E_{21}, E_{22}, E_{11} + E_{22}\}$ then the span of $\bar{S}$ would still work out to W .

2.4 linear independence

We have seen a variety of generating sets in the preceding section. In the last example I noted that if we added an additional vector $E_{11} + E_{22}$ then the same span would be created. The vector $E_{11} + E_{22}$ is **redundant** since we already had E_{11} and E_{22} . In particular, $E_{11} + E_{22}$ is a linear combination of E_{11} and E_{22} so adding it will not change the span. How can we decide if a vector is absolutely necessary for a span? In other words, if we want to span a subspace W then how do we find a **minimal spanning set**? We want a set of vectors which does not have any linear dependences. We say such vectors are linearly independent. Let me be precise¹³:

Definition 2.4.1.

Let $S \subseteq V$ a vector space over a field $\mathbb{F}$. The set of vectors S is **Linearly Independent (LI)** iff for any $\{v_1, v_2, \dots, v_k\} \subseteq S$ and $c_1, \dots, c_k \in \mathbb{F}$,

$$c_1v_1 + c_2v_2 + \dots + c_kv_k = 0 \Rightarrow c_1 = c_2 = \dots = c_k = 0.$$

Furthermore, to be clear, $\emptyset$ is LI. Moreover, if S is not LI then S is **linearly dependent**.

Example 2.4.2. Let $v = \cos^2(t)$ and $w = 1 + \cos(2t)$. Clearly v, w are linearly dependent since $w - 2v = 0$. We should remember from trigonometry $\cos^2(t) = \frac{1}{2}(1 + \cos(2t))$.

Proposition 2.4.3.

¹³if you have a sense of deja vu here, it is because I uttered many of the same words in the context of $\mathbb{R}^n$. Notice, in contrast, I now consider the abstract case. We cannot use the CCP directly here

If S is a finite set of vectors which contains the zero vector then S is linearly dependent.

Proof: Let $\{\vec{0}, v_2, \dots, v_k\} = S$ and observe that $1 \cdot \vec{0} = 0$ thus S is not linearly independent, that is, S is linearly dependent. $\square$

Proposition 2.4.4.

Let v and w be nonzero vectors.

$$v, w \text{ are linearly dependent} \Leftrightarrow \exists k \neq 0 \in \mathbb{F} \text{ such that } v = kw.$$

Proof: Let $v, w \neq 0$. Suppose $\{v, w\}$ is linearly dependent. Hence we find $c_1, c_2 \in \mathbb{F}$, not both zero, for which $c_1 v + c_2 w = 0$. Without loss of generality suppose $c_1 \neq 0$ then $v = -\frac{c_2}{c_1} w$ so identify $k = -c_2/c_1$ and $v = kw$. Conversely, if $v = kw$ with $k \neq 0$ then $v - kw = 0$ shows $\{v, w\}$ is linearly dependent. $\square$

It turns out the claim in the Proposition below does not generalize to the study of modules¹⁴ so this is why we should use the Definition above as stated. Furthermore, this definition is oft used for calculations in our study of linear independence.

Proposition 2.4.5.

$v_j \in \text{span}\{v_1, \dots, v_{j-1}, v_{j+1}, \dots, v_k\}$ iff $\{v_1, v_2, \dots, v_k\}$ is a **linearly dependent** set.

Proof: ($\Rightarrow$) if $v_j \in \text{span}\{v_1, \dots, v_{j-1}, v_{j+1}, \dots, v_k\}$ then there exist $c_i \in \mathbb{F}$ for $i \neq j$ such that $v_j = \sum_{i \neq j} c_i v_i$ thus

$$c_1 v_1 + \dots + c_{j-1} v_{j-1} - v_j + c_{j+1} v_{j+1} + \dots + c_k v_k = 0$$

therefore $\{v_1, v_2, \dots, v_k\}$ is linearly dependent as there is a nontrivial linear combination whose sum is zero.

($\Leftarrow$) Suppose $\{v_1, v_2, \dots, v_k\}$ is linearly dependent. Then there exist $c_1, \dots, c_k \in \mathbb{F}$ for which

$$c_1 v_1 + c_2 v_2 + \dots + c_k v_k = 0$$

and at least one constant, say c_j , is nonzero. Then we can divide by c_j to obtain

$$\frac{c_1}{c_j} v_1 + \frac{c_2}{c_j} v_2 + \dots + v_j + \dots + \frac{c_k}{c_j} v_k = 0 \Rightarrow v_j = -\frac{c_1}{c_j} v_1 - \frac{c_2}{c_j} v_2 - \dots - \widehat{v_j} - \dots - \frac{c_k}{c_j} v_k$$

where $\widehat{v_j}$ denotes the deletion of v_j from the list. Thus $v_j \in \text{span}\{v_1, \dots, v_{j-1}, v_{j+1}, \dots, v_k\}$. $\square$

Given a set of vectors in $\mathbb{F}^n$ the question of LI is elegantly answered by the CCP. In examples that follow in this section we leave the comfort zone and study LI in abstract vector spaces. For now we only have brute force at our disposal. In other words, I'll argue directly from the definition without the aid of the CCP from the outset.

¹⁴A **module** is like a vector space, except, we use scalars from a ring. For example, if we consider pairs from $\mathbb{Z}/6\mathbb{Z}$ we have $2(3, 0) + 3(0, 2) = 0$ hence $\{(3, 0), (0, 2)\}$ is linearly dependent. But, note we cannot solve for $(3, 0)$ as a multiple of $(0, 2)$. In short, the solving for a vector criteria is special to using scalars from a field.

Example 2.4.6. Suppose $f(x) = \cos(x)$ and $g(x) = \sin(x)$ and define $S = \{f, g\}$. Is S linearly independent with respect to the standard vector space structure on $\mathcal{F}(\mathbb{R})$? Let $c_1, c_2 \in \mathbb{R}$ and assume that

$$c_1 f + c_2 g = 0.$$

It follows that $c_1 f(x) + c_2 g(x) = 0$ for each $x \in \mathbb{R}$. In particular,

$$c_1 \cos(x) + c_2 \sin(x) = 0$$

for each $x \in \mathbb{R}$. Let $x = 0$ and we get $c_1 \cos(0) + c_2 \sin(0) = 0$ thus $c_1 = 0$. Likewise, let $x = \pi/2$ to obtain $c_1 \cos(\pi/2) + c_2 \sin(\pi/2) = 0 + c_2 = 0$ hence $c_2 = 0$. We have shown that $c_1 f + c_2 g = 0$ implies $c_1 = c_2 = 0$ thus $S = \{f, g\}$ is a linearly independent set.

Example 2.4.7. Let $f_n(t) = t^n$ for $n = 0, 1, 2, \dots$. Suppose $S = \{f_0, f_1, \dots, f_n\}$. Show S is a linearly independent subset of function space. Assume $c_0, c_1, \dots, c_n \in \mathbb{R}$ and

$$c_0 f_0 + c_1 f_1 + c_2 f_2 + \dots + c_n f_n = 0. \quad \star$$

I usually skip the expression above, but I'm including this extra step to emphasize the distinction between the function and its formula. The $\star$ equation is a function equation, it implies

$$c_0 + c_1 t + c_2 t^2 + \dots + c_n t^n = 0 \quad \star \star$$

for all $t \in \mathbb{R}$. Evaluate $\star \star$ at $t = 0$ to obtain $c_0 = 0$. Differentiate $\star \star$ and find

$$c_1 + 2c_2 t + \dots + nc_n t^{n-1} = 0 \quad \star^3$$

Evaluate $\star^3$ at $t = 0$ to obtain $c_1 = 0$. If we continue to differentiate and evaluate we will similarly obtain $c_2 = 0$, $c_3 = 0$ and so forth all the way up to $c_n = 0$. Therefore, $\star$ implies $c_0 = c_1 = \dots = c_n = 0$.

Linear dependence in function space is sometimes a source of confusion for students. The idea of evaluation doesn't help in the same way as it just has in the two examples above.

Example 2.4.8. Let $f(t) = t - 1$ and $g(t) = t + t^2$ is f linearly dependent on g ? A common mistake is to say something like $f(1) = 1 - 1 = 0$ so $\{f, g\}$ is linearly independent since it contains zero. Why is this wrong? The reason is that we have confused the value of the function with the function itself. If $f(t) = 0$ for all $t \in \mathbb{R}$ then f is the zero function which is the zero vector in function space. Many functions will be zero at a point but that doesn't make them the zero function. To prove linear dependence we must show that there exists $k \in \mathbb{R}$ such that $f = kg$, but this really means that $f(t) = kg(t)$ for all $t \in \mathbb{R}$ in the current context. I leave it to the reader to prove that $\{f, g\}$ is in fact LI. You can evaluate at $t = 1$ and $t = 0$ to obtain equations for c_1, c_2 which have a unique solution of $c_1 = c_2 = 0$.

Example 2.4.9. Let $f(t) = t^2 - 1$, $g(t) = t^2 + 1$ and $h(t) = 4t^2$. Suppose

$$c_1(t^2 - 1) + c_2(t^2 + 1) + c_3(4t^2) = 0 \quad \star$$

A little algebra reveals,

$$(c_1 + c_2 + 4c_3)t^2 - (c_1 - c_2)1 = 0$$

Using linear independence of t^2 and 1 we find

$$c_1 + c_2 + 4c_3 = 0 \quad \text{and} \quad c_1 - c_2 = 0$$

We find infinitely many solutions,

$$c_1 = c_2 \quad \text{and} \quad c_3 = -\frac{1}{4}(c_1 + c_2) = -\frac{1}{2}c_2$$

Therefore, $\star$ allows nontrivial solutions. Take $c_2 = 1$,

$$1(t^2 - 1) + 1(t^2 + 1) - \frac{1}{2}(4t^2) = 0.$$

We can write one of these functions as a linear combination of the other two,

$$f = -g + \frac{1}{2}h.$$

Once we get past the formalities of the particular vector space structure it always comes back to solving systems of linear equations.

Remark 2.4.10.

We should keep in mind that in the abstract context statements such as " v and w go in the same direction" or " u is contained in the plane spanned by v and w " are not statements about ordinary three dimensional geometry. Moreover, you **cannot** write that $u, v, w \in \mathbb{R}^n$ unless you happen to be working with that rather special vector space. I caution the reader against the common mistake of trying to use column calculation techniques for things which are not columns. This is a very popular mistake in past years.

2.5 theory of dimension

Thanks to my brother William Cook of Appalachian State University, Boone North Carolina. This section of notes is nearly identical to a handout he gives to his Linear Algebra classes. In past editions of my notes, the focus was almost entirely on finite dimensional vector spaces. Here we learn that infinite dimensional vector spaces have dimension specified by cardinality. Let's go!

We can contrast the interplay of linear independence and spanning with subsets and supersets. If $S_1 \subseteq S_2$ and S_2 is linearly independent then S_1 is linearly independent. Likewise, if $S_1 \subseteq S_2 \subseteq W$ and $\text{span}(S_1) = W$ then $\text{span}(S_2) = W$; that is supersets of spanning sets still span¹⁵

Lemma 2.5.1.

Let S be a linearly independent subset of V , $v \in V$ such that $v \notin S$. Then $S \cup \{v\}$ is linearly independent if and only if $v \notin \text{span}(S)$.

Proof: Suppose $S \cup \{v\}$ is linearly independent. Then v cannot depend linearly on S . Thus $v \notin \text{span}(S)$.

Conversely, suppose $v \notin \text{span}(S)$. Suppose $S \cup \{v\}$ were linearly dependent. Then we could write $c_1v_1 + \cdots + c_\ell v_\ell + sv = 0$ for some $v_1, \dots, v_\ell \in S$ and scalars $c_1, \dots, c_\ell, s$ not all zero. Notice that $s \neq 0$ would imply $v = s^{-1}c_1v_1 + \cdots + s^{-1}c_\ell v_\ell \in \text{span}(S)$ (contradiction) so $s = 0$. But then $c_1v_1 + \cdots + c_\ell v_\ell = 0$ with not all $c_1, \dots, c_\ell$. This means S is linearly dependent! (contradiction). Thus $S \cup \{v\}$ must be independent. $\square$

¹⁵if we drop the condition $S_2 \subseteq W$ and merely suppose $S_1 \subseteq S_2$ then we can easily prove $\text{span}(S_1) \leq \text{span}(S_2)$

Theorem 2.5.2.

Let S be a linearly independent subset of T where T spans V . Then there exists some β such that $S \subseteq \beta \subseteq T$ where β is both linearly independent and a spanning set for V . In other words, between any independent and spanning set, we can find a basis.

Proof: Consider the set $\mathcal{S} = \{S' \subseteq V \mid S \subseteq S' \subseteq T \text{ and } S' \text{ is linearly independent}\}$. Notice that S itself belongs to $\mathcal{S}$, so $\mathcal{S}$ is a non-empty set and is partially ordered by the subset relation. Suppose that we have a chain of elements in $\mathcal{S}$, say $\mathcal{C}$: this means that for all $S', S'' \in \mathcal{C}$ either $S' \subseteq S''$ or $S'' \subseteq S'$ (i.e., $\mathcal{C}$ is totally ordered by the subset relation). Consider $C = \cup \mathcal{C}$ (i.e., C is the union of all of the sets in our chain). Now every set in $\mathcal{C}$ contains S and is contained in T , so this is true of the union of such sets (i.e., $S \subseteq C \subseteq T$). Suppose $v_1, \dots, v_\ell \in C$ and $c_1, \dots, c_\ell \in \mathbb{F}$ such that $c_1 v_1 + \dots + c_\ell v_\ell = 0$. Then for each $i = 1, \dots, \ell$, we have $v_i \in S_i$ for some $S_i \in \mathcal{C}$. But $\mathcal{C}$ is a chain so we can (possibly after relabeling – without loss of generality) assume $S_1 \subseteq S_2 \subseteq \dots \subseteq S_\ell$. Thus $v_1, \dots, v_\ell \in S_\ell$. But ultimately S_ℓ is a set in $\mathcal{S}$ and so is linearly independent. Thus $c_1 = \dots = c_\ell = 0$. Therefore, C is linearly independent. Therefore, $C \in \mathcal{S}$.

All of this shows that $\mathcal{S}$ is a non-empty set such that every chain is bounded above (by something in $\mathcal{S}$). Thus Zorn's Lemma applies and provides that $\mathcal{S}$ must contain a (possibly non-unique) maximal element. Call such an element β . We have then that β contains S , is contained in T , and is linearly independent. To finish our proof we just need to establish that β spans V .

Suppose T is not contained in $\text{span}(\beta)$. This implies there is some $v \in T$ such that $v \notin \text{span}(\beta)$. So by our Lemma, $\beta' = \beta \cup \{v\}$ is linearly independent. But also $S \subseteq \beta \subseteq \beta' \subseteq T$ so $\beta' \in \mathcal{S}$. However, this is impossible because β was chosen to be maximal (i.e., it is not contained in something else in $\mathcal{S}$). Therefore, we must conclude that $T \subseteq \text{span}(\beta)$. Therefore, $V = \text{span}(T) \subseteq \text{span}(\beta) \subseteq V$ so that $\text{span}(\beta) = V$. Thus β is our desired linearly independent spanning set. $\square$

We used Zorn's lemma in a non-trivial way. Zorn's result is equivalent to the Axiom of Choice. In fact, it can be shown that our theorem above is equivalent to the Axiom of Choice (we must use choice in some way to establish our theorem in general).

A set which both spans and is linearly independent is very special. We have a name for these:

Definition 2.5.3.

If V is a vector space over the field $\mathbb{F}$ and $\beta \subseteq V$ such that β is linearly independent and $\text{span}_{\mathbb{F}}(\beta) = V$ then β is a **basis** for V .

It is difficult to stress just how important the following corollaries to above theorem are:

Corollary 2.5.4. *Basis can be created by extension of LI set or reduction of spanning set.*

Every vector space has a basis. Every linearly independent set can be extended to a basis. Every spanning set can be shrunk down to a basis.

Proof: Applying the above theorem to the case $S = \emptyset$ and $T = V$ (these are always independent and spanning sets respectively), shows that we must have at least one basis. If S is linearly independent and we let $T = V$, we get that S is contained in some basis (i.e., every independent set can be extended to a basis). Finally, if T is a spanning set and we let $S = \emptyset$, we get a basis β that is contained in T (i.e., spanning sets can be shrunk down to a basis). $\square$

All of this leads us to the following linear algebra philosophy¹⁶: Bases are small enough spanning sets (small enough to be independent). Likewise, bases are big enough independent sets (big enough to span). We get the impression that independent sets must be no larger than spanning sets. This is true but requires some proof.

Lemma 2.5.5. (*The Exchange Lemma*)

Let $\alpha = \{v_1, \dots, v_\ell\}$ be linearly independent and suppose T spans V . Then there exists some partition of $T = \{w_1, \dots, w_\ell\} \amalg T'$ (i.e., T is a disjoint union of $\{w_1, \dots, w_\ell\}$ and T') such that $T'' = \alpha \cup T'$ still spans V . In other words, if we select the right w_i 's, we can swap out $\{w_1, \dots, w_\ell\}$ with α and still span our vector space.

Proof: We proceed by induction. We can obviously swap zero vectors without difficulty, so our base case holds. Now suppose we have replaced $\{w_1, \dots, w_k\}$ by $\{v_1, \dots, v_k\}$ so that $T'' = \{v_1, \dots, v_k\} \amalg T'$ spans V . Consider v_{k+1} . Since T'' spans V , $v_{k+1} = c_1 v_1 + \dots + c_k v_k + s_1 u_1 + \dots + s_m u_m$ for some $u_1, \dots, u_m \in T'$ and $c_1, \dots, c_k, s_1, \dots, s_m \in \mathbb{F}$. Notice that if $s_1 = \dots = s_m = 0$, then $v_{k+1} \in \text{span}\{v_1, \dots, v_k\}$ implying that α is linearly dependent (contradiction). Therefore, at least one s_i is non-zero, say $s_m \neq 0$. For convenience call $s_m = s$ and $u_m = u$. Let $X = c_1 v_1 + \dots + c_k v_k + s_1 u_1 + \dots + s_{m-1} u_{m-1}$ so we have that $v_{k+1} = su + X$. Let $T' = \{u\} \amalg T_0$ (i.e., T_0 is all of T' except u). Notice that X is a linear combination of elements drawn from $\{v_1, \dots, v_k\} \cup T_0$.

Therefore, in any linear combination, if it involves u , we can replace u with $s^{-1}v_{k+1} - s^{-1}X$. We now have that any linear combination of elements drawn from $\{v_1, \dots, v_k\} \cup \{u\} \cup T_0$ can also be written as a linear combination of elements drawn from $\{v_1, \dots, v_k\} \cup \{v_{k+1}\} \cup T_0$. In other words, we can swap out $w_{k+1} = u$ for v_{k+1} and our set will still span. The theorem now follows from induction. $\square$

Corollary 2.5.6.

If α is a finite linearly independent set and T is a spanning set, then $|\alpha| \leq |T|$. Moreover, if V is spanned by a finite set, it cannot have an infinite linearly independent subset.

Proof: By the Exchange Lemma we can replace elements of T with elements of α , so T must have at least as many elements as α . Now suppose that we have a finite spanning set T , say $|T| = m < \infty$. If we had an linearly independent set α with more than m elements, the Exchange Lemma would imply that we could replace $m + 1$ elements of T with the first $m + 1$ elements of α . This is absurd. $\square$

Corollary 2.5.7.

If any basis of a vector space is finite, all bases of that space are finite. Moreover, any two finite bases must have the same size.

Proof: Our previous corollary states that we cannot have a finite basis (a finite spanning set) and also an infinite basis (an infinite linearly independent subset).

Next, suppose α and β are finite bases for V . Then since α is finite linearly independent and β is a spanning set, our previous corollary says that $|\alpha| \leq |\beta|$. But also, β is finite linearly independent and α is a spanning set, so that $|\beta| \leq |\alpha|$. Therefore, $|\alpha| = |\beta|$. $\square$

¹⁶small as possible spanning set is also known as a minimal spanning set, likewise large as possible linearly independent set is also known as a maximally linearly independent set.

Theorem 2.5.8.

Any two bases of a vector space must have the same cardinality.

Proof: Let β and β' be bases for V . We already know that either both are infinite or both are finite. If both are finite, the corollary above establishes that $|\beta| = |\beta'|$. Therefore, we suppose that both are infinite and for sake of contradiction that $|\beta| < |\beta'|$.

Consider $v \in \beta$. Then since β' spans V , there exists $v_1, \dots, v_\ell \in \beta'$ and $c_1, \dots, c_\ell \in \mathbb{F}$ such that $v = c_1 v_1 + \dots + c_\ell v_\ell$. Define $F_v = \{v_1, \dots, v_\ell\}$ and so we have $v \in \text{span}(F_v)$. Let $\alpha = \bigcup_{v \in \beta} F_v$. Then α is a subset of β' . Moreover, for each $v \in \beta$ we have $v \in \text{span}(F_v) \subseteq \text{span}(\alpha)$ so $\beta \subseteq \text{span}(\alpha)$. Thus since β spans V , we have $V = \text{span}(\beta) = \text{span}(\alpha)$.

But this is problematic. Since α is a union of finite subsets indexed by β , we have that $|\alpha| \leq |\beta| \cdot \aleph_0 = |\beta|$ (since β is infinite). So $|\alpha| = |\beta| < |\beta'|$. Therefore α must be a proper subset of β' . Thus there is some $w \in \beta'$ that does not belong to α . Since β' is linearly independent, $\alpha \amalg \{w\}$ is independent. Therefore, $w \notin \text{span}(\alpha)$. But this means that $\text{span}(\alpha) \neq V$ (contradiction). Therefore, we must have that $|\beta| = |\beta'|$. $\square$

Given all bases have the same size, we can make the following definition:

Definition 2.5.9.

Let V have a basis β . Then $\dim(V) = |\beta|$ is the *dimension* of V .

Finite dimensional spaces are special. For example, we have the following result.

Theorem 2.5.10.

Let V be a vector space of finite dimension n . Let $\beta = \{v_1, \dots, v_n\}$ be a subset of V of cardinality n . Then β is linearly independent if and only if β spans V .

Proof: If β is linearly independent, then it is contained in a basis, say γ . But $\beta \subseteq \gamma$ where $|\beta| = |\gamma| = n < \infty$ implies $\beta = \gamma$, so β is a basis. Likewise, if β spans V , then it contains a basis, say α . But $\alpha \subseteq \beta$ where $|\alpha| = |\beta| = n < \infty$ implies $\alpha = \beta$, so β is a basis. $\square$

The above theorem does not work for infinite dimensional spaces. For example, $\{1, x^2, x^4, \dots\}$ is a countably infinite, linearly independent subset of the countably infinite dimensional space $\mathbb{R}[x]$, yet it fails to span. On the other hand, $\{2, 1, x, x^2, \dots\}$ is a countably infinite, spanning set for $\mathbb{R}[x]$, yet it fails to be independent. Having one property plus the “right size” isn’t enough to force sets to be bases if we work in infinite dimensional spaces.

2.6 coordinate mappings

We begin with a useful characterization of linear independence; the *equating coefficients* property.

Proposition 2.6.1.

Let V be a vector space over a field $\mathbb{F}$ and $S \subseteq V$. S is a linearly independent set of vectors iff for any finite list of vectors $v_1, \dots, v_k \in S$ we have the property that if there exist constants $a_1, \dots, a_k, b_1, \dots, b_k \in \mathbb{F}$ such that

$$a_1v_1 + a_2v_2 + \dots + a_kv_k = b_1v_1 + b_2v_2 + \dots + b_kv_k$$

then $a_1 = b_1, a_2 = b_2, \dots, a_k = b_k$. In other words, we can equating coefficients of any finite subset of a set of vectors iff the set of vectors is a LI set.

Proof: ($\Rightarrow$) suppose S is LI and $v_1, \dots, v_k \in S$. If there exist $a_1, \dots, a_k, b_1, \dots, b_k \in \mathbb{F}$ such that

$$a_1v_1 + a_2v_2 + \dots + a_kv_k = b_1v_1 + b_2v_2 + \dots + b_kv_k$$

then

$$(a_1 - b_1)v_1 + (a_2 - b_2)v_2 + \dots + (a_k - b_k)v_k = 0$$

hence $a_j - b_j = 0$ for $j = 1, \dots, k$ by LI of S . Hence $a_j = b_j$ for $j = 1, \dots, k$.

($\Leftarrow$) suppose S has the equating coefficients property for any finite subset. Let $v_1, \dots, v_k \in S$ and suppose $c_1v_1 + \dots + c_kv_k = 0$. Notice, $c_1v_1 + \dots + c_kv_k = 0 \cdot v_1 + \dots + 0 \cdot v_k$ hence by the equating coefficients property we find $c_1 = 0, \dots, c_k = 0$. $\square$

In retrospect, partial fractions in Calculus II is based on the LI of the basic rational functions. The technique of equating coefficients only made sense because the set of functions involved was in fact LI. Likewise, if you've taken Differential Equations, you can see LI of solutions being utilized throughout the undetermined coefficients technique. The ubiquity of the role of LI in common mathematical calculation is hard to overstate.

Definition 2.6.2.

Suppose $\beta = \{f_1, f_2, \dots, f_n\}$ is a basis for vector space V over field $\mathbb{F}$. If $v \in V$ has

$$v = v_1f_1 + v_2f_2 + \dots + v_nf_n$$

then $[v]_\beta = [v_1 \ v_2 \ \dots \ v_n]^T \in \mathbb{R}^n$ is called the **coordinate vector** of v with respect to β .

Technically, the each basis considered in the course is an "ordered basis". This means the set of vectors that forms the basis has an ordering to it. This is more structure than just a plain set since basic set theory does not distinguish $\{1, 2\}$ from $\{2, 1\}$. I should always say "we have an ordered basis" but I will not (and most people do not) say that in this course. Let it be understood that when we list the vectors in a basis they are listed in order and we cannot change that order without changing the basis. For example $v = (1, 2, 3)$ has coordinate vector $[v]_{B_1} = (1, 2, 3)$ with respect to $B_1 = \{e_1, e_2, e_3\}$. On the other hand, if $B_2 = \{e_2, e_1, e_3\}$ then the coordinate vector of v with respect to B_2 is $[v]_{B_2} = (2, 1, 3)$.

Example 2.6.3. I called $\{e_1, e_2, \dots, e_n\}$ the *standard basis* of $\mathbb{F}^n$. Since $v \in \mathbb{F}^n$ can be written as

$$v = \sum_i v_i e_i$$

it follows $\mathbb{F}^n = \text{span}\{e_i \mid 1 \leq i \leq n\}$. Moreover, linear independence of $\{e_i \mid 1 \leq i \leq n\}$ follows from a simple calculation:

$$0 = \sum_i c_i e_i \Rightarrow 0 = \left[\sum_i c_i e_i \right]_k = \sum_i c_i \delta_{ik} = c_k$$

hence $c_k = 0$ for all k . Thus $\{e_i \mid 1 \leq i \leq n\}$ is a basis for $\mathbb{F}^n$, we continue to call it the **standard basis** of $\mathbb{F}^n$. The vectors e_i are also called "unit-vectors".

Example 2.6.4. Since $A \in \mathbb{F}^{m \times n}$ can be written as

$$A = \sum_{i,j} A_{ij} E_{ij}$$

it follows $\mathbb{F}^{m \times n} = \text{span}\{E_{ij} \mid 1 \leq i \leq m, 1 \leq j \leq n\}$. Moreover, linear independence of $\{E_{ij} \mid 1 \leq i \leq m, 1 \leq j \leq n\}$ follows from a simple calculation:

$$0 = \sum_{i,j} c_{ij} E_{ij} \Rightarrow 0 = \left[\sum_{i,j} c_{ij} E_{ij} \right]_{kl} = \sum_{i,j} c_{ij} \delta_{ik} \delta_{jl} = c_{kl}$$

hence $c_{kl} = 0$ for all k, l . Thus $\{E_{ij} \mid 1 \leq i \leq m, 1 \leq j \leq n\}$ is a basis for $\mathbb{F}^{m \times n}$, we continue to call it the **standard basis** of $\mathbb{F}^{m \times n}$. The matrices E_{ij} are also called "unit-matrices".

The usual ordering given to the standard basis of $\mathbb{F}^{m \times n}$ is the **lexographical ordering** which goes row-by-row from E_{11} to E_{1n} to E_{21} to E_{2n} etc. until finally reaching E_{m1} and ending at E_{mn} . Let us see how this goes for the 2×2 case.

Example 2.6.5. Let $\beta = \{E_{11}, E_{12}, E_{22}, E_{21}\}$. Observe:

$$A = \begin{bmatrix} a & b \\ c & d \end{bmatrix} = aE_{11} + bE_{12} + dE_{22} + cE_{21}.$$

Therefore, $[A]_\beta = (a, b, d, c)$.

Example 2.6.6. Consider $\beta = \{(t+1)^2, t+1, 1\}$ and calculate the coordinate vector of $f(t) = t^2$ with respect to β . I often use an adding zero trick for such a problem:

$$f(t) = t^2 = (t+1-1)^2 = (t+1)^2 - 2(t+1) + 1.$$

From the expression above we can read that $[f(t)]_\beta = (1, -2, 1)$.

Example 2.6.7. Suppose A is invertible and $Av = b$ has solution $v = (1, 2, 3, 4)$. It follows that A has 4 columns. Define,

$$\beta = \{\text{col}_4(A), \text{col}_3(A), \text{col}_2(A), \text{col}_1(A)\}$$

Given that $(1, 2, 3, 4)$ is a solution of $Av = b$ we know:

$$\text{col}_1(A) + 2\text{col}_2(A) + 3\text{col}_3(A) + 4\text{col}_4(A) = b$$

Given the above, we can deduce $[b]_\beta = (4, 3, 2, 1)$.

The three examples above were simple enough that not much calculation was needed. Understanding the definition of basis was probably the hardest part. In general, finding the coordinates of a vector with respect to a given basis is a spanning problem.

Example 2.6.8. Given that $B = \{b_1, b_2, b_3, b_4\} = \{e_1 + e_2, e_2 + e_3, e_3 + e_4, e_4\}$ is a basis for $\mathbb{R}^4$ find coordinates for $v = [1, 2, 3, 4]^T \in \mathbb{R}^4$. We can find coordinates $[x_1, x_2, x_3, x_4]^T$ such that $v = \sum_i x_i b_i$ by calculating $\text{rref}[b_1|b_2|b_3|b_4|v]$ the rightmost column will be $[v]_B$.

$$\text{rref} \left[\begin{array}{cccc|c} 1 & 0 & 0 & 0 & 1 \\ 1 & 1 & 0 & 0 & 2 \\ 0 & 1 & 1 & 0 & 3 \\ 0 & 0 & 1 & 1 & 4 \end{array} \right] = \left[\begin{array}{cccc|c} 1 & 0 & 0 & 0 & 1 \\ 0 & 1 & 0 & 0 & 1 \\ 0 & 0 & 1 & 0 & 2 \\ 0 & 0 & 0 & 1 & 2 \end{array} \right] \Rightarrow [v]_B = \begin{bmatrix} 1 \\ 1 \\ 2 \\ 2 \end{bmatrix}$$

The calculation above should be familiar. We discussed it at length in the spanning section.

Example 2.6.9. We can prove that S from Example 2.3.12 is linearly independent, thus symmetric 2×2 matrices have a S as a basis

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

thus the dimension of the vector space of 2×2 symmetric matrices is 3. (notice $\bar{S}$ from that example is not a basis because it is linearly dependent). While we're thinking about this let's find the coordinates of $A = \begin{bmatrix} 1 & 3 \\ 3 & 2 \end{bmatrix}$ with respect to S . Denote $[A]_S = [x, y, z]^T$. We calculate,

$$\begin{bmatrix} 1 & 3 \\ 3 & 2 \end{bmatrix} = x \begin{bmatrix} 1 & 0 \\ 0 & 0 \end{bmatrix} + y \begin{bmatrix} 0 & 0 \\ 0 & 1 \end{bmatrix} + z \begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix} \Rightarrow [A]_S = \begin{bmatrix} 1 \\ 2 \\ 3 \end{bmatrix}.$$

Example 2.6.10. If $V = \mathbb{R}$ and we use $\mathbb{Q}$ as the set of scalars then V can be shown to be an infinite dimensional $\mathbb{Q}$ -vector space. In contrast, $\mathbb{R}$ is a one-dimensional real vector space with basis $\{1\}$.

A given point set may have different dimension depending on the context. The last example gives a rather exotic example of that idea. What follows is probably easier to understand:

Example 2.6.11. $\mathbb{C}$ is a one-dimensional vector space over $\mathbb{C}$ with basis $\{1\}$. However, $\mathbb{C}$ is a two-dimensional vector space over $\mathbb{R}$ with basis $\{1, i\}$. Note: $z = x + iy = x(1) + y(i)$ hence $\text{span}_{\mathbb{R}}\{1, i\} = \mathbb{C}$. Moreover, if $c_1, c_2 \in \mathbb{R}$ and $c_1(1) + c_2(i) = 0$ then we obtain $c_1 = c_2 = 0$. Hence, $\{1, i\}$ is LI over $\mathbb{R}$ and we have $\dim_{\mathbb{R}}(\mathbb{C}) = 2$. On the other hand, since $z = z(1)$ for all $z \in \mathbb{C}$ and $\{1\}$ is linearly independent it follows $\dim_{\mathbb{C}}(\mathbb{C}) = 1$ thus $\{1, i\}$ is linearly dependent. Indeed, it is obvious that $i = i(1)$ thus 1 and i are linearly dependent over $\mathbb{C}$.

Example 2.6.12. Consider $V = \mathbb{C}^{2 \times 2}$ then $\beta = \{E_{11}, E_{12}, E_{21}, E_{22}\}$ gives a basis for $V(\mathbb{C})$. If we instead wish to look at V as a real vector space there are two popular choices for the basis over $\mathbb{R}$

$$Z = \left[\begin{array}{c|c} X_{11} + iY_{11} & X_{12} + iY_{12} \\ \hline X_{21} + iY_{21} & X_{22} + iY_{22} \end{array} \right] = \left[\begin{array}{c|c} X_{11} & X_{12} \\ \hline X_{21} & X_{22} \end{array} \right] + i \left[\begin{array}{c|c} Y_{11} & Y_{12} \\ \hline Y_{21} & Y_{22} \end{array} \right] = X + iY$$

where $X, Y \in \mathbb{R}^{2 \times 2}$. If we use $\gamma_1 = \beta \cup i\beta$ then

$$[Z]_{\gamma_1} = ([X]_{\beta}, [Y]_{\beta}) = (X_{11}, X_{12}, X_{21}, X_{22}, Y_{11}, Y_{12}, Y_{21}, Y_{22})$$

Whereas if we use $\gamma_2 = \{E_{11}, iE_{11}, E_{12}, iE_{12}, E_{21}, iE_{21}, E_{22}, iE_{22}\}$ then

$$[Z]_{\gamma_2} = (X_{11}, Y_{11}, X_{12}, Y_{12}, X_{21}, Y_{21}, X_{22}, Y_{22})$$

Both bases split the matrix into real and imaginary parts. However, the ordering of the bases differ and as such the coordinate maps will reveal different aspects of such a complex matrix.

When multiple fields are in use it is wise to adorn the dimension notation with the field to reduce possible confusions. Usually, we work with just one field so we omit the explicit field dependence

Remark 2.6.13.

Curvilinear coordinate systems from calculus III are in a certain sense more general than the idea of a coordinate system in linear algebra. If we focus our attention on a single point in space then a curvilinear coordinate system will produce three linearly independent vectors which are tangent to the coordinate curves. However, if we go to a different point then the curvilinear coordinate system will produce three different vectors in general. For example, in spherical coordinates the radial unit vector is $e_\rho = \langle \cos \theta \sin \phi, \sin \theta \sin \phi, \cos \phi \rangle$ and you can see that different choices for the angles θ, ϕ make e_ρ point in different directions. In contrast, in this course we work with vector spaces. Our coordinate systems have the same basis vectors over the whole space. Vector spaces are examples of flat manifolds since they allow a single global coordinate system. Vector spaces also allow for curvilinear coordinates (which are not coordinates in the sense of linear algebra). However the converse is not true; spaces with nonzero curvature do not allow for global coordinates. I digress, we may have occasion to discuss these matters more cogently in our Advanced Calculus course (Math 332)(join us)

2.6.1 how to calculate a basis for a span of row or column vectors

Given some subspace of $\mathbb{F}^n$ we would like to know how to find a basis for that space. In particular, if $V = \text{span}\{v_1, v_2, \dots, v_k\}$ then what is a basis for W ? Likewise, given some set of row vectors $W = \{w_1, w_2, \dots, w_k\} \subset \mathbb{F}^{1 \times n}$ how can we select a basis for $\text{span}(W)$. We would like to find answers to these question since most subspaces are characterized either as spans or solution sets(see the next section on $\text{Null}(A)$). We already have the tools to answer these questions, we just need to apply them to the tasks at hand.

Proposition 2.6.14.

Let $W = \text{span}\{v_1, v_2, \dots, v_k\} \subset \mathbb{F}^n$ then a basis for W can be obtained by selecting the vectors that reside in the pivot columns of $[v_1 | v_2 | \dots | v_k]$.

Proof: this is immediately obvious from Proposition 1.7.6. $\square$

The proposition that follows is also follows immediately from Proposition 1.7.6.

Proposition 2.6.15.

Let $A \in \mathbb{F}^{m \times n}$ the pivot columns of A form a basis for $\text{Col}(A)$.

Example 2.6.16. Suppose A is given as below: (I omit the details of the Gaussian elimination)

$$A = \begin{bmatrix} 1 & 2 & 3 & 4 \\ 2 & 1 & 4 & 1 \\ 0 & 0 & 0 & 3 \end{bmatrix} \quad \Rightarrow \quad \text{rref}[A] = \begin{bmatrix} 1 & 0 & 5/3 & 0 \\ 0 & 1 & 2/3 & 0 \\ 0 & 0 & 0 & 1 \end{bmatrix}.$$

Identify that columns 1,2 and 4 are pivot columns. Moreover,

$$\text{Col}(A) = \text{span}\{\text{col}_1(A), \text{col}_2(A), \text{col}_4(A)\}$$

In particular we can also read how the second column is a linear combination of the basis vectors.

$$\begin{aligned}
 \text{col}_3(A) &= \frac{5}{3}\text{col}_1(A) + \frac{2}{3}\text{col}_2(A) \\
 &= \frac{5}{3}(1, 2, 0) + \frac{2}{3}(2, 1, 0) \\
 &= (5/3, 10/3, 0) + (4/3, 2/3, 0) \\
 &= (3, 4, 0).
 \end{aligned}$$

What if we want a basis for $\text{Row}(A)$ which consists of rows in A itself?

Proposition 2.6.17.

The rows of a matrix A can be written as linear combinations of the transposes of pivot columns of A^T . Furthermore, the set of all rows of A which are transposes of pivot columns of A^T is linearly independent.

Proof: Let A be a matrix and A^T its transpose. Apply Proposition 1.7.6 to A^T to find pivot columns which we denote by $\text{col}_{i_j}(A^T)$ for $j = 1, 2, \dots, k$. The set of pivot columns for A^T are linearly independent and their span covers each column of A^T . Suppose,

$$c_1 \text{row}_{i_1}(A) + c_2 \text{row}_{i_2}(A) + \dots + c_k \text{row}_{i_k}(A) = 0.$$

Take the transpose of the equation above, use Proposition 1.3.13 to simplify:

$$c_1 (\text{row}_{i_1}(A))^T + c_2 (\text{row}_{i_2}(A))^T + \dots + c_k (\text{row}_{i_k}(A))^T = 0.$$

Recall $(\text{row}_j(A))^T = \text{col}_j(A^T)$ thus,

$$c_1 \text{col}_{i_1}(A^T) + c_2 \text{col}_{i_2}(A^T) + \dots + c_k \text{col}_{i_k}(A^T) = 0.$$

hence $c_1 = c_2 = \dots = c_k = 0$ as the pivot columns of A^T are linearly independent. This shows the corresponding rows of A are likewise linearly independent. The proof that each row of A is obtained from a span of $\{\text{row}_{i_1}(A), \text{row}_{i_2}(A), \dots, \text{row}_{i_k}(A)\}$ is similar. $\square$

Corollary 2.6.18.

Let $W = \text{span}\{w_1, w_2, \dots, w_k\} \subset \mathbb{F}^{1 \times n}$ and construct A by concatenating the row vectors in W into a matrix A :

$$A = \begin{bmatrix} w_1 \\ w_2 \\ \vdots \\ w_k \end{bmatrix}$$

A basis for W is given by the transposes of the pivot columns for A^T .

Proof: apply Proposition 2.6.17. $\square$

The proposition that follows is also follows immediately from Proposition 2.6.17. Incidentally, it is almost more important to notice what we do **not** say below; we do not say the pivot rows of A form a basis for the row space.

Proposition 2.6.19.

Let $A \in \mathbb{F}^{m \times n}$, the rows which are transposes of the pivot columns of A^T form a basis for $\text{Row}(A)$.

Example 2.6.20.

$$A^T = \begin{bmatrix} 1 & 2 & 0 \\ 2 & 1 & 0 \\ 3 & 4 & 0 \\ 4 & 1 & 3 \end{bmatrix} \Rightarrow \text{rref}[A^T] = \begin{bmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \\ 0 & 0 & 0 \end{bmatrix}.$$

Notice that each column is a pivot column in A^T thus a basis for $\text{Row}(A)$ is simply the set of all rows of A ; $\text{Row}(A) = \text{span}\{[1, 2, 3, 4], [2, 1, 4, 1], [0, 0, 0, 3]\}$ and the spanning set is linearly independent.

Example 2.6.21.

$$A = \begin{bmatrix} 1 & 1 & 1 \\ 2 & 2 & 2 \\ 3 & 4 & 0 \\ 5 & 6 & 2 \end{bmatrix} \Rightarrow A^T = \begin{bmatrix} 1 & 2 & 3 & 5 \\ 1 & 2 & 4 & 6 \\ 1 & 2 & 0 & 2 \end{bmatrix} \Rightarrow \text{rref}[A^T] = \begin{bmatrix} 1 & 2 & 0 & 2 \\ 0 & 0 & 1 & 1 \\ 0 & 0 & 0 & 0 \end{bmatrix}.$$

We deduce that rows 1 and 3 of A form a basis for $\text{Row}(A)$. Notice that $\text{row}_2(A) = 2\text{row}_1(A)$ and $\text{row}_4(A) = \text{row}_3(A) + 2\text{row}_1(A)$. We can read linear dependencies of the rows from the corresponding linear dependencies of the columns in the rref of the transpose.

The preceding examples are nice, but what should we do if we want to find both a basis for $\text{Col}(A)$ and $\text{Row}(A)$ for some given matrix? Let's pause to think about how elementary row operations modify the row and column space of a matrix. In particular, let A be a matrix and let A' be the result of performing an elementary row operation on A . It is fairly obvious that

$$\text{Row}(A) = \text{Row}(A').$$

Think about it. If we swap two rows that just switches the order of the vectors in the span that makes $\text{Row}(A)$. On the other hand if we replace one row with a nontrivial linear combination of itself and other rows then that will not change the span either. Column space is not so easy though. Notice that elementary row operations can change the column space. For example,

$$A = \begin{bmatrix} 1 & 1 \\ 1 & 1 \end{bmatrix} \Rightarrow \text{rref}[A] = \begin{bmatrix} 1 & 1 \\ 0 & 0 \end{bmatrix}$$

has $\text{Col}(A) = \text{span}\{[1, 1]^T\}$ whereas $\text{Col}(\text{rref}(A)) = \text{span}([1, 0]^T)$. We cannot hope to use columns of $\text{ref}(A)$ (or $\text{rref}(A)$) for a basis of $\text{Col}(A)$. That's no big problem though because we already have the CCP-principle which helped us pick out a basis for $\text{Col}(A)$. Let's collect our thoughts:

Proposition 2.6.22.

Let $A \in \mathbb{F}^{m \times n}$ then a basis for $\text{Col}(A)$ is given by the pivot columns in A and a basis for $\text{Row}(A)$ is given by the nonzero rows in $\text{ref}(A)$.

This means we can find a basis for $\text{Col}(A)$ and $\text{Row}(A)$ by performing the forward pass on A . We need only calculate the $\text{ref}(A)$ as the pivot columns are manifest at the end of the forward pass.

Example 2.6.23.

$$A = \begin{bmatrix} 1 & 1 & 1 \\ 1 & 1 & 1 \\ 1 & 2 & 3 \end{bmatrix} \xrightarrow{\substack{r_2 - r_1 \rightarrow r_2 \\ r_3 - r_1 \rightarrow r_3}} \begin{bmatrix} 1 & 1 & 1 \\ 0 & 0 & 0 \\ 0 & 1 & 2 \end{bmatrix} \xrightarrow{r_1 \leftrightarrow r_2} \begin{bmatrix} 1 & 1 & 1 \\ 0 & 1 & 2 \\ 0 & 0 & 0 \end{bmatrix} = \text{ref}[A]$$

We deduce that $\{[1, 1, 1], [0, 1, 2]\}$ is a basis for $\text{Row}(A)$ whereas $\{[1, 1, 1]^T, [1, 1, 2]^T\}$ is a basis for $\text{Col}(A)$. Notice that if I wanted to reveal further linear dependencies of the non-pivot columns on the pivot columns of A it would be wise to calculate $\text{rref}[A]$ by making the backwards pass on $\text{ref}[A]$.

$$\begin{bmatrix} 1 & 1 & 1 \\ 0 & 1 & 2 \\ 0 & 0 & 0 \end{bmatrix} \xrightarrow{r_1 - r_2 \rightarrow r_1} \begin{bmatrix} 1 & 0 & -1 \\ 0 & 1 & 2 \\ 0 & 0 & 0 \end{bmatrix} = \text{rref}[A]$$

From which I can read $\text{col}_3(A) = 2\text{col}_2(A) - \text{col}_1(A)$, a fact which is easy to verify.

Example 2.6.24.

$$A = \begin{bmatrix} 1 & 2 & 3 & 4 \\ 1 & 3 & 8 & 10 \\ 1 & 2 & 4 & 11 \end{bmatrix} \xrightarrow{\substack{r_2 - r_1 \rightarrow r_2 \\ r_3 - r_1 \rightarrow r_3}} \begin{bmatrix} 1 & 2 & 3 & 4 \\ 0 & 1 & 5 & 6 \\ 0 & 0 & 1 & 7 \end{bmatrix} = \text{ref}[A]$$

We find that $\text{Row}(A)$ has basis

$$\{[1, 2, 3, 4], [0, 1, 5, 6], [0, 0, 1, 7]\}$$

and $\text{Col}(A)$ has basis

$$\left\{ \begin{bmatrix} 1 \\ 1 \\ 1 \end{bmatrix}, \begin{bmatrix} 2 \\ 3 \\ 2 \end{bmatrix}, \begin{bmatrix} 3 \\ 8 \\ 4 \end{bmatrix} \right\}$$

Proposition 2.6.22 was the guide for both examples above.

2.6.2 calculating the basis of the null space of a matrix

Often a subspace is described as the solution set of some equation $Ax = 0$. How do we find a basis for $\text{Null}(A)$? If we can do that we find a basis for subspaces which are described by some equation.

Proposition 2.6.25.

Let $A \in \mathbb{F}^{m \times n}$ and define $W = \text{Null}(A)$. A basis for W is obtained from the solution set of $Ax = 0$ by writing the solution as a linear combination where the free variables appear as coefficients in the vector-sum.

Proof: $x \in W$ implies $Ax = 0$. Denote $x = [x_1, x_2, \dots, x_n]^T$. Suppose that $\text{rref}[A]$ has r -pivot columns (we must have $0 \leq r \leq n$). There will be $(m - r)$ -rows which are zero in $\text{rref}(A)$ and $(n - r)$ -columns which are not pivot columns. The non-pivot columns correspond to free-variables in the solution. Define $p = n - r$ for convenience. Suppose that $x_{i_1}, x_{i_2}, \dots, x_{i_p}$ are free whereas $x_{j_1}, x_{j_2}, \dots, x_{j_r}$ are functions of the free variables: in particular they are linear combinations of the

free variables as prescribed by $rref[A]$. There exist components or $rref(A)$, let us denote them by b_{ij} , such that

$$\begin{aligned} x_{j_1} + b_{11}x_{i_1} + b_{12}x_{i_2} + \cdots + b_{1p}x_{i_p} &= 0 \\ x_{j_2} + b_{21}x_{i_1} + b_{22}x_{i_2} + \cdots + b_{2p}x_{i_p} &= 0 \\ &\vdots \\ x_{j_r} + b_{r1}x_{i_1} + b_{r2}x_{i_2} + \cdots + b_{rp}x_{i_p} &= 0 \end{aligned}$$

Hence $x_{j_k} = -\sum_{l=1}^p b_{kl}x_{i_l}$. Thus, if $Ax = 0$ then

$$x = \sum_{k=1}^r x_{j_k} e_{j_k} + \sum_{l=1}^p x_{i_l} e_{i_l} = \sum_{k=1}^r \left(-\sum_{l=1}^p b_{kl}x_{i_l} \right) e_{j_k} + \sum_{l=1}^p x_{i_l} e_{i_l} = \sum_{l=1}^p x_{i_l} \left(e_{i_l} - \sum_{k=1}^r b_{kl} e_{j_k} \right)$$

Thus $\beta = \{e_{i_l} - \sum_{k=1}^r b_{kl} e_{j_k}\}_{l=1}^p$ serves as a spanning set for $Null(A)$. It remains to prove β is LI. I will sketch a proof¹⁷ by contradiction. If β was linearly dependent then that implies an additional dependence amongst the variables forming the solution set of $Ax = 0$. But, this is impossible since such a dependence would imply a linear dependence amongst the pivot columns of A . $\square$

Didn't follow the proof above? You might do well to sort through the example below before attempting a second reading. The examples are just the proof in action for specific cases.

Example 2.6.26. Find a basis for the null space of $A = [1, 2, 3, 4]$. This example requires no additional calculation except this; $Ax = 0$ for $x = (x_1, x_2, x_3, x_4)$ yields $x_1 = -2x_2 - 3x_3 - 4x_4$ thus:

$$x = \begin{bmatrix} -2x_2 - 3x_3 - 4x_4 \\ x_2 \\ x_3 \\ x_4 \end{bmatrix} = x_2 \begin{bmatrix} -2 \\ 1 \\ 0 \\ 0 \end{bmatrix} + x_3 \begin{bmatrix} -3 \\ 0 \\ 1 \\ 0 \end{bmatrix} + x_4 \begin{bmatrix} -4 \\ 0 \\ 0 \\ 1 \end{bmatrix}.$$

Thus, $\{(-2, 1, 0, 0), (-3, 0, 1, 0), (-4, 0, 0, 1)\}$ forms a basis for $Null(A)$.

Example 2.6.27. Find a basis for the null space of A given below,

$$A = \begin{bmatrix} 1 & 0 & 0 & 1 & 0 \\ 2 & 2 & 0 & 0 & 1 \\ 4 & 4 & 4 & 0 & 0 \end{bmatrix}$$

Gaussian elimination on the augmented coefficient matrix reveals

$$rref \begin{bmatrix} 1 & 0 & 0 & 1 & 0 \\ 2 & 2 & 0 & 0 & 1 \\ 4 & 4 & 4 & 0 & 0 \end{bmatrix} = \begin{bmatrix} 1 & 0 & 0 & 1 & 0 \\ 0 & 1 & 0 & -1 & 1/2 \\ 0 & 0 & 1 & 0 & -1/2 \end{bmatrix}$$

Denote $x = (x_1, x_2, x_3, x_4, x_5)$ in the equation $Ax = 0$ and identify from the calculation above that x_4 and x_5 are free thus solutions are of the form

$$\begin{aligned} x_1 &= -x_4 \\ x_2 &= x_4 - \frac{1}{2}x_5 \\ x_3 &= \frac{1}{2}x_5 \\ x_4 &= x_4 \\ x_5 &= x_5 \end{aligned} \Rightarrow x = \begin{bmatrix} -x_4 \\ x_4 - \frac{1}{2}x_5 \\ \frac{1}{2}x_5 \\ x_4 \\ x_5 \end{bmatrix} = x_4 \begin{bmatrix} -1 \\ 1 \\ 0 \\ 1 \\ 0 \end{bmatrix} + x_5 \begin{bmatrix} 0 \\ -\frac{1}{2} \\ \frac{1}{2} \\ 0 \\ 1 \end{bmatrix}$$

¹⁷we can prove LI of β if we find an independent argument that $\dim(Null(A)) = p$, if that dimension was known then we know a spanning set with p -vectors is necessarily a basis for $Null(A)$.

for all $x_4, x_5 \in \mathbb{R}$. We can write these results in vector form to reveal the basis for $\text{Null}(A)$, It follows that the basis for $\text{Null}(A)$ is simply $\{(-1, 1, 0, 1, 0), (0, -1/2, 1/2, 0, 1)\}$ Of course, you could multiply the second vector by 2 if you wish to avoid fractions. In fact there is a great deal of freedom in choosing a basis. We simply show one way to do it.

Example 2.6.28. To find a basis for the null space of $A = \begin{bmatrix} 1 & 1 & 1 & 1 \\ 1 & 1 & 1 & 1 \\ 1 & 1 & 1 & 1 \end{bmatrix}$ we perform Gaussian elimination to reveal

$$\text{rref} \begin{bmatrix} 1 & 1 & 1 & 1 \\ 1 & 1 & 1 & 1 \\ 1 & 1 & 1 & 1 \end{bmatrix} = \begin{bmatrix} 1 & 1 & 1 & 1 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \end{bmatrix}$$

Denote $x = (x_1, x_2, x_3, x_4)$ in the equation $Ax = 0$ and identify from the calculation above that x_2, x_3 and x_4 are free thus solutions are of the form

$$\begin{array}{l} x_1 = -x_2 - x_3 - x_4 \\ x_2 = x_2 \\ x_3 = x_3 \\ x_4 = x_4 \end{array} \Rightarrow x = \begin{bmatrix} -x_2 - x_3 - x_4 \\ x_2 \\ x_3 \\ x_4 \end{bmatrix} = x_2 \begin{bmatrix} -1 \\ 1 \\ 0 \\ 0 \end{bmatrix} + x_3 \begin{bmatrix} -1 \\ 0 \\ 1 \\ 0 \end{bmatrix} + x_4 \begin{bmatrix} -1 \\ 0 \\ 0 \\ 1 \end{bmatrix}$$

for all $x_2, x_3, x_4 \in \mathbb{R}$. We can write these results in vector form to reveal the basis for $\text{Null}(A)$, It follows that the basis for $\text{Null}(A)$ is simply $\{(-1, 1, 0, 0), (-1, 0, 1, 0), (-1, 0, 0, 1)\}$.

The process of finding the basis for the null space is a bit different than the column or row space basis problem. Both the column and row calculations allow us to read the basis of either $\text{rref}(A)$ or $\text{rref}(A^T)$ on the logical basis of the CCP. There is a matrix-theoretic way to do the same for the null space, but, it comes at a considerable cost to esthetics. See my Question and Answer on the Math Stackexchange¹⁸.

2.6.3 coordinates for an infinite dimensional vector space

When V is infinite dimensional, we can still define a notion of coordinates, but they are not nearly as useful. Before I share the construction¹⁹ for infinite dimensions, let me discuss the approach as it relates to finite dimensions.

Notice, in the finite context we can understand each coordinate vector as a function on the finite set $\{1, 2, \dots, n\}$. In particular, if $\beta = \{v_1, \dots, v_n\}$ and $v = x_1v_1 + \dots + x_nv_n$ then $[v]_\beta = (x_1, \dots, x_n)$ could be replaced with

$$f_\beta(v) : \{1, 2, \dots, n\} \rightarrow \mathbb{F}$$

defined by $f_\beta(v)(i) = x_i$ for $i = 1, \dots, n$. Of course, this would be rather inefficient since it is much simpler to just list the values of $f_\beta(v)$. Indeed, the notational question we face here was also dealt with in Calculus II where we trade functions $a : \mathbb{N} \rightarrow \mathbb{R}$ for an ordered list $\{a_n\}$. In some sense, using a sequence or coordinate vector corresponds to using the range of the coordinate vector **function** we introduce next.

¹⁸there is a link to math.stackexchange.com/q/1612616/36530 if you look at the pdf in the proper viewer

¹⁹thanks to William Cook of Appalachian State University, Boone North Carolina for the construction, it certainly is an interesting addition as of Fall 2024 Semester. I took his nice notation and uglified it for clarity.

Consider a basis β indexed by some set I (i.e., $\beta = \{v_i \mid i \in I\}$ where $v_i = v_j$ if and only if $i = j$). The set of all functions $\mathbb{F}^I = \{f \mid f : I \rightarrow \mathbb{F}\}$ can be given the structure of a vector space over $\mathbb{F}$ as follows: $f + g$ is defined to be $(f + g)(i) = f(i) + g(i)$ (we add outputs pointwise) and sf is defined by $(sf)(i) = sf(i)$ for all $i \in I$; $f, g \in \mathbb{F}^I$; and $s \in \mathbb{F}$. Then consider $\widehat{\mathbb{F}^I} = \{f : I \rightarrow \mathbb{F} \mid f(i) = 0 \text{ for all but finitely many } i \in I\}$ (i.e., the set of functions of finite support). It is easy to see that $\widehat{\mathbb{F}^I}$ is a subspace of $\mathbb{F}^I$ which corresponds with V via the correspondence²⁰ $v \mapsto f_\beta(v)$.

Definition 2.6.29.

Suppose β be a basis for vector space V over field $\mathbb{F}$ and $\beta = \{v_i \mid i \in I\}$; that is, suppose β is a basis indexed by I . Now notice that given any $v \in V$, there exists (unique up to padding out with zeros) $v_{i_1}, \dots, v_{i_\ell} \in \beta$ and $c_{i_1}, \dots, c_{i_\ell} \in \mathbb{F}$ such that $v = c_{i_1}v_{i_1} + \dots + c_{i_\ell}v_{i_\ell}$. We can then define a function $f_\beta(v) : I \rightarrow \mathbb{F}$ by $f_\beta(v)(i_j) = c_{i_j}$ for $j = 1, \dots, \ell$ and $f_\beta(v)(i) = 0$ for all other $i \in I$. Then $f_\beta(v)$ is the unique coordinate function associated with the vector in $v \in V$.

In the particular case $I = \mathbb{N}$ we can use sequential notation; $[v]_\beta = f_\beta(v)(\mathbb{N})$.

Example 2.6.30. Let $\beta = \{1, x, x^2, \dots\}$ serve as the basis for $\mathbb{F}[x]$ then if $v = a_0 + a_1x + \dots + a_nx^n$ then $[v]_\beta = (a_0, a_1, \dots, a_n, 0, \dots)$.

The example above is probably the best I can do explicitly. The funny thing is that there is a rather large distinction between knowing the existence of a basis and constructing such a basis. For instance, I know $\mathbb{R}(\mathbb{Q})$ is an infinite dimensional vector space with dimension $\aleph_1$, but I cannot write down a basis with which we can express any real number as finite $\mathbb{Q}$ -linear combination of the basis. As a practical matter, the concept of spanning tends to be modified to work in tandem with some concept of vector length (norm). Convergence replaces finiteness and with that adaptation more can be constructed. Look up Schauder basis if you wish to study this further, or study Hilbert Spaces.

2.6.4 trace based calculation of dimension

Technically, we've already established the proposition below using fancy arguments that extend well past the context of finite dimensions. Indeed, the current set of notes offer a more sophisticated take on dimension theory than I have offered in previous renditions of Math 321. That said, the following calculation is dear to me and I cannot let it go. Basically what is shown below is that if we have two finite bases then they must have the same number of vectors. If you study arguments given in other linear algebra texts you'll find arguments somewhat like I gave earlier in this chapter, except rather than using Zorn's lemma they'll most like give a proof which is anchored to a lengthy induction argument.²¹ In contrast, the proof below is a calculation.

Proposition 2.6.31.

Let V be vector space over a field $\mathbb{F}$ and suppose there exists $B = \{b_1, b_2, \dots, b_n\}$ a basis of V . Then any other basis for V also has n -elements. In other words, any two finite linearly independent generating sets for a vector space V have the same number of elements.

²⁰ $v \mapsto f_\beta(v)$ is an isomorphism from V to $\widehat{\mathbb{F}^I}$ just as $v \mapsto [v]_\beta$ is an isomorphism from V to $\mathbb{F}^n$ in the finite dimensional context. That said, when I is infinite, working with such coordinate functions isn't typically terribly useful. My apologies for this out of order footnote, we discuss isomorphism properly in the next chapter.

²¹For example, read Chapter 5 in Curtis' *Linear Algebra*. If you look at the previous version of these notes from 2019 you'll find a complete list of the results presented in Curtis.

Proof: Suppose $B = \{b_1, b_2, \dots, b_n\}$ and $F = \{f_1, f_2, \dots, f_p\}$ are both bases for a vector space V . Since F is a basis it follows $b_k \in \text{span}(F)$ for all k so there exist constants $c_{ik} \in \mathbb{F}$ such that

$$b_k = c_{1k}f_1 + c_{2k}f_2 + \dots + c_{pk}f_p$$

for $k = 1, 2, \dots, n$. Likewise, since $f_j \in \text{span}(B)$ there exist constants $d_{ij} \in \mathbb{F}$ such that

$$f_j = d_{1j}b_1 + d_{2j}b_2 + \dots + d_{nj}b_n$$

for $j = 1, 2, \dots, p$. Substituting we find

$$\begin{aligned} f_j &= d_{1j}b_1 + d_{2j}b_2 + \dots + d_{nj}b_n \\ &= d_{1j}(c_{11}f_1 + c_{21}f_2 + \dots + c_{p1}f_p) + \\ &\quad + d_{2j}(c_{12}f_1 + c_{22}f_2 + \dots + c_{p2}f_p) + \\ &\quad + \dots + d_{nj}(c_{1n}f_1 + c_{2n}f_2 + \dots + c_{pn}f_p) \\ &= (d_{1j}c_{11} + d_{2j}c_{12} + \dots + d_{nj}c_{1n})f_1 \\ &\quad + (d_{1j}c_{21} + d_{2j}c_{22} + \dots + d_{nj}c_{2n})f_2 + \\ &\quad + \dots + (d_{1j}c_{p1} + d_{2j}c_{p2} + \dots + d_{nj}c_{pn})f_p \end{aligned}$$

Suppose $j = 1$. We deduce, by the linear independence of F , that

$$d_{11}c_{11} + d_{21}c_{12} + \dots + d_{n1}c_{1n} = 1$$

from comparing coefficients of f_1 , whereas for f_2 we find,

$$d_{11}c_{21} + d_{21}c_{22} + \dots + d_{n1}c_{2n} = 0$$

likewise, for f_q with $q \neq 1$,

$$d_{11}c_{q1} + d_{21}c_{q2} + \dots + d_{n1}c_{qn} = 0$$

Notice we can rewrite all of these as

$$\delta_{q1} = c_{q1}d_{11} + c_{q2}d_{21} + \dots + c_{qn}d_{n1}$$

Similarly, for arbitrary j we'll find

$$\delta_{qj} = c_{q1}d_{1j} + c_{q2}d_{2j} + \dots + c_{qn}d_{nj}$$

If we define $C = [c_{ij}] \in \mathbb{F}^{p \times n}$ and $D = [d_{ij}] \in \mathbb{F}^{n \times p}$ then we can translate the equation above into the matrix equation that follows:

$$CD = I_p.$$

We can just as well argue that

$$DC = I_n$$

The **trace** of a matrix is the sum of the diagonal entries in the matrix; $\text{trace}(A) = \sum_{i=1}^n A_{ii}$ for $A \in \mathbb{F}^{n \times n}$. It is not difficult to show that $\text{trace}(AB) = \text{trace}(BA)$ provided the products AB and BA are both defined. Moreover, it is also easily seen $\text{tr}(I_p) = p$ and $\text{tr}(I_n) = n$. It follows that,

$$\text{tr}(CD) = \text{tr}(DC) \Rightarrow \text{tr}(I_p) = \text{tr}(I_n) \Rightarrow p = n.$$

Since the bases were arbitrary this proves any pair of bases have the same number of vectors. $\square$

The trace has use far beyond this proof. For example, to obtain an invariant over a symmetry group in physics one takes the trace of an expression to form the Lagrangian of a gauge theory. The trace is an example of a **linear functional** and it plays an important role in representation theory.

2.7 theory of subspaces

We now turn to the question of how we can produce new subspaces from a given pair.

Theorem 2.7.1.

Let V be a vector space and suppose $U \leq V$ and $W \leq V$ then $U \cap W \leq V$.

Proof: it is tempting to prove this here. But, I leave it for homework. $\square$

Examples of the Theorem in $\mathbb{R}^3$ are fun to think about. For example, one-dimensional subspaces are lines through the origin and two-dimensional subspaces are planes through the origin. The intersection of a line and plane is either the line once more or the origin. On the other hand, the intersection of two planes is either the plane once more (if the given planes are identical), or a line. Two planes cannot intersect in just one point in $\mathbb{R}^3$. In contrast, in $\mathbb{R}^4$ the planes $x_1 = x_2 = 0$ and $x_3 = x_4 = 0$ share only $(0, 0, 0, 0)$ in common. Apparently, the calculation of the intersection of two subspaces has many cases to enumerate for an arbitrary vector space. That said, we do have a nice theorem which relates the intersection to the sum of two subspaces. The following definition is quite natural:

Definition 2.7.2.

Let V be a vector space and $U \leq V$ and $W \leq V$ then define the **sum** of U and W by

$$U + W = \{x + y \mid x \in U, y \in W\}.$$

If you consider the union of two subspaces you'll find the result is only a subspace when one of the subspaces contains the other. For example, the union of the x and y -axes in $\mathbb{R}^2$ is missing the sums such as $(x, 0) + (0, y) = (x, y)$ for $x, y \neq 0$. It turns out the set defined above is the smallest subspace which contains the union of the given subspaces.

Theorem 2.7.3.

Let V be a vector space and suppose $U \leq V$ and $W \leq V$ then $U + W \leq V$. Moreover, no smaller subspace contains $U \cup W$.

Proof: it is tempting to prove this here. But, I leave it for homework. $\square$

Theorem 2.7.4.

Let V be a finite dimensional vector space with subspace W . Then $\dim(W) \leq \dim(V)$ where equality is attained only if $V = W$.

Proof: Let β be a basis for W , if β is also a basis for V then $\dim(V) = \dim(W)$ and $V = W = \text{span}(\beta)$. Otherwise, if $\text{span}(\beta) \neq V$, apply Corollary 2.5.4 and extend β to γ a basis for V . Hence, $\dim(W) = \#(\beta) < \#(\gamma) = \dim(V)$. $\square$

This result which more useful than you might first expect. In particular, suppose $U_1, U_2, \dots$ are subspaces of a finite dimensional vector space V . If $U_j \leq U_{j-1}$ for $j = 2, \dots, k$ with $U_j \neq U_{j-1}$ then we have the following nested-sequence of subsets:

$$U_k \subset U_{k-1} \subset \dots \subset U_2 \subset U_1 \subset V$$

where $\dim(U_j) < \dim(U_{j-1})$ for each $j = 1, 2, \dots$. Simple counting then reveals we cannot keep descending to smaller subspaces without end. Eventually, we obtain a smallest subspace (it might be $\{0\}$). Conversely, we could think about a sequence of subspaces which gets larger as the sequence progresses. Once again, we cannot continue without end as the dimension of V bounds the dimension of subspaces.

There is a natural relation between the dimensions of the sum and intersection of two subspaces which is related to the counting problem for two sets: if A and B are finite sets then

$$\#(A \cup B) = \#A + \#B - \#(A \cap B).$$

This rule is easy to see in a Venn Diagram. I will not omit the proof of the result below, it does require some effort.

Theorem 2.7.5.

Let $V(\mathbb{F})$ be a finite-dimensional vector space and suppose $U \leq V$ and $W \leq V$ then

$$\dim(U + W) = \dim(U) + \dim(W) - \dim(U \cap W).$$

Proof: Given $U \leq V$ and $W \leq V$ we have $U \cap W \leq V$ and $U + W \leq V$ by Theorems 2.7.3 and 2.7.3. I invite the reader to verify $U \cap W \leq U \leq U + W$ and $U \cap W \leq W \leq U + W$ (both of these assertions are simple to obtain from the Subspace Test Theorem). Observe by Corollary 2.5.4 we can find a basis $\beta_{U \cap W} = \{v_1, \dots, v_n\}$ for $U \cap W$. We count $\dim(U \cap W) = n$. Notice $\beta_{U \cap W}$ is a LI subset of U thus by Corollary 2.5.4 we can complete $\beta_{U \cap W}$ to a basis $\beta_U = \{v_1, \dots, v_n, u_1, \dots, u_d\}$ for U by adjoining vectors $u_1, \dots, u_d \in U - U \cap W$. Notice $\dim(U) = n + d$ in our current notation. Similarly, $\beta_{U \cap W}$ is a LI subset of W thus by Corollary 2.5.4 we can complete $\beta_{U \cap W}$ to a basis $\beta_W = \{v_1, \dots, v_n, w_1, \dots, w_e\}$ for W by adjoining vectors $w_1, \dots, w_e \in W - U \cap W$. We count, $\dim(W) = n + e$. We argue that $\beta_{U+W} = \{v_1, \dots, v_n, u_1, \dots, u_d, w_1, \dots, w_e\}$ forms a basis for $U + W$ hence the Theorem follows as

$$\dim(U + W) = n + d + e = (n + d) + (n + e) - n = \dim(U) + \dim(W) - \dim(U \cap W).$$

It remains to show β_{U+W} is a basis for $U + W$. Let $z \in U + W$ then there exist $x \in U$ and $y \in W$ such that $z = x + y$. However, as $\{v_1, \dots, v_n, u_1, \dots, u_d\}$ is basis for U there exist $c_i, b_j \in \mathbb{F}$ such that $x = \sum_{i=1}^n c_i v_i + \sum_{j=1}^d b_j u_j$. Also, as $\{v_1, \dots, v_n, w_1, \dots, w_e\}$ is a basis for W and $y \in W$ there exist $\alpha_i, \beta_k \in \mathbb{F}$ for which $y = \sum_{i=1}^n \alpha_i v_i + \sum_{k=1}^e \beta_k w_k$. Thus,

$$z = x + y = \sum_{i=1}^n (\alpha_i + c_i) v_i + \sum_{j=1}^d b_j u_j + \sum_{k=1}^e \beta_k w_k$$

and we find β_{U+W} is a generating set for $U + W$. Finally, we must demonstrate β_{U+W} is LI. Suppose there exist $\alpha_i, \beta_j, \gamma_k \in \mathbb{F}$ for which

$$\sum_{i=1}^n \alpha_i v_i + \underbrace{\sum_{j=1}^d \beta_j u_j}_{x \in U} + \underbrace{\sum_{k=1}^e \gamma_k w_k}_{y \in W} = 0 \quad \star$$

Recall, by construction, $u_j \in U - U \cap W$ and $w_k \in W - U \cap W$. Solve for the vector in W ,

$$\sum_{i=1}^n \alpha_i v_i + \sum_{j=1}^d \beta_j u_j = - \sum_{k=1}^e \gamma_k w_k = -y \in W$$

But, $\sum_{i=1}^n \alpha_i v_i + \sum_{j=1}^d \beta_j u_j \in U$ thus $-y \in U$ and $-y \in W$ hence $-y \in U \cap W$! Thus, there exist $\eta_1, \dots, \eta_n$ for which $\sum_{i=1}^n \alpha_i v_i + \sum_{j=1}^d \beta_j u_j = \sum_{i=1}^n \eta_i v_i$ ($\star\star$). Hence, by LI of β_U we learn $\beta_j = 0$ for $j \in \mathbb{N}_d$ by comparing coefficients²² of the LHS and RHS of $\star\star$. To complete the proof we make an entirely similar argument for $-x$ which shows $\gamma_k = 0$ for $k \in \mathbb{N}_e$. Finally, returning to $\star$ we have $\sum_{i=1}^n \alpha_i v_i = 0$ and LI of $\beta_{U \cap W} = \{v_1, \dots, v_n\}$ shows $\alpha_i = 0$ for each $i \in \mathbb{N}_n$. This completes the proof. $\square$

You might ask, what about three subspaces²³ ? Threats aside, the problem of studying the decomposition of a vector space into a finite set of subspaces is an interesting and central problem of linear algebra we will devote substantial energy towards in a later part of this course. This is just Chapter 1 of that story.

²²like the wings of an elephant, the coefficients of u_j are set to zero on the RHS of $\star\star$

²³add evil laugh to properly read this question. Or look at this question on math overflow

Chapter 3

linear transformations

For the invisible things of him from the creation of the world are clearly seen, being understood by the things that are made, even his eternal power and Godhead; so that they are without excuse: Because that, when they knew God, they glorified him not as God, neither were thankful; but became vain in their imaginations, and their foolish heart was darkened.

ROMANS 1:20-21, KJV

It would be wise to review Appendix Chapter 8 to refresh your memory set theory and functions. I do expect you be conversant in images and inverse images of sets under a function. Some of you have thought more about this than others. Of course, we review this as we go, but, you would likely profit from some preparatory reading.

The theorems on dimension also find further illumination in this chapter. We study **isomorphisms**. Roughly speaking, two vector spaces which are isomorphic are just the same set with different notation in so far as the vector space structure is concerned. Don't view this sentence as a license to trade column vectors for matrices or functions. We're not there yet. You can do that after this course, once you understand the abuse of language properly. Sort of like how certain musicians can say forbidden words since they have earned the rights through their life experience.

We also study the problem of coordinate change. Since the choice of basis is not unique the problem of comparing different pictures of vectors or transformations for abstract vector spaces requires some effort. We begin by translating our earlier work on coordinate vectors into a mapping-centered notation. Once you understand the notation properly, we can draw pictures to solve problems. This idea of **diagrammatic argument** is an important and valuable technique of modern mathematics. Modern mathematics is less concerned with equations and more concerned with functions and sets.

A theme for applications of linear algebra outside this course is the technique of **linearization**. Given a complicated set of equations we can approximate them by a simpler set of linear equations. This is the idea of Newton's Method for root-finding. Ultimately, this approximation paired with the nontrivial contraction-mapping technique provides proof of the implicit and inverse function theorems. In short these theorems say linearization works as well as you would naively hope. On the other hand, given a mapping which twists and contorts one shape into another globally may allow a rather simple description locally. The basic idea is to replace a globally nonlinear function with a local linearization. The linearization is built from a linear transformation. Much can be gleaned from the local linearization for a wide swath of problems. I really can't overstate the use of linear transformations. If you understand them then you understand more things than you know.

The conclusion of this chapter focuses on general abstract constructions which are often seen in the application of linear algebra such as dual spaces, quotient spaces and direct sum decomposition. The first isomorphism theorem is proved and applied. We also give a classification theorem for maps on finite dimensional vector spaces and discuss its connection with matrix congruence.

3.1 definition, examples and basic theory

We assume V, W are vector spaces over a field $\mathbb{F}$ in the remainder of this section.

Definition 3.1.1.

Let V, W be vector spaces over a field $\mathbb{F}$. If a function $T : V \rightarrow W$ satisfies

(1.) $T(x + y) = T(x) + T(y)$ for all $x, y \in V$; T is **additive** or T **preserves addition**

(2.) $T(cx) = cT(x)$ for all $x \in V$ and $c \in \mathbb{F}$; T is **preserves scalar multiplication**

then we say T is a **linear transformation** from V to W . The set of all linear transformations from V to W is denoted $\mathcal{L}(V, W)$. Also, $\mathcal{L}(V, V) = \mathcal{L}(V)$ and $T \in \mathcal{L}(V)$ is called a **linear transformation on V** .

I have used the terminology that (2.) is **homogeneity** of T , but, technically, T is homogeneous degree one. More generally, $T(cx) = c^k T(x)$ for all $c \in \mathbb{F}$ and $x \in V$ would make T homogeneous of degree k . In the interest of readability, let us agree that homogeneous means homogeneous of degree one. We have little use of higher homogeneity in these notes. I should also mention, if

$$T(cx + y) = cT(x) + T(y)$$

for all $x, y \in V$ and $c \in \mathbb{F}$ then it follows from $c = 1$ that T preserves addition and $y = 0$ that T preserves scalar multiplication. Thus, much like the subspace test arguments, we can combine our analysis into the simple check; does $T(cx + y) = cT(x) + T(y)$. I should also mention, other popular notations,

$$\mathcal{L}(V) = \text{End}(V) \quad \& \quad \mathcal{L}(V, W) = \text{Hom}_{\mathbb{F}}(V, W)$$

where $\text{End}(V)$ is read the **endomorphisms** of V .

3.1.1 examples of linear transformations

I'll offer several examples of functions which are **not** linear transformations then we'll conclude this subsection with a number of positive examples.

Example 3.1.2. Let $L : \mathbb{R} \rightarrow \mathbb{R}$ be defined by $L(x) = mx + b$ where $m, b \in \mathbb{R}$ and $b \neq 0$. This is often called a **linear function** in basic courses. However, this is unfortunate terminology as:

$$L(x + y) = m(x + y) + b = mx + b + my + b - b = L(x) + L(y) - b.$$

Thus L is not additive hence it is not a linear transformation. It is certainly true that $y = L(x)$ gives a line with slope m and y -intercept b . An accurate term for L is that it is an **affine function**.

Example 3.1.3. Let $f(x, y) = x^2 + y^2$ define a function from $\mathbb{R}^2$ to $\mathbb{R}$. Observe,

$$f(c(x, y)) = f(cx, cy) = (cx)^2 + (cy)^2 = c^2(x^2 + y^2) = c^2 f(x, y).$$

Clearly f does not preserve scalar multiplication hence f is not linear. (however, f is homogeneous degree 2)

Example 3.1.4. Suppose $f(t, s) = (\sqrt{t}, s^2 + t)$ note that $f(1, 1) = (1, 2)$ and $f(4, 4) = (2, 20)$. Note that $(4, 4) = 4(1, 1)$ thus we should see $f(4, 4) = f(4(1, 1)) = 4f(1, 1)$ but that fails to be true so f is not a linear transformation.

Now that we have a few examples of how not to be a linear transformation, let's take a look at some positive examples. Notice that we have many new examples to explore here now that we consider abstract vector spaces. In contrast, Math 221 focused on the case of column vectors and their maps (we review those in a future subsection).

Example 3.1.5. Define $T : \mathbb{F}^{m \times n} \rightarrow \mathbb{F}^{n \times m}$ by $T(A) = A^T$. This is clearly a linear transformation since

$$T(cA + B) = (cA + B)^T = cA^T + B^T = cT(A) + T(B)$$

for all $A, B \in \mathbb{F}^{m \times n}$ and $c \in \mathbb{F}$.

Example 3.1.6. Let V, W be a vector spaces over a field $\mathbb{F}$ and $T : V \rightarrow W$ defined by $T(x) = 0$ for all $x \in V$. This is a linear transformation known as the **trivial transformation**

$$T(x + y) = 0 = 0 + 0 = T(x) + T(y)$$

and

$$T(cx) = 0 = c0 = cT(x)$$

for all $c \in \mathbb{F}$ and $x, y \in V$.

Example 3.1.7. The identity function on a vector space¹ $V(\mathbb{F})$ is also a linear transformation. Let $Id : V \rightarrow V$ satisfy $T(x) = x$ for each $x \in V$. Observe that

$$Id(x + cy) = x + cy = x + c \cdot y = Id(x) + c \cdot Id(y)$$

for all $x, y \in V$ and $c \in \mathbb{F}$.

Example 3.1.8. Define $T : C^0(\mathbb{R}) \rightarrow \mathbb{R}$ by $L(f) = \int_0^1 f(x)dx$. Notice that L is well-defined since all continuous functions are integrable and the value of a definite integral is a number. Furthermore,

$$T(f + cg) = \int_0^1 (f + cg)(x)dx = \int_0^1 [f(x) + cg(x)]dx = \int_0^1 f(x)dx + c \int_0^1 g(x)dx = T(f) + cT(g)$$

for all $f, g \in C^0(\mathbb{R})$ and $c \in \mathbb{R}$. The definite integral is a linear transformation.

Example 3.1.9. Let $T : C^1(\mathbb{R}) \rightarrow C^0(\mathbb{R})$ be defined by $T(f)(x) = f'(x)$ for each $f \in P_3$. We know from calculus that

$$T(f + g)(x) = (f + g)'(x) = f'(x) + g'(x) = T(f)(x) + T(g)(x)$$

and

$$T(cf)(x) = (cf)'(x) = cf'(x) = cT(f)(x).$$

The equations above hold for all $x \in \mathbb{R}$ thus we find function equations $T(f + g) = T(f) + T(g)$ and $T(cf) = cT(f)$ for all $f, g \in C^1(\mathbb{R})$ and $c \in \mathbb{R}$.

Example 3.1.10. Let V be the set of smooth functions in n -real variables $x_1, x_2, \dots, x_n$ then $T = \frac{\partial}{\partial x_{i_1}} \frac{\partial}{\partial x_{i_2}} \cdots \frac{\partial}{\partial x_{i_s}}$ defines a linear mapping on the set of smooth functions on $\mathbb{R}^n$.

Example 3.1.11. Let $a \in \mathbb{R}$. The evaluation mapping $\phi_a : \mathcal{F}(\mathbb{R}) \rightarrow \mathbb{R}$ is defined by $\phi_a(f) = f(a)$. This is a linear transformation as $(f + cg)(a) = f(a) + cg(a)$ by definition of function addition and scalar multiplication. (we could also replace $\mathbb{R}$ with $\mathbb{F}$ to obtain further examples)

¹remember, $V(\mathbb{F})$ simply denotes a vector space V over the field $\mathbb{F}$

3.1.2 basic theory of linear transformations

We assume V, W are vector spaces over a field $\mathbb{F}$ in the remainder of this section.

Proposition 3.1.12.

Let $L : V \rightarrow W$ be a linear transformation,

(1.) $L(0) = 0$

(2.) $L(c_1v_1 + c_2v_2 + \cdots + c_nv_n) = c_1L(v_1) + c_2L(v_2) + \cdots + c_nL(v_n)$ for all $v_i \in V$ and $c_i \in \mathbb{F}$.

Proof: to prove of (1.) let $x \in V$ and notice that $x - x = 0$ thus

$$L(0) = L(x - x) = L(x) + L(-1x) = L(x) - L(x) = 0.$$

To prove (2.) we use induction on n . Notice the proposition is true for $n=1,2$ by definition of linear transformation. Assume inductively $L(c_1v_1 + c_2v_2 + \cdots + c_nv_n) = c_1L(v_1) + c_2L(v_2) + \cdots + c_nL(v_n)$ for all $v_i \in V$ and $c_i \in \mathbb{F}$ where $i = 1, 2, \dots, n$. Let $v_1, v_2, \dots, v_n, v_{n+1} \in V$ and $c_1, c_2, \dots, c_n, c_{n+1} \in \mathbb{F}$ and consider, $L(c_1v_1 + c_2v_2 + \cdots + c_nv_n + c_{n+1}v_{n+1}) =$

$$\begin{aligned} &= L(c_1v_1 + c_2v_2 + \cdots + c_nv_n) + c_{n+1}L(v_{n+1}) && \text{by linearity of } L \\ &= c_1L(v_1) + c_2L(v_2) + \cdots + c_nL(v_n) + c_{n+1}L(v_{n+1}) && \text{by the induction hypothesis.} \end{aligned}$$

Hence the proposition is true for $n+1$ and we conclude by the principle of mathematical induction that (2.) is true for all $n \in \mathbb{N}$. $\square$

Proposition 3.1.13.

Let $L \in L(V, W)$. If S is linearly dependent then $L(S)$ is linearly dependent.

Proof: Suppose there exists $c_1, \dots, c_k \in \mathbb{F}$ for which $v = \sum_{i=1}^k c_i v_i$ is a linear dependence in S . Calculate,

$$L(v) = L\left(\sum_{i=1}^k c_i v_i\right) = \sum_{i=1}^k c_i L(v_i)$$

which, noting $L(v), L(v_i) \in L(S)$ for all $i \in \mathbb{N}_k$, shows $L(S)$ has a linear dependence. Therefore, $L(S)$ is linearly dependent. $\square$

Just as the column and null space of a matrix are important to understand the nature of the matrix we likewise study the kernel and image of a linear transformation: (Theorem 3.1.17 shows these are subspaces)

Definition 3.1.14.

Let V, W be vector spaces over a field $\mathbb{F}$ and $T \in \mathcal{L}(V, W)$ then

$$\text{Ker}(T) = \{v \in V \mid T(v) = 0\} \quad \& \quad \text{Image}(T) = \text{Range}(T) = \{T(v) \mid v \in V\}$$

Example 3.1.15. Let $T(f(x)) = f'(x)$ for $f(x) \in P_2(\mathbb{R})$ then clearly T is a linear transformation and we express the action of this map as $T(ax^2 + bx + c) = 2ax + b$. Hence,

$$\text{Ker}(T) = \{ax^2 + bx + c \in P_2(\mathbb{R}) \mid 2ax + b = 0\} = \{c \mid c \in \mathbb{R}\} = \mathbb{R}.$$

Also, $\text{Range}(T) = \text{span}\{1, x\}$ since $T(ax^2 + bx + c) = 2ax + b$ allows arbitrary $2a, b \in \mathbb{R}$.

Notice, the statement $\text{Ker}(T) = \{0\}$ means the only solution to the equation $T(x) = 0$ is $x = 0$.

Theorem 3.1.16. *linear map is injective iff only zero maps to zero.*

$T : V \rightarrow W$ is an injective linear transformation iff $\text{Ker}(T) = \{0\}$.

Proof: this is a biconditional statement. I'll prove the converse direction to begin.

($\Leftarrow$) Suppose $T(x) = 0$ iff $x = 0$ to begin. Let $x, y \in V$ and suppose $T(x) = T(y)$. By linearity we have $T(x - y) = T(x) - T(y) = 0$ hence $x - y = 0$ therefore $x = y$ and we find T is injective.

($\Rightarrow$) Suppose T is injective. Suppose $T(x) = 0$. Note $T(0) = 0$ by linearity of T but then by 1-1 property we have $T(x) = T(0)$ implies $x = 0$ hence the unique solution of $T(x) = 0$ is the zero solution. $\square$

For a linear transformation, the image of a subspace and the inverse image of a subspace are once again subspaces. This is certainly not true for arbitrary functions. In general, a nonlinear function takes linear spaces and twists them into all sorts of nonlinear shapes. For example, $f(x) = (x, x^2)$ takes the line $\mathbb{R}$ and pastes it onto the parabola $y = x^2$ in the range. We also can observe $f^{-1}\{(0, 0)\} = \{0\}$ and yet the mapping is certainly not injective. The theorems we find for linear functions do not usually generalize to functions in general²

Theorem 3.1.17.

If $T : V \rightarrow W$ is a linear transformation

(1.) and $V_o \leq V$ then $T(V_o) \leq W$,

(2.) and $W_o \leq W$ then $T^{-1}(W_o) \leq V$.

Proof: to prove (1.) suppose $V_o \leq V$. It follows $0 \in V_o$ and hence $T(0) = 0$ implies $0 \in T(V_o)$. Suppose $T(x), T(y) \in T(V_o)$ and $c \in \mathbb{F}$. Since $x, y \in V_o$ and V_o is a subspace we have $cx + y \in V_o$ thus $T(cx + y) \in T(V_o)$ and as

$$T(cx + y) = cT(x) + T(y)$$

hence $T(V_o) \neq \emptyset$ is closed under addition and scalar multiplication. Therefore, $T(V_o) \leq W$.

To prove (2.) suppose $W_o \leq W$ and observe $0 \in W_o$ and $T(0) = 0$ implies $0 \in T^{-1}(W_o)$. Hence $T^{-1}(W_o) \neq \emptyset$. Suppose $c \in \mathbb{F}$ and $x, y \in T^{-1}(W_o)$, it follows that there exist $x_o, y_o \in W_o$ such that $T(x) = x_o$ and $T(y) = y_o$. Observe, by linearity of T ,

$$T(cx + y) = cT(x) + T(y) = cx_o + y_o \in W_o.$$

²although, perhaps it's worth noting that in advanced calculus we learn how to linearize a function at a point. Some of our results here roughly generalize locally through the linearization and what are known as the inverse and implicit function theorems

hence $cx + y \in T^{-1}(W_o)$. Therefore, by the subspace theorem, $T^{-1}(W_o) \leq V$. $\square$

The special cases $V_o = V$ and $W_o = \{0\}$ merit discussion:

Corollary 3.1.18.

If $T : V \rightarrow W$ is a linear transformation then $T(V) \leq W$ and $T^{-1}\{0\} \leq V$. In other words, $\text{Ker}(T) \leq V$ and $\text{Range}(T) \leq W$.

Proof: observe $V \leq V$ and $\{0\} \leq W$ hence by Theorem 3.1.17 the Corollary holds true. $\square$

The following definitions of rank and nullity of a linear transformation are naturally connected to our prior use of the terms. In particular, we will soon³ see that to each linear transformation we can associate a matrix and the null and column space of the associated matrix will have the same nullity and rank as the kernel and image respective.

Definition 3.1.19.

Let V, W be vector spaces. If a mapping $T : V \rightarrow W$ is a linear transformation then

$$\dim(\text{Ker}(T)) = \text{nullity}(T) \quad \& \quad \dim(\text{Range}(T)) = \text{rank}(T).$$

Thus far in this section we have studied the behaviour of a particular linear transformation. In what follows, we see how to combine given linear transformations to form new linear transformations.

Definition 3.1.20.

Suppose $T : V \rightarrow W$ and $S : V \rightarrow W$ are linear transformations then we define $T + S, T - S$ and cT for any $c \in \mathbb{F}$ by the rules

$$(T + S)(x) = T(x) + S(x). \quad (T - S)(x) = T(x) - S(x), \quad (cT)(x) = cT(x)$$

for all $x \in V$.

The proof of the proposition below as it is nearly identical to the proof of Proposition 3.1.2.

Proposition 3.1.21.

If $T, S \in \mathcal{L}(V, W)$ and $c \in \mathbb{F}$ then $T + S, cT \in \mathcal{L}(V, W)$.

Proof: I'll be greedy and prove both at once: let $x, y \in V$ and $b, c \in \mathbb{F}$,

$$\begin{aligned} (T + cS)(x + by) &= T(x + by) + (cS)(x + by) && \text{defn. of sum of transformations} \\ &= T(x + by) + cS(x + by) && \text{defn. of scalar mult. of transformations} \\ &= T(x) + bT(y) + c[S(x) + bS(y)] && \text{linearity of } S \text{ and } T \\ &= T(x) + cS(x) + b[T(y) + cS(y)] && \text{vector algebra props.} \\ &= (T + cS)(x) + b(T + cS)(y) && \text{again, defn. of sum and scal. mult. of trans.} \end{aligned}$$

Let $c = 1$ and $b = 1$ to see $T + S$ is additive. Let $c = 1$ and $x = 0$ to see $T + S$ is homogeneous. Finally, let $T = 0$ to see cS is additive ($b = 1$) and homogeneous ($x = 0$). $\square$

³Lemma 3.3.16 and Proposition 3.3.17 to be precise

Recall that function space of all functions from V to W is naturally a vector space according to the point-wise addition and scalar multiplication of functions. It follows from the subspace theorem and the proposition above that:

Proposition 3.1.22.

The set of all linear transformations from V to W forms a vector space with respect to the natural point-wise addition and scalar multiplication of functions; $\mathcal{L}(V, W) \leq \mathcal{F}(V, W)$.

Proof: If $T, S \in L(V, W)$ and $c \in \mathbb{F}$ then $T + S, cT \in L(V, W)$ hence $L(V, W)$ is closed under addition and scalar multiplication. Moreover, the trivial function $T(x) = 0$ for all $x \in V$ is clearly in $L(V, W)$ hence $L(V, W) \neq \emptyset$ and we conclude by the subspace theorem that $L(V, W) \leq \mathcal{F}(V, W)$. $\square$

Function composition in the context of abstract vector spaces is the same as it was in precalculus.

Definition 3.1.23.

Suppose $T : V \rightarrow U$ and $S : U \rightarrow W$ are linear transformations then we define $S \circ T : V \rightarrow W$ by $(S \circ T)(x) = S(T(x))$ for all $x \in V$.

The composite of linear maps is once more a linear map.

Proposition 3.1.24.

Suppose $T \in L(V, U)$ and $S \in L(U, W)$ then $S \circ T \in L(V, W)$.

Proof: Let $x, y \in V$ and $c \in \mathbb{F}$,

$$\begin{aligned}
 (S \circ T)(x + cy) &= S(T(x + cy)) && \text{defn. of composite} \\
 &= S(T(x) + cT(y)) && T \text{ is linear trans.} \\
 &= S(T(x)) + cS(T(y)) && S \text{ is linear trans.} \\
 &= (S \circ T)(x) + c(S \circ T)(y) && \text{defn. of composite}
 \end{aligned}$$

additivity follows from $c = 1$ and homogeneity of $S \circ T$ follows from $x = 0$ thus $S \circ T \in L(V, W)$. $\square$

A vector space V together with a bilinear multiplication $m : V \times V \rightarrow V$ is called an **algebra**⁴. For example, we saw before that square matrices form an algebra with respect to addition and matrix multiplication. Notice that $V = L(W, W)$ is likewise naturally an algebra with respect to function addition and composition. One of our goals in this course is to understand the interplay between the algebra of transformations and the algebra of matrices.

The theorem below says the inverse of a linear transformation is also a linear transformation.

Theorem 3.1.25.

Suppose $T \in L(V, W)$ has inverse function $S : W \rightarrow V$ then $S \in L(W, V)$.

⁴it is somewhat ironic that all too often we often neglect to define an algebra in our modern algebra courses in the US educational system. As students, you ought to demand more. See Dummit and Foote for a precise definition

Proof: suppose $T \circ S = Id_W$ and $S \circ T = Id_V$. Suppose $x, y \in W$ hence there exists $a, b \in V$ for which $T(a) = x$ and $T(b) = y$. Also, let $c \in \mathbb{F}$. Consider,

$$\begin{aligned} S(cx + y) &= S(cT(a) + T(b)). \\ &= S(T(ca + b)) : && \text{by linearity of } T \\ &= ca + b : && \text{def. of identity function} \\ &= cS(x) + S(y) : && \text{note } a = S(T(a)) = S(x) \text{ and } b = S(T(b)) = S(y). \end{aligned}$$

Therefore, S is a linear transformation. $\square$

Observe, linearity of the inverse follows automatically from linearity of the map. Furthermore, it is useful for us to characterize the behaviour of LI sets under invertible linear transformations: What about LI of sets? If S is a LI subset of V and $T \in \mathcal{L}(V, W)$ then is $T(S)$ also LI? The answer is clearly no in general. Consider the trivial transformation of Example 3.1.6. On the other extreme we have the following:

Theorem 3.1.26.

If $T : V \rightarrow W$ is an injective linear transformation then S is LI implies $T(S)$ is LI. If $L : V \rightarrow W$ is any linear transformation and if U is LI in W then $L^{-1}(U)$ is LI in V .

Proof: suppose T is an injective linear transformation from V to W and suppose $S = \{s_1, \dots, s_k\}$ is a LI subset of V . Consider $T(S) = \{T(s_1), \dots, T(s_k)\}$. In particular, suppose there exist $c_1, \dots, c_k \in \mathbb{F}$ such that $c_1T(s_1) + \dots + c_kT(s_k) = 0$ implies $T(c_1s_1 + \dots + c_ks_k) = 0$ by Proposition 3.1.12. Thus $c_1s_1 + \dots + c_ks_k \in \text{Ker}(T) = \{0\}$ by Theorem 3.1.16. Consequently, $c_1s_1 + \dots + c_ks_k = 0$ hence $c_1 = 0, \dots, c_k = 0$ by LI of S .

Conversely, if $U = \{u_1, \dots, u_k\}$ is LI in W then suppose $x_1, \dots, x_n \in L^{-1}(U)$. Assume $c_1x_1 + \dots + c_nx_n = 0$ and observe:

$$L(c_1x_1 + \dots + c_nx_n) = L(0) \Rightarrow c_1L(x_1) + \dots + c_nL(x_n) = 0$$

but, by definition of $L^{-1}(U)$ we have $L(x_1), \dots, L(x_n) \in U$ and thus by LI of U we conclude $c_1 = 0, \dots, c_n = 0$. $\square$

Injective maps preserve linear independence whereas surjective maps preserve spanning:

Theorem 3.1.27.

If $T : V \rightarrow W$ is a surjective linear map then if $\text{span}(S) = V$ then $\text{span}(T(S)) = W$.

Proof: Suppose the linear map $T : V \rightarrow W$ is surjective and $\text{span}(S) = V$. Let $y \in W$ then as T is surjective there exists $x \in V$ for which $T(x) = y$. Since $x \in \text{span}(S)$ there exist $c_1, \dots, c_k \in \mathbb{F}$ and $v_1, \dots, v_k \in S$ for which $x = c_1v_1 + \dots + c_kv_k$. Hence

$$T(x) = T(c_1v_1 + \dots + c_kv_k) = c_1T(v_1) + \dots + c_kT(v_k).$$

Since $T(v_1), \dots, T(v_k) \in T(S)$ we find $y = T(x) \in \text{span}(T(S))$ and it follows $W = T(\text{span}(S))$. $\square$

3.1.3 standard matrices and their properties

I hope this section is a review from your previous course. I include it here in the interest of these notes being complete.

Theorem 3.1.28. *fundamental theorem of linear algebra.*

$L : \mathbb{F}^n \rightarrow \mathbb{F}^m$ is a linear transformation if and only if there exists $A \in \mathbb{F}^{m \times n}$ such that $L(x) = Ax$ for all $x \in \mathbb{F}^n$.

Proof: ($\Leftarrow$) Assume there exists $A \in \mathbb{F}^{m \times n}$ such that $L(x) = Ax$ for all $x \in \mathbb{F}^n$. Observe:

$$L(x + cy) = A(x + cy) = Ax + cAy = L(x) + cL(y)$$

for all $x, y \in \mathbb{F}^n$ and $c \in \mathbb{F}$ hence L is a linear transformation.

($\Rightarrow$) Assume $L : \mathbb{F}^n \rightarrow \mathbb{F}^m$ is a linear transformation. Let e_i denote the standard basis in $\mathbb{F}^n$. If $x \in \mathbb{F}^n$ then there exist constants $x_i \in \mathbb{F}$ such that $x = x_1e_1 + x_2e_2 + \cdots + x_ne_n$ and

$$\begin{aligned} L(x) &= L(x_1e_1 + x_2e_2 + \cdots + x_ne_n) \\ &= x_1L(e_1) + x_2L(e_2) + \cdots + x_nL(e_n) \\ &= [L(e_1)|L(e_2)|\cdots|L(e_n)] \begin{bmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{bmatrix} \end{aligned}$$

by Corollary 1.3.18. Let $A = [L(e_1)|L(e_2)|\cdots|L(e_n)]$ to see $L(x) = Ax$ as desired. $\square$

The fundamental theorem of linear algebra allows us to make the following definition.

Definition 3.1.29.

Let $L : \mathbb{F}^n \rightarrow \mathbb{F}^m$ be a linear transformation. The **standard matrix** of L is defined by:

$$[L] = [L(e_1)|L(e_2)|\cdots|L(e_n)].$$

The proof of the previous theorem makes it clear that $[L]$ is the unique matrix for which $L(x) = [L]x$. We also use the notation $L_A : \mathbb{F}^n \rightarrow \mathbb{F}^m$ to denote the linear transformation defined by **left-multiplication** by $A \in \mathbb{F}^{m \times n}$.

Example 3.1.30. Given that $L(x, y, z) = (x + 2y, 3y + 4z, 5x + 6z)$ for $(x, y, z) \in \mathbb{R}^3$ find the standard matrix of L . We wish to find a 3×3 matrix such that $L(v) = Av$ for all $v = (x, y, z) \in \mathbb{R}^3$. Write $L(v)$ then collect terms with each coordinate in the domain,

$$L \left(\begin{bmatrix} x \\ y \\ z \end{bmatrix} \right) = \begin{bmatrix} x + 2y \\ 3y + 4z \\ 5x + 6z \end{bmatrix} = x \begin{bmatrix} 1 \\ 0 \\ 5 \end{bmatrix} + y \begin{bmatrix} 2 \\ 3 \\ 0 \end{bmatrix} + z \begin{bmatrix} 0 \\ 4 \\ 6 \end{bmatrix}$$

It's not hard to see that,

$$L \left(\begin{bmatrix} x \\ y \\ z \end{bmatrix} \right) = \begin{bmatrix} 1 & 2 & 0 \\ 0 & 3 & 4 \\ 5 & 0 & 6 \end{bmatrix} \begin{bmatrix} x \\ y \\ z \end{bmatrix} \Rightarrow A = [L] = \begin{bmatrix} 1 & 2 & 0 \\ 0 & 3 & 4 \\ 5 & 0 & 6 \end{bmatrix}$$

Notice that the columns in A are just as you'd expect from the proof of theorem 3.1.28.

$[L] = [L(e_1)|L(e_2)|L(e_3)]$. In future examples I will exploit this observation to save writing.

Example 3.1.31. Suppose that $L((t, x, y, z)) = (t + x + y + z, z - x, 0, 3t - z)$, find $[L]$.

$$\begin{aligned} L(e_1) &= L((1, 0, 0, 0)) = (1, 0, 0, 3) \\ L(e_2) &= L((0, 1, 0, 0)) = (1, -1, 0, 0) \\ L(e_3) &= L((0, 0, 1, 0)) = (1, 0, 0, 0) \\ L(e_4) &= L((0, 0, 0, 1)) = (1, 1, 0, -1) \end{aligned} \Rightarrow [L] = \begin{bmatrix} 1 & 1 & 1 & 1 \\ 0 & -1 & 0 & 1 \\ 0 & 0 & 0 & 0 \\ 3 & 0 & 0 & -1 \end{bmatrix}.$$

I invite the reader to check my answer here and see that $L(v) = [L]v$ for all $v \in \mathbb{R}^4$ as claimed.

Very well, let's return to the concepts of injective and surjectivity of linear mappings. How do our theorems of LI and spanning inform us about the behaviour of linear transformations? The following pair of theorems summarize the situation nicely.

Theorem 3.1.32. *linear map is injective iff only zero maps to zero.*

$L : \mathbb{F}^n \rightarrow \mathbb{F}^m$ is a linear transformation with standard matrix $[L]$ then

(1.) L is 1-1 iff the columns of $[L]$ are linearly independent,

(2.) L is onto $\mathbb{F}^m$ iff the columns of $[L]$ span $\mathbb{F}^m$.

Proof: To prove (1.) use Theorem 3.1.16:

$$L \text{ is 1-1} \Leftrightarrow \left\{ L(x) = 0 \Leftrightarrow x = 0 \right\} \Leftrightarrow \left\{ [L]x = 0 \Leftrightarrow x = 0. \right\}$$

and the last equation simply states that if a linear combination of columns of L is zero then the coefficients of that linear equation are zero so (1.) follows.

To prove (2.) recall that if $A \in \mathbb{F}^{m \times n}$, $v \in \mathbb{F}^n$ then $Av = b$ is consistent for all $b \in \mathbb{F}^m$ if and only if the columns of A span $\mathbb{F}^m$. To say L is onto $\mathbb{F}^m$ means that for each $b \in \mathbb{F}^m$ there exists $v \in \mathbb{F}^n$ such that $L(v) = b$. But, this is equivalent to saying that $[L]v = b$ is consistent for each $b \in \mathbb{F}^m$ so (2.) follows. $\square$

The standard matrix enjoys many natural formulas. The standard matrix of the sum, difference or scalar multiple of linear transformations likewise the sum, difference or scalar multiple of the standard matrices of the respective linear transformations.

Proposition 3.1.33.

Suppose $T : \mathbb{F}^n \rightarrow \mathbb{F}^m$ and $S : \mathbb{F}^n \rightarrow \mathbb{F}^m$ are linear transformations then

$$(1.) [T + S] = [T] + [S], \quad (2.) [T - S] = [T] - [S], \quad (3.) [cS] = c[S].$$

Proof: Note $(T + cS)(e_j) = T(e_j) + cS(e_j)$ hence $((T + cS)(e_j))_i = (T(e_j))_i + c(S(e_j))_i$ for all i, j hence $[T + cS] = [T] + c[S]$. Set $c = 1$ to obtain (1.). Set $c = -1$ to obtain (2.). Finally, set $T = 0$ to obtain (3.). $\square$

Example 3.1.34. Suppose $T(x, y) = (x + y, x - y)$ and $S(x, y) = (2x, 3y)$. It's easy to see that

$$[T] = \begin{bmatrix} 1 & 1 \\ 1 & -1 \end{bmatrix} \text{ and } [S] = \begin{bmatrix} 2 & 0 \\ 0 & 3 \end{bmatrix} \Rightarrow [T + S] = [T] + [S] = \begin{bmatrix} 3 & 1 \\ 1 & 2 \end{bmatrix}$$

Therefore, $(T + S)(x, y) = \begin{bmatrix} 3 & 1 \\ 1 & 2 \end{bmatrix} \begin{bmatrix} x \\ y \end{bmatrix} = \begin{bmatrix} 3x + y \\ x + 2y \end{bmatrix} = (3x + y, x + 2y)$. Naturally this is the same formula that we would obtain through direct addition of the formulas of T and S .

Matrix multiplication is naturally connected to the problem of composition of linear maps.

Proposition 3.1.35.

$S : \mathbb{F}^p \rightarrow \mathbb{F}^m$ and $T : \mathbb{F}^n \rightarrow \mathbb{F}^p$ are linear transformations then $S \circ T : \mathbb{F}^n \rightarrow \mathbb{F}^m$ is a linear transformation with standard matrix $[S][T]$; that is, $[S \circ T] = [S][T]$.

Proof: Let us denote $\mathbb{F}^n = \text{span}\{e_i \mid i = 1, \dots, n\}$. To find the matrix of the composite we need only calculate its action on the standard basis; by definition, $\text{col}_j[S \circ T] = (S \circ T)(e_j)$. Observe⁵

$$\begin{aligned} (S \circ T)(e_j) &= S(T(e_j)) && : \text{def. of composite} \\ &= S([T]e_j) && : \text{def. of } [T] \\ &= [S]([T]e_j) && : \text{def. of } [S] \\ &= ([S][T])e_j && : \text{associativity of matrix multiplication} \end{aligned}$$

Thus $\text{col}_j([S \circ T]) = \text{col}_j([S][T])$ for $j = 1, \dots, n$ hence $[S \circ T] = [S][T]$. $\square$

Think about this: **matrix multiplication was defined to make the above proposition true.** Perhaps you wondered, why don't we just multiply matrices some other way? Well, now you have an answer. If we defined matrix multiplication differently then the result we just proved would not be true. However, as the course progresses, you'll see why it is so important that this result be true. It lies at the heart of many connections between the world of linear transformations and the world of matrices. It says we can trade composition of linear transformations for multiplication of matrices.

3.2 restriction, extension, isomorphism

Another way we can create new linear transformations from a given transformation is by restriction. Recall that the restriction of a given function is simply a new function where part of the domain has been removed. Since linear transformations are only defined on vector spaces we naturally are only permitted restrictions to subspaces of a given vector space.

Definition 3.2.1.

If $T : V \rightarrow W$ is a linear transformation and $U \subseteq V$ then we define $T|_U : U \rightarrow W$ by $T|_U(x) = T(x)$ for all $x \in U$. We say $T|_U$ is the **restriction of T to U** .

⁵I've lectured this calculation a few times, but it somehow never made it to my previous version of the notes, look back if you want to see how to make this like way more complicated, I think in columns more now than when I wrote the notes originally

Proposition 3.2.2.

If $T \in L(V, W)$ and $U \leq V$ then $T|_U \in L(U, W)$.

Proof: let $x, y \in U$ and $c \in \mathbb{F}$. Since $U \leq V$ it follows $cx + y \in U$ thus

$$T|_U(cx + y) = T(cx + y) = cT(x) + T(y) = cT|_U(x) + T|_U(y)$$

where I use linearity of T for the middle equality and the definition of $T|_U$ for the outside equalities. Therefore, $T|_U \in L(U, W)$. $\square$

We can create a linear transformation on an infinity of vectors by prescribing its values on the basis alone. This is a fantastic result.

Theorem 3.2.3. *linear transformations are fixed by their action on a basis.*

Suppose β is a basis for a vector space V and suppose W is also a vector space. Furthermore, suppose $L : \beta \rightarrow W$ is a function. There exists a unique linear extension of L to V .

Proof: to begin, let us understand the final sentence. A linear extension of L to V means a function $T : V \rightarrow W$ which is a linear transformation and $T|_\beta = L$. Uniqueness requires that we show if T_1, T_2 are two such extensions then $T_1 = T_2$. With that settled, let us begin the actual proof. I'll assume V is finite dimensional since that is our main application of this theorem, however, this result holds in the infinite dimensional context.

Suppose $\beta = \{v_1, \dots, v_n\}$ if $x \in V$ then there exist **unique** $x_1, \dots, x_n \in \mathbb{F}$ for which $x = \sum_{i=1}^n x_i v_i$. Therefore, define $T : V \rightarrow W$ as follows

$$T(x) = T\left(\sum_{i=1}^n x_i v_i\right) = \sum_{i=1}^n x_i L(v_i).$$

Clearly $T|_\beta = L$. Suppose $x, y \in V$ with $x = \sum_i x_i v_i$ and $y = \sum_i y_i v_i$ then $cx + y = c \sum_i x_i v_i + \sum_i y_i v_i = \sum_i (cx_i + y_i) v_i$ thus

$$\begin{aligned} T(cx + y) &= T\left(\sum_i (cx_i + y_i) v_i\right) \\ &= \sum_i (cx_i + y_i) L(v_i) \\ &= c \sum_i x_i L(v_i) + \sum_i y_i L(v_i) \\ &= cT(x) + T(y) \end{aligned}$$

thus $T \in L(V, W)$. Suppose T_1, T_2 are two such extensions. Consider, $x = \sum_{i=1}^n x_i v_i$

$$T_1(x) = T_1\left(\sum_{i=1}^n x_i v_i\right) = \sum_{i=1}^n x_i L(v_i).$$

However, the same calculation holds for $T_2(x)$ hence $T_1(x) = T_2(x)$ for all $x \in V$ therefore the extension T is unique. I have given the proof in the finite dimensional context, however, it is a

simple exercise to prove nearly the same proof applies in the infinite dimensional context. $\square$.

When we make use of the proposition above we typically use it to simplify a definition of a given linear transformation. In practice, we may define a mapping on a basis then **extend linearly**.

We conclude this section by initiating our discussion of **isomorphism**.

Definition 3.2.4.

Vector spaces $V(\mathbb{F})$ and $W(\mathbb{F})$ are **isomorphic** if there exists an invertible linear transformation $\Psi : V \rightarrow W$. Furthermore, an invertible linear transformation is called an **isomorphism**. We write $V \cong W$ if V and W are isomorphic.

Notice that it suffices to check $\Psi : V \rightarrow W$ is linear and invertible. Linearity of Ψ^{-1} follows by Theorem 3.1.25. This is nice as it means we have less work to do when proving some given mapping is an isomorphism. In fact, isomorphisms are especially nice in their relation to bases.

Lemma 3.2.5. *image of a basis is a basis under isomorphism.*

If $\Psi : V \rightarrow W$ is an isomorphism and β is a basis for $V(\mathbb{F})$ then $\Psi(\beta)$ is a basis for $W(\mathbb{F})$.

Proof: Theorem 3.1.26 tells us injective linear transformations map LI sets to LI sets. Thus, as an isomorphism Ψ is injective and the basis β is LI we find $\Psi(\beta)$ is a LI subset of W . Likewise, Theorem 3.1.27 tells us that a surjection maps spanning sets to spanning sets. Thus, $W = \text{span}(\beta)$ implies $W = \text{span}(\Psi(\beta))$. Therefore, $\Psi(\beta)$ is a linearly independent spanning set for W ; that is, $\Psi(\beta)$ is a basis for W . $\square$

Notice the Lemma above allows the possibility that V is an infinite dimensional vector space. In fact, since the cardinality of the basis is the dimension of the vector space we find the following interesting result:

Theorem 3.2.6.

$V \cong W$ if and only if $\dim(V) = \dim(W)$.

Proof: if $\Psi : V \rightarrow W$ is an isomorphism then Ψ is a bijection. Moreover, by the Lemma above, if β is a basis for V then $\Psi(\beta)$ is a basis for W . Hence $|\beta| = |\Psi(\beta)|$ as cardinality is preserved under bijection. Thus $\dim(V) = \dim(W)$.

Conversely, if $\dim(V) = \dim(W)$ then there exist bases β for V and γ for W for which there exists a bijection $F : \beta \rightarrow \gamma$. Extending F linearly to V gives $\Psi : V \rightarrow W$ a linear map. Likewise, extending $F^{-1} : \gamma \rightarrow \beta$ linearly gives $\Phi : W \rightarrow V$ a linear map. The reader is invited to prove $\Psi^{-1} = \Phi$ and thus Ψ is an isomorphism and hence $V \cong W$. $\square$

We should beware there are many practitioners of linear algebra who have been taught a careless formalism for infinite dimensional vector spaces which flattens infinite dimensional vector spaces of differing cardinality into a single monolithic whole. So, are two infinite dimensional vector spaces isomorphic? Answer: maybe. For instance, the set of algebraic numbers⁶ $\bar{\mathbb{Q}}$ over $\mathbb{Q}$ has dimension $\aleph_0$

⁶see Dummit and Foote's 3rd edition, the set of algebraic numbers include all numbers found in a finite extension field of $\mathbb{Q}$ like $\sqrt{3}$ or $3i = \sqrt{-9}$ etc, but not things like π or e which are transcendental over $\mathbb{Q}$, something beyond finite algebra is required to reach $\mathbb{C}$ which contains $\mathbb{R}$

whereas $\mathbb{C}(\mathbb{Q})$ has dimension $\aleph_1$. Hence, $\bar{\mathbb{Q}}$ and $\mathbb{C}$ are not isomorphic as $\mathbb{Q}$ -vector spaces since they have different dimensions. On the other hand, $\mathbb{R}(\mathbb{Q})$ and $\mathbb{C}(\mathbb{Q})$ are isomorphic as their dimensions are both $\aleph_1$.

Proposition 3.2.7.

If $T : V \rightarrow U$ and $S : U \rightarrow W$ are isomorphisms then $S \circ T$ is an isomorphism. Moreover, $\cong$ is an equivalence relation on the class of all vector spaces over a given field $\mathbb{F}$.

Proof: let $T \in \mathcal{L}(V, U)$ and $S \in \mathcal{L}(U, W)$ be isomorphisms. Recall Proposition 3.1.24 gives us $S \circ T \in \mathcal{L}(V, W)$ so, by Theorem 3.1.25, all that remains is to prove $S \circ T$ is invertible. Observe that $T^{-1} \circ S^{-1}$ serves as the inverse of $S \circ T$. In particular, calculate:

$$(S \circ T)(T^{-1} \circ S^{-1})(x) = S(T(T^{-1}(S^{-1}(x)))) = S(S^{-1}(x)) = x.$$

Thus $(S \circ T) \circ (T^{-1} \circ S^{-1}) = Id_W$. Similarly, $(T^{-1} \circ S^{-1}) \circ (S \circ T) = id_V$. Therefore $S \circ T$ is invertible with inverse $T^{-1} \circ S^{-1}$.

The proof that $\cong$ is an equivalence relation is not difficult. Begin by noting that $T = Id_V$ gives an isomorphism of V to V hence $V \cong V$; that is $\cong$ is reflexive. Next, if $T : V \rightarrow W$ is an isomorphism then $T^{-1} : W \rightarrow V$ is also an isomorphism by Theorem 3.1.25 thus $V \cong W$ implies $W \cong V$; $\cong$ is symmetric. Finally, suppose $V \cong U$ and $U \cong W$ by $T \in \mathcal{L}(V, U)$ and $S \in \mathcal{L}(U, W)$ are isomorphisms. We proved that $S \circ T \in \mathcal{L}(V, W)$ is an isomorphism hence $V \cong W$; that is, $\cong$ is transitive. Therefore, $\cong$ is an equivalence relation on the class of vector spaces over $\mathbb{F}$. $\square$

3.2.1 examples of isomorphisms

In this section I give you examples of isomorphisms. I do not supply proof that these maps are in fact as claimed, but it should be straight-forward to check on my assertions.

The coordinate map is an isomorphism which allows us to trade the abstract for the concrete.

Example 3.2.8. Let V be a vector space over $\mathbb{R}$ with basis $\beta = \{f_1, \dots, f_n\}$ and define Φ_β by $\Phi_\beta(f_j) = e_j \in \mathbb{R}^n$ extended linearly. In particular,

$$\Phi_\beta(v_1 f_1 + \dots + v_n f_n) = v_1 e_1 + \dots + v_n e_n.$$

This map is a linear bijection and it follows $V \approx \mathbb{R}^n$.

The notation $V(\mathbb{R})$ indicates I intend us to consider $V(\mathbb{R})$ as a vector space over the field $\mathbb{R}$.

Example 3.2.9. Suppose $V(\mathbb{R}) = \{A \in \mathbb{C}^{2 \times 2} \mid A^T = -A\}$ find an isomorphism to $P_n \leq \mathbb{R}[x]$ for appropriate n . Note, $A_{ij} = -A_{ji}$ gives $A_{11} = A_{22} = 0$ and $A_{12} = -A_{21}$. Thus, $A \in V$ has the form:

$$A = \begin{bmatrix} 0 & a + ib \\ -a - ib & 0 \end{bmatrix}$$

I propose that $\Psi(a + bx) = \begin{bmatrix} 0 & a + ib \\ -a - ib & 0 \end{bmatrix}$ provides an isomorphism of P_1 to V .

Example 3.2.10. Let $V(\mathbb{R}) = (\mathbb{C} \times \mathbb{R})^{2 \times 2}$ and $W(\mathbb{R}) = \mathbb{C}^{2 \times 3}$. The following is an isomorphism from V to W :

$$\Psi \begin{bmatrix} (z_1, x_1) & (z_2, x_2) \\ (z_3, x_3) & (z_4, x_4) \end{bmatrix} = \begin{bmatrix} z_1 & z_2 & z_3 \\ z_4 & x_1 + ix_2 & x_3 + ix_4 \end{bmatrix}$$

Example 3.2.11. Consider $P_2(\mathbb{C}) = \{ax^2 + bx + c \mid a, b, c \in \mathbb{C}\}$ as a complex vector space. Consider the subspace of $P_2(\mathbb{C})$ defined as $V = \{f(x) \in P_2(\mathbb{C}) \mid f(i) = 0\}$. Let's find an isomorphism to $\mathbb{C}^n$ for appropriate n . Let $f(x) = ax^2 + bx + c \in V$ and calculate

$$f(i) = a(i)^2 + bi + c = -a + bi + c = 0 \Rightarrow c = a - bi$$

Thus, $f(x) = ax^2 + bx + a - bi = a(x^2 + 1) + b(x - i)$. The isomorphism from V to $\mathbb{C}^2$ is apparent from the calculation above. If $f(x) \in V$ then we can write $f(x) = a(x^2 + 1) + b(x - i)$ and

$$\Psi(f(x)) = \Psi(a(x^2 + 1) + b(x - i)) = (a, b).$$

The inverse map is also easy to find: $\Psi^{-1}(a, b) = a(x^2 + 1) + b(x - i)$

Example 3.2.12. Let $\Psi(f(x), g(x)) = f(x) + x^{n+1}g(x)$ note this defines an isomorphism of $P_n \times P_n$ and P_{2n+1} . For example, $n = 1$,

$$\Psi((ax + b, cx + d)) = ax + b + x^2(cx + d) = cx^3 + dx^2 + ax + b.$$

The reason we need $2n+1$ is just counting: $\dim(P_n) = n+1$ and $\dim(P_n \times P_n) = 2(n+1)$. However, $\dim(P_{2n+1}) = (2n+1) + 1$. Notice, we could take coefficients of P_n in $\mathbb{R}$, $\mathbb{C}$ or some other field $\mathbb{F}$ and this example is still meaningful.

Example 3.2.13. Let $V = \mathbb{F}^{m \times n}$ and $W = \mathcal{L}(\mathbb{F}^n, \mathbb{F}^m)$. Consider $\Psi : V \rightarrow W$ given by $\Psi(A) = L_A$ where $L_A(x) = Ax$ for all $x \in \mathbb{F}^n$. Then $L_A \in W$ as desired and it is a simple exercise to check $\Psi(cA + B) = c\Psi(A) + \Psi(B)$. We can also show

$$\Psi^{-1}(L) = [L]$$

thus Ψ is an isomorphism as it is a linear map with linear inverse. Notice a basis for $\mathcal{L}(\mathbb{F}^n, \mathbb{F}^m)$ can be found by mapping the matrix-units E_{ij} to $\Psi(E_{ij}) = L_{E_{ij}}$. The mappings $L_{E_{ij}} : \mathbb{F}^n \rightarrow \mathbb{F}^m$ serve as a basis for $\mathcal{L}(\mathbb{F}^n, \mathbb{F}^m)$.

Example 3.2.14. Let $V = \mathcal{L}(\mathbb{R}^n, \mathbb{R}^m)$ and $W = \mathcal{L}(\mathbb{R}^m, \mathbb{R}^n)$. Transposition gives us a natural isomorphism as follows: for each $L \in V$ there exists $A \in \mathbb{R}^{m \times n}$ for which $L = L_A$. However, to $A^T \in \mathbb{R}^{n \times m}$ there naturally corresponds $L_{A^T} : \mathbb{R}^m \rightarrow \mathbb{R}^n$. Since V and W are spaces of functions an isomorphism is conveniently given in terms $A \mapsto L_A$ isomorphism of $\mathbb{R}^{m \times n}$ and $\mathcal{L}(\mathbb{R}^n, \mathbb{R}^m)$: in particular $\Psi : V \rightarrow W$ is given by:

$$\Psi(L_A) = L_{A^T}.$$

To write this isomorphism without the use of the L_A notation requires a bit more thought. Take off your shoes and socks, put them back on, then write what follows. Let $S \in V$ and $x \in \mathbb{R}^m$,

$$(\Psi(S))(x) = (x^T[S])^T = [S]^T x = L_{[S]^T}(x).$$

Since the above holds for all $x \in \mathbb{R}^m$ it can be written as $\Psi(S) = L_{[S]^T}$.

Example 3.2.15. Let $V = P_2$ and $W = \{f(x) \in y \mid f(1) = 0\}$. By the factor theorem of algebra we know $f(x) \in W$ implies $f(x) = (x - 1)g(x)$ where $g(x) \in P_2$. Define, $\Psi(f(x)) = g(x)$ where $g(x)(x - 1) = f(x)$. We argue that Ψ is an isomorphism. Note $\Psi^{-1}(g(x)) = (x - 1)g(x)$ and it is clear that $(x - 1)g(x) \in W$ moreover, linearity of Ψ^{-1} is simply seen from the calculation below:

$$\Psi^{-1}(cg(x) + h(x)) = (x - 1)(cg(x) + h(x)) = c(x - 1)g(x) + (x - 1)g(x) = c\Psi^{-1}(g(x)) + \Psi^{-1}(h(x)).$$

Linearity of Ψ follows by Theorem 3.1.25 as $\Psi = (\Psi^{-1})^{-1}$. Thus $V \cong W$.

You might note that I found a way around using a basis in the last example. Perhaps it is helpful to see the same example done by the basis mapping technique.

Example 3.2.16. Let $V = P_2$ and $W = \{f(x) \in y \mid f(1) = 0\}$. Ignoring the fact we **know** the factor theorem, let us find a basis the hard way: if $f(x) = ax^3 + bx^2 + cx + d \in W$ then

$$f(1) = a + b + c + d = 0$$

Thus, $d = -a - b - c$ and

$$f(x) = a(x^3 - 1) + b(x^2 - 1) + c(x - 1)$$

We find basis $\beta = \{x^3 - 1, x^2 - 1, x - 1\}$ for W . Define $\phi : W \rightarrow P_2$ by linearly extending:

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

In this case, a bit of thought reveals:

$$\phi^{-1}(ax^2 + bx + c) = a(x^3 - 1) + b(x^2 - 1) + c(x - 1).$$

Again, these calculations serve to prove $W \cong P_2$.

Example 3.2.17. Consider complex numbers $\mathbb{C}$ as a real vector space and let $M_{\mathbb{C}}$ be the set of real matrices of the form: $\begin{bmatrix} a & -b \\ b & a \end{bmatrix}$. Observe that the map $\Psi(a + ib) = \begin{bmatrix} a & -b \\ b & a \end{bmatrix}$ is a linear transformation with inverse $\Psi^{-1} \cdot \left(\begin{bmatrix} a & -b \\ b & a \end{bmatrix} \right) = a + ib$. Therefore, $\mathbb{C}$ and $M_{\mathbb{C}}$ are isomorphic as real vector spaces.

Let me continue past the point of linear isomorphism. In the example above, we can show that $\mathbb{C}$ and $M_{\mathbb{C}}$ are isomorphic as **algebras** over $\mathbb{R}$. In particular, notice

$$(a + ib)(c + id) = ac - bd + i(ad + bc) \quad \& \quad \begin{bmatrix} a & -b \\ b & a \end{bmatrix} \begin{bmatrix} c & -d \\ d & c \end{bmatrix} = \begin{bmatrix} ac - bd & -(ad + bc) \\ ad + bc & ac - bd \end{bmatrix}.$$

As you can see the pattern of the multiplication is the same. To be precise,

$$\Psi(\underbrace{(a + ib)(c + id)}_{\text{complex multiplication}}) = \underbrace{\Psi(a + ib)\Psi(c + id)}_{\text{matrix multiplication}}.$$

These special 2×2 matrices form a **representation** of the complex numbers. The term **isomorphism** has wide application in mathematics. In this course, the unqualified term "isomorphism" would be more descriptively termed "linear-isomorphism". An isomorphism of $\mathbb{R}$ -algebras is a linear isomorphism which also preserves the multiplication $\star$ of the algebra; $\Psi(v \star w) = \Psi(v)\Psi(w)$.

Another related concept, a non-associative algebra on a vector space which is a generalization of the cross-product of vectors in $\mathbb{R}^3$ is known as⁷ a Lie Algebra. In short, a Lie Algebra is a vector space paired with a Lie bracket. A *Lie algebra isomorphism* is a linear isomorphism which also preserves the Lie bracket; $\Psi([v, w]) = [\Psi(v), \Psi(w)]$. Not all isomorphisms are linear isomorphisms. For example, in abstract algebra you will study *isomorphisms of groups* which are bijections between groups which preserves the group multiplication. My point is just this, the idea of isomorphism, our current endeavor, is one you will see repeated as you continue your study of mathematics.

3.3 matrix of linear transformation

I used the notation $[v]_\beta$ in the last chapter since it was sufficient. Now we need to have better notation for the coordinate maps so we can articulate the concepts clearly.

Definition 3.3.1.

Let $V(\mathbb{F})$ be a finite dimensional vector space with basis $\beta = \{v_1, v_2, \dots, v_n\}$. The coordinate map $\Phi_\beta : V \rightarrow \mathbb{F}^n$ is defined by

$$\Phi_\beta(x_1v_1 + x_2v_2 + \dots + x_nv_n) = x_1e_1 + x_2e_2 + \dots + x_ne_n = (x_1, x_2, \dots, x_n)$$

for all $v = x_1v_1 + x_2v_2 + \dots + x_nv_n \in V$.

We argued in the previous section that Φ_β is an invertible, linear transformation from V to $\mathbb{F}^n$. In other words, Φ_β is an **isomorphism** of V and $\mathbb{F}^n$. It is worthwhile to note the linear extensions of

$$\Phi_\beta(v_i) = e_i \quad \& \quad \Phi_\beta^{-1}(e_i) = v_i$$

encapsulate the action of the coordinate map and its inverse. The coordinate map is a machine which converts an abstract basis to the standard basis.

Example 3.3.2. Let $V = \mathbb{R}^{2 \times 2}$ with basis $\beta = \{E_{11}, E_{12}, E_{21}, E_{22}\}$ then

$$\Phi_\beta \left(\begin{bmatrix} a & b \\ c & d \end{bmatrix} \right) = (a, b, c, d).$$

Example 3.3.3. Let $V = \mathbb{C}^n$ as a real vector space; that is $V(\mathbb{R}) = \mathbb{C}^n$. Consider the basis $\beta = \{e_1, \dots, e_n, ie_1, \dots, ie_n\}$ of this $2n$ -dimensional vector space over $\mathbb{R}$. Observe $v \in \mathbb{C}^n$ has $v = x + iy$ where $x, y \in \mathbb{R}^n$. In particular, if $\overline{a + ib} = a - ib$ and $\bar{v} = (\bar{v}_1, \dots, \bar{v}_n)$ then the identity below shows how to construct x, y :

$$v = \underbrace{\frac{1}{2}(v + \bar{v})}_{\text{Re}(v)=x} + \underbrace{\frac{1}{2}(v - \bar{v})}_{i\text{Im}(v)=iy}$$

and it's easy to verify $\bar{\bar{x}} = x$ and $\bar{\bar{y}} = y$ hence $x, y \in \mathbb{R}^n$ as claimed. The coordinate mapping is simple enough in this notation,

$$\Phi_\beta(x + iy) = (x, y).$$

Here we abuse notation slightly. Technically, I ought to write

$$\Phi_\beta(x + iy) = (x_1, \dots, x_n, y_1, \dots, y_n).$$

⁷it is pronounced "Lee", not what Fauci did

Example 3.3.4. Let $V = P_n$ with $\beta = \{1, (x-1), (x-1)^2, \dots, (x-1)^n\}$. To find the coordinates of an n -th order polynomial in standard form $f(x) = a_n x^n + \dots + a_1 x + a_0$ requires some calculation. We've all taken calculus II so we know Taylor's Theorem.

$$f(x) = \sum_{n=0}^{\infty} \frac{f^{(n)}(1)}{n!} (x-1)^n$$

also, clearly the series truncates for the polynomial in question hence,

$$f(x) = f(1) + f'(1)(x-1) + \frac{1}{2!} f''(1)(x-1)^2 + \dots + \frac{f^{(n)}(1)}{n!} (x-1)^n$$

Therefore,

$$\Phi_{\beta}(f(x)) = \left(f(1), f'(1), \frac{1}{2!} f''(1), \dots, \frac{1}{n!} f^{(n)}(1) \right).$$

Example 3.3.5. Let $V(\mathbb{R}) = \{A = \sum_{i,j=1}^2 A_{ij} E_{ij} \mid A_{11} + A_{22} = 0, A_{12} \in P_1, A_{11}, A_{22}, A_{21} \in \mathbb{C}\}$. If $A \in V$ then we can write:

$$A = \left[\begin{array}{c|c} a+ib & ct+d \\ \hline x+iy & -a-ib \end{array} \right]$$

A natural choice for basis β is seen

$$\beta = \left\{ \begin{bmatrix} 1 & 0 \\ 0 & -1 \end{bmatrix}, \begin{bmatrix} i & 0 \\ 0 & -i \end{bmatrix}, \begin{bmatrix} 0 & t \\ 0 & 0 \end{bmatrix}, \begin{bmatrix} 0 & 1 \\ 0 & 0 \end{bmatrix}, \begin{bmatrix} 0 & 0 \\ 1 & 0 \end{bmatrix}, \begin{bmatrix} 0 & 0 \\ i & 0 \end{bmatrix} \right\}$$

The coordinate mapping $\Phi_{\beta} : V \rightarrow \mathbb{R}^6$ follows easily in the notation laid out above,

$$\Phi_{\beta}(A) = (a, b, c, d, x, y).$$

Now that we have a little experience with coordinates as mappings let us turn to the central problem of this section: *how can we associate a matrix with a given linear transformation $T : V \rightarrow W$?* It turns out we'll generally have to choose a basis for $V(\mathbb{F})$ and $W(\mathbb{F})$ in order to answer this question unambiguously. Therefore, let β once more serve as the basis for V and suppose γ is a basis for W . We assume $\#(\beta), \#(\gamma) < \infty$ throughout this discussion. The answer to the question is actually in the diagram below:

$$\begin{array}{ccc} \boxed{V(\mathbb{F})} & \xrightarrow{T} & \boxed{W(\mathbb{F})} \\ \downarrow \Phi_{\beta} & & \downarrow \Phi_{\gamma} \\ \boxed{\mathbb{F}^n} & \xrightarrow{L_{[T]_{\beta, \gamma}}} & \boxed{\mathbb{F}^m} \end{array}$$

The matrix $[T]_{\beta, \gamma}$ induces a linear transformation from $\mathbb{F}^n$ to $\mathbb{F}^m$. This means $[T]_{\beta, \gamma} \in \mathbb{F}^{m \times n}$. It is defined by the demand that the diagram above **commutes**. Directly this means $\Phi_{\gamma} \circ T = L_{[T]_{\beta, \gamma}} \circ \Phi_{\beta}$. However, since the coordinate maps are invertible we can just as well write:

$$T = \Phi_{\gamma}^{-1} \circ L_{[T]_{\beta, \gamma}} \circ \Phi_{\beta}.$$

Or, on the other hand, we could write

$$L_{[T]_{\beta,\gamma}} = \Phi_\gamma \circ T \circ \Phi_\beta^{-1}.$$

The equation above indicates how to calculate $L_{[T]_{\beta,\gamma}}$ in terms of the coordinate maps and T directly. To select the i -th column in $[T]_{\beta,\gamma}$ we simply operate on $e_i \in \mathbb{F}^m$. This reveals,

$$\text{col}_i([T]_{\beta,\gamma}) = \Phi_\gamma(T(\Phi_\beta^{-1}(e_i)))$$

However, as we mentioned at the outset of this section, $\Phi_\beta^{-1}(e_i) = v_i$ hence

$$\text{col}_i([T]_{\beta,\gamma}) = \Phi_\gamma(T(v_i)) = [T(v_i)]_\gamma$$

where I have reverted to our previous notation for coordinate vectors⁸. Stringing the columns out, we find perhaps the nicest way to look at the matrix of an abstract linear transformation:

$$[T]_{\beta,\gamma} = [[T(v_1)]_\gamma | \cdots | [T(v_n)]_\gamma]$$

Each column is a W -coordinate vector which is found in $\mathbb{F}^m$ and these are given by the n -basis vectors which generate V .

As we remarked at the outset of this discussion, **commuting of the diagram** means:

$$\Phi_\gamma \circ T = L_{[T]_{\beta,\gamma}} \circ \Phi_\beta.$$

If we feed the expression above an arbitrary vector $v \in V$ we obtain:

$$\Phi_\gamma(T(v)) = L_{[T]_{\beta,\gamma}}(\Phi_\beta(v)) \quad \Rightarrow \quad [T(v)]_\gamma = [T]_{\beta,\gamma}[v]_\beta$$

In practice, as I work to formulate $[T]_{\beta,\gamma}$ for explicit problems I find the boxed formulas convenient for calculational purposes. On the other hand, I have used each formula on this page for various theoretical purposes. Ideally, you'd like to understand these rather than memorize. I hope you are annoyed I have yet to define $[T]_{\beta,\gamma}$. Let us pick a definition for specificity of future proofs.

Definition 3.3.6.

Let $V(\mathbb{F})$ be a vector space with basis $\beta = \{v_1, \dots, v_n\}$. Let $W(\mathbb{F})$ be a vector space with basis $\gamma = \{w_1, \dots, w_m\}$. If $T : V \rightarrow W$ is a linear transformation then we define the **matrix of T with respect to β, γ as $[T]_{\beta,\gamma} \in \mathbb{F}^{m \times n}$** which is implicitly defined by

$$L_{[T]_{\beta,\gamma}} = \Phi_\gamma \circ T \circ \Phi_\beta^{-1}.$$

The discussion preceding this definition hopefully gives you some idea on what I mean by "implicitly" in the above context. In any event, we pause from our general discussion to illustrate with some explicit examples.

Example 3.3.7. Let $S : V \rightarrow W$ with $V = W = \mathbb{R}^{2 \times 2}$ are given bases $\beta = \gamma = \{E_{11}, E_{12}, E_{21}, E_{22}\}$ and $L(A) = A + A^T$. Let $A = \begin{bmatrix} a & b \\ c & d \end{bmatrix}$ and calculate,

$$S(A) = \begin{bmatrix} a & b \\ c & d \end{bmatrix} + \begin{bmatrix} a & c \\ b & d \end{bmatrix} = \begin{bmatrix} 2a & b+c \\ b+c & 2d \end{bmatrix}$$

⁸the mapping notation supplements the $[v]_\beta$ notation, I use both going forward in these notes

Observe,

$$[A]_{\beta} = (a, b, c, d) \quad \& \quad [S(A)]_{\gamma} = (2a, b + c, b + c, 2d)$$

Moreover, we need a matrix $[S]_{\beta, \gamma}$ such that $[S(A)]_{\gamma} = [S]_{\beta, \gamma}[A]_{\beta}$. Tilt head, squint, and see

$$[S]_{\beta, \gamma} = \begin{bmatrix} 2 & 0 & 0 & 0 \\ 0 & 1 & 1 & 0 \\ 0 & 1 & 1 & 0 \\ 0 & 0 & 0 & 2 \end{bmatrix}$$

Example 3.3.8. Let $V(\mathbb{R}) = P_1^{2 \times 2}$ be the set of 2×2 matrices with first order polynomials. Define $T(A(x)) = A(2)$ where $T : V \rightarrow W$ and $W = \mathbb{R}^{2 \times 2}$. Let $\gamma = \{E_{11}, E_{12}, E_{21}, E_{22}\}$ be the basis for W . Let β be the basis⁹ with coordinate mapping

$$\Phi_{\beta} \left(\left[\begin{array}{c|c} a+bx & c+dx \\ \hline e+fx & g+hx \end{array} \right] \right) = (a, b, c, d, e, f, g, h).$$

We calculate for $v = \left[\begin{array}{c|c} a+bx & c+dx \\ \hline e+fx & g+hx \end{array} \right]$ that

$$T(v) = \left[\begin{array}{c|c} a+2b & c+2d \\ \hline e+2f & g+2h \end{array} \right]$$

Therefore,

$$[T(v)]_{\gamma} = (a + 2b, c + 2d, e + 2f, g + 2h)$$

and as the coordinate vector $[v]_{\beta} = (a, b, c, d, e, f, g, h)$ the formula $[T(v)]_{\gamma} = [T]_{\beta, \gamma}[v]_{\beta}$ indicates

$$[T]_{\beta, \gamma} = \begin{bmatrix} 1 & 2 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 1 & 2 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 1 & 2 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 1 & 2 \end{bmatrix}$$

Example 3.3.9. Suppose P_3 is the set of cubic polynomials with real coefficients. Let $T : P_3 \rightarrow P_3$ be the derivative operator; $T(f(x)) = f'(x)$. Give P_3 the basis $\beta = \{1, x, x^2, x^3\}$. Calculate,

$$T(a + bx + cx^2 + dx^3) = b + 2cx + 3dx^2$$

Furthermore, note, setting $v = a + bx + cx^2 + dx^3$

$$[T(v)]_{\beta} = (b, 2c, 3d, 0) \quad \& \quad [v]_{\beta} = (a, b, c, d) \quad \Rightarrow \quad [T]_{\beta, \beta} = \begin{bmatrix} 0 & 1 & 0 & 0 \\ 0 & 0 & 2 & 0 \\ 0 & 0 & 0 & 3 \\ 0 & 0 & 0 & 0 \end{bmatrix}$$

The results of Proposition 3.1.2 and 3.1.35 naturally generalize to our current context.

⁹you should be able to find β in view of the coordinate map formula

Proposition 3.3.10.

Let V, W be n, m -dimensional vector spaces over a field $\mathbb{F}$ with bases β, γ respective. Suppose $S, T \in L(V, W)$ then for $c \in \mathbb{F}$ we have $[T \pm S]_{\beta, \gamma}, [T]_{\beta, \gamma}, [S]_{\beta, \gamma}, [cS]_{\beta, \gamma} \in \mathbb{F}^{m \times n}$ and

$$(1.) [T + S]_{\beta, \gamma} = [T]_{\beta, \gamma} + [S]_{\beta, \gamma}, \quad (2.) [T - S]_{\beta, \gamma} = [T]_{\beta, \gamma} - [S]_{\beta, \gamma}, \quad (3.) [cS]_{\beta, \gamma} = c[S]_{\beta, \gamma}.$$

Proof: the proof follows immediately from the identity below:

$$\Phi_\gamma \circ (T + cS) \circ \Phi_\beta^{-1} = \Phi_\gamma \circ T \circ \Phi_\beta^{-1} + c\Phi_\gamma \circ S \circ \Phi_\beta^{-1}.$$

This identity is true due to the linearity properties of the coordinate mappings. $\square$

The generalization of Proposition 3.1.35 is a bit more interesting.

Proposition 3.3.11.

Let U, V, W be finite-dimensional vector spaces with bases β, γ, δ respectively. If $S \in L(U, W)$ and $T \in L(V, U)$ then $[S \circ T]_{\gamma, \delta} = [S]_{\beta, \delta}[T]_{\gamma, \beta}$

Proof: Notice that $L_A \circ L_B = L_{AB}$ since $L_A(L_B(v)) = L_A(Bv) = ABv = L_{AB}(v)$ for all v . Hence,

$$\begin{aligned} L_{[S]_{\beta, \delta}[T]_{\gamma, \beta}} &= L_{[S]_{\beta, \delta}} \circ L_{[T]_{\gamma, \beta}} && \text{:set } A = [S]_{\beta, \delta} \text{ and } B = [T]_{\gamma, \beta}, \\ &= (\Phi_\delta \circ S \circ \Phi_\beta^{-1}) \circ (\Phi_\beta \circ T \circ \Phi_\gamma^{-1}) && \text{:defn. of matrix of linear transformation,} \\ &= \Phi_\delta \circ (S \circ T) \circ \Phi_\gamma^{-1} && \text{:properties of function composition,} \\ &= L_{[S \circ T]_{\gamma, \delta}} && \text{:defn. of matrix of linear transformation.} \end{aligned}$$

Thus¹⁰ $[S \circ T]_{\gamma, \delta} = [S]_{\beta, \delta}[T]_{\gamma, \beta}$ as we claimed. $\square$

If we apply the result above to a linear transformation on a vector space V where the same basis is given to the domain and codomain some nice things occur. For example:

Example 3.3.12. Continuing Example 3.3.9. Observe that $T^2(f(x)) = T(T(f(x))) = f''(x)$. Thus if $v = a + bx + cx^2 + dx^3$ then $T^2 : P_3 \rightarrow P_3$ has $T^2(v) = 2c + 6dx$ hence $[T^2(v)]_\beta = (2c, 6d, 0, 0)$ and we find

$$[T^2]_{\beta, \beta} = \begin{bmatrix} 0 & 0 & 2 & 0 \\ 0 & 0 & 0 & 6 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \end{bmatrix}.$$

You can check that $[T^2]_{\beta, \beta} = [T]_{\beta, \beta}[T]_{\beta, \beta}$. Notice, we can easily see that $[T^3]_{\beta, \beta} \neq 0$ whereas $[T^n]_{\beta, \beta} = 0$ for all $n \geq 4$. This makes $[T]_{\beta, \beta}$ a **nilpotent** matrix of **index 4**.

The following example is pretty weird.

Example 3.3.13. It might be interesting to relate the results of Example 3.2.15 and Example 3.2.16. Examining the formula for $\Psi^{-1}(g(x)) = (x - 1)g(x)$ it is evident that we should factor out

¹⁰I use a little lemma here, two left multiplication maps are equal if and only if they multiply by the same matrix; $L_M = L_N$ if and only if $M = N$. Another way to look at this, the mapping $A \mapsto L_A$ is injective.

$(x - 1)$ from our ϕ^{-1} formula to connect to the Ψ^{-1} formula,

$$\begin{aligned}\phi^{-1}(ax^2 + bx + c) &= a(x - 1)(x^2 + x + 1) + b(x - 1)(x + 1) + c(x - 1). \\ &= (x - 1)[a(x^2 + x + 1) + b(x + 1) + c] \\ &= (x - 1)[ax^2 + (a + b)x + a + b + c] \\ &= \Psi^{-1}(ax^2 + (a + b)x + a + b + c).\end{aligned}$$

Evaluating the equation above by Ψ yields $(\Psi \circ \phi^{-1})(ax^2 + bx + c) = ax^2 + (a + b)x + a + b + c$. Therefore, if $\gamma = \{x^2, x, 1\}$ then we may easily deduce

$$[\Psi \circ \phi^{-1}]_{\gamma, \gamma} = \begin{bmatrix} 1 & 0 & 0 \\ 1 & 1 & 0 \\ 1 & 1 & 1 \end{bmatrix}.$$

Example 3.3.14. Let V, W be vector spaces of dimension n over $\mathbb{F}$. In addition, suppose $T : V \rightarrow W$ is a linear transformation with inverse $T^{-1} : W \rightarrow V$. Let V have basis β whereas W has basis γ . We know that $T \circ T^{-1} = Id_W$ and $T^{-1} \circ T = Id_V$. Furthermore, I invite the reader to show that $[Id_V]_{\beta, \beta} = I \in \mathbb{F}^{n \times n}$ where $n = \dim(V)$ and similarly $[Id_W]_{\gamma, \gamma} = I \in \mathbb{F}^{n \times n}$. Apply Proposition 3.3.11 to find

$$[T^{-1} \circ T]_{\beta, \beta} = [T^{-1}]_{\gamma, \beta} [T]_{\beta, \gamma}$$

but, $[T^{-1} \circ T]_{\beta, \beta} = [Id_V]_{\beta, \beta} = I$ thus $[T^{-1}]_{\gamma, \beta} [T]_{\beta, \gamma} = I$ and we conclude $([T]_{\beta, \gamma})^{-1} = [T^{-1}]_{\gamma, \beta}$. Phew, that's a relief. Wouldn't it be strange if this weren't true? Moral of story: **the inverse matrix of the transformation is the matrix of the inverse transformation.**

Example 3.3.15. Let $A = \begin{bmatrix} 1 & 1 & 1 & 1 \\ 2 & 2 & 3 & 0 \end{bmatrix}$ find an isomorphism from $\text{Null}(A)$ to $\text{Col}(A)$. As we recall, the CCP reveals all, we can easily calculate:

$$\text{rref}(A) = \begin{bmatrix} 1 & 1 & 0 & 3 \\ 0 & 0 & 1 & -2 \end{bmatrix}$$

Null space is $x \in \mathbb{R}^4$ for which $Ax = 0$ hence $x_1 = -x_2 - 3x_4$ and $x_3 = 2x_4$ with x_2, x_4 free. Thus,

$$x = (-x_2 - 3x_4, x_2, 2x_4, x_4) = x_2(-1, 1, 0, 0) + x_4(-3, 0, 2, 1)$$

and we find $\beta_N = \{(-1, 1, 0, 0), (-3, 0, 2, 1)\}$ is basis for $\text{Null}(A)$. On the other hand $\beta_C = \{(1, 2), (1, 3)\}$ forms a basis for the column space by the CCP. Let $\Psi : \text{Null}(A) \rightarrow \text{Col}(A)$ be defined by extending

$$\Psi((-1, 1, 0, 0)) = (1, 2) \quad \& \quad \Psi((-3, 0, 2, 1)) = (1, 3)$$

linearly. In particular, if $x \in \text{Null}(A)$ then $\Psi(x) = x_2(1, 2) + x_4(1, 3)$. Fun fact, with our choice of basis the matrix $[\Psi]_{\beta_N, \beta_C} = \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix}$

The matrix of a linear transformation reveals much about the transformation.

Proposition 3.3.16.

Let $T : V \rightarrow W$ be a linear transformation where $\dim(V) = n$ and $\dim(W) = m$. Let $\Phi_\beta : V \rightarrow \mathbb{F}^n$ and $\Phi_\gamma : W \rightarrow \mathbb{F}^m$ be coordinate map isomorphisms. If β, γ are bases for V, W respectively then $[T]_{\beta, \gamma}$ satisfies the following

$$(1.) \text{Null}([T]_{\beta, \gamma}) = \Phi_\beta(\text{Ker}(T)), \quad (2.) \text{Col}([T]_{\beta, \gamma}) = \Phi_\gamma(\text{Range}(T)).$$

Proof of (1.): Let $v \in \text{Null}([T]_{\beta, \gamma})$ then there exists $x \in V$ for which $v = [x]_\beta$. By definition of nullspace, $[T]_{\beta, \gamma}[x]_\beta = 0$ hence, applying the identity $[T(x)]_\gamma = [T]_{\beta, \gamma}[x]_\beta$ we obtain $[T(x)]_\gamma = 0$ which, by injectivity of Φ_γ , yields $T(x) = 0$. Thus $x \in \text{Ker}(T)$ and it follows that $[x]_\beta \in \Phi_\beta(\text{Ker}(T))$. Therefore, $\text{Null}([T]_{\beta, \gamma}) \subseteq \Phi_\beta(\text{Ker}(T))$.

Conversely, if $[x]_\beta \in \Phi_\beta(\text{Ker}(T))$ then there exists $v \in \text{Ker}(T)$ for which $\Phi_\beta(v) = [x]_\beta$ hence, by injectivity of Φ_β , $x = v$ and $T(x) = 0$. Observe, by linearity of Φ_γ , $[T(x)]_\gamma = 0$. Recall once more, $[T(x)]_\gamma = [T]_{\beta, \gamma}[x]_\beta$. Hence $[T]_{\beta, \gamma}[x]_\beta = 0$ and we conclude $[x]_\beta \in \text{Null}([T]_{\beta, \gamma})$. Consequently, $\Phi_\beta(\text{Ker}(T)) \subseteq \text{Null}([T]_{\beta, \gamma})$.

Thus $\Phi_\beta(\text{Ker}(T)) = \text{Null}([T]_{\beta, \gamma})$. I leave the proof of (2.) to the reader. $\square$

I should caution that the results above are basis dependent in the following sense: If β_1, β_2 are bases with coordinate maps $\Phi_{\beta_1}, \Phi_{\beta_2}$ then it is not usually true that $\Phi_{\beta_1}(\text{Ker}(T)) = \Phi_{\beta_2}(\text{Ker}(T))$. It follows that $\text{Null}([T]_{\beta_1, \gamma}) \neq \text{Null}([T]_{\beta_2, \gamma})$ in general. That said, there is something which is common to all the nullspaces (and ranges); dimension. The dimension of the nullspace must match the dimension of the kernel. The dimension of the column space must match the dimension of the range. This result follows immediately from Lemma 3.2.5 and Proposition 3.3.16.

Proposition 3.3.17.

Let $T : V \rightarrow W$ be a linear transformation of finite dimensional vector spaces with basis β for V and γ for W then

$$\text{nullity}(T) = \text{nullity}([T]_{\beta, \gamma}) \quad \& \quad \text{rank}(T) = \text{rank}([T]_{\beta, \gamma}).$$

You should realize the nullity and rank on the L.H.S. and R.H.S of the above proposition are quite different quantities in concept. It required some effort on our part to connect them, but, now that they are connected, perhaps you appreciate the names. Since we already know about rank and nullity for matrices from our study of row-reduction we obtain the following result:

Theorem 3.3.18. Rank-Nullity Theorem

Let $T : V \rightarrow W$ be a linear transformation of finite dimensional vector spaces then

$$\dim(V) = \text{rank}(T) + \text{nullity}(T)$$

where $\text{rank}(T) = \dim(\text{Range}(T))$ and $\text{nullity}(T) = \dim(\text{Ker}(T))$.

Proof: choose any pair of bases for V and W respectively, say β, γ and notice from matrix theory we know $\text{nullity}([T]_{\beta, \gamma})$ is the number of non-pivotal columns whereas $\text{rank}([T]_{\beta, \gamma})$ is the number of pivot columns. But $n = \#\beta = \dim(V)$ is the number of columns in $[T]_{\beta, \gamma}$ hence

$$\dim(V) = \text{rank}([T]_{\beta, \gamma}) + \text{nullity}([T]_{\beta, \gamma}) = \text{rank}(T) + \text{nullity}(T)$$

where we applied Proposition 3.3.17 in the final equality. $\square$

There is another proof of the Rank Nullity Theorem we give in conjunction with our study of the **Straightening Theorem 3.5.1**.

3.4 coordinate change

Vectors in abstract vector spaces do not generically come with a preferred coordinate system. There are infinitely many different choices for the basis of a given vector space. Naturally, for specific examples, we tend to have a pet-basis, but this is merely a consequence of our calculational habits. We need to find a way to compare coordinate vectors for different choices of basis. Then, the same ambiguity must be faced by the matrix of a transformation. In some sense, if you understand the diagrams then you can write all the required formulas for this section. That said, we will cut the problem for mappings of column vectors a bit more finely. There are nice matrix-theoretic formulas for $\mathbb{R}^n$ that I'd like for you to know when you leave this course¹¹.

3.4.1 coordinate change of abstract vectors

Let V be a vector space with bases β and $\bar{\beta}$ over the field $\mathbb{F}$. Let $\beta = \{v_1, \dots, v_n\}$ whereas $\bar{\beta} = \{\bar{v}_1, \dots, \bar{v}_n\}$. Let $x \in V$ then there exist **column vectors** $[x]_\beta = (x_1, \dots, x_n)$ and $[x]_{\bar{\beta}} = (\bar{x}_1, \dots, \bar{x}_n) \in \mathbb{F}^n$ such that

$$x = \sum_{i=1}^n x_i v_i \quad \& \quad x = \sum_{j=1}^n \bar{x}_j \bar{v}_j$$

Or, in mapping notation, $x = \Phi_\beta^{-1}([x]_\beta)$ and $x = \Phi_{\bar{\beta}}^{-1}([x]_{\bar{\beta}})$. Of course $x = x$ hence

$$\Phi_\beta^{-1}([x]_\beta) = \Phi_{\bar{\beta}}^{-1}([x]_{\bar{\beta}}).$$

Operate by $\Phi_{\bar{\beta}}$ on both sides,

$$[x]_{\bar{\beta}} = \Phi_{\bar{\beta}}(\Phi_\beta^{-1}([x]_\beta)).$$

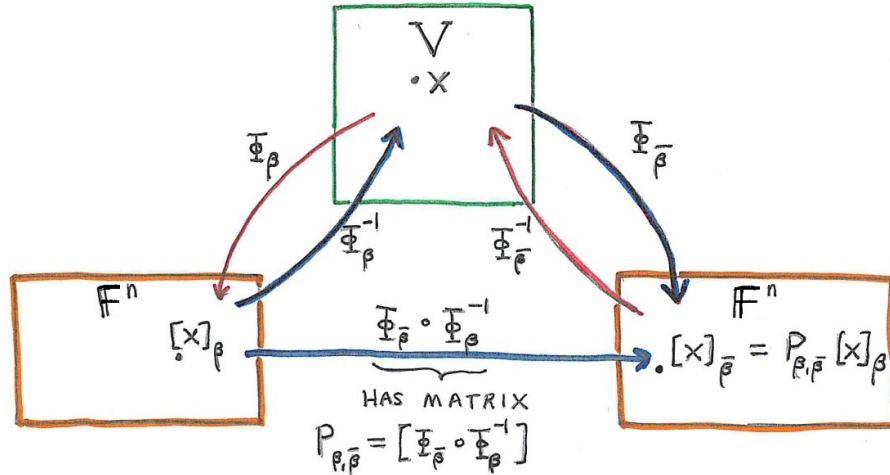
Observe that $\Phi_{\bar{\beta}} \circ \Phi_\beta^{-1} : \mathbb{F}^n \rightarrow \mathbb{F}^n$ is a linear transformation, therefore we can calculate its standard matrix. Let us collect our thoughts:

Proposition 3.4.1.

Using the notation developed in this subsection, if $P_{\beta, \bar{\beta}} = [\Phi_{\bar{\beta}} \circ \Phi_\beta^{-1}]$ then $[x]_{\bar{\beta}} = P_{\beta, \bar{\beta}}[x]_\beta$.

The diagram below contains the essential truth of the above proposition:

¹¹I mean, don't wait until then, now is a perfectly good time to learn them



Example 3.4.2. Let $V = \{A \in \mathbb{R}^{2 \times 2} \mid A_{11} + A_{22} = 0\}$. Observe $\beta = \{E_{12}, E_{21}, E_{11} - E_{22}\}$ gives a basis for V . On the other hand, $\bar{\beta} = \{E_{12} + E_{21}, E_{12} - E_{21}, E_{11} - E_{22}\}$ gives another basis. We denote $\beta = \{v_i\}$ and $\bar{\beta} = \{\bar{v}_i\}$. Let's work on finding the change of basis matrix. I can do this directly by our usual matrix theory. To find column i simply multiply by e_i . Or let the transformation act on e_i . The calculations below require a little thinking. I avoid algebra by thinking here.

$$\Phi_{\bar{\beta}}(\Phi_{\beta}^{-1}(e_1)) = \Phi_{\bar{\beta}}(E_{12}) = \Phi_{\bar{\beta}}\left(\frac{1}{2}[\bar{v}_1 + \bar{v}_2]\right) = (1/2, 1/2, 0).$$

$$\Phi_{\bar{\beta}}(\Phi_{\beta}^{-1}(e_2)) = \Phi_{\bar{\beta}}(E_{21}) = \Phi_{\bar{\beta}}\left(\frac{1}{2}[\bar{v}_1 - \bar{v}_2]\right) = (1/2, -1/2, 0).$$

$$\Phi_{\bar{\beta}}(\Phi_{\beta}^{-1}(e_3)) = \Phi_{\bar{\beta}}(E_{11} - E_{22}) = \Phi_{\bar{\beta}}(\bar{v}_3) = (0, 0, 1).$$

Admittably, if the bases considered were not so easily related we'd have some calculation to work through here. That said, we find:

$$P_{\beta, \bar{\beta}} = \begin{bmatrix} 1/2 & 1/2 & 0 \\ 1/2 & -1/2 & 0 \\ 0 & 0 & 1 \end{bmatrix}$$

Let's take it for a ride. Consider $A = \begin{bmatrix} 1 & 2 \\ 3 & -1 \end{bmatrix}$ clearly $[A]_{\beta} = (2, 3, 1)$. Calculate,

$$P_{\beta, \bar{\beta}}[A]_{\beta} = \begin{bmatrix} 1/2 & 1/2 & 0 \\ 1/2 & -1/2 & 0 \\ 0 & 0 & 1 \end{bmatrix} \begin{bmatrix} 2 \\ 3 \\ 1 \end{bmatrix} = \begin{bmatrix} 5/2 \\ -1/2 \\ 1 \end{bmatrix} = [A]_{\bar{\beta}}$$

Is this correct? Check,

$$\Phi_{\bar{\beta}}^{-1}(5/2, -1/2, 1) = \frac{5}{2} \cdot \begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix} - \frac{1}{2} \cdot \begin{bmatrix} 0 & 1 \\ -1 & 0 \end{bmatrix} + 1 \cdot \begin{bmatrix} 1 & 0 \\ 0 & -1 \end{bmatrix} = \begin{bmatrix} 1 & 2 \\ 3 & -1 \end{bmatrix} = A.$$

Yep. It works.

It is often the case we face coordinate change for mappings from $\mathbb{F}^n \rightarrow \mathbb{F}^m$. Or, even more special $m = n$. The formulas we've detailed thus far find streamlined matrix-theoretic forms in that special context. We turn our attention there now.

3.4.2 coordinate change for column vectors

Let β be a basis for $\mathbb{F}^n$. In contrast to the previous subsection, we have a standard basis with which we can compare; in particular, **the** standard basis. Huzzah!¹². Let $\beta = \{v_1, \dots, v_n\}$ and note the **matrix of** β is simply defined by concatenating the basis into an $n \times n$ invertible matrix $[\beta] = [v_1 | \dots | v_n]$. If $x \in \mathbb{F}^n$ then the coordinate vector $[x]_\beta = (y_1, \dots, y_n)$ is the column vector such that

$$x = [\beta][x]_\beta = y_1 v_1 + \dots + y_n v_n$$

here I used "y" to avoid some other more annoying notation. It is not written in stone, you could use $([x]_\beta)_i$ in place of y_i . Unfortunately, I cannot use x_i in place of y_i as the notation x_i is already reserved for the Cartesian components of x . Notice, as $[\beta]$ is invertible we can solve for the coordinate vector:

$$[x]_\beta = [\beta]^{-1}x$$

If we had another basis $\bar{\beta}$ then

$$[x]_{\bar{\beta}} = [\bar{\beta}]^{-1}x$$

Naturally, x exists independent of these bases so we find common ground at x :

$$x = [\beta][x]_\beta = [\bar{\beta}][x]_{\bar{\beta}}$$

We find the coordinate vectors are related by:

$$[x]_{\bar{\beta}} = [\bar{\beta}]^{-1}[\beta][x]_\beta$$

Let us summarize our findings in the proposition below:

Proposition 3.4.3.

Using the notation developed in this subsection and the last, if $P_{\beta, \bar{\beta}} = [\Phi_{\bar{\beta}} \circ \Phi_{\beta}^{-1}]$ then $[x]_{\bar{\beta}} = P_{\beta, \bar{\beta}}[x]_\beta$ and a simple formula to calculate the change of basis matrix is given by $P_{\beta, \bar{\beta}} = [\bar{\beta}]^{-1}[\beta]$. We also note for future convenience: $[\bar{\beta}]P_{\beta, \bar{\beta}} = [\beta]$

Example 3.4.4. Let $\beta = \{(1, 1), (1, -1)\}$ and $\gamma = \{(1, 0), (1, 1)\}$ be bases for $\mathbb{R}^2$. Find $[v]_\beta$ and $[v]_\gamma$ if $v = (2, 4)$. Let me frame the problem, we wish to solve:

$$v = [\beta][v]_\beta \quad \text{and} \quad v = [\gamma][v]_\gamma$$

where I'm using the basis in brackets to denote the matrix formed by concatenating the basis into a single matrix,

$$[\beta] = \begin{bmatrix} 1 & 1 \\ 1 & -1 \end{bmatrix} \quad \text{and} \quad [\gamma] = \begin{bmatrix} 1 & 1 \\ 0 & 1 \end{bmatrix}$$

This is the 2×2 case so we can calculate the inverse from our handy-dandy formula:

$$[\beta]^{-1} = \frac{1}{2} \begin{bmatrix} 1 & 1 \\ 1 & -1 \end{bmatrix} \quad \text{and} \quad [\gamma]^{-1} = \begin{bmatrix} 1 & -1 \\ 0 & 1 \end{bmatrix}$$

Then multiplication by inverse yields $[v]_\beta = [\beta]^{-1}v$ and $[v]_\gamma = [\gamma]^{-1}v$ thus:

$$[v]_\beta = \frac{1}{2} \begin{bmatrix} 1 & 1 \\ 1 & -1 \end{bmatrix} \begin{bmatrix} 2 \\ 4 \end{bmatrix} = \begin{bmatrix} 3 \\ -1 \end{bmatrix} \quad \text{and} \quad [v]_\gamma = \begin{bmatrix} 1 & -1 \\ 0 & 1 \end{bmatrix} \begin{bmatrix} 2 \\ 4 \end{bmatrix} = \begin{bmatrix} -2 \\ 4 \end{bmatrix}$$

¹²sorry, we visited Medieval Times over vacation and it hasn't entirely worn off just yet

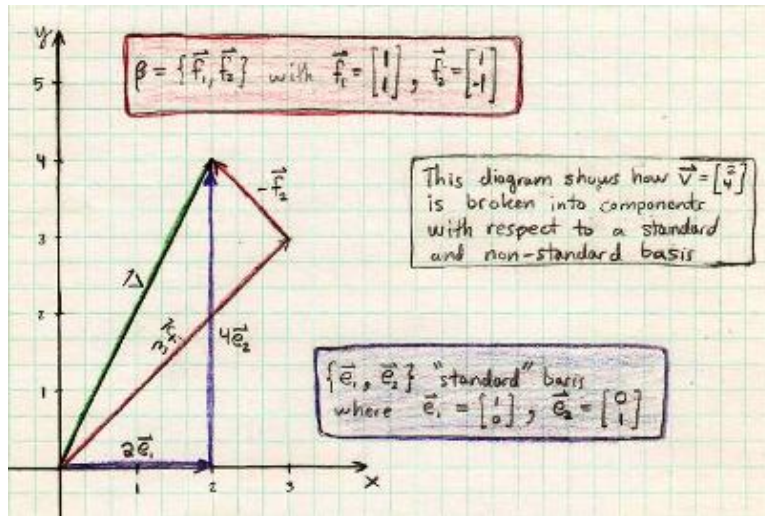
Let's verify the relation of $[v]_\gamma$ and $[v]_\beta$ relative to the change of basis matrix. In particular, we expect that if $P_{\beta,\gamma} = [\gamma]^{-1}[\beta]$ then $[v]_\gamma = P_{\beta,\gamma}[v]_\beta$. Calculate,

$$P_{\beta,\gamma} = [\gamma]^{-1}[\beta] = \begin{bmatrix} 1 & -1 \\ 0 & 1 \end{bmatrix} \begin{bmatrix} 1 & 1 \\ 1 & -1 \end{bmatrix} = \begin{bmatrix} 0 & 2 \\ 1 & -1 \end{bmatrix}$$

As the last great American president¹³, **trust, but, verify**

$$P_{\beta,\gamma}[v]_\beta = \begin{bmatrix} 0 & 2 \\ 1 & -1 \end{bmatrix} \begin{bmatrix} 3 \\ -1 \end{bmatrix} = \begin{bmatrix} -2 \\ 4 \end{bmatrix} = [v]_\gamma \quad \checkmark$$

It might be helpful to some to see a picture of just what we have calculated. Finding different coordinates for a given point (which corresponds to a vector from the origin) is just to prescribe different zig-zag paths from the origin along basis-directions to get to the point. In the picture below I illustrate the standard basis path and the β -basis path.



Now that we've seen an example, let's find $[v]_\beta$ for an arbitrary $v = (x, y)$,

$$[v]_\beta = \frac{1}{2} \begin{bmatrix} 1 & 1 \\ 1 & -1 \end{bmatrix} \begin{bmatrix} x \\ y \end{bmatrix} = \begin{bmatrix} \frac{1}{2}(x+y) \\ \frac{1}{2}(x-y) \end{bmatrix}$$

If we denote $[v]_\beta = (\bar{x}, \bar{y})$ then we can understand the calculation above as the relation between the barred and standard coordinates:

$$\bar{x} = \frac{1}{2}(x+y) \quad \bar{y} = \frac{1}{2}(x-y)$$

Conversely, we can solve these for x, y to find the inverse transformations:

$$x = \bar{x} + \bar{y} \quad y = \bar{x} - \bar{y}.$$

Similar calculations are possible with respect to the γ -basis.

¹³Upon further reflection, he wasn't so great, his approval of the bill which made drug companies free from litigation stemming from harming citizens whose lives were destroyed by dangerous vaccines has done lasting damage to this country and his compromises on immigration and failure to remove the Department of Education left America open to attack from within. Moreover, the EITC has incentivized tax fraud for decades. Finally, he paved the road for Bush and Bush Jr. to involve us in foreign wars and grow domestic surveillance openly in the name of freedom.

3.4.3 coordinate change of abstract linear transformations

In Definition 3.3.6 we saw that if $V(\mathbb{F})$ is a vector space with basis $\beta = \{v_1, \dots, v_n\}$ and $W(\mathbb{F})$ be a vector space with basis $\gamma = \{w_1, \dots, w_m\}$. Then a linear transformation $T : V \rightarrow W$ has matrix $[T]_{\beta, \gamma} \in \mathbb{F}^{m \times n}$ defined implicitly by:

$$L_{[T]_{\beta, \gamma}} = \Phi_\gamma \circ T \circ \Phi_\beta^{-1}.$$

If there was another pair of bases $\bar{\beta}$ for V and $\bar{\gamma}$ for W then we would likewise have

$$L_{[T]_{\bar{\beta}, \bar{\gamma}}} = \Phi_{\bar{\gamma}} \circ T \circ \Phi_{\bar{\beta}}^{-1}.$$

Solving for T relates the matrices with and without bars:

$$T = \Phi_\gamma^{-1} \circ L_{[T]_{\beta, \gamma}} \circ \Phi_\beta = \Phi_{\bar{\gamma}}^{-1} \circ L_{[T]_{\bar{\beta}, \bar{\gamma}}} \circ \Phi_{\bar{\beta}}.$$

From which the proposition below follows:

Proposition 3.4.5.

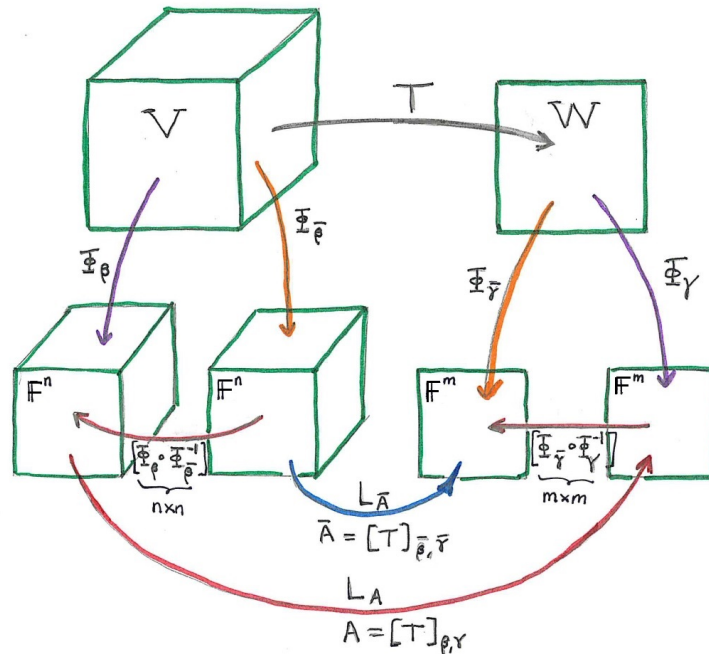
Using the notation developed in this subsection

$$[T]_{\bar{\beta}, \bar{\gamma}} = [\Phi_{\bar{\gamma}} \circ \Phi_\gamma^{-1}] [T]_{\beta, \gamma} [\Phi_\beta \circ \Phi_{\bar{\beta}}^{-1}].$$

Moreover, recalling $P_{\beta, \bar{\beta}} = [\Phi_{\bar{\beta}} \circ \Phi_\beta^{-1}]$ we find:

$$[T]_{\bar{\beta}, \bar{\gamma}} = P_{\gamma, \bar{\gamma}} [T]_{\beta, \gamma} (P_{\beta, \bar{\beta}})^{-1}.$$

Suppose $B, A \in \mathbb{F}^{m \times n}$. If there exist invertible matrices $P \in \mathbb{F}^{m \times m}, Q \in \mathbb{F}^{n \times n}$ such that $B = PAQ$ for then we say B and A are **matrix congruent**. The proposition above indicates that the matrices of a given linear transformation¹⁴ are congruent. In particular, $[T]_{\bar{\beta}, \bar{\gamma}}$ is congruent to $[T]_{\beta, \gamma}$. The picture below can be used to remember the formulas in the proposition above.



¹⁴of finite dimensional vector spaces

Example 3.4.6. Let $V = P_2$ and $W = \mathbb{C}$. Define a linear transformation $T : V \rightarrow W$ by $T(f) = f(i)$. Thus,

$$T(ax^2 + bx + c) = ai^2 + bi + c = c - a + ib.$$

Use coordinate maps given below for $\beta = \{x^2, x, 1\}$ and $\gamma = \{1, i\}$:

$$\Phi_\beta(ax^2 + bx + c) = (a, b, c) \quad \& \quad \Phi_\gamma(a + ib) = (a, b).$$

Observe $[T(ax^2 + bx + c)]_\gamma = (c - a, b)$ hence $[T]_{\beta, \gamma} = \begin{bmatrix} -1 & 0 & 1 \\ 0 & 1 & 0 \end{bmatrix}$.

Let us change the bases to

$$\bar{\beta} = \{(x-2)^2, (x-2), 1\} \quad \& \quad \bar{\gamma} = \{i, 1\}$$

Calculate, if $f(x) = ax^2 + bx + c$ then $f'(x) = 2ax + b$ and $f''(x) = 2a$. Observe, $f(2) = 4a + 2b + c$ and $f'(2) = 4a + b$ and $f''(2) = 2a$ hence, using the Taylor expansion centered at 2,

$$\begin{aligned} f(x) &= f(2) + f'(2)(x-2) + \frac{1}{2}f''(2)(x-2)^2 \\ &= 4a + 2b + c + (4a + b)(x-2) + a(x-2)^2. \end{aligned}$$

Therefore,

$$\Phi_{\bar{\beta}}(ax^2 + bx + c) = (a, 4a + b, 4a + 2b + c)$$

But, $\Phi_{\bar{\beta}}^{-1}(a, b, c) = ax^2 + bx + c$. Thus,

$$\Phi_{\bar{\beta}}(\Phi_{\bar{\beta}}^{-1}(a, b, c)) = (a, 4a + b, 4a + 2b + c) \quad \Rightarrow \quad [\Phi_{\bar{\beta}} \circ \Phi_{\bar{\beta}}^{-1}] = \begin{bmatrix} 1 & 0 & 0 \\ 4 & 1 & 0 \\ 4 & 2 & 1 \end{bmatrix}$$

Let's work out this calculation in the other direction (it's actually easier and what we need in a bit)

$$\Phi_{\beta}(a(x-2)^2 + b(x-2) + c) = \Phi_{\beta}(a(x^2 - 4x + 4) + b(x-2) + c) = (a, -4a + b, 4a - 2b + c)$$

But, $\Phi_{\bar{\beta}}^{-1}(a, b, c) = a(x-2)^2 + b(x-2) + c$ therefore:

$$\Phi_{\beta}(\Phi_{\bar{\beta}}^{-1}(a, b, c)) = (4a - 2b + c, -4a + b, a) \quad \Rightarrow \quad [\Phi_{\beta} \circ \Phi_{\bar{\beta}}^{-1}] = \begin{bmatrix} 1 & 0 & 0 \\ -4 & 1 & 0 \\ 4 & -2 & 1 \end{bmatrix}$$

On the other hand, $\Phi_{\bar{\gamma}}(a + ib) = (b, a)$. Of course, $a + ib = \Phi_{\gamma}^{-1}(a, b)$ hence $\Phi_{\bar{\gamma}}(\Phi_{\gamma}^{-1}(a, b)) = (b, a)$.

It follows that $[\Phi_{\bar{\gamma}} \circ \Phi_{\gamma}^{-1}] = \begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix}$ We'll use the change of basis proposition to find the matrix w.r.t. $\bar{\beta}$ and $\bar{\gamma}$

$$\begin{aligned} [T]_{\bar{\beta}, \bar{\gamma}} &= [\Phi_{\bar{\gamma}} \circ \Phi_{\gamma}^{-1}][T]_{\beta, \gamma}[\Phi_{\beta} \circ \Phi_{\bar{\beta}}^{-1}]. \\ &= \begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix} \begin{bmatrix} -1 & 0 & 1 \\ 0 & 1 & 0 \end{bmatrix} \begin{bmatrix} 1 & 0 & 0 \\ -4 & 1 & 0 \\ 4 & -2 & 1 \end{bmatrix} \\ &= \begin{bmatrix} 0 & 1 & 0 \\ -1 & 0 & 1 \end{bmatrix} \begin{bmatrix} 1 & 0 & 0 \\ -4 & 1 & 0 \\ 4 & -2 & 1 \end{bmatrix} \\ &= \begin{bmatrix} -4 & 1 & 0 \\ 3 & -2 & 1 \end{bmatrix}. \end{aligned}$$

Continuing, we can check this by direct calculation of the matrix. Observe

$$\begin{aligned} T(a(x-2)^2 + b(x-2) + c) &= a(i-2)^2 + b(i-2) + c \\ &= a[-1 - 4i + 4] + b(i-2) + c \\ &= 3a - 2b + c + i(-4a + b) \end{aligned}$$

Thus, $[T(a(x-2)^2 + b(x-2) + c)]_{\bar{\gamma}} = (-4a + b, 3a - 2b + c)$ hence $[T]_{\bar{\beta}, \bar{\gamma}} = \begin{bmatrix} -4 & 1 & 0 \\ 3 & -2 & 1 \end{bmatrix}$. Which agrees nicely with our previous calculation.

3.4.4 coordinate change of linear transformations of column vectors

We specialize Proposition 3.4.5 in this subsection in the case that $V = \mathbb{F}^n$ and $W = \mathbb{F}^m$. In particular, the result of Proposition 3.4.3 makes life easy; $P_{\beta, \bar{\beta}} = [\bar{\beta}]^{-1}[\beta]$ likewise, $P_{\gamma, \bar{\gamma}} = [\bar{\gamma}]^{-1}[\gamma]$

Proposition 3.4.7.

Using the notation developed in this subsection

$$[T]_{\bar{\beta}, \bar{\gamma}} = [\bar{\gamma}]^{-1}[\gamma][T]_{\beta, \gamma}[\beta]^{-1}[\bar{\beta}].$$

The standard matrix $[T]$ is related to the non-standard matrix $[T]_{\bar{\beta}, \bar{\gamma}}$ by:

$$[T]_{\bar{\beta}, \bar{\gamma}} = [\bar{\gamma}]^{-1}[T][\bar{\beta}].$$

Proof: Proposition 3.4.5 with $V = \mathbb{F}^n$ and $W = \mathbb{F}^m$ together with the result of Proposition 3.4.3 give us the first equation. The second equation follows from the observation that for standard bases β and γ we have $[\beta] = I_n$ and $[\gamma] = I_m$. $\square$

Example 3.4.8. Let $\bar{\beta} = \{(1, 0, 1), (0, 1, 1), (4, 3, 1)\}$. Furthermore, define a linear transformation $T : \mathbb{R}^3 \rightarrow \mathbb{R}^3$ by the rule $T(x, y, z) = (2x - 2y + 2z, x - z, 2x - 3y + 2z)$. Find the matrix of T with respect to the basis $\bar{\beta}$. Note first that the standard basis is read from the rule:

$$T\left(\begin{bmatrix} x \\ y \\ z \end{bmatrix}\right) = \begin{bmatrix} 2x - 2y + 2z \\ x - z \\ 2x - 3y + 2z \end{bmatrix} = \begin{bmatrix} 2 & -2 & 2 \\ 1 & 0 & -1 \\ 2 & -3 & 2 \end{bmatrix} \begin{bmatrix} x \\ y \\ z \end{bmatrix}$$

Next, use the proposition with $\bar{\beta} = \bar{\gamma}$ (omitting the details of calculating $[\bar{\beta}]^{-1}$)

$$\begin{aligned} [\bar{\beta}]^{-1}[T][\bar{\beta}] &= \begin{bmatrix} 1/3 & -2/3 & 2/3 \\ -1/2 & 1/2 & 1/2 \\ 1/6 & 1/6 & -1/6 \end{bmatrix} \begin{bmatrix} 2 & -2 & 2 \\ 1 & 0 & -1 \\ 2 & -3 & 2 \end{bmatrix} \begin{bmatrix} 1 & 0 & 4 \\ 0 & 1 & 3 \\ 1 & 1 & 1 \end{bmatrix} \\ &= \begin{bmatrix} 1/3 & -2/3 & 2/3 \\ -1/2 & 1/2 & 1/2 \\ 1/6 & 1/6 & -1/6 \end{bmatrix} \begin{bmatrix} 4 & 0 & 4 \\ 0 & -1 & 3 \\ 4 & -1 & 1 \end{bmatrix} \\ &= \begin{bmatrix} 4 & 0 & 0 \\ 0 & -1 & 0 \\ 0 & 0 & 1 \end{bmatrix} \end{aligned}$$

Therefore, in the $\bar{\beta}$ -coordinates the linear operator T takes on a particularly simple form. In particular, if $\bar{\beta} = \{f_1, f_2, f_3\}$ then¹⁵

$$\bar{T}(\bar{x}, \bar{y}, \bar{z}) = 4\bar{x}f_1 - \bar{y}f_2 + \bar{z}f_3$$

This linear transformation acts in a special way in the f_1, f_2 and f_3 directions. The basis we considered here is called an **eigenbasis** for T . We study eigenvectors and the associated problem of diagonalization later in this course.

3.5 theory of dimensions for maps

This section is yet another encounter with a **classification theorem**. Previously, we learned that vector spaces are classified by their dimension; $V \cong W$ iff $\dim(V) = \dim(W)$. In this section, we'll find a nice way to lump together many linear transformations as being essentially the same function with a change of notation. In this section, **matrix congruence** is the measure of sameness. Given $A, B \in \mathbb{F}^{m \times n}$ we say A and B are **matrix congruent** if there exist invertible matrices P, Q for which

$$B = PAQ.$$

In contrast, two square matrices $A, B \in \mathbb{F}^{n \times n}$ are said to be **similar** if there exists an invertible matrix P for which

$$B = P^{-1}AP.$$

Both similarity and congruence give equivalence relations on appropriate sets of matrices. In the context of square matrices these are not the same concept of similarity. From the viewpoint of linear transformations, we encounter matrix congruence when changing coordinates in the domain and codomain of a given linear transformation in an *independent* fashion. On the other hand, we encounter similarity transformations when changing coordinates for the matrix of a linear transformation on a given vector space where we use the same basis for both the domain and the codomain.

It turns out the problem of deciding whether two square matrices are similar is quite difficult. In contrast, the problem of deciding if two matrices are congruent is an easy corollary to the following theorem which I call the **straightening theorem**. Essentially, it means we can always change coordinates on a linear transformation to make the formula for the transformation a simple projection onto the first p -coordinates;

$$T(y_1, \dots, y_p, y_{p+1}, \dots, y_n) = (y_1, \dots, y_p, 0, \dots, 0)$$

Later in this course we'll study other problems where different types of coordinate change are allowed. When there is less freedom to modify domain and codomain coordinates it turns out the *canonical forms* of the object are greater in variety and structure. Just to jump ahead a bit, if we force $m = n$ and change coordinates in domain and codomain simultaneously¹⁶ then the **real Jordan form** captures a representative of each equivalence class of matrix up to a **similarity transformation**. On the other hand, Sylvester's Law of Inertia reveals the canonical form for the matrix of a quadratic form is simply a diagonal matrix with $\text{Diag}(D) = (-1, \dots, -1, 1, \dots, 1, 0, \dots, 0)$.

¹⁵some authors just write T , myself included, but, technically $\bar{T} = T \circ \Phi_{\bar{\beta}}^{-1}$, so... as I'm being pretty careful otherwise, it would be bad form to write the prettier, but wrong, T

¹⁶alternatively, we could study the rational canonical form, but I'm leaving that for your next course in linear algebra, it is discussed in Insel Spence and Friedberg

Quadratic forms are non-linear functions which happen to have an associated matrix. The coordinate change for the matrix of a quadratic form is quite different than what we've studied thus far. In any event, this is just a foreshadowing comment, we will return to this discussion once we study eigenvectors and quadratic forms later in this course.

I should mention, the straightening theorem is somewhat similar to the **Singular Value Decomposition (SVD) Theorem** which we may cover towards the conclusion of this course. Essentially, the Singular Value Decomposition says any linear transformation can be understood as a combination of a generalized rotation and scaling. You might find this You Tube video by Professor Pavel Grinfeld useful for gaining a quick appreciation of the SVD.

Please notice I have restated the rank-nullity theorem within the theorem below. We already gave a matrix-theoretic proof when it was originally stated in Theorem 3.3.18. In contrast, the proof given below is purely linear algebraic.

Theorem 3.5.1. *Straightening Theorem*

Let V, W be vector spaces of finite dimension over $\mathbb{F}$. In particular, suppose $\dim(V) = n$ and $\dim(W) = m$. If $T : V \rightarrow W$ be a linear transformation then

$$\dim(V) = \dim(\text{Ker}(T)) + \dim(\text{Range}(T)).$$

If $\text{rank}(T) = \dim(T(V)) = p$ then there exist bases β for V and γ for W such that:

$$[T]_{\beta, \gamma} = \left[\begin{array}{c|c} I_p & 0 \\ \hline 0 & 0 \end{array} \right].$$

Proof: Note $\text{Ker}(T) \leq V$ therefore we may select a basis $\beta_K = \{v_1, \dots, v_k\}$ for $\text{Ker}(T)$ where $\dim(\text{Ker}(T)) = k$. Apply the basis extension theorem to extend β_K to a basis for V as follows:

$$\beta = \{w_1, \dots, w_p, v_1, \dots, v_k\}$$

where $p + k = n = \dim(V)$. Notice $w_1, \dots, w_p \notin \text{span}(\beta_K)$ since otherwise β is not LI and hence not a basis. If $x \in V$ then there exist $x_i, y_j \in \mathbb{F}$ for which $x = \sum_{i=1}^p x_i w_i + \sum_{j=1}^k y_j v_j$. Calculate by linearity of T ,

$$T(x) = \sum_{i=1}^p x_i T(w_i) + \sum_{j=1}^k y_j T(v_j) = \sum_{i=1}^p x_i T(w_i)$$

since $v_1, \dots, v_k \in \text{Ker}(T)$ gives $T(v_1) = \dots = T(v_k) = 0$. Observe, it follows that the set of p vectors $\gamma' = \{T(w_1), \dots, T(w_p)\}$ serves as a spanning set for $\text{Range}(T)$. Moreover, we may argue that γ' is a LI set: suppose

$$c_1 T(w_1) + \dots + c_p T(w_p) = 0 \Rightarrow T(c_1 w_1 + \dots + c_p w_p) = 0$$

If $c_1 w_1 + \dots + c_p w_p \neq 0$ then $c_1 w_1 + \dots + c_p w_p \in \text{Ker}(T) = \text{span}(\beta_K)$ hence β is a linearly dependent set. This contradicts the LI of β hence $c_1 w_1 + \dots + c_p w_p = 0$. But, then as $\{w_1, \dots, w_p\}$ is a subset of the LI set β we find $c_1 = 0, \dots, c_p = 0$. Consequently, γ' serves as a basis for

$\text{Range}(T)$. Hence, $\dim(\text{Ker}(T)) + \dim(\text{Range}(T)) = k + p = n$. Finally, extend γ' to a basis $\gamma = \{T(w_1), \dots, T(w_p), u_{p+1}, \dots, u_m\}$ for W . Then,

$$[T]_{\beta, \gamma} = [[T(w_1)]_{\gamma}] \cdots [[T(w_p)]_{\gamma}] [T(v_1)]_{\gamma} \cdots [T(v_k)]_{\gamma} = [e_1] \cdots [e_p] [0] \cdots [0].$$

This concludes the proof of the theorem. $\square$

Corollary 3.5.2. *Matrix Congruence*

If $A, B \in \mathbb{F}^{m \times n}$ then there exists P, Q invertible matrices such that $B = PAQ$ if and only if $\text{rank}(A) = \text{rank}(B)$. In other words, two same sized matrices are matrix congruent if and only if they share the same rank.

Proof: I think I will leave the forward implication as a homework and focus on the converse here. Suppose $\text{rank}(A) = \text{rank}(B) = p$ then by the straightening theorem there exist bases β, β' for $\mathbb{F}^n$ and γ, γ' for $\mathbb{F}^m$ for which the matrices of L_A and L_B with respect to the basis are given by:

$$[L_A]_{\beta, \gamma} = \left[\begin{array}{c|c} I_p & 0 \\ \hline 0 & 0 \end{array} \right] = [L_B]_{\beta', \gamma'}$$

Since $[L_A] = A$ and $[L_B] = B$, applying Proposition 3.4.7 yields

$$[L_A]_{\beta, \gamma} = [\gamma]^{-1} A [\beta] \quad \& \quad [L_B]_{\beta', \gamma'} = [\gamma']^{-1} B [\beta']$$

Therefore, $[\gamma]^{-1} A [\beta] = [\gamma']^{-1} B [\beta']$. Solve for B to obtain $B = [\gamma'] [\gamma]^{-1} A [\beta] [\beta']^{-1}$. Let $P = [\gamma'] [\gamma]^{-1}$ and $Q = [\beta] [\beta']^{-1}$ and note both P and Q are invertible as the matrix of a basis is invertible and the product of invertible matrices is invertible. Thus $B = PAQ$ and we have shown A and B are congruent matrices. $\square$

There may be easier ways to prove the result above, my intention was to illustrate how it appears as a result flowing from the straightening theorem. The result which follows is extremely useful when it can be used.

Theorem 3.5.3.

Let V be vector spaces of finite dimension over $\mathbb{F}$. If $T : V \rightarrow V$ is a linear transformation then the following are equivalent:

- (1.) T is injective
- (2.) T is surjective
- (3.) T is an isomorphism

Proof: follows nicely from the rank nullity theorem. If $\dim(V) = n$ and $T : V \rightarrow V$ is a linear transformation then $n = \text{rank}(T) + \text{nullity}(T)$.

Suppose (1.) is true; suppose T is injective then $\text{Ker}(T) = 0$ thus $\text{nullity}(T) = 0$ and it follows $\text{rank}(T) = n$ thus T is surjective. Thus (1.) implies (2.).

Suppose (2.) is true; assume T is surjective. Then $T(V) = V$ hence $\text{rank}(T) = n$ and we find $\text{nullity}(T) = 0$ thus $\text{Ker}(T) = 0$ and hence T is injective. Thus T is an isomorphism since it is a linear transformation which is both injective and surjective. Thus (2.) implies (3.).

But then (3.) implies (1.) by the definition of isomorphism. It then follows that (1.) and (2.) and (3.) are equivalent statements.¹⁷ $\square$

Notice, we can only use this Theorem to circumvent work if we already know the function in question is indeed a linear transformation on a given vector space of finite dimension. But, if that is settled, this Theorem gives us license to claim injectivity and surjectivity after verification of either one. This should remind you of the case of maps on finite sets; if $T : S \rightarrow S$ is a function and $\#S < \infty$ then T is surjective iff T is injective. While $V(\mathbb{F})$ generally has infinitely many vectors, the basis construction brings a finiteness which is a large part of why finite-dimensional linear algebra has such simple structure.

3.5.1 a detour into matrix theory

As we noted before, the straightening theorem asserts: *there exists a choice of coordinates which makes a given linear transformation a projection onto the range*. However, the proof of the theorem did not entirely explain how to find such coordinates. We next investigate a calculational method to find β, γ for which the theorem is realized¹⁸.

Suppose $T \in L(V, W)$ where $\dim(V) = n$ and $\dim(W) = m$. Furthermore, suppose $\beta' = \{v'_1, \dots, v'_n\}$ and $\gamma' = \{w'_1, \dots, w'_m\}$ are bases for V and W respective. We define $[T]_{\beta'\gamma'}$ as usual:

$$[T]_{\beta'\gamma'} = [[T(v'_1)]_{\gamma'} | \dots | [T(v'_n)]_{\gamma'}]$$

There exists a product of elementary $m \times m$ matrices E_1 for which

$$R_1 = \text{rref}([T]_{\beta'\gamma'}) = E_1[T]_{\beta'\gamma'}$$

Let p be the number of pivot columns in R_1 . Observe that the last $(m - p)$ rows in R_1 are zero. Therefore, the last $(m - p)$ columns in R_1^T are zero. Gauss-Jordan elimination on R_1^T is accomplished by multiplication by E_2 which is formed from a product of $n \times n$ elementary matrices.

$$R_2 = \text{rref}(R_1^T) = E_2 R_1^T$$

Notice that the trivial rightmost $(m - p)$ columns stay trivial under the Gauss-Jordan elimination. Moreover, the nonzero pivot rows in R_1 become p -pivot columns in R_1^T which reduce to $e_1, \dots, e_p$ standard basis vectors in $\mathbb{R}^n$ for R_2 (the leading ones are moved to the top rows with row-swaps if necessary). In total, we find: (the subscripts indicate the size of the blocks)

$$E_2 R_1^T = [e_1 | \dots | e_p | 0 | \dots | 0] = \left[\begin{array}{c|c} I_p & 0_{p \times (m-p)} \\ \hline 0_{(n-p) \times p} & 0_{(n-p) \times (m-p)} \end{array} \right]$$

¹⁷please ask me if this is unclear, this is a common proof technique to establish equivalence of multiple statements.

¹⁸my apologies, the γ' in the discussion which follows is logically divorced from our previous use of γ' in the proof of the straightening theorem

Therefore,

$$E_2(E_1[T]_{\beta'\gamma'})^T = \left[\begin{array}{c|c} I_p & 0_{p \times (m-p)} \\ \hline 0_{(n-p) \times p} & 0_{(n-p) \times (m-p)} \end{array} \right]$$

Transposition of the above equation yields the following:

$$E_1[T]_{\beta'\gamma'} E_2^T = \left[\begin{array}{c|c} I_p & 0_{p \times (n-p)} \\ \hline 0_{(m-p) \times p} & 0_{(m-p) \times (n-p)} \end{array} \right]$$

If β, γ are bases for V and W respective then we relate the matrix $[T]_{\beta, \gamma}$ to $[T]_{\beta', \gamma'}$ as follows:

$$[T]_{\beta, \gamma} = [\Phi_{\beta'} \circ \Phi_{\beta}^{-1}][T]_{\beta', \gamma'}[\Phi_{\gamma} \circ \Phi_{\gamma'}^{-1}].$$

Therefore, we ought to define β by imposing $[\Phi_{\beta'} \circ \Phi_{\beta}^{-1}] = E_1$ and γ by $[\Phi_{\gamma} \circ \Phi_{\gamma'}^{-1}] = E_2^T$. Using $L_A(v) = Av$ notation for E_1, E_2^T ,

$$L_{E_1} = \Phi_{\beta'} \circ \Phi_{\beta}^{-1} \quad \& \quad L_{E_2^T} = \Phi_{\gamma} \circ \Phi_{\gamma'}^{-1}$$

Thus,

$$\Phi_{\beta}^{-1} = \Phi_{\beta'}^{-1} \circ L_{E_1} \quad \& \quad \Phi_{\gamma}^{-1} = \left(L_{E_2^T} \circ \Phi_{\gamma'} \right)^{-1} = \Phi_{\gamma'}^{-1} \circ L_{E_2^T}^{-1}.$$

and we construct β and γ explicitly by:

$$\beta = \{(\Phi_{\beta'}^{-1} \circ L_{E_1})(e_j)\}_{j=1}^n \quad \gamma = \{(\Phi_{\gamma'}^{-1} \circ L_{E_2^T}^{-1})(e_j)\}_{j=1}^m.$$

Note the formulas above merely use the elementary matrices and the given pair of bases. The discussion of this page shows that β and γ so constructed will give $[T]_{\beta, \gamma} = \left[\begin{array}{c|c} I_p & 0 \\ \hline 0 & 0 \end{array} \right]$.

Continuing, to implement the calculation outlined in the previous page we would like an efficient method to calculate E_1 and E_2 . We can do this much as we did for computation of the inverse. I illustrate the idea below¹⁹:

Example 3.5.4. Let $A = \begin{bmatrix} 1 & 3 & 4 \\ 1 & 4 & 5 \\ 1 & 0 & 1 \\ 1 & 2 & 3 \end{bmatrix}$. If we adjoin the identity matrix to right the matrix which

is constructed in the Gauss-Jordan elimination is the product of elementary matrices P for which $\text{rref}(A) = PA$.

$$\text{rref}[A|I_4] = \text{rref} \left[\begin{array}{ccc|cccc} 1 & 3 & 4 & 1 & 0 & 0 & 0 \\ 1 & 4 & 5 & 0 & 1 & 0 & 0 \\ 1 & 0 & 1 & 0 & 0 & 1 & 0 \\ 1 & 2 & 3 & 0 & 0 & 0 & 1 \end{array} \right] = \left[\begin{array}{ccc|cccc} 1 & 0 & 1 & 0 & 0 & 1 & 0 \\ 0 & 1 & 1 & 0 & 0 & -1/2 & 1/2 \\ 0 & 0 & 0 & 1 & 0 & 1/2 & -3/2 \\ 0 & 0 & 0 & 0 & 1 & 1 & -2 \end{array} \right]$$

We can read P for which $\text{rref}(A) = PA$ from the result above, it is simply

$$P = \begin{bmatrix} 0 & 0 & 1 & 0 \\ 0 & 0 & -1/2 & 1/2 \\ 1 & 0 & 1/2 & -3/2 \\ 0 & 1 & 1 & -2 \end{bmatrix}.$$

¹⁹see Example 2.7 on page 244 of Hefferon's Linear Algebra for a slightly different take built on explicit computation of the product of the elementary matrices needed for the reduction

Next, consider row reduction on the transpose of the reduced matrix. This corresponds to column operations on the reduced matrix.

$$\text{rref}[(\text{rref}(A))^T | I_3] = \text{rref} \left[\begin{array}{cccc|ccc} 1 & 0 & 0 & 0 & 1 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 & 1 & 0 \\ 1 & 1 & 0 & 0 & 0 & 0 & 1 \end{array} \right] = \left[\begin{array}{cccc|ccc} 1 & 0 & 0 & 0 & 0 & -1 & 1 \\ 0 & 1 & 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 1 & 1 & -1 \end{array} \right]$$

Let $Q = \begin{bmatrix} 0 & -1 & 1 \\ 0 & 1 & 0 \\ 1 & 1 & -1 \end{bmatrix}$ and define R by:

$$R^T = Q[\text{rref}(A)]^T = \begin{bmatrix} 0 & -1 & 1 \\ 0 & 1 & 0 \\ 1 & 1 & -1 \end{bmatrix} \begin{bmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 1 & 1 & 0 & 0 \end{bmatrix} = \begin{bmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 \end{bmatrix}$$

Finally, $R = (Q[\text{rref}(A)]^T)^T = \text{rref}(A)Q^T$ hence $R = PAQ^T$. In total,

$$\left[\begin{array}{cc|cc} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ \hline 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \end{array} \right] = \begin{bmatrix} 0 & 0 & 1 & 0 \\ 0 & 0 & -1/2 & 1/2 \\ 1 & 0 & 1/2 & -3/2 \\ 0 & 1 & 1 & -2 \end{bmatrix} \begin{bmatrix} 1 & 3 & 4 \\ 1 & 4 & 5 \\ 1 & 0 & 1 \\ 1 & 2 & 3 \end{bmatrix} \begin{bmatrix} 0 & 0 & 1 \\ -1 & 1 & 1 \\ 1 & 0 & -1 \end{bmatrix}$$

There is nothing terribly special about this example. We could follow the same procedure for a general matrix to find the explicit change of basis matrices which show the matrix congruence of A to $\begin{bmatrix} I_p & 0 \\ 0 & 0 \end{bmatrix}$ where $p = \text{rank}(A)$.

3.6 structure of subspaces

The first two subsections in this section flow together as a coherent narrative. The third subsection combines the material of this section with our previous work on quotient spaces.

3.6.1 independent subspaces

I will begin this section by following an elegant construction²⁰ I found in Morton L. Curtis' *Abstract Linear Algebra* pages 28-30. The results we encounter in this section prove useful in later chapters where we study eigenvectors.

Recall the construction in Example 2.1.11, this is known as the **external direct sum**. If V, W are vector spaces over $\mathbb{R}$ then $V \times W$ is given the following vector space structure:

$$(v_1, w_1) + (v_2, w_2) = (v_1 + v_2, w_1 + w_2) \quad \& \quad c(v, w) = (cv, cw).$$

In the vector space $V \times W$ the vector $(0_V, 0_W) = 0_{V \times W}$. Although, usually we just write $(0, 0) = 0$. Furthermore, if $\beta_V = \{v_1, \dots, v_n\}$ and $\beta_W = \{w_1, \dots, w_m\}$ then a basis for $V \times W$ is simply:

$$\beta = \{(v_i, 0) | i \in \mathbb{N}_n\} \cup \{(0, w_j) | j \in \mathbb{N}_m\}$$

²⁰I don't use his notation that $A \oplus B = A \times B$, I reserve $A \oplus B$ to denote internal direct sums.

I invite the reader to check LI of β . To see how β spans, please consider the calculation below:

$$\begin{aligned}(x, y) &= (x, 0) + (0, y) \\ &= (x_1v_1 + \cdots + x_nv_n, 0) + (0, y_1w_1 + \cdots + y_mw_m) \\ &= x_1(v_1, 0) + \cdots + x_n(v_n, 0) + y_1(0, w_1) + \cdots + y_m(0, w_m)\end{aligned}$$

Thus β is a basis for $V \times W$ and we can count $\#(\beta) = n+m$ hence $\dim(V \times W) = \dim(V) + \dim(W)$. This result generalizes to an s -fold cartesian product of vector spaces over $\mathbb{F}$:

Proposition 3.6.1.

If $W_1, W_2, \dots, W_s$ are vector spaces over $\mathbb{F}$ with bases $\beta_1, \beta_2, \dots, \beta_s$ respective then $W_1 \times W_2 \times \cdots \times W_s$ has basis

$$(\beta_1 \times \{0\} \times \cdots \times \{0\}) \cup (\{0\} \times \beta_2 \times \cdots \times \{0\}) \cup \cdots \cup (\{0\} \times \{0\} \times \cdots \times \beta_s)$$

hence $\dim(W_1 \times W_2 \times \cdots \times W_s) = \dim(W_1) + \dim(W_2) + \cdots + \dim(W_s)$.

Proof: left to reader. $\square$

The Example below illustrates the claim of the Proposition above:

Example 3.6.2. Find a basis for $P_2(\mathbb{R}) \times \mathbb{R}^{2 \times 2} \times \mathbb{C}$. Recall the monomial basis $\{1, x, x^2\}$ for $P_2(\mathbb{R})$ and the unit-matrix basis $\{E_{11}, E_{12}, E_{21}, E_{22}\}$ for $\mathbb{R}^{2 \times 2}$ and $\{1, i\}$ serves as the basis for $\mathbb{C}$ as a real vector space. Hence

$$\beta = \{(1, 0, 0), (x, 0, 0), (x^2, 0, 0), (0, E_{11}, 0), (0, E_{12}, 0), (0, E_{21}, 0), (0, E_{22}, 0), (0, 0, 1), (0, 0, i)\}$$

serves as a basis for the nine dimensional real vector space $P_2(\mathbb{R}) \times \mathbb{R}^{2 \times 2} \times \mathbb{C}$.

The question to ask is when is it possible to find an isomorphism between a given vector space V and some set of subspaces of V whose sum forms V . It turns out there are several ways to understand such a structure and we devote the next page or so of the notes towards exploring a number of equivalent characterizations.

Notice the notation $W_1 + W_2 = \{x_1 + x_2 \mid x_1 \in W_1, x_2 \in W_2\}$ generalizes to:

$$W_1 + W_2 + \cdots + W_k = \{x_1 + x_2 + \cdots + x_k \mid x_i \in W_i \text{ for each } i = 1, 2, \dots, k\}$$

where $W_1, W_2, \dots, W_k$ are subspaces of some vector space.

Definition 3.6.3.

If $W_1, W_2, \dots, W_k \leq V$ and $V = W_1 + W_2 + \cdots + W_k$ then we say V is the **internal direct sum** of the subspaces $W_1, W_2, \dots, W_k$ if and only if for each $x \in V$ there exist unique $x_i \in W_i$ for $i = 1, 2, \dots, k$ such that $x = x_1 + x_2 + \cdots + x_k$. When the above criteria is met we denote this by writing

$$V = W_1 \oplus W_2 \oplus \cdots \oplus W_k$$

I follow some arguments I found in Chapter 10 of Dummit and Foote's *Abstract Algebra*. I found these are a bit easier than what I've seen in undergraduate linear texts, so, I share them here:

Theorem 3.6.4.

If $W_1, \dots, W_k \leq V$ and $V = W_1 + \dots + W_k$ is finite dimensional over $\mathbb{F}$ then the following are equivalent:

- (1.) $\pi : W_1 \times \dots \times W_k \rightarrow V$ defined by $\pi(x_1, \dots, x_k) = x_1 + \dots + x_k$ is an isomorphism,
- (2.) $W_j \cap (W_1 + \dots + W_{j-1} + W_{j+1} + \dots + W_k) = \{0\}$ for each $j = 1, \dots, k$,
- (3.) $V = W_1 \oplus \dots \oplus W_k$ or, to be precise, for each $x \in V$ there exist unique $x_i \in W_i$ for $i = 1, \dots, k$ for which $x = x_1 + \dots + x_k$,
- (4.) if β_i is a basis for W_i for $i = 1, \dots, k$ then $\beta = \beta_1 \cup \dots \cup \beta_k$ is a basis for V ,
- (5.) if $x_i \in W_i$ for $i = 1, \dots, k$ and $x_1 + \dots + x_k = 0$ then $x_1 = 0, \dots, x_k = 0$.

Proof: suppose $\pi : W_1 \times \dots \times W_k \rightarrow V$ defined by $\pi(x_1, \dots, x_k) = x_1 + \dots + x_k$ is an isomorphism. Suppose $x \in W_j \cap (W_1 + \dots + W_{j-1} + W_{j+1} + \dots + W_k)$. We have $x \in W_j$ and $x \in W_1 + \dots + W_{j-1} + W_{j+1} + \dots + W_k$ thus there exist $x_i \in W_i$ for $i \neq j$ for which

$$x = x_1 + \dots + x_{j-1} + x_{j+1} + \dots + x_k \Rightarrow x_1 + \dots + x_{j-1} - x + x_{j+1} + \dots + x_k = 0$$

Since π is an isomorphism π^{-1} is a linear and $\pi^{-1}(0) = 0$ and

$$\pi^{-1}(x_1 + \dots + x_{j-1} - x + x_{j+1} + \dots + x_k) = (x_1, \dots, x_{j-1}, -x, x_{j+1}, \dots, x_k) = 0$$

hence $x = 0$. Thus $W_j \cap (W_1 + \dots + W_{j-1} + W_{j+1} + \dots + W_k) = \{0\}$ for each $j = 1, \dots, k$. This completes the proof that (1.) implies (2.).

Next, suppose $W_j \cap (W_1 + \dots + W_{j-1} + W_{j+1} + \dots + W_k) = \{0\}$ for each $j = 1, \dots, k$. Let $x \in V$ then as $V = W_1 + \dots + W_k$ there exist $x_i \in W_i$ for which $x = x_1 + \dots + x_k$. Suppose $y_i \in W_i$ for which $x = y_1 + \dots + y_k$. Notice $x_j = x - \sum_{i \neq j} x_i$ and $y_j = x - \sum_{i \neq j} y_i$ thus

$$y_j - x_j = (x - \sum_{i \neq j} x_i) - (x - \sum_{i \neq j} y_i) = - \sum_{i \neq j} (x_i + y_i) \in W_1 + \dots + W_{j-1} + W_{j+1} + \dots + W_k$$

and as $y_j - x_j \in W_j$ we find $y_j - x_j \in W_j \cap (W_1 + \dots + W_{j-1} + W_{j+1} + \dots + W_k) = \{0\}$. Thus $y_j = x_j$ for arbitrary $j = 1, \dots, k$. This completes the proof that (2.) implies (3.).

Suppose $V = W_1 \oplus \dots \oplus W_k$. Let $\beta_i = \{v_{i,j_i} \mid 1 \leq j_i \leq m_i\}$ serve as a basis for W_i for $i = 1, \dots, k$ where we've defined $\dim(W_i) = m_i$. Let $x \in V$. There exist unique $x_i \in W_i$ for $i = 1, \dots, k$ for which $x = x_1 + \dots + x_k$. Since $W_i = \text{span}(\beta_i)$ there exist $c_{i,j_i} \in \mathbb{F}$ such that

$$x_i = \sum_{j_i=1}^{m_i} c_{i,j_i} v_{i,j_i}.$$

Therefore,

$$x = \sum_{j_1=1}^{m_1} c_{1,j_1} v_{1,j_1} + \dots + \sum_{j_k=1}^{m_k} c_{k,j_k} v_{k,j_k} = \sum_{i=1}^k \sum_{j_i=1}^{m_i} c_{i,j_i} v_{i,j_i}$$

The calculation above shows $\beta = \beta_1 \cup \cdots \cup \beta_k$ is a spanning set for V . Next we prove linear independence of β . Suppose

$$\sum_{i=1}^k \sum_{j_i=1}^{m_i} c_{i,j_i} v_{i,j_i} = 0.$$

Since $0 = 0 + \cdots + 0$ by uniqueness we find

$$\sum_{j_i=1}^{m_i} c_{i,j_i} v_{i,j_i} = 0$$

and by linear independence of β_i we find $c_{i,j_i} = 0$ for $i = 1, \dots, k$ and it follows that β is linearly independent. This completes the proof that (3.) implies (4.).

Suppose if β_i is a basis for W_i for $i = 1, \dots, k$ then $\beta = \beta_1 \cup \cdots \cup \beta_k$ is a basis for V (use the same notation as in the previous portion of the proof). Let $x_i \in W_i$ for $i = 1, \dots, k$ and suppose $x_1 + \cdots + x_k = 0$. Further, suppose $c_{i,j_i} \in \mathbb{F}$ such that

$$x_i = \sum_{j_i=1}^{m_i} c_{i,j_i} v_{i,j_i}.$$

Since $x_1 + \cdots + x_k = 0$ we find

$$\sum_{j_i=1}^{m_i} c_{i,j_i} v_{i,j_i} = 0$$

thus $c_{i,j_i} = 0$ for all i, j_i and it follows $x_i = 0$ for $i = 1, \dots, k$. This completes the proof that (4.) implies (5.).

Suppose $x_i \in W_i$ for $i = 1, \dots, k$ with $x_1 + \cdots + x_k = 0$ implies $x_1 = 0, \dots, x_k = 0$. Let $\pi : W_1 \times \cdots \times W_k \rightarrow V$ be defined by $\pi(x_1, \dots, x_k) = x_1 + \cdots + x_k$. We seek to show π is an isomorphism. We begin by establishing the linearity of π . Let $x, y \in W_1 \times \cdots \times W_k$ and $c \in \mathbb{F}$ then

$$cx + y = c(x_1, \dots, x_k) + (y_1, \dots, y_k) = (cx_1 + y_1, \dots, cx_k + y_k)$$

by the definition of the vector space structure on the Cartesian product. Thus,

$$\pi(cx + y) = (cx_1 + y_1) + \cdots + (cx_k + y_k) = c(x_1 + \cdots + x_k) + y_1 + \cdots + y_k = c\pi(x) + \pi(y).$$

Next, to show π is injective consider $x \in \text{Ker}(\pi)$,

$$\pi(x) = x_1 + \cdots + x_k = 0$$

thus $x_1 = 0, \dots, x_k = 0$ and hence $x = 0$ and we find $\text{Ker}(\pi) = 0$ thus π is injective. Finally, let $x \in V$ and recall we presuppose $V = W_1 + \cdots + W_k$ at the outset of this theorem. Thus, there exist $x_i \in W_i$ for $i = 1, \dots, k$ for which $x = x_1 + \cdots + x_k$. Observe,

$$\pi(x_1, \dots, x_k) = x_1 + \cdots + x_k = x.$$

Thus π is a surjection. In summary, π is a linear bijection and is thus an isomorphism. This completes the proof of (5.) implies (1.). Logically, the equivalence of all five statement follows. $\square$

The proof above is lengthy, but it is not difficult²¹. Any subset of it would make a totally reasonable homework or test question. The neat thing is that we are now free to consider any of these as the definition of **independent subspaces**. To be honest, I usually take (2.) paired with the data $V = W_1 + \cdots + W_k$ as the definition of $V = W_1 \oplus \cdots \oplus W_k$. As I mentioned, when I wrote these notes initially I was following Morton L. Curtis' *Abstract Linear Algebra* text, so I probably chose the definition to follow that text.

The case $k = 2$ is most interesting since condition (2.) simply reads that $W_1 \cap W_2 = \{0\}$. If $V = W_1 + W_2$ and $W_1 \cap W_2 = \{0\}$ then we say that W_1, W_2 are **complementary subspaces**.

Example 3.6.5. If $\mathcal{F}(\mathbb{R})$ is the set of functions on $\mathbb{R}$ then since we have the identity:

$$f(x) = \frac{1}{2} \left(f(x) + f(-x) \right) + \frac{1}{2} \left(f(x) - f(-x) \right)$$

for all $x \in \mathbb{R}$. For example, recall $\cosh(x) = \frac{1}{2}(e^x + e^{-x})$ and $\sinh(x) = \frac{1}{2}(e^x - e^{-x})$ hence $e^x = \cosh(x) + \sinh(x)$. We note:

$$f_{\text{even}}(x) = \frac{1}{2} (f(x) + f(-x)) \quad \& \quad f_{\text{odd}}(x) = \frac{1}{2} (f(x) - f(-x))$$

satisfy $f_{\text{even}}(-x) = f_{\text{even}}(x)$ and $f_{\text{odd}}(-x) = -f_{\text{odd}}(x)$ for each $x \in \mathbb{R}$. You can easily verify the set of even functions $\mathcal{F}_e(\mathbb{R})$ and the set of odd functions $\mathcal{F}_o(\mathbb{R})$ are subspaces of $\mathcal{F}(\mathbb{R})$. Since $f = f_{\text{even}} + f_{\text{odd}}$ for any function f we find $\mathcal{F}(\mathbb{R}) = \mathcal{F}_e(\mathbb{R}) + \mathcal{F}_o(\mathbb{R})$. Moreover, it is easy to verify that $\mathcal{F}_e(\mathbb{R}) \cap \mathcal{F}_o(\mathbb{R}) = \{0\}$. Therefore, the subspaces of even and odd functions form complementary subspaces of the space of functions on $\mathbb{R}$; $\mathcal{F}(\mathbb{R}) = \mathcal{F}_e(\mathbb{R}) \oplus \mathcal{F}_o(\mathbb{R})$.

Example 3.6.6. Let A be an $n \times n$ matrix over $\mathbb{F}$ then notice that:

$$A = \frac{1}{2}(A + A^T) + \frac{1}{2}(A - A^T)$$

it follows that $\mathbb{F}^{n \times n}$ is the direct sum of the complementary subspaces of symmetric and antisymmetric matrices. I leave the details as a homework problem. I gave you the essential hint here.

A convenient notation for spans of a single element v in V a vector space over $\mathbb{R}$ is simply $v\mathbb{R}$. I utilize this notation in the examples below.

Example 3.6.7. The cartesian plane $\mathbb{R}^2 = e_1\mathbb{R} \oplus e_2\mathbb{R}$.

Example 3.6.8. The complex numbers $\mathbb{C} = \mathbb{R} \oplus i\mathbb{R}$. We could discuss how extending $i^2 = -1$ linearly gives this an algebraic structure. We have a whole course in the major to dig into this example.

Example 3.6.9. The hyperbolic numbers $\mathcal{H} = \mathbb{R} \oplus j\mathbb{R}$. We could discuss how extending $j^2 = 1$ linearly gives this an algebraic structure. This is less known, but it naturally describes problems with some hyperbolic symmetry.

Example 3.6.10. The dual numbers $\mathcal{N} = \mathbb{R} \oplus \epsilon\mathbb{R}$. We could discuss how extending $\epsilon^2 = 0$ linearly gives this an algebraic structure.

²¹I gave an in-class presentation of a proof in my Math 321 Lectures of 3-8-17 and 3-10-17, this may or may not be an improvement on those arguments

The algebraic comments above are mostly for breadth. We focus on linear algebra²² in these notes.

Naturally we should consider extending the discussion to more than two subspaces.

Example 3.6.11. *Quaternions.* $\mathbb{H} = \mathbb{R} \oplus i\mathbb{R} \oplus j\mathbb{R} \oplus k\mathbb{R}$ where $i^2 = j^2 = k^2 = -1$ and $ij = -ji = k$ and $jk = -kj = i$ and $ki = -ik = j$. Our notation for vectors in most calculus texts has a historical basis in Hamilton's quaternions.

Unfortunately, trivial pairwise intersections do not generally suffice to give direct sum decompositions for three or more subspaces. The next example illustrates this subtlety.

Example 3.6.12. Let $W_1 = (1, 1)\mathbb{R}$ and $W_2 = (1, 0)\mathbb{R}$ and $W_3 = (1, 1)\mathbb{R}$. It is not hard to verify $W_1 + W_2 + W_3 = \mathbb{R}^2$ and $W_1 \cap W_2 = W_1 \cap W_3 = W_2 \cap W_3 = \{0\}$. However, it is certainly not possible to find an isomorphism of $\mathbb{R}^2$ and the three dimensional vector space $W_1 \times W_2 \times W_3$.

3.6.2 invariant subspaces and internal direct products

We now explain the interesting relation between the direct sum decomposition of a vector space V and the block-structure for a matrix of a linear transformation on V .

Definition 3.6.13.

If $T : V \rightarrow V$ is a linear transformation on the vector space V over $\mathbb{F}$ and $W \leq V$ is a subspace of V for which $T(W) \subseteq W$ then we say W is an **T -invariant** subspace of V . We let $T_W : W \rightarrow W$ denote the map defined by $T_W(x) = T(x)$ for each $x \in W$.

Notice that $T|_W : W \rightarrow V$ whereas $T_W : W \rightarrow W$. The map T_W is only well-defined if the subspace W is T -invariant. In particular, T -invariance of W gives us that $T(x)$ is in W , that is, the map $T_W : W \rightarrow W$ is *into* W . To show a map $f : A \rightarrow B$ is well-defined we need several things. First, we need that each element a in A produces a single output $f(a)$. Second, we need that each output $f(a)$ is actually in the codomain B .

Something very interesting happens when $V = W_1 \oplus W_2 \oplus \cdots \oplus W_s$ and each W_j is T -invariant.

Theorem 3.6.14.

Suppose $T : V \rightarrow V$ is a linear transformation and $W_j \leq V$ are such that $T(W_j) \leq W_j$ for each $j = 1, \dots, s$. Also, suppose $V = W_1 \oplus \cdots \oplus W_s$ where $\dim(W_j) = d_j$ and $d_1 + \cdots + d_s = n = \dim(V)$. Then there exists a basis $\beta = \beta_1 \cup \beta_2 \cup \cdots \cup \beta_s$ for V formed by concatenating β_j basis for W_j for $j = 1, 2, \dots, s$ for which

$$[T]_{\beta, \beta} = \left[\begin{array}{c|c|c|c} M_1 & 0 & \cdots & 0 \\ \hline 0 & M_2 & \cdots & 0 \\ \hline \vdots & \vdots & \cdots & \vdots \\ \hline 0 & 0 & \cdots & M_s \end{array} \right] = \text{diag}(M_1, M_2, \dots, M_s).$$

and $M_j = [T_{W_j}]_{\beta_j, \beta_j}$ for $j = 1, 2, \dots, s$.

²²a vector space paired with a multiplication is called an algebra. The rules $i^2 = -1$, $j^2 = 1$ and $\epsilon^2 = 0$ all serve to define non-isomorphic algebraic structures on $\mathbb{R}^2$. These are isomorphic as vector spaces. I'll discuss the concept of an algebra further in the next chapter.

Proof: let β_j be a basis for W_j then $\#(\beta_j) = d_j$ for $j = 1, 2, \dots, s$. Let us denote $\beta_j = \{v_{j,1}, v_{j,2}, \dots, v_{j,d_j}\}$ for $j = 1, 2, \dots, s$. Furthermore, by 3.6.4 we have that $\beta = \beta_1 \cup \beta_2 \cup \dots \cup \beta_s$ is a basis for V . If $v_{j,i} \in \beta_j$ then

$$T(v_{j,i}) \in W_j \Rightarrow T(v_{j,i}) = c_1 v_{j,1} + \dots + c_{d_j} v_{j,d_j}$$

thus the column vector in $[T]_{\beta,\beta}$ corresponding to the basis vector $v_{j,i}$ only is nonzero in rows corresponding to the β_j part of the basis. It follows that $[T]_{\beta,\beta}$ is block-diagonal where M_j is the $d_j \times d_j$ matrix over $\mathbb{F}$ which is the matrix of T_{W_j} with respect to the β_j matrix; that is $M_j = [T_{W_j}]_{\beta_j,\beta_j}$. $\square$

A block diagonal matrix allows multiplication where the blocks behave as if they were numbers. See Section 1.4.6 where we studied how block-multiplication works. I should mention, we can add matrices of different sizes following the pattern above:

Definition 3.6.15.

If $M_j \in \mathbb{F}^{d_j \times d_j}$ for $j = 1, 2, \dots, s$ then we define $M_1 \oplus M_2 \oplus \dots \oplus M_s \in \mathbb{F}^{n \times n}$ where $n = d_1 + d_2 + \dots + d_s$ and

$$M_1 \oplus M_2 \oplus \dots \oplus M_s = \left[\begin{array}{c|c|c|c} M_1 & 0 & \dots & 0 \\ \hline 0 & M_2 & \dots & 0 \\ \hline \vdots & \vdots & \dots & \vdots \\ \hline 0 & 0 & \dots & M_s \end{array} \right].$$

Certainly this is not standard addition as $A \oplus B \neq B \oplus A$. But, it is a fun new way to make new matrices from old. To be clear, $A \oplus B$ is the **direct sum** of A and B . With this language, Theorem 3.6.14 is formulated as follows: when V is a direct sum decomposition of T -invariant subspaces then there exists a matrix for which the matrix of T is likewise formed by a direct sum of submatrices. If $V = W_1 \oplus \dots \oplus W_s$ and T is W_j -invariant for each $j = 1, \dots, s$ then if β_j is basis for W_j and $\beta = \beta_1 \cup \dots \cup \beta_s$ then

$$[T]_{\beta,\beta} = [T_{W_1}]_{\beta_1,\beta_1} \oplus [T_{W_2}]_{\beta_2,\beta_2} \oplus \dots \oplus [T_{W_s}]_{\beta_s,\beta_s}.$$

Let me conclude with the pair of examples which I began our in-class discussion on 3-8-17.

Example 3.6.16. Let $T : P_3(\mathbb{R}) \rightarrow P_3(\mathbb{R})$ be defined as $T = d/dx$. In particular $T(ax^3 + bx^2 + cx + d) = 3ax^2 + 2bx + c$. Thus, for the basis $\beta = \{1, x, x^2, x^3\}$ we find matrix:

$$[T]_{\beta,\beta} = [[T(1)]_{\beta}] [T(x)]_{\beta} [T(x^2)]_{\beta} [T(x^3)]_{\beta} = [[0]_{\beta}] [1]_{\beta} [2x]_{\beta} [3x^2]_{\beta} = \left[\begin{array}{ccc|c} 0 & 1 & 0 & 0 \\ \hline 0 & 0 & 2 & 0 \\ 0 & 0 & 0 & 3 \\ \hline 0 & 0 & 0 & 0 \end{array} \right]$$

Notice, $\text{Range}(T) = \text{span}\{1, x, x^2\}$ is an invariant subspace of T however the remaining x^3 vector has $T(x^3) = 3x^2$ so $\text{Range}(T) + \text{span}(x^3)$ is not a direct sum decomposition.

Example 3.6.17. Let $S = d^2/dx^2$ on $P_3(\mathbb{R})$. Calculate,

$$S(1) = 0, \quad S(x) = 0, \quad S(x^2) = 2, \quad S(x^3) = 6x$$

Thus (by inspection) we find invariant subspaces $W_1 = \text{span}\{1, x^2\}$ and $W_2 = \text{span}\{x, x^3\}$. Let $\gamma_1 = \{1, x^2\}$ and $\gamma_2 = \{x, x^3\}$ and $\gamma = \gamma_1 \cup \gamma_2 = \{1, x^2, x, x^3\}$ and we find

$$[S]_{\gamma, \gamma} = [[S(1)]_{\gamma} | [S(x^2)]_{\gamma} | [S(x)]_{\gamma} | [S(x^3)]_{\gamma}] = [[0]_{\gamma} | [2]_{\gamma} | [0]_{\gamma} | [6x]_{\gamma}] = \left[\begin{array}{cc|cc} 0 & 2 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ \hline 0 & 0 & 0 & 6 \\ 0 & 0 & 0 & 0 \end{array} \right]$$

We have $[S]_{\gamma, \gamma} = [S_{W_1}]_{\gamma_1, \gamma_1} \oplus [S_{W_2}]_{\gamma_2, \gamma_2} = \left[\begin{array}{cc} 0 & 2 \\ 0 & 0 \end{array} \right] \oplus \left[\begin{array}{cc} 0 & 6 \\ 0 & 0 \end{array} \right]$

There is a bit more to say here, but it requires a concept we introduce towards the final story arc of this course. See §3.7.2 for the interplay between independent subspaces and quotients.

3.7 quotient vector space

Let us begin with a discussion of how to add sets of vectors. If $S, T \subseteq V$ a vector space over $\mathbb{F}$ then we define $S + T$ as follows:

$$S + T = \{s + t \mid s \in S, t \in T\}$$

In the particular case $S = \{x\}$ it is customary to write

$$x + T = \{x + t \mid t \in T\}$$

we drop the $\{ \}$ around x in this special case.

Definition 3.7.1. Coset

Let V be a vector space with $x \in V$ and $W \leq V$ then $x + W$ is a **coset** of W .

Example 3.7.2. If $W = \text{span}\{(1, 0)\}$ is the x -axis in $\mathbb{R}^2$ then $(a, b) + W$ is the horizontal line given by equation $y = b$. You can easily see $(0, b) + W = (a, b) + W$. Each coset is obtained by translating the x -axis to a parallel horizontal line.

Example 3.7.3. If $W = \{A \in \mathbb{R}^{n \times n} \mid A^T = A\}$ then $X + W$ is a coset in the square matrices. Geometrically, it is a linear manifold of the same dimension as W . I can't picture this one directly.

Example 3.7.4. Consider $A \in \mathbb{F}^{m \times n}$ and $b \in \mathbb{F}^m$. Let S be the solution set of $Ax = b$. Recall we may express $x \in S$ as $x = c_1 v_1 + \dots + c_\nu v_\nu + x_p$ where $Av_i = 0$ for $i = 1, \dots, \nu$ where $\nu = \text{nullity}(A)$. In other words, $x = x_h + x_p$ where $Ax_h = 0$ and $Ax_p = b$. We find the solution set is a coset of the nullspace of A ; $S = x_p + \text{Null}(A)$. The solution set to a nonhomogeneous system of equations is not a subspace with respect to the standard addition of column vectors. However, the solution set is a parallel translate, or a **affine space** or **linear manifold** of dimension of the nullity of the coefficient matrix.

Example 3.7.5. Let L be a linear differential operator and g a given function then $L[y] = g$ defines a **differential equation**. If V is the set of all smooth functions then we can define the **homogeneous solution set** to be $W = \{y \in V \mid L[y] = 0\} = \text{Ker}(L)$. If L is an n -th order differential operator then we learn in the differential equations course that the general solution set of $L[y] = g$ can be expressed as

$$y = c_1 y_1 + \dots + c_n y_n + y_p$$

where $y_1, \dots, y_n \in W$ are linearly independent so-called **homogeneous solutions** of the differential equation and y_p is so-called **particular solution** for which $L[y_p] = g$. In the language of cosets, we see y is a general solution implies $y \in y_p + W$ where $L[y_p] = g$.

Quotient space of V by W is the set of all such **cosets** of W . We now work towards motivating the definition of quotient space. In particular, we need to show how it has a natural vector space structure induced from V .

Proposition 3.7.6.

Let V be vector space over $\mathbb{F}$ and $W \leq V$. Then $x + W = y + W$ iff $x - y \in W$.

Proof: Suppose $x + W = y + W$. If $p \in x + W$ then it follows there exists $w_1 \in W$ for which $p = x + w_1$. However, as $x + W \subseteq y + W$ we find $x + w_1 \in y + W$ and thus there exists $w_2 \in W$ for which $x + w_1 = y + w_2$. Therefore, $y - x = w_1 - w_2 \in W$ as W is a subspace of V .

Conversely, suppose $x, y \in V$ and $x - y \in W$. Thus, there exists $w \in W$ for which $x - y = w$ and so for future reference $x = y + w$ or $y = x - w$. Let $p \in x + W$ hence there exists $w_1 \in W$ for which $p = x + w_1$. Furthermore, as W is a subspace we know $w, w_1 \in W$ implies $w + w_1 \in W$. Consider then, $p = x + w_1 = y + w + w_1 \in y + W$. Therefore, $x + W \subseteq y + W$. A similar argument shows $y + W \subseteq x + W$ hence $x + W = y + W$. $\square$

Proposition 3.7.7.

Let V be vector space over $\mathbb{F}$ and $W \leq V$. Then $x + W = W$ iff $x \in W$.

Proof: if $x + W = W$ then $x + w \in W$ for some w hence $x + w = w_1$. But, it follows $x = w_1 - w$ which makes clear that $x \in W$ as $W \leq V$.

Conversely, if $x \in W$ then consider $p = x + w_1 \in x + W$ and note $x + w_1 \in W$ hence $p \in W$ and we find $x + W \subseteq W$. Likewise, if $w \in W$ then note $w = x + w - x$ and $w - x \in W$ thus $w \in x + W$ and we find $W \subseteq x + W$. Therefore, $x + W = W$. $\square$

Observe that Proposition 3.7.6 can be reformulated to say $x + W$ is the same as $y + W$ if $y = x + w$ for some $w \in W$. We say that x and y are coset representatives of the same coset iff $x + W = y + W$. Suppose $x_1 + W = x_2 + W$ and $y_1 + W = y_2 + W$; that is, suppose x_1, x_2 are representatives of the same coset and suppose y_1, y_2 are representatives of the same coset.

Remark 3.7.8. *a contrast between cosets and subspaces*

Notice the difference between a coset and subspace. Two points in a subspace when added are once more in the subspace. However, the sum of two points in a coset need not be in the coset. Perhaps it is an interesting homework question to consider; if $x, y \in z + W$ then when is $x + y \in z + W$? In view of the above proposition, it should be an easy question.

Proposition 3.7.9.

Let V be vector space over $\mathbb{F}$ and $W \leq V$. If $x_1 + W = x_2 + W$ and $y_1 + W = y_2 + W$ and $c \in \mathbb{F}$ then $x_1 + y_1 + W = x_2 + y_2 + W$ and $cx_1 + W = cx_2 + W$.

Proof: Suppose $x_1 + W = x_2 + W$ and $y_1 + W = y_2 + W$ then by Proposition 3.7.6 we find $x_2 - x_1 = w_x$ and $y_2 - y_1 = w_y$ for some $w_x, w_y \in W$. Consider

$$(x_2 + y_2) - (x_1 + y_1) = x_2 - x_1 + y_2 - y_1 = w_x + w_y.$$

However, $w_x, w_y \in W$ implies $w_x + w_y \in W$ hence by Proposition 3.7.6 we find $x_1 + y_1 + W = x_2 + y_2 + W$. I leave proof that $cx_1 + W = cx_2 + W$ as an exercise to the reader. $\square$

The preceding triple of propositions serves to show that the definitions given below are independent of the choice of coset representative. That is, while a particular coset representative is used to make the definition, the choice is immaterial to the outcome.

Definition 3.7.10.

We define V/W to be the **quotient space of V by W** . In particular, we define:

$$V/W = \{x + W \mid x \in V\}$$

and for all $x + W, y + W \in V/W$ and $c \in \mathbb{F}$ we define:

$$(x + W) + (y + W) = x + y + W \quad \& \quad c(x + W) = cx + W.$$

Note, we have argued thus far that addition and scalar multiplication defined on V/W are well-defined functions. Let us complete the thought:

Theorem 3.7.11.

If $W \leq V$ a vector space over $\mathbb{F}$ then V/W is a vector space over $\mathbb{F}$.

Proof: if $x + W, y + W \in V/W$ note $(x + W) + (y + W)$ and $c(x + W)$ are single elements of V/W thus Axioms 9 and 10 of Definition 2.1.1 are true. Axiom 1: by commutativity of addition in V we obtain commutativity in V/W :

$$(x + W) + (y + W) = x + y + W = y + x + W = (y + W) + (x + W).$$

Axiom 2: associativity of addition follows from associativity of V ,

$$\begin{aligned} (x + W) + [(y + W) + (z + W)] &= x + W + [(y + z) + W] && \text{defn. of } + \text{ in } V/W \\ &= x + (y + z) + W && \text{defn. of } + \text{ in } V/W \\ &= (x + y) + z + W && \text{associativity of } + \text{ in } V \\ &= [(x + y) + W] + (z + W) && \text{defn. of } + \text{ in } V/W \\ &= [(x + W) + (y + W)] + (z + W) && \text{defn. of } + \text{ in } V/W. \end{aligned}$$

Axiom 3: note that $0 + W = W$ and it follows that W serves as the additive identity in the quotient:

$$(x + W) + (0 + W) = x + 0 + W = x + W.$$

Axiom 4: the additive inverse of $x + W$ is simply $-x + W$ as $(x + W) + (-x + W) = W$.

Axiom 5: observe that

$$1(x + W) = 1 \cdot x + W = x + W.$$

I leave verification of Axioms 6,7 and 8 for V/W to the reader. I hope you can see these will easily transfer of the Axioms 6,7 and 8 for V itself. $\square$

The notation $x + W$ is at times tiresome. An alternative notation is given below:

$$[x] = x + W$$

then the vector space operations on V/W are

$$[x] + [y] = [x + y] \quad \& \quad c[x] = [cx].$$

Naturally, the disadvantage of this notation is that it hides the particular subspace by which the quotient is formed. For a given vector space V many different subspaces are typically available and hence a wide variety of quotients may be constructed.

Example 3.7.12. Suppose $V = \mathbb{R}^3$ and $W = \text{span}\{(0, 0, 1)\}$. Let $[(a, b, c)] \in V/W$ note

$$[(a, b, c)] = \{(a, b, z) \mid z \in \mathbb{R}\}$$

thus a point in V/W is actually a line in V . The parameters a, b fix the choice of line so we expect V/W is a two dimensional vector space with basis $\{[(1, 0, 0)], [(0, 1, 0)]\}$.

Example 3.7.13. Suppose $V = \mathbb{R}^3$ and $W = \text{span}\{(1, 0, 0), (0, 1, 0)\}$. Let $[(a, b, c)] \in V/W$ note

$$[(a, b, c)] = \{(x, y, c) \mid x, y \in \mathbb{R}\}$$

thus a point in V/W is actually a plane in V . In this case, each plane is labeled by a single parameter c so we expect V/W is a one-dimensional vector space with basis $\{[(0, 0, 1)]\}$.

Our claims about constructing a basis for the quotient space of the preceding pair of examples are easily affirmed by applying the following proposition:

Proposition 3.7.14.

If $W \leq V$ and V is finite-dimensional then $\dim(V/W) = \dim(V) - \dim(W)$. If β is a basis for V and $\beta_W \subseteq \beta$ serves as a basis for W then if $\gamma = \beta - \beta_W$ then $\gamma + W$ serves as basis for V/W .

Proof: Suppose $\dim(V) = n$ and $\dim(W) = k$. Let $\beta = \{w_1, \dots, w_k\}$ be a basis for W . Extend β to $\gamma = \{w_1, \dots, w_k, v_1, \dots, v_{n-k}\}$ a basis for V . Observe that $w_j + W = W$ for $j = 1, \dots, k$ as $w_j \in W$ for each j . Since $O_{V/W} = W$ we certainly cannot form the basis for V/W with β . However, we can show $\{v_i + W\}_{i=1}^{n-k}$ serves as a basis for V/W . Suppose

$$c_1(v_1 + W) + c_2(v_2 + W) + \dots + c_{n-k}(v_{n-k} + W) = O_{V/W} = 0 + W$$

thus, by the definitions of coset addition and scalar multiplication,

$$(c_1v_1 + c_2v_2 + \dots + c_{n-k}v_{n-k}) + W = W$$

it follows $c_1v_1 + c_2v_2 + \dots + c_{n-k}v_{n-k} \in W$. But, this must be the zero vector since by construction the vectors $v_1, \dots, v_{n-k}$ are outside $\text{span}(\beta)$. Thus $c_1v_1 + c_2v_2 + \dots + c_{n-k}v_{n-k} = 0$ and hence by linear independence of γ we find $c_1 = c_2 = \dots = c_{n-k} = 0$. Suppose $x + W \in V/W$ then there exist $x_j, y_i \in \mathbb{F}$ for which $x = \sum_{j=1}^k x_j w_j + \sum_{i=1}^{n-k} y_i v_i$ thus

$$x + W = \sum_{j=1}^k x_j w_j + \sum_{i=1}^{n-k} y_i v_i + W = \sum_{i=1}^{n-k} y_i v_i + W = \sum_{i=1}^{n-k} y_i (v_i + W).$$

Thus, $\text{span}\{v_1 + W, \dots, v_{n-k} + W\} = V/W$. It follows $\{v_1 + W, \dots, v_{n-k} + W\}$ is a basis for V/W and we count $\dim(V/W) = n - k = \dim(V) - \dim(W)$. $\square$

Notice the proof outlines a method to derive a basis for the quotient space; first find a basis for the subspace forming the quotient. Second, extend to a basis for the total space. Then the basis for the quotient is simply the cosets represented by the extension.

Example 3.7.15. Let $V = \mathbb{R}[x]$ and let $W = \mathbb{R}$ the set of constant polynomials.

$$[a_0 + a_1x + \dots + a_nx^n] = \{c + a_1x + \dots + a_nx^n \mid c \in \mathbb{R}\}$$

Perhaps, more to the point,

$$[a_0 + a_1x + \dots + a_nx^n] = [a_1x + \dots + a_nx^n]$$

In this quotient space, we identify polynomials which differ by a constant.

We could also form quotients of $\mathcal{F}(\mathbb{R})$ or P_n or $C^\infty(\mathbb{R})$ by $\mathbb{R}$ and it would have the same meaning; if we quotient by constant functions then $[f] = [f + c]$.

The quotient space construction allows us to modify a given transformation such that its reformulation is injective. For example, consider the problem of inverting the derivative operator $D = d/dx$.

$$D(f) = f' \quad \& \quad D(f + c) = f'$$

thus D is not injective. However, if we instead look at the derivative operator on²³ a quotient space of differentiable functions of a connected domain where $[f] = [f + c]$ then defining $D([f]) = f'$ proves to be injective. Suppose $D([f]) = D([g])$ hence $f' = g'$ so $f - g = c$ and $[f] = [g]$. We generalize this example in the next subsection.

3.7.1 the first isomorphism theorem

I'll begin by defining an standard linear map attached the quotient construction:

Definition 3.7.16. Let V be a vector space with $W \leq V$.

The **quotient map** $\pi : V \rightarrow V/W$ is defined by $\pi(x) = x + W$ for each $x \in V$.

We observe π is a linear transformation.

Proposition 3.7.17.

The quotient map $\pi : V \rightarrow V/W$ is a linear transformation.

Proof: suppose $x, y \in V$ and $c \in \mathbb{F}$. Consider

$$\pi(cx + y) = (cx + y) + W = (cx + W) + (y + W) = c(x + W) + (y + W) = c\pi(x) + \pi(y). \quad \square$$

When W is formed as the kernel of a linear transformation the mapping π takes on a special significance. The π map allows us to create isomorphisms as described in the theorem below:

²³to be careful, I only modify the domain of the derivative operator here, note the output of D is not an equivalence class. Furthermore, perhaps a different symbol like $\bar{D}$ should be used to write $\bar{D}([f]) = f'$ as $D \neq \bar{D}$

Theorem 3.7.18. *First Isomorphism Theorem of Linear Algebra*

If $T : V \rightarrow U$ is a linear transformation and $W = \text{Ker}(T)$ has quotient map π then the mapping $\bar{T} : V/W \rightarrow U$ defined implicitly by $\bar{T} \circ \pi = T$ is an injection. In particular, $\bar{T}(x + \text{ker}(T)) = T(x)$ for each $x + \text{ker}(T) \in V/W$. Moreover, $V/\text{ker}(T) \cong T(V)$.

Proof: to begin we show $\bar{T}$ is well-defined. Of course, $T(x) \in U$ for each $x \in V$ hence $\bar{T}$ is **into** U . Is $\bar{T}$ single-valued? Suppose $x + \text{Ker}(T) = y + \text{Ker}(T)$ then $y - x \in \text{Ker}(T)$ hence $T(y - x) = 0$ which gives $T(x) = T(y)$. Thus, $\bar{T}(x + \text{Ker}(T)) = T(x) = T(y) = \bar{T}(y + \text{Ker}(T))$. Therefore, $\bar{T}$ is single-valued.

Next, we show $\bar{T}$ is a linear transformation. Let $x, y \in V$ and $c \in \mathbb{F}$. Consider,

$$\bar{T}(c(x + W) + (y + W)) = \bar{T}(cx + y + W) = T(cx + y) = cT(x) + T(y) = c\bar{T}(x + W) + \bar{T}(y + W).$$

We find linearity of $\bar{T}$ follows naturally from the definition of V/W as a vector space and the linearity of T .

We now turn to the question of injectivity of $\bar{T}$. Let $x + W, y + W \in V/W$ where $W = \text{ker}(T)$ and suppose $\bar{T}(x + W) = \bar{T}(y + W)$. It follows that $T(x) = T(y)$ thus $T(x - y) = 0$ and we find $x - y \in W = \text{ker}(T)$ which proves $x + \text{ker}(T) = y + \text{ker}(T)$. We have shown $\bar{T}$ is injective.

The isomorphism of $V/\text{ker}(T)$ and $T(V)$ is given by $T' : V/\text{ker}(T) \rightarrow T(V)$ where $T'(x + \text{ker}(T)) = \bar{T}(x + \text{ker}(T)) = T(x)$. If $y = T(x) \in T(V)$ then clearly $T'(x + \text{ker}(T)) = y$ hence T' is a surjection and hence an isomorphism as we have injectivity from our work on $\bar{T}$. $\square$

The last paragraph simply says that the injective map $\bar{T}$ can be made into a surjection by reducing its codomain to its range. This is not surprising. What may be surprising is how this theorem can be used to see isomorphisms in a terribly efficient manner:

Example 3.7.19. Consider $V \times W/(\{0\} \times W)$ and V . To show these are isomorphic we consider $T(v, w) = v$. It is simple to verify that $T : V \times W \rightarrow V$ is a linear surjection. Moreover, $\text{Ker}(T) = \{(0, w) \mid w \in W\} = \{0\} \times W$. The first isomorphism theorem reveals $V \times W/(\{0\} \times W) \approx V$.

Example 3.7.20. Consider $S : \mathbb{R}^{n \times n} \rightarrow \mathbb{R}^{n \times n}$ defined by $S(A) = A + A^T$. Notice that the range of $S(A)$ is simply symmetric matrices as $(S(A))^T = (A + A^T)^T = A^T + (A^T)^T = A + A^T = S(A)$. Moreover, if $A^T = -A$ then clearly $S(A) = 0$ hence S is onto the symmetric matrices. What is the kernel of S ? Suppose $S(A) = 0$ and note:

$$A + A^T = 0 \quad \Rightarrow \quad A^T = -A.$$

Thus $\text{Ker}(S)$ is the set of antisymmetric matrices. Therefore,

$$S'([A]) = A + A^T$$

is an isomorphism from $\mathbb{R}^{n \times n}/\text{Ker}(S)$ to the set of symmetric $n \times n$ matrices.

Example 3.7.21. Consider $D : P \rightarrow P$ defined by $D(f(x)) = df/dx$. Here I denote $P = \mathbb{R}[x]$, the set of all polynomials with real coefficients. Notice

$$\text{Ker}(D) = \{f(x) \in P \mid df/dx = 0\} = \{f(x) \in P \mid f(x) = c\}.$$

In this case D is already a surjection since we work with all polynomials hence:

$$\overline{D}([f(x)]) = f'(x)$$

is an isomorphism. Just to reiterate in this case:

$$\overline{D}([f(x)]) = \overline{D}([g(x)]) \Rightarrow f'(x) = g'(x) \Rightarrow f(x) = g(x) + c \Rightarrow [f(x)] = [g(x)].$$

Essentially, $\overline{D}$ is just d/dx on equivalence classes of polynomials. Notice that $\overline{D}^{-1} : P \rightarrow P/\text{Ker}(D)$ is a mapping you have already studied for several months! In particular,

$$\overline{D}^{-1}(f(x)) = \{F(x) \mid dF/dx = f(x)\}$$

Just to be safe, let's check that my formula for the inverse is correct:

$$\overline{D}^{-1}(\overline{D}([f(x)])) = \overline{D}^{-1}(df/dx) = \{F(x) \mid dF/dx = df/dx\} = \{f(x) + c \mid c \in \mathbb{R}\} = [f(x)].$$

Conversely, for $f(x) \in P$,

$$\overline{D}(\overline{D}^{-1}(f(x))) = \overline{D}(\{F(x) \mid dF/dx = f(x)\}) = f(x).$$

Perhaps if I use a different notation to discuss the preceding example then you will see what is happening: we usually call $\overline{D}^{-1}(f(x)) = \int f(x)dx$ and $\overline{D} = d/dx$ then

$$\frac{d}{dx} \int f dx = f \quad \& \quad \int \frac{d}{dx}(f + c_1) dx = f + c_2$$

In fact, if your calculus instructor was careful, then he should have told you that when we calculate the indefinite integral of a function the answer is not a function. Rather, $\int f(x) dx = \{g(x) \mid g'(x) = f(x)\}$. However, nobody wants to write a set of functions every time they integrate so we instead make the custom to write $g(x) + c$ to indicate the non-uniqueness of the answer. Antidifferentiation of f is finding a specific function F for which $F'(x) = f(x)$. Indefinite integration of f is the process of finding the set of all functions $\int f dx$ for which $\frac{d}{dx} \int f dx = f$. In any event, I hope you see that we can claim that differentiation and integration are inverse operations, however, this is in the understanding that we work on a quotient space of functions where two functions which differ by a constant are considered the same function. In that context, $f + c_1 = f + c_2$.

Example 3.7.22. Consider $D : P_2 \rightarrow P_2$ defined by

$$D(ax^2 + bx + c) = 2ax + b$$

Observe $\overline{D}([ax^2 + bx + c]) = 2ax + b$ defines a natural isomorphism from $P_2/\mathbb{R}$ to P_1 where I denote $\text{Ker}(D) = \mathbb{R}$. In other words, when I write the quotient by $\mathbb{R}$ I am identifying the set of constant polynomials with the set of real numbers.

Example 3.7.23. Consider $\mathcal{F}(\mathbb{R})$ the set of all functions on $\mathbb{R}$. Observe, any function can be written as a sum of an even and odd function:

$$f(x) = \frac{1}{2} \left(f(x) + f(-x) \right) + \frac{1}{2} \left(f(x) - f(-x) \right)$$

Furthermore, if we denote the subspaces of even and odd functions as $\mathcal{F}_{\text{even}} \leq \mathcal{F}(\mathbb{R})$ and $\mathcal{F}_{\text{odd}} \leq \mathcal{F}(\mathbb{R})$ and note $\mathcal{F}_{\text{even}} \cap \mathcal{F}_{\text{odd}} = \{0\}$ hence $\mathcal{F}(\mathbb{R}) = \mathcal{F}_{\text{even}} \oplus \mathcal{F}_{\text{odd}}$. Consider the projection $T : \mathcal{F}(\mathbb{R}) \rightarrow \mathcal{F}_{\text{even}}$ clearly $\text{Null}(T) = \mathcal{F}_{\text{odd}}$ hence by the first isomorphism theorem, $\mathcal{F}(\mathbb{R})/\mathcal{F}_{\text{odd}} \approx \mathcal{F}_{\text{even}}$.

Example 3.7.24. Continuing Example 3.7.5, this example will be most meaningful for students of differential equations, however, there is something here for everyone to learn. An n -th order linear differential equation can be written as $L[y] = g$. Here y and g are functions on a connected interval $I \subseteq \mathbb{R}$. There is an existence theorem for such problems which says that **any** solution can be written as

$$y = y_h + y_p$$

where $L[y_h] = 0$ and $L[y_p] = g$. The so-called **homogeneous solution** y_h is generally formed from a linear combination of n -LI fundamental solutions $y_1, y_2, \dots, y_n$ as

$$y_h = c_1 y_1 + c_2 y_2 + \dots + c_n y_n.$$

Here $L[y_i] = 0$ for $i = 1, 2, \dots, n$. It follows that $\text{Null}(L)$ is n -dimensional and the fundamental solution set forms a basis for this null-space. On the other hand the particular solution y_p can be formed through a technique known as **variation of parameters**. Without getting into the technical details, the point is there is an explicit, although tedious, method to calculate y_p once we know the fundamental solution set and g . Techniques for finding the fundamental solution set vary from problem to problem. For the constant coefficient case or Cauchy Euler problems it is as simple as factoring the characteristic polynomial and writing down the homogeneous solutions. Enough about that, let's think about this problem in view of quotient spaces.

The differential equation $L[y] = g$ can be instead thought of as a function which takes g as an input and produces y as an output. Of course, given the infinity of possible homogeneous solutions this would not really be a function, it's not single-valued. If we instead associate with the differential equation a function $H : V \rightarrow V/\text{Null}(L)$ where $H(g) = y + \text{Null}(L)$ then the formula can be compactly written as $H(g) = [y_p]$. For convenience, suppose $V = C^0(\mathbb{R})$ then $\text{dom}(H) = V$ as variation of parameters only requires integration of the forcing function g . Thus $H^{-1} : V/\text{Null}(L) \rightarrow V$ is an isomorphism. In short, the mathematics I outline here shows us there is a one-one correspondance between forcing functions and solutions modulo homogeneous terms. Linear differential equations have this beautiful feature; the net-response of a system L to inputs $g_1, \dots, g_k$ is nothing more than the sum of the responses to each forcing term. This is the principal of superposition which makes linear differential equations comparatively easy to understand.

3.7.2 on direct sums and quotients

Proposition 3.7.25.

If $V = A \oplus B$ then $V/A \cong B$.

Proof: Since $V \cong A \times B$ under $\eta : A \times B \rightarrow V$ with $\eta(a, b) = a + b$ it follows for each $v \in V$ there exists a unique pair (a, b) such that $v = a + b$. Given this decomposition of each vector in V we can define a projection onto B as follows: define $\pi_B : V \rightarrow B$ by $\pi_B(a + b) = b$. It is clear π_B is linear and $\text{Ker}(\pi_B) = A$ thus the first isomorphism theorem gives $V/A \cong B$. $\square$

It is interesting to study how the matrix of T and the matrix of $\bar{T}$ are related. This is part of a larger story which I tell now²⁴.

²⁴here I follow pages 231-233 of Charles Curtis' text

Recall²⁵, if V permits a direct sum decomposition in terms of invariant subspaces $W_1, \dots, W_k$ then there exists a basis β for V formed by concatenating the bases $\beta_1, \dots, \beta_k$ for $W_1, \dots, W_k$ respective. Moreover $[T]_{\beta, \beta}$ is in block-diagonal form where each block is simply the matrix of the restriction of $T|_{W_i} : W_i \rightarrow W_i$ with respect to β_i . What follows below is a bit different since we only assume that U is a T -invariant subspace.

Proposition 3.7.26.

Let V be finite a finite dimensional vector space over $\mathbb{F}$. If $T : V \rightarrow V$ is a linear transformation and $U \leq V$ for which $T(U) \leq U$ and if β_U is a basis of U and $\beta_U \cup \beta_2$ is a basis for V then

$$[T]_{\beta, \beta} = \begin{bmatrix} A & B \\ 0 & C \end{bmatrix}$$

where $A = [T|_U]_{\beta_U, \beta_U}$ and $C = [T_{V/U}]_{\beta_{V/U}, \beta_{V/U}}$ where $\beta_{V/U} = \{v_l + U \mid v_l \in \beta_2\}$.

Proof: let T and U be as in the statement of the proposition. The fact that $A = [T|_U]_{\beta_U, \beta_U}$ follows from $T(U) \leq U$. Denote $\beta_U = \{u_1, \dots, u_k\}$ and $\beta_2 = \{v_1, \dots, v_{n-k}\}$. Notice, $T(v_j) = \sum_{i=1}^k B_{ij}u_i + \sum_{l=1}^{n-k} C_{lj}v_l$. We define $T_{V/U} : V/U \rightarrow V/U$ by $T_{V/U}(x + U) = T(x) + U$ (this is well-defined since we assumed $T(U) \leq U$). Notice,

$$\bar{T}(v_j + U) = T(v_j) = \sum_{i=1}^k B_{ij}u_i + \sum_{l=1}^{n-k} C_{lj}v_l + U = \sum_{l=1}^{n-k} C_{lj}v_l + U. \quad \square$$

It is interesting to use the result above paired with Proposition 3.7.25. If $V = V_1 \oplus V_2$ and $T : V \rightarrow V$ is a linear transformation for which $T(V_1) \leq V_1$ and $T(V_2) \leq V_2$. We know from Theorem 3.6.14 that for the V_1, V_2 concatenated basis $\beta = \beta_1 \cup \beta_2$ we have $[T]_{\beta, \beta} = \begin{bmatrix} A_1 & 0 \\ 0 & A_2 \end{bmatrix}$. It follows, omitting explicit basis dependence, from Proposition 3.7.26 that

$$A_1 = [T_{V/V_2}] = [T|_{V_1}] \quad \& \quad A_2 = [T_{V/V_1}] = [T|_{V_2}].$$

In other words, given a block decomposition we can either view the blocks being attached to the restriction of the map to particular subspaces, or, we can see the blocks in terms of induced maps on quotients. Similar comments can be made for direct sums of more than two subspaces.

²⁵see Definition 3.6.13 and Theorem .

3.8 dual space

Definition 3.8.1.

Let V be a vector space over a field $\mathbb{F}$ then the **dual space** $V^* = L(V, \mathbb{F})$.

In the case that $\dim(V) = \infty$ this **algebraic dual space** is quite large and it is common to replace it with the set of bounded linear functionals. That said, our focus will be on the case $\dim(V) < \infty$. Roman's *Advanced Linear Algebra* is a good place to read more about the infinite dimensional case, or, most functional analysis texts. Let it be understood that $\dim(V) = n$ in the remainder of this section.

We should recall $L(V, \mathbb{F})$ is a vector space over $\mathbb{F}$ hence V^* is also a vector space over $\mathbb{F}$. The definition which follows is a natural next step:

Definition 3.8.2.

Let V be a vector space over a field $\mathbb{F}$ then the **double dual space** $V^{**} = L(V^*, \mathbb{F})$.

We can exchange V , V^* and V^{**} in a given application of linear algebra. In the finite dimensional case these are all isomorphic it is often possible to exchange one of these for the other.

Theorem 3.8.3.

Let V be a finite dimensional vector space over a field $\mathbb{F}$ then $V \cong V^*$ and $V \cong V^{**}$

Proof: observe $\dim(L(V, \mathbb{F})) = \dim(\mathbb{F}^{1 \times n}) = n = \dim(V)$ thus $V^* \cong V$. Next, since $V^{**} = (V^*)^*$ by construction we have $V^* \cong V^{**}$. Transitivity of isomorphism yields $V \cong V^{**}$. $\square$

Next, let us explore the explicit isomorphisms which support the result above.

Definition 3.8.4. an evaluation map:

For $x \in V$, let $\text{eval}_x : V^* \rightarrow \mathbb{F}$ be defined by $\text{eval}_x(\alpha) = \alpha(x)$ for each $\alpha \in V^*$.

I invite the reader to confirm that $\text{eval}_x \in L(V^*, \mathbb{F})$ hence $\text{eval}_x \in V^{**}$. Furthermore, the assignment $x \mapsto \text{eval}_x$ defines an explicit isomorphism of V and V^{**} . In other words, $\Psi : V \rightarrow V^{**}$ defined by

$$(\Psi(x))(\alpha) = \text{eval}_x(\alpha) = \alpha(x)$$

gives a bijective linear transformation from V to V^{**} . In appropriate contexts, we sometimes just write $x = \text{eval}_x$. If I make this abuse, I'll warn you²⁶.

The isomorphism $x \mapsto \text{eval}_x$ is *natural* in the sense that we could describe it without reference to a choice of basis. This is also possible for V^* if a metric²⁷ is given with V . Naturality aside, we can find an explicit isomorphism via the use a particular basis for V .

²⁶in these notes that is... In §2.6 of Insel, Spence and Friedberg, the notation $\hat{x}$ is used for eval_x .

²⁷a metric on a real vector space is a nondegenerate symmetric bilinear map on V

Definition 3.8.5.

Let $\beta = \{v_1, v_2, \dots, v_n\}$ form a basis for V over $\mathbb{F}$. For each $i = 1, 2, \dots, n$, define $v^i : V \rightarrow \mathbb{F}$ by linearly extending the formula $v^i(v_j) = \delta_{ij}$ for all $j = 1, 2, \dots, n$. We say $\beta^* = \{v^1, v^2, \dots, v^n\}$ is the **dual basis** to β .

The position of the indices up²⁸ or down²⁹ indicates how the given quantity transforms when we change coordinates. This notation is fairly popular in certain sectors of abstract math. We write for $x \in V$ that the coordinates with respect to $\beta = \{v_1, \dots, v_n\}$ are $x^1, \dots, x^n$ in the sense that:

$$x = \sum_{i=1}^n x^i v_i = x^1 v_1 + \dots + x^n v_n.$$

I know this notation is a bit weird the first time you see it. Just keep in mind the upper-indices are not powers. We ought to confirm that the dual basis is not wrongly labeled. You know, just because I call something a basis it doesn't make it so.

Proposition 3.8.6.

If $\beta^* = \{v^1, \dots, v^n\}$ is dual to basis $\beta = \{v_1, \dots, v_n\}$ for V then:

- (i.) if $x = \sum_{i=1}^n x^i v_i$ then $e^i(x) = x^i$.
- (ii.) β^* is a linearly independent subset of V^*
- (iii.) $\text{span}(\beta^*) = V^*$ and if $\alpha = \sum_{i=1}^n \alpha_i v^i$ then $\alpha(v_i) = \alpha_i$ for each $\alpha \in V^*$

Proof: most of these claims follow from the defining formula $v^i(v_j) = \delta_{ij}$. Suppose $x = \sum_{j=1}^n x^j v_j$ and calculate: by linearity of $v^i : V \rightarrow \mathbb{F}$ we have:

$$v^i(x) = v^i \left(\sum_{j=1}^n x^j v_j \right) = \sum_{j=1}^n x^j v^i(v_j) = \sum_{j=1}^n x^j \delta_{ij} = x^i.$$

To prove (ii.) suppose $\sum_{i=1}^n c_i v^i = 0$. Evaluate on v_j ,

$$\left(\sum_{i=1}^n c_i v^i \right) (v_j) = 0(v_j) \Rightarrow \sum_{i=1}^n c_i v^i(v_j) = 0 \Rightarrow \sum_{i=1}^n c_i \delta_{ij} = 0 \Rightarrow c_j = 0.$$

Hence β^* is a LI subset of V^* . By construction, it is clear that $\beta^* \subseteq V^*$ hence $\text{span}(\beta^*) \subseteq V^*$. Conversely, suppose $\alpha \in V^*$. Calculate: for $x = \sum_{j=1}^n x^j v_j$, by linearity of α and (i.)

$$\alpha(x) = \alpha \left(\sum_{j=1}^n x^j v_j \right) = \sum_{j=1}^n x^j \alpha(v_j) = \sum_{j=1}^n \alpha(v_j) v^j(x) \Rightarrow \alpha = \sum_{j=1}^n \alpha(v_j) v^j$$

²⁸up indices are called contravariant indices in physics

²⁹down indices are called covariant indices in physics

notice we have the desired formula as claimed in (iii.). We've shown $V^* \subseteq \text{span}(\beta^*)$ thus $V^* = \text{span}(\beta^*)$. $\square$

With the Proposition above in hand it is easy to provide an isomorphism of V and V^* . Simply define $\Psi : V \rightarrow V^*$ by linearly extending the formula $\Psi(v_i) = v^i$ for $i = 1, \dots, n$. This makes Ψ an isomorphism. Let me pause to include a few simple examples:

Example 3.8.7. Let $\alpha : \mathbb{F}^{n \times n} \rightarrow \mathbb{F}$ be defined by $\alpha(A) = \text{trace}(A)$. Since the trace is a linear map we have $\alpha \in (\mathbb{F}^{n \times n})^*$.

Example 3.8.8. Let $V = \mathcal{F}(\mathbb{R})$ and $x_o \in \mathbb{R}$. Define $\alpha(f) = f(x_o)$. Notice, for $f, g \in V$ and $c \in \mathbb{R}$,

$$\alpha(cf + g) = (cf + g)(x_o) = cf(x_o) + g(x_o) = c\alpha(f) + \alpha(g)$$

thus $\alpha \in V^*$.

Example 3.8.9. Let $v \in \mathbb{R}^n$ and define $\alpha(x) = x \cdot v$. It is easy to see α is linear hence $\alpha \in (\mathbb{R}^n)^*$.

Example 3.8.10. If $V = C^0[0, 1]$ then define $\alpha(f) = \int_0^1 xf(x) dx$. Observe

$$\alpha(cf + g) = \int_0^1 x(cf(x) + g(x)) dx = c \int_0^1 xf(x) dx + \int_0^1 xg(x) dx = c\alpha(f) + \alpha(g).$$

Thus $\alpha \in V^*$.

Definition 3.8.11.

If $W \leq V$ then define the **annihilator** of W in V^* by $\text{ann}(W) = \{\alpha \in V^* \mid \alpha(W) = 0\}$.

The condition $\alpha(W) = 0$ means $\alpha(w) = 0$ for each $w \in W$. There are many interesting theorems we can state for annihilators. For instance, if $V_1 \leq V_2$ then $\text{ann}(V_2) \leq \text{ann}(V_1)$. We explore some such theorems in the homework.

Example 3.8.12. Consider $V = \mathbb{R}^4$. Let $W = \text{span}\{(1, 1, 1, 1), (1, 0, 0, 0)\}$ then $\text{ann}(W) = \{\alpha \in (\mathbb{R}^4)^* \mid \alpha(1, 1, 1, 1) = 0 \text{ \& } \alpha(1, 0, 0, 0) = 0\}$. Thus, $\alpha \in \text{ann}(W)$ has

$$\alpha_1 = 0 \quad \& \quad \alpha_1 + \alpha_2 + \alpha_3 + \alpha_4 = 0 \quad \Rightarrow \quad \alpha_2 = -\alpha_3 - \alpha_4$$

we deduce $\text{ann}(W) = \text{span}\{e^3 - e^2, e^4 - e^2\}$

Example 3.8.13. Consider $V = \mathbb{R}^{2 \times 2}$ and let W be the symmetric 2×2 matrices. If $\alpha \in \text{ann}(W)$ then for $A = A^T$ we have $\alpha(A) = 0$. In fact, the basis for W is given by $\{E_{11}, E_{12} + E_{21}, E_{22}\}$. We can extend this to a basis for V by adjoining $E_{12} - E_{21}$. The dual basis to the standard matrix basis is given by $E^{ij}(E_{kl}) = \delta_{ik}\delta_{jl}$. Notice, $\alpha = E^{12} - E^{21}$ is in the annihilator since:

$$\alpha(A) = (E^{12} - E^{21})(A) = (E^{12} - E^{21})(A_{11}E_{11} + A_{12}E_{12} + A_{21}E_{21} + A_{22}E_{22}) = A_{12} - A_{21} = 0.$$

In fact, $\text{span}\{E^{12} - E^{21}\} = \text{ann}(W)$.

3.8.1 transpose of linear map and the up-down notation

The notation I used in Definition 3.3.6 is a perfectly good notation. However, the following notation is popular in many good books:

Definition 3.8.14.

Let $V(\mathbb{F})$ be a vector space with basis $\beta = \{v_1, \dots, v_n\}$. Let $W(\mathbb{F})$ be a vector space with basis $\gamma = \{w_1, \dots, w_m\}$. If $T : V \rightarrow W$ is a linear transformation then we define the **matrix of T with respect to β, γ** as $[T]_{\beta}^{\gamma} \in \mathbb{F}^{m \times n}$ which is implicitly defined by

$$L_{[T]_{\beta}^{\gamma}} = \Phi_{\gamma} \circ T \circ \Phi_{\beta}^{-1} \quad \text{or} \quad [T]_{\beta}^{\gamma} = [[T(v_1)]_{\gamma} | \cdots | [T(v_n)]_{\gamma}].$$

Let's examine how we can express the formula for a linear transformation in terms of the dual basis and basis. Note we have identities $x = \sum_{i=1}^n v^i(x) v_i$ for any $x \in V$ and $y = \sum_{j=1}^m w^j(y) w_j$ for any $y \in W$.

$$\begin{aligned} T(x) &= T\left(\sum_{i=1}^n v^i(x) v_i\right) \\ &= \sum_{i=1}^n v^i(x) T(v_i) \\ &= \sum_{i=1}^n \left(\sum_{j=1}^m w^j(T(v_i)) w_j\right) v^i(x) \\ &= \left(\sum_{i=1}^n \sum_{j=1}^m \left([T]_{\beta}^{\gamma}\right)_i^j w_j v^i\right)(x) \end{aligned}$$

Thus

$$T = \sum_{i=1}^n \sum_{j=1}^m \left([T]_{\beta}^{\gamma}\right)_i^j w_j v^i.$$

In this formalism, the row-index j is written as a superscript whereas the column-index i is written as a subscript. The formula

$$w^j(T(v_i)) = \left([T]_{\beta}^{\gamma}\right)_i^j$$

can be useful to derive formulas of coordinate change. Another rather neat aspect of this notation is seen in the construction of the **transpose of a linear transformation**.

Theorem 3.8.15.

Suppose V and W are finite dimensional vector spaces with bases β and γ respective. Further suppose $L : V \rightarrow W$ is a linear transformation. Then $L^t : W^* \rightarrow V^*$ given by

$$(L^t(\alpha))(v) = \alpha(L(v))$$

for each $\alpha \in W^*$ and $v \in V$ defines a linear transformation such that $[L^t]_{\gamma^*}^{\beta^*} = \left([L]_{\beta}^{\gamma}\right)^T$

Proof: suppose $\alpha, \sigma \in W^*$ and $x \in V$ and $c \in \mathbb{F}$ then calculate

$$\begin{aligned} (L^t(c\alpha + \sigma))(x) &= (c\alpha + \sigma)(L(x)) \\ &= c\alpha(L(x)) + \sigma(L(x)) \\ &= (cL^t(\alpha))(x) + (L^t(\sigma))(x) \\ &= (cL^t(\alpha) + L^t(\sigma))(x). \end{aligned}$$

Therefore $L^t(c\alpha + \sigma) = cL^t(\alpha) + L^t(\sigma)$. Similarly, we can verify $(L^t(\alpha))(cx + y) = c(L^t(\alpha))(x) + (L^t(\alpha))(y)$ hence L^t is indeed a mapping from W^* to V^* and L^t is itself a linear map. Suppose $\beta = \{v_1, \dots, v_n\}$ and $\gamma = \{w_1, \dots, w_m\}$ are bases for V and W respectively and their dual basis are denoted by $\beta^* = \{v^1, \dots, v^n\}$ and $\gamma^* = \{w^1, \dots, w^m\}$. The matrix of $L^t : W^* \rightarrow V^*$ with respect to these bases is calculated via:

$$[L^t]_{\gamma^*}^{\beta^*} = [[L^t(w^1)]_{\beta^*} | \dots | [L^t(w^m)]_{\beta^*}]$$

Thus

$$\begin{aligned} \text{col}_i([L^t]_{\gamma^*}^{\beta^*}) &= [L^t(w^i)]_{\beta^*} \\ &= (L^t(w^i)(v_1), \dots, L^t(w^i)(v_n)) \\ &= (w^i(L(v_1)), \dots, w^i(L(v_n))) \end{aligned}$$

which provides the j -th component of the i -th column is:

$$([L^t]_{\gamma^*}^{\beta^*})_i^j = w^i(L(v_j)) = ([L]_{\beta}^{\gamma})_j^i$$

Therefore, $[L^t]_{\gamma^*}^{\beta^*} = ([L]_{\beta}^{\gamma})^T$ where T denotes the transpose of the matrix. $\square$

The up/down index notation is cute, but, I think I will forego this notation elsewhere. This section is partly based on §2.6 of Insel, Spence and Friedberg. I believe there are a number of worthwhile exercises to deepen our understanding of dual space. I will probably assign some of those.

Chapter 4

Jordan Form

He that believeth on the Son hath everlasting life: and he that believeth not the Son shall not see life; but the wrath of God abideth on him. JOHN 3:36, KJV

Given two linear transformations on a finite dimensional vector space $V(\mathbb{F})$ when is it the case that they share the same formula ? In other words, given $T : V \rightarrow V$ and $S : V \rightarrow V$ both linear, when does there exist bases β and γ for V such that $[T]_{\beta,\beta} = [S]_{\gamma,\gamma}$. Notice coordinate change gives that

$$[S]_{\gamma,\gamma} = P^{-1}[S]_{\beta,\beta}P.$$

If T and S have the same formula in different coordinate systems then there will exist an invertible matrix P for which

$$[T]_{\beta,\beta} = P^{-1}[S]_{\beta,\beta}P.$$

Notice that if this is true for basis β then it is also true for any other basis δ since

$$[T]_{\beta,\beta} = Q^{-1}[T]_{\delta,\delta}Q \quad \& \quad [S]_{\beta,\beta} = Q^{-1}[S]_{\delta,\delta}Q$$

for some change of basis matrix Q . Thus

$$Q^{-1}[T]_{\delta,\delta}Q = P^{-1}Q^{-1}[S]_{\delta,\delta}QP.$$

which yields the following if we solve for $[T]_{\delta,\delta}$

$$[T]_{\delta,\delta} = QP^{-1}Q^{-1}[S]_{\delta,\delta}QPQ^{-1} = (QPQ^{-1})^{-1}[S]_{\delta,\delta}(QPQ^{-1}).$$

Consequently, if the matrix of T and the matrix of S are similar in view of one basis choice then their matrices will be similar in any other choice of basis. We find that deciding whether two linear transformations on V have the same formula is equivalent to deciding whether a given pair of matrices $A, B \in \mathbb{F}^{n \times n}$ are **similar**. Recall, A and B are similar if there exists an invertible matrix R for which

$$B = R^{-1}AR.$$

Similarity is an equivalence relation on $\mathbb{F}^{n \times n}$ and our goal in this chapter is to understand how to categorize the equivalence classes of this relation. It turns out, over $\mathbb{C}$ the answer is elegantly given by the **Jordan form**. However, over $\mathbb{R}$, we must use the **real Jordan form**. For a more abstract field, the Jordan form may or may not be helpful since it requires a certain polynomial equation to have solutions in the field. I should mention, another competitor in this space of concepts is the **rational canonical form**. It does not require the polynomials which define it to have roots in

the field, thus it is a more general construction. That said, the Jordan form is more often seen in applications and it will provide sufficient challenge for this course. I leave the rational canonical form for your next course in Linear Algebra.¹

4.1 eigenvectors and diagonalization

Let us begin with the definition of a diagonalizable transformation and matrix:

Definition 4.1.1.

Let $T : V \rightarrow V$ be a linear transformation over a vector space V over $\mathbb{F}$. If there exists a basis β for which $[T]_{\beta, \beta}$ is diagonal then T is said to be **diagonalizable**. Likewise, if $A \in \mathbb{F}^{n \times n}$ is similar to a diagonal matrix then A is diagonalizable.

Notice that the matrix of a diagonalizable transformation is necessarily diagonalizable. If $\beta = \{v_1, \dots, v_n\}$ is such that $[T]_{\beta, \beta}$ is diagonal then there exist scalars $\lambda_1, \dots, \lambda_n \in \mathbb{F}$, possibly repeated, such that

$$[T]_{\beta, \beta} = \begin{bmatrix} \lambda_1 & & \\ & \ddots & \\ & & \lambda_n \end{bmatrix} = [\lambda_1 e_1 | \dots | \lambda_n e_n].$$

Thus $[T(v_i)]_{\beta} = \lambda_i e_i$ and we find $T(v_i) = \lambda_i v_i$ for $i = 1, \dots, n$.

Definition 4.1.2.

If $T : V \rightarrow V$ is a linear transformation on a vector space V over $\mathbb{F}$ then $v \neq 0$ is an **eigenvector** with **eigenvalue** $\lambda \in \mathbb{F}$ if $T(v) = \lambda v$. Likewise, a matrix $A \in \mathbb{F}^{n \times n}$ has eigenvector $x \neq 0$ with **eigenvalue** $\lambda \in \mathbb{F}$ if $Ax = \lambda x$. A basis of eigenvectors is known as an **eigenbasis**.

Notice if $T(v) = \lambda v$ then $[T]_{\beta, \beta}[v]_{\beta} = \lambda[v]_{\beta}$. This calculation shows that the coordinate vector of an eigenvector is itself an eigenvector of the matrix of the transformation. It follows we may study the problem of finding eigenvectors for a given transformation by selecting a basis and working out the eigenvectors for the matrix with respect to the basis. We should also note, diagonalizability of a matrix or transformation amounts to deciding whether or not there exists an eigenbasis for the matrix or map. We will soon see examples which demonstrate that not all transformations are diagonalizable. I often motivate the quest for the Jordan form as a method to deal with pesky non-diagonalizable objects.

Theorem 4.1.3.

Let $A \in \mathbb{F}^{n \times n}$, then $\lambda \in \mathbb{F}$ is an **eigenvalue** of A if and only if $\det(A - \lambda I) = 0$.

Proof: if $\lambda \in \mathbb{F}$ is an eigenvalue of A then there exists $x \neq 0$ for which $Ax = \lambda x$ thus $(A - \lambda I)x = 0$. Thus $A - \lambda I$ is a noninvertible matrix with $\det(A - \lambda I) = 0$.

Conversely, if $\det(A - \lambda I) = 0$ then the matrix $A - \lambda I$ is non invertible so the columns of $A - \lambda I$ are linearly dependent which implies there exists $x \neq 0$ for which $(A - \lambda I)x = 0$. Thus $Ax = \lambda x$

¹it can reasonably be covered in Math 422 where we have more machinery for the theory of polynomial factoring at our disposal

which proves $\lambda \in \mathbb{F}$ is an eigenvalue of A . $\square$

If $T : V \rightarrow V$ is a linear transformation on a finite dimensional vector space $V(\mathbb{F})$ with basis β then $\det(T) = \det([T]_{\beta,\beta})$. Thus,

$$\det(T - \lambda Id_V) = \det([T - \lambda Id_V]_{\beta,\beta}) = \det([T]_{\beta,\beta} - \lambda I)$$

as $[Id_V]_{\beta,\beta} = I$. Therefore we find the natural corollary to the above theorem:

Corollary 4.1.4.

If $T : V \rightarrow V$ is a linear transformation on a finite dimensional vector space $V(\mathbb{F})$ with eigenvalue $\lambda \in \mathbb{F}$ if and only if $\det(T - \lambda Id_V) = 0$.

Definition 4.1.5.

The **characteristic polynomial** of A is given by $p(s) = \det(A - sI)$. Likewise, the **characteristic polynomial** of T is given by $p(s) = \det(T - sId)$

Notice the characteristic polynomial defined has the form:

$$p(t) = (-1)^n t^n + c_{n-1} t^{n-1} + \cdots + c_1 t + c_0$$

Polynomials of operators are understood by interpreting multiplication as composition.

Definition 4.1.6.

If $T : V \rightarrow V$ is a linear transformation on a vector space V over $\mathbb{F}$ then $T^0 = Id_V$ and $T^1 = T$ and $T^{k+1} = T \circ T^k$ for all $k \in \mathbb{N}$. If $f(x) = c_0 + c_1 x + \cdots + c_k x^k$ then we define:

$$f(T) = c_0 + c_1 T + \cdots + c_k T^k.$$

We use $Id_v = 1$ and $cId_v = c$ for convenience of exposition.

Example 4.1.7. Let $T(x, y) = (-y, x)$ for all $(x, y) \in \mathbb{R}^2$ then $A = [T] = \begin{pmatrix} 0 & -1 \\ 1 & 0 \end{pmatrix}$ then

$$\det(T - sI) = \det \begin{pmatrix} -s & -1 \\ 1 & -s \end{pmatrix} = s^2 + 1.$$

Since $s^2 + 1 \neq 0$ for all $s \in \mathbb{R}$ we find T has no eigenvalues and hence T is not diagonalizable. However, it may be interesting to note $p(s) = s^2 + 1$ has $p(T) = T^2 + Id$ and $T \circ T(x, y) = T(T(x, y)) = T(-y, x) = (-x, -y) = -Id(x, y)$ thus $T^2 = -Id$ and we find $p(T) \equiv 0$ (this is read $p(T)$ is identically zero). This exemplifies the **Cayley Hamilton Theorem**² which states that the a linear transformation solves its own characteristic equation. Another way to appreciate this is via the matrix formalism:

$$p(A) = A^2 + I = \begin{pmatrix} 0 & -1 \\ 1 & 0 \end{pmatrix}^2 + \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix} = 0.$$

²I'll hopefully properly state and prove this result later in this chapter

Remark 4.1.8. *dependence on field.*

Notice that $s^2 + 1 = (s + i)(s - i)$ over $\mathbb{C}$. If the previous example was defined with domain $\mathbb{C}^2$ then $T(z, w) = (-w, z)$ would define T which is diagonalizable with eigenvalues $\lambda = \pm i$. Our inability to diagonalize T in the previous example can be remedied by enlarging our field of scalars. This is not always possible as the next example will show.

Example 4.1.9. Let $V = P_2(\mathbb{R})$ the space of polynomials of degree at most 2 with real coefficients. Define $T = d/dx$ by

$$T(ax^2 + bx + c) = 2ax + b$$

Notice $T^2(ax^2 + bx + c) = 2a$ and $T^3 = 0$. Use basis $\beta = \{x^2, x, 1\}$ and let $[T]_{\beta, \beta} = A$ we find

$$A = [[T(x^2)]_{\beta} | [T(x)]_{\beta} | [T(1)]_{\beta}] = [[2x]_{\beta} | [1]_{\beta} | [0]_{\beta}] = \begin{bmatrix} 0 & 0 & 0 \\ 2 & 0 & 0 \\ 0 & 1 & 0 \end{bmatrix}.$$

Thus,

$$p(s) = \det \begin{bmatrix} -s & 0 & 0 \\ 2 & -s & 0 \\ 0 & 1 & -s \end{bmatrix} = -s^3$$

We find the eigenvalues of T are $\lambda = 0$ with an **algebraic multiplicity** of three. What are the eigenvectors of T ? We seek $v = ax^2 + bx + c$ for which $T(v) = 0v = 0$ which gives $2ax + b = 0$ and hence $a = 0$ and $b = 0$. Thus $v = c$ is the form of an eigenvector with eigenvalue zero; that is, the constant polynomials are the only eigenvectors of the differentiation operator. We find T is not diagonalizable since it is not possible to form an eigenbasis for T . Note $\dim(V) = 3$ yet we cannot even find two linearly independent eigenvectors for T .

Apparently it is fairly easy to find maps and matrices which are not diagonalizable. Our first example was a rotation in the plane by $\pi/2$ radians and our second example is mere differentiation. If your world of linear transformations only goes up to diagonalizable objects then clearly we're missing a big part of the overarching story. That said, for the remainder of this section we'll look at some important properties of eigenvectors and some indirect methods to test diagonalizability.

Proposition 4.1.10.

Let $A \in \mathbb{F}^{n \times n}$, then zero is an eigenvalue if and only if A^{-1} does not exist.

Proof: if $\lambda = 0$ is an eigenvalue then there exists $x \neq 0$ for which $Ax = 0$. Suppose A^{-1} existed then $Ax = 0$ implies $A^{-1}Ax = A^{-1}0$ hence $x = 0$ which is a contradiction. Therefore, A^{-1} does not exist.

Conversely, if A^{-1} does not exist then there must be a linear dependence amongst the columns of A . Otherwise, the columns of A form a basis for $\mathbb{F}^n$ and so $Av_i = e_i$ has a solution for $i = 1, \dots, n$ and we may form $A^{-1} = [v_1 | \dots | v_n]$ which is a contradiction. Thus A has a linear dependence amongst its columns and so there exists $x \neq 0$ for which $Ax = 0$. Consequently, $\lambda = 0$ is an eigenvalue for A . $\square$

I tried to give relatively complete arguments for the proposition above. If we remembered more matrix theory from the previous course the proof could be considerably shorter. If in doubt about the detail you've shown in a proof, you can ask me whether or not you're on target. Sometimes I can give a helpful hint or at least let you know your current argument is circular etc.

Proposition 4.1.11.

Let $A \in \mathbb{R}^{n \times n}$ where n is odd, then there exists at least one real eigenvalue.
 Let $A \in \mathbb{C}^{n \times n}$ then there exist n -eigenvalues, possibly repeated.

Proof: notice $p(t)$ is an n -th order polynomial. Over $\mathbb{R}$, any polynomial of odd degree has at least one real zero. We can prove this assertion carefully by an application of the intermediate value theorem on a sufficiently large interval. Similarly, over $\mathbb{C}$, we know the Fundamental Theorem of Algebra states an n -th order polynomial over $\mathbb{C}$ has n -zeros possibly repeated. $\square$

Proposition 4.1.12.

If $A \in \mathbb{F}^{n \times n}$ then A has n eigenvalues $\lambda_1, \lambda_2, \dots, \lambda_n$ then $\det(A) = \lambda_1 \lambda_2 \cdots \lambda_n$.

Proof: If $A \in \mathbb{F}^{n \times n}$ then A has n eigenvalues $\lambda_1, \lambda_2, \dots, \lambda_n \in \mathbb{F}$ then the characteristic polynomial $p(t) = \det(A - tI)$ factors over $\mathbb{F}$:

$$p(t) = (-1)^n (t - \lambda_1)(t - \lambda_2) \cdots (t - \lambda_n)$$

since $p(\lambda_i) = 0$ for $i = 1, 2, \dots, n$. Notice $p(0) = \det(A)$. However, we also know $p(0)$ is the constant term in the standard form of $p(t)$. Thus,

$$(-1)^n (-\lambda_1)(-\lambda_2) \cdots (-\lambda_n) = \lambda_1 \lambda_2 \cdots \lambda_n = \det(A). \quad \square$$

Proposition 4.1.13.

If $A \in \mathbb{F}^{n \times n}$ has e-vector v with eigenvalue λ then v is a e-vector of A^k with e-value λ^k .

Proof: let $A \in \mathbb{F}^{n \times n}$ have e-vector v with eigenvalue λ . Consider,

$$A^k v = A^{k-1} A v = A^{k-1} \lambda v = \lambda A^{k-2} A v = \lambda^2 A^{k-2} v = \cdots = \lambda^k v.$$

The $\cdots$ is properly replaced by a formal induction argument. $\square$.

Proposition 4.1.14.

Let A be a upper or lower triangular matrix then the eigenvalues of A are the diagonal entries of the matrix.

Proof: follows immediately from Proposition ?? since the diagonal entries of $A - \lambda I$ are of the form $A_{ii} - \lambda$ hence the characteristic equation has the form $\det(A - \lambda I) = (A_{11} - \lambda)(A_{22} - \lambda) \cdots (A_{nn} - \lambda)$ which has solutions $\lambda = A_{ii}$ for $i = 1, 2, \dots, n$. $\square$

Likewise, as a diagonal matrix is both upper and lower triangular we find the diagonal entries are the eigenvalues of such a matrix.

Proposition 4.1.15.

Let $A \in \mathbb{C}^{2 \times 2}$. The eigenvalues determine the $\det(A)$ and $\text{trace}(A)$:

$$\det(A) = \lambda_1 \lambda_2 \quad \& \quad \text{trace}(A) = \lambda_1 + \lambda_2.$$

Proof: we know Proposition 4.1.12 yields $\det(A) = \lambda_1\lambda_2$. If $A = \begin{bmatrix} a & b \\ c & d \end{bmatrix}$ then

$$p(t) = \det \begin{bmatrix} a-t & b \\ c & d-t \end{bmatrix} = (t-a)(t-d) - bc.$$

Algebra reveals $p(t) = t^2 - (a+d)t + ad - bc$ and completing the square yields:

$$\lambda_{\pm} = \frac{a+d \pm \sqrt{(a+d)^2 + 4bc}}{2}$$

Let $\lambda_1 = \lambda_+$ and $\lambda_2 = \lambda_-$. Observe $\lambda_1 + \lambda_2 = a + d = \text{trace}(A)$. $\square$

In fact, the proposition above also applies to $A \in \mathbb{C}^{n \times n}$. We can show³, $\text{trace}(A) = \sum_{j=1}^n \lambda_j$ where λ_j are eigenvalues of A and Proposition 4.1.12 explained why the determinant is the product of all the e-values in such a case. In fact, this is also true for $A \in \mathbb{F}^{n \times n}$ provided the characteristic polynomial of A factors into a product of possibly repeated linear factors.

Proposition 4.1.16.

If $A \in \mathbb{F}^{n \times n}$ has e-vector v_1 with e-value λ_1 and e-vector v_2 with e-value λ_2 such that $\lambda_1 \neq \lambda_2$ then $\{v_1, v_2\}$ is linearly independent.

Proof: Let v_1, v_2 have e-values λ_1, λ_2 respective and assume towards a contradiction that $v_2 = kv_1$ for some nonzero constant k . Multiply by the matrix A ,

$$Av_1 = A(kv_1) \Rightarrow \lambda_1 v_1 = k\lambda_2 v_1$$

But we can replace v_1 on the l.h.s. with kv_1 hence,

$$\lambda_1 kv_1 = k\lambda_2 v_1 \Rightarrow k(\lambda_1 - \lambda_2)v_1 = 0$$

Note, $k \neq 0$ and $v_1 \neq 0$ by assumption thus the equation above indicates $\lambda_1 - \lambda_2 = 0$ therefore $\lambda_1 = \lambda_2$ which is a contradiction. Therefore there does not exist such a k and the vectors are linearly independent. $\square$

A direct argument is also possible. Suppose $\{v_1, v_2\}$ is a set of nonzero vectors with $Av_1 = \lambda_1 v_1$ and $Av_2 = \lambda_2 v_2$ suppose $c_1 v_1 + c_2 v_2 = 0$. Multiply by $A - \lambda_1 I$,

$$c_1(A - \lambda_1 I)v_1 + c_2(A - \lambda_1 I)v_2 = 0 \Rightarrow c_2(\lambda_2 - \lambda_1)v_2 = 0$$

as $\lambda_2 - \lambda_1 \neq 0$ and $v_2 \neq 0$ hence $c_2 = 0$. Multiplication by $A - \lambda_2 I$ likewise reveals $c_1 = 0$. Therefore, $\{v_1, v_2\}$ is LI. You can choose which proof you think is best.

Proposition 4.1.17.

If $A \in \mathbb{F}^{n \times n}$ has eigenvectors $v_1, v_2, \dots, v_k$ with eigenvalues $\lambda_1, \lambda_2, \dots, \lambda_k \in \mathbb{F}$ such that $\lambda_i \neq \lambda_j$ for all $i \neq j$ then $\{v_1, v_2, \dots, v_k\}$ is linearly independent.

³this is simple to prove once we have established the Jordan form, and we could prove this for diagonalizable matrices, I've asked it on tests before, it's not too hard a problem

Proof: I begin with a direct proof. Suppose $v_1, v_2, \dots, v_k$ are e-vectors with e-values $\lambda_1, \lambda_2, \dots, \lambda_k \in \mathbb{F}$ such that $\lambda_i \neq \lambda_j$ for all $i \neq j$. Suppose $c_1 v_1 + c_2 v_2 + \dots + c_k v_k = 0$. Multiply by $\prod_{i=1}^{k-1} (A - \lambda_i I)$,

$$c_1 \prod_{i=1}^{k-1} (A - \lambda_i I) v_1 + \dots + c_{k-1} \prod_{i=1}^{k-1} (A - \lambda_i I) v_{k-1} + c_k \prod_{i=1}^{k-1} (A - \lambda_i I) v_k = 0 \star$$

Consider that the terms in the product commute as:

$$(A - \lambda_i I)(A - \lambda_j I) = A^2 - (\lambda_i + \lambda_j)A + \lambda_i \lambda_j I = (A - \lambda_j I)(A - \lambda_i I).$$

It follows that we can bring $(A - \lambda_j I)$ to the right of the product multiplying the j -th summand:

$$c_1 \prod_{i \neq 1}^{k-1} (A - \lambda_i I)(A - \lambda_1 I) v_1 + \dots + c_{k-1} \prod_{i \neq k-1}^{k-1} (A - \lambda_i I)(A - \lambda_{k-1} I) v_{k-1} + c_k \prod_{i=1}^{k-1} (A - \lambda_i I) v_k = 0 \star^2$$

Notice, for $i \neq j$, $(A - \lambda_j I)v_i = \lambda_i v_i - \lambda_j v_i = (\lambda_i - \lambda_j)v_i \neq 0$ as $\lambda_i \neq \lambda_j$ and $v_i \neq 0$. On the other hand, if $i = j$ then $(A - \lambda_i I)v_i = \lambda_i v_i - \lambda_i v_i = 0$. Therefore, in $\star$ we find that terms with coefficients $c_1, c_2, \dots, c_{k-1}$ all vanish. All that remains is:

$$c_k \prod_{i=1}^{k-1} (A - \lambda_i I) v_k = 0 \star^3$$

We calculate,

$$\begin{aligned} c_k \prod_{i=1}^{k-1} (A - \lambda_i I) v_k &= \prod_{i=1}^{k-2} (A - \lambda_i I)(A - \lambda_{k-1} I) v_k = c_k (\lambda_k - \lambda_{k-1}) \prod_{i=1}^{k-2} (A - \lambda_i I) v_k \\ &= c_k (\lambda_k - \lambda_{k-1}) (\lambda_k - \lambda_{k-2}) \prod_{i=1}^{k-3} (A - \lambda_i I) v_k \\ &= c_k (\lambda_k - \lambda_{k-1}) (\lambda_k - \lambda_{k-2}) \dots (\lambda_k - \lambda_1) v_k. \end{aligned}$$

However, as $v_k \neq 0$ and $\lambda_k \neq \lambda_i$ for $i = 1, \dots, k-1$ it follows from the identity above that $\star^3$ implies $c_k = 0$. Next, we repeat the argument, except only multiply $\star$ by $\prod_{i=1}^{k-2} (A - \lambda_i I)$ which yields $c_{k-1} = 0$. We continue in this fashion until we have shown $c_1 = c_2 = \dots = c_k = 0$. Hence $\{v_1, \dots, v_k\}$ is linearly independent as claimed. $\square$

I am fond of the argument which was just offered. Technically, it could be improved by including explicit induction arguments in place of $\dots$. The next argument is similar to our initial argument for two vectors.

Proof: Let e-vectors $v_1, v_2, \dots, v_k$ have e-values $\lambda_1, \lambda_2, \dots, \lambda_k$. Let us prove the claim by induction on k . Note $k = 1$ and $k = 2$ we have already shown in previous work. Suppose inductively the claim is true for $k-1$. Consider, towards a contradiction, that there is some vector v_j which is a nontrivial linear combination of the other vectors:

$$v_j = c_1 v_1 + c_2 v_2 + \dots + \widehat{c_j v_j} + \dots + c_k v_k$$

Multiply by A ,

$$A v_j = c_1 A v_1 + c_2 A v_2 + \dots + \widehat{c_j A v_j} + \dots + c_k A v_k$$

Which yields,

$$\lambda_j v_j = c_1 \lambda_1 v_1 + c_2 \lambda_2 v_2 + \cdots + \widehat{c_j \lambda_j v_j} + \cdots + c_k \lambda_k v_k$$

But, we can replace v_j on the l.h.s with the linear combination of the other vectors. Hence

$$\lambda_j [c_1 v_1 + c_2 v_2 + \cdots + \widehat{c_j v_j} + \cdots + c_k v_k] = c_1 \lambda_1 v_1 + c_2 \lambda_2 v_2 + \cdots + \widehat{c_j \lambda_j v_j} + \cdots + c_k \lambda_k v_k$$

Consequently,

$$c_1(\lambda_j - \lambda_1)v_1 + c_2(\lambda_j - \lambda_2)v_2 + \cdots + c_j(\lambda_j - \lambda_j)v_j + \cdots + c_k(\lambda_j - \lambda_k)v_k = 0$$

However, this is a set of $k - 1$ e-vectors with distinct e-values linearly combined to give zero. It follows from the induction claim that each coefficient is trivial. As $\lambda_j \neq \lambda_i$ for $i \neq j$ it is thus necessary that $c_1 = c_2 = \cdots = c_k = 0$. But, this implies $v_j = 0$ which contradicts $v_j \neq 0$ as is known since v_j was assumed an e-vector. Hence $\{v_1, \dots, v_k\}$ is LI as claimed and by induction on $k \in \mathbb{N}$ we find the proposition is true. $\square$

Further structure theory of eigenvectors is considered later in this chapter. I'll conclude by explaining why our proofs for the matrix case also inform our study of linear transformations.

Proposition 4.1.18.

If $T \in L(V, V)$ for a vector space V over $\mathbb{F}^{n \times n}$ has eigenvectors $v_1, v_2, \dots, v_k$ with eigenvalues $\lambda_1, \lambda_2, \dots, \lambda_k \in \mathbb{F}$ such that $\lambda_i \neq \lambda_j$ for all $i \neq j$ then $\{v_1, v_2, \dots, v_k\}$ is linearly independent.

Proof: pick a basis for V , set $A = [T]_{\beta, \beta}$ and observe $T(v) = \lambda v$ implies $[T(v)]_{\beta} = \lambda[v]_{\beta}$. Thus, as $[T(v)]_{\beta} = [T]_{\beta, \beta}[v]_{\beta}$ we find $A[v]_{\beta} = \lambda[v]_{\beta}$. That is, the coordinate vector of each eigenvector in V is a column vector which is an eigenvector of $A = [T]_{\beta, \beta}$ with the same eigenvalue. Apply Proposition 4.1.17 to find a set of linearly independent e-vectors of A . Finally, use the isomorphism $\Phi_{\beta}^{-1} : \mathbb{F}^n \rightarrow V$ to show the set of e-vectors in V are LI. $\square$

4.2 eigenbases and eigenspaces

If we have a basis of eigenvectors then it is called an *eigenbasis*. For a linear transformation:

Definition 4.2.1. *eigenbasis for linear transformation*

Let $T : V \rightarrow V$ be a linear transformation on a vector space V over $\mathbb{F}$. If there exists a basis $\beta = \{v_1, v_2, \dots, v_n\}$ of V such that $T(v_j) = \lambda_j v_j$ for some constant $\lambda_j \in \mathbb{F}$ then we say β is an **eigenbasis** of T .

Recall, a **diagonal** matrix D is one for which $D_{ij} = 0$ for $i \neq j$. The matrix of a linear transformation with respect to an eigenbasis will be diagonal with e-values as the diagonal entries:

Proposition 4.2.2.

If $T : V \rightarrow V$ is a linear transformation and T has an eigenbasis $\beta = \{f_1, \dots, f_n\}$ where f_j is an eigenvector with eigenvalue λ_j for $j = 1, \dots, n$ then

$$[T]_{\beta, \beta} = \begin{bmatrix} \lambda_1 & 0 & \cdots & 0 \\ 0 & \lambda_2 & \cdots & 0 \\ \vdots & \vdots & \cdots & \vdots \\ 0 & 0 & \cdots & \lambda_n \end{bmatrix}.$$

Proof: In general, $[T]_{\beta,\beta} = [[T(f_1)]_\beta \mid [T(f_2)]_\beta \mid \cdots \mid [T(f_n)]_\beta]$. However, as v_j is an eigenvector we have $T(v_j) = \lambda_j v_j$. Moreover, by definition of β coordinates, $[f_j]_\beta = e_j \in \mathbb{F}^n$ hence:

$$\begin{aligned} [T]_{\beta,\beta} &= [[\lambda_1 f_1]_\beta \mid [\lambda_2 f_2]_\beta \mid \cdots \mid [\lambda_n f_n]_\beta] \\ &= [\lambda_1 e_1 \mid \lambda_2 e_2 \mid \cdots \mid \lambda_n e_n]. \end{aligned}$$

Thus, $[T]_{\beta,\beta}$ is diagonal with $\lambda_1, \lambda_2, \dots, \lambda_n$ on the diagonal as claimed. $\square$

Now would be a good time to read Example 3.4.8 again. There we found the matrix of a linear transformation $T : \mathbb{R}^3 \rightarrow \mathbb{R}^3$ is diagonal with respect to an eigenbasis. As we have seen, there exist linear transformations which can not be diagonalized. However, even for those transformations, we may still be able to find a basis which partially diagonalizes the matrix. In particular, this brings us to the definition of the λ_j -**eigenspace**. We will soon see that the restriction of the linear transformation to this space will be diagonal.

Definition 4.2.3. *eigenspace and geometric vs. algebraic multiplicity*

Let $T : V \rightarrow V$ be a linear transformation. We define the set of all eigenvectors of T with eigenvalue λ_j adjoined the zero-vector is the λ_j -**eigenspace** denoted by $\mathcal{E}_{\lambda_j}$ or simply $\mathcal{E}_j$. The dimension of $\mathcal{E}_{\lambda_j}$ is known as the **geometric multiplicity** of λ_j . The **algebraic multiplicity** of λ_j is the largest $m \in \mathbb{N}$ for which number of times $(t - \lambda_j)^m$ appears as a factor of the characteristic polynomial $p(t)$.

I will provide examples once we focus on the matrix analog of the definition above. For the moment, we just have a few more theoretical items to clarify:

Proposition 4.2.4.

If $T : V \rightarrow V$ is a linear transformation and $\mathcal{E}_\lambda$ is an eigenspace of T then $\mathcal{E}_\lambda \leq V$.

Proof: notice $T - \lambda \cdot Id$ is a linear transformation on V and $\mathcal{E}_\lambda = \text{Ker}(T - \lambda \cdot Id) \leq V$. $\square$

Often I drop the Id and simply write $T - \lambda$ in the place of $T - \lambda \cdot Id$.

Proposition 4.2.5.

If $T : V \rightarrow V$ is a linear transformation and $\mathcal{E}_\lambda$ is an eigenspace of T then $T|_{\mathcal{E}_\lambda} = \lambda Id|_{\mathcal{E}_\lambda}$. Moreover, if β is a basis for $\mathcal{E}_\lambda$ then $[T|_{\mathcal{E}_\lambda}]_{\beta,\beta} = \lambda I$.

Proof: if $w \in \mathcal{E}_\lambda$ then $T(w) = \lambda w = \lambda Id_{\mathcal{E}_\lambda}(w)$ hence $T|_{\mathcal{E}_\lambda} = \lambda Id_{\mathcal{E}_\lambda}$. The fact that $[T|_{\mathcal{E}_\lambda}]_{\beta,\beta} = \lambda I$ follows from the same argument as was given in Proposition 4.2.2. $\square$

Suppose the characteristic polynomial of a given linear transformation completely factors over the given field $\mathbb{F}$. If the algebraic and geometric multiplicities are equal for each eigenvalue of a transformation then the transformation is diagonalizable.

Theorem 4.2.6.

If $T : V \rightarrow V$ is a linear transformation with **distinct** e-values $\lambda_1, \lambda_2, \dots, \lambda_k$ with geometric multiplicities $g_1, g_2, \dots, g_k$ and algebraic multiplicities $a_1, a_2, \dots, a_k$ respective such that $a_j = g_j$ for all $j \in \mathbb{N}_k$ and $a_1 + a_2 + \dots + a_k = \dim(V)$. Then $V = \mathcal{E}_1 \oplus \mathcal{E}_2 \oplus \dots \oplus \mathcal{E}_k$ where $\mathcal{E}_j = \{x \in V \mid T(x) = \lambda_j x\}$. Moreover, the matrix of T with respect to a basis $\beta = \beta_1 \cup \beta_2 \cup \dots \cup \beta_k$ where β_j is basis for $\mathcal{E}_j$ from $j = 1, 2, \dots, k$ is diagonal with:

$$\text{Diag}([T]_{\beta, \beta}) = (\underbrace{\lambda_1, \dots, \lambda_1}_{g_1}, \underbrace{\lambda_2, \dots, \lambda_2}_{g_2}, \dots, \underbrace{\lambda_k, \dots, \lambda_k}_{g_k}).$$

Proof: Suppose the presuppositions of the theorem are true. Since $\dim(\mathcal{E}_j) = g_j$ there exist basis β_j with $|\beta_j| = g_j$ for $j = 1, 2, \dots, k$. Suppose $x_j \in \mathcal{E}_j$ for $j = 1, 2, \dots, k$ and

$$x_1 + x_2 + \dots + x_k = 0$$

If there exists $x_i \neq 0$ then

$$x_i = - \sum_{j \neq i} x_j \in \mathcal{E}_i$$

Let $L = \prod_{l \neq i} (T - \lambda_l)$ notice $T(x_i) = \lambda_i x_i$ and thus $L(x_i) = \prod_{l \neq i} (T - \lambda_l)(x_i) = \prod_{l \neq i} (\lambda_i - \lambda_l)x_i$. Observe $\prod_{l \neq i} (\lambda_i - \lambda_l) \neq 0$ since we assume $\lambda_1, \dots, \lambda_k$ are distinct. Note also,

$$L \left(- \sum_{j \neq i} x_j \right) = - \sum_{j \neq i} \prod_{l \neq i} (T - \lambda_l) x_j$$

Without loss of generality, for convenience of exposition, let us suppose $i = 1$ hence

$$L \left(- \sum_{j \neq i} x_j \right) = - \prod_{l \neq 1} (T - \lambda_l) x_2 - \prod_{l \neq 1} (T - \lambda_l) x_3 - \dots - \prod_{l \neq 1} (T - \lambda_l) x_k$$

But, then, as the operators $T - \lambda_1, \dots, T - \lambda_k$ commute with each other we are free to rearrange the products above such that $T - \lambda_j$ appears as the rightmost factor operating on x_j ,

$$L \left(- \sum_{j \neq i} x_j \right) = - \prod_{l \neq 1, 2} (T - \lambda_l) (T - \lambda_2) x_2 - \prod_{l \neq 1, 3} (T - \lambda_l) (T - \lambda_3) x_3 - \dots - \prod_{l \neq 1, k} (T - \lambda_l) (T - \lambda_k) x_k$$

But, by definition of eigenspace, $(T - \lambda_j)x_j = 0$ for each j hence $L \left(- \sum_{j \neq i} x_j \right) = 0$. But, we previously calculated that $L(x_i) = \prod_{l \neq i} (T - \lambda_l)(x_i) = \prod_{l \neq i} (\lambda_i - \lambda_l)x_i \neq 0$ which is a contradiction. Consequently, $x_i = 0$ for $i = 1, \dots, k$. Therefore, criteria (5.) of Theorem 3.6.4 is met and we find $V = \mathcal{E}_1 \oplus \mathcal{E}_2 \oplus \dots \oplus \mathcal{E}_k$. The remaining claim of the theorem is immediate upon application of Proposition 4.2.2. $\square$

In retrospect, the proof of Proposition 4.1.17 is based on a very similar calculation. Notice we can prove Proposition 4.1.17 by taking the case that all the geometric multiplicities are simply one. I could edit the notes to remove the earlier proof, but I leave it since I think the detail in the previous proof may help better prepare the reader for the proof given in this section.

The geometric multiplicity cannot be larger than the algebraic multiplicity.

Proposition 4.2.7.

If $T : V \rightarrow V$ is a linear transformation with eigenvalue λ with algebraic multiplicity a and geometric multiplicity g then $g \leq a$.

Proof: Suppose g is the geometric multiplicity of λ . Then there exists a basis $\{v_1, v_2, \dots, v_g\}$ for $\mathcal{E}_\lambda \subseteq V$. Extend this to a basis $\beta = \{v_1, \dots, v_g, v_{g+1}, \dots, v_n\}$ for V . Observe,

$$\begin{aligned} T\left(\sum_{i=1}^n x_i v_i\right) &= \sum_{i=1}^g x_i T(v_i) + \sum_{i=g+1}^n x_i T(v_i) \\ &= \sum_{i=1}^g \lambda x_i v_i + \sum_{i=g+1}^n x_i T(v_i). \end{aligned} \quad (4.1)$$

Recall, $[T]_{\beta, \beta} = [[T(v_1)]_\beta] \cdots [[T(v_n)]_\beta]$. Our calculation above implies that first g columns are given as follows:

$$[T]_{\beta, \beta} = [\lambda e_1 | \cdots | \lambda e_g | [T(v_{g+1})]_\beta | \cdots | [T(v_n)]_\beta].$$

Thus, the matrix of T with respect to basis β has the following block-structure:

$$[T]_{\beta, \beta} = \left[\begin{array}{c|c} \lambda I_g & B \\ \hline 0 & C \end{array} \right]$$

We calculate the characteristic polynomial in x by an identity of the determinant: the determinant of an upper-block-triangular matrix is the product of the determinants of the blocks on the diagonal

$$\det([T]_{\beta, \beta} - xI) = \det(\lambda I_g - xI_g) \det(C - xI_{n-g}) = (\lambda - x)^g \det(C - xI_{n-g}).$$

Thus there are at least g factors of $(x - \lambda)$ in $p(x)$ hence $a \geq g$. $\square$

Theorem 4.1.3 implies that for each eigenvalue λ there exists at least one eigenvector $v \in \text{Null}(T - \lambda Id)$. This fact together with the proposition above shows that for each eigenvalue λ_j of a linear transformation T we have $1 \leq g_j \leq a_j$. We saw in the previous section that if T is diagonalizable then the basis for which $[T]_{\beta, \beta}$ is diagonal is an eigenbasis for T . Conversely, if T has an eigenbasis then T is diagonalizable. The calculations and theory of this section show that T is diagonalizable if and only if the characteristic polynomial completely factors and each eigenspace has dimension which matches the algebraic multiplicity of the eigenvalue. In the section after next we deal with the case that the geometric multiplicity is smaller than the algebraic multiplicity. In other words, we'll deal with the non-diagonalizable case in the future section on Jordan forms. But, first, we should look at some of the theoretical developments which underly the Jordan form.

4.3 Invariant Subspaces and the Cayley Hamilton Theorem

I am following parts of §5.4 of Friedberg, Insel and Spence's *Linear Algebra*.

Definition 4.3.1.

Let $T : V \rightarrow V$ be a linear transformation. A subspace $W \leq V$ is a **T -invariant** subspace if $T(W) \subseteq W$. For such T we define $T_W : W \rightarrow W$ by $T_W(x) = T(x)$ for each $x \in W$.

Notice the **restriction** of a linear map to a subspace is necessarily a linear map. In particular, if $S : V \rightarrow V$ is a linear transformation and $U \leq V$ then $S|_U : U \rightarrow V$ is a linear map given by $S|_U(x) = S(x)$ for each $x \in U$. Notice that $T_W : W \rightarrow W$ is a linear transformation on W . We can think of T_W as the restriction to W where the codomain V has been replaced by W . This construction is only reasonable if the map is T -invariant. In fact, there are many such cases:

$$0, \quad V, \quad \text{Ker}(T), \quad T(V)$$

are all T -invariant subspaces for a linear map on V .

An interesting way to create a T -invariant subspace which contains a given vector $x \in V$ is to consider

$$W = \text{span}\{x, T(x), T^2(x), T^3(x), \dots, T^k(x), \dots\}$$

where $T^k(x) = T(T^{k-1}(x))$ for each $k \in \mathbb{N}$. This subspace of V is known as the T -cyclic subspace generated by x .

Proposition 4.3.2.

Let T be a linear map on V and $x \in V$. Then

$$W = \text{span}\{T^k(x) \mid k \in \mathbb{N} \cup \{0\}\} = \langle x \rangle$$

is a T -invariant subspace of V . We call $\langle x \rangle$ the T -cyclic subspace generated by x under T .

Proof: Note W is a subspace since it is formed by a span. If $y \in W$ then there exist $c_k \in \mathbb{F}$, with $c_k \neq 0$ for only finitely many choices of k and

$$y = \sum_{k=0}^{\infty} c_k T^k(x)$$

Operate on y by T to obtain

$$T(y) = T\left(\sum_{k=0}^{\infty} c_k T^k(x)\right) = \sum_{k=0}^{\infty} c_k T(T^k(x)) = \sum_{k=0}^{\infty} c_k T^{k+1}(x) \in W$$

since $T(y)$ is a finite linear combination of powers of T acting on x . $\square$

We could ask, given a subspace $W \leq V$, does there exist $x \in V$ for which $W = \langle x \rangle$? What do you think, is it always possible? Certainly if we pick the wrong x , like say $x = 0$ then the subspace generated by x is merely the zero subspace; $\langle 0 \rangle = 0$. What about $x \neq 0$? Notice that $\{x, T(x), \dots\}$ then contains at least one nonzero vector hence $\dim \langle x \rangle \geq 1$. It seems like there is some hope to express a subspace in the form $W = \langle x \rangle$. However, not every subspace is T -invariant. Therefore, not every subspace is T -cyclic. That said, the following proposition gives many T -invariant subspaces including eigenspaces and generalized eigenspaces.

Proposition 4.3.3.

Let $f(t) \in \mathbb{F}[t]$ be a polynomial with coefficients in the field $\mathbb{F}$. If $T : V \rightarrow V$ is a linear map then $\text{Ker}(f(T))$ is a T -invariant subspace.

Proof: suppose $f(x) = a_0 + a_1x + \cdots + a_nx^n$ then $f(T) = a_0 + a_1T + \cdots + a_nT^n$. Notice $f(T)$ is a linear transformation and $\text{Ker}(f(T))$ is a subspace since the kernel of a linear map is a subspace. It remains to prove $\text{Ker}(f(T))$ is a T -invariant subspace. Suppose $x \in \text{Ker}(f(T))$. Notice that $f(T)T = Tf(T)$ since T commutes with powers of T and scalar multiplication. Here multiplication of operators is understood as repeated composition. Hence,

$$f(T)(T(x)) = T(f(T)(x)) = T(0) = 0$$

Therefore, $T(x) \in \text{Ker}(f(T))$ and this proves $T(\text{Ker}(f(T))) \subseteq \text{Ker}(f(T))$. $\square$

There are many applications of the above proposition. Notice that eigenspace $\mathcal{E}_\lambda = \text{Ker}(T - \lambda)$ is T -invariant since it is of the form $\text{Ker}(f(T))$ where $f(x) = x - \lambda$.

Theorem 4.3.4.

Let T be a linear map on a finite dimensional vector space V with invariant subspace W . The characteristic polynomial of T_W divides the characteristic polynomial of T .

Proof: let β_W be a basis for W . Let $\dim(W) = k$ and $\dim(V) = n$. Extend β_W by adjoining vectors in $\beta' \subseteq V - W$ such that $\beta = \beta_W \cup \beta'$ is a basis for V . Since $T(W) \subseteq W$ we find $T(\beta_W) \subseteq \beta_W$ thus the matrix for T has the form

$$[T]_{\beta, \beta} = \begin{bmatrix} A & B \\ 0 & C \end{bmatrix}$$

Moreover, by construction, $[T]_{\beta_W, \beta_W} = A$. Let $P(t)$ and $P_W(t)$ be the characteristic polynomials of T and T_W respective. Observe,

$$P(t) = \det \begin{bmatrix} A - tI_k & B \\ 0 & C - tI_{n-k} \end{bmatrix} = \det(A - tI_k) \det(C - tI_{n-k}) = P_W(t) \det(C - tI_{n-k}). \quad \square$$

What follows is Theorem 5.21 from §5.4 of Friedberg, Insel and Spence's *Linear Algebra*.

Theorem 4.3.5.

Let T be a linear map on a finite dimensional vector space V , and let $W = \langle x \rangle$ where $\dim(W) = k$. Then,

- (a.) $\{x, T(x), \dots, T^{k-1}(x)\}$ is a basis for W
- (b.) If $a_0x + a_1T(x) + \cdots + a_{k-1}T^{k-1}(x) + T^k(x) = 0$ then the characteristic polynomial of T_W is given by $P_W(t) = (-1)^k(a_0 + a_1t + \cdots + a_{k-1}t^{k-1} + t^k)$.

Proof: see §5.4 of Friedberg, Insel and Spence's *Linear Algebra*. The argument for (a.) is neat. $\square$

Example 6 see §5.4 of Friedberg, Insel and Spence's *Linear Algebra* on page 315 illustrates how the above theorem allows for calculation of characteristic polynomials without use of determinants. This is a curious calculation, but certainly determinants are a far more clear path in most instances.

Theorem 4.3.6. *Cayley-Hamilton Theorem*

Let T be a linear map on a finite dimensional vector space $V(\mathbb{F})$ with characteristic polynomial $P(t)$. Then $P(T) = 0$.

Proof: if $x = 0$ then $P(T)(x) = 0$ since $P(T)$ is a linear map. Suppose $x \neq 0$. Let $W = \langle x \rangle$ with $\dim(W) = k$. By part (a.) of Theorem 4.3.5 there exist $a_0, a_1, \dots, a_{k-1} \in \mathbb{F}$ for which

$$a_0x + a_1T(x) + \dots + a_{k-1}T^{k-1}(x) + T^k(x) = 0.$$

Moreover, the characteristic polynomial of T_W has form

$$P_W(t) = (-1)^k(a_0 + a_1t + \dots + a_{k-1}t^{k-1} + t^k).$$

Thus,

$$P_W(T)(x) = (-1)^k(a_0 + a_1T + \dots + a_{k-1}T^{k-1} + T^k)(x) = (-1)^k(a_0x + a_1T(x) + \dots + a_{k-1}T^{k-1}(x) + T^k(x)) = 0.$$

Theorem 4.3.4 implies $P(T) = f(T)P_W(T)$ for some polynomial $f(t)$. Thus,

$$P(T)(x) = f(T)(P_W(T)(x)) = f(T)(0) = 0$$

Since $P(T)(x) = 0$ for all $x \in V$ we find $P(T) = 0$. $\square$

4.4 Theory of the Jordan Form

Friedberg, Insel and Spence develop the full theory of the Jordan form in Chapter 7 of *Linear Algebra*. I merely quote a few especially interesting high points of their development with a focus on their use of polynomial arguments. I think the role polynomial algebra plays is somewhat surprising and it brings some new calculational methods which may be new to most students. That said, if you wish the full logical development then please read the text carefully.

Definition 4.4.1.

A **generalized eigenspace** of eigenvalue λ for a linear transformation $T : V \rightarrow V$ is denoted K_λ . We define $x \in K_\lambda$ if there exists a positive integer k such that

$$(T - \lambda)^k x = 0.$$

Theorem 7.1 from page 478 of Friedberg, Insel and Spence's 5-th Ed. of *Linear Algebra* states:

Theorem 4.4.2.

Let T be a linear operator on a vector space V , and let λ be an eigenvalue of T . Then

- (a.) K_λ is T -invariant subspace of V and $\mathcal{E}_\lambda \leq K_\lambda$,
- (b.) For any eigenvalue μ of T such that $\mu \neq \lambda$ we find $K_\lambda \cap K_\mu = \{0\}$.
- (c.) For any scalar $\mu \neq \lambda$, the restriction of $T - \mu I$ to K_λ is one-to-one and onto.

See page 478-479 for the proof. Part (b.) is particularly interesting in the use of polynomial algebra. I would like to discuss that part in lecture.

Theorem 4.4.3.

Let T be a linear operator on a finite dimensional vector space V such that the characteristic polynomial splits. Suppose that λ is an eigenvalue with algebraic multiplicity m . Then

- (a.) $\dim(K_\lambda) \leq m$,
- (b.) $K_\lambda = \text{Ker}((T - \lambda I)^m)$.

See page 479 for the proof, it relies in part on the Cayley Hamilton Theorem.

Theorem 4.4.4.

Let T be a linear operator on a finite dimensional vector space V such that the characteristic polynomial splits, and let $\lambda_1, \lambda_2, \dots, \lambda_k$ be the distinct eigenvalues of T . Then, for every $x \in V$, there exist unique vectors $v_i \in K_{\lambda_i}$, for $i = 1, 2, \dots, k$, such that

$$x = v_1 + v_2 + \dots + v_k.$$

Proof is on page 480 and it involves interesting polynomial algebra once more.

Theorem 4.4.5.

Let T be a linear operator on a finite dimensional vector space V such that the characteristic polynomial $(t - \lambda_1)^{m_1}(t - \lambda_2)^{m_2} \dots (t - \lambda_k)^{m_k}$ splits. For $i = 1, 2, \dots, k$ let β_i be an ordered basis for K_{λ_i} . Then

- (a.) $\beta_i \cap \beta_j = \emptyset$ for $i \neq j$.
- (b.) $\beta = \beta_1 \cup \beta_2 \cup \dots \cup \beta_k$ is an ordered basis for V .
- (c.) $\dim(K_{\lambda_i}) = m_i$ for all i .

The following appear as a Corollary to the result above:

Corollary 4.4.6.

Let T be a linear operator on a finite dimensional vector space V such that the characteristic polynomial of T splits. Then T is diagonalizable if and only if $\mathcal{E}_\lambda = K_\lambda$ for every eigenvalue of T .

What remains of §7.1 in Friedberg, Insel and Spence's 5-th Ed. of *Linear Algebra* is covered in some sense by the next section. I suspect it would be wise to read these notes first then read §7.1 and §7.2 to get some sense of how my assertions here are proved. I'll try to present some of the proofs in lecture, but we will not have time for the full body of the argument. The dot-diagrams in §7.2 give a strategy for calculation which I don't really explain in these notes. On the other hand, Friedberg, Insel and Spence's 5-th Ed. of *Linear Algebra* does not explain the idea of complexification or the real Jordan form as I do here. It's best to read both.

4.5 generalized eigenvectors and Jordan blocks

We begin again with the definition as it applies to a linear transformation.

Definition 4.5.1.

A **generalized eigenvector** of order k with eigenvalue λ with respect to a linear transformation $T : V \rightarrow V$ is a nonzero vector v such that

$$(T - \lambda Id)^k v = 0 \quad \& \quad (T - \lambda Id)^{k-1} v \neq 0.$$

The existence of a generalized eigenvector of order k with eigenvalue λ implies the null space $\text{Null}[(T - \lambda Id)^{k-1}] \neq 0$. However, if $k \geq 2$, this also implies $\text{Null}[(T - \lambda Id)^{k-2}] \neq 0$. Indeed, if there exists a single generalized eigenvector of order k it follows that:

$$(T - \lambda Id)^{k-1}, (T - \lambda Id)^{k-2}, \dots, T - \lambda Id$$

all have nontrivial null spaces. This claim is left to the reader as an exercise. If you would like more complete exposition of this topic you can read Insel Spence and Friedberg. I am trying to get to the point without too much detail here.

Definition 4.5.2.

A **k -chain with eigenvalue λ** of a linear transformation $T : V \rightarrow V$ is set of k nonzero vectors $v_1, v_2, \dots, v_k$ such that $(T - \lambda Id)(v_j) = v_{j-1}$ for $j = 2, \dots, k$ and v_1 is an eigenvector with eigenvalue λ ; $(T - \lambda Id)(v_1) = 0$.

Of course, the reason we care about the chain is what follows:

Theorem 4.5.3.

A k -chain with e-value λ for $T : V \rightarrow V$ is a set of LI generalized e-vectors order $1, \dots, k$.

Proof: Let $\{v_1, \dots, v_k\}$ be a k -chain with e-value λ for T . By definition $(T - \lambda Id)(v_1) = 0$. Consider:

$$(T - \lambda Id)(v_2) = v_1 \Rightarrow (T - \lambda Id)^2(v_2) = (T - \lambda Id)(v_1) = 0.$$

Thus v_2 is a generalized e-vector of order 2. Next, observe

$$(T - \lambda Id)(v_3) = v_2 \Rightarrow (T - \lambda Id)^3(v_3) = (T - \lambda Id)^2(v_2) = 0.$$

Thus v_3 is a generalized e-vector of order 3. We continue in this fashion until we reach the k -th vector in the chain:

$$(T - \lambda Id)(v_k) = v_{k-1} \Rightarrow (T - \lambda Id)^k(v_k) = (T - \lambda Id)^{k-1}(v_{k-1}) = 0.$$

Thus v_k is a generalized e-vector of order k . To prove LI of the chain suppose that:

$$c_1 v_1 + c_2 v_2 + \dots + c_k v_k = 0.$$

Operate successively by $(T - \lambda Id)^j$ for $j = k-1, k-2, \dots, 2, 1$ to derive first $c_k = 0$ then $c_{k-1} = 0$ then continuing until we reach $c_2 = 0$ and finally $c_1 = 0$. $\square$.

It turns out that we can always choose generalized eigenvectors such that they line-up into chains. The details of the proof of the theorem that follow can be found in Insel Spence and Friedberg's *Linear Algebra* and most graduate linear algebra texts. They introduce an organizational tool known as **dot-diagrams** to see how to arrange the chains. We can go a long way by just finding e-vectors and building chains from there, but, there is more to explain about chains.

Theorem 4.5.4. *Jordan basis theorem*

Let V be a vector space over $\mathbb{F}$. If $T : V \rightarrow V$ is a linear transformation with eigenvalues in $\mathbb{F}$ then there exists a basis for V formed from chains of generalized e-vectors. Such a basis is a **Jordan basis**. Moreover, up to ordering of the chains, the matrix of T is unique and is called the **Jordan form** of T

Proof: see Chapter 7 of Insel Spence and Friedberg's fifth edition of *Linear Algebra*. $\square$

The matrix of T with respect to a Jordan basis will be block-diagonal and each block will be a **Jordan block**. For brevity of exposition⁴ consider $T : V \rightarrow V$ which has a single k -chain as it a basis for V , $\beta = \{v_1, v_2, \dots, v_k\}$ is a k -chain with e-value λ for T :

$$T(v_1) = \lambda v_1, \quad T(v_2) = \lambda v_2 + v_1, \quad \dots, \quad T(v_k) = \lambda v_k + v_{k-1}$$

Thus the matrix of T has the form:

$$[T]_{\beta, \beta} = \begin{bmatrix} \lambda & 1 & 0 & \cdots & 0 & 0 \\ 0 & \lambda & 1 & \cdots & 0 & 0 \\ \vdots & \vdots & \vdots & \cdots & \vdots & \vdots \\ 0 & 0 & 0 & \cdots & \lambda & 1 \\ 0 & 0 & 0 & \cdots & 0 & \lambda \end{bmatrix}$$

To be clear, all the diagonal entries are λ and there is a string of 1's along the superdiagonal. All other entries are zero. In some other texts, for example Hefferon, it should be noted the Jordan block has 1's right below the diagonal. This stems from a different formulation of the chains.

Definition 4.5.5.

Let $N = E_{12} + E_{23} + \cdots + E_{d-1,d} \in \mathbb{F}^{d \times d}$ be the matrix which is everywhere zero except where it is one on its superdiagonal. We define the $d \times d$ -Jordan block by $J_d(\lambda) = \lambda_d I + N$. A matrix $J \in \mathbb{F}^{n \times n}$ is said to be in **Jordan Form** if it is a block-diagonal matrix with Jordan blocks on the diagonal; $J = J_{d_1}(\lambda_1) \oplus \cdots \oplus J_{d_k}(\lambda_k)$.

To be explicit,

$$N = E_{12} + E_{23} + \cdots + E_{d-1,d} = \begin{bmatrix} 0 & 1 & 0 & 0 & \cdots & 0 \\ 0 & 0 & 1 & 0 & \cdots & 0 \\ \vdots & \vdots & \vdots & \vdots & \cdots & \vdots \\ 0 & 0 & 0 & 0 & \cdots & 1 \\ 0 & 0 & 0 & 0 & \cdots & 0 \end{bmatrix}$$

and you can verify $N^d = 0$ yet $N^{d-1} = E_{1d} \neq 0$. This is the quintessential **nilpotent** matrix of order d . Let me give a few examples before we continue our study of chains. All of the matrices below are in Jordan form⁵:

$$A_1 = \begin{bmatrix} 1+2i & 0 & 0 & 0 \\ 0 & 3+4i & 0 & 0 \\ 0 & 0 & 5+6i & 0 \\ 0 & 0 & 0 & 7+8i \end{bmatrix} \quad A_2 = \begin{bmatrix} 2 & 1 & 0 & 0 \\ 0 & 2 & 0 & 0 \\ 0 & 0 & 3 & 0 \\ 0 & 0 & 0 & 4 \end{bmatrix}$$

⁴you doubt this?

⁵a good exercise is formulating these directly in terms of the block-notation; for instance $A_4 = J_2(0) \oplus J_1(2)$.

$$A_3 = \begin{bmatrix} 6 & 1 & 0 \\ 0 & 6 & 1 \\ 0 & 0 & 6 \end{bmatrix} \quad A_4 = \begin{bmatrix} 0 & 1 & 0 \\ 0 & 0 & 0 \\ 0 & 0 & 2 \end{bmatrix}, \quad A_5 = \text{diag}(A_1, A_2, A_3, A_4) \in \mathbb{C}^{14 \times 14}.$$

In particular, every diagonal matrix is in Jordan form where the Jordan-Blocks are all 1×1 . Generally, the Jordan form is the next best thing to a diagonal matrix. In fact, we could prove A_2, A_3, A_4 and A_5 are not diagonalizable. Now we return to our study of chains.

Perhaps you wonder why even look at chains? Of course, the Jordan basis theorem is reason enough, but another reason is that they appear somewhat naturally in differential equations. Let's examine how in a simple example.

Example 4.5.6. Consider $T = D$ on $P_2 = \text{span}\{1, x, x^2\}$. Clearly $T(1) = 0$ hence $v_1 = 1$ is an eigenvector with eigenvalue $\lambda = 0$ for T . Furthermore, as $T(x) = 1$ and $T(x^2) = 2x$ it follows

$$[T]_{\beta, \beta} = \begin{bmatrix} 0 & 1 & 0 \\ 0 & 0 & 2 \\ 0 & 0 & 0 \end{bmatrix}. \text{ Thus } T \text{ has only zero as an e-value and its algebraic multiplicity is three.}$$

If we consider $\gamma = \{1, x, x^2/2\}$ then this is a 3-chain with e-value $\lambda = 0$. Note:

$$T(1) = 0, T(x) = 1, T(x^2/2) = x \Rightarrow [T]_{\gamma, \gamma} = \begin{bmatrix} 0 & 1 & 0 \\ 0 & 0 & 1 \\ 0 & 0 & 0 \end{bmatrix}.$$

There are more exciting reasons attached to the study of the *matrix exponential*, see Chapter 4.10. It's *deja vu* all over again.

Definition 4.5.7.

A **generalized eigenvector** of order k with eigenvalue $\lambda \in \mathbb{F}$ with respect to a $A \in \mathbb{F}^{n \times n}$ is a nonzero vector v such that

$$(A - \lambda I)^k v = 0 \quad \& \quad (A - \lambda I)^{k-1} v \neq 0.$$

Naturally, the chains are also of interest in the matrix case:

Definition 4.5.8.

A **k -chain with eigenvalue λ** of $A \in \mathbb{F}^{n \times n}$ is a set of k nonzero vectors $v_1, v_2, \dots, v_k$ such that $(A - \lambda I)v_j = v_{j-1}$ for $j = 1, 2, \dots, k$ and v_1 is an eigenvector with eigenvalue λ ; $(A - \lambda I)v_1 = 0$.

The analog of Theorem 4.5.3 is true for the matrix case. However, perhaps this special case with the contradiction-based proof will add some insight for the reader.

Proposition 4.5.9.

Suppose $A \in \mathbb{F}^{n \times n}$ has e-value λ and e-vector v_1 then if $(A - \lambda I)v_2 = v_1$ has a solution then v_2 is a generalized e-vector of order 2 which is linearly independent from v_1 .

Proof: Suppose $(A - \lambda I)v_2 = v_1$ is consistent then multiply by $(A - \lambda I)$ to find $(A - \lambda I)^2 v_2 = (A - \lambda I)v_1$. However, we assumed v_1 was an e-vector hence $(A - \lambda I)v_1 = 0$ and we find v_2 is a generalized e-vector of order 2. Suppose $v_2 = kv_1$ for some nonzero k then $Av_2 = Akv_1 = k\lambda v_1 = \lambda v_2$ hence $(A - \lambda I)v_2 = 0$ but this contradicts the construction of v_2 as the solution to $(A - \lambda I)v_2 = v_1$. Consequently, v_2 is linearly independent from v_1 by virtue of its construction. $\square$.

Example 4.5.10. Let's return to Example 4.8.8 and look for a generalized e -vector of order 2. Recall $A = \begin{bmatrix} 1 & 1 \\ 0 & 1 \end{bmatrix}$ and we found a repeated e -value of $\lambda_1 = 1$ and single e -vector $u_1 = \begin{bmatrix} 1 \\ 0 \end{bmatrix}$ (fix $u = 1$ for convenience). Let's complete the chain: find $v_2 = [u, v]^T$ such that $(A - I)u_2 = u_1$,

$$\begin{bmatrix} 0 & 1 \\ 0 & 0 \end{bmatrix} \begin{bmatrix} u \\ v \end{bmatrix} = \begin{bmatrix} 1 \\ 0 \end{bmatrix} \Rightarrow v = 1, u \text{ is free}$$

Any choice of u will do, in this case we can even set $u = 0$ to find

$$u_2 = \begin{bmatrix} 0 \\ 1 \end{bmatrix}$$

Notice, $\{u_1, u_2\}$ is not an eigenbasis for A , however, it is a **Jordan basis** for A .

Theorem 4.5.11.

Any matrix $A \in \mathbb{F}^{n \times n}$ with eigenvalues all in $\mathbb{F}$ can be transformed to Jordan form J by a similarity transformation based on conjugation by the matrix $[\beta]$ of a Jordan basis β . That is, there exists Jordan basis β for $\mathbb{F}^n$ such that $[\beta]^{-1}A[\beta] = J$

Proof: apply Theorem 4.5.4 to the linear transformation $T = L_A : \mathbb{R}^n \rightarrow \mathbb{R}^n$. $\square$

The nicest examples are those which are already in Jordan form at the beginning:

Example 4.5.12. Suppose $A = \begin{bmatrix} 1 & 1 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 1 \\ 0 & 0 & 0 & 1 \end{bmatrix}$ it is not hard to show that $\det(A - \lambda I) = (\lambda - 1)^4 = 0$. We have a quadruple e -value $\lambda_1 = \lambda_2 = \lambda_3 = \lambda_4 = 1$.

$$0 = (A - I)\vec{u} = \begin{bmatrix} 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 \\ 0 & 0 & 0 & 0 \end{bmatrix} \vec{u} \Rightarrow \vec{u} = \begin{bmatrix} s_1 \\ 0 \\ s_3 \\ 0 \end{bmatrix}$$

Any nonzero choice of s_1 or s_3 gives us an e -vector. Let's define two e -vectors which are clearly linearly independent, $\vec{u}_1 = [1, 0, 0, 0]^T$ and $\vec{u}_2 = [0, 0, 1, 0]^T$. We'll find a generalized e -vector to go with each of these. There are two length two chains to find here. In particular,

$$(A - I)\vec{u}_3 = \vec{u}_1 \Rightarrow \begin{bmatrix} 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 \\ 0 & 0 & 0 & 0 \end{bmatrix} \begin{bmatrix} s_1 \\ s_2 \\ s_3 \\ s_4 \end{bmatrix} = \begin{bmatrix} 1 \\ 0 \\ 0 \\ 0 \end{bmatrix} \Rightarrow s_2 = 1, s_4 = 0, s_1, s_3 \text{ free}$$

I choose $s_1 = 0$ and $s_3 = 1$ since I want a new vector, define $\vec{u}_3 = [0, 0, 1, 0]^T$. Finally solving $(A - I)\vec{u}_4 = \vec{u}_2$ for $\vec{u}_4 = [s_1, s_2, s_3, s_4]^T$ yields conditions $s_4 = 1, s_2 = 0$ and s_1, s_3 free. I choose $\vec{u}_4 = [0, 0, 0, 1]^T$. To summarize we have four linearly independent vectors which form two chains:

$$(A - I)\vec{u}_3 = \vec{u}_1, (A - I)\vec{u}_1 = 0 \quad (A - I)\vec{u}_4 = \vec{u}_2, (A - I)\vec{u}_2 = 0$$

The matrix above was in an example of a matrix in Jordan form. When the matrix is in Jordan form then the each element of then standard basis is an e-vector or generalized e-vector.

Example 4.5.13.

$$A = \begin{bmatrix} 2 & 1 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 2 & 1 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 2 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 3 & 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 3 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 3 & 1 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 3 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 4 \end{bmatrix}$$

Here we have the chain $\{e_1, e_2, e_3\}$ with e-value $\lambda_1 = 2$, the chain $\{e_4, e_5, e_6, e_7\}$ with e-value $\lambda_2 = 3$ and finally a lone e-vector e_8 with e-value $\lambda_3 = 4$

I do not attempt to give a full account of the calculational method to compute the Jordan form, but I will give some further guidelines based on a simple observation:

$$\text{Ker}(S) \leq \text{Ker}(S^2) \leq \text{Ker}(S^3) \leq \dots$$

In particular, if we apply the general result above to the case $S = T - \lambda$ we have

$$\text{Ker}(T - \lambda) \leq \text{Ker}(T - \lambda)^2 \leq \text{Ker}(T - \lambda)^3 \leq \dots$$

This means $\mathcal{E}_\lambda \leq \text{Ker}(T - \lambda)^2 \leq \dots \leq K_\lambda$. Notice that K_λ includes all eigenvectors, generalized eigenvectors of order two, three, etc. In short, K_λ contains all possible k -chains of eigenvalue λ for T . Some things cannot happen. For instance, we cannot have $\dim(\text{Ker}(T - \lambda)^j) < \dim(\text{Ker}(T - \lambda)^{j+1})$. The dimensions of the kernels of descending powers of $T - \lambda$ must decrease.

Example 4.5.14. Suppose $T : V \rightarrow V$ is a linear transformation with

$$\begin{aligned} \dim(\text{Ker}(T - 3)) &= 2, \\ \dim(\text{Ker}(T - 3)^2) &= 4, \\ \dim(\text{Ker}(T - 3)^3) &= 5, \\ \dim(\text{Ker}(T - 3)^4) &= 6, \\ \dim(\text{Ker}(T - 3)^5) &= 6. \end{aligned}$$

This tells me there are 2-chains of generalized eigenvectors with $\lambda = 3$. One of these chains is a 2-chain, the other is apparently a 4-chain. If $\{v_1, v_2, v_3, v_4\}$ is the 4-chain and $\{w_1, w_2\}$ is the 2-chain then

$$\begin{aligned} \mathcal{E}_3 &= \text{Ker}(T - 3) = \text{span}\{w_1, v_1\}, \\ \text{Ker}((T - 3)^2) &= \text{span}\{w_1, w_2, v_1, v_2\}, \\ \text{Ker}((T - 3)^3) &= \text{span}\{w_1, w_2, v_1, v_2, v_3\}, \\ K_3 &= \text{Ker}((T - 3)^4) = \text{span}\{w_1, w_2, v_1, v_2, v_3, v_4\}, \\ \text{Ker}((T - 3)^5) &= \text{span}\{w_1, w_2, v_1, v_2, v_3, v_4\}, \end{aligned}$$

Continuing, suppose

$$\begin{aligned} \dim(\text{Ker}(T + 1)) &= 3, \\ \dim(\text{Ker}(T + 1)^2) &= 6, \\ \dim(\text{Ker}(T + 1)^3) &= 8, \\ \dim(\text{Ker}(T + 1)^4) &= 9, \\ \dim(\text{Ker}(T + 1)^5) &= 10. \end{aligned}$$

From this I deduce there are three $\lambda = -1$ chains of lengths 2, 3 and 5. Using $\{x_1, x_2\}$, $\{y_1, y_2, y_3\}$ and $\{z_1, z_2, z_3, z_4, z_5\}$ to denote these chains,

$$\begin{aligned} \mathcal{E}_{-1} &= \text{Ker}(T + 1) = \text{span}\{x_1, y_1, z_1\}, \\ \text{Ker}((T + 1)^2) &= \text{span}\{x_1, x_2, y_1, y_2, z_1, z_2\}, \\ \text{Ker}((T + 1)^3) &= \text{span}\{x_1, x_2, y_1, y_2, y_3, z_1, z_2, z_3\}, \\ \text{Ker}((T + 1)^4) &= \text{span}\{x_1, x_2, y_1, y_2, y_3, z_1, z_2, z_3, z_4\}, \\ K_{-1} &= \text{Ker}((T + 1)^5) = \text{span}\{x_1, x_2, y_1, y_2, y_3, z_1, z_2, z_3, z_4, z_5\} \end{aligned}$$

Without further information, for all we know, I am wrong and K_{-1} includes further vectors. However, if we know $\dim(V) = 16$ then we can be sure that K_{-1} is as claimed above and $V = K_3 \oplus K_{-1}$. Moreover,

$$\beta = \{w_1, w_2, v_1, v_2, v_3, v_4, x_1, x_2, y_1, y_2, y_3, z_1, z_2, z_3, z_4, z_5\}$$

is a Jordan basis for which

$$[T]_{\beta, \beta} = J_2(3) \oplus J_4(3) \oplus J_2(-1) \oplus J_3(-1) \oplus J_5(-1).$$

Example 4.5.15. Suppose $\dim(V) = 9$ and the linear transformation $T : V \rightarrow V$ has

$$\begin{aligned} \dim(\text{Ker}(T - 7)) &= 7, \\ \dim(\text{Ker}(T - 7)^2) &= 8, \\ \dim(\text{Ker}(T - 7)^3) &= 9 \end{aligned}$$

Then there exists β for which $[T]_{\beta} = \text{Diag}(7, 7, 7, 7, 7, 7) \oplus J_3(7)$.

Example 4.5.16. Suppose the linear transformation $T : V \rightarrow V$ has

$$\begin{aligned} \dim(\text{Ker}(T - 4)) &= 3, \\ \dim(\text{Ker}(T - 4)^2) &= 2. \end{aligned}$$

This is impossible. Since $\text{Ker}(T - 4) \leq \text{Ker}((T - 4)^2)$ we find a 3-dimensional subspace of a 2-dimensional space. This cannot be.

Example 4.5.17. Suppose $\dim(V) = 12$ and the linear transformation $T : V \rightarrow V$ has

$$\begin{aligned} \dim(\text{Ker}(T - 2)) &= 4, \\ \dim(\text{Ker}(T - 2)^2) &= 8, \\ \dim(\text{Ker}(T - 2)^3) &= 12 \end{aligned}$$

Then there exists β for which $[T]_{\beta} = J_3(2) \oplus J_3(2) \oplus J_3(2) \oplus J_3(2)$.

4.6 complexification

In this section we study how a given real vector space is naturally extended to a vector space over $\mathbb{C}$. It is interesting to note the construction of the complexification of V as a particular structure on $V \times V$ is the same in essence as Gauss' construction of the complex numbers from $\mathbb{R}^2$.

4.6.1 concerning matrices and vectors with complex entries

To begin, we denote the complex numbers by $\mathbb{C}$. As a two-dimensional real vector space we can decompose the complex numbers into the direct sum of the real and pure-imaginary numbers; $\mathbb{C} = \mathbb{R} \oplus i\mathbb{R}$. In other words, any complex number $z \in \mathbb{C}$ can be written as $z = a + ib$ where $a, b \in \mathbb{R}$. It is convenient to define

$$\boxed{\text{If } \lambda = \alpha + i\beta \in \mathbb{C} \text{ for } \alpha, \beta \in \mathbb{R} \text{ then } \operatorname{Re}(\lambda) = \alpha, \operatorname{Im}(\lambda) = \beta}$$

The projections onto the real or imaginary part of a complex number are actually linear transformations from $\mathbb{C}$ to $\mathbb{R}$; $\operatorname{Re} : \mathbb{C} \rightarrow \mathbb{R}$ and $\operatorname{Im} : \mathbb{C} \rightarrow \mathbb{R}$. Next, complex vectors are simply n -tuples of complex numbers:

$$\boxed{\mathbb{C}^n = \{(z_1, z_2, \dots, z_n) \mid z_j \in \mathbb{C}\}}.$$

Definitions of scalar multiplication and vector addition follow the obvious rules: if $z, w \in \mathbb{C}^n$ and $c \in \mathbb{C}$ then

$$(z + w)_j = z_j + w_j \quad (cz)_j = cz_j$$

for each $j = 1, 2, \dots, n$. The complex n -space is again naturally decomposed into the direct sum of two n -dimensional real spaces; $\mathbb{C}^n = \mathbb{R}^n \oplus i\mathbb{R}^n$. In particular, any complex n -vector can be written uniquely as the sum of real vectors are known as the real and imaginary vector components:

$$\boxed{\text{If } v = a + ib \in \mathbb{C}^n \text{ for } a, b \in \mathbb{R}^n \text{ then } \operatorname{Re}(v) = a, \operatorname{Im}(v) = b.}$$

Recall $z = x + iy \in \mathbb{C}$ has complex conjugate $z^* = x - iy$. Let $v \in \mathbb{C}^n$ we define the complex conjugate of the vector v to be v^* which is the vector of complex conjugates;

$$(v^*)_j = (v_j)^*$$

for each $j = 1, 2, \dots, n$. For example, $[1 + i, 2, 3 - i]^* = [1 - i, 2, 3 + i]$. It is easy to verify the following properties for complex conjugation of numbers and vectors:

$$(v + w)^* = v^* + w^*, \quad (cv)^* = c^* v^*, \quad v^{**} = v.$$

Complex matrices $\mathbb{C}^{m \times n}$ can be added, subtracted, multiplied and scalar multiplied in precisely the same ways as real matrices in $\mathbb{R}^{m \times n}$. However, we can also identify them as $\mathbb{C}^{m \times n} = \mathbb{R}^{m \times n} \oplus i\mathbb{R}^{m \times n}$ via the real and imaginary part maps $(\operatorname{Re}(Z))_{ij} = \operatorname{Re}(Z_{ij})$ and $(\operatorname{Im}(Z))_{ij} = \operatorname{Im}(Z_{ij})$ for all i, j . There is an obvious isomorphism $\mathbb{C}^{m \times n} \cong \mathbb{R}^{2m \times 2n}$ which follows from stringing out all the real and imaginary parts. Again, complex conjugation is also defined component-wise: $((X + iY)^*)_{ij} = X_{ij} - iY_{ij}$. It's easy to verify that

$$(Z + W)^* = Z^* + W^*, \quad (cZ)^* = c^* Z^*, \quad (ZW)^* = Z^* W^*$$

for appropriately sized complex matrices Z, W and $c \in \mathbb{C}$. Conjugation gives us a natural operation to characterize the *reality* of a variable. Let $c \in \mathbb{C}$ then c is **real** iff $c^* = c$. Likewise, if $v \in \mathbb{C}^n$ then we say that v is **real** iff $v^* = v$. If $Z \in \mathbb{C}^{m \times n}$ then we say that Z is **real** iff $Z^* = Z$. In short, an object is real if all its imaginary components are zero. Finally, while there is of course much more to say we will stop here for now.

4.6.2 the complexification

Suppose V is a vector space over $\mathbb{R}$, we seek to construct a new vector space $V_{\mathbb{C}}$ which is a natural extension of V . In particular, define:

$$V_{\mathbb{C}} = \{(x, y) \mid x, y \in V\}$$

Suppose $(x, y), (v, w) \in V_{\mathbb{C}}$ and $a + ib \in \mathbb{C}$ where $a, b \in \mathbb{R}$. Define:

$$(x, y) + (v, w) = (x + v, y + w) \quad \& \quad (a + ib) \cdot (x, y) = (ax - by, ay + bx).$$

I invite the reader to verify that $V_{\mathbb{C}}$ given the addition and scalar multiplication above forms a vector space over $\mathbb{C}$. In particular we may argue $(0, 0)$ is the zero in $V_{\mathbb{C}}$ and $1 \cdot (x, y) = (x, y)$. Moreover, as $x, y \in V$ and $a, b \in \mathbb{R}$ the fact that V is a real vector space yields $ax - by, ay + bx \in V$. The other axioms all follow from transferring the axioms over $\mathbb{R}$ for V to $V_{\mathbb{C}}$. Our current notation for $V_{\mathbb{C}}$ is a bit tiresome. Note:

$$(1 + 0i) \cdot (x, y) = (x, y) \quad \& \quad (0 + i) \cdot (x, y) = (-y, x).$$

Since $\mathbb{R} \subset \mathbb{C}$ the fact that $V_{\mathbb{C}}$ is a complex vector space automatically makes $V_{\mathbb{C}}$ a real vector space. Moreover, with respect to the real vector space structure of $V_{\mathbb{C}}$, there are two natural subspaces of $V_{\mathbb{C}}$ which are isomorphic to V .

$$W_1 = (1 + i0) \cdot V = V \times \{0\} \quad \& \quad W_2 = (0 + i) \cdot V = \{0\} \times V$$

Note $W_1 + W_2 = V_{\mathbb{C}}$ and $W_1 \cap W_2 = \{(0, 0)\}$ hence $V_{\mathbb{C}} = W_1 \oplus W_2$. Here $\oplus$ could be denoted $\oplus_{\mathbb{R}}$ to emphasize it is a direct sum with respect to the real vector space structure of $V_{\mathbb{C}}$. Moreover, it is convenient to simply write $V_{\mathbb{C}} = V \oplus iV$. Another notation for this is $V_{\mathbb{C}} = \mathbb{C} \otimes V$ where $\otimes$ is the **tensor** product. This is perhaps the simplest way to think of the complexification:

To find the complexification of $V(\mathbb{R})$ we simply consider $V(\mathbb{C})$. In other words, replace the real scalars with complex scalars.

This slogan is just a short-hand for the explicit construction outlined thus far in this section.

Example 4.6.1. If $V = \mathbb{R}$ then $V_{\mathbb{C}} = \mathbb{R} \oplus i\mathbb{R} = \mathbb{C}$.

Example 4.6.2. If $V = \mathbb{R}^n$ then $V_{\mathbb{C}} = \mathbb{R}^n \oplus i\mathbb{R}^n = \mathbb{C}^n$.

Example 4.6.3. If $V = \mathbb{R}^{m \times n}$ then $V_{\mathbb{C}} = \mathbb{R}^{m \times n} \oplus i\mathbb{R}^{m \times n} = \mathbb{C}^{m \times n}$.

We might notice a simple result about the basis of $V_{\mathbb{C}}$ which is easy to verify in the examples given thus far: if $\text{span}_{\mathbb{R}}(\beta) = V$ then $\text{span}_{\mathbb{C}}(\beta) = V_{\mathbb{C}}$.

Proposition 4.6.4.

Suppose V over $\mathbb{R}$ has basis $\beta = \{v_1, \dots, v_n\}$ then β is also a basis for $V_{\mathbb{C}}$

Proof: let $\beta = \{v_1, \dots, v_n\}$ be a basis for V over $\mathbb{R}$. Notice, $\beta \subset V_{\mathbb{C}}$ under the usual identification of $V \leq V_{\mathbb{C}}$ as described in this section. Let $z \in V_{\mathbb{C}}$ then there exist $x, y \in V$ for which $z = x + iy$. Moreover, as $x, y \in \text{span}_{\mathbb{R}}(\beta)$ there exists $x_j, y_j \in \mathbb{R}$ for which $x = \sum_{j=1}^n x_j v_j$ and $y = \sum_{j=1}^n y_j v_j$. Thus,

$$z = x + iy = \sum_{j=1}^n x_j v_j + i \sum_{j=1}^n y_j v_j = \sum_{j=1}^n (x_j + iy_j) v_j \in \text{span}_{\mathbb{C}}(\beta)$$

Therefore, β is a generating set for $V_{\mathbb{C}}$. To prove linear independence of β over $\mathbb{C}$ suppose $c_j = a_j + ib_j$ are complex constants with real parts $a_j \in \mathbb{R}$ and imaginary coefficients $b_j \in \mathbb{R}$. Consider,

$$\sum_{j=1}^n c_j v_j = 0 \Rightarrow \sum_{j=1}^n (a_j + ib_j) v_j = \sum_{j=1}^n a_j v_j + i \left(\sum_{j=1}^n b_j v_j \right) = 0$$

Therefore, both $\sum_{j=1}^n a_j v_j = 0$ and $\sum_{j=1}^n b_j v_j = 0$. By the real LI of β we find $a_j = 0$ and $b_j = 0$ for all $j \in \mathbb{N}_n$ hence $c_j = a_j + ib_j = 0$ for all $j \in \mathbb{N}_n$ and we conclude β is linearly independent in $V_{\mathbb{C}}$ and thus β is a basis for the complex vector space $V_{\mathbb{C}}$. $\square$

If we view $V_{\mathbb{C}}$ as real vector space and if β is a basis for V then $\beta \cup i\beta$ is a natural basis for $V_{\mathbb{C}}$. Although, it is often useful to order the real basis for $V_{\mathbb{C}}$ as follows: given $\beta = \{v_1, v_2, \dots, v_n\}$ construct $\beta_{\mathbb{C}}$ as

$$\beta_{\mathbb{C}} = \{v_1, iv_1, v_2, iv_2, \dots, v_n, iv_n\}$$

Example 4.6.5. If $V = \mathbb{R}[t]$ then $V_{\mathbb{C}} = \mathbb{R}[t] \oplus i\mathbb{R}[t] = \mathbb{C}[t]$. Likewise for polynomials of limited degree. For example $W = P_2$ is given by $\text{span}_{\mathbb{R}}\{1, t, t^2\}$ whereas $W_{\mathbb{C}} = \text{span}_{\mathbb{R}}\{1, i, t, it, t^2, it^2\}$.

From a purely complex perspective viewing an n -complex-dimensional space as a $2n$ -dimensional real vector space is awkward. However, in the application we are most interested, the complex vector space viewed as a real vector space yields data of interest to our study. We are often interested in solving real problems, but a complexification of the problem at times yields a simpler problem which is easily solved. Once the complexification has served its purpose of solvability then we have to drop back to the context of real vector spaces. This is the game plan, and the reason we are spending some effort to discuss the complexification technique.

Example 4.6.6. If $V = \mathcal{L}(U, W)$ then $V_{\mathbb{C}} = \mathcal{L}(U, W) \oplus i\mathcal{L}(U, W)$. If $T \in V_{\mathbb{C}}$ then $T = L_1 + iL_2$ for some $L_1, L_2 \in V$. However, if β is a basis for U then β is a complex basis for $U_{\mathbb{C}}$ thus T extends uniquely to a complex linear map $T_{\mathbb{C}} : U_{\mathbb{C}} \rightarrow W_{\mathbb{C}}$. Therefore, we find $V_{\mathbb{C}} = \mathcal{L}_{\mathbb{C}}(U_{\mathbb{C}}, W_{\mathbb{C}})$

Example 4.6.7. As a particular application of the discussion in the last example: if $V = \mathcal{L}(\mathbb{R}^n, \mathbb{R}^m)$ then $V_{\mathbb{C}} = \mathcal{L}_{\mathbb{C}}(\mathbb{C}^n, \mathbb{C}^m)$. Note that isomorphism and complexification intertwine nicely: $V \cong \mathbb{R}^{m \times n}$ and $\mathbb{C} \otimes V \cong \mathbb{C} \otimes \mathbb{R}^{m \times n}$ as $V_{\mathbb{C}} \cong \mathbb{C}^{m \times n}$.

The last example brings us to the main-point of this discussion. If we consider $T : \mathbb{R}^n \rightarrow \mathbb{R}^m$ and we extend to $T_{\mathbb{C}} : \mathbb{C}^n \rightarrow \mathbb{C}^m$ then this simply amounts to allowing the matrix of T be complex. Also, conversely, if we allow the matrix to be complex then it implies we have extended to a complex domain. The formula which defines the complexified version of a real linear transformation is simply:

$$T_{\mathbb{C}}(x + iy) = T(x) + iT(y)$$

for all $x, y \in V$. This idea is at times tacitly used without any explicit mention of the complexification. In view of our discussion in this section that omission is not too dangerous. Indeed, that is why in other courses I at times just *allow* the variable to be complex. This amounts to the complexification procedure defined in this section.

4.6.3 complexification for 2nd order constant coefficient problem

I'll illustrate the technique of complexification in the study of differential equations. To solve $ay'' + by' + cy = 0$ we try to use the **real** solution $y = e^{\lambda t}$. Since $y' = \lambda e^{\lambda t}$ and $y'' = \lambda^2 e^{\lambda t}$

$$ay'' + by' + cy = 0 \Rightarrow a\lambda^2 e^{\lambda t} + b\lambda e^{\lambda t} + ce^{\lambda t} = 0$$

then as $e^{\lambda t} \neq 0$ we can divide by it to reveal the⁶ **auxillary equation** of $a\lambda^2 + b\lambda + c = 0$. If the solution to the auxillary equation is real then we get at least one solution. For example, $y'' + 3y' + 2y = 0$ gives $\lambda^2 + 3\lambda + 2 = (\lambda + 1)(\lambda + 2) = 0$ this $\lambda_1 = -1$ and $\lambda_2 = -2$ and the general solution is simply $y = c_1 e^{-t} + c_2 e^{-2t}$.

On the other hand, if we try the same real approach to solve $y'' + y = 0$ then we face $\lambda^2 + 1 = 0$ which has no real solutions. Therefore, we complexify the problem and study $z'' + z = 0$ where $z = x + iy$ and x, y are real-valued functions of t . Conveniently, if $\lambda = \alpha + i\beta$ is complex then $z = e^{\lambda t} = e^{\alpha t}(\cos \beta t + i \sin \beta t)$ and we can derive $z' = \lambda e^{\lambda t}$ and $z'' = \lambda^2 e^{\lambda t}$ hence, by the same argument as in the real case, we look for (possibly complex) solutions of $a\lambda^2 + b\lambda + c = 0$. Returning to our $y'' + y = 0$ example, we now solve $\lambda^2 + 1 = 0$ to obtain $\lambda = \pm i$. It follow that $z_1 = e^{it}$ and $z_2 = e^{-it}$ are complex solutions for $z'' + z = 0$. Notice, $z = x + iy$ has $z' = x' + iy'$ and $z'' = x'' + iy''$ thus $z'' + z = 0$ implies $x'' + iy'' + x + iy = 0$ thus $x'' + x + i(y'' + y) = 0$ (call this $\star$). Notice the real and imaginary components of $\star$ must both be zero hence $x'' + x = 0$ and $y'' + y = 0$. Notice, for $z = e^{it} = \cos t + i \sin t$ we have $x = \cos t$ and $y = \sin t$. It follows that $\cos t$ and $\sin t$ are **real** solutions to $y'' + y = 0$. Indeed, the general solution to $y'' + y = 0$ is $y = c_1 \cos t + c_2 \sin t$. To summarize: we take the given real problem, extend to a corresponding complex problem, solve the complex problem using the added algebraic flexibility of $\mathbb{C}$, then extract a pair of real solutions from the complex solution. You might notice, we didn't need to use e^{-it} since $e^{-it} = \cos t - i \sin t$ only has data we already found in e^{it} .

The case $y'' + 2y' + y = 0$ is more troubling. Here $\lambda^2 + 2\lambda + 1 = (\lambda + 1)^2$ hence $\lambda = -1$ twice. We only get $y = e^{-t}$ as a solution. The other solution $y = te^{-t}$ arises from a generalized eigenvector for a system of differential equations which corresponds to $y'' + 2y' + y = 0$. The reason for the t is subtle. We will discuss this further when we study systems of differential equations.

4.7 complex eigenvalues and vectors for real linear maps

By now it should be clear that as we consider problems of real vector spaces the general results, especially those algebraic in nature, invariably involve some complex case. However, technically it usually happens that the construction from which the complex algebra arose is no longer valid if the algebra requires complex solutions. The technique to capture data in the complex cases of the real problems is to **complexify** the problem. What this means is we replace the given vector spaces with their complexifications and we extend the linear transformations of interest in the same fashion. It turns out that solutions to the complexification of the problem reveal both the real solutions of the original problem as well as complex solutions which, while not real solutions, still yield useful data for unwrapping the general real problem. If this all seems a little vague, don't worry, we will get into all the messy details for the eigenvector problem.

Definition 4.7.1.

If $T : V \rightarrow V$ is a linear transformation over $\mathbb{R}$ then the **complexification of T** is the natural extension of T to $T_{\mathbb{C}} : V_{\mathbb{C}} \rightarrow V_{\mathbb{C}}$ where $V_{\mathbb{C}} = V \oplus iV$ given by:

$$T_{\mathbb{C}}(x + iy) = T(x) + iT(y)$$

for all $x + iy \in V_{\mathbb{C}}$. If $v \in V_{\mathbb{C}}$ is a nonzero vector and $\lambda \in \mathbb{C}$ for which $T_{\mathbb{C}}(v) = \lambda v$ then we say v is a **complex eigenvector with eigenvalue λ for T** .

⁶I usually call it the characteristic equation, but, I'd rather not at the moment

Example 4.7.2. Consider $T = D$ where $D = d/dx$. If $\lambda = \alpha + i\beta$ then $e^{\lambda x} = e^{\alpha x}(\cos(\beta x) + i \sin(\beta x))$ by definition of the complex exponential. It is first semester calculus to show $D_{\mathbb{C}}(e^{\lambda x}) = \lambda e^{\lambda x}$. Thus $e^{\lambda x}$ is a complex e-vector of $T_{\mathbb{C}}$ with complex e-value λ . In other words, $e^{\lambda x}$ for complex λ are complex eigenfunctions of the differentiation operator.

Suppose $\beta = \{f_1, \dots, f_n\}$ is a basis for V ; $\text{span}_{\mathbb{R}}(\beta) = V$. Then, Proposition 4.6.4 showed us that β also serves as a complex basis for $V_{\mathbb{C}}$, $\text{span}_{\mathbb{C}}(\beta) = V_{\mathbb{C}}$. It follows that the matrix of $T_{\mathbb{C}}$ with respect to β over $\mathbb{C}$ is the same as the matrix of T with respect to β over $\mathbb{R}$. In particular:

$$[T_{\mathbb{C}}(f_i)]_{\beta} = [T(f_i)]_{\beta}.$$

Suppose v is a complex e-vector with e-value λ then note $T_{\mathbb{C}}(v) = \lambda v$ implies $[T_{\mathbb{C}}]_{\beta, \beta}[v]_{\beta} = \lambda[v]_{\beta}$ where $[v]_{\beta} \in \mathbb{C}^n$. However, $[T_{\mathbb{C}}]_{\beta, \beta} = [T]_{\beta, \beta}$. Conversely, if $[T]_{\beta, \beta}$ viewed as a matrix in $\mathbb{C}^{n \times n}$ has complex e-vector w with e-value λ then $v = \Phi_{\beta}^{-1}(w)$ is a complex e-vector for $T_{\mathbb{C}}$ with e-value λ . My point is simply this: we can exchange the problem of complex e-vectors of T for the associated problem of finding complex e-vectors of $[T]_{\beta, \beta}$. Just as we found in the case of real e-vectors it suffices to study the matrix problem.

Definition 4.7.3.

Let $A \in \mathbb{R}^{n \times n}$. If $v \in \mathbb{C}^n$ is nonzero and $Av = \lambda v$ for some $\lambda \in \mathbb{C}$ then we say v is a **complex eigenvector** with **complex eigenvalue** λ of the real matrix A .

The solutions of the characteristic equation are eigenvalues.

Proposition 4.7.4.

Let $A \in \mathbb{R}^{n \times n}$ then $\lambda \in \mathbb{C}$ is an eigenvalue of A iff $\det(A - \lambda I) = 0$. We say $P(\lambda) = \det(A - \lambda I)$ the **characteristic polynomial** and $\det(A - \lambda I) = 0$ is the **characteristic equation**.

Proof: a complex e-vector of A is an e-vector of linear transformation $L_A : \mathbb{C}^n \rightarrow \mathbb{C}^n$ and the possible eigenvalues of L_A are solutions of $\det(A - \lambda I) = 0$ since $[L_A] = A$. $\square$

The complex case is different than the real case for one main reason: the complex numbers are an algebraically closed field. In particular we have the Fundamental Theorem of Algebra⁷

Theorem 4.7.5.

Fundamental Theorem of Algebra: if $P(x)$ is an n -th order polynomial complex coefficients then the equation $P(x) = 0$ has n -solutions where some of the solutions may be repeated. Moreover, if $P(x)$ is an n -th order polynomial with real coefficients then complex solutions to $P(x) = 0$ come in conjugate pairs. It follows that any polynomial with real coefficients can be factored into a unique product of repeated real and irreducible quadratic factors.

A proof of this theorem would take us far of topic here⁸. I state it here to remind us of the possibilities for solutions of the characteristic equation $P(\lambda) = \det(A - \lambda I) = 0$ which is simply an n -th order polynomial equation in λ .

⁷sometimes this is stated as "there exists at least one complex solution to an n -th order complex polynomial equation" then the factor theorem repeated applied leads to the theorem I quote here.

⁸there is a nice proof which can be given in our complex variables course

Proposition 4.7.6.

If $A \in \mathbb{C}^{n \times n}$ then A has n eigenvalues, however, some may be repeated and/or complex.

Proof: observe $P(\lambda) = \det(A - \lambda I) = 0$ is an n -th order polynomial equation in λ . $\square$

In the case $A \in \mathbb{R}^{n \times n}$ the complex e-vectors have special structure⁹.

Proposition 4.7.7.

If $A \in \mathbb{R}^{n \times n}$ has complex eigenvalue λ and complex eigenvector v then λ^* is likewise a complex eigenvalue with complex eigenvector v^* for A .

Proof: We assume $Av = \lambda v$ for some $\lambda \in \mathbb{C}$ and $v \in \mathbb{C}^{n \times 1}$ with $v \neq 0$. Take the complex conjugate of $Av = \lambda v$ to find $A^*v^* = \lambda^*v^*$. But, $A \in \mathbb{R}^{n \times n}$ thus $A^* = A$ and we find $Av^* = \lambda^*v^*$. Moreover, if $v \neq 0$ then $v^* \neq 0$. Therefore, v^* is an e-vector with e-value λ^* . $\square$

This is a useful proposition. It means that once we calculate one complex e-vectors we almost automatically get a second e-vector merely by taking the complex conjugate.

Proposition 4.7.8.

If $A \in \mathbb{R}^{m \times n}$ has complex e-value $\lambda = \alpha + i\beta$ such that $\beta \neq 0$ and e-vector $v = a + ib \in \mathbb{C}^{n \times 1}$ such that $a, b \in \mathbb{R}^n$ then $\lambda^* = \alpha - i\beta$ is a complex e-value with e-vector $v^* = a - ib$ and $\{v, v^*\}$ is a linearly independent set of vectors over $\mathbb{C}$.

Proof: Proposition 4.7.7 showed that v^* is an e-vector with e-value $\lambda^* = \alpha - i\beta$. Notice that $\lambda \neq \lambda^*$ since $\beta \neq 0$. Therefore, v and v^* are e-vectors with distinct e-values. Note that Proposition 4.1.17 is equally valid for complex e-values and e-vectors. Hence, $\{v, v^*\}$ is linearly independent since these are complex e-vectors with distinct complex e-values. $\square$

Proposition 4.7.9.

If $A \in \mathbb{R}^{m \times n}$ has complex e-value $\lambda = \alpha + i\beta$ such that $\beta \neq 0$ and e-vector $v = a + ib \in \mathbb{C}^{n \times 1}$ such that $a, b \in \mathbb{R}^n$ then $a \neq 0$ and $b \neq 0$.

Proof: Expand $Av = \lambda v$ into the real components,

$$\lambda v = (\alpha + i\beta)(a + ib) = \alpha a - \beta b + i(\beta a + \alpha b)$$

and

$$Av = A(a + ib) = Aa + iAb$$

Equating real and imaginary components yeilds two real matrix equations,

$$Aa = \alpha a - \beta b \quad \text{and} \quad Ab = \beta a + \alpha b$$

Suppose $a = 0$ towards a contradiction, note that $0 = -\beta b$ but then $b = 0$ since $\beta \neq 0$ thus $v = 0 + i0 = 0$ but this contradicts v being an e-vector. Likewise if $b = 0$ we find $\beta a = 0$ which

⁹notice, in Lecture on March 25 of 2016, I presented these results for a linear transformation and I think the arguments I gave there are an improvement on those offered here

implies $a = 0$ and again $v = 0$ which contradicts v being an e-vector. Therefore, $a, b \neq 0$. $\square$

Let T be a linear transformation on a $\mathbb{R}^2$ such that $v = a + ib$ is a complex eigenvector with $\lambda = \alpha + i\beta$. The calculations above make it clear that if we set $\gamma = \{a, b\}$ then

$$[T]_{\gamma, \gamma} = [[T(a)]_{\gamma} \mid [T(b)]_{\gamma}] = \begin{bmatrix} \alpha & \beta \\ -\beta & \alpha \end{bmatrix}.$$

Of course, to be careful, we should prove $\{a, b\}$ is a LI before are certain γ is a basis.

Proposition 4.7.10.

If $A \in \mathbb{R}^{n \times n}$ and $\lambda = \alpha + i\beta \in \mathbb{C}$ with $\alpha, \beta \in \mathbb{R}$ and $\beta \neq 0$ is an e-value with e-vector $v = a + ib \in \mathbb{C}^{n \times 1}$ and $a, b \in \mathbb{R}^n$ then $\{a, b\}$ is a linearly independent set of real vectors.

Proof: Add and subtract the equations $v = a + ib$ and $v^* = a - ib$ to deduce

$$a = \frac{1}{2}(v + v^*) \quad \text{and} \quad b = \frac{1}{2i}(v - v^*)$$

Let $c_1, c_2 \in \mathbb{R}$ then consider,

$$\begin{aligned} c_1 a + c_2 b = 0 &\Rightarrow c_1 \left[\frac{1}{2}(v + v^*) \right] + c_2 \left[\frac{1}{2i}(v - v^*) \right] = 0 \\ &\Rightarrow [c_1 - ic_2]v + [c_1 + ic_2]v^* = 0 \end{aligned}$$

But, $\{v, v^*\}$ is linearly independent hence $c_1 - ic_2 = 0$ and $c_1 + ic_2 = 0$. Adding these equations gives $2c_1 = 0$. Subtracting yields $2ic_2 = 0$. Thus $c_1 = c_2 = 0$ and we conclude $\{a, b\}$ is linearly independent set of real vectors. $\square$

Proposition 4.7.11.

If $A \in \mathbb{R}^{n \times n}$ and $\lambda_j = \alpha_j + i\beta_j \in \mathbb{C}$ with $\alpha_j, \beta_j \in \mathbb{R}$ and $\beta_j \neq 0$ is an e-value with e-vector $v_j = a_j + ib_j \in \mathbb{C}^n$ and $a_j, b_j \in \mathbb{R}^n$ for $j = 1, 2, \dots, k$ then $\{a_1, b_1, a_2, b_2, \dots, a_k, b_k\}$ is a linearly independent set of real vectors.

Proof: should be similar to that of Proposition 4.7.10. I leave the details to the reader. $\square$

4.8 examples of real and complex eigenvectors

And now, the examples! Note, we should see all the propositions exhibited in these examples.

4.8.1 characteristic equations

Example 4.8.1. Let $A = \begin{bmatrix} 3 & 0 \\ 8 & -1 \end{bmatrix}$. Find the eigenvalues of A from the characteristic equation:

$$\det(A - \lambda I) = \det \begin{bmatrix} 3 - \lambda & 0 \\ 8 & -1 - \lambda \end{bmatrix} = (3 - \lambda)(-1 - \lambda) = (\lambda + 1)(\lambda - 3) = 0$$

Hence the eigenvalues are $\lambda_1 = -1$ and $\lambda_2 = 3$. Notice this is precisely the factor of 3 we saw scaling the vector in the first example of this chapter.

Example 4.8.2. Let $A = \begin{bmatrix} \frac{1}{2} & \frac{\sqrt{3}}{2} \\ -\frac{\sqrt{3}}{2} & \frac{1}{2} \end{bmatrix}$. Find the eigenvalues of A from the characteristic equation:

$$\det(A - \lambda I) = \det \begin{bmatrix} \frac{1}{2} - \lambda & \frac{\sqrt{3}}{2} \\ -\frac{\sqrt{3}}{2} & \frac{1}{2} - \lambda \end{bmatrix} = \left(\frac{1}{2} - \lambda\right)^2 + \frac{3}{4} = \left(\lambda - \frac{1}{2}\right)^2 + \frac{3}{4} = 0$$

Well how convenient is that? The determinant completed the square for us. We find: $\lambda = \frac{1}{2} \pm i\frac{\sqrt{3}}{2}$. It would seem that elliptical orbits somehow arise from complex eigenvalues

Example 4.8.3. Given A below, find the eigenvalues. Since the matrix is upper triangular the determinant is calculated as the product of the diagonal entries:

$$A = \begin{bmatrix} 2 & 3 & 4 \\ 0 & 5 & 6 \\ 0 & 0 & 7 \end{bmatrix} \Rightarrow \det(A - \lambda I) = \begin{bmatrix} 2 - \lambda & 3 & 4 \\ 0 & 5 - \lambda & 6 \\ 0 & 0 & 7 - \lambda \end{bmatrix} = (2 - \lambda)(5 - \lambda)(7 - \lambda)$$

Therefore, $\lambda_1 = 2, \lambda_2 = 5$ and $\lambda_3 = 7$.

Remark 4.8.4. *eigenwarning*

Calculation of eigenvalues does not need to be difficult. However, I urge you to be careful in solving the characteristic equation. More often than not I see students make a mistake in calculating the eigenvalues. If you do that wrong then the eigenvector calculations will often turn into inconsistent equations. This should be a clue that the eigenvalues were wrong, but often I see what I like to call the "principle of minimal calculation" take over and students just adhoc set things to zero, hoping against all logic that I will somehow not notice this. Don't be this student. If the eigenvalues are correct, the eigenvector equations are consistent and you will be able to find nonzero eigenvectors. And don't forget, the eigenvectors must be nonzero.

4.8.2 real eigenvector examples

Example 4.8.5. Let $A = \begin{bmatrix} 3 & 1 \\ 3 & 1 \end{bmatrix}$ find the e -values and e -vectors of A .

$$\det(A - \lambda I) = \det \begin{bmatrix} 3 - \lambda & 1 \\ 3 & 1 - \lambda \end{bmatrix} = (3 - \lambda)(1 - \lambda) - 3 = \lambda^2 - 4\lambda = \lambda(\lambda - 4) = 0$$

We find $\lambda_1 = 0$ and $\lambda_2 = 4$. Now find the e -vector with e -value $\lambda_1 = 0$, let $u_1 = [u, v]^T$ denote the e -vector we wish to find. Calculate,

$$(A - 0I)u_1 = \begin{bmatrix} 3 & 1 \\ 3 & 1 \end{bmatrix} \begin{bmatrix} u \\ v \end{bmatrix} = \begin{bmatrix} 3u + v \\ 3u + v \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \end{bmatrix}$$

Obviously the equations above are redundant and we have infinitely many solutions of the form $3u + v = 0$ which means $v = -3u$ so we can write, $u_1 = \begin{bmatrix} u \\ -3u \end{bmatrix} = u \begin{bmatrix} 1 \\ -3 \end{bmatrix}$. In applications we often make a choice to select a particular e -vector. Most modern graphing calculators can calculate e -vectors. It is customary for the e -vectors to be chosen to have length one. That is a useful choice for certain applications as we will later discuss. If you use a calculator it would likely give

$u_1 = \frac{1}{\sqrt{10}} \begin{bmatrix} 1 \\ -3 \end{bmatrix}$ although the $\sqrt{10}$ would likely be approximated unless your calculator is smart.

Continuing we wish to find eigenvectors $u_2 = [u, v]^T$ such that $(A - 4I)u_2 = 0$. Notice that u, v are disposable variables in this context, I do not mean to connect the formulas from the $\lambda = 0$ case with the case considered now.

$$(A - 4I)u_1 = \begin{bmatrix} -1 & 1 \\ 3 & -3 \end{bmatrix} \begin{bmatrix} u \\ v \end{bmatrix} = \begin{bmatrix} -u + v \\ 3u - 3v \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \end{bmatrix}$$

Again the equations are redundant and we have infinitely many solutions of the form $v = u$. Hence, $u_2 = \begin{bmatrix} u \\ u \end{bmatrix} = u \begin{bmatrix} 1 \\ 1 \end{bmatrix}$ is an eigenvector for any $u \in \mathbb{R}$ such that $u \neq 0$.

Remark 4.8.6.

It was obvious the equations were redundant in the example above. However, we need not rely on pure intuition. The problem of calculating all the e-vectors is precisely the same as finding all the vectors in the null space of a matrix. We already have a method to do that without ambiguity. We find the rref of the matrix and the general solution falls naturally from that matrix. I don't bother with the full-blown theory for simple examples because there is no need. However, with 3×3 examples it may be advantageous to keep our earlier null space calculational scheme in mind.

Example 4.8.7. Let $A = \begin{bmatrix} 0 & 0 & -4 \\ 2 & 4 & 2 \\ 2 & 0 & 6 \end{bmatrix}$ find the e-values and e-vectors of A .

$$\begin{aligned} 0 = \det(A - \lambda I) &= \det \begin{bmatrix} -\lambda & 0 & -4 \\ 2 & 4 - \lambda & 2 \\ 2 & 0 & 6 - \lambda \end{bmatrix} \\ &= (4 - \lambda)[- \lambda(6 - \lambda) + 8] \\ &= (4 - \lambda)[\lambda^2 - 6\lambda + 8] \\ &= -(\lambda - 4)(\lambda - 4)(\lambda - 2) \end{aligned}$$

Thus we have a repeated e-value of $\lambda_1 = \lambda_2 = 4$ and $\lambda_3 = 2$. Let's find the eigenvector $u_3 = [u, v, w]^T$ such that $(A - 2I)u_3 = 0$, we find the general solution by row reduction

$$\text{rref} \left[\begin{array}{ccc|c} -2 & 0 & -4 & 0 \\ 2 & 2 & 2 & 0 \\ 2 & 0 & 4 & 0 \end{array} \right] = \left[\begin{array}{ccc|c} 1 & 0 & 2 & 0 \\ 0 & 1 & -1 & 0 \\ 0 & 0 & 0 & 0 \end{array} \right] \Rightarrow \begin{array}{l} u + 2w = 0 \\ v - w = 0 \end{array} \Rightarrow u_3 = w \begin{bmatrix} -2 \\ 1 \\ 1 \end{bmatrix}$$

Next find the e-vectors with e-value 4. Let $u_1 = [u, v, w]^T$ satisfy $(A - 4I)u_1 = 0$. Calculate,

$$\text{rref} \left[\begin{array}{ccc|c} -4 & 0 & -4 & 0 \\ 2 & 0 & 2 & 0 \\ 2 & 0 & 2 & 0 \end{array} \right] = \left[\begin{array}{ccc|c} 1 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \end{array} \right] \Rightarrow u + w = 0$$

Notice this case has two free variables, we can use v, w as parameters in the solution,

$$u_1 = \begin{bmatrix} u \\ v \\ w \end{bmatrix} = \begin{bmatrix} -w \\ v \\ w \end{bmatrix} = v \begin{bmatrix} 0 \\ 1 \\ 0 \end{bmatrix} + w \begin{bmatrix} -1 \\ 0 \\ 1 \end{bmatrix} \Rightarrow u_1 = v \begin{bmatrix} 0 \\ 1 \\ 0 \end{bmatrix} \text{ and } u_2 = w \begin{bmatrix} -1 \\ 0 \\ 1 \end{bmatrix}$$

I have boxed two linearly independent eigenvectors u_1, u_2 . These vectors will be linearly independent for any pair of nonzero constants v, w .

You might wonder if it is always the case that repeated e-values get multiple e-vectors. In the preceding example the e-value 4 had *algebraic multiplicity* two and there were likewise two linearly independent e-vectors. The next example shows that is not the case.

Example 4.8.8. Let $A = \begin{bmatrix} 1 & 1 \\ 0 & 1 \end{bmatrix}$ find the e-values and e-vectors of A .

$$\det(A - \lambda I) = \det \begin{bmatrix} 1 - \lambda & 1 \\ 0 & 1 - \lambda \end{bmatrix} = (1 - \lambda)(1 - \lambda) = 0$$

Hence we have a repeated e-value of $\lambda_1 = 1$. Find all e-vectors for $\lambda_1 = 1$, let $u_1 = [u, v]^T$,

$$(A - I)u_1 = \begin{bmatrix} 0 & 1 \\ 0 & 0 \end{bmatrix} \begin{bmatrix} u \\ v \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \end{bmatrix} \Rightarrow v = 0 \Rightarrow \boxed{u_1 = u \begin{bmatrix} 1 \\ 0 \end{bmatrix}}$$

We have only one e-vector for this system.

I should mention, if λ is an eigenvalue for $A \in \mathbb{F}^{n \times n}$ or $T : V \rightarrow V$ then there exists at least one non-trivial eigenvector with eigenvalue λ . We can expect there is at least one eigenvector if our alleged eigenvalue is correct. In fact, finding otherwise means you need to return to the characteristic equation and try to solve it more carefully.

Example 4.8.9. Let $A = \begin{bmatrix} 1 & 2 & 3 \\ 4 & 5 & 6 \\ 7 & 8 & 9 \end{bmatrix}$ find the e-values and e-vectors of A .

$$\begin{aligned} 0 = \det(A - \lambda I) &= \det \begin{bmatrix} 1 - \lambda & 2 & 3 \\ 4 & 5 - \lambda & 6 \\ 7 & 8 & 9 - \lambda \end{bmatrix} \\ &= (1 - \lambda)[(5 - \lambda)(9 - \lambda) - 48] - 2[4(9 - \lambda) - 42] + 3[32 - 7(5 - \lambda)] \\ &= -\lambda^3 + 15\lambda^2 + 18\lambda \\ &= -\lambda(\lambda^2 - 15\lambda - 18) \end{aligned}$$

Therefore, using the quadratic equation to factor the ugly part,

$$\lambda_1 = 0, \quad \lambda_2 = \frac{15 + 3\sqrt{33}}{2}, \quad \lambda_3 = \frac{15 - 3\sqrt{33}}{2}$$

The e-vector for e-value zero is not too hard to calculate. Find $u_1 = [u, v]^T$ such that $(A - 0I)u_1 = 0$. This amounts to row reducing A itself:

$$\text{rref} \left[\begin{array}{ccc|c} 1 & 2 & 3 & 0 \\ 4 & 5 & 6 & 0 \\ 7 & 8 & 9 & 0 \end{array} \right] = \left[\begin{array}{ccc|c} 1 & 0 & -1 & 0 \\ 0 & 1 & 2 & 0 \\ 0 & 0 & 0 & 0 \end{array} \right] \Rightarrow \begin{array}{l} u - w = 0 \\ v + 2w = 0 \end{array} \Rightarrow \boxed{u_1 = w \begin{bmatrix} 1 \\ -2 \\ 1 \end{bmatrix}}$$

The e-vectors corresponding e-values λ_2 and λ_3 are hard to calculate without numerical help. Let's discuss Texas Instrument calculator output. To my knowledge, TI-85 and higher will calculate both

e-vectors and *e*-values. For example, my ancient TI-89, displays the following if I define our matrix $A = \text{mat2}$,

$$\text{eigVl}(\text{mat2}) = \{16.11684397, -1.11684397, 1.385788954e - 13\}$$

Calculators often need a little interpretation, the third entry is really zero in disguise. The *e*-vectors will be displayed in the same order, they are given from the "eigVc" command in my TI-89,

$$\text{eigVc}(\text{mat2}) = \begin{bmatrix} .2319706872 & .7858302387 & .4082482905 \\ .5253220933 & .0867513393 & -.8164965809 \\ .8186734994 & -.6123275602 & .4082482905 \end{bmatrix}$$

From this we deduce that eigenvectors for λ_1, λ_2 and λ_3 are

$$u_1 = \begin{bmatrix} .2319706872 \\ .5253220933 \\ .8186734994 \end{bmatrix} \quad u_2 = \begin{bmatrix} .7858302387 \\ .0867513393 \\ -.6123275602 \end{bmatrix} \quad u_3 = \begin{bmatrix} .4082482905 \\ -.8164965809 \\ .4082482905 \end{bmatrix}$$

Notice that $1/\sqrt{6} \approx 0.408248905$ so you can see that u_3 above is simply the $u = 1/\sqrt{6}$ case for the family of *e*-vectors we calculated by hand already. The calculator chooses *e*-vectors so that the vectors have length one.

While we're on the topic of calculators, perhaps it is worth revisiting the example where there was only one *e*-vector. How will the calculator respond in that case? Can we trust the calculator?

Example 4.8.10. Recall Example 4.8.8. We let $A = \begin{bmatrix} 1 & 1 \\ 0 & 1 \end{bmatrix}$ and found a repeated *e*-value of $\lambda_1 = 1$ and single *e*-vector $u_1 = u \begin{bmatrix} 1 \\ 0 \end{bmatrix}$. Hey now, it's time for technology, let $A = a$,

$$\text{eigVl}(a) = \{1, 1\} \quad \text{and} \quad \text{eigVc}(a) = \begin{bmatrix} 1. & -1. \\ 0. & 1.e - 15 \end{bmatrix}$$

Behold, the calculator has given us two alleged *e*-vectors. The first column is the genuine *e*-vector we found previously. The second column is the result of machine error. The calculator was tricked by round-off error into claiming that $[-1, 0.000000000000001]$ is a distinct *e*-vector for A . It is not. Moral of story? When using calculator you must first master the theory or else you'll stay mired in ignorance as prescribed by your robot masters.

Finally, I should mention that TI-calculators may or may not deal with complex *e*-vectors adequately. There are doubtless many web resources for calculating *e*-vectors/values. I would wager if you Googled it you'd find an online calculator that beats any calculator. Many of you have a laptop with wireless so there is almost certainly a way to check your answers if you just take a minute or two. I don't mind you checking your answers. If I assign it in homework then I do want you to work it out without technology. Otherwise, you could get a false confidence before the test. Technology is to supplement not replace calculation.

Remark 4.8.11.

I would also remind you that there are oodles of examples beyond these lecture notes in the homework solutions from previous year(s). If these notes do not have enough examples on some topic then you should seek additional examples elsewhere, ask me, etc... Do not suffer in silence, ask for help. Thanks.

4.8.3 complex eigenvector examples

Example 4.8.12. Let $A = \begin{bmatrix} 0 & 1 \\ -1 & 0 \end{bmatrix}$ and find the e -values and e -vectors of the matrix. Observe that $\det(A - \lambda I) = \lambda^2 + 1$ hence the eigenvalues are $\lambda = \pm i$. Find $u_1 = [u, v]^T$ such that $(A - iI)u_1 = 0$

$$0 = \begin{bmatrix} -i & 1 \\ -1 & -i \end{bmatrix} \begin{bmatrix} u \\ v \end{bmatrix} = \begin{bmatrix} -iu + v \\ -u - iv \end{bmatrix} \Rightarrow \begin{cases} -iu + v = 0 \\ -u - iv = 0 \end{cases} \Rightarrow v = iu \Rightarrow \boxed{u_1 = u \begin{bmatrix} 1 \\ i \end{bmatrix}}$$

We find infinitely many complex eigenvectors, one for each nonzero complex constant u . In applications, it may be convenient to set $u = 1$ so we can write, $u_1 = \begin{bmatrix} 1 \\ 0 \end{bmatrix} + i \begin{bmatrix} 0 \\ 1 \end{bmatrix}$

Let's generalize the last example.

Example 4.8.13. Let $\theta \in \mathbb{R}$ and define $A = \begin{bmatrix} \cos \theta & \sin \theta \\ -\sin \theta & \cos \theta \end{bmatrix}$ and find the e -values and e -vectors of the matrix. Observe

$$\begin{aligned} 0 = \det(A - \lambda I) &= \det \begin{bmatrix} \cos \theta - \lambda & \sin \theta \\ -\sin \theta & \cos \theta - \lambda \end{bmatrix} \\ &= (\cos \theta - \lambda)^2 + \sin^2 \theta \\ &= \cos^2 \theta - 2\lambda \cos \theta + \lambda^2 + \sin^2 \theta \\ &= \lambda^2 - 2\lambda \cos \theta + 1 \\ &= (\lambda - \cos \theta)^2 - \cos^2 \theta + 1 \\ &= (\lambda - \cos \theta)^2 + \sin^2 \theta \end{aligned}$$

Thus $\lambda = \cos \theta \pm i \sin \theta = e^{\pm i\theta}$. Find $u_1 = [u, v]^T$ such that $(A - e^{i\theta} I)u_1 = 0$

$$0 = \begin{bmatrix} -i \sin \theta & \sin \theta \\ -\sin \theta & -i \sin \theta \end{bmatrix} \begin{bmatrix} u \\ v \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \end{bmatrix} \Rightarrow -iu \sin \theta + v \sin \theta = 0$$

If $\sin \theta \neq 0$ then we divide by $\sin \theta$ to obtain $v = iu$ hence $u_1 = [u, iu]^T = u[1, i]^T$ which is precisely what we found in the preceding example. However, if $\sin \theta = 0$ we obtain no condition what-so-ever on the matrix. That special case is not complex. Moreover, if $\sin \theta = 0$ it follows $\cos \theta = 1$ and in fact $A = I$ in this case. The identity matrix has the repeated eigenvalue of $\lambda = 1$ and every vector in $\mathbb{R}^{2 \times 1}$ is an e -vector.

Example 4.8.14. Let $A = \begin{bmatrix} 1 & 1 & 0 \\ -1 & 1 & 0 \\ 0 & 0 & 3 \end{bmatrix}$ find the e -values and e -vectors of A .

$$\begin{aligned} 0 = \det(A - \lambda I) &= \det \begin{bmatrix} 1 - \lambda & 1 & 0 \\ -1 & 1 - \lambda & 0 \\ 0 & 0 & 3 - \lambda \end{bmatrix} \\ &= (3 - \lambda)[(1 - \lambda)^2 + 1] \end{aligned}$$

Hence $\lambda_1 = 3$ and $\lambda_2 = 1 \pm i$. We have a pair of complex e -values and one real e -value. Notice that for any $n \times n$ matrix we must have at least one real e -value since all odd polynomials possess

at least one zero. Let's begin with the real e -value. Find $u_1 = [u, v, w]^T$ such that $(A - 3I)u_1 = 0$:

$$\text{rref} \left[\begin{array}{ccc|c} -2 & 1 & 0 & 0 \\ -1 & -2 & 0 & 0 \\ 0 & 0 & 0 & 0 \end{array} \right] = \left[\begin{array}{ccc|c} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 \end{array} \right] \Rightarrow \boxed{u_1 = w \begin{bmatrix} 0 \\ 0 \\ 1 \end{bmatrix}}$$

Next find e -vector with $\lambda_2 = 1 + i$. We wish to find $u_2 = [u, v, w]^T$ such that $(A - (1 + i)I)u_2 = 0$:

$$\left[\begin{array}{ccc|c} -i & 1 & 0 & 0 \\ -1 & -i & 0 & 0 \\ 0 & 0 & -1-i & 0 \end{array} \right] \xrightarrow[r_2 + ir_1 \rightarrow r_2]{\frac{1}{-1-i}r_3 \rightarrow r_3} \left[\begin{array}{ccc|c} -i & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 \end{array} \right]$$

One more row-swap and a rescaling of row 1 and it's clear that

$$\text{rref} \left[\begin{array}{ccc|c} -i & 1 & 0 & 0 \\ -1 & -i & 0 & 0 \\ 0 & 0 & -1-i & 0 \end{array} \right] = \left[\begin{array}{ccc|c} 1 & i & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 \end{array} \right] \Rightarrow \begin{matrix} u + iv = 0 \\ w = 0 \end{matrix} \Rightarrow \boxed{u_2 = v \begin{bmatrix} i \\ 1 \\ 0 \end{bmatrix}}$$

I chose the free parameter to be v . Any choice of a nonzero complex constant v will yield an e -vector with e -value $\lambda_2 = 1 + i$. For future reference, it's worth noting that if we choose $v = 1$ then we find

$$u_2 = \begin{bmatrix} 0 \\ 1 \\ 0 \end{bmatrix} + i \begin{bmatrix} 1 \\ 0 \\ 0 \end{bmatrix}$$

We identify that $\text{Re}(u_2) = e_2$ and $\text{Im}(u_2) = e_1$

Example 4.8.15. Let $B = \begin{bmatrix} 0 & 1 \\ -1 & 0 \end{bmatrix}$ and let $C = \begin{bmatrix} \frac{1}{2} & \frac{\sqrt{3}}{2} \\ -\frac{\sqrt{3}}{2} & \frac{1}{2} \end{bmatrix}$. Define A to be the **block matrix**

$$A = \left[\begin{array}{c|c} B & 0 \\ \hline 0 & C \end{array} \right] = \left[\begin{array}{cc|cc} 0 & 1 & 0 & 0 \\ -1 & 0 & 0 & 0 \\ \hline 0 & 0 & \frac{1}{2} & \frac{\sqrt{3}}{2} \\ 0 & 0 & -\frac{\sqrt{3}}{2} & \frac{1}{2} \end{array} \right]$$

find the e -values and e -vectors of the matrix. Block matrices have nice properties: the blocks behave like numbers. Of course there is something to prove here, and I have yet to discuss block multiplication in these notes.

$$\det(A - \lambda I) = \det \begin{bmatrix} B - \lambda I & 0 \\ 0 & C - \lambda I \end{bmatrix} = \det(B - \lambda I) \det(C - \lambda I)$$

Notice that both B and C are rotation matrices. B is the rotation matrix with $\theta = \pi/2$ whereas C is the rotation by $\theta = \pi/3$. We already know the e -values and e -vectors for each of the blocks if we ignore the other block. It would be nice if a block matrix allowed for analysis of each block one at a time. This turns out to be true, I can tell you without further calculation that we have e -values $\lambda_1 = \pm i$ and $\lambda_2 = \frac{1}{2} \pm i\frac{\sqrt{3}}{2}$ which have complex e -vectors

$$u_1 = \begin{bmatrix} 1 \\ i \\ 0 \\ 0 \end{bmatrix} = e_1 + ie_2 \quad u_2 = \begin{bmatrix} 0 \\ 0 \\ 1 \\ i \end{bmatrix} = e_3 + ie_4$$

I invite the reader to check my results through explicit calculation. Technically, this is bad form as I have yet to prove anything about block matrices. Perhaps this example gives you a sense of why we should talk about the blocks at some point.

Finally, you might wonder are there matrices which have a repeated complex e-value. And if so are there always as many complex e-vectors as there are complex e-values? The answer: sometimes.

Take for instance $A = \left[\begin{array}{c|c} B & 0 \\ \hline 0 & B \end{array} \right]$ (where B is the same B as in the preceding example) this matrix will have a repeated e-value of $\lambda = \pm i$ and you'll be able to calculate $u_1 = e_1 \pm ie_2$ and $u_2 = e_3 \pm ie_4$ are linearly independent e-vectors for A . However, there are other matrices for which only one complex e-vector is available despite a repeat of the e-value.

Example 4.8.16. Let $A = \begin{bmatrix} 2 & 3 & 1 & 0 \\ -3 & 2 & 0 & 1 \\ 0 & 0 & 2 & 3 \\ 0 & 0 & -3 & 2 \end{bmatrix}$ you can calculate $\lambda = 2 \pm 3i$ is repeated and yet

there are only two LI complex eigenvectors for A . In particular, $v = a + ib$ for $\lambda = 2 + 3i$ and v^* for $\lambda^* = 2 - 3i$. From this pair, or just one of the complex eigenvectors, we find just two LI real vectors: $\{a, b\}$. Naturally, if we wish to associate some basis of $\mathbb{R}^4$ with A then we are missing two vectors. We return to this mystery in the next section. Note:

$$A = \begin{bmatrix} 2 & 3 \\ -3 & 2 \end{bmatrix} \otimes \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix} + \begin{bmatrix} 0 & 1 \\ 0 & 0 \end{bmatrix} \otimes \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix}.$$

The $\otimes$ is the tensor product. Can you see how it is defined?

4.9 real Jordan form

Consider $A \in \mathbb{R}^{n \times n}$. It may not have a Jordan form. Why? We must account for the possibility of complex eigenvalues for A . We continue the work we began in Section 4.7 here.

Theorem 4.9.1.

If V is an n -dimensional real vector space and $T : V \rightarrow V$ is a linear transformation then T has n -complex e-values. Furthermore, if the geometric multiplicity of the complexification of T matches the algebraic multiplicity for each complex e-value then the complexification is diagonalizable; in particular, $T_{\mathbb{C}} : V_{\mathbb{C}} \rightarrow V_{\mathbb{C}}$ permits a complex eigenbasis β for $V_{\mathbb{C}} = V \oplus iV$ such that $[T]_{\beta, \beta} \in \mathbb{C}^{n \times n}$ is diagonal with the complex e-values on the diagonal. If the geometric multiplicity of the complexification does not match the algebraic multiplicity for some complex eigenvalue(s) then it is possible to find a basis of generalized complex e-vectors for $V_{\mathbb{C}}$ for which the matrix of the complexified T has complex Jordan form. Furthermore, up to the ordering of the chains of complex generalized e-vectors the Jordan form of the complexification of T is unique.

Proof: if V is a vector space over $\mathbb{R}$ then $T \in L(V, V)$ has complexification $T_{\mathbb{C}} : V_{\mathbb{C}} \rightarrow V_{\mathbb{C}}$. Since $\mathbb{C}$ is algebraically closed the characteristic equation for $T_{\mathbb{C}}$ has n -complex e-values (allowing repeats). Thus, Theorems 4.2.6 and 4.5.4 apply. $\square$

Diagonalization of $T : V_{\mathbb{C}} \rightarrow V_{\mathbb{C}}$ is interesting, but, we are mostly interested in what the diagonalization reveals about $T : V \rightarrow V$. The simplest case is two-dimensional.

Theorem 4.9.2.

If V is an 2-dimensional real vector space and $T : V \rightarrow V$ is a linear transformation with complex eigenvalue $\lambda = \alpha + i\beta$ where $\beta \neq 0$ with complex eigenvector $v = a + ib \in V_{\mathbb{C}}$ then the matrix of T with respect to $\gamma = \{a, b\}$ is $[T]_{\gamma, \gamma} = \begin{bmatrix} \alpha & \beta \\ -\beta & \alpha \end{bmatrix}$

Proof: If T has complex eigenvalue $\lambda = \alpha + i\beta$ where $\beta \neq 0$ corresponding to complex eigenvector $v = a + ib$ for $a, b \in V$. We assume $T(v) = \lambda v$ hence:

$$T(a + ib) = (\alpha + i\beta)(a + ib)$$

thus, by definition of the complexification,

$$T(a) + iT(b) = \alpha a - \beta b + i(\beta a + \alpha b) \quad \star$$

Then, by a modification of the arguments for Proposition 4.7.10 to the abstract context, we have that $\{a, b\}$ forms a LI set of vectors for V . Since $\dim(V) = 2$ it follows $\gamma = \{a, b\}$ forms a basis. Moreover, from $\star$ we obtain:

$$T(a) = \alpha a - \beta b \quad \& \quad T(b) = \beta a + \alpha b.$$

Recall, the matrix $[T]_{\gamma, \gamma} = [[T(a)]_{\gamma} | [T(b)]_{\gamma}]$. Therefore, the theorem follows as $[T(a)]_{\gamma} = (\alpha, -\beta)$ and $[T(b)]_{\gamma} = (\beta, \alpha)$ are clear from the equations above. $\square$

It might be instructive to note the complexification has a different complex matrix than the real matrix we just exhibited. The key equations are $T(v) = \lambda v$ and $T(v^*) = \lambda^* v^*$ thus if $\delta = \{v, v^*\}$ is a basis for $V_{\mathbb{C}} = V \oplus iV$ then the complexification $T : V_{\mathbb{C}} \rightarrow V_{\mathbb{C}}$ has matrix:

$$[T]_{\delta, \delta} = \begin{bmatrix} \alpha + i\beta & 0 \\ 0 & \alpha - i\beta \end{bmatrix}.$$

The matrix above is complex, but it clearly contains information about the linear transformation T of the real vector space V . Next, we study a repeated complex eigenvalue where the complexification is not complex diagonalizable.

Theorem 4.9.3.

If V is an 4-dimensional real vector space and $T : V \rightarrow V$ is a linear transformation with repeated complex eigenvalue $\lambda = \alpha + i\beta$ where $\beta \neq 0$ with complex eigenvector $v_1 = a_1 + ib_1 \in V_{\mathbb{C}}$ and generalized complex eigenvector $v_2 = a_2 + ib_2$ where $(T - \lambda Id)(v_2) = v_1$ then the matrix of T with respect to $\gamma = \{a_1, b_1, a_2, b_2\}$ is $[T]_{\gamma, \gamma} = \begin{bmatrix} \alpha & \beta & 1 & 0 \\ -\beta & \alpha & 0 & 1 \\ 0 & 0 & \alpha & \beta \\ 0 & 0 & -\beta & \alpha \end{bmatrix}$

Proof: we are given $T(v_1) = \lambda v_1$ and $T(v_2) = \lambda v_2 + v_1$. We simply need to extract real equations from this data: note $v_1 = a_1 + ib_1$ and $v_2 = a_2 + ib_2$ where $a_1, a_2, b_1, b_2 \in V$ and $\lambda = \alpha + i\beta$. Set $\gamma = \{a_1, b_1, a_2, b_2\}$. The first two columns follow from the same calculation as in the proof of Theorem 4.9.2. Calculate,

$$T(a_2 + ib_2) = (\alpha + i\beta)(a_2 + ib_2) + (a_1 + ib_1) = \alpha a_2 - \beta b_2 + a_1 i(\beta a_2 + \alpha b_2 + b_1).$$

Note $T(a_2 + ib_2) = T(a_2) + iT(b_2)$. Thus $T(a_2) = a_1 + \alpha a_2 - \beta b_2$ hence $[T(a_2)]_\gamma = (1, 0, \alpha, -\beta)$. Also, $T(b_2) = b_1 + \beta a_2 + \alpha b_2$ from which it follows $[T(b_2)]_\gamma = (0, 1, \beta, \alpha)$. The theorem follows. $\square$

Once more, I write the matrix of the complexification of T for the linear transformation considered above. Let $\delta = \{v_1, v_2, v_1^*, v_2^*\}$ then

$$[T]_{\delta, \delta} = \left[\begin{array}{cc|cc} \alpha + i\beta & 1 & 0 & 0 \\ 0 & \alpha + i\beta & 0 & 0 \\ \hline 0 & 0 & \alpha - i\beta & 1 \\ 0 & 0 & 0 & \alpha - i\beta \end{array} \right]$$

The next case would be a complex eigenvalue repeated three times. If $\delta = \{v_1, v_2, v_3, v_1^*, v_2^*, v_3^*\}$ where $(T - \lambda)(v_3) = v_2$, $(T - \lambda)(v_2) = v_1$ and $(T - \lambda)(v_1) = 0$. The complex Jordan matrix would have the form:

$$[T]_{\delta, \delta} = \left[\begin{array}{ccc|ccc} \lambda & 1 & 0 & 0 & 0 & 0 \\ 0 & \lambda & 1 & 0 & 0 & 0 \\ 0 & 0 & \lambda & 0 & 0 & 0 \\ \hline 0 & 0 & 0 & \lambda^* & 1 & 0 \\ 0 & 0 & 0 & 0 & \lambda^* & 1 \\ 0 & 0 & 0 & 0 & 0 & \lambda^* \end{array} \right].$$

In this case, if we use the real and imaginary components of v_1, v_2, v_3 as the basis $\gamma = \{a_1, b_1, a_2, b_2, a_3, b_3\}$ then the matrix of $T : V \rightarrow V$ will be formed as follows:

$$[T]_{\gamma, \gamma} = \left[\begin{array}{cccccc} \alpha & \beta & 1 & 0 & 0 & 0 \\ -\beta & \alpha & 0 & 1 & 0 & 0 \\ 0 & 0 & \alpha & \beta & 1 & 0 \\ 0 & 0 & -\beta & \alpha & 0 & 1 \\ 0 & 0 & 0 & 0 & \alpha & \beta \\ 0 & 0 & 0 & 0 & -\beta & \alpha \end{array} \right]. \quad (4.2)$$

The proof is essentially the same as we already offered for the repeated complex eigenvalue case. In Example 4.8.16 we encountered a matrix with a repeated complex eigenvalue with geometric multiplicity of one. I observed a particular formula in terms of the tensor product. I think it warrants further comment here. In particular, we can write an analogous formula here for the 6×6 matrix above:

$$[T]_{\gamma, \gamma} = \begin{bmatrix} \alpha & \beta \\ -\beta & \alpha \end{bmatrix} \otimes \begin{bmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{bmatrix} + \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix} \otimes \begin{bmatrix} 0 & 1 & 0 \\ 0 & 0 & 1 \\ 0 & 0 & 0 \end{bmatrix}$$

If T has a 4-chain of generalized complex e-vectors then we expect the pattern continues to:

$$[T]_{\gamma, \gamma} = \begin{bmatrix} \alpha & \beta \\ -\beta & \alpha \end{bmatrix} \otimes \begin{bmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \end{bmatrix} + \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix} \otimes \begin{bmatrix} 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \\ 0 & 0 & 0 & 0 \end{bmatrix}.$$

The term built from tensoring with the superdiagonal matrix will be **nilpotent**. Perhaps we will explore this in the exercises. Hefferon or Damiano and Little etc. has a section if you wish a second opinion on all this.

Remark 4.9.4.

I'll abstain from writing the general real Jordan form of a matrix. Sufficient to say, it is block diagonal where each block is either formed as discussed thus far in this section or it is a Jordan block. Any real matrix A is similar to a unique matrix in real Jordan form up to the ordering of the blocks.

Example 4.9.5. To begin let's try an experiment using the e -vector and complex e -vectors for found in Example 4.8.14. We'll perform a similarity transformation based on this complex basis: $\beta = \{(i, 1, 0), (-i, 1, 0), (0, 0, 1)\}$. Notice that

$$[\beta] = \begin{bmatrix} i & -i & 0 \\ 1 & 1 & 0 \\ 0 & 0 & 1 \end{bmatrix} \Rightarrow [\beta]^{-1} = \frac{1}{2} \begin{bmatrix} -i & 1 & 0 \\ i & 1 & 0 \\ 0 & 0 & 2 \end{bmatrix}$$

Then, we can calculate that

$$[\beta]^{-1}A[\beta] = \frac{1}{2} \begin{bmatrix} -i & 1 & 0 \\ i & 1 & 0 \\ 0 & 0 & 2 \end{bmatrix} \begin{bmatrix} 1 & 1 & 0 \\ -1 & 1 & 0 \\ 0 & 0 & 3 \end{bmatrix} \begin{bmatrix} i & -i & 0 \\ 1 & 1 & 0 \\ 0 & 0 & 1 \end{bmatrix} = \begin{bmatrix} 1+i & 0 & 0 \\ 0 & 1-i & 0 \\ 0 & 0 & 3 \end{bmatrix}$$

Note that A is complex-diagonalizable in this case. Furthermore, A is already in real Jordan form.

We should take a moment to appreciate the significance of the 2×2 complex blocks in the real Jordan form of a matrix. It turns out there is a simple interpretation:

Example 4.9.6. Suppose $b \neq 0$ and $C = \begin{bmatrix} a & -b \\ b & a \end{bmatrix}$. We can calculate that $\det(A - \lambda I) = (a - \lambda)^2 + b^2 = 0$ hence we have complex eigenvalues $\lambda = a \pm ib$. Denoting $r = \sqrt{a^2 + b^2}$ (the modulus of $a + ib$). We can work out that

$$C = \begin{bmatrix} a & -b \\ b & a \end{bmatrix} = r \begin{bmatrix} a/r & -b/r \\ b/r & a/r \end{bmatrix} = \begin{bmatrix} r & 0 \\ 0 & r \end{bmatrix} \begin{bmatrix} \cos(\beta) & -\sin(\beta) \\ \sin(\beta) & \cos(\beta) \end{bmatrix}$$

Therefore, a 2×2 matrix with complex eigenvalue will factor into a dilation by the modulus of the e -value $|\lambda|$ times a rotation by the argument of the eigenvalue. If we write $\lambda = r \exp(i\beta)$ then we can identify that $r > 0$ is the modulus and β is an argument (there is degeneracy here because angle is multiply defined).

Transforming a given matrix by a similarity transformation into real Jordan form is a generally difficult calculation. On the other hand, reading the eigenvalues as well as geometric and algebraic multiplicities is a simple matter given an explicit matrix in real Jordan form.

Example 4.9.7. Suppose $A = \begin{bmatrix} 2 & 3 & 0 & 0 \\ -3 & 2 & 0 & 0 \\ 0 & 0 & 5 & 1 \\ 0 & 0 & 0 & 5 \end{bmatrix}$. I can read $\lambda_1 = 2 + 3i$ with geometric and algebraic multiplicity one and $\lambda_2 = 5$ with geometric multiplicity one and algebraic multiplicity two. Of course, $\lambda = 2 - 3i$ is also an e -value as complex e -values come in conjugate pairs.

Example 4.9.8. Suppose $A = \begin{bmatrix} 0 & 3 & 1 & 0 & 0 & 0 \\ -3 & 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 3 & 0 & 0 \\ 0 & 0 & -3 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 5 & 0 \\ 0 & 0 & 0 & 0 & 0 & 5 \end{bmatrix}$. I read $\lambda_1 = 3i$ with geometric multiplicity one and algebraic multiplicity two. Also $\lambda_2 = 5$ with geometric multiplicity and algebraic multiplicity two.

Let us conclude with introducing a standard notation for real Jordan blocks¹⁰.

Definition 4.9.9.

Suppose $\alpha + i\beta \in \mathbb{C}$ with $\beta \neq 0$ then

$$R_2(\alpha + i\beta) = \begin{bmatrix} \alpha & \beta \\ -\beta & \alpha \end{bmatrix} \quad \& \quad R_4(\alpha + i\beta) = \begin{bmatrix} \alpha & \beta & 1 & 0 \\ -\beta & \alpha & 0 & 1 \\ 0 & 0 & \alpha & \beta \\ 0 & 0 & -\beta & \alpha \end{bmatrix}$$

and the matrix given in Equation 4.2 defines $R_6(\alpha + i\beta)$. Generally,

$$R_{2k}(\alpha + i\beta) = R_2(\alpha + i\beta) \otimes I_k + I_k \otimes N_k.$$

4.10 systems of differential equations

Systems of differential equations are found at the base of many nontrivial questions in physics, math, biology, chemistry, nuclear engineering, economics, etc... Consider this, anytime a problem is described by several quantities which depend on time and each other it is likely that a simple conservation of mass, charge, population, particle number,... force linear relations between the time-rates of change of the quantities involved. This means, we get a system of differential equations. To be specific, Newton's Second Law is a system of differential equations. Maxwell's Equations are a system of differential equations. Now, generally, the methods we discover in this chapter will not allow solutions to problems I allude to above. Many of those problems are **nonlinear**. There are researchers who spend a good part of their career just unraveling the structure of a particular partial differential equation. That said, once simplifying assumptions are made and the problem is linearized one often faces the problem we solve in this chapter.

We show how to solve **any** system of first order differential equations with constant coefficients. This is accomplished by the application of Jordan basis for the matrix of the system to the **matrix exponential**. I'm not sure the exact history of the method I show in this chapter. In my opinion, the manner in which the chains of generalized eigenvectors tame the matrix exponential are reason enough to study them. I have left some redundant definitions in this chapter, I hope it makes this more readable. I would default to the earlier definition if a definition given in this chapter seems in disagreement on some point. If in doubt, please ask.

¹⁰please forgive me for not giving a proper definition of the Kronecker product $\otimes$, I hope the pattern is clear from the examples in this section

I should mention, the results of this Chapter allow generalization. We could develop theorems for calculus of an $\mathbb{F}$ -valued function of a real variable. But, we content ourselves to focus on $\mathbb{R}$ and $\mathbb{C}$ as is convenient to the applications of interest.

4.10.1 calculus of matrices

A more apt title would be "calculus of matrix-valued functions of a real variable".

Definition 4.10.1.

A matrix-valued function of a real variable is a function from $I \subseteq \mathbb{R}$ to $\mathbb{R}^{m \times n}$. Suppose $A : I \subseteq \mathbb{R} \rightarrow \mathbb{R}^{m \times n}$ is such that $A_{ij} : I \subseteq \mathbb{R} \rightarrow \mathbb{R}$ is differentiable for each i, j then we define

$$\frac{dA}{dt} = \left[\frac{dA_{ij}}{dt} \right]$$

which can also be denoted $(A')_{ij} = A'_{ij}$. We likewise define $\int A dt = [\int A_{ij} dt]$ for A with integrable components. Definite integrals and higher derivatives are also defined component-wise.

Example 4.10.2. Suppose $A(t) = \begin{bmatrix} 2t & 3t^2 \\ 4t^3 & 5t^4 \end{bmatrix}$. I'll calculate a few items just to illustrate the definition above. calculate; to differentiate a matrix we differentiate each component one at a time:

$$A'(t) = \begin{bmatrix} 2 & 6t \\ 12t^2 & 20t^3 \end{bmatrix} \quad A''(t) = \begin{bmatrix} 0 & 6 \\ 24t & 60t^2 \end{bmatrix} \quad A'(0) = \begin{bmatrix} 2 & 0 \\ 0 & 0 \end{bmatrix}$$

Integrate by integrating each component:

$$\int A(t) dt = \begin{bmatrix} t^2 + c_1 & t^3 + c_2 \\ t^4 + c_3 & t^5 + c_4 \end{bmatrix} \quad \int_0^2 A(t) dt = \begin{bmatrix} t^2|_0^2 & t^3|_0^2 \\ t^4|_0^2 & t^5|_0^2 \end{bmatrix} = \begin{bmatrix} 4 & 8 \\ 16 & 32 \end{bmatrix}$$

Proposition 4.10.3.

Suppose A, B are matrix-valued functions of a real variable, f is a function of a real variable, c is a constant, and C is a constant matrix then

- (1.) $(AB)' = A'B + AB'$ (product rule for matrices)
- (2.) $(AC)' = A'C$
- (3.) $(CA)' = CA'$
- (4.) $(fA)' = f'A + fA'$
- (5.) $(cA)' = cA'$
- (6.) $(A + B)' = A' + B'$

where each of the functions is evaluated at the same time t and I assume that the functions and matrices are differentiable at that value of t and of course the matrices A, B, C are such that the multiplications are well-defined.

Proof: Suppose $A(t) \in \mathbb{R}^{m \times n}$ and $B(t) \in \mathbb{R}^{n \times p}$ consider,

$$\begin{aligned}
 (AB)'_{ij} &= \frac{d}{dt}((AB)_{ij}) && \text{defn. derivative of matrix} \\
 &= \frac{d}{dt}(\sum_k A_{ik} B_{kj}) && \text{defn. of matrix multiplication} \\
 &= \sum_k \frac{d}{dt}(A_{ik} B_{kj}) && \text{linearity of derivative} \\
 &= \sum_k \left[\frac{dA_{ik}}{dt} B_{kj} + A_{ik} \frac{dB_{kj}}{dt} \right] && \text{ordinary product rules} \\
 &= \sum_k \frac{dA_{ik}}{dt} B_{kj} + \sum_k A_{ik} \frac{dB_{kj}}{dt} && \text{algebra} \\
 &= (A'B)_{ij} + (AB')_{ij} && \text{defn. of matrix multiplication} \\
 &= (A'B + AB')_{ij} && \text{defn. matrix addition}
 \end{aligned}$$

this proves (1.) as i, j were arbitrary in the calculation above. The proof of (2.) and (3.) follow quickly from (1.) since C constant means $C' = 0$. Proof of (4.) is similar to (1.):

$$\begin{aligned}
 (fA)'_{ij} &= \frac{d}{dt}((fA)_{ij}) && \text{defn. derivative of matrix} \\
 &= \frac{d}{dt}(fA_{ij}) && \text{defn. of scalar multiplication} \\
 &= \frac{df}{dt}A_{ij} + f \frac{dA_{ij}}{dt} && \text{ordinary product rule} \\
 &= \left(\frac{df}{dt}A + f \frac{dA}{dt} \right)_{ij} && \text{defn. matrix addition} \\
 &= \left(\frac{df}{dt}A + f \frac{dA}{dt} \right)_{ij} && \text{defn. scalar multiplication.}
 \end{aligned}$$

The proof of (5.) follows from taking $f(t) = c$ which has $f' = 0$. I leave the proof of (6.) as an exercise for the reader. $\square$.

To summarize: the calculus of matrices is the same as the calculus of functions with the small qualifier that we must respect the rules of matrix algebra. The noncommutativity of matrix multiplication is the main distinguishing feature.

Since we're discussing this type of differentiation perhaps it would be worthwhile for me to insert a comment about complex functions here. Differentiation of functions from $\mathbb{R}$ to $\mathbb{C}$ is defined by splitting a given function into its real and imaginary parts then we just differentiate with respect to the real variable one component at a time. For example:

$$\begin{aligned}
 \frac{d}{dt}(e^{2t} \cos(t) + ie^{2t} \sin(t)) &= \frac{d}{dt}(e^{2t} \cos(t)) + i \frac{d}{dt}(e^{2t} \sin(t)) \\
 &= (2e^{2t} \cos(t) - e^{2t} \sin(t)) + i(2e^{2t} \sin(t) + e^{2t} \cos(t)) \\
 &= e^{2t}(2 + i)(\cos(t) + i \sin(t)) = (2 + i)e^{(2+i)t}.
 \end{aligned}$$

where I have made use of the identity¹¹ $e^{x+iy} = e^x(\cos(y) + i \sin(y))$. We just saw that $\frac{d}{dt}e^{\lambda t} = \lambda e^{\lambda t}$ which seems obvious enough until you appreciate that we just proved it for $\lambda = 2 + i$. We make use of this calculation in the next section in the case we have complex e-values.

4.10.2 the normal form and theory for systems

A system of ODEs in normal form is a finite collection of first order ODEs which share dependent variables and a single independent variable.

$$1. \ (n = 1) \ \frac{dx}{dt} = A_{11}x + f$$

¹¹or definition, depending on how you choose to set-up the complex exponential, I take this as the definition in calculus II

2. ($n = 2$) $\frac{dx}{dt} = A_{11}x + A_{12}y + f_1$ and $\frac{dy}{dt} = A_{21}x + A_{22}y + f_2$ we can express this in **matrix normal form** as follows, use $x = x_1$ and $y = x_2$,

$$\begin{bmatrix} \frac{dx_1}{dt} \\ \frac{dx_2}{dt} \end{bmatrix} = \begin{bmatrix} A_{11} & A_{12} \\ A_{21} & A_{22} \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \end{bmatrix} + \begin{bmatrix} f_1 \\ f_2 \end{bmatrix}$$

This is nicely abbreviated by writing $d\vec{x}/dt = A\vec{x} + \vec{f}$ where $\vec{x} = (x_1, x_2)$ and $\vec{f} = (f_1, f_2)$ whereas the 2×2 matrix A is called the **coefficient matrix** of this system.

3. ($n = 3$) The matrix normal form is simply

$$\begin{bmatrix} \frac{dx_1}{dt} \\ \frac{dx_2}{dt} \\ \frac{dx_3}{dt} \end{bmatrix} = \begin{bmatrix} A_{11} & A_{12} & A_{13} \\ A_{21} & A_{22} & A_{23} \\ A_{31} & A_{32} & A_{33} \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \\ x_3 \end{bmatrix} + \begin{bmatrix} f_1 \\ f_2 \\ f_3 \end{bmatrix}$$

Expanded into **scalar normal form** we have $\frac{dx_1}{dt} = A_{11}x_1 + A_{12}x_2 + A_{13}x_3 + f_1$ and $\frac{dx_2}{dt} = A_{21}x_1 + A_{22}x_2 + A_{23}x_3 + f_2$ and $\frac{dx_3}{dt} = A_{31}x_1 + A_{32}x_2 + A_{33}x_3 + f_3$.

Generally an n -th order system of ODEs in normal form on an interval $I \subseteq \mathbb{R}$ can be written as $\frac{dx_i}{dt} = \sum_{j=1}^n A_{ij}x_j + f_i$ for **coefficient functions** $A_{ij} : I \subseteq \mathbb{R} \rightarrow \mathbb{R}$ and **forcing functions** $f_i : I \subseteq \mathbb{R} \rightarrow \mathbb{R}$. You might consider the problem of solving a system of k -first order differential equations in n -dependent variables where $n \neq k$, however, we do not discuss such over or underdetermined problems in these notes. That said, the concept of a system of differential equations in normal form is perhaps more general than you expect. Let me illustrate this by example. I'll start with a single second order ODE:

Example 4.10.4. Consider $ay'' + by' + cy = f$. We define $x_1 = y$ and $x_2 = y'$. Observe that

$$x'_1 = x_2 \quad \& \quad x'_2 = y'' = -\frac{1}{a}(f - by' - cy) = \frac{1}{a}(f - bx_2 - cx_1)$$

Thus,

$$\begin{bmatrix} x'_1 \\ x'_2 \end{bmatrix} = \begin{bmatrix} 0 & 1 \\ -c/a & -b/a \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \end{bmatrix} + \begin{bmatrix} 0 \\ f/a \end{bmatrix}$$

The matrix $\begin{bmatrix} 0 & 1 \\ -c/a & -b/a \end{bmatrix}$ is called the **companion matrix** of the second order ODE $ay'' + by' + cy = f$.

The example above nicely generalizes to the general n -th order linear ODE.

Example 4.10.5. Consider $a_0y^{(n)} + a_1y^{(n-1)} + \cdots + a_{n-1}y' + a_ny = f$. Introduce variables to reduce the order:

$$x_1 = y, \quad x_2 = y', \quad x_3 = y'', \quad \dots \quad x_n = y^{(n-1)}$$

From which is clear that $x'_1 = x_2$ and $x'_2 = x_3$ continuing up to $x'_{n-1} = x_n$ and $x'_n = y^{(n)}$. Hence,

$$x'_n = -\frac{a_1}{a_0}x_n - \cdots - \frac{a_{n-1}}{a_0}x_2 - \frac{a_n}{a_0}x_1 + \frac{f}{a_0}$$

Once again the matrix below is called the **companion matrix** of the given n -th order ODE.

$$\begin{bmatrix} x'_1 \\ x'_2 \\ \vdots \\ x'_{n-1} \\ x'_n \end{bmatrix} = \begin{bmatrix} 0 & 1 & 0 & \cdots & 0 & 0 \\ 0 & 0 & 1 & \cdots & 0 & 0 \\ \vdots & \vdots & \vdots & \cdots & \vdots & \vdots \\ 0 & 0 & 0 & \cdots & 0 & 1 \\ -\frac{a_n}{a_0} & -\frac{a_{n-1}}{a_0} & -\frac{a_{n-2}}{a_0} & \cdots & -\frac{a_2}{a_0} & -\frac{a_1}{a_0} \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \\ \vdots \\ x_{n-1} \\ x_n \end{bmatrix} + \begin{bmatrix} 0 \\ 0 \\ \vdots \\ 0 \\ \frac{f}{a_0} \end{bmatrix}$$

The problem of many higher order ODEs is likewise confronted by introducing variables to reduce the order.

Example 4.10.6. Consider $y'' + 3x' = \sin(t)$ and $x'' + 6y' - x = e^t$. We begin with a system of two second order differential equations. Introduce new variables:

$$x_1 = x, \quad x_2 = y, \quad x_3 = x', \quad x_4 = y'$$

It follows that $x'_3 = x''$ and $x'_4 = y''$ whereas $x'_1 = x_3$ and $x'_2 = x_4$. We convert the given differential equations to first order ODEs:

$$x'_4 + 3x_3 = \sin(t) \quad \& \quad x'_3 + 6x_4 - x_1 = e^t$$

Let us collect these results as a matrix problem:

$$\begin{bmatrix} x'_1 \\ x'_2 \\ x'_3 \\ x'_4 \end{bmatrix} = \begin{bmatrix} 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \\ 1 & 0 & 0 & 6 \\ 0 & 0 & -3 & 0 \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \\ x_3 \\ x_4 \end{bmatrix} + \begin{bmatrix} 0 \\ 0 \\ e^t \\ \sin(t) \end{bmatrix}$$

Generally speaking the order of the normal form corresponding to a system of higher order ODE will simply be the sum of the orders of the systems (assuming the given system has no redundancies; for example $x'' + y'' = x$ and $x'' - x = -y''$ are redundant). I will not prove the following assertion, however, it should be fairly clear why it is true given the examples thus far discussed:

Proposition 4.10.7. *linear systems have a normal form.*

A given systems of linear ODEs may be converted to an equivalent system of first order ODEs in normal form.

For this reason the first order problem will occupy the majority of our time. That said, the method of the next section is applicable to any order.

Since normal forms are essentially general it is worthwhile to state the theory which will guide our work. I do not offer all the proof here, but you can find proof in many texts. For example, in Nagel Saff and Snider these theorems are given in §9.4 and are proved in Chapter 13.

Definition 4.10.8. *linear independence of vector-valued functions*

Suppose $\vec{v}_j : I \subseteq \mathbb{R} \rightarrow \mathbb{R}^n$ is a function for $j = 1, 2, \dots, k$ then we say that $\{\vec{v}_1, \vec{v}_2, \dots, \vec{v}_k\}$ is linearly independent on I iff $\sum_{j=1}^k c_j \vec{v}_j(t) = 0$ for all $t \in I$ implies $c_j = 0$ for $j = 1, 2, \dots, k$.

We can use the determinant to test LI of a set of n -vectors which are all n -dimensional vectors. It is true that $\{\vec{v}_1, \vec{v}_2, \dots, \vec{v}_n\}$ is LI on I iff $\det[\vec{v}_1(t)|\vec{v}_2(t)|\dots|\vec{v}_n(t)] \neq 0$ for all $t \in I$.

Definition 4.10.9. *wronskian for vector-valued functions of a real variable.*

Suppose $\vec{v}_j : I \subseteq \mathbb{R} \rightarrow \mathbb{R}^n$ is differentiable for $j = 1, 2, \dots, n$. The **Wronskian** is defined by $W(\vec{v}_1, \vec{v}_2, \dots, \vec{v}_n; t) = \det[\vec{v}_1|\vec{v}_2|\dots|\vec{v}_n]$ for each $t \in I$.

Theorems for wronskians of solutions sets mirror those already discussed for the n -th order problem.

Definition 4.10.10. *solution and homogeneous solutions of $d\vec{x}/dt = A\vec{x} + \vec{f}$*

Let $A : I \rightarrow \mathbb{R}^{n \times n}$ and $\vec{f} : I \rightarrow \mathbb{R}^n$ be continuous. A **solution** of $d\vec{v}/dt = A\vec{v} + \vec{f}$ on $I \subseteq \mathbb{R}$ is a vector-valued function $\vec{x} : I \rightarrow \mathbb{R}^n$ such that $d\vec{x}/dt = A\vec{x} + \vec{f}$ for all $t \in I$. A **homogeneous solution** on $I \subseteq \mathbb{R}$ is a solution of $d\vec{v}/dt = A\vec{v}$.

In the example below we see three LI homogeneous solutions and a single particular solution.

Example 4.10.11. Suppose $x' = x - 1$, $y' = 2y - 2$ and $z' = 3z - 3$. In matrix normal form we face:

$$\begin{bmatrix} x' \\ y' \\ z' \end{bmatrix} = \begin{bmatrix} 1 & 0 & 0 \\ 0 & 2 & 0 \\ 0 & 0 & 3 \end{bmatrix} \begin{bmatrix} x \\ y \\ z \end{bmatrix} + \begin{bmatrix} -1 \\ -2 \\ -3 \end{bmatrix}$$

It is easy to show by separately solving the the DEqns that $x = c_1 e^t + 1$, $y = c_2 e^{2t} + 2$ and $z = c_3 e^{3t} + 3$. In vector notation the solution is

$$\vec{x}(t) = \begin{bmatrix} c_1 e^t + 1 \\ c_2 e^{2t} + 2 \\ c_3 e^{3t} + 3 \end{bmatrix} = c_1 \begin{bmatrix} e^t \\ 0 \\ 0 \end{bmatrix} + c_2 \begin{bmatrix} 0 \\ e^{2t} \\ 0 \end{bmatrix} + c_3 \begin{bmatrix} 0 \\ 0 \\ e^{3t} \end{bmatrix} + \begin{bmatrix} 1 \\ 2 \\ 3 \end{bmatrix}$$

I invite the reader to show that $S = \{\vec{x}_1, \vec{x}_2, \vec{x}_3\}$ is LI on $\mathbb{R}$ where $\vec{x}_1(t) = \langle e^t, 0, 0 \rangle$, $\vec{x}_2(t) = \langle 0, e^{2t}, 0 \rangle$ and $\vec{x}_3(t) = \langle 0, 0, e^{3t} \rangle$. On the other hand, $\vec{x}_p = \langle 1, 2, 3 \rangle$ is a particular solution to the given problem.

In truth, any choice of c_1, c_2, c_3 with at least one nonzero constant will produce a homogeneous solution. To obtain the solutions I pointed out in the example you can choose $c_1 = 1, c_2 = 0, c_3 = 0$ to obtain $\vec{x}_1(t) = \langle e^t, 0, 0 \rangle$ or $c_1 = 0, c_2 = 1, c_3 = 0$ to obtain $\vec{x}_2(t) = \langle 0, e^{2t}, 0 \rangle$ or $c_1 = 0, c_2 = 0, c_3 = 1$ to obtain $\vec{x}_3(t) = \langle 0, 0, e^{3t} \rangle$.

Definition 4.10.12. *fundamental solution set of a linear system $d\vec{x}/dt = A\vec{x} + \vec{f}$*

Let $A : I \rightarrow \mathbb{R}^{n \times n}$ and $\vec{f} : I \rightarrow \mathbb{R}^n$ be continuous. A **fundamental solution set** on $I \subseteq \mathbb{R}$ is a set of n -homogeneous solutions of $d\vec{v}/dt = A\vec{v} + \vec{f}$ for which $\{\vec{x}_1, \vec{x}_2, \dots, \vec{x}_n\}$ is a LI set on I . A **solution matrix** on $I \subseteq \mathbb{R}$ is a matrix X is a matrix for which each column is a homogeneous solution on I . A **fundamental matrix** on $I \subseteq \mathbb{R}$ is an invertible solution matrix.

Example 4.10.13. Continue Example 4.10.11. Note that $S = \{\vec{x}_1, \vec{x}_2, \vec{x}_3\}$ is a fundamental solution set. The fundamental solution matrix is found by concatenating $\vec{x}_1, \vec{x}_2$ and $\vec{x}_3$:

$$X = [\vec{x}_1 | \vec{x}_2 | \vec{x}_3] = \begin{bmatrix} e^t & 0 & 0 \\ 0 & e^{2t} & 0 \\ 0 & 0 & e^{3t} \end{bmatrix}$$

Observe $\det(X) = e^t e^{2t} e^{3t} = e^{6t} \neq 0$ on $\mathbb{R}$ hence X is invertible on $\mathbb{R}$.

Example 4.10.14. Let $A = \begin{bmatrix} 0 & 0 & -4 \\ 2 & 4 & 2 \\ 2 & 0 & 6 \end{bmatrix}$ define the system of DEqns $\frac{d\vec{x}}{dt} = A\vec{x}$. I claim that the

matrix $X(t) = \begin{bmatrix} 0 & -e^{4t} & -2e^{2t} \\ e^{4t} & 0 & e^{2t} \\ 0 & e^{4t} & e^{2t} \end{bmatrix}$ is a solution matrix. Calculate,

$$AX = \begin{bmatrix} 0 & 0 & -4 \\ 2 & 4 & 2 \\ 2 & 0 & 6 \end{bmatrix} \begin{bmatrix} 0 & -e^{4t} & -2e^{2t} \\ e^{4t} & 0 & e^{2t} \\ 0 & e^{4t} & e^{2t} \end{bmatrix} = \begin{bmatrix} 0 & -4e^{4t} & -4e^{2t} \\ 4e^{4t} & 0 & 2e^{2t} \\ 0 & 4e^{4t} & 2e^{2t} \end{bmatrix}.$$

On the other hand, differentiation yields $X' = \begin{bmatrix} 0 & -4e^{4t} & -4e^{2t} \\ 4e^{4t} & 0 & 2e^{2t} \\ 0 & 4e^{4t} & 2e^{2t} \end{bmatrix}$. Therefore $X' = AX$.

Notice that if we express X in terms of its columns $X = [\vec{x}_1 | \vec{x}_2 | \vec{x}_3]$ then it follows that $X' = [\vec{x}_1' | \vec{x}_2' | \vec{x}_3']$ and $AX = A[\vec{x}_1 | \vec{x}_2 | \vec{x}_3] = [A\vec{x}_1 | A\vec{x}_2 | A\vec{x}_3]$ hence

$$\vec{x}_1' = A\vec{x}_1 \quad \& \quad \vec{x}_2' = A\vec{x}_2 \quad \& \quad \vec{x}_3' = A\vec{x}_3$$

We find that $\vec{x}_1(t) = \langle 0, e^{4t}, 0 \rangle$, $\vec{x}_2(t) = \langle -e^{4t}, 0, e^{4t} \rangle$ and $\vec{x}_3(t) = \langle -2e^{2t}, e^{2t}, e^{2t} \rangle$ form a fundamental solution set for the given system of DEqns.

Theorem 4.10.15. Let $A : I \rightarrow \mathbb{R}^{n \times n}$ and $\vec{f} : I \rightarrow \mathbb{R}^n$ be continuous.

- (1.) there exists a fundamental solution set $\{\vec{x}_1, \vec{x}_2, \dots, \vec{x}_n\}$ on I
- (2.) if $t_o \in I$ and $\vec{x}_o$ is a given initial condition vector then there exists a unique solution $\vec{x}$ on I such that $\vec{x}(t_o) = \vec{x}_o$
- (3.) the **general solution** has the form $\vec{x} = \vec{x}_h + \vec{x}_p$ where $\vec{x}_p$ is a **particular solution** and $\vec{x}_h$ is the **homogeneous solution** is formed by a real linear combination of the fundamental solution set ($\vec{x}_h = c_1\vec{x}_1 + c_2\vec{x}_2 + \dots + c_n\vec{x}_n$)

The term *general solution* is intended to indicate that the formula given includes all possible solutions to the problem. Part (2.) of the theorem indicates that there must be some 1-1 correspondance between a given initial condition and the choice of the constants $c_1, c_2, \dots, c_n$ with respect to a given fundamental solution set. Observe that if we define $\vec{c} = [c_1, c_2, \dots, c_n]^T$ and the fundamental matrix $X = [\vec{x}_1 | \vec{x}_2 | \dots | \vec{x}_n]$ we can express the homogeneous solution via a matrix-vector product:

$$\vec{x}_h = X\vec{c} = c_1\vec{x}_1 + c_2\vec{x}_2 + \dots + c_n\vec{x}_n \quad \Rightarrow \quad \boxed{\vec{x}(t) = X(t)\vec{c} + \vec{x}_p(t)}$$

Further suppose that we wish to set $\vec{x}(t_o) = \vec{x}_o$. We need to solve for $\vec{c}$:

$$\vec{x}_o = X(t_o)\vec{c} + \vec{x}_p(t_o) \quad \Rightarrow \quad X(t_o)\vec{c} = \vec{x}_o - \vec{x}_p(t_o)$$

Since $X^{-1}(t_o)$ exists we can multiply by the inverse on the right and find

$$\vec{c} = X^{-1}(t_o)[\vec{x}_o - \vec{x}_p(t_o)]$$

Next, place the result above back in the general solution to derive

$$\boxed{\vec{x}(t) = X(t)X^{-1}(t_o)[\vec{x}_o - \vec{x}_p(t_o)] + \vec{x}_p(t)}$$

We can further simplify this general formula in the constant coefficient case, or in the study of variation of parameters for systems. Note that in the homogeneous case this gives us a clean formula to calculate the constants to fit initial data:

$$\boxed{\vec{x}(t) = X(t)X^{-1}(t_o)\vec{x}_o} \quad (\text{homogeneous case})$$

Example 4.10.16. We found $x' = -y$ and $y' = x$ had solutions $x(t) = c_1 \cos(t) + c_2 \sin(t)$ and $y(t) = c_1 \sin(t) - c_2 \cos(t)$. It follows that $X(t) = \begin{bmatrix} \cos(t) & \sin(t) \\ \sin(t) & -\cos(t) \end{bmatrix}$. Calculate that $\det(X) = -1$ to see that $X^{-1}(t) = \begin{bmatrix} \cos(t) & \sin(t) \\ \sin(t) & -\cos(t) \end{bmatrix}$. Suppose we want the solution through (a, b) at time t_o then the solution is given by

$$\vec{x}(t) = X(t)X^{-1}(t_o)\vec{x}_o = \begin{bmatrix} \cos(t) & \sin(t) \\ \sin(t) & -\cos(t) \end{bmatrix} \begin{bmatrix} \cos(t_o) & \sin(t_o) \\ \sin(t_o) & -\cos(t_o) \end{bmatrix} \begin{bmatrix} a \\ b \end{bmatrix}.$$

This concludes our brief tour of the theory for systems of ODEs. Clearly we have two main goals past this point (1.) find the fundamental solution set (2.) find the particular solution.

4.10.3 solutions by eigenvector

We narrow our focus at this point: our goal is to find nontrivial¹² solutions to the homogeneous constant coefficient problem $\frac{d\vec{x}}{dt} = A\vec{x}$ where $A \in \mathbb{R}^{n \times n}$. A reasonable ansatz for this problem is that the solution should have the form $\vec{x} = e^{\lambda t}\vec{u}$ for some constant scalar λ and some constant vector $\vec{u}$. If such solutions exist then what conditions must we place on λ and $\vec{u}$? To begin clearly $\vec{u} \neq 0$ since we are seeking nontrivial solutions. Differentiate,

$$\frac{d}{dt}[e^{\lambda t}\vec{u}] = [e^{\lambda t}]\vec{u} = \lambda e^{\lambda t}\vec{u}$$

Hence $\frac{d\vec{x}}{dt} = A\vec{x}$ implies $\lambda e^{\lambda t}\vec{u} = Ae^{\lambda t}\vec{u}$. However, $e^{\lambda t} \neq 0$ hence we find $\lambda\vec{u} = A\vec{u}$. We can write the vector $\lambda\vec{u}$ as a matrix product with identity matrix I ; $\lambda\vec{u} = \lambda I\vec{u}$. Therefore, we find

$$(A - \lambda I)\vec{u} = 0$$

to be a necessary condition for the solution. Note that the system of linear equations defined by $(A - \lambda I)\vec{u} = 0$ is consistent since 0 is a solution. It follows that for $\vec{u} \neq 0$ to be a solution we must have that the matrix $(A - \lambda I)$ is singular. It follows that we find

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

a necessary condition for our solution. Moreover, for a given matrix A this is nothing more than an n -th order polynomial in λ hence there are at most n -distinct solutions for λ . The equation $\det(A - \lambda I) = 0$ is called the **characteristic equation** of A and the solutions are called **eigenvalues**. The nontrivial vector $\vec{u}$ such that $(A - \lambda I)\vec{u} = 0$ is called the **eigenvector** with **eigenvalue** λ . We often abbreviate these by referring to "e-vectors" or "e-values". Many interesting theorems are known for eigenvectors, see a linear algebra text or my linear notes for elaboration on this point.

Definition 4.10.17. eigenvalues and eigenvectors

Suppose A is an $n \times n$ matrix then we say $\lambda \in \mathbb{C}$ which is a solution of $\det(A - \lambda I) = 0$ is an **eigenvalue of A** . Given such an eigenvalue λ a nonzero vector $\vec{u}$ such that $(A - \lambda I)\vec{u} = 0$ is called an **eigenvector** with eigenvalue λ .

¹²nontrivial simply means the solution is not identically zero. The zero solution does exist, but it is not the solution we are looking for...

Example 4.10.18. Problem: find the fundamental solutions of the system $x' = -4x - y$ and $y' = 5x + 2y$

Solution: we seek to solve $\frac{d\vec{x}}{dt} = A\vec{x}$ where $A = \begin{bmatrix} -4 & -1 \\ 5 & 2 \end{bmatrix}$. Consider the characteristic equation:

$$\begin{aligned} \det(A - \lambda I) &= \det \begin{bmatrix} -4 - \lambda & -1 \\ 5 & 2 - \lambda \end{bmatrix} \\ &= (-4 - \lambda)(2 - \lambda) + 5 \\ &= \lambda^2 + 2\lambda - 3 \\ &= (\lambda + 3)(\lambda - 1) \\ &= 0 \end{aligned}$$

We find $\lambda_1 = 1$ and $\lambda_2 = -3$. Next calculate the e-vectors for each e-value. We seek $\vec{u}_1 = [u, v]^T$ such that $(A - I)\vec{u}_1 = 0$ thus solve:

$$\begin{bmatrix} -5 & -1 \\ 5 & 1 \end{bmatrix} \begin{bmatrix} u \\ v \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \end{bmatrix} \Rightarrow -5u - v = 0 \Rightarrow v = -5u, u \neq 0 \Rightarrow \vec{u}_1 = \begin{bmatrix} u \\ -5u \end{bmatrix}$$

Naturally we can write $\vec{u}_1 = u[1, -5]^T$ and for convenience we set $u = 1$ and find $\vec{u}_1 = [1, -5]^T$ which gives us the fundamental solution $\boxed{\vec{x}_1(t) = e^t[1, -5]^T}$. Continue¹³ to the next e-value $\lambda_2 = -3$ we seek $\vec{u}_2 = [u, v]^T$ such that $(A + 3I)\vec{u}_2 = 0$.

$$\begin{bmatrix} -1 & -1 \\ 5 & 5 \end{bmatrix} \begin{bmatrix} u \\ v \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \end{bmatrix} \Rightarrow -u - v = 0 \Rightarrow v = -u, u \neq 0 \Rightarrow \vec{u}_2 = \begin{bmatrix} u \\ -u \end{bmatrix}$$

Naturally we can write $\vec{u}_2 = u[1, -1]^T$ and for convenience we set $u = 1$ and find $\vec{u}_2 = [1, -1]^T$ which gives us the fundamental solution $\boxed{\vec{x}_2(t) = e^{-3t}[1, -1]^T}$. The fundamental solution set is given by $\{\vec{x}_1, \vec{x}_2\}$ and the domains of these solution clearly extend to all of $\mathbb{R}$.

We can assemble the general solution as a linear combination of the fundamental solutions $\vec{x}(t) = c_1\vec{x}_1 + c_2\vec{x}_2$. In particular this yields

$$\vec{x}(t) = c_1\vec{x}_1 + c_2\vec{x}_2 = c_1e^t \begin{bmatrix} 1 \\ -5 \end{bmatrix} + c_2e^{-3t} \begin{bmatrix} 1 \\ -1 \end{bmatrix} = \begin{bmatrix} c_1e^t + c_2e^{-3t} \\ -5c_1e^t - c_2e^{-3t} \end{bmatrix}$$

Thus the system $x' = -4x - y$ and $y' = 5x + 2y$ has **scalar** solutions $x(t)c_1e^t + c_2e^{-3t}$ and $y(t) = -5c_1e^t - c_2e^{-3t}$. Finally, a fundamental matrix for this problem is given by

$$X = [\vec{x}_1 | \vec{x}_2] = \begin{bmatrix} e^t & e^{-3t} \\ -5e^t & -e^{-3t} \end{bmatrix}.$$

¹³the upcoming u, v are not the same as those I just worked out, I call these letters disposable variables because I like to reuse them in several ways in a particular example where we repeat the e-vector calculation over several e-values.

Example 4.10.19. Problem: find the fundamental solutions of the system $x' = -3x$ and $y' = -3y$

Solution: we seek to solve $\frac{d\vec{x}}{dt} = A\vec{x}$ where $A = \begin{bmatrix} -3 & 0 \\ 0 & -3 \end{bmatrix}$. Consider the characteristic equation:

$$\det(A - \lambda I) = \det \begin{bmatrix} -3 - \lambda & 0 \\ 0 & -3 - \lambda \end{bmatrix} = (\lambda + 3)^2 = 0$$

We find $\lambda_1 = -3$ and $\lambda_2 = -3$. Finding the eigenvectors here offers an unusual algebra problem; to find $\vec{u}$ with e -value $\lambda = -3$ we should find nontrivial solutions of $(A + 3I)\vec{u} = \begin{bmatrix} 0 & 0 \\ 0 & 0 \end{bmatrix} \begin{bmatrix} u \\ v \end{bmatrix} = 0$.

We find no condition on $\vec{u}$. It follows that **any** nonzero vector is an eigenvector of A . Indeed, note that $A = -3I$ and $A\vec{u} = -3I\vec{u}$ hence $(A + 3I)\vec{u} = 0$. Convenient choices for $\vec{u}$ are $[1, 0]^T$ and $[0, 1]^T$ hence we find fundamental solutions:

$$\vec{x}_1(t) = e^{-3t} \begin{bmatrix} 1 \\ 0 \end{bmatrix} = \begin{bmatrix} e^{-3t} \\ 0 \end{bmatrix} \quad \& \quad \vec{x}_2(t) = e^{-3t} \begin{bmatrix} 0 \\ 1 \end{bmatrix} = \begin{bmatrix} 0 \\ e^{-3t} \end{bmatrix}.$$

We can assemble the general solution as a linear combination of the fundamental solutions

$\vec{x}(t) = c_1 \begin{bmatrix} e^{-3t} \\ 0 \end{bmatrix} + c_2 \begin{bmatrix} 0 \\ e^{-3t} \end{bmatrix}$. Thus the system $x' = -3x$ and $y' = -3y$ has **scalar** solutions $x(t) = c_1 e^{-3t}$ and $y(t) = c_2 e^{-3t}$. Finally, a fundamental matrix for this problem is given by

$$X = [\vec{x}_1 | \vec{x}_2] = \begin{bmatrix} e^{-3t} & 0 \\ 0 & e^{-3t} \end{bmatrix}.$$

Example 4.10.20. Problem: find the fundamental solutions of the system $x' = 3x + y$ and $y' = -4x - y$

Solution: we seek to solve $\frac{d\vec{x}}{dt} = A\vec{x}$ where $A = \begin{bmatrix} 3 & 1 \\ -4 & -1 \end{bmatrix}$. Consider the characteristic equation:

$$\begin{aligned} \det(A - \lambda I) &= \det \begin{bmatrix} 3 - \lambda & 1 \\ -4 & -1 - \lambda \end{bmatrix} \\ &= (\lambda - 3)(\lambda + 1) + 4 \\ &= \lambda^2 - 2\lambda + 1 \\ &= (\lambda - 1)^2 \\ &= 0 \end{aligned}$$

We find $\lambda_1 = 1$ and $\lambda_2 = 1$. Let us find the e -vector $\vec{u}_1 = [u, v]^T$ such that $(A - I)\vec{u}_1 = 0$

$$\begin{bmatrix} 2 & 1 \\ -4 & -2 \end{bmatrix} \begin{bmatrix} u \\ v \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \end{bmatrix} \Rightarrow 2u + v = 0 \Rightarrow v = -2u, \quad u \neq 0 \Rightarrow \vec{u}_1 = \begin{bmatrix} u \\ -2u \end{bmatrix}$$

We choose $u = 1$ for convenience and thus find the fundamental solution $\vec{x}_1(t) = e^t \begin{bmatrix} 1 \\ -2 \end{bmatrix}$.

Remark 4.10.21.

In the previous example the **algebraic multiplicity** of the e-value $\lambda = 1$ was 2. However, we found only one LI e-vector. This means the **geometric multiplicity** for $\lambda = 1$ is only 1. This means we are missing a vector and hence a fundamental solution. Where is $\vec{x}_2$ which is LI from the $\vec{x}_1$ we just found? This question is ultimately answered via the matrix exponential.

Example 4.10.22. Problem: find the fundamental solutions of the system $x' = -y$ and $y' = 4x$

Solution: we seek to solve $\frac{d\vec{x}}{dt} = A\vec{x}$ where $A = \begin{bmatrix} 0 & -1 \\ 4 & 0 \end{bmatrix}$. Consider the characteristic equation:

$$\det(A - \lambda I) = \det \begin{bmatrix} -\lambda & -1 \\ 4 & -\lambda \end{bmatrix} = \lambda^2 + 4 = 0 \Rightarrow \lambda = \pm 2i.$$

This e-value is a **pure imaginary** number which is a special type of **complex number** where there is no real part. Careful review of the arguments that framed the e-vector solution reveal that the same calculations apply when either λ or $\vec{u}$ are complex. With this in mind we seek the e-vector for $\lambda = 2i$: let us find the e-vector $\vec{u}_1 = [u, v]^T$ such that $(A - 2iI)\vec{u}_1 = 0$

$$\begin{bmatrix} -2i & -1 \\ 4 & -2i \end{bmatrix} \begin{bmatrix} u \\ v \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \end{bmatrix} \Rightarrow 2iu - v = 0 \Rightarrow v = 2iu, u \neq 0 \Rightarrow \vec{u}_1 = \begin{bmatrix} u \\ 2iu \end{bmatrix}$$

Let $u = 1$ for convenience and find $\vec{u}_1 = [1, 2i]^T$. We find the **fundamental complex solution** $\vec{x}$:

$$\vec{x} = e^{2it} \begin{bmatrix} 1 \\ 2i \end{bmatrix} = (\cos(2t) + i \sin(2t)) \begin{bmatrix} 1 \\ 2i \end{bmatrix} = \begin{bmatrix} \cos(2t) + i \sin(2t) \\ 2i \cos(2t) - 2 \sin(2t) \end{bmatrix}$$

Note: if $\vec{x} = \text{Re}(\vec{x}) + i\text{Im}(\vec{x})$ then it follows that the real and imaginary parts of the complex solution are themselves real solutions. Why? Because differentiation with respect to t is defined such that:

$$\frac{d\vec{x}}{dt} = \frac{d\text{Re}(\vec{x})}{dt} + i \frac{d\text{Im}(\vec{x})}{dt}$$

and $A\vec{x} = A[\text{Re}(\vec{x}) + i\text{Im}(\vec{x})] = A\text{Re}(\vec{x}) + iA\text{Im}(\vec{x})$. However, we know $d\vec{x}/dt = A\vec{x}$ hence we find, equating real parts and imaginary parts separately that:

$$\frac{d\text{Re}(\vec{x})}{dt} = A\text{Re}(\vec{x}) \quad \& \quad \frac{d\text{Im}(\vec{x})}{dt} = A\text{Im}(\vec{x})$$

Hence $\vec{x}_1 = \text{Re}(\vec{x})$ and $\vec{x}_2 = \text{Im}(\vec{x})$ give a solution set for the given system. In particular we find the fundamental solution set

$$\vec{x}_1(t) = \begin{bmatrix} \cos(2t) \\ -2 \sin(2t) \end{bmatrix} \quad \& \quad \vec{x}_2(t) = \begin{bmatrix} \sin(2t) \\ 2 \cos(2t) \end{bmatrix}.$$

We can assemble the general solution as a linear combination of the fundamental solutions

$\vec{x}(t) = c_1 \begin{bmatrix} \cos(2t) \\ -2 \sin(2t) \end{bmatrix} + c_2 \begin{bmatrix} \sin(2t) \\ 2 \cos(2t) \end{bmatrix}$. Thus the system $x' = -y$ and $y' = 4x$ has **scalar** solutions

$x(t) = c_1 \cos(2t) + c_2 \sin(2t)$ and $y(t) = -2c_1 \sin(2t) + 2c_2 \cos(2t)$. Finally, a fundamental matrix for this problem is given by

$$X = [\vec{x}_1 | \vec{x}_2] = \begin{bmatrix} \cos(2t) & \sin(2t) \\ -2 \sin(2t) & 2 \cos(2t) \end{bmatrix}.$$

Example 4.10.23. Problem: find the fundamental solutions of the system $x' = 2x - y$ and $y' = 9x + 2y$

Solution: we seek to solve $\frac{d\vec{x}}{dt} = A\vec{x}$ where $A = \begin{bmatrix} 2 & -1 \\ 9 & 2 \end{bmatrix}$. Consider the characteristic equation:

$$\det(A - \lambda I) = \det \begin{bmatrix} 2 - \lambda & -1 \\ 9 & 2 - \lambda \end{bmatrix} = (\lambda - 2)^2 + 9 = 0.$$

Thus $\lambda = 2 \pm 3i$. Consider $\lambda = 2 + 3i$, we seek the e -vector subject to $(A - (2 + 3i)I)\vec{u} = 0$. Solve:

$$\begin{bmatrix} -3i & -1 \\ 9 & -3i \end{bmatrix} \begin{bmatrix} u \\ v \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \end{bmatrix} \Rightarrow -3iu - v = 0 \Rightarrow v = -3iu, u \neq 0 \Rightarrow \vec{u}_1 = \begin{bmatrix} u \\ -3iu \end{bmatrix}$$

We choose $u = 1$ for convenience and thus find the fundamental complex solution

$$\vec{x}(t) = e^{(2+3i)t} \begin{bmatrix} 1 \\ -3i \end{bmatrix} = e^{2t}(\cos(3t) + i \sin(3t)) \begin{bmatrix} 1 \\ -3i \end{bmatrix} = e^{2t} \begin{bmatrix} \cos(3t) + i \sin(3t) \\ -3i \cos(3t) + 3 \sin(3t) \end{bmatrix}$$

Therefore, using the discussion of the last example, we find fundamental real solutions of the system by selecting real and imaginary parts of the complex solution above:

$$\vec{x}_1(t) = \begin{bmatrix} e^{2t} \cos(3t) \\ 3e^{2t} \sin(3t) \end{bmatrix} \quad \& \quad \vec{x}_2(t) = \begin{bmatrix} e^{2t} \sin(3t) \\ -3e^{2t} \cos(3t) \end{bmatrix}.$$

We can assemble the general solution as a linear combination of the fundamental solutions

$\vec{x}(t) = c_1 \begin{bmatrix} e^{2t} \cos(3t) \\ 3e^{2t} \sin(3t) \end{bmatrix} + c_2 \begin{bmatrix} e^{2t} \sin(3t) \\ -3e^{2t} \cos(3t) \end{bmatrix}$. Thus the system $x' = 2x - y$ and $y' = 9x + 2y$ has **scalar** solutions $x(t) = c_1 e^{2t} \cos(3t) + c_2 e^{2t} \sin(3t)$ and $y(t) = 3c_1 e^{2t} \sin(3t) - 3c_2 e^{2t} \cos(3t)$. Finally, a fundamental matrix for this problem is given by

$$X = [\vec{x}_1 | \vec{x}_2] = \begin{bmatrix} e^{2t} \cos(3t) & e^{2t} \sin(3t) \\ 3e^{2t} \sin(3t) & -3e^{2t} \cos(3t) \end{bmatrix}.$$

Example 4.10.24. Problem: we seek to solve $\frac{d\vec{x}}{dt} = A\vec{x}$ where $A = \begin{bmatrix} 2 & 0 & 0 \\ -1 & -4 & -1 \\ 0 & 5 & 2 \end{bmatrix}$.

Solution: Begin by solving the characteristic equation:

$$\begin{aligned} 0 = \det(A - \lambda I) &= \det \begin{bmatrix} 2 - \lambda & 0 & 0 \\ -1 & -4 - \lambda & -1 \\ 0 & 5 & 2 - \lambda \end{bmatrix} \\ &= (2 - \lambda)[(\lambda - 2)(\lambda + 4) + 5] \\ &= (2 - \lambda)(\lambda - 1)(\lambda + 3). \end{aligned}$$

Thus $\lambda_1 = 1$, $\lambda_2 = 2$ and $\lambda_3 = -3$. We seek $\vec{u}_1 = [u, v, w]^T$ such that $(A - I)\vec{u}_1 = 0$:

$$\begin{bmatrix} 1 & 0 & 0 \\ -1 & -5 & -1 \\ 0 & 5 & 1 \end{bmatrix} \begin{bmatrix} u \\ v \\ w \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \\ 0 \end{bmatrix} \Rightarrow \begin{matrix} u = 0 \\ 5v + w = 0 \end{matrix} \Rightarrow \begin{matrix} u = 0 \\ w = -5v \end{matrix} \Rightarrow \vec{u}_1 = v \begin{bmatrix} 0 \\ 1 \\ -5 \end{bmatrix}.$$

Choose $v = 1$ for convenience and find $\vec{u}_1 = [0, 1, -5]^T$. Next, seek $\vec{u}_2 = [u, v, w]^T$ such that $(A - 2I)\vec{u}_2 = 0$:

$$\begin{bmatrix} 0 & 0 & 0 \\ -1 & -6 & -1 \\ 0 & 5 & 0 \end{bmatrix} \begin{bmatrix} u \\ v \\ w \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \\ 0 \end{bmatrix} \Rightarrow \begin{array}{l} -u - 6v - w = 0 \\ v = 0 \end{array} \Rightarrow \begin{array}{l} v = 0 \\ w = -u \end{array} \Rightarrow \vec{u}_2 = u \begin{bmatrix} 1 \\ 0 \\ -1 \end{bmatrix}.$$

Choose $u = 1$ for convenience and find $\vec{u}_2 = [1, 0, -1]^T$. Last, seek $\vec{u}_3 = [u, v, w]^T$ such that $(A + 3I)\vec{u}_3 = 0$:

$$\begin{bmatrix} 5 & 0 & 0 \\ -1 & -1 & -1 \\ 0 & 5 & 5 \end{bmatrix} \begin{bmatrix} u \\ v \\ w \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \\ 0 \end{bmatrix} \Rightarrow \begin{array}{l} 5u = 0 \\ 5v + 5w = 0 \end{array} \Rightarrow \begin{array}{l} u = 0 \\ w = -v \end{array} \Rightarrow \vec{u}_3 = v \begin{bmatrix} 0 \\ 1 \\ -1 \end{bmatrix}.$$

Choose $v = 1$ for convenience and find $\vec{u}_3 = [0, 1, -1]^T$. The general solution follows:

$$\vec{x}(t) = c_1 e^t \begin{bmatrix} 0 \\ 1 \\ -5 \end{bmatrix} + c_2 e^{2t} \begin{bmatrix} 1 \\ 0 \\ -1 \end{bmatrix} + c_3 e^{-3t} \begin{bmatrix} 0 \\ 1 \\ -1 \end{bmatrix}.$$

The fundamental solutions and the fundamental matrix for the system above are given as follows:

$$\vec{x}_1(t) = \begin{bmatrix} 0 \\ e^t \\ -5e^t \end{bmatrix}, \quad \vec{x}_2(t) = \begin{bmatrix} e^{2t} \\ 0 \\ -e^{2t} \end{bmatrix}, \quad \vec{x}_3(t) = \begin{bmatrix} 0 \\ e^{-3t} \\ -e^{-3t} \end{bmatrix}, \quad X(t) = \begin{bmatrix} 0 & e^{2t} & 0 \\ e^t & 0 & e^{-3t} \\ -5e^t & -e^{2t} & -e^{-3t} \end{bmatrix}.$$

Example 4.10.25. we seek to solve $\frac{d\vec{x}}{dt} = A\vec{x}$ where $A = \begin{bmatrix} 2 & 0 & 0 \\ 0 & 2 & 0 \\ 1 & 0 & 3 \end{bmatrix}$.

Solution: Begin by solving the characteristic equation:

$$\det(A - \lambda I) = \det \begin{bmatrix} 2 - \lambda & 0 & 0 \\ 0 & 2 - \lambda & 0 \\ 1 & 0 & 3 - \lambda \end{bmatrix} = (2 - \lambda)(2 - \lambda)(3 - \lambda) = 0.$$

Thus $\lambda_1 = 2$, $\lambda_2 = 2$ and $\lambda_3 = 3$. We seek $\vec{u}_1 = [u, v, w]^T$ such that $(A - 2I)\vec{u}_1 = 0$:

$$\begin{bmatrix} 0 & 0 & 0 \\ 0 & 0 & 0 \\ 1 & 0 & 1 \end{bmatrix} \begin{bmatrix} u \\ v \\ w \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \\ 0 \end{bmatrix} \Rightarrow \begin{array}{l} u + w = 0 \\ v \text{ free} \end{array} \Rightarrow \begin{array}{l} v \text{ free} \\ w = -u \end{array} \Rightarrow \vec{u}_1 = \begin{bmatrix} u \\ v \\ -u \end{bmatrix}.$$

There are two free variables in the solution above and it follows we find two e -vectors. A convenient choice is $u = 1$ and $v = 0$ or $u = 0$ and $v = 1$; $\vec{u}_1 = [1, 0, -1]^T$ and $\vec{u}_2 = [0, 1, 0]^T$. Next, seek $\vec{u}_3 = [u, v, w]^T$ such that $(A - 3I)\vec{u}_3 = 0$

$$\begin{bmatrix} -1 & 0 & 0 \\ 0 & -1 & 0 \\ 1 & 0 & 0 \end{bmatrix} \begin{bmatrix} u \\ v \\ w \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \\ 0 \end{bmatrix} \Rightarrow \begin{array}{l} u = 0 \\ v = 0 \\ w \text{ free} \end{array} \Rightarrow \vec{u}_3 = w \begin{bmatrix} 0 \\ 0 \\ 1 \end{bmatrix}.$$

Choose $w = 1$ for convenience to find $\vec{u}_3 = [0, 0, 1]^T$. The general solution follows:

$$\vec{x}(t) = c_1 e^{2t} \begin{bmatrix} 1 \\ 0 \\ -1 \end{bmatrix} + c_2 e^{2t} \begin{bmatrix} 0 \\ 1 \\ 0 \end{bmatrix} + c_3 e^{-3t} \begin{bmatrix} 0 \\ 0 \\ 1 \end{bmatrix}.$$

The fundamental solutions and the fundamental matrix for the system above are given as follows:

$$\vec{x}_1(t) = \begin{bmatrix} e^{2t} \\ 0 \\ -e^{2t} \end{bmatrix}, \quad \vec{x}_2(t) = \begin{bmatrix} 0 \\ e^{2t} \\ 0 \end{bmatrix}, \quad \vec{x}_3(t) = \begin{bmatrix} 0 \\ 0 \\ e^{-3t} \end{bmatrix}, \quad X(t) = \begin{bmatrix} e^{2t} & 0 & 0 \\ 0 & e^{2t} & 0 \\ -e^{2t} & 0 & e^{-3t} \end{bmatrix}.$$

Example 4.10.26. we seek to solve $\frac{d\vec{x}}{dt} = A\vec{x}$ where $A = \begin{bmatrix} 2 & 1 & 0 \\ 0 & 2 & 0 \\ 1 & -1 & 3 \end{bmatrix}$.

Solution: Begin by solving the characteristic equation:

$$\det \begin{bmatrix} 2-\lambda & 1 & 0 \\ 0 & 2-\lambda & 0 \\ 1 & -1 & 3-\lambda \end{bmatrix} = (2-\lambda)^2(3-\lambda) = 0.$$

Thus $\lambda_1 = 2$, $\lambda_2 = 2$ and $\lambda_3 = 3$. We seek $\vec{u}_1 = [u, v, w]^T$ such that $(A - 2I)\vec{u}_1 = 0$:

$$\begin{bmatrix} 0 & 1 & 0 \\ 0 & 0 & 0 \\ 1 & -1 & 1 \end{bmatrix} \begin{bmatrix} u \\ v \\ w \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \\ 0 \end{bmatrix} \Rightarrow \begin{matrix} v = 0 \\ u + w = 0 \end{matrix} \Rightarrow \begin{matrix} v = 0 \\ w = -u \end{matrix} \Rightarrow \vec{u}_1 = \begin{bmatrix} u \\ 0 \\ -u \end{bmatrix}.$$

Choose $u = 1$ to select $\vec{u}_1 = [1, 0, -1]^T$. Next find $\vec{u}_2$ such that $(A - 3I)\vec{u}_2 = 0$

$$\begin{bmatrix} -1 & 1 & 0 \\ 0 & -1 & 0 \\ 1 & -1 & 0 \end{bmatrix} \begin{bmatrix} u \\ v \\ w \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \\ 0 \end{bmatrix} \Rightarrow \begin{matrix} -u + v = 0 \\ -v = 0 \\ w \text{ free} \end{matrix} \Rightarrow \begin{matrix} u = 0 \\ v = 0 \\ w \text{ free} \end{matrix} \Rightarrow \vec{u}_2 = \begin{bmatrix} 0 \\ 0 \\ w \end{bmatrix}.$$

Choose $w = 1$ to find $\vec{u}_2 = [0, 0, 1]^T$. We find two fundamental solutions from the e -vector method:

$$\vec{x}_1(t) = e^{2t} \begin{bmatrix} 1 \\ 0 \\ -1 \end{bmatrix} \quad \& \quad \vec{x}_2(t) = e^{3t} \begin{bmatrix} 0 \\ 0 \\ 1 \end{bmatrix}.$$

We cannot solve the system at this juncture since we are missing the third fundamental solution $\vec{x}_3$. In the next section we will find the missing solution via the generalized e -vector/ matrix exponential method.

Example 4.10.27. we seek to solve $\frac{d\vec{x}}{dt} = A\vec{x}$ where $A = \begin{bmatrix} 7 & 0 & 0 \\ 0 & 7 & 0 \\ 0 & 0 & 7 \end{bmatrix}$.

Solution: Begin by solving the characteristic equation:

$$\det \begin{bmatrix} 7-\lambda & 0 & 0 \\ 0 & 7-\lambda & 0 \\ 0 & 0 & 7-\lambda \end{bmatrix} = (7-\lambda)^3 = 0.$$

Thus $\lambda_1 = 7$, $\lambda_2 = 7$ and $\lambda_3 = 7$. The e -vector equation in this case is easy to solve; since $A - 7I = 7I - 7I = 0$ it follows that $(A - 7I)\vec{u} = 0$ for all $\vec{u} \in \mathbb{R}^3$. Therefore, any nontrivial vector is an eigenvector with e -value 7. A natural choice is $\vec{u}_1 = [1, 0, 0]^T$, $\vec{u}_2 = [0, 1, 0]^T$ and $\vec{u}_3 = [0, 0, 1]^T$. Thus,

$$\boxed{\vec{x}(t) = c_1 e^{7t} \begin{bmatrix} 1 \\ 0 \\ 0 \end{bmatrix} + c_2 e^{7t} \begin{bmatrix} 0 \\ 1 \\ 0 \end{bmatrix} + c_3 e^{7t} \begin{bmatrix} 0 \\ 0 \\ 1 \end{bmatrix} = e^{7t} \begin{bmatrix} c_1 \\ c_2 \\ c_3 \end{bmatrix}}.$$

Example 4.10.28. we seek to solve $\frac{d\vec{x}}{dt} = A\vec{x}$ where $A = \begin{bmatrix} -2 & 0 & 0 \\ 4 & -2 & 0 \\ 1 & 0 & -2 \end{bmatrix}$.

Solution: Begin by solving the characteristic equation:

$$\det(A - \lambda I) = \det \begin{bmatrix} -2 - \lambda & 0 & 0 \\ 4 & -2 - \lambda & 0 \\ 1 & 0 & -2 - \lambda \end{bmatrix} = -(\lambda + 2)^3 = 0.$$

Thus $\lambda_1 = -2$, $\lambda_2 = -2$ and $\lambda_3 = -2$. We seek $\vec{u}_1 = [u, v, w]^T$ such that $(A + 2I)\vec{u}_1 = 0$:

$$\begin{bmatrix} 0 & 0 & 0 \\ 4 & 0 & 0 \\ 1 & 0 & 0 \end{bmatrix} \begin{bmatrix} u \\ v \\ w \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \\ 0 \end{bmatrix} \Rightarrow \begin{array}{l} u = 0 \\ v \text{ free} \\ w \text{ free} \end{array} \Rightarrow \vec{u}_1 = \begin{bmatrix} 0 \\ v \\ w \end{bmatrix}.$$

Choose $v = 1, w = 0$ to select $\vec{u}_1 = [0, 1, 0]^T$ and $v = 0, w = 1$ to select $\vec{u}_2 = [0, 0, 1]^T$. Thus we find fundamental solutions:

$$\vec{x}_1(t) = e^{-2t} \begin{bmatrix} 0 \\ 1 \\ 0 \end{bmatrix} \quad \& \quad \vec{x}_2(t) = e^{-2t} \begin{bmatrix} 0 \\ 0 \\ 1 \end{bmatrix}.$$

We cannot solve the system at this juncture since we are missing the third fundamental solution $\vec{x}_3$. In the next section we will find the missing solution via the generalized e -vector/ matrix exponential method.

Example 4.10.29. we seek to solve $\frac{d\vec{x}}{dt} = A\vec{x}$ where $A = \begin{bmatrix} 2 & 1 & -1 \\ -3 & -1 & 1 \\ 9 & 3 & -4 \end{bmatrix}$.

Solution: Begin by solving the characteristic equation:

$$\begin{aligned} \det(A - \lambda I) &= \det \begin{bmatrix} 2 - \lambda & 1 & -1 \\ -3 & -1 - \lambda & 1 \\ 9 & 3 & -4 - \lambda \end{bmatrix} \\ &= (2 - \lambda)[(\lambda + 1)(\lambda + 4) - 3] - [3(\lambda + 4) - 9] - [-9 + 9(\lambda + 1)] \\ &= (2 - \lambda)[\lambda^2 + 5\lambda + 1] - 3\lambda - 3 - 9\lambda \\ &= -\lambda^3 - 5\lambda^2 - \lambda + 2\lambda^2 + 10\lambda + 2 - 12\lambda - 3 \\ &= -\lambda^3 - 3\lambda^2 - 3\lambda - 1 \\ &= -(\lambda + 1)^3 \end{aligned}$$

Thus $\lambda_1 = -1$, $\lambda_2 = -1$ and $\lambda_3 = -1$. We seek $\vec{u}_1 = [u, v, w]^T$ such that $(A + I)\vec{u}_1 = 0$:

$$\begin{bmatrix} 3 & 1 & -1 \\ -3 & 0 & 1 \\ 9 & 3 & -3 \end{bmatrix} \begin{bmatrix} u \\ v \\ w \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \\ 0 \end{bmatrix} \Rightarrow \begin{array}{l} 3u + v - w = 0 \\ -3u + w = 0 \end{array} \Rightarrow \begin{array}{l} w = 3u \\ v = w - 3u = 0 \end{array} \Rightarrow \vec{u}_1 = \begin{bmatrix} u \\ 0 \\ 3u \end{bmatrix}.$$

Choose $u = 1$ to select $\vec{u}_1 = [1, 0, 3]^T$. We find just one fundamental solution: $\vec{x}_1 = e^{-t}[1, 0, 3]^T$. We cannot solve the problem in it's entirety with our current methods. In the section that follows we find the missing pair of solutions.

Example 4.10.30. we seek to solve $\frac{d\vec{x}}{dt} = A\vec{x}$ where $A = \begin{bmatrix} 0 & 1 & 0 \\ 0 & 0 & 1 \\ 1 & -1 & 1 \end{bmatrix}$.

Solution: Begin by solving the characteristic equation:

$$\begin{aligned} \det(A - \lambda I) &= \det \begin{bmatrix} -\lambda & 1 & 0 \\ 0 & -\lambda & 1 \\ 1 & -1 & 1 - \lambda \end{bmatrix} \\ &= -\lambda[\lambda(\lambda - 1) + 1] + 1 \\ &= -\lambda^3 + \lambda^2 - \lambda + 1 \\ &= -\lambda^2(\lambda - 1) - (\lambda - 1) \\ &= (1 - \lambda)(\lambda^2 + 1) \end{aligned}$$

Thus $\lambda_1 = 1$, $\lambda_2 = i$ and $\lambda_3 = -i$. We seek $\vec{u}_1 = [u, v, w]^T$ such that $(A - I)\vec{u}_1 = 0$:

$$\begin{bmatrix} -1 & 1 & 0 \\ 0 & -1 & 1 \\ 1 & -1 & 0 \end{bmatrix} \begin{bmatrix} u \\ v \\ w \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \\ 0 \end{bmatrix} \Rightarrow \begin{array}{l} -u + v = 0 \\ -v + w = 0 \end{array} \Rightarrow \begin{array}{l} v = u \\ w = v \end{array} \Rightarrow \vec{u}_1 = \begin{bmatrix} u \\ u \\ u \end{bmatrix}.$$

Choose $u = 1$ thus select $\vec{u}_1 = [1, 1, 1]^T$. Now seek $\vec{u}_2$ such that $(A - iI)\vec{u}_2 = 0$

$$\begin{bmatrix} -i & 1 & 0 \\ 0 & -i & 1 \\ 1 & -1 & 1 - i \end{bmatrix} \begin{bmatrix} u \\ v \\ w \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \\ 0 \end{bmatrix} \Rightarrow \begin{array}{l} v = iu \\ w = iv = i(iu) = -u \\ (i - 1)w = u - v \end{array} \Rightarrow \vec{u}_2 = \begin{bmatrix} u \\ iu \\ -u \end{bmatrix}.$$

Set $u = 1$ to select the following complex solution:

$$\vec{x}(t) = e^{it} \begin{bmatrix} 1 \\ i \\ -1 \end{bmatrix} = \begin{bmatrix} e^{it} \\ ie^{it} \\ -e^{it} \end{bmatrix} = \begin{bmatrix} \cos(t) + i\sin(t) \\ i\cos(t) - \sin(t) \\ -\cos(t) - i\sin(t) \end{bmatrix} = \begin{bmatrix} \cos(t) \\ -\sin(t) \\ -\cos(t) \end{bmatrix} + i \begin{bmatrix} \sin(t) \\ \cos(t) \\ -\sin(t) \end{bmatrix}.$$

We select the second and third solutions by taking the real and imaginary parts of the above complex solution; $\vec{x}_2(t) = \text{Re}(\vec{x}(t))$ and $\vec{x}_3(t) = \text{Im}(\vec{x}(t))$. The general solution follows:

$$\boxed{\vec{x}(t) = c_1 e^t \begin{bmatrix} 1 \\ 1 \\ 1 \end{bmatrix} + c_2 \begin{bmatrix} \cos(t) \\ -\sin(t) \\ -\cos(t) \end{bmatrix} + c_3 \begin{bmatrix} \sin(t) \\ \cos(t) \\ -\sin(t) \end{bmatrix}}.$$

The fundamental solution set and fundamental matrix of the example above are simply:

$$\vec{x}_1 = \begin{bmatrix} e^t \\ e^t \\ e^t \end{bmatrix}, \quad \vec{x}_2 = \begin{bmatrix} \cos(t) \\ -\sin(t) \\ -\cos(t) \end{bmatrix}, \quad \vec{x}_3 = \begin{bmatrix} \sin(t) \\ \cos(t) \\ -\sin(t) \end{bmatrix} \quad \& \quad X = \begin{bmatrix} e^t & \cos(t) & \sin(t) \\ e^t & -\sin(t) & \cos(t) \\ e^t & -\cos(t) & -\sin(t) \end{bmatrix}$$

4.10.4 solutions by matrix exponential

Recall the Maclaurin series for the exponential is given by:

$$e^t = \sum_{j=0}^{\infty} \frac{t^j}{j!} = 1 + t + \frac{1}{2}t^2 + \frac{1}{6}t^3 + \dots$$

This provided the inspiration for the definition given below¹⁴

Definition 4.10.31. *matrix exponential*

Suppose A is an $n \times n$ matrix then we define the **matrix exponential of A** by:

$$e^A = \sum_{j=0}^{\infty} \frac{A^j}{j!} = I + A + \frac{1}{2}A^2 + \frac{1}{6}A^3 + \dots$$

Suppose $A = 0$ is the zero matrix. Note that

$$e^0 = I + 0 + \frac{1}{2}0^2 + \dots = I.$$

Furthermore, as $(-A)^j = (-1)^j A^j$ it follows that $e^{-A} = I - A + \frac{1}{2}A^2 - \frac{1}{6}A^3 + \dots$. Hence,

$$\begin{aligned} e^A e^{-A} &= \left(I + A + \frac{1}{2}A^2 + \frac{1}{6}A^3 + \dots \right) \left(I - A + \frac{1}{2}A^2 - \frac{1}{6}A^3 + \dots \right) \\ &= I - A + \frac{1}{2}A^2 - \frac{1}{6}A^3 + \dots + A \left(I - A + \frac{1}{2}A^2 + \dots \right) + \frac{1}{2}A^2 \left(I - A + \dots \right) + \frac{1}{6}A^3 I + \dots \\ &= I + A - A + \frac{1}{2}A^2 - A^2 + \frac{1}{2}A^2 - \frac{1}{6}A^3 + \frac{1}{2}A^3 - \frac{1}{2}A^3 + \frac{1}{6}A^3 + \dots \\ &= I. \end{aligned}$$

I have only shown the result up to the third-order in A , but you can verify higher orders if you wish. We find an interesting result:

$$(e^A)^{-1} = e^{-A} \quad \Rightarrow \quad \det(e^A) \neq 0 \quad \Rightarrow \quad \text{columns of } A \text{ are LI.}$$

Noncommutativity of matrix multiplication spoils the usual law of exponents. Let's examine how this happens. Suppose A, B are square matrices. Calculate e^{A+B} to second order in A, B :

$$e^{A+B} = I + (A+B) + \frac{1}{2}(A+B)^2 + \dots = I + A + B + \frac{1}{2}(A^2 + AB + BA + B^2) + \dots$$

On the other hand, calculate the product $e^A e^B$ to second order in A, B ,

$$e^A e^B = \left(I + A + \frac{1}{2}A^2 + \dots \right) \left(I + B + \frac{1}{2}B^2 + \dots \right) = I + A + B + \frac{1}{2}(A^2 + 2AB + B^2) + \dots$$

¹⁴the concept of an exponential actually extends in much more generality than this, we could derive this from more basic and general principles, but that has little to do with this course so we behave. In addition, the reason the series of matrices below converges is not immediately obvious, see my linear notes for a sketch of the analysis needed here

We find that, to second order, $e^A e^B - e^{A+B} = \frac{1}{2}(AB - BA)$. Define the **commutator** $[A, B] = AB - BA$ and note (after a short calculation)

$$\boxed{e^A e^B = e^{A+B+\frac{1}{2}[A,B]+\dots}}$$

When A, B are **commuting** matrices the commutator $[A, B] = AB - BA = AB - AB = 0$ hence the usual algebra $e^A e^B = e^{A+B}$ applies. It turns out that the higher-order terms in the boxed formula above can be written as nested-commutators of A and B . This formula is known as the Baker-Campbell-Hausdorff, it is the essential calculation in the theory of matrix Lie groups (which is the math used to formulate important symmetry aspects of modern physics).

Let me pause¹⁵ to give a better proof that $AB = BA$ implies $e^A e^B = e^{A+B}$. The heart of the argument follows from the fact the binomial theorem holds for $(A + B)^k$ in this context. I seek to prove by mathematical induction on k that $(A + B)^k = \sum_{n=0}^k \binom{k}{n} A^{k-n} B^n$. Note $k = 1$ is clearly true as $\binom{1}{0} = \binom{1}{1} = 1$ and $(A + B)^1 = A + B$. Assume inductively the binomial theorem holds for k and seek to prove $k + 1$ true:

$$\begin{aligned} (A + B)^{k+1} &= (A + B)^k (A + B) \\ &= \left(\sum_{n=0}^k \binom{k}{n} A^{k-n} B^n \right) (A + B) \quad : \text{by induction hypothesis} \\ &= \sum_{n=0}^k \binom{k}{n} A^{k-n} B^n A + \sum_{n=0}^k \binom{k}{n} A^{k-n} B^n B \\ &= \sum_{n=0}^k \binom{k}{n} A^{k-n} AB^n + \sum_{n=0}^k \binom{k}{n} A^{k-n} B^{n+1} \quad : AB = BA \text{ implies } B^n A = AB^n \\ &= \sum_{n=0}^k \binom{k}{n} A^{k+1-n} B^n + \sum_{n=0}^k \binom{k}{n} A^{k-n} B^{n+1} \end{aligned}$$

Continuing,

$$\begin{aligned} (A + B)^{k+1} &= A^{k+1} + \sum_{n=1}^k \binom{k}{n} A^{k+1-n} B^n + \sum_{n=0}^{k-1} \binom{k}{n} A^{k-n} B^{n+1} + B^{k+1} \\ &= A^{k+1} + \sum_{n=1}^k \binom{k}{n} A^{k+1-n} B^n + \sum_{n=1}^k \binom{k}{n-1} A^{k+1-n} B^n + B^{k+1} \\ &= A^{k+1} + \sum_{n=1}^k \left[\binom{k}{n} + \binom{k}{n-1} \right] A^{k+1-n} B^n + B^{k+1} \\ &= A^{k+1} + \sum_{n=1}^k \binom{k+1}{n} A^{k+1-n} B^n + B^{k+1} \quad : \text{by Pascal's Triangle} \\ &= \sum_{n=0}^{k+1} \binom{k+1}{n} A^{k+1-n} B^n \end{aligned}$$

¹⁵you may skip ahead if you are not interested in how to make arguments precise, in fact, even this argument has gaps, but I include it to give the reader some idea about what is missing when we resort to $+\dots$ -style induction

Which completes the induction step and we find by mathematical induction the binomial theorem for commuting matrices holds for all $k \in \mathbb{N}$. Consider the matrix exponential formula in light of the binomial theorem, also recall $\binom{k+1}{n} = \frac{k!}{n!(k-n)!}$,

$$\begin{aligned} e^{A+B} &= \sum_{k=0}^{\infty} \frac{1}{k!} (A+B)^k \\ &= \sum_{k=0}^{\infty} \sum_{n=0}^k \frac{1}{k!} \frac{k!}{n!(k-n)!} A^{k-n} B^n \\ &= \sum_{k=0}^{\infty} \sum_{n=0}^k \frac{1}{n!} \frac{1}{(k-n)!} A^{k-n} B^n \\ &= \sum_{k=0}^{\infty} \sum_{n=0}^k \frac{1}{n!} \frac{1}{(k-n)!} A^{k-n} B^n \end{aligned}$$

On the other hand, if we compute the product of e^A with e^B we find:

$$e^A e^B = \sum_{j=0}^{\infty} \frac{1}{j!} A^j \sum_{n=0}^{\infty} \frac{1}{n!} B^n = \sum_{j=0}^{\infty} \sum_{n=0}^{\infty} \frac{1}{n!} \frac{1}{j!} A^j B^n$$

It follows¹⁶ that $e^A e^B = e^{A+B}$. We use this result implicitly in much of what follows in this section.

Suppose A is a constant $n \times n$ matrix. Calculate¹⁷

$$\begin{aligned} \frac{d}{dt} [\exp(tA)] &= \frac{d}{dt} \left[\sum_{k=0}^{\infty} \frac{1}{k!} t^k A^k \right] && \text{defn. of matrix exponential} \\ &= \sum_{k=0}^{\infty} \frac{d}{dt} \left[\frac{1}{k!} t^k A^k \right] && \text{since matrix exp. uniformly conv.} \\ &= \sum_{k=0}^{\infty} \frac{k}{k!} t^{k-1} A^k && A^k \text{ constant and } \frac{d}{dt}(t^k) = kt^{k-1} \\ &= A \sum_{k=1}^{\infty} \frac{1}{(k-1)!} t^{k-1} A^{k-1} && \text{since } k! = k(k-1)! \text{ and } A^k = AA^{k-1}. \\ &= A \exp(tA) && \text{defn. of matrix exponential.} \end{aligned}$$

I suspect the following argument is easier to follow:

$$\begin{aligned} \frac{d}{dt} (\exp(tA)) &= \frac{d}{dt} (I + tA + \frac{1}{2}t^2 A^2 + \frac{1}{3!}t^3 A^3 + \cdots) \\ &= \frac{d}{dt}(I) + \frac{d}{dt}(tA) + \frac{1}{2} \frac{d}{dt}(t^2 A^2) + \frac{1}{3 \cdot 2} \frac{d}{dt}(t^3 A^3) + \cdots \\ &= A + tA^2 + \frac{1}{2}t^2 A^3 + \cdots \\ &= A(I + tA + \frac{1}{2}t^2 A^2 + \cdots) \\ &= A \exp(tA). \end{aligned}$$

□

¹⁶after some analytical arguments beyond this course; what is missing is an explicit examination of the infinite limits at play here, the doubly infinite limits seem to reach the same terms but the structure of the sums differs

¹⁷the term-by-term differentiation theorem for power series extends to a matrix power series, the proof of this involves real analysis

Whichever notation you prefer, the calculation above completes the proof of the following central theorem for this section:

Theorem 4.10.32.

Suppose $A \in \mathbb{R}^{n \times n}$. The matrix exponential e^{tA} gives a fundamental matrix for $\frac{d\vec{x}}{dt} = A\vec{x}$.

Proof: we have already shown that (1.) e^{tA} is a solution matrix ($\frac{d}{dt}[e^{tA}] = Ae^{tA}$) and (2.) $(e^{tA})^{-1} = e^{-tA}$ thus the columns of e^{tA} are LI. $\square$

It follows that the general solution of $\frac{d\vec{x}}{dt} = A\vec{x}$ is simply $\vec{x}(t) = e^{tA}\vec{c}$ where $\vec{c} = [c_1, c_2, \dots, c_n]^T$ determines the initial conditions of the solution. In theory this is a great formula, we've solved most of the problems we set-out to solve. However, more careful examination reveals this result is much like the result from calculus; any continuous function is integrable. Ok, so f continuous on an interval I implies F exists on I and $F' = f$, but... how do you actually calculate the antiderivative F ? It's possible in principle, but in practice the computation may fall outside the computation scope of the techniques covered in calculus¹⁸.

Example 4.10.33. Suppose $x' = x, y' = 2y, z' = 3z$ then in matrix form we have:

$$\begin{bmatrix} x \\ y \\ z \end{bmatrix}' = \begin{bmatrix} 1 & 0 & 0 \\ 0 & 2 & 0 \\ 0 & 0 & 3 \end{bmatrix} \begin{bmatrix} x \\ y \\ z \end{bmatrix}$$

The coefficient matrix is diagonal which makes the k -th power particularly easy to calculate,

$$\begin{aligned} A^k &= \begin{bmatrix} 1 & 0 & 0 \\ 0 & 2 & 0 \\ 0 & 0 & 3 \end{bmatrix}^k = \begin{bmatrix} 1 & 0 & 0 \\ 0 & 2^k & 0 \\ 0 & 0 & 3^k \end{bmatrix} \\ \Rightarrow \exp(tA) &= \sum_{k=0}^{\infty} \frac{t^k}{k!} \begin{bmatrix} 1 & 0 & 0 \\ 0 & 2^k & 0 \\ 0 & 0 & 3^k \end{bmatrix} = \begin{bmatrix} \sum_{k=0}^{\infty} \frac{t^k}{k!} 1^k & 0 & 0 \\ 0 & \sum_{k=0}^{\infty} \frac{t^k}{k!} 2^k & 0 \\ 0 & 0 & \sum_{k=0}^{\infty} \frac{t^k}{k!} 3^k \end{bmatrix} \\ \Rightarrow \exp(tA) &= \begin{bmatrix} e^t & 0 & 0 \\ 0 & e^{2t} & 0 \\ 0 & 0 & e^{3t} \end{bmatrix} \end{aligned}$$

Thus we find three solutions to $x' = Ax$,

$$x_1(t) = \begin{bmatrix} e^t \\ 0 \\ 0 \end{bmatrix} \quad x_2(t) = \begin{bmatrix} 0 \\ e^{2t} \\ 0 \end{bmatrix} \quad x_3(t) = \begin{bmatrix} 0 \\ 0 \\ e^{3t} \end{bmatrix}$$

In turn these vector solutions amount to the solutions $x = e^t, y = 0, z = 0$ or $x = 0, y = e^{2t}, z = 0$ or $x = 0, y = 0, z = e^{3t}$. It is easy to check these solutions.

Of course the example above is very special. In order to unravel the mystery of just how to calculate the matrix exponential for less trivial matrices we return to the construction of the previous section.

¹⁸for example, $\int \frac{\sin(x)dx}{x}$ or $\int e^{-x^2}dx$ are known to be incalculable in terms of elementary functions

Let's see what happens when we calculate $e^{tA}\vec{u}$ for $\vec{u}$ and e-vector with e-value λ .

$$\begin{aligned}
e^{tA}\vec{u} &= e^{t(A-\lambda I+\lambda I)}\vec{u} && \text{: added zero anticipating use of } (A-\lambda I)\vec{u} = 0, \\
&= e^{t\lambda I+t(A-\lambda I)}\vec{u} \\
&= e^{t\lambda I}e^{t(A-\lambda I)}\vec{u} && \text{: noted that } t\lambda I \text{ commutes with } t(A-\lambda I), \\
&= e^{t\lambda I}e^{t(A-\lambda I)}\vec{u} && \text{: a short exercise shows } e^{t\lambda I} = e^{t\lambda}I. \\
&= e^{t\lambda}\left(I + t(A-\lambda I) + \frac{t^2}{2}(A-\lambda I)^2 + \cdots\right)\vec{u} \\
&= e^{t\lambda}\left(I\vec{u} + t(A-\lambda I)\vec{u} + \frac{t^2}{2}(A-\lambda I)^2\vec{u} + \cdots\right) \\
&= e^{t\lambda}\vec{u} && \text{: as it was given } (A-\lambda I)\vec{u} = 0 \text{ hence all but the first term vanishes.}
\end{aligned}$$

The fact that this is a solution of $\vec{x}' = A\vec{x}$ was already known to us, however, it is nice to see it arise from the matrix exponential. Moreover the calculation above reveals the central formula that guides the technique of this section. The **magic formula**. For any square matrix and possibly constant λ we find:

$$e^{tA} = e^{t\lambda}\left(I + t(A-\lambda I) + \frac{t^2}{2}(A-\lambda I)^2 + \cdots\right) = e^{t\lambda}\sum_{k=0}^{\infty} \frac{t^k}{k!}(A-\lambda I)^k.$$

When we choose λ as an e-value and multiply this formula by the corresponding e-vector then this infinite series truncates nicely to reveal $e^{\lambda t}\vec{u}$. It follows that we should define vectors which truncate the series at higher order, this is the natural next step:

Definition 4.10.34. *generalized eigenvectors and chains of generalized e-vectors*

Given an eigenvalue λ a nonzero vector $\vec{u}$ such that $(A-\lambda I)^p\vec{u} = 0$ and $(A-\lambda I)^{p-1}\vec{u} \neq 0$ is called an **generalized eigenvector of order p** with eigenvalue λ . If $\vec{u}_1, \vec{u}_2, \dots, \vec{u}_p$ are nonzero vectors such that $(A-\lambda I)\vec{u}_j = \vec{u}_{j-1}$ for $j = 2, 3, \dots, p$ and $\vec{u}_1$ is an e-vector with e-value λ then we say $\{\vec{u}_1, \vec{u}_2, \dots, \vec{u}_p\}$ forms a **chain of generalized e-vectors of length p** .

In the notation of the definition above, it is true that $\vec{u}_k$ is a generalized e-vector of order k with e-value λ . Let's examine $k = 2$,

$$(A-\lambda I)\vec{u}_2 = \vec{u}_1 \quad \Rightarrow \quad (A-\lambda I)^2\vec{u}_2 = (A-\lambda I)\vec{u}_1 = 0.$$

Then suppose inductively the claim is true for k which means $(A-\lambda I)^k\vec{u}_k = 0$, consider $k+1$

$$(A-\lambda I)\vec{u}_{k+1} = \vec{u}_k \quad \Rightarrow \quad (A-\lambda I)^{k+1}\vec{u}_{k+1} = (A-\lambda I)^k\vec{u}_k = 0.$$

Hence, in terms of the notation in the definition above, we have shown by mathematical induction that $\vec{u}_k$ is a generalized e-vector of order k with e-value λ .

I do not mean to claim this is true for all $k \in \mathbb{N}$. In practice for an $n \times n$ matrix we cannot find a chain longer than length n . However, up to that bound such chains are possible for an arbitrary matrix.

Example 4.10.35. The matrices below are in **Jordan form** which means the vectors $e_1 = [1, 0, 0, 0, 0]^T$ etc... $e_5 = [0, 0, 0, 0, 1]^T$ are (generalized)- e -vectors:

$$A = \begin{bmatrix} 2 & 1 & 0 & 0 & 0 \\ 0 & 2 & 1 & 0 & 0 \\ 0 & 0 & 2 & 0 & 0 \\ 0 & 0 & 0 & 3 & 1 \\ 0 & 0 & 0 & 0 & 3 \end{bmatrix} \quad \& \quad B = \begin{bmatrix} 4 & 1 & 0 & 0 & 0 \\ 0 & 4 & 0 & 0 & 0 \\ 0 & 0 & 5 & 0 & 0 \\ 0 & 0 & 0 & 6 & 0 \\ 0 & 0 & 0 & 0 & 6 \end{bmatrix}$$

You can easily calculate $(A - 2I)e_1 = 0$, $(A - 2I)e_2 = e_1$, $(A - 2I)e_3 = e_2$ or $(A - 3I)e_4 = 0$, $(A - 2I)e_5 = e_4$. On the other hand, $(B - 4I)e_1 = 0$, $(B - 4I)e_2 = e_1$ and $(A - 5I)e_3 = 0$ and $(A - 6I)e_4 = 0$, $(A - 6I)e_5 = 0$. The matrix B needs only one generalized e -vector whereas the matrix A has 3 generalized e -vectors.

Let's examine why chains are nice for the magic formula:

Example 4.10.36. Problem: Suppose A is a 3×3 matrix with a chain of generalized e -vector $\vec{u}_1, \vec{u}_2, \vec{u}_3$ with respect to e -value $\lambda = 2$. Solve $\frac{d\vec{x}}{dt} = A\vec{x}$ in view of these facts.

Solution: we are given $(A - 2I)\vec{u}_1 = 0$ and $(A - 2I)\vec{u}_2 = \vec{u}_1$ and $(A - 2I)\vec{u}_3 = \vec{u}_2$. It is easily shown that $(A - 2I)^2\vec{u}_2 = 0$ and $(A - 2I)^3\vec{u}_3 = 0$. It is also possible to prove $\{\vec{u}_1, \vec{u}_2, \vec{u}_3\}$ is a LI set. Apply the magic formula with $\lambda = 2$ to derive the following results:

1. $\vec{x}_1(t) = e^{tA}\vec{u}_1 = e^{2t}\vec{u}_1$ (we've already shown this in general earlier in this section)
2. $\vec{x}_2(t) = e^{tA}\vec{u}_2 = e^{2t}(I\vec{u}_2 + t(A - 2I)\vec{u}_2 + \frac{t^2}{2}(A - 2I)^2\vec{u}_2 + \dots) = e^{2t}(\vec{u}_2 + t\vec{u}_1)$.
3. note that $(A - 2I)^2\vec{u}_3 = (A - 2I)\vec{u}_2 = \vec{u}_1$ hence:

$$\vec{x}_3(t) = e^{tA}\vec{u}_3 = e^{2t}(I\vec{u}_3 + t(A - 2I)\vec{u}_3 + \frac{t^2}{2}(A - 2I)^2\vec{u}_3 + \dots) = e^{2t}(\vec{u}_3 + t\vec{u}_2 + \frac{t^2}{2}\vec{u}_1).$$

Therefore, $\boxed{\vec{x}(t) = c_1 e^{2t}\vec{u}_1 + c_2 e^{2t}(\vec{u}_2 + t\vec{u}_1) + c_3 e^{2t}(\vec{u}_3 + t\vec{u}_2 + \frac{t^2}{2}\vec{u}_1)}$ is the general solution.

Perhaps it is interesting to calculate $e^{tA}[\vec{u}_1|\vec{u}_2|\vec{u}_3]$ in view of the calculations in the example above. Observe:

$$e^{tA}[\vec{u}_1|\vec{u}_2|\vec{u}_3] = [e^{tA}\vec{u}_1|e^{tA}\vec{u}_2|e^{tA}\vec{u}_3] = e^{2t}\left[\vec{u}_1\left|\vec{u}_2 + t\vec{u}_1\right|\vec{u}_3 + t\vec{u}_2 + \frac{t^2}{2}\vec{u}_1\right].$$

I suppose we could say more about this formula, but let's get back on task: we seek to complete the solution of the unsolved problems of the previous section. It is our hope that we can find generalized e -vector solutions to complete the fundamental solution sets in Examples 4.10.20, 4.10.26, 4.10.28 and 4.10.29.

Example 4.10.37. Problem: (returning to Example 4.10.20) solve $\frac{d\vec{x}}{dt} = A\vec{x}$ where $A = \begin{bmatrix} 3 & 1 \\ -4 & -1 \end{bmatrix}$

Solution: we found $\lambda_1 = 1$ and $\lambda_2 = 1$ and a single e -vector $\vec{u}_1 = \begin{bmatrix} 1 \\ -2 \end{bmatrix}$. Now seek a generalized e -vector $\vec{u}_2 = [u, v]^T$ such that $(A - I)\vec{u}_2 = \vec{u}_1$,

$$\begin{bmatrix} 2 & 1 \\ -4 & -2 \end{bmatrix} \begin{bmatrix} u \\ v \end{bmatrix} = \begin{bmatrix} 1 \\ -2 \end{bmatrix} \Rightarrow 2u + v = 1 \Rightarrow v = 1 - 2u, u \neq 0 \Rightarrow \vec{u}_1 = \begin{bmatrix} u \\ 1 - 2u \end{bmatrix}$$

We choose $u = 0$ for convenience and thus find $\vec{u}_2 = [0, 1]^T$ hence the fundamental solution

$$\vec{x}_2(t) = e^{tA}\vec{u}_2 = e^t(I + t(A - I) + \cdots)\vec{u}_2 = e^t(\vec{u}_2 + t\vec{u}_1) = e^t \begin{bmatrix} t \\ 1 - 2t \end{bmatrix}.$$

Therefore, we find $\vec{x}(t) = c_1 e^t \begin{bmatrix} 1 \\ -2 \end{bmatrix} + c_2 e^t \begin{bmatrix} t \\ 1 - 2t \end{bmatrix}.$

Example 4.10.38. Problem: (returning to Example 4.10.26) solve $\frac{d\vec{x}}{dt} = A\vec{x}$ where

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

Solution: we found $\lambda_1 = 2$, $\lambda_2 = 2$ and $\lambda_3 = 3$ and we also found e-vector $\vec{u}_1 = [1, 0, -1]^T$ with e-value 2 and e-vector $\vec{u}_2 = [0, 0, 1]^T$. Seek $\vec{u}_3$ such that $(A - 2I)\vec{u}_3 = \vec{u}_1$ since we are missing a solution paired with $\lambda_2 = 2$.

$$\begin{bmatrix} 0 & 1 & 0 \\ 0 & 0 & 0 \\ 1 & -1 & 1 \end{bmatrix} \begin{bmatrix} u \\ v \\ w \end{bmatrix} = \begin{bmatrix} 1 \\ 0 \\ -1 \end{bmatrix} \Rightarrow \begin{matrix} v = 1 \\ u - 1 + w = -1 \end{matrix} \Rightarrow \begin{matrix} v = 1 \\ w = -u \end{matrix} \Rightarrow \vec{u}_1 = \begin{bmatrix} u \\ 1 \\ -u \end{bmatrix}.$$

Choose $u = 0$ to select $\vec{u}_1 = [0, 1, 0]^T$. It follows from the magic formula that $\vec{x}_3(t) = e^{tA}\vec{u}_3 = e^{2t}(\vec{u}_3 + t\vec{u}_1)$. Hence, the general solution is

$$\vec{x}(t) = c_1 e^{2t} \begin{bmatrix} 1 \\ 0 \\ -1 \end{bmatrix} + c_2 e^{3t} \begin{bmatrix} 0 \\ 0 \\ 1 \end{bmatrix} + c_3 e^{2t} \begin{bmatrix} t \\ 1 \\ -t \end{bmatrix}.$$

Once more we found a generalized e-vector of order two to complete the solution set and find $\vec{x}_3$ in the example above. You might notice that had we replaced the choice $u = 0$ in both of the last examples with some nonzero u then we would have added a copy of $\vec{x}_1$ to the generalized e-vector solution. This is permissible since the sum of solutions to the system $\vec{x}' = A\vec{x}$ is once more a solution. This freedom works hand-in-hand with the ambiguity of the generalized e-vector problem.

Example 4.10.39. Problem: (returning to Example 4.10.28) we seek to solve $\frac{d\vec{x}}{dt} = A\vec{x}$

where $A = \begin{bmatrix} -2 & 0 & 0 \\ 4 & -2 & 0 \\ 1 & 0 & -2 \end{bmatrix}.$

Solution: We already found $\lambda_1 = -2$, $\lambda_2 = -2$ and $\lambda_3 = -2$ and a pair of e-vectors $\vec{u}_1 = [0, 1, 0]^T$ and $v = 0, w = 1$ to select $\vec{u}_2 = [0, 0, 1]^T$. We face a dilemma, should we look for a chain that ends with $\vec{u}_1 = [0, 1, 0]^T$ or $\vec{u}_2 = [0, 0, 1]^T$? Generally it may not be possible to do either. Thus, we set aside the chain condition and instead look for directly for solutions of $(A + 2I)^2\vec{u}_3 = 0$.

$$(A + 2I)^2 = \begin{bmatrix} 0 & 0 & 0 \\ 4 & 0 & 0 \\ 1 & 0 & 0 \end{bmatrix} \begin{bmatrix} 0 & 0 & 0 \\ 4 & 0 & 0 \\ 1 & 0 & 0 \end{bmatrix} = \begin{bmatrix} 0 & 0 & 0 \\ 0 & 0 & 0 \\ 0 & 0 & 0 \end{bmatrix}$$

Since we seek $\vec{u}_3$ which forms a LI set with $\vec{u}_1, \vec{u}_2$ it is natural to choose $\vec{u}_3 = [1, 0, 0]^T$. Calculate,

$$\begin{aligned}\vec{x}_3(t) &= e^{tA}\vec{u}_3 = e^{-2t}(I\vec{u}_3 + t(A + 2I)\vec{u}_3 + \frac{t^2}{2}(A + 2I)^2\vec{u}_3 + \cdots) \\ &= e^{-2t} \left[\begin{bmatrix} 1 \\ 0 \\ 0 \end{bmatrix} + t \begin{bmatrix} 0 & 0 & 0 \\ 4 & 0 & 0 \\ 1 & 0 & 0 \end{bmatrix} \begin{bmatrix} 1 \\ 0 \\ 0 \end{bmatrix} \right] \\ &= e^{-2t} \begin{bmatrix} 1 \\ 4t \\ t \end{bmatrix}\end{aligned}\tag{4.3}$$

Thus we find the general solution:

$$\vec{x}(t) = c_1 e^{-2t} \begin{bmatrix} 0 \\ 1 \\ 0 \end{bmatrix} + c_2 e^{-2t} \begin{bmatrix} 0 \\ 0 \\ 1 \end{bmatrix} + c_3 e^{-2t} \begin{bmatrix} 1 \\ 4t \\ t \end{bmatrix}.$$

Example 4.10.40. Problem: (returning to Example 4.10.29) we seek to solve $\frac{d\vec{x}}{dt} = A\vec{x}$

where $A = \begin{bmatrix} 2 & 1 & -1 \\ -3 & -1 & 1 \\ 9 & 3 & -4 \end{bmatrix}$.

Solution: we found $\lambda_1 = -1$, $\lambda_2 = -1$ and $\lambda_3 = -1$ and a single e-vector $\vec{u}_1 = [1, 0, 3]^T$. Seek $\vec{u}_2$ such that $(A + I)\vec{u}_2 = \vec{u}_1$,

$$\begin{bmatrix} 3 & 1 & -1 \\ -3 & 0 & 1 \\ 9 & 3 & -3 \end{bmatrix} \begin{bmatrix} u \\ v \\ w \end{bmatrix} = \begin{bmatrix} 1 \\ 0 \\ 3 \end{bmatrix} \Rightarrow \begin{aligned} 3u + v - w &= 1 \\ -3u + w &= 0 \end{aligned} \Rightarrow \begin{aligned} w &= 3u \\ v &= w - 3u + 1 \end{aligned} \Rightarrow \vec{u}_2 = \begin{bmatrix} 0 \\ 1 \\ 0 \end{bmatrix}.$$

where we set $u = 0$ for convenience. Continuing, we seek $\vec{u}_3$ where $(A + I)\vec{u}_3 = \vec{u}_2$,

$$\begin{bmatrix} 3 & 1 & -1 \\ -3 & 0 & 1 \\ 9 & 3 & -3 \end{bmatrix} \begin{bmatrix} u \\ v \\ w \end{bmatrix} = \begin{bmatrix} 0 \\ 1 \\ 0 \end{bmatrix} \Rightarrow \begin{aligned} 3u + v - w &= 0 \\ -3u + w &= 1 \end{aligned} \Rightarrow \begin{aligned} w &= 1 + 3u \\ v &= w - 3u \end{aligned} \Rightarrow \begin{aligned} w &= 1 + 3u \\ v &= 1 \end{aligned}$$

Choose $u = 0$ to select $\vec{u}_3 = [0, 1, 1]^T$. Given the algebra we've completed we know that

$$(A + I)\vec{u}_1 = (A + I)^2\vec{u}_2 = (A + I)^3\vec{u}_3 = 0, \quad (A + I)\vec{u}_2 = \vec{u}_1, \quad (A + I)\vec{u}_3 = \vec{u}_2, \quad (A + I)^2\vec{u}_3 = \vec{u}_1$$

These identities paired with the magic formula with $\lambda = -1$ yield:

$$e^{tA}\vec{u}_1 = e^{-t}\vec{u}_1 \quad \& \quad e^{tA}\vec{u}_2 = e^{-t}(\vec{u}_2 + t\vec{u}_1) \quad \& \quad e^{tA}\vec{u}_3 = e^{-t}(\vec{u}_3 + t\vec{u}_2 + \frac{t^2}{2}\vec{u}_1)$$

Therefore, we find general solution:

$$\vec{x}(t) = c_1 e^{-t} \begin{bmatrix} 1 \\ 0 \\ 3 \end{bmatrix} + c_2 e^{-t} \begin{bmatrix} t \\ 1 \\ 3t \end{bmatrix} + c_3 e^{-t} \begin{bmatrix} 0 \\ 1 + t + \frac{t^2}{2} \\ 1 + \frac{t^2}{2} \end{bmatrix}.$$

The method we've illustrated extends naturally to the case of repeated complex e-values where there are insufficient e-vectors to form the general solution.

Example 4.10.41. Problem: Suppose A is a 6×6 matrix with a chain of generalized e -vector $\vec{u}_1, \vec{u}_2, \vec{u}_3$ with respect to e -value $\lambda = 2 + i$. Solve $\frac{d\vec{x}}{dt} = A\vec{x}$ in view of these facts.

Solution: we are given $(A - (2 + i)I)\vec{u}_1 = 0$ and $(A - (2 + i)I)\vec{u}_2 = \vec{u}_1$ and $(A - (2 + i)I)\vec{u}_3 = \vec{u}_2$. It is easily shown that $(A - (2 + i)I)^2\vec{u}_2 = 0$ and $(A - (2 + i)I)^3\vec{u}_3 = 0$. It is also possible to prove $\{\vec{u}_1, \vec{u}_2, \vec{u}_3\}$ is a LI set. Apply the magic formula with $\lambda = (2 + i)$ to derive the following results:

1. $\vec{x}_1(t) = e^{tA}\vec{u}_1 = e^{(2+i)t}\vec{u}_1$ (we've already shown this in general earlier in this section)
2. $\vec{x}_2(t) = e^{tA}\vec{u}_2 = e^{(2+i)t}(I\vec{u}_2 + t(A - (2 + i)I)\vec{u}_2 + \frac{t^2}{2}(A - (2 + i)I)^2\vec{u}_2 + \dots) = e^{(2+i)t}(\vec{u}_2 + t\vec{u}_1)$.
3. note that $(A - (2 + i)I)^2\vec{u}_3 = (A - (2 + i)I)\vec{u}_2 = \vec{u}_1$ hence:

$$\vec{x}_3(t) = e^{tA}\vec{u}_3 = e^{(2+i)t}(I\vec{u}_3 + t(A - (2 + i)I)\vec{u}_3 + \frac{t^2}{2}(A - (2 + i)I)^2\vec{u}_3 + \dots) = e^{(2+i)t}(\vec{u}_3 + t\vec{u}_2 + \frac{t^2}{2}\vec{u}_1).$$

The solutions $\vec{x}_1(t)$, $\vec{x}_2(t)$ and $\vec{x}_3(t)$ are complex-valued solutions. To find the real solutions we select the real and imaginary parts to form the fundamental solution set

$$\{\operatorname{Re}(\vec{x}_1), \operatorname{Im}(\vec{x}_1), \operatorname{Re}(\vec{x}_2), \operatorname{Im}(\vec{x}_2), \operatorname{Re}(\vec{x}_3), \operatorname{Im}(\vec{x}_3)\}$$

I leave the explicit formulas to the reader, it is very similar to the case we treated in the last section for the complex e -vector problem.

Suppose A is idempotent or order k then $A^{k-1} \neq I$ and $A^k = I$. In this case the matrix exponential simplifies:

$$e^{tA} = I + tA + \frac{t^2}{2}A^2 + \dots + \frac{t^{k-1}}{(k-1)!}A^{k-1} + \left(\frac{t^k}{k!} + \frac{t^{k+1}}{(k+1)!} + \dots\right)I$$

However, $\frac{t^k}{k!} + \frac{t^{k+1}}{(k+1)!} + \dots = e^t - 1 - t - \frac{t^2}{2} - \dots - \frac{t^{k-1}}{(k-1)!}$ hence we can calculate e^{tA} nicely in such a case. On the other hand, if the matrix A is nilpotent of order k then $A^{k-1} \neq 0$ and $A^k = 0$. Once again, the matrix exponential simplifies:

$$e^{tA} = I + tA + \frac{t^2}{2}A^2 + \dots + \frac{t^{k-1}}{(k-1)!}A^{k-1}$$

Therefore, if A is nilpotent then we can calculate the matrix exponential directly without too much trouble... of course this means we can solve $\vec{x}' = A\vec{x}$ without use of the generalized e -vector method.

Example 4.10.42. Problem: solve the system given in Example 4.10.29) by applying the

Cayley Hamilton Theorem to $A = \begin{bmatrix} 2 & 1 & -1 \\ -3 & -1 & 1 \\ 9 & 3 & -4 \end{bmatrix}$.

Solution: we found $p(\lambda) = -(\lambda - 1)^3 = 0$ hence $-(A - I)^3 = 0$. Consider the magic formula:

$$e^{tA} = e^t(I + t(A - I) + \frac{t^2}{2}(A - I)^2 + \frac{t^3}{3!}(A - I)^3 + \dots) = e^t(I + t(A - I) + \frac{t^2}{2}(A - I)^2)$$

Calculate,

$$A - I = \begin{bmatrix} 1 & 1 & -1 \\ -3 & -2 & 1 \\ 9 & 3 & -5 \end{bmatrix} \quad \& \quad (A - I)^2 = \begin{bmatrix} -11 & -4 & 5 \\ 12 & 2 & -4 \\ -45 & -12 & 19 \end{bmatrix}$$

Therefore,

$$e^{tA} = e^t \begin{bmatrix} 1 + t - \frac{11t^2}{2} & t - 2t^2 & -t + \frac{5t^2}{2} \\ -3t + 6t^2 & 1 - 2t + t^2 & t - 2t^2 \\ 9t - \frac{45t^2}{2} & 3t - 6t^2 & 1 - 5t - \frac{19t^2}{2} \end{bmatrix}$$

The general solution is given by $\vec{x}(t) = e^{tA}\vec{c}$.

Example 4.10.43. Consider for example, the system

$$x' = x + y, \quad y' = 3x - y$$

We can write this as the matrix problem

$$\underbrace{\begin{bmatrix} x' \\ y' \end{bmatrix}}_{d\vec{x}/dt} = \underbrace{\begin{bmatrix} 1 & 1 \\ 3 & -1 \end{bmatrix}}_A \underbrace{\begin{bmatrix} x \\ y \end{bmatrix}}_{\vec{x}}$$

It is easily calculated that A has eigenvalue $\lambda_1 = -2$ with e -vector $\vec{u}_1 = (-1, 3)$ and $\lambda_2 = 2$ with e -vectors $\vec{u}_2 = (1, 1)$. The general solution of $d\vec{x}/dt = A\vec{x}$ is thus

$$\vec{x}(t) = c_1 e^{-2t} \begin{bmatrix} -1 \\ 3 \end{bmatrix} + c_2 e^t \begin{bmatrix} 1 \\ 1 \end{bmatrix} = \begin{bmatrix} -c_1 e^{-2t} + c_2 e^{2t} \\ 3c_1 e^{-2t} + c_2 e^{2t} \end{bmatrix}$$

So, the **scalar solutions** are simply $x(t) = -c_1 e^{-2t} + c_2 e^{2t}$ and $y(t) = 3c_1 e^{-2t} + c_2 e^{2t}$.

Thus far I have simply told you how to solve the system $d\vec{x}/dt = A\vec{x}$ with e -vectors, it is interesting to see what this means geometrically. For the sake of simplicity we'll continue to think about the preceding example. In it's given form the DEqn is **coupled** which means the equations for the derivatives of the dependent variables x, y cannot be solved one at a time. We have to solve both at once. In the next example I solve the same problem we just solved but this time using a change of variables approach.

Example 4.10.44. Suppose we change variables using the diagonalization idea: introduce new variables $\bar{x}, \bar{y}$ by $P(\bar{x}, \bar{y}) = (x, y)$ where $P = [\vec{u}_1 | \vec{u}_2]$. Note $(\bar{x}, \bar{y}) = P^{-1}(x, y)$. We can diagonalize A by the similarity transformation by P ; $D = P^{-1}AP$ where $\text{Diag}(D) = (-2, 2)$. Note that $A = PDP^{-1}$ hence $d\vec{x}/dt = A\vec{x} = PDP^{-1}\vec{x}$. Multiply both sides by P^{-1} :

$$P^{-1} \frac{d\vec{x}}{dt} = P^{-1} P D P^{-1} \vec{x} \Rightarrow \frac{d(P^{-1}\vec{x})}{dt} = D(P^{-1}\vec{x}).$$

You might not recognize it but the equation above is decoupled. In particular, using the notation $(\bar{x}, \bar{y}) = P^{-1}(x, y)$ we read from the matrix equation above that

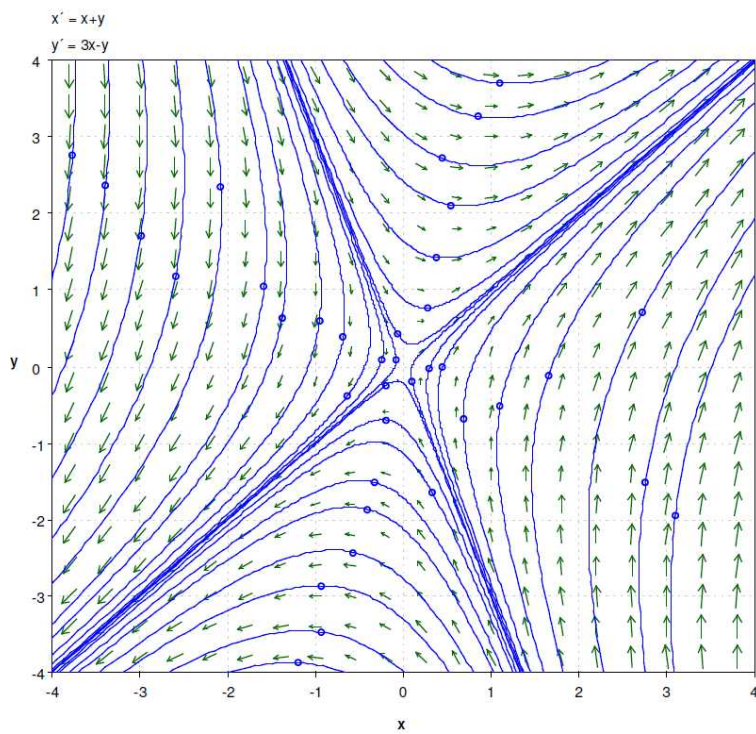
$$\frac{d\bar{x}}{dt} = -2\bar{x}, \quad \frac{d\bar{y}}{dt} = 2\bar{y}.$$

Separation of variables and a little algebra yields that $\bar{x}(t) = c_1 e^{-2t}$ and $\bar{y}(t) = c_2 e^{2t}$. Finally, to find the solution back in the original coordinate system we multiply $P^{-1}\vec{x} = (c_1 e^{-2t}, c_2 e^{2t})$ by P to isolate $\vec{x}$,

$$\vec{x}(t) = \begin{bmatrix} -1 & 1 \\ 3 & 1 \end{bmatrix} \begin{bmatrix} c_1 e^{-2t} \\ c_2 e^{2t} \end{bmatrix} = \begin{bmatrix} -c_1 e^{-2t} + c_2 e^{2t} \\ 3c_1 e^{-2t} + c_2 e^{2t} \end{bmatrix}.$$

This is the same solution we found in the last example. Usually linear algebra texts present this solution because it shows more interesting linear algebra, however, from a pragmatic viewpoint the first method is clearly faster.

Finally, we can better appreciate the solutions we found if we plot the direction field $(x', y') = (x+y, 3x-y)$ via the "ppplane" tool in Matlab. I have clicked on the plot to show a few representative trajectories (solutions):



Chapter 5

Linear Algebra with Geometry

Who hath measured the waters in the hollow of his hand, and meted out heaven with the span, and comprehended the dust of the earth in a measure, and weighed the mountains in scales, and the hills in a balance?
ISAIAH 40:12, KJV

5.1 analytic geometry for vector spaces

The foundation of analytic geometry rests on the concepts of distance between points and angles between rays. In this section we define an abstraction of the dot-product known as the *inner product* and an abstraction of the geometric vector length known as the *norm*. Angle between vectors are also defined since the Cauchy Schwarz inequality allows us to define angle algebraically even in contexts where direct visualization is impossible. Perhaps a better way to understand what we are doing in this chapter is this:

We are learning to see geometry through the lense of algebra.

Please note, we consider only vector spaces over $\mathbb{R}$ or $\mathbb{C}$ in this Chapter.

Definition 5.1.1. *inner product*

Let V be a vector space over $\mathbb{F}$ (either $\mathbb{F} = \mathbb{R}$ or $\mathbb{F} = \mathbb{C}$) then an **inner product** on V over $\mathbb{F}$ is a mapping $\langle \cdot, \cdot \rangle : V \times V \rightarrow \mathbb{F}$ such that

- (i.) $\langle x + y, z \rangle = \langle x, z \rangle + \langle y, z \rangle$
- (ii.) $\langle cx, y \rangle = c\langle x, y \rangle$
- (iii.) $\overline{\langle x, y \rangle} = \langle y, x \rangle$
- (iv.) If $x \neq 0$ then $\langle x, x \rangle \in (0, \infty)$

Given the above structure, we say $(V, \langle \cdot, \cdot \rangle)$ forms an inner product space.

Here $\overline{a + ib} = a - ib$ denotes complex conjugation¹. In the case $\mathbb{F} = \mathbb{R}$ axiom (iii.) simply reads $\langle x, y \rangle = \langle y, x \rangle$ which is known as **symmetry**. It would likely be wise to settle some complex notation up front here since it will streamline what follows. My apologies for the break in flow here.

¹my apologies, it seems I have used $(a + ib)^* = a - ib$ in previous chapters. I am trying to follow Insel Spence and Friedberg's Chapter 6 for this current work, so I am going to try to mirror their notation here, the problem is we soon introduce the adjoint which uses $*$ in a related, but distinct, fashion

Definition 5.1.2. *Hermitian adjoint of matrix*

Given $A \in \mathbb{C}^{m \times n}$ we define $\bar{A} \in \mathbb{C}^{m \times n}$ by $(\bar{A})_{ij} = \overline{A_{ij}}$ for all $1 \leq i \leq m$ and $1 \leq j \leq n$. Likewise, the **Hermitian adjoint** of A is denoted A^* and is defined by $A^* = \bar{A}^T$. In other words, A^* is the **conjugate transpose** of A .

In the Physics literature it is common to see a dagger used to denote the Hermitian adjoint of a matrix or operator; $A^* = A^\dagger$. This is a nice option if confusion with $*$ is a problem.

I'll break from my usual format here since I think an ordered list is helpful here. I'll forego much of the proof that these definitions support the axioms of an inner product. I encourage the reader to verify axioms (i.), (ii.), (iii.) and (iv.) hold in each case.

- (1.) Let $x, y \in \mathbb{R}^n$ then define $\langle x, y \rangle = x \bullet y = \sum_{i=1}^n x_i y_i = x^T y$.
- (2.) Let $x, y \in \mathbb{C}^n$ then define $\langle x, y \rangle = x \bullet \bar{y} = \sum_{i=1}^n x_i \bar{y}_i = x^T \bar{y}$.
- (3.) Let $A, B \in \mathbb{R}^{m \times n}$ then $\langle A, B \rangle = \text{tr}(AB^T)$ defines the Frobenius inner product. Notice, this formula is simply a slick way of writing out the sum $A_{11}B_{11} + A_{12}B_{12} + \cdots + A_{mn}B_{mn}$. Let's see why:

$$\langle A, B \rangle = \text{tr}(AB^T) = \sum_{i=1}^n (AB^T)_{ii} = \sum_{i=1}^n \sum_{j=1}^n A_{ij} (B^T)_{ji} = \sum_{i=1}^n \sum_{j=1}^n A_{ij} B_{ij}$$

- (4.) Let $X, Y \in \mathbb{C}^{m \times n}$ then $\langle X, Y \rangle = \text{tr}(X\bar{Y}^T) = \text{tr}(XY^*)$. There are other ways to write this formula. Notice $\text{tr}(M^T) = \text{tr}(M)$ hence

$$\langle X, Y \rangle = \text{tr}((X\bar{Y}^T)^T) = \text{tr}((\bar{Y}^T)^T X^T) = \text{tr}(\bar{Y} X^T).$$

If we write $X = A + iB$ where $A, B \in \mathbb{R}^{m \times n}$ then note

$$\langle X, X \rangle = \sum_{i,j=1}^n X_{ij} \bar{X}_{ij} = \sum_{i,j=1}^n (A_{ij} + iB_{ij})(A_{ij} - iB_{ij}) = \sum_{i,j=1}^n (A_{ij}^2 + B_{ij}^2) = \sum_{i,j=1}^n |X_{ij}|^2$$

where $|X_{ij}| = \sqrt{A_{ij}^2 + B_{ij}^2}$ is the usual notation for the length of a complex number. Notice the calculation above also indicates that $\langle X, X \rangle = \langle A, A \rangle + \langle B, B \rangle$.

- (5.) Consider $C[0, 1]$ the vector space of all continuous real-valued functions on $[0, 1]$. A natural inner product is given by the definite integral: if $f, g \in C[0, 1]$ then

$$\langle f, g \rangle = \int_0^1 f(t)g(t) dt.$$

I'll explain why axiom (iv.) holds. Notice, if f is continuous on $[0, 1]$ and if there exists $p \in [0, 1]$ for which $f(p) \neq 0$ then there exists a neighborhood N containing p for which $f(x) \neq 0$ for each $x \in N$. Thus, $f \neq 0$ implies

$$\langle f, f \rangle = \int_0^1 f(t)^2 dt > 0$$

since $f(t)^2 > 0$ for each $t \in N$ and as $f(t)^2 \geq 0$ we find the integral cannot be non-positive.

(6.) Consider $H = \{f : [0, 2\pi] \rightarrow \mathbb{C} \mid f \text{ continuous}\}$.

$$\langle f, g \rangle = \frac{1}{2\pi} \int_0^{2\pi} f(t) \overline{g(t)} dt.$$

Here, to be clear, the calculus of functions from $\mathbb{R}$ with values in $\mathbb{C}$ is simply done component-wise: if $f = f_1 + if_2$ and $g = g_1 + ig_2$ where $f_1, f_2, g_1, g_2 \in C[0, 2\pi] = \{h : [0, 2\pi] \rightarrow \mathbb{R} \mid h \text{ continuous}\}$ then

$$\frac{d}{dt}(f_1 + if_2) = \frac{df_1}{dt} + i\frac{df_2}{dt} \quad \& \quad \int_0^{2\pi} (f_1(t) + if_2(t))dt = \int_0^{2\pi} f_1(t)dt + i \int_0^{2\pi} f_2(t)dt.$$

Notice $\overline{f_1 + if_2} = f_1 - if_2$ and thus $f\bar{f} = (f_1 + if_2)(f_1 - if_2) = f_1^2 + f_2^2 = |f_1 + if_2|^2$. Hence,

$$\langle f, f \rangle = \frac{1}{2\pi} \int_0^{2\pi} |f(t)|^2 dt.$$

Once more, if $f(t) \neq 0$ for some $t_0 \in [0, 2\pi]$ then $|f(t)|^2 > 0$ for all t close to t_0 and as $|f(t)|^2 \geq 0$ we find axiom (iv.) holds. Observe that

$$\begin{aligned} \overline{\int_0^{2\pi} (f_1(t) + if_2(t))dt} &= \overline{\int_0^{2\pi} f_1(t)dt + i \int_0^{2\pi} f_2(t)dt} \\ &= \int_0^{2\pi} f_1(t)dt - i \int_0^{2\pi} f_2(t)dt \\ &= \int_0^{2\pi} (f_1(t) - if_2(t))dt \end{aligned}$$

Thus $\overline{\int_0^{2\pi} f(t)dt} = \int_0^{2\pi} \overline{f(t)}dt$. Consequently, we find axiom (iii.) holds:

$$\overline{\langle f, g \rangle} = \overline{\frac{1}{2\pi} \int_0^{2\pi} f(t) \overline{g(t)} dt} = \frac{1}{2\pi} \int_0^{2\pi} \overline{f(t) \overline{g(t)}} dt = \frac{1}{2\pi} \int_0^{2\pi} g(t) \overline{f(t)} dt = \langle g, f \rangle.$$

(7.) Let V be a real vector space with basis $\beta = \{v_1, \dots, v_n\}$. Let $x, y \in V$ then define

$$\langle x, y \rangle = [x]_\beta \bullet [y]_\beta$$

In other words, if $x = x_1v_1 + \dots + x_nv_n$ and $y = y_1v_1 + \dots + y_nv_n$ then

$$\langle x, y \rangle = x_1y_1 + \dots + x_ny_n.$$

This construction shows that every finite dimensional vector space may be assigned the structure of an inner product space. However, there is no reason the geometry implicit within this construction has much to do with any intuitive geometric structure in V . In any event, this is an important example as it means we can use geometric techniques even rather abstract contexts.

(8.) Let V be a complex vector space with basis $\beta = \{v_1, \dots, v_n\}$. If $z, w \in V$ then define

$$\langle z, w \rangle = [z]_\beta^T \overline{[w]_\beta}.$$

Once more, this construction is general. It shows there are many ways to construct an inner product on a complex vector space.

- (9.) It should be fairly evident from the previous pair of examples that the choice of inner product on a vector space is far from unique. This example further illustrates such freedom. Given an inner product space $(V, \langle \cdot, \cdot \rangle)$ and any positive constant r we define

$$\langle x, y \rangle_r = r^2 \langle x, y \rangle.$$

It is easy to verify this is an inner product for V . We will explain the reason for using r^2 a little later in this section.

- (10.) Given a complex inner product space $(V, \langle \cdot, \cdot \rangle)$ we may regard V as a real vector space. In fact, $\langle x, y \rangle_{\mathbb{R}} = \operatorname{Re}(\langle x, y \rangle)$ defines a real inner product on V as a real vector space. Since for all $z, w \in \mathbb{C}$ and $c \in \mathbb{R}$, $\operatorname{Re}(z + w) = \operatorname{Re}(z) + \operatorname{Re}(w)$ and $\operatorname{Re}(cw) = c\operatorname{Re}(w)$ we can calculate

$$\langle x + y, z \rangle_{\mathbb{R}} = \operatorname{Re}(\langle x + y, z \rangle) = \operatorname{Re}(\langle x, z \rangle) + \operatorname{Re}(\langle y, z \rangle) = \langle x, z \rangle_{\mathbb{R}} + \langle y, z \rangle_{\mathbb{R}}.$$

and

$$\langle cx, y \rangle_{\mathbb{R}} = \operatorname{Re}(\langle cx, y \rangle) = \operatorname{Re}(c\langle x, y \rangle) = c\operatorname{Re}(\langle x, y \rangle) = c\langle x, y \rangle_{\mathbb{R}}.$$

I leave proof of axioms (iii.) and (iv.) to the reader. They're simple exercises².

Proposition 5.1.3.

Let $(V, \langle \cdot, \cdot \rangle)$ be an inner product space then

(1.) $\langle \sum_{i=1}^n c_i v_i, y \rangle = \sum_{i=1}^n c_i \langle v_i, y \rangle$ for all $c_i \in \mathbb{F}$ and $v_i, y \in V$.

(2.) $\langle x, cy \rangle = \bar{c} \langle x, y \rangle$ for all $c \in \mathbb{F}$ and $x, y \in V$.

(3.) $\langle x, y + z \rangle = \langle x, y \rangle + \langle x, z \rangle$ for all $x, y, z \in V$.

(4.) $\langle x, \sum_{i=1}^n c_i v_i \rangle = \sum_{i=1}^n \bar{c}_i \langle x, v_i \rangle$ for all $c_i \in \mathbb{F}$ and $x, v_i \in V$.

(5.) $\langle x, 0 \rangle = \langle 0, x \rangle = 0$ for any $x \in V$.

(6.) For any $x \in V$, $\langle x, x \rangle = 0$ iff $x = 0$.

(7.) If $\langle x, y \rangle = \langle x, z \rangle$ for all $x \in V$ then $y = z$.

Proof: the proof of (1.) is by induction. Observe $\langle c_1 v_1, y \rangle = c_1 \langle v_1, y \rangle$ by axiom (ii.). Thus (1.)

²here I assume you understand complex numbers and the algebra needed to analyze the real part of a complex number

holds for $n = 1$. Suppose (1.) holds for n and consider

$$\begin{aligned}
 \left\langle \sum_{i=1}^{n+1} c_i v_i, y \right\rangle &= \left\langle c_{n+1} v_{n+1} + \sum_{i=1}^n c_i v_i, y \right\rangle && \text{:definition of } \sum \\
 &= c_{n+1} \langle v_{n+1}, y \rangle + \left\langle \sum_{i=1}^n c_i v_i, y \right\rangle && \text{:by axioms (i.) and (ii.)} \\
 &= c_{n+1} \langle v_{n+1}, y \rangle + \sum_{i=1}^n c_i \langle v_i, y \rangle && \text{:induction hypothesis} \\
 &= \sum_{i=1}^{n+1} c_i \langle v_i, y \rangle && \text{:definition of } \sum
 \end{aligned}$$

Therefore, (1.) holds by proof by mathematical induction. To prove (2.) with $\mathbb{F} = \mathbb{C}$ simply recall $\bar{\bar{z}} = z$ and use axiom (iii.) in what follows:

$$\langle x, cy \rangle = \overline{\overline{\langle x, cy \rangle}} = \overline{\langle cy, x \rangle} = \overline{c \langle y, x \rangle} = \bar{c} \overline{\langle y, x \rangle} = \bar{c} \langle x, y \rangle$$

since $\overline{\bar{z}w} = \bar{z}\bar{w}$ is a property of complex conjugation. In the case $\mathbb{F} = \mathbb{R}$ we have $\langle x, y \rangle = \langle y, x \rangle$ and hence (2.) simplifies to $\langle x, cy \rangle = \langle cy, x \rangle = c \langle y, x \rangle = c \langle x, y \rangle$. Next, we prove (3.), using axiom (i.) in the third equality, in the case $\mathbb{F} = \mathbb{C}$,

$$\langle x, y + z \rangle = \overline{\overline{\langle x, y + z \rangle}} = \overline{\langle y + z, x \rangle} = \overline{\langle y, x \rangle + \langle z, x \rangle} = \overline{\langle y, x \rangle} + \overline{\langle z, x \rangle} = \langle x, y \rangle + \langle x, z \rangle.$$

since $\overline{\bar{z} + \bar{w}} = \bar{z} + \bar{w}$ is a property of complex conjugation. A similar, but easier proof can be given for (3.) in the case $\mathbb{F} = \mathbb{R}$. I leave the proof of (4.) to the reader. The key to the proof of (5.) is $0 = 0 + 0$. Let $x \in V$ and consider, using the additivity proved in (3.)

$$\langle x, 0 \rangle = \langle x, 0 + 0 \rangle = \langle x, 0 \rangle + \langle x, 0 \rangle \Rightarrow \langle x, 0 \rangle = 0.$$

We conclude the proof of (5.) by nearly the same argument using axiom (i.)

$$\langle 0, x \rangle = \langle 0 + 0, x \rangle = \langle 0, x \rangle + \langle 0, x \rangle \Rightarrow \langle 0, x \rangle = 0.$$

To prove (6.) begin by noting $x = 0$ implies $\langle 0, 0 \rangle = 0$ by (5.). Conversely, suppose $x \in V$ and $\langle x, x \rangle = 0$. If $x \neq 0$ then by axiom (iv.) we have $\langle x, x \rangle \in (0, \infty)$. Thus $x = 0$ as $0 \notin (0, \infty)$. This proves (6.). Suppose $\langle x, y \rangle = \langle x, z \rangle$ for all $x \in V$. Consider,

$$0 = \langle x, y \rangle - \langle x, z \rangle = \langle x, y \rangle + \langle x, -z \rangle = \langle x, y + (-z) \rangle = \langle x, y - z \rangle$$

If $y - z \neq 0$ then we find a contradiction since $x = y - z$ would yield $\langle x, x \rangle = \langle x, y - z \rangle \neq 0$ yet $0 = \langle x, y - z \rangle$ by the calculation above. Therefore, $y - z = 0$ which is to say $y = z$ and we have shown (7.) is true. $\square$

Definition 5.1.4. norm

Let V be a vector space over $\mathbb{F}$ (either $\mathbb{F} = \mathbb{R}$ or $\mathbb{F} = \mathbb{C}$) then an **norm** on V over $\mathbb{F}$ is a mapping $\| \cdot \| : V \rightarrow [0, \infty)$ such that

(i.) $\|cx\| = |c|\|x\|$ for all $c \in \mathbb{F}$ and $x \in V$,

(ii.) $\|x + y\| \leq \|x\| + \|y\|$ for all $x, y \in V$,

(iii.) $\|x\| = 0$ if and only if $x = 0$.

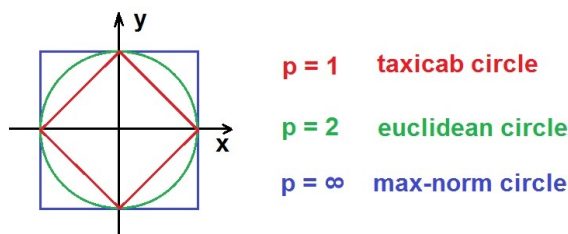
Given the above structure, we say $(V, \| \cdot \|)$ forms a normed linear space.

The term **norm** could be replaced with **length**. The value of $\|x\|$ is the length of x . We can also define distance in a normed linear space.

Definition 5.1.5. *distance*

Let $(V, \|\cdot\|)$ be a normed linear space then the **distance** from $x \in V$ to $y \in V$ is denoted $d(x, y)$ and is defined by $d(x, y) = \|y - x\|$.

Abstractly, a distance function can be defined on spaces with less structure than a vector space. If S is a set and $d : S \times S \rightarrow [0, \infty)$ has $d(x, y) = d(y, x)$ for all $x, y \in S$ and $d(x, y) = 0$ iff $x = y$ and $d(x, z) \leq d(x, y) + d(y, z)$ for all $x, y, z \in S$ then (S, d) is a **metric space**. A circle in S centered at p with radius R is the collection of all $x \in S$ for which $d(x, p) = R$. For example, in $\mathbb{R}^2$ we can define $d_p(x, y) = \sqrt[p]{(y_1 - x_1)^p + (y_2 - x_2)^p}$ and if $p \neq 2$ then the unit-circle is probably not what you expect:



Analyzing metric spaces and their structure is a typical topic in an introductory analysis class. Let us return to Linear Algebra once more.

Definition 5.1.6. *induced norm*

Let $(V, \langle \cdot, \cdot \rangle)$ be an inner product space then $\|x\| = \sqrt{\langle x, x \rangle}$ defines the **induced norm**.

Thankfully the induced norm is in fact a norm.

Proposition 5.1.7.

Let $(V, \langle \cdot, \cdot \rangle)$ be an inner product space and define $\|x\| = \sqrt{\langle x, x \rangle}$ for each $x \in V$.

- (1.) $\|\cdot\|$ defines a norm on V
- (2.) Cauchy Schwarz inequality: $|\langle x, y \rangle| \leq \|x\| \|y\|$ for all $x, y \in V$.

Proof: suppose $c \in \mathbb{F}$ and $x \in V$ then

$$\|cx\| = \sqrt{\langle cx, cx \rangle} = \sqrt{c\bar{c}\langle x, x \rangle} = \sqrt{|c|^2\langle x, x \rangle} = \sqrt{|c|^2}\sqrt{\langle x, x \rangle} = |c|\|x\|.$$

Likewise, if $x = 0$ then $\langle x, x \rangle = \langle 0, 0 \rangle = 0$ hence $\|x\| = \sqrt{0} = 0$. Conversely, suppose $\|x\| = \sqrt{\langle x, x \rangle} = 0$ then $\langle x, x \rangle = 0$ hence $x = 0$ by (6.) of Proposition 5.1.3. To prove the Cauchy Schwarz inequality we take a somewhat indirect approach. Let $x, y \in V$. If $\langle x, y \rangle = 0$ then $|\langle x, y \rangle| = 0 \leq \|x\| \|y\|$ hence the Cauchy Schwarz inequality holds. Hence assume $\langle x, y \rangle \neq 0$ in what follows. Note there exists λ with $|\lambda| = 1$ such that $\lambda\langle x, y \rangle \in \mathbb{R}$. Construct a function $Q : \mathbb{R} \rightarrow \mathbb{R}$ defined by $Q(t) = \|t\lambda x + y\|^2$ for all $t \in \mathbb{R}$. Notice $Q(t) \geq 0$ since $\|t\lambda x + y\|^2 = \langle t\lambda x + y, t\lambda x + y \rangle \in [0, \infty)$ by

axiom (iv.) of the inner product. Moreover,

$$\begin{aligned} Q(t) &= \langle t\lambda x, t\lambda x \rangle + \langle t\lambda x, y \rangle + \langle y, t\lambda x \rangle + \langle y, y \rangle \\ &= t^2 \lambda \bar{\lambda} \langle x, x \rangle + t\lambda \langle x, y \rangle + t\bar{\lambda} \langle y, x \rangle + \langle y, y \rangle \\ &= t^2 \underbrace{\|x\|^2}_A + t \underbrace{(\lambda \langle x, y \rangle + \bar{\lambda} \langle y, x \rangle)}_B + \underbrace{\|y\|^2}_C. \end{aligned}$$

We should verify $Q(t)$ has real coefficients. Identify $A = \|x\|^2$ and $C = \|y\|^2$ are clearly real. To simplify B it is helpful to write $\langle x, y \rangle = re^{i\theta}$ for some $r, \theta \in \mathbb{R}$. Hence $\lambda = e^{-i\theta}$, notice that $\lambda \langle x, y \rangle = e^{-i\theta} re^{i\theta} = r \in \mathbb{R}$ with $|\lambda| = |e^{-i\theta}| = \sqrt{\cos^2(-\theta) + \sin^2(\theta)} = 1$ as claimed. Likewise, calculate $\langle y, x \rangle = \overline{\langle x, y \rangle} = \overline{re^{i\theta}} = re^{-i\theta}$. These details will allow us to further simplify B ,

$$B = \lambda \langle x, y \rangle + \bar{\lambda} \langle y, x \rangle = e^{-i\theta} e^{i\theta} r + e^{i\theta} re^{-i\theta} = 2r = 2|\langle x, y \rangle|.$$

Since the quadratic function $Q(t) = At^2 + Bt + C$ in t has non-negative values the discriminant $B^2 - 4AC \leq 0$. Thus,

$$4|\langle x, y \rangle|^2 - 4\|x\|^2\|y\|^2 \leq 0 \Rightarrow |\langle x, y \rangle| \leq \|x\|\|y\|.$$

Finally, we work towards proving Axiom (ii.), the triangle inequality. Suppose $x, y \in V$ and observe

$$\begin{aligned} \|x + y\|^2 &= \langle x + y, x + y \rangle \\ &= \langle x, x \rangle + \langle x, y \rangle + \langle y, x \rangle + \langle y, y \rangle \\ &\leq \|x\|^2 + |\langle x, y \rangle| + |\langle y, x \rangle| + \|y\|^2 \quad (\text{triangle inequality for } \mathbb{C}) \\ &\leq \|x\|^2 + 2|\langle x, y \rangle| + \|y\|^2 \\ &\leq \|x\|^2 + 2\|x\|\|y\| + \|y\|^2 \quad (\text{Cauchy Schwarz inequality}) \\ &= (\|x\| + \|y\|)^2 \end{aligned}$$

thus $\|x + y\| \leq \|x\| + \|y\|$. $\square$

Once we have established the Cauchy Schwarz inequality we are free to define the angle between nonzero vectors. There are two cases to consider:

- If $\mathbb{F} = \mathbb{R}$ then the real inner product space V has $|\langle x, y \rangle| \leq \|x\|\|y\|$. Thus, for $x, y \neq 0$ we find $\frac{|\langle x, y \rangle|}{\|x\|\|y\|} \leq 1$ and hence $-1 \leq \frac{\langle x, y \rangle}{\|x\|\|y\|} \leq 1$. It is a fact of trigonometry that there exists a unique $\theta \in [0, \pi]$ for which $\cos \theta = \frac{\langle x, y \rangle}{\|x\|\|y\|}$.
- If $\mathbb{F} = \mathbb{C}$ then the complex inner product space V has $|\langle x, y \rangle| \leq \|x\|\|y\|$ where $|\langle x, y \rangle|$ denotes the length of the complex number $\langle x, y \rangle$. Once more, if $x, y \neq 0$ we deduce that $\frac{|\langle x, y \rangle|}{\|x\|\|y\|} \leq 1$. However, we cannot meaningfully write an inequality involving $\langle x, y \rangle$ alone since it is a complex quantity in this context. In this case, we have to content ourselves with the fact that there is a unique $\hat{\theta} \in [0, \pi/2]$ for which $\cos \hat{\theta} = \frac{|\langle x, y \rangle|}{\|x\|\|y\|}$.

Definition 5.1.8. *angle in inner product space*

Let $(V, \langle \cdot, \cdot \rangle)$ be an inner product space over $\mathbb{F}$ with nonzero vectors x, y then we define for

- (i.) $\mathbb{F} = \mathbb{R}$ the $\angle(x, y) = \theta$ to be the unique $\theta \in [0, \pi]$ for which $\cos \theta = \frac{\langle x, y \rangle}{\|x\| \|y\|}$,
- (ii.) $\mathbb{F} = \mathbb{C}$ the $\tilde{\angle}(x, y) = \tilde{\theta}$ to be the unique $\tilde{\theta} \in [0, \pi/2]$ for which $\cos \tilde{\theta} = \frac{|\langle x, y \rangle|}{\|x\| \|y\|}$.

I usually refer to θ as the **real angle** and $\tilde{\theta}$ as the **complex angle**. However, please note, these are both real numbers in the context in which they are defined. In the case of a complex inner product space we have opportunity to contrast the meaning of real and complex angle in an overlapping context.

Let V be a complex inner product space with inner product $\langle \cdot, \cdot \rangle$. If $x, y \in V$ are nonzero then there are at least two competing ideas to describe the angle between these vectors:

- (1.) the **complex angle** defined by $\cos \tilde{\theta} = \frac{|\langle x, y \rangle|}{\|x\| \|y\|}$ where $\tilde{\theta} \in [0, \pi/2]$.
- (2.) or, viewing V as a real vector space with inner product given by $\langle x, y \rangle_{\mathbb{R}} = \operatorname{Re} \langle x, y \rangle$, the **real angle** between x, y is given by $\cos \theta = \frac{\langle x, y \rangle_{\mathbb{R}}}{\|x\|_{\mathbb{R}} \|y\|_{\mathbb{R}}}$ where $\|x\|_{\mathbb{R}} = \sqrt{\langle x, x \rangle_{\mathbb{R}}}$ and $\theta \in [0, \pi]$.

I'll illustrate with an explicit example.

Example 5.1.9. Consider $x = (1, 1 + i)$ and $y = (1 - i, 2)$ in $\mathbb{C}^2$. Calculate,

$$\langle x, y \rangle = (1, 1 + i) \cdot (1 - i, 2) = 1 + i + (1 + i)(2) = 3 + 3i \Rightarrow \langle x, y \rangle_{\mathbb{R}} = 3$$

$$\langle x, x \rangle = (1, 1 + i) \cdot (1, 1 - i) = 1 + (1 + i)(1 - i) = 3 \Rightarrow \langle x, x \rangle_{\mathbb{R}} = 3$$

$$\langle y, y \rangle = (1 - i, 2) \cdot (1 + i, 2) = (1 - i)(1 + i) + 4 = 6 \Rightarrow \langle y, y \rangle_{\mathbb{R}} = 6$$

Hence $\|x\|_{\mathbb{R}} = \|x\| = \sqrt{3}$ and $\|y\|_{\mathbb{R}} = \|y\| = \sqrt{6}$. Thus

$$\cos \tilde{\theta} = \frac{|3 + 3i|}{\sqrt{3}\sqrt{6}} = \frac{\sqrt{9+9}}{\sqrt{18}} = 1 \Rightarrow \tilde{\theta} = 0.$$

whereas

$$\cos \theta = \frac{3}{\sqrt{3}\sqrt{6}} = \frac{1}{\sqrt{2}} \Rightarrow \theta = \frac{\pi}{4}.$$

Notice the angle between vectors $(1, 0, 1, 1)$ and $(1, -1, 2, 0)$ is also $\pi/4$. There is an isomorphism $\Psi : \mathbb{C}^2 \rightarrow \mathbb{R}^4$ at work here; $\Psi(a + ib, c + id) = (a, b, c, d)$. In fact, it is more than an isomorphism, it also preserves the angle between vectors and vector length. We call such a map an **isometry**. Perhaps this comment helps make sense of the real angle measured in $\mathbb{C}^2$, it is just a complex notation for $\mathbb{R}^4$ with the usual Euclidean inner product. The complex angle has its own role as well. Observe:

$$(1 - i)x = (1 - i)(1, 1 + i) = (1 - i, (1 - i)(1 + i)) = (1 - i, 2) = y$$

thus $\{x, y\}$ is linearly dependent in $\mathbb{C}^2$ as a complex vector space. In contrast, there does not exist a real constant c for which $cx = y$ hence $\{x, y\}$ is linearly independent in $\mathbb{C}^2$ as a real vector space.

Consider $x, -x \in V$ where $x \neq 0$ in a complex inner product space. Observe $\|x\| = \|-x\|$ and $|\langle x, -x \rangle| = |-\langle x, x \rangle| = |\langle x, x \rangle| = \|x\|^2$ thus both $\frac{|\langle x, x \rangle|}{\|x\|\|x\|}$ and $\frac{|\langle x, -x \rangle|}{\|x\|\|-x\|}$ both reduce to 1 and hence

$$\tilde{Z}(x, x) = 0 = \tilde{Z}(x, -x).$$

In contrast, for a real inner product space we can work out that for $x \neq 0$,

$$\angle(x, x) = 0 \quad \& \quad \angle(x, -x) = \pi.$$

Obviously the concept of complex angle is not as intuitive as the real angle. You will not find the complex angle discussed nearly as much as the real angle. In any event, I suppose I ought to define linear isometry for reference since I just used it in passing in the example above.

Definition 5.1.10. *linear isometry*

Let $(V, \langle \cdot, \cdot \rangle_V)$ and $(W, \langle \cdot, \cdot \rangle_W)$ be an inner product spaces of the same dimension over $\mathbb{F}$ then $\Psi : V \rightarrow W$ is a **linear isometry** if Ψ is an isomorphism for which

$$\langle x, y \rangle_V = \langle \Psi(x), \Psi(y) \rangle_W$$

for all $x, y \in V$.

I invite the reader to verify that a linear isometry preserves the angle between vectors and the length of each vector under its action; $\angle_V(x, y) = \angle_W(\Psi(x), \Psi(y))$ and $\|x\|_V = \|\Psi(x)\|_W$. It turns out, if we study all maps which preserve distance between arbitrary pairs of points in $\mathbb{R}^n$ then after some work we can show such a map is a bijection which is formed by the composition of a linear Euclidean isometry and a translation. Such maps are called **rigid motions**³.

Definition 5.1.11. *orthogonal and orthonormal*

Let $(V, \langle \cdot, \cdot \rangle)$ be an inner product space over $\mathbb{F}$ then we say $x, y \in V$ are **orthogonal** if $\langle x, y \rangle = 0$. In this case we write $x \perp y$. If $S \subseteq V$ and if $x, y \in S$ with $x \neq y$ implies $\langle x, y \rangle = 0$ then we say S is an **orthogonal** subset of V . If $S \subseteq V$ is orthogonal and if each $x \in S$ has $\|x\| = 1$ then S is an **orthonormal** subset of V . An **orthonormal basis** of V is a basis of V which is orthonormal. Likewise, an **orthogonal basis** of V is a basis of V which is orthogonal.

Orthogonality does depend on the choice of base field when there is a choice to be made. Notice $\langle x, y \rangle = a + ib = 0$ only if both $a = 0$ and $b = 0$. However, $\langle x, y \rangle_{\mathbb{R}} = \operatorname{Re}(a + ib) = a$ only needs $a = 0$ to obtain the orthogonal condition $\langle x, y \rangle_{\mathbb{R}} = 0$. We see orthogonality in the complex sense implies orthogonality in the real sense. However, orthogonality in the real sense need not imply orthogonality in the complex sense. For instance, 1 and i are orthogonal in $\mathbb{C}$ with respect to the real Euclidean geometry of $\mathbb{C} = \mathbb{R}^2$. However, 1 and i are not orthogonal in the complex sense as $\langle 1, i \rangle = 1(-i) = -i \neq 0$.

Proposition 5.1.12.

Suppose $z = x + y$ where $x \perp y$ then $\|z\|^2 = \|x\|^2 + \|y\|^2$.

Proof: Suppose $z = x + y$ where $\langle x, y \rangle = 0$. Note $\overline{\langle y, x \rangle} = \langle x, y \rangle = 0$ hence $\langle y, x \rangle = 0$. Consider,

$$\|z\|^2 = \langle z, z \rangle = \langle x + y, x + y \rangle = \langle x, x \rangle + \langle x, y \rangle + \langle y, x \rangle + \langle y, y \rangle = \|x\|^2 + \|y\|^2. \quad \square$$

³I usually prove this in our Abstract Algebra course, the proof is not entirely obvious, ask if interested

5.2 orthogonality

In this section we study the properties and creation of orthogonal and orthonormal sets. Many of the standard bases are orthonormal.

Example 5.2.1. For $\mathbb{R}^n$ the standard basis $e_1, \dots, e_n$ has $\langle e_i, e_j \rangle = e_i \cdot e_j = \delta_{ij} = \begin{cases} 1 & i = j \\ 0 & i \neq j \end{cases}$. This means dot-product of distinct standard basis vectors is zero and $\|e_i\| = \sqrt{e_i \cdot e_i} = \sqrt{1} = 1$. Thus the standard basis of $\mathbb{R}^n$ is an orthonormal basis.

Example 5.2.2. For $\mathbb{C}^n$ the standard basis $e_1, \dots, e_n$ has $\langle e_i, e_j \rangle = e_i^T \overline{e_j} = \delta_{ij}$ and it follows $\{e_i\}_{i=1}^n$ is an orthonormal basis for the complex vector space $\mathbb{C}^n$.

Example 5.2.3. For $\mathbb{R}^{m \times n}$ the standard basis of matrix units E_{ij} defined by $(E_{ij})_{kl} = \delta_{ik}\delta_{jl}$ can be constructed by products of $f_i \in \mathbb{R}^{m \times 1}$ and $e_j \in \mathbb{R}^{n \times 1}$ via $E_{ij} = f_i e_j^T$ then

$$(E_{ij})_{kl} = f_k^T f_i e_j^T e_l = \delta_{ki} \delta_{jl}$$

where I have used that $f_k^T f_i = f_k \cdot f_i = \delta_{ki}$ and $e_j^T e_l = e_j \cdot e_l = \delta_{jl}$. Then calculate,

$$E_{ij} E_{kl}^T = f_i e_j^T (f_k e_l^T)^T = f_i e_j^T e_l f_k^T = \delta_{jl} F_{ik}$$

where we introduce $F_{ik} \in \mathbb{R}^{m \times m}$ as the $m \times m$ matrix unit with 1 in the (i, k) -th component and zeros elsewhere. Consider then,

$$\langle E_{ij}, E_{kl} \rangle = \text{tr}(E_{ij} E_{kl}^T) = \text{tr}(\delta_{jl} F_{ik}) = \sum_{a=1}^m \delta_{jl} (F_{ik})_{aa} = \sum_{a=1}^m \delta_{jl} \delta_{ia} \delta_{ka} = \delta_{ik} \delta_{jl}.$$

Thus $\{E_{ij} \mid 1 \leq i \leq m, 1 \leq j \leq n\}$ serves as an orthonormal basis for $\mathbb{R}^{m \times n}$ with the usual Frobenius inner product. Note, very similar arguments prove $\{E_{ij} \mid 1 \leq i \leq m, 1 \leq j \leq n\}$ serves as an orthonormal basis for $\mathbb{C}^{m \times n}$ with the inner product $\langle A, B \rangle = \text{tr}(AB^*)$.

Not every natural basis is orthonormal.

Example 5.2.4. Consider $P_2(\mathbb{R})$ with basis $\beta = \{1, x, x^2\}$. Use inner product $\langle f, g \rangle = \int_{-1}^1 f(x)g(x) dx$.

$$\langle 1, 1 \rangle = \int_{-1}^1 dx = 2, \quad \& \quad \langle x, x \rangle = \int_{-1}^1 x^2 dx = 2/3, \quad \& \quad \langle x^2, x^2 \rangle = \int_{-1}^1 x^4 dx = 2/5.$$

thus $\|1\| = \sqrt{2}$ and $\|x\| = \sqrt{2/3}$ and $\|x^2\| = \sqrt{2/5}$ hence β is not normalized. Notice,

$$\langle x, 1 \rangle = \int_{-1}^1 x dx = 0 \quad \& \quad \langle x, x^2 \rangle = \int_{-1}^1 x^3 dx = 0$$

thus $x \perp 1$ and $x \perp x^2$. However,

$$\langle 1, x^2 \rangle = \int_{-1}^1 x^2 dx = 2/3 \neq 0$$

hence β is not orthogonal.

Example 5.2.5. Consider $H = \{f : [0, 2\pi] \rightarrow \mathbb{C} \mid f \text{ continuous}\}$.

$$\langle f, g \rangle = \frac{1}{2\pi} \int_0^{2\pi} f(t) \overline{g(t)} dt.$$

Let $S = \{e^{int} \mid n \in \mathbb{N}\}$. Notice $\frac{d}{dt}e^{int} = ine^{int}$ hence $\int e^{int} dt = \frac{1}{in}e^{int} + c$. Thus, for $m, n \in \mathbb{N}$ with $m \neq n$,

$$\langle e^{imt}, e^{int} \rangle = \frac{1}{2\pi} \int_0^{2\pi} e^{imt} \overline{e^{int}} dt = \frac{1}{2\pi} \int_0^{2\pi} e^{i(m-n)t} dt = \frac{1}{2\pi i(m-n)} (e^{2\pi(m-n)i} - 1) = 0.$$

Likewise,

$$\langle e^{int}, e^{int} \rangle = \frac{1}{2\pi} \int_0^{2\pi} e^{int} \overline{e^{int}} dt = \frac{1}{2\pi} \int_0^{2\pi} dt = \frac{2\pi}{2\pi} = 1.$$

Thus S is an orthonormal subset of H .

Proposition 5.2.6.

Let $(V, \langle \cdot, \cdot \rangle)$ be an inner product space and suppose $S = \{v_1, \dots, v_k\}$ is a set of nonzero orthogonal vectors.

(1.) S is linearly independent,

(2.) If $x \in \text{span}(S)$ then $x = \sum_{i=1}^k \frac{\langle x, v_i \rangle}{\langle v_i, v_i \rangle} v_i$.

(3.) If S is orthonormal and $x \in \text{span}(S)$ then $x = \sum_{i=1}^k \langle x, v_i \rangle v_i$.

Proof: S orthogonal means $\langle v_i, v_j \rangle = \delta_{ij} \langle v_i, v_i \rangle$. Suppose $x = \sum_{j=1}^k c_j v_j$ then

$$\langle x, v_i \rangle = \left\langle \sum_{j=1}^k c_j v_j, v_i \right\rangle \Rightarrow \langle x, v_i \rangle = \sum_{j=1}^k c_j \langle v_j, v_i \rangle = \sum_{j=1}^k c_j \delta_{ij} \langle v_i, v_i \rangle = c_i \langle v_i, v_i \rangle.$$

Notice $\langle v_i, v_i \rangle = \|v_i\|^2 \neq 0$ as $v_i \neq 0$ for each i . Therefore, $c_i = \frac{\langle x, v_i \rangle}{\langle v_i, v_i \rangle}$ and (2.) and (3.) follow.

Notice, if we suppose $x = 0 = \sum_{j=1}^k c_j v_j$ then $c_i = \frac{\langle 0, v_i \rangle}{\langle v_i, v_i \rangle} = 0$ thus S is LI hence (1.) is true. $\square$

The Gram Schmidt Algorithm allows us to replace a linearly independent subset of an inner product space with an orthogonal set which spans the same subspace as the given set. In other words, this algorithm means we are free to create an orthonormal basis in the context of an inner product space. Moreover, since we showed any finite dimensional vector space can be given an inner product this means are free to create an orthonormal basis in any finite dimensional context. However, it may or may not be possible to maintain other structures jointly with the desired orthonormality.

Theorem 5.2.7. *Gram Schmidt Algorithm*

Let $(V, \langle \cdot, \cdot \rangle)$ be an inner product space and suppose $S = \{v_1, \dots, v_n\}$ is a set of linearly independent vectors. The Gram Schmidt Algorithm is as follows:

- Let $v'_1 = v_1$
- Let $v'_2 = v_2 - \frac{\langle v_2, v'_1 \rangle}{\langle v'_1, v'_1 \rangle} v'_1$
- Let $v'_3 = v_3 - \frac{\langle v_3, v'_1 \rangle}{\langle v'_1, v'_1 \rangle} v'_1 - \frac{\langle v_3, v'_2 \rangle}{\langle v'_2, v'_2 \rangle} v'_2$
- In summary, $v'_k = v_k - \sum_{i=1}^{k-1} \frac{\langle v_k, v'_i \rangle}{\langle v'_i, v'_i \rangle} v'_i$ for $k = 1, \dots, n$.

Then $S' = \{v'_1, \dots, v'_n\}$ has $\text{span}(S) = \text{span}(S')$ and S' is orthogonal. Let $v''_i = \frac{1}{\|v'_i\|} v'_i$ for $i = 1, \dots, n$ and $S'' = \{v''_1, \dots, v''_n\}$ then S'' is orthonormal with $\text{span}(S) = \text{span}(S'')$.

Proof: let $S_k = \{v_1, \dots, v_k\}$ define $S'_k = \{v'_1, \dots, v'_k\}$ where v'_i are defined as described above. Notice $S'_1 = \{v_1\}$ thus S'_1 is an orthogonal set of nonzero vectors⁴ with $\text{span}(S'_1) = \text{span}(S_1)$. Inductively suppose S'_k is an orthogonal set of nonzero vectors with $\text{span}(S'_k) = \text{span}(S_k)$. Consider $S'_{k+1} = S'_k \cup \{v'_{k+1}\}$ and suppose towards a contradiction that $v'_{k+1} = 0$. Then, by definition of v'_{k+1} ,

$$v'_{k+1} = v_{k+1} - \sum_{i=1}^k \frac{\langle v_{k+1}, v'_i \rangle}{\langle v'_i, v'_i \rangle} v'_i = 0 \Rightarrow v_{k+1} = \sum_{i=1}^k \frac{\langle v_{k+1}, v'_i \rangle}{\langle v'_i, v'_i \rangle} v'_i \in \text{span}(S_k)$$

thus S_{k+1} is linearly dependent, but we assumed from the outset that $\{v_1, \dots, v_{k+1}\} = S_{k+1}$ is linearly independent. Thus $v'_{k+1} \neq 0$ and we find S'_{k+1} is a set of nonzero vectors. It remains to show S'_{k+1} is orthogonal. By induction hypothesis we know S'_k is orthogonal hence any pair of vectors taken from $S'_k \subset S'_{k+1}$ is orthogonal. Consider for $1 \leq i \leq k$, using the formula for v'_{k+1} once more,

$$\begin{aligned} \langle v'_{k+1}, v'_i \rangle &= \langle v_{k+1}, v'_i \rangle - \sum_{j=1}^k \frac{\langle v_{k+1}, v'_j \rangle}{\langle v'_j, v'_j \rangle} \langle v'_j, v'_i \rangle \\ &= \langle v_{k+1}, v'_i \rangle - \sum_{j=1}^k \frac{\langle v_{k+1}, v'_j \rangle}{\langle v'_j, v'_j \rangle} \delta_{ij} \langle v'_i, v'_i \rangle \\ &= \langle v_{k+1}, v'_i \rangle - \langle v_{k+1}, v'_i \rangle \\ &= 0 \end{aligned}$$

where the simplification $\langle v'_j, v'_i \rangle = \delta_{ij} \langle v'_i, v'_i \rangle$ is the orthogonality of S'_k given by the induction hypothesis. Hence S'_{k+1} is an orthogonal set of nonzero vectors. Therefore, by proof by mathematical induction we find S'_n is a nonzero orthogonal set provided S_n is LI. Moreover, we proved a nonzero orthogonal set of n -vectors is LI hence S'_n is LI and as S_n and S'_n are both LI sets where $S'_n \subset \text{span}(S_n)$ by construction it follows that $\text{span}(S_n) = \text{span}(S'_n)$. The remaining claims about S''_n we leave to the reader. $\square$

⁴orthogonality is automatic for a set with one vector

Remark 5.2.8.

It is also possible to normalize as you go with the Gram Schmidt Algorithm. In particular, with $S = \{v_1, \dots, v_n\}$ linearly independent subset of $(V, \langle \cdot, \cdot \rangle)$ we set $v_1'' = \frac{1}{\|v_1\|} v_1$. Then $v_2' = v_2 - \langle v_2, v_1'' \rangle v_1''$ and $v_2'' = \frac{1}{\|v_2'\|} v_2'$. Next, $v_3' = v_3 - \langle v_3, v_1'' \rangle v_1'' - \langle v_3, v_2'' \rangle v_2''$ and $v_3'' = \frac{1}{\|v_3'\|} v_3'$. Etc. In practice, I make more errors with this method, so I've made a habit of separating the normalization and the orthogonalization in this course.

Example 5.2.9. Consider $\mathbb{R}^4$ with the inner product given by the dot-product. Suppose $v_1 = (1, 0, 1, 0)$, $v_2 = (1, 1, 1, 1)$ and $v_3 = (0, 1, 2, 1)$. Then

$$v_2' = v_2 - \frac{v_2 \cdot v_1}{v_1 \cdot v_1} v_1 = (1, 1, 1, 1) - \frac{2}{2}(1, 0, 1, 0) = (0, 1, 0, 1).$$

and

$$v_3' = v_3 - \frac{v_3 \cdot v_1}{v_1 \cdot v_1} v_1 - \frac{v_3 \cdot v_2'}{v_2' \cdot v_2'} v_2' = (0, 1, 2, 1) - \frac{2}{2}(1, 0, 1, 0) - \frac{2}{2}(0, 1, 0, 1) = (-1, 0, 1, 0).$$

Thus $S' = \{(1, 0, 1, 0), (0, 1, 0, 1), (-1, 0, 1, 0)\}$ and normalization yields

$$S'' = \left\{ \frac{1}{\sqrt{2}}(1, 0, 1, 0), \frac{1}{\sqrt{2}}(0, 1, 0, 1), \frac{1}{\sqrt{2}}(-1, 0, 1, 0) \right\}.$$

Next we return to Example 5.2.4 to find an orthonormal basis for $P_2(\mathbb{R})$.

Example 5.2.10. Consider $P_2(\mathbb{R}) = \text{span}(\beta)$ where $\beta = \{1, x, x^2\}$. We use the inner product $\langle f, g \rangle = \int_{-1}^1 f(x)g(x) dx$. Apply the Gram Schmidt algorithm to β . We set $v_1' = 1$ and note $\langle 1, x \rangle = 0$ thus $v_2' = x - \langle 1, x \rangle 1 = x$. Consider,

$$v_3' = x^2 - \frac{\langle x^2, 1 \rangle}{\langle 1, 1 \rangle} 1 - \frac{\langle x^2, x \rangle}{\langle x, x \rangle} x = x^2 - \frac{2/3}{2} 1 = x^2 - \frac{1}{3}.$$

Normalization requires some calculation,

$$\langle v_3', v_3' \rangle = \int_{-1}^1 \left(x^2 - \frac{1}{3} \right)^2 dx = \int_{-1}^1 \left(x^4 - \frac{2}{3}x^2 + \frac{1}{9} \right) dx = \frac{2}{5} - \frac{4}{9} + \frac{2}{9} = \frac{18 - 10}{45} = \frac{8}{45}$$

Thus,

$$v_3'' = \sqrt{\frac{45}{8}} \left(x^2 - \frac{1}{3} \right) = \sqrt{\frac{5}{8}} (3x^2 - 1).$$

Hence, $S'' = \left\{ \frac{1}{\sqrt{2}}, \sqrt{\frac{3}{2}}x, \sqrt{\frac{5}{8}}(3x^2 - 1) \right\}$ is an orthonormal basis for $P_2(\mathbb{R})$. These are examples of **Legendre polynomials**.

Legendre polynomials are a particular instance of a set of orthogonal polynomials. There is a vast literature of how orthogonal polynomials provide solutions to famous problems of mathematical physics. For instance, the Legendre polynomials appear in the multi-pole expansion of electricity and magnetism. Hermite polynomials appear in the solution of Schrodinger's Equation for the Hydrogen atom. There are lengthy texts focused on detailing the gory formulas for all such problems. You can also find these formulas using modern CAS software.

Proposition 5.2.11.

Let $(V, \langle \cdot, \cdot \rangle)$ be a finite dimensional inner product space. Then there exists an orthonormal basis $\beta = \{w_1, \dots, w_n\}$ and $x = \sum_{i=1}^n \langle x, w_i \rangle w_i$ for any $x \in V$.

Proof: choose a basis for V and apply the Gram Schmid algorithm to create an orthonormal basis for V . Apply part (3.) of Proposition 5.2.6 to complete the proof. $\square$

There is also a nice formula for the matrix of a linear transformation given an orthonormal basis.

Proposition 5.2.12.

Let $(V, \langle \cdot, \cdot \rangle)$ be an inner product space with orthonormal basis $\beta = \{v_1, \dots, v_n\}$ and $T : V \rightarrow V$ is a linear transformation then if $[T]_{\beta, \beta} = A$ then $A_{ij} = \langle T(v_j), v_i \rangle$.

Proof: if $\beta = \{v_1, \dots, v_n\}$ is an orthonormal basis then $(\Phi_\beta(x))_i = \langle x, v_i \rangle$. If $T : V \rightarrow V$ then $[T]_{\beta, \beta} = [[T(v_1)]_\beta] \cdots [[T(v_n)]_\beta]$ hence $[T]_{\beta, \beta} = A$ implies $A_{ij} = ([T(v_j)]_\beta)_i = \langle T(v_j), v_i \rangle$. $\square$

Definition 5.2.13. *orthogonal complement*

Let $(V, \langle \cdot, \cdot \rangle)$ be an inner product space. If $S \subseteq V$ and $S \neq \emptyset$ then the **orthogonal complement** of S is $S^\perp$ which is defined by $S^\perp = \{x \in V \mid \langle x, y \rangle = 0 \text{ for all } y \in S\}$.

Notice $\langle 0, y \rangle = 0$ for all $y \in V$ hence $0^\perp = V$. Likewise, as $\langle x, y \rangle = 0$ for all $y \in V$ implies $x = 0$ implies $V^\perp = \{0\}$. Both of these simple claims hold for any inner product space V .

Example 5.2.14. Consider $\mathbb{R}^3$ with the dot-product,

$$\{e_3\}^\perp = \text{span}\{e_1, e_2\}, \quad \{e_2\}^\perp = \text{span}\{e_1, e_3\}, \quad \{e_1\}^\perp = \text{span}\{e_2, e_3\}$$

and

$$\{e_1, e_2\}^\perp = \text{span}\{e_3\}, \quad \{e_1, e_3\}^\perp = \text{span}\{e_2\}, \quad \{e_2, e_3\}^\perp = \text{span}\{e_1\}$$

Example 5.2.15. Consider $\mathbb{R}^{2 \times 2}$ with the Frobenius inner product. Then

$$\{E_{11}, E_{22}, E_{12} + E_{21}\}^\perp = \text{span}(E_{12} - E_{21}).$$

Notice the proposition below holds for infinite dimensional V and W .

Proposition 5.2.16. *complement of subspace and complement of basis coincide*

Let $(V, \langle \cdot, \cdot \rangle)$ be an inner product space. If $W = \text{span}(\beta)$ then $W^\perp = \beta^\perp$.

Proof: Let $x \in W^\perp$ then $\langle x, y \rangle = 0$ for all $y \in W$. Hence $\langle x, v_i \rangle$ for each $v_i \in \beta \subset W$ and we find $x \in \beta^\perp$. Thus $W^\perp \subseteq \beta^\perp$. Next, suppose $x \in \beta^\perp$ then $\langle x, v \rangle = 0$ for each $v \in \beta$. Let $w \in W = \text{span}(\beta)$ then there exist $c_i \in \mathbb{F}$ and $v_i \in \beta$ for which $w = \sum_{i=1}^k c_i v_i$. Observe,

$$\langle x, w \rangle = \sum_{i=1}^k \overline{c_i} \langle x, v_i \rangle = \sum_{i=1}^k \overline{c_i} (0) = 0$$

thus $x \in W^\perp$. Hence, $\beta^\perp \subseteq W^\perp$ and we conclude $W^\perp = \beta^\perp$. $\square$

Proposition 5.2.17. *subspace and its complement*

Let $(V, \langle \cdot, \cdot \rangle)$ be a finite dimensional over $\mathbb{F}$ with subspace W then $V = W \oplus W^\perp$.

Proof: let $\dim(W) = k$ and suppose $\beta_W = \{v_1, \dots, v_k\}$ is an orthonormal basis for W . Extend β_W to a basis for V and orthonormalize that basis to create the orthonormal basis $\beta = \beta_W \cup \beta' = \{v_1, \dots, v_n\}$ for V . Notice $\beta_W \subseteq \beta$ by the Gram Schmidt algorithm. Notice $W^\perp$ can be shown to be a subspace by the subspace test. Note, $0 \in W^\perp$ since $\langle 0, w \rangle = 0$ for all $w \in W$ hence $W^\perp \neq \emptyset$. Let $x, y \in W^\perp$ and $c \in \mathbb{F}$ then $\langle x, w \rangle = 0$ and $\langle y, w \rangle = 0$ for all $w \in W$ by definition of $W^\perp$. Hence,

$$\langle cx + y, w \rangle = c\langle x, w \rangle + \langle y, w \rangle = c(0) + 0 = 0.$$

Therefore, $cx + y \in W^\perp$ and we conclude $W^\perp \leq V$. Suppose $x \in W \cap W^\perp$ then $x \in W$ and $x \in W^\perp$ hence $\langle x, x \rangle = 0$ thus $x = 0$. It follows $W \cap W^\perp = \{0\}$. Notice $\beta' = \{v_{k+1}, \dots, v_n\} \in W^\perp$ since $\langle v_j, v_i \rangle = 0$ as $i \neq j$ when $1 \leq i \leq k$ and $k+1 \leq j \leq n$ and once more we apply Proposition 5.2.16. Notice $\beta' \subset W^\perp$ implies $n - k \leq \dim(W^\perp)$. Recall Theorem 2.7.5 gives that

$$\dim(W + W^\perp) = \dim(W) + \dim(W^\perp) - \dim(W \cap W^\perp)$$

hence as $W + W^\perp \leq V$ we have $\dim(W + W^\perp) \leq n$ we deduce

$$\dim(W + W^\perp) = k + \dim(W^\perp) \leq n \Rightarrow \dim(W^\perp) \leq n - k.$$

Thus $n - k \leq \dim(W^\perp) \leq n - k$ and we have shown $\dim(W^\perp) = n - k$. Hence $W + W^\perp$ has $\dim(W + W^\perp) = \dim(W) + \dim(W^\perp) = k + n - k = n$ and we conclude $W + W^\perp = V$ and thus $W \oplus W^\perp = V$. $\square$

The proof above can be shortened by showing $W^\perp \leq V$ separately. It may well be possible to prove $W + W^\perp = V$ without appealing to Theorem 2.7.5, I decided to use it here to illustrate its power. We should also note that $W \cap W^\perp = \{0\}$ can be shown even in the infinite dimensional context. In what follows, notice we only assume W is finite dimensional, it could be the case that V is of infinite dimension.

Proposition 5.2.18. *subspace complement*

Let $(V, \langle \cdot, \cdot \rangle)$ be an inner product space and suppose W is a finite dimensional subspace of V . Let $x \in V$ then there exist unique vectors $y \in W$ and $z \in W^\perp$ such that $x = y + z$.
Moreover, if $\{w_1, \dots, w_k\}$ is an orthonormal basis for W then $y = \sum_{j=1}^k \langle x, w_j \rangle w_j$.

Proof: apply the Gram Schmidt algorithm to create an orthonormal basis $\beta = \{v_1, \dots, v_k\}$ for W . We have $\langle v_i, v_j \rangle = \delta_{ij}$. Let $x \in V$ and construct $y = \sum_{j=1}^k \langle x, v_j \rangle v_j$. Observe $y \in \text{span}(\beta) = W$.

Next, construct $z = x - y$ hence $x = y + z$. Next we show $z \in W^\perp$,

$$\begin{aligned}
 \langle z, v_j \rangle &= \langle x - y, v_j \rangle \\
 &= \langle x, v_j \rangle - \langle y, v_j \rangle \\
 &= \langle x, v_j \rangle - \left\langle \sum_{i=1}^k \langle x, v_i \rangle v_i, v_j \right\rangle \\
 &= \langle x, v_j \rangle - \sum_{i=1}^k \langle x, v_i \rangle \langle v_i, v_j \rangle \\
 &= \langle x, v_j \rangle - \sum_{i=1}^k \langle x, v_i \rangle \delta_{ij} \\
 &= \langle x, v_j \rangle - \langle x, v_j \rangle \\
 &= 0.
 \end{aligned}$$

Thus $\langle v_j, z \rangle = 0$ for all $j = 1, \dots, k$. Therefore, applying Proposition 5.2.16, we find $z \in W^\perp$. Suppose $y' \in W$ and $z' \in W^\perp$ such that $x = y' + z'$ then $y + z = y' + z'$ and $y' - y = z - z'$. Observe $y, y' \in W$ implies $y' - y \in W$. Likewise, $z, z' \in W^\perp$ implies $z - z' \in W^\perp$. Hence $y' - y, z - z' \in W \cap W^\perp = \{0\}$ thus $y' - y = 0 = z - z'$ hence $y = y'$ and $z = z'$ which proves the desired uniqueness for y and z . $\square$

Proposition 5.2.19. *closest vector*

Let $(V, \langle \cdot, \cdot \rangle)$ be an inner product space and suppose W is a finite dimensional subspace of V . Let $x \in V$ then if $y \in W$ and $z \in W^\perp$ such that $x = y + z$ then y is the vector in W closest to x and z is the vector in $W^\perp$ closest to x . Here we define **closest vector** in $S \leq V$ to $x \in V$ to be the vector in $s_0 \in S$ for which $\|x - s_0\| \leq \|x - s\|$ for all $s \in S$.

Proof: by Proposition 5.2.18 if $x \in V$ then there exist unique $y \in W$ and $z \in W^\perp$ for which $x = y + z$. Let $w \in W$ and consider

$$\|x - w\|^2 = \|y + z - w\|^2 = \|y - w\|^2 + \|z\|^2$$

since $y, w \in W$ implies $y - w \in W$ and $z \in W^\perp$ gives $(y - w) \perp z$ hence the Pythagorean identity holds. Observe $\|x - w\|^2 \geq \|z\|^2$ with equality in the case $y = w$. In other words, y is the closest point in W to x and the distance from x to y is given to be $\|x - y\| = \|z\|$ which means the distance⁵ from x to W is given by $\|z\|$. Likewise, if we consider $z' \in W^\perp$ then

$$\|x - z'\|^2 = \|y + z - z'\|^2 = \|y\|^2 + \|z - z'\|^2$$

as $y \perp (z - z')$ and hence $z \in W^\perp$ is the point closest to x . In fact, the distance from x to z is seen to be $\|y\|$ as $\|x - z\| = \|y\|$. $\square$

We define $Proj_W$ and $Orth_W$ based on the decomposition $V = W \oplus W^\perp$.

⁵distance from a point to an extended object is sometimes understood as the smallest possible distance which is attained between points in the object and the given point. Or, as may be the case in analysis, the infimum of all such possible distances since outside our context the distance of closest approach may not actually be attained except in some limiting sense. For example, if we take $y = |x|$ and delete the origin from its graph and ask how far it is from $(0, -1)$ the answer is 1 even though there is technically no point on the origin-deleted $y = |x|$ graph which is distance 1 from $(0, -1)$.

Definition 5.2.20. *projection*

Let V be a finite dimensional inner product space with subspace W . If $\{w_1, \dots, w_k\}$ is an orthonormal basis for W then we define: $Proj_W : V \rightarrow W$ and $Orth_W : V \rightarrow W^\perp$ by

$$Proj_W(x) = \sum_{j=1}^k \langle x, w_j \rangle w_j \quad \& \quad Orth_W(x) = x - \sum_{j=1}^k \langle x, w_j \rangle w_j$$

for each $x \in V$.

In view of Proposition 5.2.19 we see that the projection and orthogonal projection with respect to W give us formulas to select closest points in W or $W^\perp$.

Proposition 5.2.21.

Let $W \leq V$ where V is a finite dimensional inner product space then $Proj_W : V \rightarrow W$ and $Orth_W : V \rightarrow W^\perp$ are linear transformations and $Proj_W \circ Proj_W = Proj_W$ and $Orth_W \circ Orth_W = Orth_W$. Moreover, $Proj_{W^\perp} = Orth_W$ and $Orth_{W^\perp} = Proj_W$. Finally,

$$Id_V = Proj_W + Orth_W.$$

Proof: I leave this for the reader. It is a good exercise. You ought to be able to prove it using techniques which have been used to prove the other results in this section. $\square$

We conclude this section with a theorem which is nearly a restatement of Proposition 5.2.17. Once more I leave the proof to the reader, but we have shown most of this already in this section.

Proposition 5.2.22.

Let $S = \{v_1, \dots, v_k\}$ be an orthonormal set in $(V, \langle \cdot, \cdot \rangle)$ with $\dim(V) = n$. Then,

- (i.) S can be extended to an orthonormal basis $\{v_1, \dots, v_k, v_{k+1}, \dots, v_n\}$ for V ,
- (ii.) If $W = \text{span}(S)$ then $\{v_{k+1}, \dots, v_n\}$ is an orthonormal basis for $S^\perp = W^\perp$,
- (iii.) For any subspace W of V , $\dim(V) = \dim(W) + \dim(W^\perp)$.

5.3 theory of adjoints

Let V be a finite dimensional inner product space with orthonormal basis $\beta = \{v_1, \dots, v_n\}$. Let $T : V \rightarrow V$ be a linear transformation then $([T]_{\beta, \beta})_{ij} = \langle T(v_j), v_i \rangle$. Suppose we define $S : V \rightarrow V$ to be the linear transformation on V for which $[S]_{\beta, \beta}^* = [T]_{\beta, \beta}$; that is, $\overline{([S]_{\beta, \beta})_{ji}} = ([T]_{\beta, \beta})_{ij}$. Thus,

$$\overline{([S]_{\beta, \beta})_{ji}} = \overline{\langle S(v_i), v_j \rangle} = \langle v_j, S(v_i) \rangle = \langle T(v_j), v_i \rangle$$

We find $\langle T(v_j), v_i \rangle = \langle v_j, S(v_i) \rangle$ for all i, j . It follows by properties of the inner product⁶ that

$$\langle T(x), y \rangle = \langle x, S(y) \rangle$$

⁶this is a good exercise for the reader, it also makes a nice test question

for all $x, y \in V$. Suppose $R : V \rightarrow V$ is linear and $\langle T(x), y \rangle = \langle x, R(y) \rangle$ for all $x, y \in V$ then

$$\langle x, R(y) \rangle = \langle T(x), y \rangle = \langle x, S(y) \rangle$$

for all $x, y \in V$. Thus it follows⁷ $R(y) = S(y)$ for all $y \in V$. Thus, $R = S$ and we deduce there is a unique linear transformation S for which $\langle x, S(y) \rangle = \langle x, T(y) \rangle$ for all $x, y \in V$. Moreover, $[S]_{\beta, \beta}^* = [T]_{\beta, \beta}$. But, as we know $(A^*)^* = A$ hence $([S]_{\beta, \beta}^*)^* = [T]_{\beta, \beta}^*$ thus $[T]_{\beta, \beta}^* = [S]_{\beta, \beta}$.

The arguments above support the existence of the adjoint of a linear transformation on a finite dimensional inner product space. In the infinite dimensional context it may or may not be the case that the adjoint of a linear transformation exists.

Definition 5.3.1. *adjoint of a linear transformation*

Let V be an inner product space and let $T : V \rightarrow V$ be a linear transformation. Then we call the unique linear transformation $T^* : V \rightarrow V$ for which $\langle T(x), y \rangle = \langle x, T^*(y) \rangle$ for all $x, y \in V$ the **adjoint** of T provided such a transformation exists.

Proposition 5.3.2.

Suppose $T : V \rightarrow V$ is a linear transformation on a finite dimensional inner product space with orthonormal basis β then

- (i.) T^* exists and $[T]_{\beta, \beta}^* = [T^*]_{\beta, \beta}$,
- (ii.) $\langle x, T(y) \rangle = \langle T^*(x), y \rangle$ for all $x, y \in V$.

Proof: the proof of (i.) is found at the beginning of this section.

To prove (ii.) consider, since $\langle T(y), x \rangle = \langle y, T^*(x) \rangle$ for all $x, y \in V$ we find $\overline{\langle T(y), x \rangle} = \overline{\langle y, T^*(x) \rangle}$ thus, by axiom (iii.) of the inner product, $\langle x, T(y) \rangle = \langle T^*(x), y \rangle$ for all $x, y \in V$. $\square$

Remark 5.3.3.

For a real matrix, $A^* = \overline{A}^T = A^T$. This simplifies the calculation of the adjoint in the context of real inner product spaces.

Example 5.3.4. Consider $D : P_2(\mathbb{R}) \rightarrow P_2(\mathbb{R})$ where $D(ax^2 + bx + c) = 2ax + b$. Let $\beta = \left\{ \frac{1}{\sqrt{2}}, \sqrt{\frac{3}{2}}x, \sqrt{\frac{5}{8}}(3x^2 - 1) \right\}$ and recall from Example 5.2.10 we know that β is an orthonormal basis for $P_2(\mathbb{R})$ with respect to the inner product $\langle f, g \rangle = \int_{-1}^1 f(x)g(x) dx$. Calculate,

$$D\left(\frac{1}{\sqrt{2}}\right) = 0 \Rightarrow \text{col}_1([D]_{\beta, \beta}) = 0.$$

and

$$D\left(\sqrt{\frac{3}{2}}x\right) = \sqrt{\frac{3}{2}} = \sqrt{3}\left(\frac{1}{\sqrt{2}}\right) \Rightarrow \text{col}_2([D]_{\beta, \beta}) = \begin{bmatrix} \sqrt{3} \\ 0 \\ 0 \end{bmatrix}.$$

⁷again, a good exercise, we didn't quite prove this, we have proved $\langle x, y \rangle = \langle z, y \rangle$ for all $y \in V$ implies $x = z$, but this isn't quite that pattern

and

$$D\left(\sqrt{\frac{5}{8}}(3x^2 - 1)\right) = 6x\sqrt{\frac{5}{8}} = \left(\sqrt{\frac{3}{2}}x\right)\sqrt{15} \Rightarrow \text{col}_3([D]_{\beta,\beta}) = \begin{bmatrix} 0 \\ \sqrt{15} \\ 0 \end{bmatrix}.$$

Thus $[D]_{\beta,\beta} = \begin{bmatrix} 0 & \sqrt{3} & 0 \\ 0 & 0 & \sqrt{15} \\ 0 & 0 & 0 \end{bmatrix}$. Hence $[D^*]_{\beta,\beta} = \begin{bmatrix} 0 & 0 & 0 \\ \sqrt{3} & 0 & 0 \\ 0 & \sqrt{15} & 0 \end{bmatrix}$ which means

$$D^*\left(a\frac{1}{\sqrt{2}} + b\sqrt{\frac{3}{2}}x + c\sqrt{\frac{5}{8}}(3x^2 - 1)\right) = a\sqrt{3}\sqrt{\frac{3}{2}}x + b\sqrt{15}\sqrt{\frac{5}{8}}(3x^2 - 1).$$

Example 5.3.5. Consider $H = \{f : [0, 2\pi] \rightarrow \mathbb{C} \mid f \text{ continuous}\}$.

$$\langle f, g \rangle = \frac{1}{2\pi} \int_0^{2\pi} f(t)\overline{g(t)} dt.$$

Form $V = \text{span}\{e^{it}, e^{2it}, e^{3it}\} \leq H$ and recall from Example 5.2.5 that $\beta = \{e^{it}, e^{2it}, e^{3it}\}$ forms an orthonormal basis for V with respect to the given inner product. Define $T : V \rightarrow V$ by

$$T(ae^{it} + be^{2it} + ce^{3it}) = (a + ib)e^{it} + (b + ic)e^{2it} + (c + ia)e^{3it}$$

Hence,

$$[T]_{\beta,\beta} = \begin{bmatrix} 1 & i & 0 \\ 0 & 1 & i \\ i & 0 & 1 \end{bmatrix} \Rightarrow [T^*]_{\beta,\beta} = [T]_{\beta,\beta}^* = \begin{bmatrix} 1 & 0 & -i \\ -i & 1 & 0 \\ 0 & -i & 1 \end{bmatrix}$$

which implies

$$T^*(ae^{it} + be^{2it} + ce^{3it}) = (a - ic)e^{it} + (b - ia)e^{2it} + (c - ib)e^{3it}.$$

It is also possible to calculate the adjoint without explicit use of an orthonormal basis.

Example 5.3.6. Consider $T : \mathbb{R}^{2 \times 2} \rightarrow \mathbb{R}^{2 \times 2}$ defined by $T(A) = MA$ where M is a given 2×2 matrix. We wish to find T^* for which $\langle T(A), B \rangle = \langle A, T^*(B) \rangle$. Using the Frobenius norm, we need:

$$\text{tr}(MAB^T) = \text{tr}(A(T^*(B))^T)$$

Notice $\text{tr}(MAB^T) = \text{tr}(AB^T M)$. Thus we find $(T^*(B))^T = B^T M$. Hence, $T^*(B) = M^T B$.

Proposition 5.3.7.

Suppose V is an inner product space and $S, T \in \mathcal{L}(V)$ for which $S^*, T^* \in \mathcal{L}(V)$. Let $c \in \mathbb{F}$. Then, $S + T, cT, ST, T^*, Id$ all have adjoints and

- (i.) $(S + T)^* = S^* + T^*$,
- (ii.) $(cT)^* = \bar{c}T^*$,
- (iii.) $(ST)^* = T^*S^*$,
- (iv.) $(T^*)^* = T$,
- (v.) $Id^* = Id$.

Proof: I'll prove (iii.) and leave the rest as exercises. Let $x, y \in V$ and observe:

$$\langle (ST)(x), y \rangle = \langle S(T(x)), y \rangle = \langle T(x), S^*(y) \rangle = \langle x, T^*(S^*(y)) \rangle = \langle x, (T^*S^*)(y) \rangle.$$

Therefore, ST has an adjoint and $(ST)^* = T^*S^*$. $\square$

We should note there are corresponding results for matrices and their adjoints.

Proposition 5.3.8.

Let $A, B \in \mathbb{F}^{n \times n}$ and $c \in \mathbb{F}$ then $L_A : \mathbb{F}^n \rightarrow \mathbb{F}^n$ defined by $L_A(x) = Ax$ for all $x \in \mathbb{F}^n$ has $(L_A)^* = L_{A^*}$ and

- (i.) $(A + B)^* = A^* + B^*$,
- (ii.) $(cA)^* = \bar{c}A^*$,
- (iii.) $(AB)^* = B^*A^*$,
- (iv.) $(A^*)^* = A$,
- (v.) $I^* = I$.

Proof: If $L_A : \mathbb{F}^n \rightarrow \mathbb{F}^n$ is defined by $L_A(x) = Ax$ for all $x \in \mathbb{F}^n$ then observe

$$\langle L_A(x), y \rangle = \langle Ax, y \rangle = x^T A^T \bar{y} = x^T \overline{A^T y} = x^T \overline{A^* y} = \langle x, A^* y \rangle = \langle x, L_{A^*}(y) \rangle$$

Therefore, $(L_A)^* = L_{A^*}$. Observe $L_{A+B} = L_A + L_B$ and $L_{cA} = cL_A$ and $L_AL_B = L_{AB}$ and $L_I = Id$. Thus,

$$L_{(A+B)^*} = (L_{A+B})^* = (L_A + L_B)^* = (L_A)^* + (L_B)^* = L_{A^*} + L_{B^*} = L_{A^*+B^*}$$

and as $L_M = L_N$ if and only if $M = N$ we find $(A + B)^* = A^* + B^*$. Now, you might ask yourself, is there an easier way to prove (i.). Sure. Note $(A + B)^T = A^T + B^T$ and $\overline{M + N} = \overline{M} + \overline{N}$ thus

$$(A + B)^* = \overline{(A + B)^T} = \overline{A^T + B^T} = \overline{A^T} + \overline{B^T} = A^* + B^*.$$

Likewise,

$$(AB)^* = \overline{(AB)^T} = \overline{B^T A^T} = \overline{B^T} \overline{A^T} = B^* A^*.$$

Proofs of the remaining items can be proved either by exploiting the correspondence with the previous proposition or via explicit calculations based on properties of transposition and complex conjugation. $\square$

Remark 5.3.9.

We could discuss the method of least squares at this point. However, I have already covered this topic in Math 221 so I place the emphasis here elsewhere.

5.4 diagonalization in inner product spaces

Proposition 5.4.1.

Let $(V, \langle \cdot, \cdot \rangle)$ inner product of finite dimension and suppose $T : V \rightarrow V$ is a linear map. If T has an eigenvector with eigenvalue λ then T^* has an eigenvector with eigenvalue $\bar{\lambda}$.

Proof: Suppose $T(v) = \lambda v$ for $v \neq 0$. For any $x \in V$ we calculate:

$$0 = \langle 0, x \rangle = \langle (T - \lambda)(v), x \rangle = \langle v, (T - \lambda)^*(x) \rangle.$$

Thus $v \perp (T - \lambda)^*(x)$ for all $x \in V$. Consequently, $v \in \text{Range}(T^* - \bar{\lambda})^\perp$. We know that $\text{Range}(T^* - \bar{\lambda})^\perp \oplus \text{Range}(T^* - \bar{\lambda}) = V$. Since $v \neq 0$ we find $\text{Range}(T^* - \bar{\lambda}) \neq V$. Since $T^* - \bar{\lambda}$ is a linear map on V it satisfies the rank nullity theorem:

$$\dim(V) = \dim(\text{Ker}(T^* - \bar{\lambda})) + \dim(\text{Range}(T^* - \bar{\lambda}))$$

and we deduce $\dim(\text{Ker}(T^* - \bar{\lambda})) \geq 1$. Therefore, there exists $y \in V$ for which $(T^* - \bar{\lambda})(y) = 0$ hence y is an eigenvector with eigenvalue $\bar{\lambda}$ for T^* . $\square$

Definition 5.4.2. split polynomial

Let $f(x) \in \mathbb{F}[x]$ then $f(x)$ is **split** over $\mathbb{F}$ if all factors of $f(x)$ are linear;

$$f(x) = a(x - \lambda_1)(x - \lambda_2) \cdots (x - \lambda_n).$$

Theorem 5.4.3. Schur's Theorem:

Let $(V, \langle \cdot, \cdot \rangle)$ inner product of finite dimension and suppose $T : V \rightarrow V$ is a linear map such that $P_T(x) = \det(T - x)$ splits. Then there exists an orthonormal basis γ for V such that $[T]_{\gamma, \gamma}$ is upper-triangular.

Proof (sketch): since $P_T(x)$ splits it follows all the eigenvalues of T are in $\mathbb{F}$ hence a Jordan basis for T exists. Note, the Jordan form is upper triangular and if we apply the Gram Schmidt algorithm then it processes the Jordan basis from first to last and the process preserves the upper-triangular shape of the matrix of T . $\square$

Insel Spence and Friedberg, 5th edition, give a proof on page 367 which is based on weaving together a handful of results shown in their homework exercises.

Definition 5.4.4. normal operator

Let V be an inner product space and $T \in \mathcal{L}(V)$. We say T is **normal** if $TT^* = T^*T$. Likewise, $A \in \mathbb{F}^{n \times n}$ is **normal** if $AA^* = A^*A$.

Proposition 5.4.5.

If V is a finite dimensional inner product space with orthonormal basis β and $T \in \mathcal{L}(V)$ is normal then $[T]_{\beta, \beta}$ is normal.

Proof: let $\beta = \{v_1, \dots, v_n\}$ be an orthonormal basis for V and suppose $T : V \rightarrow V$ is normal. Consider, since T is normal,

$$[TT^*]_{\beta, \beta} = [T^*T]_{\beta, \beta}$$

However, the matrix of a composition is the product of the matrices for the transformations composed. That is,

$$[T]_{\beta,\beta}[T^*]_{\beta,\beta} = [T^*]_{\beta,\beta}[T]_{\beta,\beta}.$$

Hence $[T]_{\beta,\beta}$ is normal. $\square$

Example 5.4.6. Let $A = \begin{bmatrix} \cos \theta & -\sin \theta \\ \sin \theta & \cos \theta \end{bmatrix}$ then $AA^T = I = A^T A$ and as $A^* = A^T$ in the context of $\mathbb{R}^{n \times n}$ we find A is normal. Since A is the matrix of L_A with respect to the orthonormal standard basis for $\mathbb{R}^n$ we deduce L_A is normal.

Notice L_A in the above example is the rotation in the plane by angle θ . This is not diagonalizable as it is clear such a rotation has no eigenvector (unless $\theta = 0$ or $\theta = \pi$). Normal transformations need not be diagonalizable.

Example 5.4.7. If $A \in \mathbb{R}^{n \times n}$ and $A^T = A$ then $AA^T = AA = A^T A$ hence A is normal. Likewise, if $B \in \mathbb{R}^{n \times n}$ and $B^T = -B$ then $B^T B = -BB = B(-B) = BB^T$ hence B is normal. Furthermore, L_A and L_B are normal transformations on $\mathbb{R}^n$.

Proposition 5.4.8.

Let V be an inner product space and $T \in \mathcal{L}(V)$ normal. Then,

- (i.) $\|T(x)\| = \|T^*(x)\|$ for all $x \in V$,
- (ii.) $T - c$ is normal for every $c \in \mathbb{F}$,
- (iii.) if $T(x) = \lambda x$ for $x \neq 0$ then $T^*(x) = \bar{\lambda}x$,
- (iv.) if $\lambda_1 \neq \lambda_2$ are eigenvalues of T with eigenvectors x_1, x_2 , then $x_1 \perp x_2$.

Proof: Suppose $T : V \rightarrow V$ is normal. Let $x \in V$,

$$\|T(x)\|^2 = \langle T(x), T(x) \rangle = \langle x, T^*(T(x)) \rangle = \langle x, T(T^*(x)) \rangle = \langle T^*(x), T^*(x) \rangle = \|T^*(x)\|^2.$$

Thus $\|T(x)\| = \|T^*(x)\|$ for all $x \in V$ and (i.) is true. Let $c \in \mathbb{F}$ and note $(T - c)^* = T^* - \bar{c}$. Thus,

$$(T - c)^*(T - c) = (T^* - \bar{c})(T - c) = T^*T - cT^* - \bar{c}T + c\bar{c}.$$

Likewise,

$$(T - c)(T - c)^* = (T - c)(T^* - \bar{c}) = TT^* - cT^* - \bar{c}T + c\bar{c}.$$

Recall T normal means $T^*T = TT^*$ hence $(T - c)^*(T - c) = (T - c)(T - c)^*$ and we conclude $T - c$ is normal and the proof of (ii.) is complete. Suppose there exists $x \neq 0$ and $\lambda \in \mathbb{F}$ for which $T(x) = \lambda x$. Apply (i.) to see that

$$\|(T - \lambda)(x)\| = \|(T - \lambda)^*(x)\|.$$

However, $(T - \lambda)(x) = T(x) - \lambda x = 0$ hence $\|(T - \lambda)(x)\| = 0$ and thus $\|(T - \lambda)^*(x)\| = 0$. Therefore, $(T - \lambda)^*(x) = 0$ which yields $T^*(x) = \bar{\lambda}x$ which completes the proof of (iii.). Let $\lambda_1, \lambda_2 \in \mathbb{F}$ with

$\lambda_1 \neq \lambda_2$ and suppose there exist $x_1, x_2 \neq 0$ for which $T(x_1) = \lambda_1 x_1$ and $T(x_2) = \lambda_2 x_2$. Consider,

$$\begin{aligned}\lambda_1 \langle x_1, x_2 \rangle &= \langle \lambda_1 x_1, x_2 \rangle \\ &= \langle T(x_1), x_2 \rangle \\ &= \langle x_1, T^*(x_2) \rangle \\ &= \langle x_1, \overline{\lambda_2} x_2 \rangle \\ &= \overline{\lambda_2} \langle x_1, x_2 \rangle \\ &= \lambda_2 \langle x_1, x_2 \rangle.\end{aligned}$$

Thus $(\lambda_2 - \lambda_1) \langle x_1, x_2 \rangle = 0$ and since $\lambda_1 \neq \lambda_2$ we find $\langle x_1, x_2 \rangle = 0$ thus $x_1 \perp x_2$. $\square$

Theorem 5.4.9. *In complex vector space, orthonormally diagonalizable iff normal.*

If V is a finite dimensional complex inner product space and $T \in \mathcal{L}(V)$ then T is normal if and only if there exists an orthonormal eigenbasis for T .

Proof: suppose T is normal over $V(\mathbb{C})$ then $P_T(x) = \det(T - x)$ splits and thus by Schur's Theorem there exists an orthonormal basis $\beta = \{v_1, \dots, v_n\}$ for which $A = [T]_{\beta, \beta}$ is upper-triangular. Notice $T(v_1) = A_{11}v_1$ thus v_1 is an eigenvector with eigenvalue $\lambda_1 = A_{11}$. Inductively suppose $v_1, \dots, v_{k-1}$ are eigenvectors of T with eigenvalues $\lambda_1, \dots, \lambda_{k-1}$. Consider $1 \leq j \leq k-1$ and note upper-triangularity of A implies

$$T(v_k) = A_{1k}v_1 + A_{2k}v_2 + \dots + A_{jk}v_j + \dots + A_{kk}v_k = \sum_{i=1}^k A_{ik}v_i$$

Observe,

$$\langle T(v_k), v_j \rangle = \sum_{i=1}^k A_{ik} \langle v_i, v_j \rangle = \sum_{i=1}^k A_{ik} \delta_{ij} = A_{jk}.$$

However, by definition of $[T]_{\beta, \beta}$ since β is orthonormal; $A_{jk} = \langle T(v_k), v_j \rangle = \langle v_k, T^*(v_j) \rangle$. Next, by the induction hypothesis and orthonormality of β ,

$$A_{jk} = \langle v_k, \overline{\lambda_j} v_j \rangle = \lambda_j \langle v_k, v_j \rangle = 0.$$

Consequently, we find $T(v_k) = A_{kk}v_k$ and so $\lambda_k = A_{kk}$. Thus, by induction, we find β is an orthonormal eigenbasis for T .

Assume there exists an orthonormal eigenbasis $\beta = \{v_1, \dots, v_n\}$ for T . In particular, $T(v_i) = \lambda_i v_i$ for $i = 1, \dots, n$ and $\langle v_i, v_j \rangle = \delta_{ij}$. Let $A = [T]_{\beta, \beta}$ and calculate

$$A_{ij} = \langle T(v_j), v_i \rangle = \langle \lambda_j v_j, v_i \rangle = \lambda_j \langle v_j, v_i \rangle = \lambda_j \delta_{ji}.$$

Consequently, $(A^*)_{ij} = \overline{\lambda_j} \delta_{ij}$. Calculate,

$$\begin{aligned}(A^*A)_{ij} &= \sum_{k=1}^n (A^*)_{ik} A_{kj} = \sum_{k=1}^n \overline{\lambda_k} \delta_{ik} \lambda_j \delta_{kj} = \overline{\lambda_i} \lambda_j \delta_{ji} \\ (AA^*)_{ij} &= \sum_{k=1}^n A_{ik} (A^*)_{kj} = \sum_{k=1}^n \lambda_k \delta_{ki} \overline{\lambda_j} \delta_{kj} = \lambda_j \overline{\lambda_i} \delta_{ij} = \overline{\lambda_i} \lambda_j \delta_{ji}\end{aligned}$$

Thus $A^*A = AA^*$ and thus T is normal since $[T]_{\beta, \beta} = A$ is normal. $\square$

Example 5.4.10. Consider $H = \{f : [0, 2\pi] \rightarrow \mathbb{C} \mid f \text{ continuous}\}$.

$$\langle f, g \rangle = \frac{1}{2\pi} \int_0^{2\pi} f(t) \overline{g(t)} dt.$$

Following arguments in Example 5.2.5 we can show $\beta = \{e^{int} \mid n \in \mathbb{Z}\}$ forms an orthonormal subset of H . Notice $V = \text{span}(\beta)$ form an interesting infinite dimensional subspace of H . In fact, $\dim(V) = \aleph_0$ since it has a countably infinite basis. Define $R : V \rightarrow V$ by $R(e^{int}) = e^{i(n+1)t}$ and $L : V \rightarrow V$ by $L(e^{int}) = e^{i(n-1)t}$. These are right and left shift operators on the basis $\beta = \{\dots, e^{-2it}, e^{-it}, 1, e^{it}, e^{2it}, \dots\}$. It can be shown that R and L are normal, however, neither R nor L are orthonormally diagonalizable. This illustrates the fact that the assumption of finite dimensionality in the theorem above is a necessary one.

Insel Spence and Friedberg, 5th edition, give arguments on pages 369-370 which flesh out the claims of the example above.

Definition 5.4.11. *Hermitian or self-adjoint operator*

Let V be an inner product space and $T \in \mathcal{L}(V)$. We say T is **self-adjoint** or **Hermitian** if $T = T^*$. Likewise, $A \in \mathbb{F}^{n \times n}$ is **self-adjoint** or **Hermitian** if $A = A^*$.

Proposition 5.4.12.

If V is a finite dimensional inner product space and assume $T \in \mathcal{L}(V)$ is self-adjoint. Then,

- (i.) every eigenvalue of T is real
- (ii.) if V is a real inner product space then $P_T(x) = \det(T - x)$ splits.

Proof: let $T : V \rightarrow V$ be a linear map on the finite dimensional inner product space V and suppose $T^* = T$. If there exists $x \neq 0$ and $\lambda \in \mathbb{F}$ for which $T(x) = \lambda x$ then

$$\lambda \langle x, x \rangle = \langle \lambda x, x \rangle = \langle T(x), x \rangle = \langle x, T^*(x) \rangle = \langle x, T(x) \rangle = \langle x, \lambda x \rangle = \bar{\lambda} \langle x, x \rangle.$$

However, $x \neq 0$ hence $\langle x, x \rangle \neq 0$ and we find $\lambda = \bar{\lambda}$ and we have shown (i.) is true. Next, suppose V is a real inner product space with orthonormal basis β . Let $[T]_{\beta, \beta} = A \in \mathbb{R}^{n \times n}$ and note $T^* = T$ implies $A^* = A$. Define $L_A : \mathbb{C}^n \rightarrow \mathbb{C}^n$ by $L_A(x) = Ax$ for each $x \in \mathbb{C}^n$. Note, $\det(L_A - x)$ is split over $\mathbb{C}$. However, L_A is self-adjoint hence every eigenvalue of L_A is real. Recall T and $[T]_{\beta, \beta}$ have the same characteristic polynomial hence $\det(T - x)$ is split over $\mathbb{R}$. $\square$

Theorem 5.4.9 is very similar to what follows. The essential difference is that normality implies the characteristic polynomial splits over $\mathbb{C}$ whereas we need the self-adjoint criteria to be sure the eigenvalues are real in the context of a real inner product space. That said, the proofs of both theorems are rather similar.

Theorem 5.4.13. *Orthonormally diagonalizable iff self-adjoint.*

If V is a finite dimensional real inner product space and $T \in \mathcal{L}(V)$ then T is self-adjoint if and only if there exists an orthonormal eigenbasis for T .

Proof: Let V be a finite dimensional real inner product space. Suppose $T \in \mathcal{L}(V)$ is self-adjoint. Then Proposition 5.4.12 shows the characteristic polynomial of T is split. Then we can repeat the proof given for Theorem 5.4.9 to see there exists an orthonormal eigenbasis for T .

Conversely, assume there exists a real orthonormal eigenbasis $\beta = \{v_1, \dots, v_n\}$ for T . In particular, $T(v_i) = \lambda_i v_i$ for $i = 1, \dots, n$ and $\langle v_i, v_j \rangle = \delta_{ij}$ and $\lambda_i \in \mathbb{R}$. Let $A = [T]_{\beta, \beta}$ and calculate

$$A_{ij} = \langle T(v_j), v_i \rangle = \langle \lambda_j v_j, v_i \rangle = \lambda_j \langle v_j, v_i \rangle = \lambda_j \delta_{ji} = \lambda_i \delta_{ij} = A_{ji}.$$

thus $A = A^T$ and so $A = A^*$ as $A^T = A^*$ for real matrices. Thus $T = T^*$. $\square$

Chapter 6

Multilinear Algebra

Praise the LORD from the earth, ye dragons, and all deeps: PSALMS 148:7, KJV

We study multilinear maps. Bilinear and trilinear maps are described in detail as well as their components. A bilinear form on an n -dimensional vector space is essentially given by an $n \times n$ matrix. However, a multilinear map with more than two inputs corresponds to an object with n -indices. We place indices up or down to reflect the nature of the index as it relates to either V or its dual V^* . These so-called *contravariant* and *covariant* indices change coordinates inversely much the same as components and bases. A basis for the set of multilinear maps is provided by forming $\otimes$ (tensor) products of basis and dual bases. We then introduce symmetric and antisymmetric multilinear maps and we find there is a natural connection between the determinant and antisymmetric maps. This leads us to group certain collections of tensor products into so-called $\wedge$ -products. At this stage, the $\wedge$ -product of dual vectors gives us a method to construct a particular antisymmetric map on V . Our initial section is very heavily computational and coordinate-based. Our second exposure to $\wedge$ later in the chapter has a more abstract algebraic flavor.

Next we turn to the problem of describing an *algebra*. In short, an algebra is simply a vector space paired with a multiplication. This leads us to our penultimate § 6.4 on wedge product viewed as an abstract algebra. This section does not view wedge products of vectors or dual vectors as multilinear maps. Instead, we begin with the algebraic properties of $\wedge$ and attempt to derive and define concepts from that algebraic basis. Of course, all the algebraic properties are present when we construct the wedge product from particular multilinear maps, but we can think of the multilinear maps as just a particular model which exhibits the structure of the exterior algebra. In fact, we could also build the exterior algebra on a subset of matrices if we were so inclined¹. We find a new method to study determinants in terms of the structure of the exterior algebra. This leads us to find an elegant proof that $\det(AB) = \det(A)\det(B)$ as well as to the characterization of linear dependence of less than n -vectors in an n -dimensional context.

Finally, we conclude by studying linear geometry. A vector space paired with a metric is a linear geometry. A metric is a bilinear, symmetric, nondegenerate form. Every inner product is a metric. However, not every metric is an inner product. For instance, we study the Minkowski metric which underlies the spacetime of Special Relativity. The concept of the Reisz vector generalizes here to the so-called *musical morphisms*. We describe how to change a vector to a dual vector or how to change a dual vector to a vector. These musical morphisms give natural isomorphisms from

¹if you are interested, ask, I can make a homework of this

tensors of differing types. Physicists make wide use of such isomorphisms in the formulation of Physical laws where a scalar is typically formed by the contraction of covariant and contravariant indices². The set of all metric-preserving maps is also of interest. These *isometries* give us a better understanding of the meaning of a particular geometry. For example, the origin-preserving isometries of Euclidean geometry are simply the orthogonal transformations. In contrast, the Lorentz transformations appear as origin preserving isometries of the Minkowski metric.

6.1 multilinearity and the tensor product

A multilinear mapping is a function of a Cartesian product of vector spaces which is linear with respect to each "slot". The goal of this section is to explain what that means. It turns out the set of all multilinear mappings on a particular set of vector spaces forms a vector space and we'll show how the tensor product can be used to construct an explicit basis by tensoring a bases which are dual to the bases in the domain. We also examine the concepts of symmetric and antisymmetric multilinear mappings, these form interesting subspaces of the set of all multilinear mappings. Our approach in this section is to treat the case of bilinearity in depth then transition to the case of multilinearity. Naturally this whole discussion demands a familiarity with the preceding section.

6.1.1 bilinear maps

Definition 6.1.1.

Suppose V_1, V_2 are vector spaces then $b : V_1 \times V_2 \rightarrow \mathbb{R}$ is a **bilinear mapping** on $V_1 \times V_2$ iff for all $x, y \in V_1, z, w \in V_2$ and $c \in \mathbb{R}$:

- (1.) $b(cx + y, z) = cb(x, z) + b(y, z)$ (linearity in the first slot)
- (2.) $b(x, cz + w) = cb(x, z) + b(x, w)$ (linearity in the second slot).

bilinear maps on $V \times V$

When $V_1 = V_2 = V$ we simply say that $b : V \times V \rightarrow \mathbb{R}$ is a **bilinear mapping on V** . The **set of all bilinear maps of V is denoted $T_0^2 V$** . You can show that $T_0^2 V$ forms a vector space under the usual point-wise defined operations of function addition and scalar multiplication³. Hopefully you are familiar with the example below.

Example 6.1.2. Define $b : \mathbb{R}^n \times \mathbb{R}^n \rightarrow \mathbb{R}$ by $b(x, y) = x \cdot y$ for all $x, y \in \mathbb{R}^n$. Linearity in each slot follows easily from properties of dot-products:

$$b(cx + y, z) = (cx + y) \cdot z = cx \cdot z + y \cdot z = cb(x, z) + b(y, z)$$

$$b(x, cy + z) = x \cdot (cy + z) = cx \cdot y + x \cdot z = cb(x, y) + b(x, z).$$

We can use matrix multiplication to generate a large class of examples with ease.

Example 6.1.3. Suppose $A \in \mathbb{R}^{n \times n}$ and define $b : \mathbb{R}^n \times \mathbb{R}^n \rightarrow \mathbb{R}$ by $b(x, y) = x^T A y$ for all $x, y \in \mathbb{R}^n$. Observe that, by properties of matrix multiplication,

$$b(cx + y, z) = (cx + y)^T A z = (cx^T + y^T) A z = cx^T A z + y^T A z = cb(x, z) + b(y, z)$$

²beyond this, there are also spinorial equations, I would like to give a talk on spinors some time, ask if interested

³sounds like homework

$$b(x, cy + z) = x^T A(cy + z) = cx^T Ay + x^T Az = cb(x, y) + b(x, z)$$

for all $x, y, z \in \mathbb{R}^n$ and $c \in \mathbb{R}$. It follows that b is bilinear on $\mathbb{R}^n$.

Suppose $b : V \times V \rightarrow \mathbb{R}$ is bilinear and suppose $\beta = \{e_1, e_2, \dots, e_n\}$ is a basis for V whereas $\beta^* = \{e^1, e^2, \dots, e^n\}$ is a basis of V^* with $e^j(e_i) = \delta_{ij}$

$$\begin{aligned} b(x, y) &= b\left(\sum_{i=1}^n x^i e_i, \sum_{j=1}^n y^j e_j\right) \\ &= \sum_{i,j=1}^n b(x^i e_i, y^j e_j) \\ &= \sum_{i,j=1}^n x^i y^j b(e_i, e_j) \\ &= \sum_{i,j=1}^n b(e_i, e_j) e^i(x) e^j(y) \end{aligned} \tag{6.1}$$

Therefore, if we define $b_{ij} = b(e_i, e_j)$ then we may compute $b(x, y) = \sum_{i,j=1}^n b_{ij} x^i y^j$. The calculation above also indicates that b is a linear combination of certain basic bilinear mappings. In particular, b can be written a linear combination of a tensor product of dual vectors on V .

Definition 6.1.4.

Suppose V is a vector space with dual space V^* . If $\alpha, \beta \in V^*$ then we define $\alpha \otimes \beta : V \times V \rightarrow \mathbb{R}$ by $(\alpha \otimes \beta)(x, y) = \alpha(x)\beta(y)$ for all $x, y \in V$.

Given the notation⁴ preceding this definition, we note $(e^i \otimes e^j)(x, y) = e^i(x)e^j(y)$ hence for all $x, y \in V$ we find:

$$b(x, y) = \sum_{i,j=1}^n b(e_i, e_j) (e^i \otimes e^j)(x, y) \quad \text{therefore,} \quad b = \sum_{i,j=1}^n b(e_i, e_j) e^i \otimes e^j$$

We find⁵ that $T_0^2 V = \text{span}\{e^i \otimes e^j\}_{i,j=1}^n$. Moreover, it can be argued⁶ that $\{e^i \otimes e^j\}_{i,j=1}^n$ is a linearly independent set, therefore $\{e^i \otimes e^j\}_{i,j=1}^n$ forms a basis for $T_0^2 V$. We can count there are n^2 vectors in $\{e^i \otimes e^j\}_{i,j=1}^n$ hence $\dim(T_0^2 V) = n^2$.

If $V = \mathbb{R}^n$ and if $\{e^i\}_{i=1}^n$ denotes the standard dual basis, then there is a standard notation for the set of coefficients found in the summation for b . In particular, we denote $B = [b]$ where $B_{ij} = b(e_i, e_j)$ hence, following Equation 6.1,

$$b(x, y) = \sum_{i,j=1}^n x^i y^j b(e_i, e_j) = \sum_{i=1}^n \sum_{j=1}^n x^i B_{ij} y^j = x^T B y$$

⁴perhaps you would rather write $(e^i \otimes e^j)(x, y)$ as $e^i \otimes e^j(x, y)$, that is also fine.

⁵with the help of your homework where you will show $\{e^i \otimes e^j\}_{i,j=1}^n \subseteq T_0^2 V$

⁶yes, again, in your homework

Definition 6.1.5.

Suppose $b : V \times V \rightarrow \mathbb{R}$ is a bilinear mapping then we say:

1. b is **symmetric** iff $b(x, y) = b(y, x)$ for all $x, y \in V$
2. b is **antisymmetric** iff $b(x, y) = -b(y, x)$ for all $x, y \in V$

Any bilinear mapping on V can be written as the sum of a symmetric and antisymmetric bilinear mapping, this claim follows easily from the calculation below:

$$b(x, y) = \underbrace{\frac{1}{2} \left(b(x, y) + b(y, x) \right)}_{\text{symmetric}} + \underbrace{\frac{1}{2} \left(b(x, y) - b(y, x) \right)}_{\text{antisymmetric}}.$$

We say S_{ij} is **symmetric** in i, j iff $S_{ij} = S_{ji}$ for all i, j . Likewise, we say A_{ij} is **antisymmetric** in i, j iff $A_{ij} = -A_{ji}$ for all i, j . If S is a symmetric bilinear mapping and A is an antisymmetric bilinear mapping then the components of S are symmetric and the components of A are antisymmetric. Why? Simply note:

$$S(e_i, e_j) = S(e_j, e_i) \Rightarrow S_{ij} = S_{ji}$$

and

$$A(e_i, e_j) = -A(e_j, e_i) \Rightarrow A_{ij} = -A_{ji}.$$

You can prove that the sum or scalar multiple of an (anti)symmetric bilinear mapping is once more (anti)symmetric therefore the set of antisymmetric bilinear maps $\Lambda^2(V)$ and the set of symmetric bilinear maps $ST_2^0 V$ are subspaces of $T_2^0 V$. The notation $\Lambda^2(V)$ is part of a larger discussion on the wedge product, we will return to it in a later section.

Finally, if we consider the special case of $V = \mathbb{R}^n$ once more we find that a bilinear mapping $b : \mathbb{R}^n \times \mathbb{R}^n \rightarrow \mathbb{R}$ has a symmetric matrix $[b]^T = [b]$ iff b is symmetric whereas it has an antisymmetric matrix $[b]^T = -[b]$ iff b is antisymmetric.

bilinear maps on $V^* \times V^*$

Suppose $h : V^* \times V^* \rightarrow \mathbb{R}$ is bilinear then we say $h \in T_0^2 V$. In addition, suppose $\beta = \{e_1, e_2, \dots, e_n\}$ is a basis for V whereas $\beta^* = \{e^1, e^2, \dots, e^n\}$ is a basis of V^* with $e^j(e_i) = \delta_{ij}$. Let $\alpha, \beta \in V^*$

$$\begin{aligned} h(\alpha, \beta) &= h\left(\sum_{i=1}^n \alpha_i e^i, \sum_{j=1}^n \beta_j e^j\right) \\ &= \sum_{i,j=1}^n h(\alpha_i e^i, \beta_j e^j) \\ &= \sum_{i,j=1}^n \alpha_i \beta_j h(e^i, e^j) \\ &= \sum_{i,j=1}^n h(e^i, e^j) \alpha(e_i) \beta(e_j) \end{aligned} \tag{6.2}$$

Therefore, if we define $h^{ij} = h(e^i, e^j)$ then we find the nice formula $h(\alpha, \beta) = \sum_{i,j=1}^n h^{ij} \alpha_i \beta_j$. To further refine the formula above we need a new concept.

The dual of the dual is called the double-dual and it is denoted V^{**} . For a finite dimensional vector space there is a canonical isomorphism of V and V^{**} . In particular, $\Phi : V \rightarrow V^{**}$ is defined by $\Phi(v)(\alpha) = \alpha(v)$ for all $\alpha \in V^*$. It is customary to replace V with V^{**} wherever the context allows. For example, to define the tensor product of two vectors $x, y \in V$ as follows:

Definition 6.1.6.

Suppose V is a vector space with dual space V^* . We define the tensor product of vectors x, y as the mapping $x \otimes y : V^* \times V^* \rightarrow \mathbb{R}$ by $(x \otimes y)(\alpha, \beta) = \alpha(x)\beta(y)$ for all $x, y \in V$.

We could just as well have defined $x \otimes y = \Phi(x) \otimes \Phi(y)$ where Φ is once more the canonical isomorphism of V and V^{**} . It's called *canonical* because it has no particular dependence on the coordinates used on V . In contrast, the isomorphism of $\mathbb{R}^n$ and $(\mathbb{R}^n)^*$ was built around the dot-product and the standard basis.

All of this said, note that $\alpha(e_i)\beta(e_j) = e_i \otimes e_j(\alpha, \beta)$ thus,

$$h(\alpha, \beta) = \sum_{i,j=1}^n h(e^i, e^j) e_i \otimes e_j(\alpha, \beta) \Rightarrow h = \sum_{i,j=1}^n h(e^i, e^j) e_i \otimes e_j$$

We argue that $\{e_i \otimes e_j\}_{i,j=1}^n$ is a basis⁷

Definition 6.1.7.

Suppose $h : V^* \times V^* \rightarrow \mathbb{R}$ is a bilinear mapping then we say:

1. h is **symmetric** iff $h(\alpha, \beta) = h(\beta, \alpha)$ for all $\alpha, \beta \in V^*$
2. h is **antisymmetric** iff $h(\alpha, \beta) = -h(\beta, \alpha)$ for all $\alpha, \beta \in V^*$

The discussion of the preceding subsection transfers to this context, we simply have to switch some vectors to dual vectors and move some indices up or down. I leave this to the reader.

bilinear maps on $V \times V^*$

Suppose $H : V \times V^* \rightarrow \mathbb{R}$ is bilinear, we say $H \in T_1^1 V$ (or, if the context demands this detail $H \in T_1^{-1} V$). We define $\alpha \otimes x \in T_1^{-1}(V)$ by the natural rule; $(\alpha \otimes x)(y, \beta) = \alpha(x)\beta(y)$ for all $(y, \beta) \in V \times V^*$. We find, by calculations similar to those already given in this section,

$$H(y, \beta) = \sum_{i,j=1}^n H_i^j y^i \beta_j \quad \text{and} \quad H = \sum_{i,j=1}^n H_i^j e^i \otimes e_j$$

where we defined $H_i^j = H(e_i, e^j)$.

⁷ $T_0^2 V$ is a vector space and we've shown $T_0^2(V) \subseteq \text{span}\{e_i \otimes e_j\}_{i,j=1}^n$ but we should also show $e_i \otimes e_j \in T_0^2$ and check for LI of $\{e_i \otimes e_j\}_{i,j=1}^n$.

bilinear maps on $V^* \times V$

Suppose $G : V^* \times V \rightarrow \mathbb{R}$ is bilinear, we say $G \in T_1^1 V$ (or, if the context demands this detail $G \in T_1^1 V$). We define $x \otimes \alpha \in T_1^1 V$ by the natural rule; $(x \otimes \alpha)(\beta, y) = \beta(x)\alpha(y)$ for all $(\beta, y) \in V^* \times V$. We find, by calculations similar to those already given in this section,

$$G(\beta, y) = \sum_{i,j=1}^n G^i_j \beta_i y^j \quad \text{and} \quad G = \sum_{i,j=1}^n G^i_j e_i \otimes e^j$$

where we defined $G^i_j = G(e^i, e_j)$.

6.1.2 trilinear maps

Definition 6.1.8.

Suppose V_1, V_2, V_3 are vector spaces then $T : V_1 \times V_2 \times V_3 \rightarrow \mathbb{R}$ is a **trilinear mapping** on $V_1 \times V_2 \times V_3$ iff for all $u, v \in V_1$, $w, x \in V_2$, $y, z \in V_3$ and $c \in \mathbb{R}$:

- (1.) $T(cu + v, w, y) = cT(u, w, y) + T(v, w, y)$ (linearity in the first slot)
- (2.) $T(u, cw + x, y) = cT(u, w, y) + T(u, x, y)$ (linearity in the second slot).
- (3.) $T(u, w, cy + z) = cT(u, w, y) + T(u, w, z)$ (linearity in the third slot).

If $T : V \times V \times V \rightarrow \mathbb{R}$ is trilinear on $V \times V \times V$ then we say T is a **trilinear mapping on V** and we denote the set of all such mappings $T_3^0 V$. The tensor product of three dual vectors is defined much in the same way as it was for two,

$$(\alpha \otimes \beta \otimes \gamma)(x, y, z) = \alpha(x)\beta(y)\gamma(z)$$

Let $\{e_i\}_{i=1}^n$ is a basis for V with dual basis $\{e^i\}_{i=1}^n$ for V^* . If T is trilinear on V it follows

$$T(x, y, z) = \sum_{i,j,k=1}^n T_{ijk} x^i y^j z^k \quad \text{and} \quad T = \sum_{i,j,k=1}^n T_{ijk} e^i \otimes e^j \otimes e^k$$

where we defined $T_{ijk} = T(e_i, e_j, e_k)$ for all $i, j, k \in \mathbb{N}_n$.

Generally suppose that V_1, V_2, V_3 are possibly distinct vector spaces. Moreover, suppose V_1 has basis $\{e_i\}_{i=1}^{n_1}$, V_2 has basis $\{f_j\}_{j=1}^{n_2}$ and V_3 has basis $\{g_k\}_{k=1}^{n_3}$. Denote the dual bases for V_1^*, V_2^*, V_3^* in the usual fashion: $\{e^i\}_{i=1}^{n_1}$, $\{f^j\}_{j=1}^{n_2}$, $\{g^k\}_{k=1}^{n_3}$. With this notation, we can write a trilinear mapping on $V_1 \times V_2 \times V_3$ as follows: (where we define $T_{ijk} = T(e_i, f_j, g_k)$)

$$T(x, y, z) = \sum_{i=1}^{n_1} \sum_{j=1}^{n_2} \sum_{k=1}^{n_3} T_{ijk} x^i y^j z^k \quad \text{and} \quad T = \sum_{i=1}^{n_1} \sum_{j=1}^{n_2} \sum_{k=1}^{n_3} T_{ijk} e^i \otimes f^j \otimes g^k$$

However, if V_1, V_2, V_3 happen to be related by duality then it is customary to use up/down indices. For example, if $T : V \times V \times V^* \rightarrow \mathbb{R}$ is trilinear then we write⁸

$$T = \sum_{i,j,k=1}^n T_{ij}{}^k e^i \otimes e^j \otimes e_k$$

⁸we identify e_k with its double-dual hence this tensor product is already defined, but to be safe let me write it out in this context $e^i \otimes e^j \otimes e_k(x, y, \alpha) = e^i(x)e^j(y)\alpha(e_k)$.

and say $T \in T_2^{-1}V$. On the other hand, if $S : V^* \times V^* \times V$ is trilinear then we'd write

$$T = \sum_{i,j,k=1}^n S_{ijk}^{ij} e_i \otimes e_j \otimes e^k$$

and say $T \in T_1^2V$. I'm not sure that I've ever seen this notation elsewhere, but perhaps it could be useful to denote the set of trilinear maps $T : V \times V^* \times V \rightarrow \mathbb{R}$ as $T_1^1 V$. Hopefully we will not need such silly notation in what we consider this semester.

There was a natural correspondence between bilinear maps on $\mathbb{R}^n$ and square matrices. For a trilinear map we would need a three-dimensional array of components. In some sense you could picture $T : \mathbb{R}^n \times \mathbb{R}^n \times \mathbb{R}^n \rightarrow \mathbb{R}$ as multiplication by a cube of numbers. Don't think too hard about these silly comments, we actually already wrote the useful formulae for dealing with trilinear objects. Let's stop to look at an example.

Example 6.1.9. Define $T : \mathbb{R}^3 \times \mathbb{R}^3 \times \mathbb{R}^3 \rightarrow \mathbb{R}$ by $T(x, y, z) = \det(x|y|z)$. You may not have learned this in your linear algebra course⁹ but a nice formula¹⁰ for the determinant is given by the Levi-Civita symbol,

$$\det(A) = \sum_{i,j,k=1}^3 \epsilon_{ijk} A_{i1} A_{j2} A_{k3}$$

note that $\text{col}_1(A) = [A_{i1}]$, $\text{col}_2(A) = [A_{i2}]$ and $\text{col}_3(A) = [A_{i3}]$. It follows that

$$T(x, y, z) = \sum_{i,j,k=1}^3 \epsilon_{ijk} x^i y^j z^k$$

Multilinearity follows easily from this formula. For example, linearity in the third slot:

$$\begin{aligned} T(x, y, cz + w) &= \sum_{i,j,k=1}^3 \epsilon_{ijk} x^i y^j (cz + w)^k \\ &= \sum_{i,j,k=1}^3 \epsilon_{ijk} x^i y^j (cz^k + w^k) \\ &= c \sum_{i,j,k=1}^3 \epsilon_{ijk} x^i y^j z^k + \sum_{i,j,k=1}^3 \epsilon_{ijk} x^i y^j w^k \\ &= cT(x, y, z) + T(x, y, w). \end{aligned}$$

Observe that by properties of determinants, or the Levi-Civita symbol if you prefer, swapping a pair of inputs generates a minus sign, hence:

$$T(x, y, z) = -T(y, x, z) = T(y, z, x) = -T(z, y, x) = T(z, x, y) = -T(x, z, y).$$

If $T : V \times V \times V \rightarrow \mathbb{R}$ is a trilinear mapping such that

$$T(x, y, z) = -T(y, x, z) = T(y, z, x) = -T(z, y, x) = T(z, x, y) = -T(x, z, y)$$

⁹maybe you haven't even taken linear yet!

¹⁰actually, I take this as the definition in linear algebra, it does take considerable effort to recover the expansion by minors formula which I use for concrete examples

for all $x, y, z \in V$ then we say T is **antisymmetric**. Likewise, if $S : V \times V \times V \rightarrow \mathbb{R}$ is a trilinear mapping such that

$$S(x, y, z) = -S(y, x, z) = S(y, z, x) = -S(z, y, x) = S(z, x, y) = -S(x, z, y).$$

for all $x, y, z \in V$ then we say T is **symmetric**. Clearly the mapping defined by the determinant is antisymmetric. In fact, many authors define the determinant of an $n \times n$ matrix as the antisymmetric n -linear mapping which sends the identity matrix to 1. It turns out these criteria uniquely define the determinant. That is the motivation behind my Levi-Civita symbol definition. That formula is just the nuts and bolts of complete antisymmetry.

You might wonder, can every trilinear mapping can be written as a the sum of a symmetric and antisymmetric mapping? The answer is no. Consider $T : V \times V \times V \rightarrow \mathbb{R}$ defined by $T = e^1 \otimes e^2 \otimes e^3$. Is it possible to find constants a, b such that:

$$e^1 \otimes e^2 \otimes e^3 = ae^{[1} \otimes e^2 \otimes e^3] + be^{(1} \otimes e^2 \otimes e^3)}$$

where $[\dots]$ denotes complete antisymmetrization of 1, 2, 3 and $(\dots)$ complete symmetrization:

$$e^{[1} \otimes e^2 \otimes e^3] = \frac{1}{6} [e^{123} + e^{231} + e^{312} - e^{321} - e^{213} - e^{132}]$$

For the symmetrization we also have to include all possible permutations of (1, 2, 3) but all with +:

$$e^{(1} \otimes e^2 \otimes e^3) = \frac{1}{6} [e^{123} + e^{231} + e^{312} + e^{321} + e^{213} + e^{132}]$$

As you can see:

$$ae^{[1} \otimes e^2 \otimes e^3] + be^{(1} \otimes e^2 \otimes e^3) = \frac{a+b}{6}(e^{123} + e^{231} + e^{312}) + \frac{b-a}{6}(e^{321} + e^{213} + e^{132})$$

There is no way for these to give back only $e^1 \otimes e^2 \otimes e^3$. I leave it to the reader to fill the gaps in this argument. Generally, the decomposition of a multilinear mapping into more basic types is a problem which requires much more thought than we intend here. Representation theory does address this problem: how can we decompose a tensor product into irreducible pieces. Their idea of tensor product is not precisely the same as ours, however algebraically the problems are quite intertwined. I'll leave it at that unless you'd like to do an independent study on representation theory. Ideally you'd already have linear algebra and abstract algebra complete before you attempt that study.

6.1.3 multilinear maps

Definition 6.1.10.

Suppose $V_1, V_2, \dots, V_k$ are vector spaces then $T : V_1 \times V_2 \times \dots \times V_k \rightarrow \mathbb{R}$ is a **k -multilinear mapping** on $V_1 \times V_2 \times \dots \times V_k$ iff for each $c \in \mathbb{R}$ and $x_1, y_1 \in V_1, x_2, y_2 \in V_2, \dots, x_k, y_k \in V_k$

$$T(x_1, \dots, cx_j + y_j, \dots, x_k) = cT(x_1, \dots, x_j, \dots, x_k) + T(x_1, \dots, y_j, \dots, x_k)$$

for $j = 1, 2, \dots, k$. In other words, we assume T is linear in each of its k -slots. If T is multilinear on $V^r \times (V^*)^s$ then we say that $T \in T_r^s V$ and we say T is a **type (r, s) tensor on V** .

The definition above makes a dual vector a type $(1, 0)$ tensor whereas a double dual of a vector a type $(0, 1)$ tensor, a bilinear mapping on V is a type $(2, 0)$ tensor and a bilinear mapping on V^* is a type $(0, 2)$ tensor with respect to V .

We are free to define tensor products in this context in the same manner as we have previously. Suppose $\alpha_1 \in V_1^*, \alpha_2 \in V_2^*, \dots, \alpha_k \in V_k^*$ and $v_1 \in V_1, v_2 \in V_2, \dots, v_k \in V_k$ then

$$\alpha_1 \otimes \alpha_2 \otimes \dots \otimes \alpha_k(v_1, v_2, \dots, v_k) = \alpha_1(v_1)\alpha_2(v_2) \cdots \alpha_k(v_k)$$

It is easy to show the tensor product of k -dual vectors as defined above is indeed a k -multilinear mapping. Moreover, the set of all k -multilinear mappings on $V_1 \times V_2 \times \dots \times V_k$ clearly forms a vector space of dimension $\dim(V_1)\dim(V_2) \cdots \dim(V_k)$ since it naturally takes the tensor product of the dual bases for $V_1^*, V_2^*, \dots, V_k^*$ as its basis. In particular, suppose for $j = 1, 2, \dots, k$ that V_j has basis $\{E_{ji}\}_{i=1}^{n_j}$ which is dual to $\{E_j^i\}_{i=1}^{n_j}$ the basis for V_j^* . Then we can derive that a k -multilinear mapping can be written as

$$T = \sum_{i_1=1}^{n_1} \sum_{i_2=1}^{n_2} \cdots \sum_{i_k=1}^{n_k} T_{i_1 i_2 \dots i_k} E_1^{i_1} \otimes E_2^{i_2} \otimes \cdots \otimes E_k^{i_k}$$

If T is a type (r, s) tensor on V then there is no need for the ugly double indexing on the basis since we need only tensor a basis $\{e_i\}_{i=1}^n$ for V and its dual $\{e^i\}_{i=1}^n$ for V^* in what follows:

$$T = \sum_{i_1, \dots, i_r=1}^n \sum_{j_1, \dots, j_s=1}^n T_{i_1 i_2 \dots i_r j_1 j_2 \dots j_s} e^{i_1} \otimes e^{i_2} \otimes \cdots \otimes e^{i_r} \otimes e_{j_1} \otimes e_{j_2} \otimes \cdots \otimes e_{j_s}.$$

permutations

Before I define symmetric and antisymmetric for k -linear mappings on V I think it is best to discuss briefly some ideas from the theory of permutations.

Definition 6.1.11.

A permutation on $\{1, 2, \dots, p\}$ is a bijection on $\{1, 2, \dots, p\}$. We define the set of permutations on $\{1, 2, \dots, p\}$ to be Σ_p . Further, define the sign of a permutation to be $\text{sgn}(\sigma) = 1$ if σ is the product of an even number of transpositions whereas $\text{sgn}(\sigma) = -1$ if σ is the product of an odd number of transpositions.

Let us consider the set of permutations on $\{1, 2, 3, \dots, n\}$, this is called S_n the symmetric group, its order is $n!$ if you were wondering. Let me remind¹¹ you how the cycle notation works since it allows us to explicitly present the number of transpositions contained in a permutation,

$$\sigma = \begin{pmatrix} 1 & 2 & 3 & 4 & 5 & 6 \\ 2 & 1 & 5 & 4 & 6 & 3 \end{pmatrix} \iff \sigma = (12)(356) = (12)(36)(35) \quad (6.3)$$

recall the cycle notation is to be read right to left. If we think about inputting 5 we can read from the matrix notation that we ought to find $5 \mapsto 6$. Clearly that is the case for the first version of σ written in cycle notation; (356) indicates that $5 \mapsto 6$ and nothing else messes with 6 after that. Then consider feeding 5 into the version of σ written with just two-cycles (a.k.a. transpositions), first we note (35) indicates $5 \mapsto 3$, then that 3 hits (36) which means $3 \mapsto 6$, finally the cycle (12)

¹¹or perhaps, more likely, introduce you to this notation

doesn't care about 6 so we again have that $\sigma(5) = 6$. Finally we note that $\text{sgn}(\sigma) = -1$ since it is made of 3 transpositions.

It is always possible to write any permutation as a product of transpositions, such a decomposition is not unique. However, if the number of transpositions is even then it will remain so no matter how we rewrite the permutation. Likewise if the permutation is an product of an odd number of transpositions then any other decomposition into transpositions is also comprised of an odd number of transpositions. This is why we can **define** an **even** permutation is a permutation comprised by an even number of transpositions and an **odd** permutation is one comprised of an odd number of transpositions.

Example 6.1.12. Sample cycle calculations: *we rewrite as product of transpositions to determine if the given permutation is even or odd,*

$$\sigma = (12)(134)(152) = (12)(14)(13)(12)(15) \implies \text{sgn}(\sigma) = -1$$

$$\lambda = (1243)(3521) = (13)(14)(12)(31)(32)(35) \implies \text{sgn}(\lambda) = 1$$

$$\gamma = (123)(45678) = (13)(12)(48)(47)(46)(45) \implies \text{sgn}(\gamma) = 1$$

We will not actually write down permutations in the calculations the follow this part of the notes. I merely include this material as to give a logically complete account of antisymmetry. In practice, if you understood the terms as they apply to the bilinear and trilinear case it will usually suffice for concrete examples. Now we are ready to define symmetric and antisymmetric.

Definition 6.1.13.

A k -linear mapping $L : V \times V \times \cdots \times V \rightarrow \mathbb{R}$ is **completely symmetric** if

$$L(x_1, \dots, x, \dots, y, \dots, x_k) = L(x_1, \dots, y, \dots, x, \dots, x_k)$$

for all possible $x, y \in V$. Conversely, if a k -linear mapping on V has

$$L(x_1, \dots, x, \dots, y, \dots, x_p) = -L(x_1, \dots, y, \dots, x, \dots, x_p)$$

for all possible pairs $x, y \in V$ then it is said to be **completely antisymmetric or alternating**. Equivalently a k -linear mapping L is alternating if for all $\pi \in \Sigma_k$

$$L(x_{\pi_1}, x_{\pi_2}, \dots, x_{\pi_k}) = \text{sgn}(\pi)L(x_1, x_2, \dots, x_k)$$

The set of alternating multilinear mappings on V is denoted ΛV , the set of k -linear alternating maps on V is denoted $\Lambda^k V$. Often an alternating k -linear map is called a **k -form**. Moreover, we say the **degree** of a k -form is k .

Similar terminology applies to the components of tensors. We say $T_{i_1 i_2 \dots i_k}$ is completely symmetric in $i_1, i_2, \dots, i_k$ iff $T_{i_1 i_2 \dots i_k} = T_{i_{\sigma(1)} i_{\sigma(2)} \dots i_{\sigma(k)}}$ for all $\sigma \in \Sigma_k$. On the other hand, $T_{i_1 i_2 \dots i_k}$ is completely antisymmetric in $i_1, i_2, \dots, i_k$ iff $T_{i_1 i_2 \dots i_k} = \text{sgn}(\sigma) T_{i_{\sigma(1)} i_{\sigma(2)} \dots i_{\sigma(k)}}$ for all $\sigma \in \Sigma_k$. It is a simple exercise to show that a completely (anti)symmetric tensor¹² has completely (anti)symmetric components.

¹²in this context a tensor is simply a multilinear mapping, in physics there is more attached to the term

The tensor product is an interesting construction to discuss at length. To summarize, it is associative and distributive across addition. Scalars factor out and it is not generally commutative. For a given vector space V we can in principle generate by tensor products multilinear mappings of arbitrarily high order. This tensor algebra is infinite dimensional. In contrast, the space ΛV of forms on V is a finite-dimensional subspace of the tensor algebra. We discuss this next.

6.2 wedge product

We assume V is a vector space with basis $\{e_i\}_{i=1}^n$ throughout this section. The dual basis is denoted $\{e^i\}_{i=1}^n$ as is our usual custom. Our goal is to find a basis for the alternating maps on V and explore the structure implicit within its construction. This will lead us to call ΩV the **exterior algebra** of V after the discussion below is complete. I should mention, the approach here is inelegant and concrete. There are far more efficient, slick, and incomprehensible ways to build ΩV using abstract algebra. I encourage the reader to seek those out once they know more abstract algebra. To give a sketch, basically any algebraic object you wish to construct can be built from forming an appropriate quotient where you divide by whatever you wish to treat as zero. We examine this concept of $\otimes$ or $\wedge$ defined via a formal quotient in Math 422 if anywhere.

6.2.1 wedge product of dual basis generates basis for ΛV

Suppose $b : V \times V \rightarrow \mathbb{R}$ is antisymmetric and $b = \sum_{i,j=1}^n b_{ij} e^i \otimes e^j$, it follows that $b_{ij} = -b_{ji}$ for all $i, j \in \mathbb{N}_n$. Notice this implies that $b_{ii} = 0$ for $i = 1, 2, \dots, n$. For a given pair of indices i, j either $i < j$ or $j < i$ or $i = j$ hence,

$$\begin{aligned}
 b &= \sum_{i < j} b_{ij} e^i \otimes e^j + \sum_{j < i} b_{ij} e^i \otimes e^j + \sum_{i=j} b_{ij} e^i \otimes e^j \\
 &= \sum_{i < j} b_{ij} e^i \otimes e^j + \sum_{j < i} b_{ij} e^i \otimes e^j \\
 &= \sum_{i < j} b_{ij} e^i \otimes e^j - \sum_{j < i} b_{ji} e^i \otimes e^j \\
 &= \sum_{k < l} b_{kl} e^k \otimes e^l - \sum_{k < l} b_{kl} e^l \otimes e^k \\
 &= \sum_{k < l} b_{kl} (e^k \otimes e^l - e^l \otimes e^k).
 \end{aligned} \tag{6.4}$$

Therefore, $\{e^k \otimes e^l - e^l \otimes e^k | l, k \in \mathbb{N}_n \text{ and } l < k\}$ spans the set of antisymmetric bilinear maps on V . Moreover, you can show this set is linearly independent hence it is a basis for $\Lambda^2 V$. We define the wedge product of $e^k \wedge e^l = e^k \otimes e^l - e^l \otimes e^k$. With this notation we find that the alternating bilinear form b can be written as

$$b = \sum_{k < l} b_{kl} e^k \wedge e^l = \sum_{i,j=1}^n \frac{1}{2} b_{ij} e^i \wedge e^j$$

where the summation on the r.h.s. is over all indices¹³. Notice that $e^i \wedge e^j$ is an antisymmetric bilinear mapping because $e^i \wedge e^j(x, y) = -e^i \wedge e^j(y, x)$, however, there is more structure here than

¹³yes there is something to work out here, probably in your homework

just that. It is also true that $e^i \wedge e^j = -e^j \wedge e^i$. This is a conceptually different antisymmetry, it is the antisymmetry of the wedge product $\wedge$.

Suppose $b : V \times V \times V \rightarrow \mathbb{R}$ is antisymmetric and $b = \sum_{i,j,k=1}^n b_{ijk} e^i \otimes e^j \otimes e^k$, it follows that $b_{ijk} = b_{jki} = b_{kij}$ and $b_{ijk} = -b_{kji} = -b_{jik} = b_{ikj}$ for all $i, j, k \in \mathbb{N}_n$. Notice this implies that $b_{iii} = 0$ for $i = 1, 2, \dots, n$. A calculation similar to the one just offered for the case of a bilinear map reveals that we can write b as follows:

$$b = \sum_{i < j < k} b_{ijk} \left(e^i \otimes e^j \otimes e^k + e^j \otimes e^k \otimes e^i + e^k \otimes e^i \otimes e^j - e^k \otimes e^j \otimes e^i - e^j \otimes e^i \otimes e^k - e^i \otimes e^k \otimes e^j \right) \quad (6.5)$$

Define $e^i \wedge e^j \wedge e^k = e^i \otimes e^j \otimes e^k + e^j \otimes e^k \otimes e^i + e^k \otimes e^i \otimes e^j - e^k \otimes e^j \otimes e^i - e^j \otimes e^i \otimes e^k - e^i \otimes e^k \otimes e^j$ thus

$$b = \sum_{i < j < k} b_{ijk} e^i \wedge e^j \wedge e^k = \sum_{i,j,k=1}^n \frac{1}{3!} b_{ijk} e^i \wedge e^j \wedge e^k \quad (6.6)$$

and it is clear that $\{e^i \wedge e^j \wedge e^k \mid i, j, k \in \mathbb{N}_n \text{ and } i < j < k\}$ forms a basis for the set of alternating trilinear maps on V .

Following the patterns above, we define the wedge product of p dual basis vectors,

$$e^{i_1} \wedge e^{i_2} \wedge \dots \wedge e^{i_p} = \sum_{\pi \in \Sigma_p} \text{sgn}(\pi) e^{i_{\pi(1)}} \otimes e^{i_{\pi(2)}} \otimes \dots \otimes e^{i_{\pi(p)}} \quad (6.7)$$

If $x, y \in V$ we would like to show that

$$e^{i_1} \wedge e^{i_2} \wedge \dots \wedge e^{i_p}(\dots, x, \dots, y, \dots) = -e^{i_1} \wedge e^{i_2} \wedge \dots \wedge e^{i_p}(\dots, y, \dots, x, \dots) \quad (6.8)$$

follows from the complete antisymmetrization in the definition of the wedge product. Before we give the general argument, let's see how this works in the trilinear case. Consider, $e^i \wedge e^j \wedge e^k =$

$$= e^i \otimes e^j \otimes e^k + e^j \otimes e^k \otimes e^i + e^k \otimes e^i \otimes e^j - e^k \otimes e^j \otimes e^i - e^j \otimes e^i \otimes e^k - e^i \otimes e^k \otimes e^j.$$

Calculate, noting that $e^i \otimes e^j \otimes e^k(x, y, z) = e^i(x)e^j(y)e^k(z) = x^i y^j z^k$ hence

$$e^i \wedge e^j \wedge e^k(x, y, z) = x^i y^j z^k + x^j y^k z^i + x^k y^i z^j - x^k y^j z^i - x^j y^i z^k - x^i y^k z^j$$

Thus,

$$e^i \wedge e^j \wedge e^k(x, z, y) = x^i z^j y^k + x^j z^k y^i + x^k z^i y^j - x^k z^j y^i - x^j z^i y^k - x^i z^k y^j$$

and you can check that $e^i \wedge e^j \wedge e^k(x, y, z) = -e^i \wedge e^j \wedge e^k(x, z, y)$. Similar tedious calculations prove antisymmetry of the the interchange of the first and second or the first and third slots. Therefore, $e^i \wedge e^j \wedge e^k$ is an alternating trilinear map as it is clearly trilinear since it is built from the sum of tensor products which we know are likewise trilinear.

The multilinear case follows essentially the same argument, note

$$e^{i_1} \wedge e^{i_2} \wedge \dots \wedge e^{i_p}(\dots, x_j, \dots, x_k, \dots) = \sum_{\pi \in \Sigma_p} \text{sgn}(\pi) x_1^{i_{\pi(1)}} \dots x_j^{i_{\pi(j)}} \dots x_k^{i_{\pi(k)}} \dots x_p^{i_{\pi(p)}} \quad (6.9)$$

whereas,

$$e^{i_1} \wedge e^{i_2} \wedge \cdots \wedge e^{i_p}(\dots, x_k, \dots, x_j, \dots) = \sum_{\sigma \in \Sigma_p} \text{sgn}(\sigma) x_1^{i_{\sigma(1)}} \cdots x_k^{i_{\sigma(k)}} \cdots x_j^{i_{\sigma(j)}} \cdots x_p^{i_{\sigma(p)}}. \quad (6.10)$$

Suppose we take each permutation σ and substitute $\delta \in \Sigma_p$ such that $\sigma(j) = \delta(k)$ and $\sigma(k) = \delta(j)$ and otherwise δ and σ agree. In cycle notation, $\delta(jk) = \sigma$. Substitution δ into Equation 6.10:

$$\begin{aligned} e^{i_1} \wedge e^{i_2} \wedge \cdots \wedge e^{i_p}(\dots, x_k, \dots, x_j, \dots) &= \sum_{\delta \in \Sigma_p} \text{sgn}(\delta(jk)) x_1^{i_{\delta(1)}} \cdots x_k^{i_{\delta(j)}} \cdots x_j^{i_{\delta(k)}} \cdots x_p^{i_{\delta(p)}} \\ &= - \sum_{\delta \in \Sigma_p} \text{sgn}(\delta) x_1^{i_{\delta(1)}} \cdots x_j^{i_{\delta(k)}} \cdots x_k^{i_{\delta(j)}} \cdots x_p^{i_{\delta(p)}} \\ &= -e^{i_1} \wedge e^{i_2} \wedge \cdots \wedge e^{i_p}(\dots, x_j, \dots, x_k, \dots) \end{aligned} \quad (6.11)$$

Here the sgn of a permutation σ is $(-1)^N$ where N is the number of cycles in σ . We observed that $\delta(jk)$ has one more cycle than δ hence $\text{sgn}(\delta(jk)) = -\text{sgn}(\delta)$. Therefore, we have shown that $e^{i_1} \wedge e^{i_2} \wedge \cdots \wedge e^{i_p} \in \Lambda^p V$.

Recall that $e^i \wedge e^j = -e^j \wedge e^i$ in the $p = 2$ case. There is a generalization of that result to the $p > 2$ case. In words, the wedge product is antisymmetric with respect the interchange of any two dual vectors. For $p = 3$ we have the following identities for the wedge product:

$$e^i \wedge e^j \wedge e^k = - \underbrace{e^j \wedge e^i}_{\text{swapped}} \wedge e^k = e^j \wedge \underbrace{e^k \wedge e^i}_{\text{swapped}} = - \underbrace{e^k \wedge e^j}_{\text{swapped}} \wedge e^i = e^k \wedge \underbrace{e^i \wedge e^j}_{\text{swapped}} = - \underbrace{e^i \wedge e^k}_{\text{swapped}} \wedge e^j$$

I've indicated how these signs are consistent with the $p = 2$ antisymmetry. Any permutation of the dual vectors can be thought of as a combination of several transpositions. In any event, it is sometimes useful to just know that the wedge product of three elements is invariant under **cyclic** permutations of the dual vectors,

$$e^i \wedge e^j \wedge e^k = e^j \wedge e^k \wedge e^i = e^k \wedge e^i \wedge e^j$$

and changes by a sign for **anticyclic** permutations of the given object,

$$e^i \wedge e^j \wedge e^k = -e^j \wedge e^i \wedge e^k = -e^k \wedge e^j \wedge e^i = -e^i \wedge e^k \wedge e^j$$

Generally we can argue that, for any permutation $\pi \in \Sigma_p$:

$$\boxed{e^{i_1} \wedge e^{i_2} \wedge \cdots \wedge e^{i_p} = \text{sgn}(\pi) e^{i_{\pi(1)}} \wedge e^{i_{\pi(2)}} \wedge \cdots \wedge e^{i_{\pi(p)}}}$$

This is just a slick formula which says the wedge product generates a minus whenever you flip two dual vectors which are wedged.

6.2.2 the exterior algebra

The careful reader will realize we have yet to define wedge products of anything except for the dual basis. But, naturally you must wonder if we can take the wedge product of other dual vectors or morer generally alternating tensors. The answer is yes. Let us define the general wedge product:

Definition 6.2.1. Suppose $\alpha \in \Lambda^p V$ and $\beta \in \Lambda^q V$. We define $\mathcal{I}_p$ to be the set of all increasing lists of p -indices, this set can be empty if $\dim(V)$ is not sufficiently large. Moreover, if $I = (i_1, i_2, \dots, i_p)$ then introduce notation $e^I = e^{i_1} \wedge e^{i_2} \wedge \dots \wedge e^{i_p}$ hence:

$$\alpha = \sum_{i_1, i_2, \dots, i_p=1}^n \frac{1}{p!} \alpha_{i_1 i_2 \dots i_p} e^{i_1} \wedge e^{i_2} \wedge \dots \wedge e^{i_p} = \sum_I \frac{1}{p!} \alpha_I e^I = \sum_{I \in \mathcal{I}_p} \alpha_I e^I$$

and

$$\beta = \sum_{j_1, j_2, \dots, j_q=1}^n \frac{1}{q!} \beta_{j_1 j_2 \dots j_q} e^{j_1} \wedge e^{j_2} \wedge \dots \wedge e^{j_q} = \sum_J \frac{1}{q!} \beta_J e^J = \sum_{J \in \mathcal{I}_q} \beta_J e^J$$

Naturally, $e^I \wedge e^J = e^{i_1} \wedge e^{i_2} \wedge \dots \wedge e^{i_p} \wedge e^{j_1} \wedge e^{j_2} \wedge \dots \wedge e^{j_q}$ and we defined this carefully in the preceding subsection. Define $\alpha \wedge \beta \in \Lambda^{p+q} V$ as follows:

$$\alpha \wedge \beta = \sum_I \sum_J \frac{1}{p!q!} \alpha_I \beta_J e^I \wedge e^J.$$

Again, but with less slick notation:

$$\alpha \wedge \beta = \sum_{i_1, i_2, \dots, i_p=1}^n \sum_{j_1, j_2, \dots, j_q=1}^n \frac{1}{p!q!} \alpha_{i_1 i_2 \dots i_p} \beta_{j_1 j_2 \dots j_q} e^{i_1} \wedge e^{i_2} \wedge \dots \wedge e^{i_p} \wedge e^{j_1} \wedge e^{j_2} \wedge \dots \wedge e^{j_q}$$

All the definition above really says is that we extend the wedge product on the basis to distribute over the addition of dual vectors. What this means calculationally is that the wedge product obeys the usual laws of addition and scalar multiplication. The one feature that is perhaps foreign is the antisymmetry of the wedge product. We must take care to maintain the order of expressions since the wedge product is not generally commutative.

Proposition 6.2.2.

Let α, β, γ be forms on V and $c \in \mathbb{R}$ then

- | | | |
|-------|---|------------------------------------|
| (i) | $(\alpha + \beta) \wedge \gamma = \alpha \wedge \gamma + \beta \wedge \gamma$ | distributes across vector addition |
| (ii) | $\alpha \wedge (\beta + \gamma) = \alpha \wedge \beta + \alpha \wedge \gamma$ | distributes across vector addition |
| (iii) | $(c\alpha) \wedge \beta = \alpha \wedge (c\beta) = c(\alpha \wedge \beta)$ | scalars factor out |
| (iv) | $\alpha \wedge (\beta \wedge \gamma) = (\alpha \wedge \beta) \wedge \gamma$ | associativity |

I leave the proof of this proposition to the reader.

Proposition 6.2.3. *graded commutivity of homogeneous forms.*

Let α, β be forms on V of degree p and q respectively then

$$\alpha \wedge \beta = -(-1)^{pq} \beta \wedge \alpha$$

Proof: suppose $\alpha = \sum_I \frac{1}{p!} e^I$ is a p -form on V and $\beta = \sum_J \frac{1}{q!} e^J$ is a q -form on V . Calculate:

$$\begin{aligned}
 \alpha \wedge \beta &= \sum_I \sum_J \frac{1}{p!q!} \alpha_I \beta_J e^I \wedge e^J && \text{by defn. of } \wedge, \\
 &= \sum_I \sum_J \frac{1}{p!q!} \beta_J \alpha_I e^I \wedge e^J && \text{coefficients are scalars,} \\
 &= (-1)^{pq} \sum_I \sum_J \frac{1}{p!q!} \beta_J \alpha_I e^J \wedge e^I && \text{(details on sign given below)} \\
 &= (-1)^{pq} \beta \wedge \alpha
 \end{aligned}$$

Let's expand in detail why $e^J \wedge e^I = (-1)^{pq} e^I \wedge e^J$. Suppose $I = (i_1, i_2, \dots, i_p)$ and $J = (j_1, j_2, \dots, j_q)$, the key is that every swap of dual vectors generates a sign:

$$\begin{aligned}
 e^I \wedge e^J &= e^{i_1} \wedge e^{i_2} \wedge \dots \wedge e^{i_p} \wedge e^{j_1} \wedge e^{j_2} \wedge \dots \wedge e^{j_q} \\
 &= (-1)^q e^{i_1} \wedge e^{i_2} \wedge \dots \wedge e^{i_{p-1}} \wedge e^{j_1} \wedge e^{j_2} \wedge \dots \wedge e^{j_q} \wedge e^{i_p} \\
 &= (-1)^q (-1)^q e^{i_1} \wedge e^{i_2} \wedge \dots \wedge e^{i_{p-2}} \wedge e^{j_1} \wedge e^{j_2} \wedge \dots \wedge e^{j_q} \wedge e^{i_{p-1}} \wedge e^{i_p} \\
 &\quad \vdots \quad \quad \quad \vdots \quad \quad \quad \vdots \\
 &= \underbrace{(-1)^q (-1)^q \dots (-1)^q}_{p\text{-factors}} e^{j_1} \wedge e^{j_2} \wedge \dots \wedge e^{j_q} \wedge e^{i_1} \wedge e^{i_2} \wedge \dots \wedge e^{i_p} \\
 &= (-1)^{pq} e^J \wedge e^I.
 \end{aligned}$$

□

Example 6.2.4. Let α be a 2-form defined by

$$\alpha = ae^1 \wedge e^2 + be^2 \wedge e^3$$

And let β be a 1-form defined by

$$\beta = 3e^1$$

Consider then,

$$\begin{aligned}
 \alpha \wedge \beta &= (ae^1 \wedge e^2 + be^2 \wedge e^3) \wedge (3e^1) \\
 &= (3ae^1 \wedge e^2 \wedge e^1 + 3be^2 \wedge e^3 \wedge e^1) \\
 &= 3be^1 \wedge e^2 \wedge e^3.
 \end{aligned} \tag{6.12}$$

whereas,

$$\begin{aligned}
 \beta \wedge \alpha &= 3e^1 \wedge (ae^1 \wedge e^2 + be^2 \wedge e^3) \\
 &= (3ae^1 \wedge e^1 \wedge e^2 + 3be^1 \wedge e^2 \wedge e^3) \\
 &= 3be^1 \wedge e^2 \wedge e^3.
 \end{aligned} \tag{6.13}$$

so this agrees with the proposition, $(-1)^{pq} = (-1)^2 = 1$ so we should have found that $\alpha \wedge \beta = \beta \wedge \alpha$. This illustrates that although the wedge product is antisymmetric on the basis, it is not always antisymmetric, in particular it is commutative for even forms.

The graded commutivity rule $\alpha \wedge \beta = -(-1)^{pq} \beta \wedge \alpha$ has some suprising implications. This rule is ultimately the reason ΛV is finite dimensional. Let's see how that happens.

Proposition 6.2.5. *linear dependent one-forms wedge to zero:*

If $\alpha, \beta \in V^*$ and $\alpha = c\beta$ for some $c \in \mathbb{R}$ then $\alpha \wedge \beta = 0$.

Proof: to begin, note that $\beta \wedge \beta = -\beta \wedge \beta$ hence $2\beta \wedge \beta = 0$ and it follows that $\beta \wedge \beta = 0$. Note:

$$\alpha \wedge \beta = c\beta \wedge \beta = c(0) = 0$$

therefore the proposition is true. $\square$

Proposition 6.2.6.

Suppose that $\alpha_1, \alpha_2, \dots, \alpha_p$ are linearly dependent 1-forms then

$$\alpha_1 \wedge \alpha_2 \wedge \dots \wedge \alpha_p = 0.$$

Proof: by assumption of linear dependence there exist constants $c_1, c_2, \dots, c_p$ not all zero such that

$$c_1\alpha_1 + c_2\alpha_2 + \dots + c_p\alpha_p = 0.$$

Suppose that c_k is a nonzero constant in the sum above, then we may divide by it and consequently we can write α_k in terms of all the other 1-forms,

$$\alpha_k = \frac{-1}{c_k} \left(c_1\alpha_1 + \dots + c_{k-1}\alpha_{k-1} + c_{k+1}\alpha_{k+1} + \dots + c_p\alpha_p \right)$$

Insert this sum into the wedge product in question,

$$\begin{aligned} \alpha_1 \wedge \alpha_2 \wedge \dots \wedge \alpha_p &= \alpha_1 \wedge \alpha_2 \wedge \dots \wedge \alpha_k \wedge \dots \wedge \alpha_p \\ &= (-c_1/c_k)\alpha_1 \wedge \alpha_2 \wedge \dots \wedge \alpha_1 \wedge \dots \wedge \alpha_p \\ &\quad + (-c_2/c_k)\alpha_1 \wedge \alpha_2 \wedge \dots \wedge \alpha_2 \wedge \dots \wedge \alpha_p + \dots \\ &\quad + (-c_{k-1}/c_k)\alpha_1 \wedge \alpha_2 \wedge \dots \wedge \alpha_{k-1} \wedge \dots \wedge \alpha_p \\ &\quad + (-c_{k+1}/c_k)\alpha_1 \wedge \alpha_2 \wedge \dots \wedge \alpha_{k+1} \wedge \dots \wedge \alpha_p + \dots \\ &\quad + (-c_p/c_k)\alpha_1 \wedge \alpha_2 \wedge \dots \wedge \alpha_p \wedge \dots \wedge \alpha_p \\ &= 0. \end{aligned} \tag{6.14}$$

We know all the wedge products are zero in the above because in each there is at least one 1-form repeated, we simply permute the wedge products till they are adjacent and by the previous proposition the term vanishes. The proposition follows. $\square$

Let us pause to reflect on the meaning of the proposition above for a n -dimensional vector space V . The dual space V^* is likewise n -dimensional, this is a general result which applies to all finite-dimensional vector spaces¹⁴. Thus, any set of more than n dual vectors is necessarily linearly dependent. Consequently, using the proposition above, we find the wedge product of more than n one-forms is trivial. Therefore, while it is possible to construct $\Lambda^k V$ for $k > n$ we should understand that this space only contains zero. The highest degree of a nontrivial form over a vector space of dimension n is an n -form.

¹⁴however, in infinite dimensions, the story is not so simple

Moreover, we can use the proposition to deduce the dimension of a basis for $\Lambda^p V$, it must consist of the wedge product of distinct linearly independent one-forms. The number of ways to choose p distinct objects from a list of n distinct objects is precisely "n choose p",

$$\binom{n}{p} = \frac{n!}{(n-p)!p!} \quad \text{for } 0 \leq p \leq n. \quad (6.15)$$

Proposition 6.2.7.

If V is an n -dimensional vector space then $\Lambda^k V$ is an $\binom{n}{k}$ -dimensional vector space of k -forms. Moreover, the direct sum of all forms over V has the structure

$$\Omega V = \mathbb{R} \oplus \Lambda^1 V \oplus \cdots \Lambda^{n-1} V \oplus \Lambda^n V$$

and is a vector space of dimension 2^n

Proof: define $\Lambda^0 V = \mathbb{R}$ then it is clear $\Lambda^k V$ forms a vector space for $k = 0, 1, \dots, n$. Moreover, $\Lambda^j V \cap \Lambda^k V = \{0\}$ for $j \neq k$ hence the term "direct sum" is appropriate. It remains to show $\dim(\Lambda V) = 2^n$ where $\dim(V) = n$. A natural basis β for ΩV is found from taking the union of the bases for each subspace of k -forms,

$$\beta = \{1, e^{i_1}, e^{i_1} \wedge e^{i_2}, \dots, e^{i_1} \wedge e^{i_2} \wedge \cdots \wedge e^{i_n} \mid 1 \leq i_1 < i_2 < \cdots < i_n \leq n\}$$

But, we can count the number of vectors N in the set above as follows:

$$N = 1 + n + \binom{n}{2} + \cdots + \binom{n}{n-1} + \binom{n}{n}$$

Recall the binomial theorem states

$$(a+b)^n = \sum_{k=0}^n \binom{n}{k} a^{n-k} b^k = a^n + na^{n-1}b + \cdots + nab^{n-1} + b^n.$$

Recognize that $N = (1+1)^n$ and the proposition follows. $\square$

We should note that in the basis above the space of n -forms is one-dimensional because there is only one way to choose a strictly increasing list of n integers in $\mathbb{N}_n$. In particular, it is useful to note $\Lambda^n V = \text{span}\{e^1 \wedge e^2 \wedge \cdots \wedge e^n\}$. The form $e^1 \wedge e^2 \wedge \cdots \wedge e^n$ is sometimes called the *top-form*¹⁵.

Example 6.2.8. exterior algebra of $\mathbb{R}^2$ Let us begin with the standard dual basis $\{e^1, e^2\}$. By definition we take the $p = 0$ case to be the field itself; $\Lambda^0 V \equiv \mathbb{R}$, it has basis 1. Next, $\Lambda^1 V = \text{span}(e^1, e^2) = V^*$ and $\Lambda^2 V = \text{span}(e^1 \wedge e^2)$ is all we can do here. This makes ΛV a $2^2 = 4$ -dimensional vector space with basis

$$\{1, e^1, e^2, e^1 \wedge e^2\}.$$

Example 6.2.9. exterior algebra of $\mathbb{R}^3$ Let us begin with the standard dual basis $\{e^1, e^2, e^3\}$. By definition we take the $p = 0$ case to be the field itself; $\Lambda^0 V \equiv \mathbb{R}$, it has basis 1. Next, $\Lambda^1 V = \text{span}(e^1, e^2, e^3) = V^*$. Now for something a little more interesting,

$$\Lambda^2 V = \text{span}(e^1 \wedge e^2, e^1 \wedge e^3, e^2 \wedge e^3)$$

¹⁵or *volume form* for reasons we will explain later, other authors begin the discussion of forms from the consideration of volume, see Chapter 4 in Bernard Schutz' *Geometrical methods of mathematical physics*

and finally,

$$\Lambda^3 V = \text{span}(e^1 \wedge e^2 \wedge e^3).$$

This makes ΛV a $2^3 = 8$ -dimensional vector space with basis

$$\{1, e^1, e^2, e^3, e^1 \wedge e^2, e^1 \wedge e^3, e^2 \wedge e^3, e^1 \wedge e^2 \wedge e^3\}$$

it is curious that the number of independent one-forms and 2-forms are equal.

Example 6.2.10. exterior algebra of $\mathbb{R}^4$ Let us begin with the standard dual basis $\{e^1, e^2, e^3, e^4\}$. By definition we take the $p = 0$ case to be the field itself; $\Lambda^0 V \equiv \mathbb{R}$, it has basis 1. Next, $\Lambda^1 V = \text{span}(e^1, e^2, e^3, e^4) = V^*$. Now for something a little more interesting,

$$\Lambda^2 V = \text{span}(e^1 \wedge e^2, e^1 \wedge e^3, e^1 \wedge e^4, e^2 \wedge e^3, e^2 \wedge e^4, e^3 \wedge e^4)$$

and three forms,

$$\Lambda^3 V = \text{span}(e^1 \wedge e^2 \wedge e^3, e^1 \wedge e^2 \wedge e^4, e^1 \wedge e^3 \wedge e^4, e^2 \wedge e^3 \wedge e^4).$$

and $\Lambda^4 V = \text{span}(e^1 \wedge e^2 \wedge e^3 \wedge e^4)$. Thus ΛV a $2^4 = 16$ -dimensional vector space. Note that, in contrast to $\mathbb{R}^3$, we do not have the same number of independent one-forms and two-forms over $\mathbb{R}^4$.

Let's explore how this algebra fits with calculations we already know about determinants.

Example 6.2.11. Suppose $A = [A_1 | A_2]$. I propose the determinant of A is given by the top-form on $\mathbb{R}^2$ via the formula $\det(A) = (e^1 \wedge e^2)(A_1, A_2)$. Suppose $A = \begin{bmatrix} a & b \\ c & d \end{bmatrix}$ then $A_1 = (a, c)$ and $A_2 = (b, d)$. Thus,

$$\begin{aligned} \det \begin{bmatrix} a & b \\ c & d \end{bmatrix} &= (e^1 \wedge e^2)(A_1, A_2) \\ &= (e^1 \otimes e^2 - e^2 \otimes e^1)((a, c), (b, d)) \\ &= e^1(a, c)e^2(b, d) - e^2(a, c)e^1(b, d) \\ &= ad - bc. \end{aligned}$$

I hope this is not surprising!

Example 6.2.12. Suppose $A = [A_1 | A_2 | A_3]$. I propose the determinant of A is given by the top-form on $\mathbb{R}^3$ via the formula $\det(A) = (e^1 \wedge e^2 \wedge e^3)(A_1, A_2, A_3)$. Let's see if we can find the expansion by cofactors. By the definition we have $e^1 \wedge e^2 \wedge e^3 =$

$$\begin{aligned} &= e^1 \otimes e^2 \otimes e^3 + e^2 \otimes e^3 \otimes e^1 + e^3 \otimes e^1 \otimes e^2 - e^3 \otimes e^2 \otimes e^1 - e^2 \otimes e^1 \otimes e^3 - e^1 \otimes e^3 \otimes e^2 \\ &= e^1 \otimes (e^2 \otimes e^3 - e^3 \otimes e^2) - e^2 \otimes (e^1 \otimes e^3 - e^3 \otimes e^1) + e^3 \otimes (e^1 \otimes e^2 - e^2 \otimes e^1) \\ &= e^1 \otimes (e^2 \wedge e^3) - e^2 \otimes (e^1 \wedge e^3) + e^3 \otimes (e^1 \wedge e^2). \end{aligned}$$

I submit to the reader that this is precisely the cofactor expansion formula with respect to the first

column of A . Suppose $A = \begin{bmatrix} a & b & c \\ d & e & f \\ g & h & i \end{bmatrix}$ then $A_1 = (a, d, g)$, $A_2 = (b, e, h)$ and $A_3 = (c, f, i)$.

Calculate,

$$\begin{aligned} \det(A) &= e^1(A_1)(e^2 \wedge e^3)(A_2, A_3) - e^2(A_1)(e^1 \wedge e^3)(A_2, A_3) + e^3(A_1)(e^1 \wedge e^2)(A_2, A_3) \\ &= a(e^2 \wedge e^3)(A_2, A_3) - d(e^1 \wedge e^3)(A_2, A_3) + g(e^1 \wedge e^2)(A_2, A_3) \\ &= a(ei - fh) - d(bi - ch) + g(bf - ce) \end{aligned}$$

which is precisely my claim.

6.2.3 connecting vectors and forms in $\mathbb{R}^3$

There are a couple ways to connect vectors and forms in $\mathbb{R}^3$. Mainly we need the following maps:

Definition 6.2.13.

Given $v = \langle a, b, c \rangle \in \mathbb{R}^3$ we can construct a corresponding one-form $\omega_v = ae^1 + be^2 + ce^3$ or we can construct a corresponding two-form $\Phi_v = ae^2 \wedge e^3 + be^3 \wedge e^1 + ce^1 \wedge e^2$

Recall that $\dim(\Lambda^1 \mathbb{R}^3) = \dim(\Lambda^2 \mathbb{R}^3) = 3$ hence the space of vectors, one-forms, and also two-forms are isomorphic as vector spaces. It is not difficult to show that $\omega_{v_1+cv_2} = \omega_{v_1} + c\omega_{v_2}$ and $\Phi_{v_1+cv_2} = \Phi_{v_1} + c\Phi_{v_2}$ for all $v_1, v_2 \in \mathbb{R}^3$ and $c \in \mathbb{R}$. Moreover, $\omega_v = 0$ iff $v = 0$ and $\Phi_v = 0$ iff $v = 0$ hence $\ker(\omega) = \{0\}$ and $\ker(\Phi) = \{0\}$ but this means that ω and Φ are injective and since the dimensions of the domain and codomain are 3 and these are linear transformations¹⁶ it follows ω and Φ are isomorphisms.

Example 6.2.14. Suppose $v = \langle 2, 0, 3 \rangle$ and $w = \langle 0, 1, 2 \rangle$ then $\omega_v = 2e^1 + 3e^3$ and $\omega_w = e^2 + 2e^3$. Calculate the wedge product,

$$\begin{aligned}
 \omega_v \wedge \omega_w &= (2e^1 + 3e^3) \wedge (e^2 + 2e^3) \\
 &= 2e^1 \wedge (e^2 + 2e^3) + 3e^3 \wedge (e^2 + 2e^3) \\
 &= 2e^1 \wedge e^2 + 4e^1 \wedge e^3 + 3e^3 \wedge e^2 + 6e^3 \wedge e^3 \\
 &= -3e^2 \wedge e^3 - 4e^3 \wedge e^1 + 2e^1 \wedge e^2 \\
 &= \Phi_{\langle -3, -4, 2 \rangle} \\
 &= \Phi_{v \times w}
 \end{aligned} \tag{6.16}$$

Coincidence? Nope.

Proposition 6.2.15.

Suppose $v, w \in \mathbb{R}^3$ then $\omega_v \wedge \omega_w = \Phi_{v \times w}$ where $v \times w$ denotes the cross-product which is defined by $v \times w = \sum_{i,j,k=1}^3 \epsilon_{ijk} v_i w_j e_k$.

Proof: sounds like an enjoyable homework problem. I leave it for the reader. $\square$.

6.3 an introduction to algebra

Intuitively, I view an algebra as a set of generalized numbers which allow addition and multiplication in a way which is similar to our common number systems such as $\mathbb{R}$ or $\mathbb{C}$. However, the general definition of an algebra probably forbids the phrase *similar* for many practioners of mathematics since in general algebras include multiplications which are non-associative and noncommutative. Let us state a formal definition for our reference:

¹⁶this is not generally true, note $f(x) = x^2$ has $f(x) = 0$ iff $x = 0$ and yet f is not injective. The linearity is key.

Definition 6.3.1.

An **algebra** is a vector space over a field $\mathbb{F}$ which has a **multiplication** $\star : \mathcal{A} \times \mathcal{A} \rightarrow \mathcal{A}$. We assume the multiplication satisfies distributive properties:

$$x \star (y + z) = x \star y + x \star z \quad \& \quad (x + y) \star z = x \star z + y \star z$$

for all $x, y, z \in \mathcal{A}$. We also assume scalar multiplication and $\star$ interact in the following fashion:

$$c(x \star y) = (cx) \star y = x \star (cy)$$

for all $c \in \mathbb{F}$ and $x, y \in \mathcal{A}$. If there exists $\mathbb{I} \in \mathcal{A}$ for which $\mathbb{I} \star x = x = x \star \mathbb{I}$ for each $x \in \mathcal{A}$ then $\mathcal{A}$ is **unital** with multiplicative identity¹⁷ $\mathbb{I}$. If $x \star (y \star z) = (x \star y) \star z$ for all $x, y, z \in \mathcal{A}$ then $\mathcal{A}$ is an **associative algebra**. If $x \star y = y \star x$ for all $x, y \in \mathcal{A}$ then $\mathcal{A}$ is a **commutative algebra**.

In practice, we often use juxtaposition to denote the multiplication in an algebra.

Example 6.3.2. Let $\mathbb{F}^{n \times n}$ with multiplication given by the usual matrix multiplication. Then $\mathbb{F}^{n \times n}$ is an associative unital algebra where the multiplicative identity is the identity matrix.

Example 6.3.3. Complex numbers $\mathbb{C}$ form an algebra with multiplication defined by

$$(x + iy)(a + ib) = xa - yb + i(xb + ya).$$

Here $i^2 = -1$ and this is a commutative, associative algebra with multiplicative identity 1. Since every nonzero element has a multiplicative identity, the complex number system is an example of a field.

Example 6.3.4. Hyperbolic numbers $\mathcal{H}$ form an algebra with multiplication defined by

$$(x + jy)(a + jb) = xa + yb + j(xb + ya)$$

Here $j^2 = 1$ and this is a commutative, associative algebra with multiplicative identity 1. Notice $(j + 1)(j - 1) = j^2 - 1 = 0$ yet $j + 1, j - 1 \neq 0$. Thus $j \pm 1$ cannot have a multiplicative inverse. Elements like $j \pm 1$ are **zero divisors**. It turns out that $z = x + jy$ with $y = \pm x$ give all the zero divisors in $\mathcal{H}$. The hyperbolic numbers are not a field.

Example 6.3.5. Null numbers Γ form an algebra with multiplication defined by

$$(x + \varepsilon y)(a + \varepsilon b) = xa + \varepsilon(xb + ya)$$

Here $\varepsilon^2 = 0$ and this is a commutative, associative algebra with multiplicative identity 1. Null numbers are not a field since $\varepsilon \neq 0$ yet $\varepsilon^2 = 0$ means ε is a zero-divisor. There are many zero divisors, I leave it for the reader to find a complete description.

Example 6.3.6. Let $\mathcal{A} = \mathbb{F}^n$ where

$$(x_1, \dots, x_n) \star (y_1, \dots, y_n) = (x_1 y_1, \dots, x_n y_n).$$

This is known as the **direct product algebra**. It defines an associative, commutative algebra where the multiplicative identity is given by $(1, \dots, 1)$.

If we study $\mathbb{R}^2$ with the direct product algebra then we can show it is isomorphic to an earlier example. Which do you think it is? I suppose I should define *isomorphic* in this context. In general the term isomorphism has different meanings in different context, but it always means a structure-preserving bijection. The question is, *which structure*.

Definition 6.3.7.

Two algebras $\mathcal{A}$ and $\mathcal{B}$ are said to be **isomorphic** if they are isomorphic as vector spaces with an isomorphism $\Psi : \mathcal{A} \rightarrow \mathcal{B}$ for which $\Psi(x \star y) = \Psi(x) \star \Psi(y)$ for all $x, y \in \mathcal{A}$. To be more pedantic, if $\mathcal{A}$ has multiplication $\star_A$ and $\mathcal{B}$ has multiplication $\star_B$ then we require $\Psi(x \star_A y) = \Psi(x) \star_B \Psi(y)$

Here is an example of an algebra isomorphism we have already studied this term.

Example 6.3.8. Let V be a vector space of dimension over $\mathbb{F}$. Recall $L(V)$ is the set of all linear transformations on V . If $T, S \in L(V)$ then notice $T \circ S \in L(V)$ and

$$(T_1 + T_2) \circ S = T_1 \circ S + T_2 \circ S \quad \& \quad T \circ (S_1 + S_2) = T \circ S_1 + T \circ S_2$$

Moreover, $T \circ (S \circ R) = (T \circ S) \circ R$ hence $\circ$ is associative. Composition is unital since $Id_V : V \rightarrow V$ where $Id_V(x) = x$ has $T \circ Id_V = T = Id_V \circ T$. If $\dim(V) = n$ then there exists a basis $\beta = \{v_1, \dots, v_n\}$ for V and the mapping $\Psi(T) = [T]_{\beta, \beta}$ defines an isomorphism of $L(V)$ and $\mathbb{F}^{n \times n}$ as associative algebras over $\mathbb{F}$. Likewise, $\Psi : L(\mathbb{F}^n) \rightarrow \mathbb{F}^{n \times n}$ given by $\Psi(T) = [T]$ gives an isomorphism.

There is much more to say about algebras, but I will restrain myself to stop here. Feel free to ask me more in office hours if you are interested. There is a later chapter in Dummit and Foote where some deeper theory can be found, but probably that will make better sense after you've worked through Math 421.

I'm mostly including this brief introduction as a set-up for the exterior algebra which is our primary interest. Essentially, the exterior algebra is equivalent to the theory of determinants. Since you saw the low-tech theory of determinants in the elementary matrix course, I think it is interesting to attack determinants here in a vary different fashion. I got the idea for this chapter from Mortin Curtis' *Abstract Linear Algebra* which is a delightful little book which I might use as the required textbook some future term.

6.4 wedge product and determinants

Remark: this section does not treat the wedge product as multilinear maps. Here we describe ΩV as a formal algebraic system which we use to ultimately define the determinant. In the early part of this section we are still using past knowledge of determinant to see what to expect from the exterior algebra. However, by the conclusion of this section we've shown how the abstract algebraic construction of the wedge-product allows us to implicitly define the determinant. We find a particularly easy proof that $\det(AB) = \det(A)\det(B)$ as a result. If you have a sense of deja vu as you study this, that is to be expected. There is some redundancy between this section and the past section.

Given a vector space V of dimension n over $\mathbb{F}$ there exists an associative algebraic structure known as the **exterior algebra** of V . We use the notation $\wedge$ to denote the multiplication. In general,

$$\Omega V = \bigoplus_{k=0}^n \Lambda_k V = \Lambda_0 V \oplus \Lambda_1 V \oplus \cdots \oplus \Lambda_n V.$$

where $\Lambda_0 V = \mathbb{F}$, $\Lambda_1 V = V$ and generally $\Lambda_k V$ consists of sums of k -fold wedge products of vectors. Let us make the exterior algebra explicit for $\mathbb{F}^2$ and $\mathbb{F}^3$. Our ultimate goal is to use this algebra to define the determinant and study its properties.

6.4.1 exterior algebra over $\mathbb{F}^2$

Recall $\mathbb{F}^2 = \text{span}\{e_1, e_2\}$ then we define $e_i \wedge e_j = -e_j \wedge e_i$ for $1 \leq i, j \leq 2$. Thus $e_i \wedge e_i = -e_i \wedge e_i$ hence $e_i \wedge e_i = 0$. Anytime we have a repeated element under the wedge product the result will be zero. The wedge product satisfies the usual distributive laws. In particular,

$$\begin{aligned} (ae_1 + be_2) \wedge (ce_1 + de_2) &= ae_1 \wedge (ce_1 + de_2) + be_2 \wedge (ce_1 + de_2) \\ &= ace_1 \wedge e_1 + ade_1 \wedge e_2 + bce_2 \wedge e_1 + bde_2 \wedge e_2 \\ &= (ad - bc)e_1 \wedge e_2. \end{aligned}$$

Here $\Omega(\mathbb{F}^2) = \text{span}\{1, e_1, e_2, e_1 \wedge e_2\}$ is a $2^2 = 4$ -dimensional vector space. We say 1 is a zero-vector and $e_1 \wedge e_2$ is a two-vector. Notice the appearance of the determinant in the formula above. In particular, if $A = \begin{bmatrix} a & c \\ b & d \end{bmatrix}$ then $\det(A) = ad - bc$. Note $Ae_1 = ae_1 + be_2$ and $Ae_2 = ce_1 + de_2$ thus

$$\boxed{Ae_1 \wedge Ae_2 = \det(A)e_1 \wedge e_2.}$$

6.4.2 exterior algebra over $\mathbb{F}^3$

Note $\mathbb{F}^3 = \text{span}\{e_1, e_2, e_3\}$ and suppose $e_i \wedge e_j = -e_j \wedge e_i$ for $1 \leq i, j \leq 3$. We find $e_1 \wedge e_1 = 0$, $e_2 \wedge e_2 = 0$ and $e_3 \wedge e_3 = 0$. In contrast, $e_i \wedge e_j \neq 0$ for $1 \leq i < j \leq 3$. It is convenient to use $\{e_2 \wedge e_3, e_3 \wedge e_1, e_1 \wedge e_2\}$ as the basis for $\Lambda_2(\mathbb{F}^3)$. On the other hand, $\Lambda_3(\mathbb{F}^3) = \text{span}\{e_1 \wedge e_2 \wedge e_3\}$.

Consider $A = \begin{bmatrix} A_{11} & A_{12} & A_{13} \\ A_{21} & A_{22} & A_{23} \\ A_{31} & A_{32} & A_{33} \end{bmatrix}$ then $Ae_1 = \begin{bmatrix} A_{11} \\ A_{21} \\ A_{31} \end{bmatrix}$ and $Ae_2 = \begin{bmatrix} A_{12} \\ A_{22} \\ A_{32} \end{bmatrix}$ and $Ae_3 = \begin{bmatrix} A_{13} \\ A_{23} \\ A_{33} \end{bmatrix}$.

Then

$$\begin{aligned} Ae_1 \wedge Ae_2 &= (A_{11}e_1 + A_{21}e_2 + A_{31}e_3) \wedge (A_{12}e_1 + A_{22}e_2 + A_{32}e_3) \\ &= (A_{11}A_{22} - A_{21}A_{12})e_1 \wedge e_2 + (A_{11}A_{32} - A_{31}A_{12})e_1 \wedge e_3 + (A_{21}A_{32} - A_{31}A_{22})e_2 \wedge e_3 \\ &= ae_1 \wedge e_2 + be_1 \wedge e_3 + ce_2 \wedge e_3 \end{aligned}$$

The expression above only involves the first two columns of A and we've introduced a, b, c for brevity of calculations below. Next we take the wedge product with the third column;

$$\begin{aligned} Ae_1 \wedge Ae_2 \wedge Ae_3 &= [ae_1 \wedge e_2 + be_1 \wedge e_3 + ce_2 \wedge e_3] \wedge [A_{13}e_1 + A_{23}e_2 + A_{33}e_3] \\ &= aA_{33}e_1 \wedge e_2 \wedge e_3 + bA_{23}e_1 \wedge e_3 \wedge e_2 + cA_{13}e_2 \wedge e_3 \wedge e_1 \\ &= (aA_{33} - bA_{23} + cA_{13})e_1 \wedge e_2 \wedge e_3 \\ &= ([A_{11}A_{22} - A_{21}A_{12}]A_{33} - [A_{11}A_{32} - A_{31}A_{12}]A_{23} + [A_{21}A_{32} - A_{31}A_{22}]A_{13})e_1 \wedge e_2 \wedge e_3 \end{aligned}$$

You might recognize the coefficient of $e_1 \wedge e_2 \wedge e_3$ as the determinant of A formed by the cofactor expansion of the third column of A . If we did the algebra differently then we would naturally have derived other expansions by minors. For example, if we had calculated $Ae_2 \wedge Ae_3$ first then calculated $Ae_1 \wedge (Ae_2 \wedge Ae_3)$ then the resulting expression would have produced the expansion of the determinant with respect to the first column.

$$\boxed{Ae_1 \wedge Ae_2 \wedge Ae_3 = \det(A)e_1 \wedge e_2 \wedge e_3.}$$

6.4.3 the exterior algebra over a vector space

Given a vector space V over $\mathbb{F}$ we define $\Lambda_0 V = \mathbb{F}$ and $\Lambda_1 = V$. Then

$$\Lambda_k V = \text{span}\{v_1 \wedge v_2 \wedge \cdots \wedge v_k \mid v_1, \dots, v_k \in V\}$$

elements of $\Lambda_k V$ are known as **k -vectors** of $\Lambda_k V$. The **wedge product** gives a multiplication for which the product of a p -vector and q -vector gives a $(p+q)$ -vector; $\wedge : \Lambda_p V \times \Lambda_q V \rightarrow \Lambda_{p+q} V$. The external direct sum of k -vectors ranging from $k = 0, \dots, n$ is known as the **exterior algebra** of V and it is denoted by

$$\Omega V = \bigoplus_{k=0}^n \Lambda_k V = \Lambda_0 V \oplus \Lambda_1 V \oplus \cdots \oplus \Lambda_n V.$$

The properties of the wedge product include linearity:

$$(cx + y) \wedge z = cx \wedge z + y \wedge z \quad \& \quad x \wedge (cz + w) = cx \wedge z + x \wedge w$$

for each $c \in \mathbb{F}$ and $x, y, z, w \in \Omega V$. The wedge product is also associative:

$$x \wedge (y \wedge z) = (x \wedge y) \wedge z.$$

If $x \in \Lambda_p V$ and $y \in \Lambda_q V$ then

$$x \wedge y = (-1)^{pq} y \wedge x$$

describes the **graded-commutativity** of the wedge product. If p is a non-negative integer then $x \in \Lambda_p V$ is an **even-vector** if p is even whereas x is an **odd-vector** if p is odd. Notice even vectors commute with every vector in ΩV whereas odd vectors **anticommute** with other odd vectors¹⁸.

Remark 6.4.1. k -forms verses k -vectors

If we use the dual space V^* and form wedge products of dual vectors then such objects are commonly called k -forms. For example, a dual vector is also called a 1-form. However, to be clear, there is a more sophisticated use of the term k -form which might be better known as a **k -form field**. In the same way a vector is different than a vector field a form-field is different than a k -form. We could study the smooth assignment of a k -form at each point in space. Such an object is known as differential form. There is a calculus of differential forms which is typically studied in a course on manifold theory. I cover the theory of differential forms on Euclidean space in my Advanced Calculus course. Come join us.

6.4.4 definition of determinant

The notation below is very helpful at times.

Definition 6.4.2.

$\epsilon_{i_1 \dots i_n}$ is the **completely antisymmetric symbol** which is defined to be antisymmetric in any exchange of indices and $\epsilon_{1 \dots n} = 1$.

Example 6.4.3. $\epsilon_{12} = 1$ and $\epsilon_{21} = -1$ whereas $\epsilon_{11} = \epsilon_{22} = 0$.

¹⁸ x and y anticommute if $x \wedge y = -y \wedge x$. Likewise, x and y commute if $x \wedge y = y \wedge x$

Example 6.4.4. $\epsilon_{123} = \epsilon_{231} = \epsilon_{312} = 1$ and $\epsilon_{321} = \epsilon_{213} = \epsilon_{132} = -1$ whereas the other 21-values of ϵ_{ijk} are zero.

Generally, $\epsilon_{i_1, \dots, i_n}$ has $n!$ nonzero values of which half are 1 and half are -1 . The antisymmetric symbol could be used to formulate the wedge product. The following lemma will be central to connecting the wedge product to the determinants (we've not defined determinant yet in these notes, but I know you saw them in your previous course on linear algebra)

Lemma 6.4.5.

$$e_{j_1} \wedge \cdots \wedge e_{j_n} = \epsilon_{j_1 \dots j_n} e_1 \wedge \cdots \wedge e_n.$$

Proof: if any pair of indices is repeated then both sides of the equation in the lemma are zero. If $j_1, \dots, j_n$ are distinct then there exists a rearrangement to the list $1, \dots, n$ to $j_1, \dots, j_n$ by swapping entries as needed. If the net number of swaps is even then $\epsilon_{j_1 \dots j_n} = 1$ and

$$e_{j_1} \wedge \cdots \wedge e_{j_n} = e_1 \wedge \cdots \wedge e_n.$$

whereas if the net number of swaps is odd then $\epsilon_{j_1 \dots j_n} = -1$ and

$$e_{j_1} \wedge \cdots \wedge e_{j_n} = -e_1 \wedge \cdots \wedge e_n. \quad \square$$

For an n -dimensional vector space the maximum number of vectors whose wedge product is non-trivial is n . I'll prove the result for $\mathbb{F}^n$ as that is our primary application.

Proposition 6.4.6.

If $v_1, \dots, v_n \in \mathbb{F}^n$ then there exists a unique $D \in \mathbb{F}$ for which

$$v_1 \wedge \cdots \wedge v_n = D e_1 \wedge \cdots \wedge e_n.$$

Proof: suppose $v_i = \sum_{j_i} A_{j_i i} e_{j_i}$ for $1 \leq i \leq n$ then calculate

$$\begin{aligned} v_1 \wedge \cdots \wedge v_n &= \left(\sum_{j_1} A_{j_1 1} e_{j_1} \right) \wedge \cdots \wedge \left(\sum_{j_n} A_{j_n n} e_{j_n} \right) \\ &= \sum_{j_1, \dots, j_n} A_{j_1 1} \cdots A_{j_n n} e_{j_1} \wedge \cdots \wedge e_{j_n} \\ &= \sum_{j_1, \dots, j_n} \epsilon_{j_1 \dots j_n} A_{j_1 1} \cdots A_{j_n n} e_1 \wedge \cdots \wedge e_n. \end{aligned}$$

Let $D = \sum_{j_1, \dots, j_n} \epsilon_{j_1 \dots j_n} A_{j_1 1} \cdots A_{j_n n}$ and the proposition follows. $\square$

Using the notation above we note $A = [v_1 | \cdots | v_n]$ hence we make the following definition:

Definition 6.4.7.

Let $A \in \mathbb{F}^{n \times n}$ then $\det(A) \in \mathbb{F}$ is the unique constant for which

$$\text{col}_1(A) \wedge \cdots \wedge \text{col}_n(A) = \det(A) e_1 \wedge \cdots \wedge e_n.$$

In view of the proof of Proposition 6.4.6 we find the following formula for the determinant:

$$\det(A) = \sum_{j_1, \dots, j_n} \epsilon_{j_1 \dots j_n} A_{j_1 1} \cdots A_{j_n n}.$$

Notice $I = [e_1 | \cdots | e_n]$ hence

$$\text{col}_1(I) \wedge \cdots \wedge \text{col}_n(I) = e_1 \wedge \cdots \wedge e_n$$

thus $\det(I) = 1$. Next we derive $\det(cA) = c^n \det(A)$ for $A \in \mathbb{F}^{n \times n}$ and $c \in \mathbb{F}$,

$$\text{col}_1(cA) \wedge \cdots \wedge \text{col}_n(cA) = \det(cA) e_1 \wedge \cdots \wedge e_n$$

and as $\text{col}_i(cA) = c \text{col}_i(A)$ for $i = 1, \dots, n$ we find

$$c^n \text{col}_1(A) \wedge \cdots \wedge \text{col}_n(A) = \det(cA) e_1 \wedge \cdots \wedge e_n$$

thus

$$c^n \det(A) e_1 \wedge \cdots \wedge e_n = \det(cA) e_1 \wedge \cdots \wedge e_n$$

therefore $c^n \det(A) = \det(cA)$. Let me record this result for future reference.

Proposition 6.4.8.

If $A \in \mathbb{F}^{n \times n}$ and $c \in \mathbb{F}$ then $\det(cA) = c^n \det(A)$.

Remark 6.4.9. *determinant of transpose*

In view of the development thus far I suspect there is a simple proof that $\det(A) = \det(A^T)$. But, I've not found it yet and the proof in the other approach is somewhat technical.

6.4.5 wedge product proofs for determinants

Proposition 6.4.10.

If $T : \mathbb{F}^n \times \cdots \times \mathbb{F}^n \rightarrow \mathbb{F}$ is defined by

$$T(v_1, \dots, v_n) = \det(v_1 | \cdots | v_n)$$

then T is an antisymmetric n -linear map for which $T(e_1, \dots, e_n) = 1$. Moreover, T is the unique n -linear antisymmetric map which maps $(e_1, \dots, e_n)$ to 1.

Proof: we already showed $\det(e_1 | \cdots | e_n) = 1$. Suppose $v_1, \dots, v_n, w_j \in \mathbb{F}^n$ and $c \in \mathbb{F}$,

$$\begin{aligned} \det(v_1 | \cdots | cv_j + w_j | \cdots | v_n) e_1 \wedge \cdots \wedge e_n &= \\ &= v_1 \wedge \cdots \wedge (cv_j + w_j) \wedge \cdots \wedge v_n \\ &= cv_1 \wedge \cdots \wedge v_j \wedge \cdots \wedge v_n + v_1 \wedge \cdots \wedge w_j \wedge \cdots \wedge v_n \\ &= (c \det(v_1 | \cdots | v_j | \cdots | v_n) + \det(v_1 | \cdots | w_j | \cdots | v_n)) e_1 \wedge \cdots \wedge e_n \end{aligned} \tag{6.17}$$

Thus equating coefficients of $e_1 \wedge \cdots \wedge e_n$ we find

$$\det(v_1 | \cdots | cv_j + w_j | \cdots | v_n) = c \det(v_1 | \cdots | v_j | \cdots | v_n) + \det(v_1 | \cdots | w_j | \cdots | v_n)$$

Thus T is an n -linear map. Furthermore, since

$$v_1 \wedge \cdots \wedge v_i \wedge \cdots \wedge v_j \wedge \cdots \wedge v_n = -v_1 \wedge \cdots \wedge v_j \wedge \cdots \wedge v_i \wedge \cdots \wedge v_n$$

we find

$$\det(v_1 | \cdots | v_i | \cdots | v_j | \cdots | v_n) = -\det(v_1 | \cdots | v_j | \cdots | v_i | \cdots | v_n)$$

thus T is antisymmetric.

Suppose S is an antisymmetric n -linear map and $S(e_1, \dots, e_n) = 1$. If $v_1, \dots, v_n \in \mathbb{F}^n$ and suppose $v_i = \sum_{j_i} A_{j_i i} e_{j_i}$ for $1 \leq i \leq n$ then $A = [v_1 | \cdots | v_n]$ and

$$\begin{aligned} S(v_1, \dots, v_n) &= S\left(\sum_{j_1} A_{j_1 1} e_{j_1}, \dots, \sum_{j_n} A_{j_n n} e_{j_n}\right) \\ &= \sum_{j_1, \dots, j_n} A_{j_1 1} \cdots A_{j_n n} S(e_{j_1}, \dots, e_{j_n}) \\ &= \sum_{j_1, \dots, j_n} A_{j_1 1} \cdots A_{j_n n} \epsilon_{j_1 \dots j_n} S(e_1, \dots, e_n) \\ &= \sum_{j_1, \dots, j_n} A_{j_1 1} \cdots A_{j_n n} \epsilon_{j_1 \dots j_n} \\ &= \det[v_1 | \cdots | v_n] \\ &= T(v_1, \dots, v_n). \quad \square \end{aligned}$$

Matrix multiplication naturally extends to k -vectors.

Definition 6.4.11.

Let $A \in \mathbb{F}^{n \times n}$ then for $v_1, \dots, v_k \in \mathbb{F}^{n \times n}$ we define

$$A(v_1 \wedge \cdots \wedge v_k) = Av_1 \wedge \cdots \wedge Av_k$$

Since matrix multiplication is associative we derive the identity below

$$\begin{aligned} (AB)(v_1 \wedge \cdots \wedge v_k) &= (AB)v_1 \wedge \cdots \wedge (AB)v_k \\ &= A(Bv_1) \wedge \cdots \wedge A(Bv_k) \\ &= A((Bv_1) \wedge \cdots \wedge (Bv_k)) \\ &= A(B(v_1 \wedge \cdots \wedge v_k)). \end{aligned}$$

Proposition 6.4.12.

Let $A, B \in \mathbb{F}^{n \times n}$ then $\det(AB) = \det(A)\det(B)$.

Proof: by definition of the determinant we note

$$\text{col}_1(AB) \wedge \cdots \wedge \text{col}_n(AB) = \det(AB)e_1 \wedge \cdots \wedge e_n.$$

Likewise, $\text{col}_1(B) \wedge \cdots \wedge \text{col}_n(B) = \det(B)e_1 \wedge \cdots \wedge e_n$. Multiply by A and apply the definition of matrix multiplication on a n -vector,

$$A\text{col}_1(B) \wedge \cdots \wedge A\text{col}_n(B) = \det(B)Ae_1 \wedge \cdots \wedge Ae_n.$$

Since $ABe_i = Acol_i(B) = col_i(AB)$ for $1 \leq i \leq n$ the equation above implies

$$col_1(AB) \wedge \cdots \wedge col_n(AB) = det(B)Ae_1 \wedge \cdots \wedge Ae_n.$$

Consequently, by the definition of the determinant,

$$det(AB)e_1 \wedge \cdots \wedge e_n = det(B)det(A)e_1 \wedge \cdots \wedge e_n$$

thus $det(AB) = det(B)det(A)$ which completes the proof. $\square$

The traditional proof offered for the product identity of determinants is based on a casewise analysis of elementary matrices. The proof above is far easier, but it does require some preparation.

6.5 linear dependence and the wedge product

You hopefully recall from your previous course that a set of n -vectors in $\mathbb{F}^n$ is linearly independent if and only if the determinant of the matrix formed by gluing the vectors together is nonzero. Likewise, a zero determinant signals a linear dependence amongst the columns forming such a matrix. Short of some clever modification, it seems the determinant only helps us with a particular type of question. For example, the determinant wouldn't seem to say anything directly about linear dependence of a set of two vectors in $\mathbb{F}^n$ where $n \geq 3$.

Proposition 6.5.1.

Suppose $x, y \in \mathbb{F}^n$. If $\{x, y\}$ is linearly dependent then $x \wedge y = 0$.

Proof: If $\{x, y\}$ is linearly dependent then there exists $k \in \mathbb{F}$ for which $y = kx$ thus $x \wedge y = x \wedge (kx) = kx \wedge x = 0$ since $x \wedge x = -x \wedge x$ implies¹⁹ $x \wedge x = 0$. $\square$

The contrapositive of the above result says if $x \wedge y \neq 0$ then $\{x, y\}$ is linearly independent.

Proposition 6.5.2.

Suppose $x_1, \dots, x_k \in \mathbb{F}^n$. If $\{x_1, \dots, x_k\}$ is linearly dependent then $x_1 \wedge \cdots \wedge x_k = 0$.

Proof: if $\{x_1, \dots, x_k\}$ is linearly dependent then there exists x_j for which $x_j = \sum_{i \neq j} c_i x_i$ then

$$\begin{aligned} x_1 \wedge \cdots \wedge x_k &= x_1 \wedge \cdots \wedge \left(\sum_{i \neq j} c_i x_i \right) \wedge \cdots \wedge x_k \\ &= \sum_{i \neq j} c_i x_1 \wedge \cdots \wedge x_i \wedge \cdots \wedge x_k \\ &= \sum_{i \neq j} \pm c_i x_i \wedge x_i \wedge x_1 \wedge \cdots \wedge x_k. \end{aligned}$$

where in the last line $x_1 \wedge \cdots \wedge x_k$ does not include x_i and the $\pm$ depends on how many wedge products needed to be swapped in order to move $x_i \wedge x_i$ to the start of the expression. Then,

¹⁹technically, I am assuming our field is one in which $2 \neq 0$, this is often the case, but the binary number system and its extensions have $1 + 1 = 0$.

$x_i \wedge x_i = 0$ for all $i \neq j$ hence $x_1 \wedge \cdots \wedge x_k = 0$. $\square$

Notice the contrapositive of the proposition gives us the statement: if $x_1 \wedge \cdots \wedge x_k \neq 0$ then $\{x_1, \dots, x_k\}$ is linearly independent.

Remark 6.5.3. *meaning of the wedge product as it pertains to generalized volume.*

The magnitude of a wedge product says something about volume in the generalized sense. For example, $x \wedge y$ has a magnitude which corresponds to the parallelogram with edges x and y . Likewise, $x \wedge y \wedge z$ in some sense represents the oriented volume with edges x , y and z . Notice, in $\mathbb{R}^3$ we have the identity

$$x \wedge y \wedge z = \det(x|y|z)e_1 \wedge e_2 \wedge e_3$$

and $\det(x|y|z) = x \bullet (y \times z)$ is the **triple product** which gives the signed-volume of the parallel-piped with edges x , y and z . In general, $x_1 \wedge \cdots \wedge x_k$ has a magnitude which corresponds to the signed k -volume of the convex-hull generated by $x_1, \dots, x_k$. Sometimes people say that determinants are volumes. That is a shorthand for saying determinants allow the calculation of volume. Likewise, wedge products are volumes. However, the wedge product in $\mathbb{R}^n$ naturally allows the calculation of volume for any sub-dimensional object. In contrast, the determinant only directly allows for calculation of n -volume in $\mathbb{R}^n$. Of course, determinants are more than that, you may recall Cramer's Rule or the classical adjoint formula for the inverse of a matrix. Many interesting objects can be formulated with determinants.

Another interesting rabbit to chase here is the connection between k and $(n - k)$ -volumes. If $W \subseteq \mathbb{R}^n$ is k -dimensional then $W^\perp$ is $(n - k)$ -dimensional. Geometrically, it is intuitively clear that W fixes $W^\perp$ and vice-versa. I can specify a plane through the origin either by the span of its tangent vectors or via the normal line. That reciprocity extends far beyond lines and planes. An algebraic aspect of this duality is seen in the correspondence between k -vectors and $(n - k)$ -vectors. It turns out there is an isomorphism between $\Lambda_k \mathbb{R}^n$ and $\Lambda_{n-k} \mathbb{R}^n$. In fact, this isomorphism also preserves lengths of multivectors. Perhaps we will return to this vague paragraph later in the course.

Finally, we should recall the application of the determinant to invertibility of matrices.

Theorem 6.5.4.

Let $A \in \mathbb{F}^{n \times n}$ then

- (1.) $\det(A) \neq 0$ if and only if A^{-1} exists,
- (2.) $\det(A) = 0$ if and only if $Ax = 0$ has a nonzero solution.

Proof: let $A = [v_1 | \cdots | v_n]$.

(1.) Suppose $\det(A) \neq 0$ then

$$v_1 \wedge \cdots \wedge v_n = \det(A)e_1 \wedge \cdots \wedge e_n \neq 0$$

thus $v_1, \dots, v_n$ are linearly independent. Therefore, $v_1, \dots, v_n$ serves as a basis for $\mathbb{F}^n$ and we find $Aw_i = e_i$ has a unique solution for each $i = 1, \dots, n$ and hence $A^{-1} = [w_1 | \cdots | w_n]$. Conversely, if

A^{-1} exists then $A^{-1}A = I$ and thus $\det(A^{-1})\det(A) = \det(I) = 1$ thus $\det(A) \neq 0$.

(2.) Suppose $\det(A) = 0$ then if $Ax = 0$ implies $x = 0$ we find

$$x_1v_1 + \cdots + x_nv_n = 0$$

implies $x_1 = 0, \dots, x_n = 0$ thus $v_1, \dots, v_n$ is linearly independent hence

$$v_1 \wedge \cdots \wedge v_n = \det(A)e_1 \wedge \cdots \wedge e_n \neq 0$$

thus $\det(A) \neq 0$ which a contradiction. Thus there exists $x \neq 0$ for which $Ax = 0$. Conversely, if $Ax = 0$ for some $x \neq 0$ then we find a linear dependence amongst the columns of A which is to say $v_1, \dots, v_n$ are linearly dependent and

$$v_1 \wedge \cdots \wedge v_n = \det(A)e_1 \wedge \cdots \wedge e_n = 0$$

hence $\det(A) = 0$. $\square$

The usual proof I offer in Math 221 is accomplished via some subtle elementary matrix arguments.

6.6 bilinear forms and geometry, metric duality

The concept of a metric goes beyond the familiar case of the dot-product. This is a bit more general than an inner-product. Every inner product is a metric. However, not every metric is an inner-product. The general theory of metrics is far more complicated than that of inner-products. We only give an introduction here.

6.6.1 metric geometry

A **geometry** is a vector space paired with a metric. For example, if we pair $\mathbb{R}^n$ with the dot-product you get Euclidean space. However, if we pair $\mathbb{R}^4$ with the Minkowski metric then we obtain Minkowski space.

Definition 6.6.1.

If V is a vector space and $g : V \times V \rightarrow \mathbb{R}$ is

(i.) **bilinear**: $g \in T_2^0V$,

(ii.) **symmetric**: $g(x, y) = g(y, x)$ for all $x, y \in V$,

(iii.) **nondegenerate**: $g(x, y) = 0$ for all $x \in V$ implies $y = 0$.

then we call g a **metric** on V .

If $V = \mathbb{R}^n$ then we can write $g(x, y) = x^T G y$ where $[g] = G$. Moreover, $g(x, y) = g(y, x)$ implies $G^T = G$. Nondegenerate means that $g(x, y) = 0$ for all $y \in \mathbb{R}^n$ iff $x = 0$. It follows that $Gy = 0$ has no non-trivial solutions hence G^{-1} exists.

Example 6.6.2. Suppose $g(x, y) = x^T y$ for all $x, y \in \mathbb{R}^n$. This defines a metric for $\mathbb{R}^n$, it is just the dot-product. Note that $g(x, y) = x^T y = x^T I y$ hence we see $[g] = I$ where I denotes the identity matrix in $\mathbb{R}^{n \times n}$.

Example 6.6.3. Suppose $v = (v^0, v^1, v^2, v^3), w = (w^0, w^1, w^2, w^3) \in \mathbb{R}^4$ then define the **Minkowski product** of v and w as follows:

$$g(v, w) = -v^0w^0 + v^1w^1 + v^2w^2 + v^3w^3$$

It is useful to write the Minkowski product in terms of a matrix multiplication. Observe that for $x, y \in \mathbb{R}^4$,

$$g(x, y) = -x^0y^0 + x^1y^1 + x^2y^2 + x^3y^3 = \begin{pmatrix} x^0 & x^1 & x^2 & x^3 \end{pmatrix} \begin{pmatrix} -1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \end{pmatrix} \begin{pmatrix} y^0 \\ y^1 \\ y^2 \\ y^3 \end{pmatrix} \equiv x^t \eta y$$

where we have introduced η the matrix of the Minkowski product. Notice that $\eta^T = \eta$ and $\det(\eta) = -1 \neq 0$ hence $g(x, y) = x^t \eta y$ makes g a symmetric, nondegenerate bilinear form on $\mathbb{R}^4$. The formula is clearly related to the dot-product. Suppose $\vec{v} = (v^0, \vec{v})$ and $\vec{w} = (w^0, \vec{w})$ then note

$$g(v, w) = -v^0w^0 + \vec{v} \cdot \vec{w}$$

For vectors with zero in the zeroth slot this Minkowski product reduces to the dot-product. However, for vectors which have nonzero entries in both the zeroth and later slots much differs. Recall that any vector's dot-product with itself gives the square of the vectors length. Of course this means that $\vec{x} \cdot \vec{x} = 0$ iff $\vec{x} = 0$. Contrast that with the following: if $v = (1, 1, 0, 0)$ then

$$g(v, v) = -1 + 1 = 0$$

Yet $v \neq 0$. Why study such a strange generalization of length? The answer lies in physics. I'll give you a brief account by defining a few terms: Let $v = (v^0, v^1, v^2, v^3) \in \mathbb{R}^4$ then we say

1. v is a timelike vector if $\langle v, v \rangle < 0$
2. v is a lightlike vector if $\langle v, v \rangle = 0$
3. v is a spacelike vector if $\langle v, v \rangle > 0$

If we consider the trajectory of a massive particle in $\mathbb{R}^4$ that begins at the origin then at any later time the trajectory will be located at a timelike vector. If we consider a light beam emitted from the origin then at any future time it will be located at the tip of a lightlike vector. Finally, spacelike vectors point to points in $\mathbb{R}^4$ which cannot be reached by the motion of physical particles that pass throughout the origin. We say that massive particles are confined within their light cones, this means that they are always located at timelike vectors relative to their current position in space time. If you'd like to know more I can recommend a few books.

At this point you might wonder if there are other types of metrics beyond these two examples. Surprisingly, in a certain sense, no. A rather old theorem of linear algebra due to Sylvester states that we can change coordinates so that the metric more or less resembles either the dot-product or something like it with some sign-flips. Perhaps I will prove this result in lecture. Another interesting result to consider is Hurwitz Theorem that the only real normed division algebras are $\mathbb{R}$, $\mathbb{C}$ the four dimensional quaternions and finally the eight dimensional octonions. Such mathematics is the beginning of many interesting stories we ought to tell.

Chapter 7

Appendix on Modular Arithmetic

And I gave my heart to know wisdom, and to know madness and folly: I perceived that this also is vexation of spirit. ECCLESIASTES 1:17, KJV

I do not expect you to be able to prove the results proved in this chapter. However, I would like you to be able to do modular arithmetic and be able to decide whether or not a given modular integer has a multiplicative inverse. So, this chapter is a bit overkill, but I err on the side of curiosity here.

7.1 $\mathbb{Z}$ -Basics

Let's start at the very beginning, it is a good place to start.

Definition 7.1.1. *The integers $\mathbb{Z}$ are the set of natural numbers $\mathbb{N}$ together with 0 and the negatives of $\mathbb{N}$. It is possible to concretely construct (we will not) these from sets and set-operations.*

From the construction of $\mathbb{Z}$ it is clear (we assume these to be true)

1. the sum of integers is an integer
2. the product of integers is an integer
3. the usual rules of arithmetic hold for $\mathbb{Z}$

Much is hidden in (3.): let me elaborate, we assume for all $a, b, c \in \mathbb{Z}$,

$$\begin{aligned}a + b &= b + a \\ab &= ba \\a(b + c) &= ab + ac \\(a + b)c &= ac + bc \\(a + b) + c &= a + (b + c) \\(ab)c &= a(bc) \\a + 0 &= 0 + a = a \\1a &= a1.\end{aligned}$$

Where we assume the **order of operations** is done multiplication then addition; so, for example, $ab + ac$ means to first multiply a with b and a with c then you add the result.

Let me comment briefly about our standard conventions for the presentation of numbers. If I write 123 then we understand this is the **base-ten** representation. In particular,

$$123 = 1 \times 10^2 + 2 \times 10 + 3.$$

On the other hand, $1 \cdot 2 \cdot 3$ denotes the product of 1, 2 and 3 and $1 \cdot 2 \cdot 3 = 6$. By default, algebraic variables juxtaposed denote multiplication; xy denotes x multiplied by y . If we wish for symbolic variables to denote digits in a number then we must explain this explicitly. For example, to study all numbers between 990 and 999 I could analyze $99x$ where $x \in \{0, 1, \dots, 9\}$. But, to be clear I ought to preface such analysis by a statement like: let $99x$ be the base-ten representation of a number where x represents the 1's digit.

7.1.1 division algorithm

Division is repeated subtraction. For example, consider $11/3$. Notice repeated subtraction of the dividing number¹ 3 gives:

$$11 - 3 = 8 \quad 8 - 3 = 5 \quad 5 - 3 = 2$$

then we cannot subtract anymore. We were able to subtract 3 copies of 3 from 11. Then we stopped at 2 since $2 < 3$. To summarize,

$$\boxed{11 = 3(3) + 2}$$

We say 2 is the **remainder**; the remainder is the part which is too small to subtract for the given *dividing number*. Divide the boxed equation by the divisor to see:

$$\frac{11}{3} = 3 + \frac{2}{3}.$$

The generalization of the boxed equation for an arbitrary pair of natural numbers is known as the **division algorithm**.

Theorem 7.1.2. positive division algorithm: *If $a, b \in \mathbb{Z}$ with $b > 0$ then there is a unique quotient $q \in \mathbb{Z}$ and remainder $r \in \mathbb{Z}$ for which $a = qb + r$ and $0 \leq r < b$.*

Proof (existence): suppose $a, b \in \mathbb{Z}$ and $b > 0$. Construct $R = \{a - nb \mid q \in \mathbb{Z}, a - nb \geq 0\}$. The set R comprises all non-negative integers which are reached from a by integer multiples of b . Explicitly,

$$R = \{a, a \pm b, a \pm 2b, \dots\} \cap \{0, 1, 2, \dots\}.$$

To prove R is non-empty we consider $n = -|a| \in \mathbb{Z}$ yields $a - nb = a + |a|b$. If $a \geq 0$ then clearly $a + |a|b \geq 0$. If $a < 0$ then $|a| = -a$ hence $a + |a|b = -|a| + |a|b = |a|(b - 1)$ but $b \in \mathbb{N}$ by assumption hence $b \geq 1$ and we find $a + |a|b \geq 0$. Therefore, as R is a non-empty subset of the non-negative integers. We apply the **Well-Ordering-Principle** to deduce there exists a smallest element $r \in R$.

Suppose r is the smallest element in R and $r \geq b$. In particular, $r = a - nb$ for some $n \in \mathbb{Z}$. Thus $a - nb \geq b$ hence $r' = a - (n + 1)b \geq 0$ hence $r' \in R$ and $r' < r$. But $r' < r$ contradicts r being the smallest element. Thus, using proof by contradiction, we find $r < b$.

¹my resident Chinese scholar tells me in Chinese a/b has the "dividing" number b and the "divided" number a . I am tempted to call b the divisor, but the term "divisor" has a precise meaning, if b is a divisor of a then $a = mb$ for some $n \in \mathbb{Z}$. In our current discussion, to say b is a divisor assumes the remainder is zero.

Proof (uniqueness): assume $q, q' \in \mathbb{Z}$ and $r, r' \in \mathbb{Z}$ such that $a = qb + r$ and $a = q'b + r'$ where $0 \leq r, r' < b$. We have $qb + r = q'b + r'$ hence $(q - q')b = r - r'$. Suppose towards a contradiction $q \neq q'$. Since $q, q' \in \mathbb{Z}$ the inequality of q and q' implies $|q - q'| \geq 1$ and thus $|r - r'| = |(q - q')b| \geq |b| = b$. However, $r, r' \in [0, b)$ thus the distance² between r and r' cannot be larger than or equal to b . This is a contradiction, therefore, $q = q'$. Finally, $qb + r = q'b + r'$ yields $r = r'$. $\square$

We can say more about q and r in the case $b > 0$. We have

$$\frac{a}{b} = q + \frac{r}{b} \quad \& \quad q = \lfloor a/b \rfloor$$

That is q is the greatest integer which is below a/b . The function $x \mapsto \lfloor x \rfloor$ is the **floor function**. For example,

$$\lfloor -0.4 \rfloor = -1, \quad \lfloor \pi \rfloor = 3, \quad \lfloor n + \varepsilon \rfloor = n$$

for all $n \in \mathbb{Z}$ provided $0 \leq \varepsilon < 1$. It is easy to calculate the floor function of x when x is presented in decimal form. For example,

$$\frac{324}{11} = 29.4545\ldots \Rightarrow \frac{324}{11} = 29 + 0.4545\ldots \Rightarrow 324 = 29(11) + (0.4545\ldots)(11)$$

We can calculate, $0.4545 \cdot 11 = 4.9995$. From this we find

$$324 = 29(11) + 5$$

In other words, $\frac{324}{11} = 29 + \frac{5}{11}$. The decimal form of numbers and the floor function provides a simple way to find quotients and remainders.

Consider $456/(-10) = -45.6 = -45 - 0.6$ suggests $456 = (-10)(-45) + 6$. In the case of a negative divisor ($b < 0$) the division algorithm needs a bit of modification:

Theorem 7.1.3. nonzero division algorithm: *If $a, b \in \mathbb{Z}$ with $b \neq 0$ then there is a unique quotient $q \in \mathbb{Z}$ and remainder $r \in \mathbb{Z}$ for which*

$$a = qb + r \quad \& \quad 0 \leq r < |b|.$$

Proof: Theorem 7.1.2 covers case $b > 0$. Thus, assume $b < 0$ hence $b' = -b > 0$. Apply Theorem 7.1.2 to $a, b' \in \mathbb{Z}$ to find q', r' such that $a = q'b' + r'$ with $0 \leq r' < b'$. However, $b' = -b = |b|$ as $b < 0$. Thus,

$$a = -q'b + r'$$

with $0 \leq r' < |b|$. Identify $q = -q'$ and $r = r'$ in the case $b < 0$. Uniqueness is clear from the equations which define q and r from the uniquely given q' and r' . This concludes the proof as $b \neq 0$ means either $b < 0$ or $b > 0$. $\square$

The selection of the quotient in the negative divisor case is given by the **ceiling** function $x \mapsto \lceil x \rceil$. The notation $\lceil x \rceil$ indicates the next integer which is greater than or equal to x . For example,

$$\lceil 456/(-10) \rceil = -45, \quad \lceil 3.7 \rceil = 4, \quad \lceil n - \varepsilon \rceil = n$$

for all $n \in \mathbb{Z}$ given $0 \leq \varepsilon < 1$.

²for a non-geometric argument here: note $0 \leq r < b$ and $0 \leq r' < b$ imply $-r' < r - r' < b - r' \leq b$. But, $r' < b$ gives $-b < -r'$ hence $-b < r - r' < b$. Thus $|r - r'| < b$. Indeed, the distance between r and r' is less than b .

Remark 7.1.4. The division algorithm proves an assertion of elementary school arithmetic. For example, consider the **improper fraction** $10/3$ we can write it as the sum of 3 and $1/3$. When you write $3\frac{1}{3}$ what is truly meant is $3 + \frac{1}{3}$. In fact, the truth will set you free of a myriad of errors which arise from the poor notation $3\frac{1}{3}$. With this example in mind, let $a, b \in \mathbb{N}$. The division algorithm simply says for a/b there exists $q, r \in \mathbb{N} \cup \{0\}$ such that $a = qb + r$ hence $a/b = q + r/b$ where $0 \leq r < b$. This is merely the statement that any improper fraction can be reduced to the sum of a whole number and a proper fraction. In other words, you already knew the division algorithm. However, thinking of it without writing fractions is a bit of an adjustment for some of us.

7.1.2 divisibility in $\mathbb{Z}$

Consider $105 = 3 \cdot 5 \cdot 7$. We say 3 is a *factor* or *divisor* of 105. Also, we say 35 *divides* 105. Furthermore, 105 is a *multiple* of 3. Indeed, 105 is also a multiple of 5, 7 and even 21 or 35. Examples are nice, but, definitions are crucial:

Definition 7.1.5. Let $a, b \in \mathbb{Z}$ then we say b **divides** a if there exists $c \in \mathbb{Z}$ such that $a = bc$. If b divides a then we also say b is a **factor** of a and a is a **multiple** of b .

The notation $b \mid a$ means b divides a . If b does not divide a then we write $b \nmid a$. The divisors of a given number are not unique. For example, $105 = 7(15) = (3)(35) = (-1)(-105)$. However, the prime divisors are unique up to reordering: $105 = (3)(5)(7)$. Much of number theory is centered around the study of primes. We ought to give a proper definition:

Definition 7.1.6. If $p \in \mathbb{N}$ such that $n \mid p$ implies $n = p$ or $n = 1$ then we say p is **prime**.

In words: a prime is a positive integer whose only divisors are 1 and itself.

There are many interesting features of divisibility. Notice, every number $b \in \mathbb{Z}$ divides 0 as $0 = b \cdot 0$. Furthermore, $b \mid b$ for all $b \in \mathbb{Z}$ as $b = b \cdot 1$. In related news, 1 is a factor of every integer and every integer is a multiple of 1^3

Proposition 7.1.7. Let $a, b, c, d, m \in \mathbb{Z}$. Then,

- (i.) if $a \mid b$ and $b \mid c$ then $a \mid c$,
- (ii.) if $a \mid b$ and $c \mid d$ then $ac \mid bd$,
- (iii.) if $m \neq 0$, then $ma \mid mb$ if and only if $a \mid b$
- (iv.) if $d \mid a$ and $a \neq 0$ then $|d| \leq |a|$.

Proof (i.) : suppose $a \mid b$ and $b \mid c$. By the definition of divisibility there exist $m, n \in \mathbb{Z}$ such that $b = ma$ and $c = nb$. Hence $c = n(ma) = (nm)a$. Therefore, $a \mid c$ as $nm \in \mathbb{Z}$.

Proof (ii.) : suppose $a \mid b$ and $c \mid d$. By the definition of divisibility there exist $m, n \in \mathbb{Z}$ such that $b = ma$ and $d = nc$. Substitution yields $bd = (ma)(nc) = mn(ac)$. But, $mn \in \mathbb{Z}$ hence we have shown $ac \mid bd$.

³I should mention, I am partly following the excellent presentation of Jones and Jones *Elementary Number Theory* which I almost used as the text for Math 307 in Spring 2015. We're on page 4.

Proof (iii.) : left to the reader.

Proof (iv.) : if $d \mid a$ and $a \neq 0$ then $a = md$ for some $m \in \mathbb{Z}$. Suppose $m = 0$ then $a = (0)d = 0$ which contradicts $a \neq 0$. Therefore, $m \neq 0$. Recall that the absolute value function is multiplicative; $|md| = |m||d|$. As $m \neq 0$ we have $|m| \geq 1$ thus $|a| = |m||d| \geq |d|$. $\square$

I hope you see these proofs are not too hard. You ought to be able to reproduce them without much effort.

Theorem 7.1.8. *Let $a_1, \dots, a_k, c \in \mathbb{Z}$. Then,*

- (i.) *if $c \mid a_i$ for $i = 1, \dots, k$ then $c \mid (u_1a_1 + \dots + u_ka_k)$ for all $u_1, \dots, u_k \in \mathbb{Z}$,*
- (ii.) *$a \mid b$ and $b \mid a$ if and only if $a = \pm b$.*

Proof (i.): suppose $c \mid a_1, c \mid a_2, \dots, c \mid a_k$. It follows there exist $m_1, m_2, \dots, m_k \in \mathbb{Z}$ such that $a_1 = cm_1, a_2 = cm_2$ and $a_k = cm_k$. Let $u_1, u_2, \dots, u_k \in \mathbb{Z}$ and consider,

$$u_1a_1 + \dots + u_ka_k = u_1(cm_1) + \dots + u_k(cm_k) = c(u_1m_1 + \dots + u_km_k).$$

Notice $u_1m_1 + \dots + u_km_k \in \mathbb{Z}$ thus the equation above shows $c \mid (u_1a_1 + \dots + u_ka_k)$.

Proof (ii.): suppose $a \mid b$ and $b \mid a$. If $a = 0$ then $a \mid b$ implies there exists $m \in \mathbb{Z}$ such that $b = m(0) = 0$ hence $b = 0$. Observe $a = \pm b = 0$. Continuing, we suppose $a \neq 0$ which implies $b \neq 0$ by the argument above. Notice $a \mid b$ and $b \mid a$ imply there exist $m, n \in \mathbb{Z} - \{0\}$ such that $a = mb$ and $b = na$. Multiply $a = mb$ by $n \neq 0$ to find $na = mnb$. But, $b = na$ hence $na = mn(na)$ which implies $1 = mn$. Thus, $m = n = 1$ or $m = n = -1$. These cases yield $a = b$ and $a = -b$ respectively hence $a = \pm b$. $\square$

The proof above is really not much more difficult than those we gave for Proposition 7.1.7. The most important case of the Theorem above is when $k = 2$ in part (i.).

Corollary 7.1.9. *If $c \mid x$ and $c \mid y$ then $c \mid (ax + by)$ for all $a, b \in \mathbb{Z}$.*

The result above is used repeatedly as we study the structure of common divisors.

Definition 7.1.10. *If $d \mid a$ and $d \mid b$ then d is a **common divisor** of a and b .*

Proposition 7.1.7 part (iv.) shows that a divisor cannot have a larger magnitude than its multiple. It follows that the largest a common divisor could be is $\max\{|a|, |b|\}$. Furthermore, 1 is a divisor of all nonzero integers. If both a and b are not zero then $\max\{|a|, |b|\} \geq 1$. Therefore, if both a and b are not zero then there must be a largest number between 1 and $\max\{|a|, |b|\}$ which divides both a and b . Thus, the definition to follow is reasonable:

Definition 7.1.11. *If $a, b \in \mathbb{Z}$, not both zero, then the **greatest common divisor** of a and b is denoted $\gcd(a, b)$.*

The method to find the greatest common divisor which served me well as a child was simply to a and b in their prime factorization. Then to find the gcd I just selected all the primes which I could pair in both numbers.

Example 7.1.12.

$$\gcd(105, 90) = \gcd(\underline{3} \cdot \underline{5} \cdot 7, 2 \cdot 3 \cdot \underline{3} \cdot \underline{5}) = 3 \cdot 5 = 15.$$

The method above faces several difficulties as we attempt to solve non-elementary problems.

1. it is not an easy problem to find the prime factorization of a given integer. Indeed, this difficulty is one of the major motivations RSA cryptography.
2. it is not so easy to compare lists and select all the common pairs. Admittedly, this is not as serious a problem, but even with the simple example above I had to double-check.

Thankfully, there is a better method to find the gcd. It's old, but, popular. Euclid (yes, the same one with the parallel lines and all that) gave us the **Euclidean Algorithm**. We prove a Lemma towards developing Euclid's Algorithm.

Lemma 7.1.13. *Let $a, b, q, r \in \mathbb{Z}$. If $a = qb + r$ then $\gcd(a, b) = \gcd(b, r)$.*

Proof: by Corollary 7.1.9 we see a divisor of both b and r is also a divisor of a . Likewise, as $r = a - qb$ we see any common divisor of a and b is also a divisor of r . It follows that a, b and b, r share the same divisors. Hence, $\gcd(a, b) = \gcd(b, r)$. $\square$

We now work towards Euclid's Algorithm. Let $a, b \in \mathbb{Z}$, not both zero. Our goal is to calculate $\gcd(a, b)$. If $a = 0$ and $b \neq 0$ then $\gcd(a, b) = |b|$. Likewise, if $a \neq 0$ and $b = 0$ then $\gcd(a, b) = |a|$. Note $\gcd(a, a) = |a|$ hence we may assume $a \neq b$ in what follows. Furthermore,

$$\gcd(a, b) = \gcd(-a, b) = \gcd(a, -b) = \gcd(-a, -b).$$

Therefore, suppose $a, b \in \mathbb{N}$ with $a > b$ ⁴. Apply the division algorithm (Theorem 7.1.2) to select q_1, r_1 such that

$$a = q_1b + r_1 \quad \text{such that} \quad 0 \leq r_1 < b.$$

If $r_1 = 0$ then $a = q_1b$ hence $b \mid a$ and as b is the largest divisor of b we find $\gcd(a, b) = b$. If $r_1 \neq 0$ then we continue to apply the division algorithm once again to select q_2, r_2 such that

$$b = q_2r_1 + r_2 \quad \text{such that} \quad 0 \leq r_2 < r_1.$$

If $r_2 = 0$ then $r_1 \mid b$ and clearly $\gcd(b, r_1) = r_1$. However, as $a = q_1b + r_1$ allows us to apply Lemma 7.1.13 to obtain $\gcd(a, b) = \gcd(b, r_1) = r_1$. Continuing, we suppose $r_2 \neq 0$ with $r_1 > r_2$ hence we may select q_3, r_3 for which:

$$r_1 = q_3r_2 + r_3 \quad \text{such that} \quad 0 \leq r_3 < r_2.$$

Once again, if $r_3 = 0$ then $r_2 \mid r_1$ hence it is clear $\gcd(r_1, r_2) = r_2$. However, as $b = q_2r_1 + r_2$ gives $\gcd(b, r_1) = \gcd(r_1, r_2)$ and $a = q_1b + r_1$ gives $\gcd(a, b) = \gcd(b, r_1)$ we find that $\gcd(a, b) = r_2$. This process continues. It cannot go on forever as we have the conditions:

$$0 < \cdots < r_3 < r_2 < r_1 < b.$$

There must exist some $n \in \mathbb{N}$ for which $r_{n+1} = 0$ yet $r_n \neq 0$. All together we have:

$$\begin{aligned} a &= q_1b + r_1, \\ b &= q_2r_1 + r_2, \\ r_1 &= q_3r_2 + r_3, \dots, \\ r_{n-2} &= q_nr_{n-1} + r_n, \\ r_{n-1} &= q_{n+1}r_n. \end{aligned}$$

⁴the equation above shows we can cover all other cases once we solve the problem for positive integers.

The last condition yields $r_n \mid r_{n-1}$ hence $\gcd(r_{n-1}, r_n) = r_n$. Furthermore, we find, by repeated application of Lemma 7.1.13 the following string of equalities

$$\gcd(a, b) = \gcd(b, r_1) = \gcd(r_1, r_2) = \gcd(r_2, r_3) = \cdots = \gcd(r_{n-1}, r_n) = r_n.$$

In summary, we have shown that repeated division of remainders into remainder gives a strictly decreasing sequence of positive integers whose last member is precisely $\gcd(a, b)$.

Theorem 7.1.14. Euclidean Algorithm: *suppose $a, b \in \mathbb{N}$ with $a > b$ and form the finite sequence $\{b, r_1, r_2, \dots, r_n\}$ for which $r_{n+1} = 0$ and $b, r_1, \dots, r_n$ are defined as discussed above. Then $\gcd(a, b) = r_n$.*

Example 7.1.15. *Let me show you how the euclidean algorithm works for a simple example. Consider $a = 100$ and $b = 44$. Euclid's algorithm will allow us to find $\gcd(100, 44)$.*

1. $100 = 44(2) + 12$ divided 100 by 44 got remainder of 12
2. $44 = 12(3) + 8$ divided 44 by 12 got remainder of 8
3. $12 = 8(1) + \boxed{4}$ divided 12 by 8 got remainder of 4
4. $8 = 4(2) + 0$ divided 4 by 1 got remainder of zero

The last nonzero remainder will always be the gcd when you play the game we just played. Here we find $\boxed{\gcd(100, 44) = 4}$. Moreover, we can write 4 as a $\mathbb{Z}$ -linear combination of 100 and 44. This can be gleaned from the calculations already presented by working backwards from the gcd:

3. $4 = 12 - 8$
2. $8 = 44 - 12(3)$ implies $4 = 12 - (44 - 12(3)) = 4(12) - 44$
1. $12 = 100 - 44(2)$ implies $4 = 4(100 - 44(2)) - 44 = 4(100) - 9(44)$

I call this a " $\mathbb{Z}$ -linear combination of 100 and 44 since $4, -9 \in \mathbb{Z}$. We find $\boxed{4(100) - 9(44) = 4}$.

The fact that we can always work euclid's algorithm backwards to find how the $\gcd(a, b)$ is written as $ax + by = \gcd(a, b)$ for some $x, y \in \mathbb{Z}$ is remarkable. I continue to showcase this side-benefit of the Euclidean Algorithm as we continue. We will give a general argument after the examples. I now shift to a less verbose presentation:

Example 7.1.16. *Find $\gcd(62, 626)$*

$$626 = 10(62) + 6$$

$$62 = 10(6) + 2$$

$$6 = 3(2) + 0$$

From the E.A. I deduce $\gcd(62, 626) = 2$. Moreover,

$$2 = 62 - 10(6) = 62 - 10[626 - 10(62)] = 101(62) - 10(626)$$

Example 7.1.17. Find $\gcd(240, 11)$.

$$240 = 11(21) + 9$$

$$11 = 9(1) + 2$$

$$9 = 2(4) + 1$$

$$2 = 1(2)$$

Thus, by E.A., $\gcd(240, 11) = 1$. Moreover,

$$1 = 9 - 2(4) = 9 - 4(11 - 9) = -4(11) + 5(9) = -4(11) + 5(240 - 11(21))$$

That is,

$$\boxed{1 = -109(11) + 5(240)}$$

Example 7.1.18. Find $\gcd(4, 20)$. This example is a bit silly, but I include it since it is an exceptional case in the algorithm. The algorithm works, you just need to interpret the instructions correctly.

$$20 = 4(5) + 0$$

Since there is only one row to go from we identify 4 as playing the same role as the last non-zero remainder in most examples. Clearly, $\gcd(4, 20) = 4$. Now, what about working backwards? Since we do not have the gcd appearing by itself in the next to last equation (as we did in the last example) we are forced to solve the given equation for the gcd,

$$20 = 4(4 + 1) = 4(4) + 4 \implies \boxed{20 - 4(4) = 4}$$

The following result also follows from the discussion before Theorem 7.1.14. I continue to use the notational set-up given there.

Theorem 7.1.19. Bezout's Identity: if $a, b \in \mathbb{Z}$, not both zero, then there exist $x, y \in \mathbb{Z}$ such that $ax + by = \gcd(a, b)$.

Proof: we have illustrated the proof in the examples. Basically we just back-substitute the division algorithms. For brevity of exposition, I assume $r_3 = \gcd(a, b)$. It follows that:

$$a = q_1b + r_1 \implies r_1 = a - q_1b$$

$$b = q_2r_1 + r_2 \implies r_2 = b - q_2r_1$$

$$r_1 = q_3r_2 + r_3 \implies r_3 = r_1 - q_3r_2$$

where $\gcd(a, b) = r_3$. Moreover, $r_2 = b - q_2(a - q_1b)$ implies $r_3 = r_1 - q_3[b - q_2(a - q_1b)]$. Therefore,

$$\gcd(a, b) = a - q_1b - q_3[b - q_2(a - q_1b)] = a - (q_1 - q_3[1 - q_2(a - q_1)])b.$$

Identify $x = 1$ and $y = q_1 - q_3[1 - q_2(a - q_1)]$. $\square$

We should appreciate that x, y in the above result are far from unique. However, as we have shown, the method at least suffices to find a solution of the equation $ax + by = \gcd(a, b)$.

7.2 On modular arithmetic and groups

In this section we assume $n \in \mathbb{N}$ throughout. In summary, we develop a careful model for $\mathbb{Z}_n$ in this section.

Remark 7.2.1. I use some notation in this section which we can omit elsewhere for the sake of brevity. In particular, in the middle of this section I might use the notation $[2]$ or $\bar{2}$ for $2 \in \mathbb{Z}_n$ whereas in later work we simply use 2 with the understanding that we are working in the context of modular arithmetic. I have a bit more to say about this notational issue and the deeper group theory it involves at the conclusion of this section.

Definition 7.2.2. $a \equiv b \pmod{n}$ if and only if $n \mid (b - a)$.

The definition above is made convenient by the simple equivalent criteria below:

Theorem 7.2.3. Let $a, b \in \mathbb{Z}$ then we say a is **congruent** to $b \pmod{n}$ and write $a \equiv b \pmod{n}$ if a and b have the same remainder when divided by n .

Proof: Suppose $a \equiv b \pmod{n}$ then a and b share the same remainder after division by n . By the Division Algorithm, there exist $q_1, q_2 \in \mathbb{Z}$ for which $a = q_1n + r$ and $b = q_2n + r$. Observe, $b - a = (q_2n + r) - (q_1n + r) = (q_2 - q_1)n$. Therefore, $n \mid (b - a)$.

Conversely, suppose $n \mid (b - a)$ then there exists $q \in \mathbb{Z}$ for which $b - a = qn$. Apply the Division Algorithm to find q_1, q_2 and r_1, r_2 such that: $a = q_1n + r_1$ and $b = q_2n + r_2$ with $0 \leq r_1 < n$ and $0 \leq r_2 < n$. We should pause to note $|r_2 - r_1| < n$. Observe,

$$b - a = qn = (q_2n + r_2) - (q_1n + r_1) = (q_2 - q_1)n + r_2 - r_1.$$

Therefore, solving for the difference of the remainders and taking the absolute value,

$$|q - q_2 + q_1|n = |r_2 - r_1|$$

Notice $|q - q_2 + q_1| \in \mathbb{N} \cup \{0\}$ and $|r_2 - r_1| < n$. It follows $|q - q_2 + q_1| = 0$ hence $|r_2 - r_1| = 0$ and we conclude $r_1 = r_2$. $\square$

Congruence has properties you might have failed to notice as a child.

Proposition 7.2.4. Let n be a positive integer, for all $x, y, z \in \mathbb{Z}$,

- (i.) $x \equiv x \pmod{n}$,
- (ii.) $x \equiv y \pmod{n}$ implies $y \equiv x \pmod{n}$,
- (iii.) if $x \equiv y \pmod{n}$ and $y \equiv z \pmod{n}$ then $x \equiv z \pmod{n}$.

Proof: we use Definition 7.2.2 throughout what follows.

- (i.) Let $x \in \mathbb{Z}$ then $x - x = 0 = 0 \cdot n$ hence $n \mid (x - x)$ and we find $x \equiv x \pmod{n}$.
- (ii.) Suppose $x \equiv y \pmod{n}$. Observe $n \mid (x - y)$ indicates $x - y = nk$ for some $k \in \mathbb{Z}$. Hence $y - x = n(-k)$ where $-k \in \mathbb{Z}$. Therefore, $n \mid (y - x)$ and we find $y \equiv x \pmod{n}$.
- (iii.) Suppose $x \equiv y \pmod{n}$ and $y \equiv z \pmod{n}$. Thus $n \mid (y - x)$ and $n \mid (z - y)$. Corollary 7.1.9 indicates n also divides the sum of two integers which are each divisible by n . Thus, $n \mid [(y - x) + (z - y)]$ hence $n \mid (z - x)$ which shows $x \equiv z \pmod{n}$. $\square$

I referenced the Corollary to prove part (iii.) to remind you how our current discussion fits naturally with our previous discussion.

Corollary 7.2.5. *Let $n \in \mathbb{N}$. Congruence modulo n forms an equivalence relation on $\mathbb{Z}$.*

This immediately informs us of an interesting **partition** of the integers. Recall, a **partition** of a set S is a family of subsets $U_\alpha \subseteq S$ where $\alpha \in \Lambda$ is some index set such that $U_\alpha \cap U_\beta = \emptyset$ for $\alpha \neq \beta$ and $\cup_{\alpha \in \Lambda} U_\alpha = S$. A partition takes a set and parses it into disjoint pieces which cover the whole set. The partition induced from an equivalence relation is simply formed by the **equivalence classes** of the relation. Let me focus on $\mathbb{Z}$ with the equivalence relation of congruence modulo a positive integer n . We define:⁵:

Definition 7.2.6. equivalence classes of $\mathbb{Z}$ modulo $n \in \mathbb{N}$:

$$[x] = \{y \in \mathbb{Z} \mid y \equiv x \pmod{n}\}$$

Observe, there are several ways to characterize such sets:

$$[x] = \{y \in \mathbb{Z} \mid y \equiv x \pmod{n}\} = \{y \in \mathbb{Z} \mid y - x = nk \text{ for some } k \in \mathbb{Z}\} = \{x + nk \mid k \in \mathbb{Z}\}.$$

I find the last presentation of $[x]$ to be useful in practical computations.

Example 7.2.7. *Congruence $\pmod{2}$ partitions $\mathbb{Z}$ into even and odd integers:*

$$[0] = \{2k \mid k \in \mathbb{Z}\} \quad \& \quad [1] = \{2k + 1 \mid k \in \mathbb{Z}\}$$

Example 7.2.8. *Congruence $\pmod{4}$ partitions $\mathbb{Z}$ into four classes of numbers:*

$$\begin{aligned} [0] &= \{4k \mid k \in \mathbb{Z}\} = \{\dots, -8, -4, 0, 4, 8, \dots\} \\ [1] &= \{4k + 1 \mid k \in \mathbb{Z}\} = \{\dots, -7, -3, 1, 5, 9, \dots\} \\ [2] &= \{4k + 2 \mid k \in \mathbb{Z}\} = \{\dots, -6, -2, 2, 6, 10, \dots\} \\ [3] &= \{4k + 3 \mid k \in \mathbb{Z}\} = \{\dots, -5, -1, 3, 7, 11, \dots\} \end{aligned}$$

The patterns above are interesting, there is something special about $[0]$ and $[2]$ in comparison to $[1]$ and $[3]$. Patterns aside, the notation of the previous two example can be improved. Let me share a natural notation which helps us understand the structure of congruence classes.

Definition 7.2.9. Coset Notation: *Let $n \in \mathbb{N}$ and $a \in \mathbb{Z}$ we define:*

$$n\mathbb{Z} = \{nk \mid k \in \mathbb{Z}\} \quad a + n\mathbb{Z} = \{a + nk \mid k \in \mathbb{Z}\}.$$

Observe, in the notation just introduced, we have

$$\boxed{[a] = a + n\mathbb{Z}}$$

Example 7.2.10. *Congruence $\pmod{2}$ partitions $\mathbb{Z}$ into even and odd integers:*

$$[0] = 2\mathbb{Z} \quad \& \quad [1] = 1 + 2\mathbb{Z}.$$

Example 7.2.11. *Congruence $\pmod{4}$ partitions $\mathbb{Z}$ into four classes of numbers:*

$$[0] = 4\mathbb{Z}, \quad [1] = 1 + 4\mathbb{Z}, \quad [2] = 2 + 4\mathbb{Z}, \quad [3] = 3 + 4\mathbb{Z}.$$

⁵there are other notations, the concept here is far more important than the notation we currently employ

We should pause to appreciate a subtle aspect of the notation. It is crucial to note $[x] = [y]$ does **not** imply $x = y$. For example, modulo 2:

$$[1] = [3] = [7] = [1000037550385987987987971] \quad \& \quad [2] = [-2] = [-42].$$

Or, modulo 9:

$$[1] = [10] = [-8], \quad \& \quad [3] = [12] = [-6], \quad \& \quad [0] = [90] = [-9].$$

Yet, modulo 9, $[1] \neq [3]$. Of course, I just said $[1] = [3]$. How can this be? Well, context matters. In some sense, the notation $[x]$ is dangerous and $[x]_n$ would be better. We could clarify that $[1]_2 = [3]_2$ whereas $[1]_9 \neq [3]_9$. I don't recall such notation used in any text. What is more common is to use the *coset notation* to clarify:

$$1 + 2\mathbb{Z} = 3 + 2\mathbb{Z} \quad \text{whereas} \quad 1 + 9\mathbb{Z} \neq 3 + 9\mathbb{Z}.$$

I'm not entirely sure the Proposition below is necessary.

Proposition 7.2.12. *Let $n \in \mathbb{N}$. We have $[x] = [y]$ if and only if $x \equiv y \pmod{n}$. Or, in the coset notation $x + n\mathbb{Z} = y + n\mathbb{Z}$ if and only if $y - x \in n\mathbb{Z}$.*

Proof: Observe $x \in [x]$. If $[x] = [y]$ then $x \in [y]$ hence there exists $k \in \mathbb{Z}$ for which $x = y + nk$ hence $x - y = nk$ and we find $x \equiv y \pmod{n}$. Conversely, if $x \equiv y \pmod{n}$ then there exists $k \in \mathbb{Z}$ such that $y - x = nk$ thus $x = y - nk$ and $y = x + nk$. Suppose $a \in [x]$ then there exists $j \in \mathbb{Z}$ for which $a = nj + x$ hence $a = nj + y - nk = n(j - k) + y \in [y]$. We have shown $[x] \subseteq [y]$. Likewise, if $b \in [y]$ then there exists $j \in \mathbb{Z}$ for which $b = nj + y$ hence $b = nj + x + nk = n(j + k) + x \in [x]$. Thus $[y] \subseteq [x]$ and we conclude $[x] = [y]$. $\square$

Notice the proposition above allows us to calculate as follows: for $n \in \mathbb{N}$

$$na + b + n\mathbb{Z} = b + n\mathbb{Z} \quad \text{or} \quad [na + b] = [b]$$

for $a, b \in \mathbb{Z}$. There is more.

Proposition 7.2.13. *Let $n \in \mathbb{N}$. If $[x] = [x']$ and $[y] = [y']$ then*

$$(i.) \quad [x + y] = [x' + y'],$$

$$(ii.) \quad [xy] = [x'y']$$

$$(iii.) \quad [x - y] = [x' - y']$$

Proof: Suppose $[x] = [x']$ and $[y] = [y']$. It follows there exists $j, k \in \mathbb{Z}$ such that $x' = nj + x$ and $y' = nk + y$. Notice $x' \pm y' = nj + x \pm (nk + y) = n(j \pm k) + x \pm y$. Therefore, $x \pm y \equiv x' \pm y' \pmod{n}$ and by Proposition 7.2.12 we find $[x \pm y] = [x' \pm y']$. This proves (i.) and (iii.). Next, consider:

$$x'y' = (nj + x)(nk + y) = n(jkn + jy + xk) + xy$$

thus $x'y' \equiv xy \pmod{n}$ we apply Proposition 7.2.12 once more to find $[xy] = [x'y']$. $\square$

We ought to appreciate the content of the proposition above as it applies to congruence modulo n . In fact, the assertions below all appear in the proof above.

Corollary 7.2.14. *Let $n \in \mathbb{N}$. If $x \equiv x'$ and $y \equiv y'$ modulo n then*

- (i.) $x + y \equiv x' + y' \pmod{n}$,
- (ii.) $xy \equiv x'y' \pmod{n}$,
- (iii.) $x - y \equiv x' - y' \pmod{n}$,

Example 7.2.15. *Suppose $x + y \equiv 3$ and $x - y \equiv 1$ modulo 4. Then, by Corollary 7.2.14 we add and subtract the given congruences to obtain:*

$$2x \equiv 4 \quad 2y \equiv 2$$

There are 4 cases to consider. Either $x \in [0]$, $x \in [1]$, $x \in [2]$ or $x \in [3]$. Observe,

$$\begin{array}{ll} 2(0) \equiv 0 \equiv 4, & 2(0) \not\equiv 2 \\ 2(1) \equiv 2 \not\equiv 4, & 2(1) \equiv 2 \\ 2(2) \equiv 4, & 2(2) \equiv 4 \not\equiv 2 \\ 2(3) \equiv 2 \not\equiv 4, & 2(3) \equiv 2. \end{array}$$

It follows that $x \in [0] \cup [2]$ and $y \in [1] \cup [3]$ forms the solution set of this system of congruences.

The method I used to solve the above example was not too hard since there were just 4 cases to consider. I suppose, if we wished to solve the same problem modulo 42 we probably would like to learn a better method.

Proposition 7.2.13 justifies that the definition below does give a **binary operation** on the set of equivalence classes modulo n . Recall, a *binary operation* on a set S is simply a *function* from $S \times S$ to S . It is a single-valued assignment of pairs of S -elements to S -elements.

Definition 7.2.16. modular arithmetic: *let $n \in \mathbb{N}$, define*

$$[x] + [y] = [x + y] \quad \& \quad [x][y] = [xy]$$

for all $x, y \in \mathbb{Z}$. Or, if we denote the set of all equivalence classes modulo n by $\mathbb{Z}/n\mathbb{Z}$ then write: for each $x + n\mathbb{Z}, y + n\mathbb{Z} \in \mathbb{Z}/n\mathbb{Z}$

$$(x + n\mathbb{Z}) + (y + n\mathbb{Z}) = x + y + n\mathbb{Z} \quad \& \quad (x + n\mathbb{Z})(y + n\mathbb{Z}) = xy + n\mathbb{Z}.$$

Finally, we often use the notation $\mathbb{Z}_n = \mathbb{Z}/n\mathbb{Z}$.

Notice the operation defined above is a binary operation on $\mathbb{Z}/n\mathbb{Z}$ (not $\mathbb{Z}$). Many properties of integer arithmetic transfer to $\mathbb{Z}/n\mathbb{Z}$:

$$\begin{aligned} [a] + [b] &= [b] + [a] \\ [a][b] &= [b][a] \\ [a]([b] + [c]) &= [a][b] + [a][c] \\ ([a] + [b])[c] &= [a][c] + [b][c] \\ ([a] + [b]) + [c] &= [a] + ([b] + [c]) \\ ([a][b])[c] &= [a]([b][c]) \\ [a] + [0] &= [0] + [a] = [a] \\ [1][a] &= [a][1]. \end{aligned}$$

Furthermore, for $k \in \mathbb{N}$,

$$\begin{aligned}[a_1] + [a_2] + \cdots + [a_k] &= [a_1 + a_2 + \cdots + a_k] \\ [a_1][a_2] \cdots [a_k] &= [a_1 a_2 \cdots a_k] \\ [a]^k &= [a^k].\end{aligned}$$

Example 7.2.17. Simplify $[1234]$ modulo 5. Notice,

$$1234 = 1 \times 10^3 + 2 \times 10^2 + 3 \times 10 + 4.$$

However, $10 = 2(5)$ thus,

$$1234 = 1 \times 2^3 5^3 + 2 \times 2^2 5^2 + 3 \times 2 \cdot 5 + 4.$$

Note, $[5] = [0]$ hence $[5^k] = [0]$ for $k \in \mathbb{N}$. By the properties of modular arithmetic it is clear that the 10's, 100's and 1000's digits are irrelevant to the result. Only the first digit matters, $[1234] = [4]$.

It is not hard to see the result of the example above equally well applies to larger numbers; if $a_k, a_{k-1}, \dots, a_2, a_1$ are the digits in a decimal representation of an integer then $[a_k a_{k-1} \cdots a_2 a_1] = [a_1] \bmod(5)$.

Example 7.2.18. Calculate the cube of 51 modulo 7.

$$[51^3] = [51][51][51] = [51]^3 = [49 + 2]^3 = [2]^3 = [8].$$

Of course, you can also denote the same calculation via congruence:

$$51^3 = 51 \cdot 51 \cdot 51 \equiv 2 \cdot 2 \cdot 2 = 8 \Rightarrow [51^3] = [8].$$

The next example is a cautionary tale:

Example 7.2.19. Simplify 7^{100} modulo 6. Consider,

$$[7^{100}] = [7]^{100} = [1]^{100} = [1^{100}] = [1].$$

or, (incorrectly !)

$$[7^{100}] = [7^{[100]}] = [7^{6(16)+4}] = [7^4] = [28] = [4].$$

The point is this: it is **not** true that $[a^k] = [a^{[k]}]$.

Naturally, as we discuss $\mathbb{Z}_n$ it is convenient to have a particular choice of representative for this set of residues. Two main choices: the *set of least non-negative residues*

$$\mathbb{Z}_n = \{[0], [1], [2], \dots, [n-1]\}$$

alternatively, *set of least absolute value residues* or simply *least absolute residues*

$$\mathbb{Z}_n = \{[0], [\pm 1], [\pm 2], \dots\}$$

where the details depend on if n is even or odd. For example,

$$\mathbb{Z}_5 = \{[0], [1], [2], [3], [4]\} = \{[-2], [-1], [0], [1], [2]\}$$

or,

$$\mathbb{Z}_4 = \{[0], [1], [2], [3]\} = \{[-2], [-1], [0], [1]\}$$

Honestly, if we work in the particular context of $\mathbb{Z}_n$ then there is not much harm in dropping the $[\cdot]$ -notation. Sometimes, I use $[x] = \bar{x}$. Whichever notation we choose, we must be careful to not fall into the trap of assuming the usual properties of $\mathbb{Z}$ when calculating in the specific context of modular arithmetic. The example that follows would be very clumsy to write in the $[\cdot]$ -notation.

Definition 7.2.20. group of units: let $n \in \mathbb{N}$ define $U(n) \subseteq \mathbb{Z}_n$ by

$$U(n) = \{x \in \mathbb{Z}_n \mid \text{there exists } y \in \mathbb{Z}_n \text{ for which } xy = 1\}$$

In other words, each element of $U(n)$ has a **multiplicative inverse**.

A well-known theorem gives a simple method to ascertain if $[x] \in U(n)$; simply this, $[x] \in U(n)$ if and only if $\gcd(x, n) = 1$. This theorem is nearly immediate from Bezout's Theorem.

Example 7.2.21. In Example 7.1.16 we found $\gcd(62, 626) = 2$. This shows 62 does not have a multiplicative inverse modulo 626. Also, it shows 626 does not have a multiplicative inverse modulo 62.

Example 7.2.22. In Example 7.1.17 we found $\gcd(11, 240) = 1$ and $1 = -109(11) + 5(240)$. From this we may read several things:

$$[-109]^{-1} = [11] \text{ mod}(240) \quad \& \quad [-109]^{-1} = [11] \text{ mod}(5)$$

and,

$$[5]^{-1} = [240] \text{ mod}(11) \quad \& \quad [5]^{-1} = [240] \text{ mod}(109).$$

In terms of least positive residues the last statement reduces to $[5]^{-1} = [22]$. Of course, we can check this; $[5][22] = [110] = [1]$.

Remark 7.2.23. At this point our work on the model $\mathbb{Z}/n\mathbb{Z}$ for $\mathbb{Z}_n$ comes to an end. From this point forward, we return to the less burdensome notation

$$\mathbb{Z}_n = \{0, 1, 2, \dots, n-1\}$$

as a default. However, we are open-minded, if you wish, you can define

$$\mathbb{Z}_n = \{1, 2, \dots, n\}$$

where n serves as the additive identity of the group. Furthermore, we wish to allow calculations such as: working modulo 5 we have:

$$3(17) = 3(15 + 2) = 3(2) = 6 = 1$$

thus $17^{-1} = 2$ or $3^{-1} = 2$ etc. Technically, if we define $G_1, G_2, G_3 \subseteq \mathbb{Z}$ where

$$G_1 = \{0, 1, 2\}, \quad G_2 = \{1, 2, 3\}, \quad G_3 = \{10, 11, 12\}$$

and addition is defined modulo 3 then G_1, G_2 and G_3 are distinct **point sets** and hence are different groups. However, these are all models of $\mathbb{Z}_3$. In fact, G_1, G_2, G_3 are all **isomorphic**⁶. Algebraists will say things like, all of the sets G_1, G_2, G_3 are $\mathbb{Z}_3$. What they mean by that is that these set are all to the intuition of a group theorist the same thing. What we define $\mathbb{Z}_n$ to be is largely a matter of convenience. Again, the two main models:

- (1.) $\mathbb{Z}_n = \mathbb{Z}/n\mathbb{Z} = \{k + n\mathbb{Z} \mid k \in \mathbb{Z}\}$ makes $\mathbb{Z}_n$ a set of sets of integers, or, a set of cosets of $\mathbb{Z}$.
- (2.) $\mathbb{Z}_n = \{0, 1, \dots, n-1\}$ where $\mathbb{Z}_n \subset \mathbb{Z}$.

In either case, $\mathbb{Z}_n$ is not a subgroup of $\mathbb{Z}$, but, for slightly different reasons. In case (1.) the set of cosets is not a subset of $\mathbb{Z}$ so it fails the subset criterion for subgroup. In case (2.) while it is a subset it fails to have the same additive operation as $\mathbb{Z}$.

⁶we will define this carefully in due time, for Fall 2016, this is after Test 1

7.3 matrices of modular integers

Matrices with entries in $\mathbb{Z}_n$ are multiplied and added in the usual fashion. In particular,

$$(A + B)_{ij} = A_{ij} + B_{ij}, \quad (cA)_{ij} = cA_{ij}, \quad (XY)_{ij} = \sum_{k=1}^r X_{ik}Y_{kj}$$

where $A, B \in \mathbb{Z}_n^{p \times q}$, $X \in \mathbb{Z}_n^{p \times r}$ and $Y \in \mathbb{Z}_n^{r \times q}$. We can show $I_{ij} = \delta_{ij}$ has $XI = IX = X$ for any matrix. Naturally, the addition and multiplications above are all done modulo n . This has some curious side-effects:

$$\underbrace{A + A + \cdots + A}_{n\text{-summands}} = nA = 0$$

A square $M \in \mathbb{Z}_n^{p \times p}$ is invertible only if $\det(M) \in U(n)$. To prove this we could go through all the usual linear algebra simply replacing regular addition and multiplication with the modular equivalent. In particular, we can show the classical adjoint of M satisfies

$$M \operatorname{adj}(M) = \det(M)I$$

If $\det(M) \in U(n)$ then there exists $\det(M)^{-1} \in U(n)$ for which $\det(M)^{-1} \cdot \det(M) = 1$. Multiplying the classical adjoint equation we derive

$$M(\det(M)^{-1} \cdot \operatorname{adj}(M)) = \det(M)^{-1} \cdot \det(M)I = I.$$

Thus,

$$\boxed{M^{-1} = \det(M)^{-1} \cdot \operatorname{adj}(M).}$$

Calculation of the inverse in 3×3 or larger cases requires some calculation, but, for the 2×2 case we have a simple formula: if $\begin{bmatrix} a & b \\ c & d \end{bmatrix} \in \mathbb{Z}_n^{2 \times 2}$ and $ad - bc \in U(n)$ then

$$\boxed{\begin{bmatrix} a & b \\ c & d \end{bmatrix}^{-1} = (ad - bc)^{-1} \begin{bmatrix} d & -b \\ -c & a \end{bmatrix}}$$

In the case $n = p$ a prime, $U(p) = \mathbb{Z}_p^\times$ and the inverse of a matrix M over $\mathbb{Z}_p$ exists whenever $\det(M) \neq 0$. In fact, we can define the general linear group over matrices even when the entries are not taken from a field.

Definition 7.3.1. *The general linear group of $p \times p$ matrices over $\mathbb{Z}_n$ is defined by:*

$$GL(p, \mathbb{Z}_n) = \{A \in \mathbb{Z}_n^{p \times p} \mid \det(A) \in U(n)\}.$$

Moreover, $GL(p, \mathbb{Z}) = \{A \in \mathbb{Z}^{p \times p} \mid \det(A) = \pm 1\}$.

I will forego proof that the general linear groups are indeed groups at the moment. In fact, we can define the general linear group for matrices built over any ring in a similar fashion. These suffice for our current purposes. (my apologies as I have yet to define *group* in this course, rest assured, we'll not miss that in Math 421).

Chapter 8

Appendix on Sets and Functions

This brief appendix collects the basic definitions for set theory as well as set-theoretic constructions involving functions. Ideally these concepts would be explored in the introduction to proofs course. Probably much of this was covered, but I include it here for review. I'll begin with basic set theoretic definitions then we'll proceed to discuss how functions and sets interact.

8.1 set theory

Definition 8.1.1.

We write $x \in S$ to mean x is an element of the set S . A set is often described using **roster notation**

$$S = \{x \mid \text{condition on } x\}$$

We also use notation $S = \{x \in U \mid \text{condition on } x\}$ to indicate that $x \in S$ has $x \in U$ subject to the given condition. Equality of sets is a basic concept:

Definition 8.1.2.

Let A and B be sets. We say $A = B$ if and only if A and B have the same elements. In other words, $x \in A$ implies $x \in B$ and $x \in B$ implies $x \in A$.

Likewise the concept of subset and superset is of central importance to the application of the set-concept to mathematics:

Definition 8.1.3.

Let B be a set then we write $A \subseteq B$ if for any $x \in A$ we have $x \in B$. In this case we say that A is a **subset** of B and that B is a **superset** of A . Moreover, we say B **contains** A .

The set containing no elements is known as the **empty set** we denote it by $\{\}$ or $\emptyset$. Notice it is vacuously true that $\emptyset \subseteq A$ for any set A . If $A \subseteq B$ and $A \neq B$ then A is said to be a **proper** subset of B . Notice that set equality can also be understood in terms of subsets; $A = B$ if and only if $A \subseteq B$ and $B \subseteq A$. Often we prove equality of sets by establishing both the containment $A \subseteq B$ and $B \subseteq A$ (so-called **double containment**).

Definition 8.1.4.

Suppose A and B are sets then we define $A \cup B = \{x \mid x \in A \text{ or } x \in B\}$ and we define $A \cap B = \{x \mid x \in A \text{ and } x \in B\}$. We say $A \cup B$ is the **union** of A and B whereas $A \cap B$ is the **intersection** of A and B .

Notice that $A \cup \emptyset = A$ and $A \cap \emptyset = \emptyset$ for any set A . Extending unions and intersections to families of sets which are indexed by some set is also interesting:

Definition 8.1.5.

If Λ is a set and A_α is a set for each $\alpha \in \Lambda$ then we define the **union** of A_α over Λ by

$$\bigcup_{\alpha \in \Lambda} A_\alpha = \{x \mid \text{there exists } \alpha \in \Lambda \text{ for which } x \in A_\alpha \}.$$

Likewise, the **intersection** of A_α over Λ is given by

$$\bigcap_{\alpha \in \Lambda} A_\alpha = \{x \mid x \in A_\alpha \text{ for each } \alpha \in \Lambda \}.$$

In the case our index set $\Lambda = \mathbb{N}$ then the notations

$$\bigcup_{i=1}^{\infty} A_i \quad \& \quad \bigcap_{i=1}^{\infty} A_i$$

can be used for the union and intersection over a **countable**¹ collection of sets $A_1, A_2, \dots$

Definition 8.1.6.

If A and B are sets then $A - B = \{x \in A \mid x \notin B\}$. We say $A - B$ is the **complement** of B in A or, as the notation suggests, the **set difference** of A by B .

Notice $A - A = \emptyset$ and $A - \emptyset = A$ and $\emptyset - A = \emptyset$ for any set A . In some sense the empty set behaves like zero, but the analogy cannot be stretched too far².

There are many properties of sets to which we should be aware. I'll forego most proofs here in the interest of saving these as exercises.

Proposition 8.1.7. *Let A, B, C be sets then we have*

1. *set equality by double containment: $A \subseteq B$ and $B \subseteq A$ if and only if $A = B$,*
2. *commutative properties: $A \cup B = B \cup A$ and $A \cap B = B \cap A$,*
3. *associative properties: $(A \cup B) \cup C = A \cup (B \cup C)$ and $(A \cap B) \cap C = A \cap (B \cap C)$,*
4. *distributive properties: $A \cup (B \cap C) = (A \cup B) \cap (A \cup C)$ and $A \cap (B \cup C) = (A \cap B) \cup (A \cap C)$,*
5. *idempotent laws: $A \cup A = A$ and $A \cap A = A$,*

¹more on what this means a bit later in this appendix

²I will not attempt it here, but it is possible to build a model of the natural numbers $\mathbb{N}$ via a construction with the empty set and sets of the empty set.

6. *laws of absorption:* $A \cup (A \cap B) = A$ and $A \cap (A \cup B) = A$,

7. $A \cap B = A - (A - B)$,

8. *transitivity of inclusion:* if $A \subseteq B$ and $B \subseteq C$ then $A \subseteq C$

Proof: suppose $x \in A \cup B$ then $x \in A$ or $x \in B$ by definition of $A \cup B$. Hence $x \in B$ or $x \in A$ and so $x \in B \cup A$ by definition of union. Thus $A \cup B \subseteq B \cup A$. Likewise, if $x \in B \cup A$ then $x \in B$ or $x \in A$ by definition of $A \cup B$. Hence $x \in A$ or $x \in B$ and so $x \in A \cup B$ by definition of union. Thus $B \cup A \subseteq A \cup B$. Therefore, $A \cup B = B \cup A$ by (1.) of this proposition. Proofs of other parts left for homework. $\square$

If U is the **universal set**³ and $X \subseteq U$,

$$\overline{X} = U - X$$

De Morgan's Laws relate complements of unions and intersections. If A, B are sets in a common universe of discourse (meaning $A, B \subseteq U$ and $\overline{A} = U - A$ etc.) then

$$\overline{A \cup B} = \overline{A} \cap \overline{B} \quad \& \quad \overline{A \cap B} = \overline{A} \cup \overline{B}$$

Likewise, if $A_\alpha, B_\alpha \subseteq U$ where $\alpha \in \Lambda$ then De Morgan's Law's naturally extend:

$$\overline{\bigcup_{\alpha \in \Lambda} A_\alpha} = \bigcap_{\alpha \in \Lambda} \overline{A_\alpha} \quad \& \quad \overline{\bigcap_{\alpha \in \Lambda} B_\alpha} = \bigcup_{\alpha \in \Lambda} \overline{B_\alpha}$$

To be explicit, in terms of set-difference we would express De Morgan's Laws as:

$$U - \bigcup_{\alpha \in \Lambda} A_\alpha = \bigcap_{\alpha \in \Lambda} (U - A_\alpha) \quad \& \quad U - \bigcap_{\alpha \in \Lambda} B_\alpha = \bigcup_{\alpha \in \Lambda} (U - B_\alpha).$$

8.2 functions

I'll begin with the definition of a **relation** since it may be helpful to place the concept of a function in a larger context.

Definition 8.2.1.

Let A and B be sets. Then a **relation** from A to B is a particular subset $r \subseteq A \times B$. We write $y = r(x)$ to mean $(x, y) \in r$. The set A is called the **domain** of r and the set B is called the **codomain** of r . When $A = B$ then we simply say r is a **relation on A** .

Notice a given $x \in A$ could map to many pairs in r .

Example 8.2.2. If $r = \{(1, 1), (1, 2), (1, 3), (2, 2), (3, 3)\}$ describes a relation r then $r(1) = 1$ and $r(1) = 2$ and $r(3) = 3$ whereas $r(2) = 2$ and $r(3) = 3$. It would probably be better to write $r(1) = 1, 2, 3$ as to avoid errors which stem from conflating the different values of 1 under r .

The behavior of the relation in the example above is annoying. To avoid such pathology we introduce a special kind of relation: *the function*

³this is a matter of convention and context

Definition 8.2.3.

Let A and B be sets. Then a **function** from A to B is a relation from A to B for which $(x, y_1), (x, y_2) \in f$ then $y_1 = y_2$. We write $y = f(x)$ to mean $(x, y) \in f$. The set A is called the **domain** of f and the set B is called the **codomain** of f . The notation $f : A \rightarrow B$ indicates that f is a function with domain A and codomain B . We also call a function a **map** and the notation $x \mapsto y$ implicitly indicates a function f which **maps** x to y ($y = f(x)$).

Notice the notation $y = f(x)$ is no longer ambiguous for a function f since the **output** $f(x)$ from the **input** x is a **single-value**. In other words, a function is a **single-valued** relation. Sometimes in higher mathematics or engineering you will find literature which refers to a **multiply-valued function**. That would seem to be a contradiction in terms, and I suppose technically it is. But, in defense of such speech I would point out that the study of such object predates the pedantic clarity of the modern function.

Example 8.2.4. Let $r = \{(x, y) \in \mathbb{R}^2 \mid y^2 = x\} = \{(x, y) \mid y = \pm\sqrt{x} \ x \in [0, \infty)\}$. This is not a function since most inputs produce two distinct outputs of $\sqrt{x}$ and $-\sqrt{x}$.

Definition 8.2.5.

Let A, B be sets and $f : A \rightarrow B$. We define the **image** of $S \subseteq A$ under f by:

$$f(S) = \{f(x) \mid x \in S\}.$$

If $f(A) = B$ then we say f is an **onto** or **surjective** function. We call $f(A)$ the **range** or **image** of f . Likewise, we define the inverse image of $T \subseteq B$ by:

$$f^{-1}(T) = \{x \in A \mid f(x) \in T\}.$$

If $f(x) = f(y)$ implies $x = y$ for any $x, y \in A$ then f is a **one-to-one** or **injective** function. If f is both surjective and injective then we say f is **bijective** or a **one-to-one correspondence** of A and B .

A **singleton** is a set which contains a single element such as $\{a\}$. We can characterize injectivity with inverse images of singletons. Suppose $f : A \rightarrow B$ is a function and $b \in B$. By definition:

$$f^{-1}\{b\} = \{x \in A \mid f(x) \in \{b\}\} = \{x \in A \mid f(x) = b\}$$

It may be the case that $f^{-1}\{b\} = \emptyset$. However, if $x, y \in f^{-1}\{b\}$ then $f(x) = b$ and $f(y) = b$ hence $f(x) = f(y)$. Therefore, if f is a one-to-one function then the inverse image of a singleton is either $\emptyset$ or a singleton.

Definition 8.2.6.

If $f : A \rightarrow B$ and $b \in B$ then $f^{-1}\{b\}$ is the **inverse image** or **fiber** of f over b .

Proposition 8.2.7. Suppose A, B are sets and $f : A \rightarrow B$.

1. f is an injection if every non-empty fiber is a singleton
2. if for any $x, y \in A$ we define $x \sim y$ whenever $f(x) = f(y)$ then $\sim$ defines an equivalence relation on A whose equivalence classes are the non-empty fibers of f .

Proof: the proof of (1.) was given in the discussion above the definition of fiber. To prove (2.) let us recall the definition of equivalence relation. An equivalence relation on A is a relation on A which is **reflexive**, **symmetric** and **transitive**. Since $f(x) = f(x)$ for each $x \in A$ we find $x \sim x$ for each $x \in A$ hence $\sim$ is reflexive. If $x \sim y$ then $f(x) = f(y)$ thus $f(y) = f(x)$ and $y \sim x$ so $\sim$ is symmetric. Likewise, if $x \sim y$ and $y \sim z$ then $f(x) = f(y) = f(z)$ thus $x \sim z$ and we deduce $\sim$ is transitive. Finally, the equivalence class under $\sim$ which contains $x \in A$ can be denoted $[x]$ and

$$[x] = \{y \in A \mid y \sim x\} = \{y \in A \mid f(y) = f(x)\}$$

Note $y \in f^{-1}\{f(x)\}$ means $f(y) \in \{f(x)\}$ thus $f(y) = f(x)$. Thus

$$[x] = f^{-1}\{f(x)\}.$$

Definition 8.2.8.

Let A, B, C be sets and $f : A \rightarrow B$ a surjection and $g : B \rightarrow C$ a function. Then the **composite** of f and g is denoted $g \circ f : A \rightarrow C$ and is defined by $(g \circ f)(x) = g(f(x))$ for each $x \in A$.

Let A be a set, we denote the **identity function on A** by $Id_A : A \rightarrow A$ and we define $Id_A(x) = x$ for each $x \in A$ ⁴.

Definition 8.2.9.

Let $f : A \rightarrow B$ be a bijection then $f^{-1} : B \rightarrow A$ is the **inverse function** of f defined by $f(x) = y$ if and only if $f^{-1}(y) = x$.

Notice $f \circ f^{-1} : B \rightarrow B$ and $f^{-1} \circ f : A \rightarrow A$. In fact, given $y = f(x)$ for some $x \in A$ note

$$(f \circ f^{-1})(y) = f(f^{-1}(y)) = f(x) = y$$

hence $f \circ f^{-1} = Id_B$. Likewise, if $y = f(x)$ then

$$(f^{-1} \circ f)(x) = f^{-1}(f(x)) = f^{-1}(y) = x$$

hence $f^{-1} \circ f = Id_A$. Often a function is not invertible unless we modify its domain.

Definition 8.2.10.

Let A, B be sets and $f : A \rightarrow B$. If $U \subseteq A$ then $f|_U : U \rightarrow B$ is the **restriction** of f to U and we define $f|_U(x) = f(x)$ for each $x \in U$. If $f|_U$ is injective then the **local inverse** of f on U is the inverse function of $g : U \rightarrow f(U)$.

Notice we replaced B with $f(U)$ so that g would necessarily be both an injection and a surjection given the condition that $f|_U$ is one-to-one. Notable examples of the local inverse construction include the even root functions $x \mapsto \sqrt[n]{x}$ which serve as local inverses for $x \mapsto x^{2n}$ for any $n \in \mathbb{N}$. The restriction for the even power functions is customarily assigned to be the non-negative real numbers. Other examples include the inverse trigonometric functions. For instance, $\tan^{-1}$ is the local inverse of $\tan$ with respect to $U = (-\pi/2, \pi/2)$. Likewise, $\sin^{-1}$ is the local inverse of $\sin$ with respect to $U = [-\pi/2, \pi/2]$. Also, $\cos^{-1}$ is the local inverse of $\cos$ with respect to $U = [0, \pi]$.

⁴It is unfortunate that the concept of identity has become far more complicated in our modern degenerate society. You would be driven to believe identity has little to do with reality. In math, the identity function of A is unabiguously specified by A , no matter what the life experience of A might be.

In contrast, $\cosh^{-1}$ is the local inverse of $\cosh$ with respect to $U = [0, \infty)$.

To construct a local inverse we must somehow select a subset of the domain of the given function which includes at most one element in each fiber. This is always possible thanks to the axiom of choice. We can always select some subset of the domain on which just one element of each fiber is found. Consequently, it is always possible (in-principle at least) to construct a local inverse of a given function.

Something beautiful happens in Linear Algebra as we study functions on vector spaces which are linear⁵. It turns out that the fibers of a given linear transformation all have the same size and that is determined by what is known as the kernel of the map. When we construct the quotient space by dividing by the kernel this brings the representatives of the quotient vector space in one-to-one correspondence with the fibers of the linear map. It follows we are able to create a bijection by replacing the domain of the linear map with the quotient space by the kernel. This induced map is related, but different, than the local inverse construction. As you study various branches of mathematics this story plays out again and again in different contexts as we study bijections in the context of the given branch.

Once more let us state a few results without proof:

Proposition 8.2.11. *Suppose A, B are sets and $f : A \rightarrow B$ and $A_1, A_2 \subseteq A$ whereas $B_1, B_2 \subseteq B$*

1. $f(A_1 \cap A_2) \subseteq f(A_1) \cap f(A_2)$,
2. $f(A_1 \cup A_2) = f(A_1) \cup f(A_2)$,
3. $f^{-1}(B_1 \cup B_2) = f^{-1}(B_1) \cup f^{-1}(B_2)$,
4. $f^{-1}(B_1 \cap B_2) = f^{-1}(B_1) \cap f^{-1}(B_2)$,
5. $f(A) - f(A_1) \subseteq f(A - A_1)$,
6. $f^{-1}(B_2 - B_1) = f^{-1}(B_2) - f^{-1}(B_1)$.

Proof: I'll give part of the proof of (6.). If $x \in f^{-1}(B_2 - B_1)$ then $f(x) \in B_2 - B_1$. Thus $f(x) \in B_2$ and $f(x) \notin B_1$. Therefore, $x \in f^{-1}(B_2)$ and $x \notin f^{-1}(B_1)$. Hence $x \in f^{-1}(B_2) - f^{-1}(B_1)$ and we find $f^{-1}(B_2 - B_1) \subseteq f^{-1}(B_2) - f^{-1}(B_1)$. It remains to show the reverse inclusion to establish the proof of (6.). I leave proof of the remaining claims as exercises. $\square$

It might be worthwhile to give an example which shows why equality is not found in (1.) for arbitrary f . Let $A_1 = [0, \infty)$ and $A_2 = (-\infty, 0)$ then $A_1 \cap A_2 = \emptyset$. Let $f(x) = x^2$ for all $x \in \mathbb{R}$ then $f(A_1) = [0, \infty)$ and $f(A_2) = (0, \infty)$ hence $f(A_1) \cap f(A_2) = (0, \infty)$. However, $f(A_1 \cap A_2) = f(\emptyset) = \emptyset$ hence $f(A_1) \cap f(A_2) \not\subseteq f(A_1 \cap A_2)$.

⁵worry not if this paragraph at first does not make sense, come back to it later