

Chapter 10

Applications

10.1 Graphs and Networks

Over the years I have seen one model so often, and I found it so basic and useful, that I always put it first. The model consists of **nodes connected by edges**. This is called a **graph**.

Graphs of the usual kind display functions $f(x)$. Graphs of this node-edge kind lead to matrices. This section is about the **incidence matrix** of a graph—which tells how the n nodes are connected by the m edges. Normally $m > n$, there are more edges than nodes.

For any m by n matrix there are two fundamental subspaces in $\mathbf{R}^n$ and two in $\mathbf{R}^m$. They are the row spaces and nullspaces of A and A^T . Their *dimensions* $r, n - r$ and $r, m - r$ come from the most important theorem in linear algebra. The second part of that theorem is the *orthogonality* of the row space and nullspace. Our goal is to show how examples from graphs illuminate this Fundamental Theorem of Linear Algebra.

When I construct a **graph** and its **incidence matrix**, the subspace dimensions will be easy to discover. But we want the subspaces themselves—and orthogonality helps. It is essential to connect the subspaces to the graph they come from. By specializing to incidence matrices, **the laws of linear algebra become Kirchhoff's laws**. Please don't be put off by the words "current" and "voltage." These rectangular matrices are the best.

Every entry of an incidence matrix is 0 or 1 or -1 . This continues to hold during elimination. All pivots and multipliers are ± 1 . Therefore both factors in $A = LU$ also contain 0, 1, -1 . So do the nullspace matrices! All four subspaces have basis vectors with these exceptionally simple components. The matrices are not concocted for a textbook, they come from a model that is absolutely essential in pure and applied mathematics.

The Incidence Matrix

Figure 10.1 displays a graph with $m = 6$ edges and $n = 4$ nodes. The 6 by 4 matrix A tells which nodes are connected by which edges. The first row $-1, 1, 0, 0$ shows that the first edge goes *from node 1 to node 2* (-1 for node 1 because the arrow goes out, $+1$ for node 2 with arrow in).

Row numbers in A are edge numbers, column numbers 1, 2, 3, 4 are node numbers!

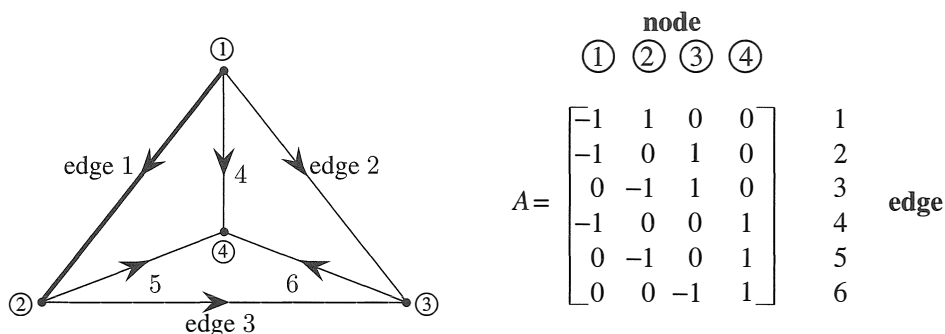


Figure 10.1: Complete graph with $m=6$ edges and $n=4$ nodes: 6 by 4 incidence matrix A .

You can write down the matrix by looking at the graph. The second graph has the same four nodes but only three edges. Its incidence matrix B is 3 by 4.

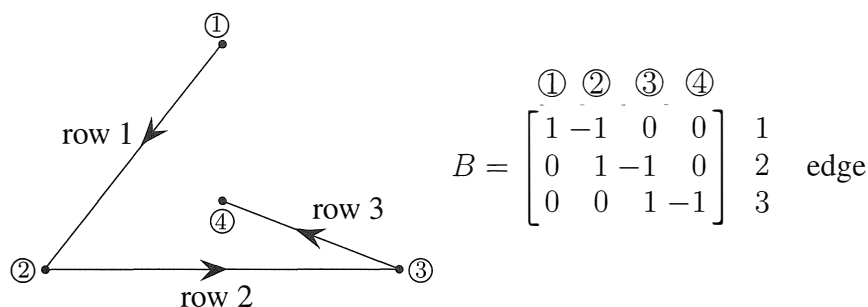


Figure 10.1*: Tree with 3 edges and 4 nodes and no loops. Then B has independent rows.

The first graph is *complete*—every pair of nodes is connected by an edge. The second graph is a *tree*—the graph has *no closed loops*. Those are the two extremes. The maximum number of edges is $\frac{1}{2}n(n-1) = 6$ and the minimum to stay connected is $n-1 = 3$.

Elimination reduces every graph to a tree. Loops produce dependent rows in A and zero rows in the echelon forms U and R . Look at the large loop from edges 1, 2, 3 in the first graph, which leads to a zero row in U :

$$\begin{bmatrix} -1 & 1 & 0 & 0 \\ -1 & 0 & 1 & 0 \\ 0 & -1 & 1 & 0 \end{bmatrix} \rightarrow \begin{bmatrix} -1 & 1 & 0 & 0 \\ 0 & -1 & 1 & 0 \\ 0 & -1 & 1 & 0 \end{bmatrix} \rightarrow \begin{bmatrix} -1 & 1 & 0 & 0 \\ 0 & -1 & 1 & 0 \\ 0 & 0 & 0 & 0 \end{bmatrix}$$

Those steps are typical. When edges 1 and 2 share node 1, elimination produces the “short-cut edge” without node 1. If the graph already has this shortcut edge making a loop, then elimination gives a row of zeros. When the dust clears we have a tree.

An idea suggests itself: **Rows are dependent when edges form a loop.** Independent rows come from trees. This is the key to the row space. We are assuming that the graph is connected, and the arrows could go either way. On each edge, *flow with the arrow is “positive.”* Flow in the opposite direction counts as negative. The flow might be a current or a signal or a force—or even oil or gas or water.

When x_1, x_2, x_3, x_4 are voltages at the nodes, Ax gives voltage differences:

$$Ax = \begin{bmatrix} -1 & 1 & 0 & 0 \\ -1 & 0 & 1 & 0 \\ 0 & -1 & 1 & 0 \\ -1 & 0 & 0 & 1 \\ 0 & -1 & 0 & 1 \\ 0 & 0 & -1 & 1 \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \\ x_3 \\ x_4 \end{bmatrix} = \begin{bmatrix} x_2 - x_1 \\ x_3 - x_1 \\ x_3 - x_2 \\ x_4 - x_1 \\ x_4 - x_2 \\ x_4 - x_3 \end{bmatrix}. \quad (1)$$

Let me say that again. The incidence matrix A is a difference matrix. The input vector x gives voltages, the output vector Ax gives voltage differences (along edges 1 to 6). If the voltages are equal, the differences are zero. This tells us the nullspace of A .

1 The **nullspace** contains the solutions to $Ax = 0$. All six voltage differences are zero. This means: *All four voltages are equal*. Every x in the nullspace is a **constant vector**: $x = (c, c, c, c)$. The nullspace of A is a line in $\mathbf{R}^n$ —its dimension is $n - r = 1$.

The second incidence matrix B has the same nullspace. It contains $(1, 1, 1, 1)$:

$$\text{1-dimensional nullspace: same for the tree} \quad Bx = \begin{bmatrix} -1 & 1 & 0 & 0 \\ 0 & -1 & 1 & 0 \\ 0 & 0 & -1 & 1 \end{bmatrix} \begin{bmatrix} 1 \\ 1 \\ 1 \\ 1 \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \\ 0 \end{bmatrix}.$$

We can raise or lower all voltages by the same amount c , without changing the differences. There is an “arbitrary constant” in the voltages. Compare this with the same statement for functions. We can raise or lower a function by C , without changing its derivative.

Calculus adds “+ C ” to indefinite integrals. Graph theory adds (c, c, c, c) to the vector x . Linear algebra adds any vector x_n in the nullspace to one particular solution of $Ax = b$.

The “+ C ” disappears in calculus when a definite integral starts at a known point. Similarly the nullspace disappears when we fix $x_4 = 0$. The unknown x_4 is removed and so are the fourth columns of A and B (those columns multiplied x_4). Electrical engineers would say that node 4 has been “grounded.”

2 The **row space** contains all combinations of the six rows. Its dimension is certainly not 6. The equation $r + (n - r) = n$ must be $3 + 1 = 4$. The rank is $r = 3$, as we saw from elimination. After 3 edges, we start forming loops! The new rows are not independent.

How can we tell if $v = (v_1, v_2, v_3, v_4)$ is in the row space? The slow way is to combine rows. The quick way is by orthogonality:

v is in the row space if and only if it is perpendicular to $(1, 1, 1, 1)$ in the nullspace.

The vector $v = (0, 1, 2, 3)$ fails this test—its components add to 6. The vector $(-6, 1, 2, 3)$ is in the row space: $-6 + 1 + 2 + 3 = 0$. That vector equals 6(row 1) + 5(row 3) + 3(row 6).

Each row of A adds to zero. This must be true for every vector in the row space.

3 The *column space* contains all combinations of the four columns. We expect three independent columns, since there were three independent rows. The first three columns of A are independent (so are any three). But the four columns add to the zero vector, which says again that $(1, 1, 1, 1)$ is in the nullspace. *How can we tell if a particular vector b is in the column space of an incidence matrix?*

First answer Try to solve $Ax = b$. That misses all the insight. As before, orthogonality gives a better answer. We are now coming to Kirchhoff's two famous laws of circuit theory—the voltage law and current law (**KVL** and **KCL**). Those are natural expressions of “laws” of linear algebra. It is especially pleasant to see the key role of the left nullspace.

Second answer Ax is the vector of voltage differences $x_i - x_j$. If we add differences around a closed loop in the graph, they cancel to leave zero. Around the big triangle formed by edges 1, 3, -2 (the arrow goes backward on edge 2) the differences cancel:

$$\text{Sum of differences is 0} \quad (x_2 - x_1) + (x_3 - x_2) - (x_3 - x_1) = 0.$$

Kirchhoff's Voltage Law: The components of $Ax = b$ add to zero around every loop.

$$\text{Around the big triangle:} \quad b_1 + b_3 - b_2 = 0.$$

By testing each loop, the Voltage Law decides whether b is in the column space. $Ax = b$ can be solved exactly when the components of b satisfy all the same dependencies as the rows of A . Then elimination leads to $0 = 0$, and $Ax = b$ is consistent.

4 The *left nullspace* contains the solutions to $A^T y = 0$. Its dimension is $m - r = 6 - 3$:

$$\text{Current Law} \quad A^T y = \begin{bmatrix} -1 & -1 & 0 & -1 & 0 & 0 \\ 1 & 0 & -1 & 0 & -1 & 0 \\ 0 & 1 & 1 & 0 & 0 & -1 \\ 0 & 0 & 0 & 1 & 1 & 1 \end{bmatrix} \begin{bmatrix} y_1 \\ y_2 \\ y_3 \\ y_4 \\ y_5 \\ y_6 \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \\ 0 \\ 0 \end{bmatrix}. \quad (2)$$

The true number of equations is $r = 3$ and not $n = 4$. Reason: The four equations add to $0 = 0$. The fourth equation follows automatically from the first three.

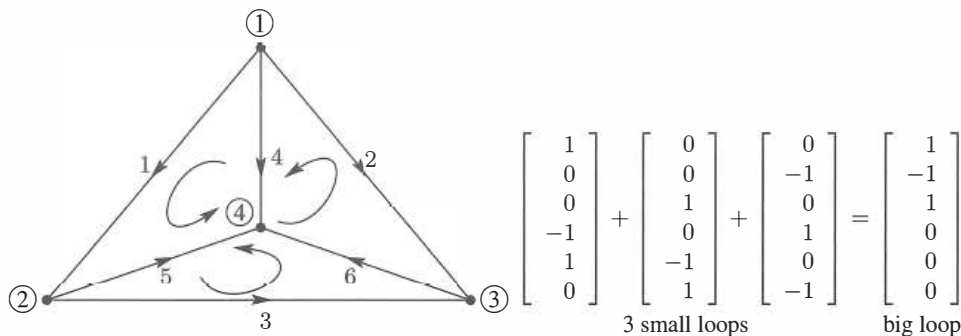
What do the equations mean? The first equation says that $-y_1 - y_2 - y_4 = 0$. *The net flow into node 1 is zero.* The fourth equation says that $y_4 + y_5 + y_6 = 0$. *Flow into node 4 minus flow out is zero.* The equations $A^T y = 0$ are famous and fundamental:

$$\text{Kirchhoff's Current Law: } A^T y = 0 \quad \text{Flow in equals flow out at each node.}$$

This law deserves first place among the equations of applied mathematics. It expresses “conservation” and “continuity” and “balance.” Nothing is lost, nothing is gained. When currents or forces are balanced, the equation to solve is $A^T y = 0$. Notice the beautiful fact that the matrix in this balance equation is the transpose of the incidence matrix A .

What are the actual solutions to $A^T \mathbf{y} = \mathbf{0}$? The currents must balance themselves. The easiest way is to **flow around a loop**. If a unit of current goes around the big triangle (forward on edge 1 and 3, backward on 2), the six currents are $\mathbf{y} = (1, -1, 1, 0, 0, 0)$. This satisfies $A^T \mathbf{y} = \mathbf{0}$. *Every loop current is a solution to the Current Law*. Flow in equals flow out at every node. A smaller loop goes forward on edge 1, forward on 5, back on 4. Then $\mathbf{y} = (1, 0, 0, -1, 1, 0)$ is also in the left nullspace.

We expect three independent $\mathbf{y}$'s: $m - r = 6 - 3 = 3$. The three small loops in the graph are independent. The big triangle seems to give a fourth $\mathbf{y}$, but that flow is the sum of flows around the small loops. *Flows around the 3 small loops are a basis for the left nullspace.*



The incidence matrix A comes from a connected graph with n nodes and m edges. The row space and column space have dimensions $r = n - 1$. The nullspaces of A and A^T have dimensions 1 and $m - n + 1$:

$N(A)$ The constant vectors $(c, c, \dots, c)$ make up the nullspace of A : $\dim = 1$.

$C(A^T)$ The edges of any tree give r independent rows of A : $r = n - 1$.

$C(A)$ **Voltage Law**: The components of $A\mathbf{x}$ add to zero around all loops: $\dim = n - 1$.

$N(A^T)$ **Current Law**: $A^T \mathbf{y} = (\text{flow in}) - (\text{flow out}) = \mathbf{0}$ is solved by loop currents.

There are $m - r = m - n + 1$ independent small loops in the graph.

For every graph in a plane, linear algebra yields **Euler's formula**: Theorem 1 in topology!

$$(\text{number of nodes}) - (\text{number of edges}) + (\text{number of small loops}) = 1.$$

This is $(n) - (m) + (m - n + 1) = 1$. The graph in our example has $4 - 6 + 3 = 1$.

A single triangle has $(3 \text{ nodes}) - (3 \text{ edges}) + (1 \text{ loop})$. On a 10-node tree with 9 edges and no loops, Euler's count is $10 - 9 + 0$. All planar graphs lead to the answer 1.

The next figure shows a network with a current source. Kirchhoff's Current Law changes from $A^T \mathbf{y} = \mathbf{0}$ to $A^T \mathbf{y} = \mathbf{f}$, to balance the source $\mathbf{f}$ from outside. *Flow into each node still equals flow out*. The six edges would have conductances $c_1, \dots, c_6$, and the current source goes into node 1. The source comes out from node 4 to keep the overall balance (**in** = **out**). The problem is: **Find the currents $y_1, \dots, y_6$ on the six edges.**

Flows in networks now lead us from the incidence matrix A to the Laplacian matrix $A^T A$.

Voltages and Currents and $A^T A x = f$

We started with voltages $x = (x_1, \dots, x_n)$ at the nodes. So far we have Ax to find voltage differences $x_i - x_j$ along edges. And we have the Current Law $A^T y = 0$ to find edge currents $y = (y_1, \dots, y_m)$. If all resistances in the network are 1, Ohm's Law will match $y = Ax$. Then $A^T y = A^T Ax = 0$. We are close but not quite there.

Without any sources, the solution to $A^T Ax = 0$ will just be no flow: $x = 0$ and $y = 0$. I can see three ways to produce $x \neq 0$ and $y \neq 0$.

- 1 Assign fixed voltages x_i to one or more nodes.
- 2 Add batteries (voltage sources) in one or more edges.
- 3 Add current sources going into one or more nodes. See Figure 10.2

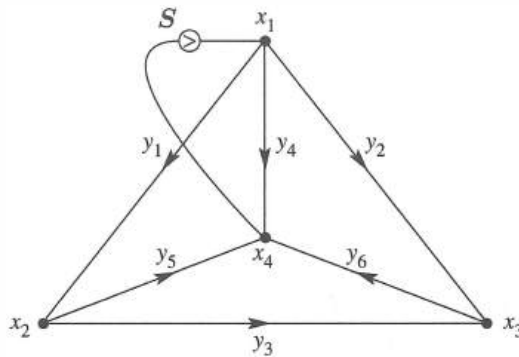


Figure 10.2: The currents y_1 to y_6 in a network with a source S from node 4 to node 1.

Example Figure 10.2 includes a current source S from node 4 to node 1. That current will trickle back through the network to node 4. Some current y_4 will go directly on edge 4. Other current will go the long way from node 1 to 2 to 4, or 1 to 3 to 4. By symmetry I expect no current ($y_3 = 0$) from node 2 to node 3. Solving the network equations will confirm this. **The matrix in those equations is $A^T A$, the graph Laplacian matrix:**

$$\begin{bmatrix} -1 & -1 & 0 & -1 & 0 & 0 \\ 1 & 0 & -1 & 0 & -1 & 0 \\ 0 & 1 & 1 & 0 & 0 & -1 \\ 0 & 0 & 0 & 1 & 1 & 1 \end{bmatrix}
 \begin{bmatrix} -1 & 1 & 0 & 0 \\ -1 & 0 & 1 & 0 \\ 0 & -1 & 1 & 0 \\ -1 & 0 & 0 & 1 \\ 0 & -1 & 0 & 1 \\ 0 & 0 & -1 & 1 \end{bmatrix}
 =
 \begin{bmatrix} 3 & -1 & -1 & -1 \\ -1 & 3 & -1 & -1 \\ -1 & -1 & 3 & -1 \\ -1 & -1 & -1 & 3 \end{bmatrix}$$

$A^T A$

That Laplacian matrix is not invertible! We cannot solve for all four potentials because $(1, 1, 1, 1)$ is in the nullspace of A and $A^T A$. *One node has to be grounded.* Setting $x_4 = 0$ removes the fourth row and column, and this leaves a 3 by 3 invertible matrix. Now we solve $A^T A x = f$ for the unknown potentials x_1, x_2, x_3 , with source S into node 1:

$$\begin{array}{l} \text{Voltages} \\ A^T A x = f \end{array} \quad \begin{bmatrix} 3 & -1 & -1 \\ -1 & 3 & -1 \\ -1 & -1 & 3 \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \\ x_3 \end{bmatrix} = \begin{bmatrix} S \\ 0 \\ 0 \end{bmatrix} \quad \text{gives} \quad \begin{bmatrix} x_1 \\ x_2 \\ x_3 \end{bmatrix} = \begin{bmatrix} S/2 \\ S/4 \\ S/4 \end{bmatrix}.$$

$$\begin{array}{l} \text{Currents} \\ y = -Ax \end{array} \quad \begin{bmatrix} y_1 \\ y_2 \\ y_3 \\ y_4 \\ y_5 \\ y_6 \end{bmatrix} = - \begin{bmatrix} -1 & 1 & 0 & 0 \\ -1 & 0 & 1 & 0 \\ 0 & -1 & 1 & 0 \\ -1 & 0 & 0 & 1 \\ 0 & -1 & 0 & 1 \\ 0 & 0 & -1 & 1 \end{bmatrix} \begin{bmatrix} S/2 \\ S/4 \\ S/4 \\ 0 \end{bmatrix} = \begin{bmatrix} S/4 \\ S/4 \\ 0 \\ S/2 \\ S/4 \\ S/4 \end{bmatrix}.$$

Half the current goes directly on edge 4. That is $y_4 = S/2$. No current crosses from node 2 to node 3. Symmetry indicated $y_3 = 0$ and now the solution proves it.

Admission of error I remembered that current flows from high voltage to low voltage. That produces the minus sign in $y = -Ax$. And the correct form of Ohm's Law will be $Ry = -Ax$ when the resistances on the edges are not all 1. *Conductances* are neater than resistances: $C = R^{-1}$ = diagonal matrix. **We now present Ohm's Law $y = -CAx$.**

Networks and $A^T C A$

In a real network, the current y along an edge is the product of two numbers. One number is the difference between the potentials x at the ends of the edge. This voltage difference is Ax and it drives the flow. The other number c is the “conductance”—which measures how easily flow gets through.

In physics and engineering, c is decided by the material. For electrical currents, c is high for metal and low for plastics. For a superconductor, c is nearly infinite. If we consider elastic stretching, c might be low for metal and higher for plastics. In economics, c measures the capacity of an edge or its cost.

To summarize, the graph is known from its incidence matrix A . This tells the node-edge connections. A **network** goes further, and assigns a conductance c to each edge. These numbers $c_1, \dots, c_m$ go into the “conductance matrix” C —which is diagonal.

For a network of resistors, the conductance is $c = 1/(\text{resistance})$. In addition to Kirchhoff's Laws for the whole system of currents, we have Ohm's Law for each current. Ohm's Law connects the current y_1 on edge 1 to the voltage difference $x_2 - x_1$:

Ohm's Law: Current along edge = conductance times voltage difference.

Ohm's Law for all m currents is $y = -CAx$. The vector Ax gives the potential differences, and C multiplies by the conductances. Combining Ohm's Law with Kirchhoff's

Current Law $A^T \mathbf{y} = \mathbf{0}$, we get $A^T C A \mathbf{x} = \mathbf{0}$. This is *almost* the central equation for network flows. The only thing wrong is the zero on the right side! The network needs power from outside—a voltage source or a current source—to make something happen.

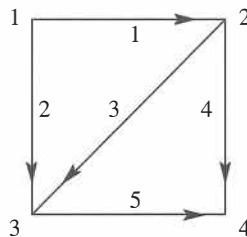
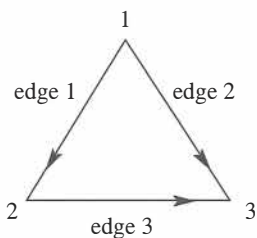
Note about signs In circuit theory we change from $A\mathbf{x}$ to $-A\mathbf{x}$. The flow is from higher potential to lower potential. There is (positive) current from node 1 to node 2 when $x_1 - x_2$ is positive—whereas $A\mathbf{x}$ was constructed to yield $x_2 - x_1$. The minus sign in physics and electrical engineering is a plus sign in mechanical engineering and economics. $A\mathbf{x}$ versus $-A\mathbf{x}$ is a general headache but unavoidable.

Note about applied mathematics Every new application has its own form of Ohm's Law. For springs it is Hooke's Law. The stress $\mathbf{y}$ is (elasticity C) times (stretching $A\mathbf{x}$). For heat conduction, $A\mathbf{x}$ is a temperature gradient. For oil flows it is a pressure gradient. For least squares regression in statistics (Chapter 12) C^{-1} is the covariance matrix.

My textbooks *Introduction to Applied Mathematics* and *Computational Science and Engineering* (Wellesley-Cambridge Press) are practically built on $A^T C A$. This is the key to equilibrium in matrix equations and also in differential equations. Applied mathematics is more organized than it looks! *In new problems I have learned to watch for $A^T C A$.*

Problem Set 10.1

Problems 1–7 and 8–14 are about the incidence matrices for these graphs.



- 1 Write down the 3 by 3 incidence matrix A for the triangle graph. The first row has -1 in column 1 and $+1$ in column 2. What vectors (x_1, x_2, x_3) are in its nullspace? How do you know that $(1, 0, 0)$ is not in its row space?
- 2 Write down A^T for the triangle graph. Find a vector $\mathbf{y}$ in its nullspace. The components of $\mathbf{y}$ are currents on the edges—how much current is going around the triangle?
- 3 Eliminate x_1 and x_2 from the third equation to find the echelon matrix U . What tree corresponds to the two nonzero rows of U ?

$$-x_1 + x_2 = b_1$$

$$-x_1 + x_3 = b_2$$

$$-x_2 + x_3 = b_3.$$

- 4 Choose a vector (b_1, b_2, b_3) for which $Ax = b$ can be solved, and another vector b that allows no solution. How are those b 's related to $y = (1, -1, 1)$?
- 5 Choose a vector (f_1, f_2, f_3) for which $A^T y = f$ can be solved, and a vector f that allows no solution. How are those f 's related to $x = (1, 1, 1)$? The equation $A^T y = f$ is Kirchhoff's _____ law.
- 6 Multiply matrices to find $A^T A$. Choose a vector f for which $A^T Ax = f$ can be solved, and solve for x . Put those potentials x and the currents $y = -Ax$ and current sources f onto the triangle graph. Conductances are 1 because $C = I$.
- 7 With conductances $c_1 = 1$ and $c_2 = c_3 = 2$, multiply matrices to find $A^T CA$. For $f = (1, 0, -1)$ find a solution to $A^T CAx = f$. Write the potentials x and currents $y = -CAx$ on the triangle graph, when the current source f goes into node 1 and out from node 3.
- 8 Write down the 5 by 4 incidence matrix A for the square graph with two loops. Find one solution to $Ax = 0$ and two solutions to $A^T y = 0$.
- 9 Find two requirements on the b 's for the five differences $x_2 - x_1, x_3 - x_1, x_3 - x_2, x_4 - x_2, x_4 - x_3$ to equal b_1, b_2, b_3, b_4, b_5 . You have found Kirchhoff's _____ law around the two _____ in the graph.
- 10 Reduce A to its echelon form U . The three nonzero rows give the incidence matrix for what graph? You found one tree in the square graph—find the other seven trees.
- 11 Multiply matrices to find $A^T A$ and guess how its entries come from the graph:
 - (a) The diagonal of $A^T A$ tells how many _____ into each node.
 - (b) The off-diagonals -1 or 0 tell which pairs of nodes are _____.
- 12 Why is each statement true about $A^T A$? *Answer for $A^T A$ not A .*
 - (a) Its nullspace contains $(1, 1, 1, 1)$. Its rank is $n - 1$.
 - (b) It is positive semidefinite but not positive definite.
 - (c) Its four eigenvalues are real and their signs are _____.
- 13 With conductances $c_1 = c_2 = 2$ and $c_3 = c_4 = c_5 = 3$, multiply the matrices $A^T CA$. Find a solution to $A^T CAx = f = (1, 0, 0, -1)$. Write these potentials x and currents $y = -CAx$ on the nodes and edges of the square graph.
- 14 The matrix $A^T CA$ is not invertible. What vectors x are in its nullspace? Why does $A^T CAx = f$ have a solution if and only if $f_1 + f_2 + f_3 + f_4 = 0$?
- 15 A connected graph with 7 nodes and 7 edges has how many loops?
- 16 For the graph with 4 nodes, 6 edges, and 3 loops, add a new node. If you connect it to one old node, Euler's formula becomes $() - () + () = 1$. If you connect it to two old nodes, Euler's formula becomes $() - () + () = 1$.

- 17** Suppose A is a 12 by 9 incidence matrix from a connected (but unknown) graph.
- (a) How many columns of A are independent?
 - (b) What condition on $\mathbf{f}$ makes it possible to solve $A^T \mathbf{y} = \mathbf{f}$?
 - (c) The diagonal entries of $A^T A$ give the number of edges into each node. What is the sum of those diagonal entries?
- 18** Why does a complete graph with $n = 6$ nodes have $m = 15$ edges? A tree connecting 6 nodes has _____ edges.

Note The *stoichiometric matrix* in chemistry is an important “generalized” incidence matrix. Its entries show how much of each chemical species (each column) goes into each reaction (each row).

10.2 Matrices in Engineering

This section will show how engineering problems produce symmetric matrices K (often K is positive definite). The “linear algebra reason” for symmetry and positive definiteness is their form $K = A^T A$ and $K = A^T C A$. The “physical reason” is that the expression $\frac{1}{2} \mathbf{u}^T K \mathbf{u}$ represents *energy*—and energy is never negative. The matrix C , often diagonal, contains positive physical constants like conductance or stiffness or diffusivity.

Our best examples come from mechanical and civil and aeronautical engineering. K is the **stiffness matrix**, and $K^{-1} \mathbf{f}$ is the structure’s response to forces $\mathbf{f}$ from outside. Section 10.1 turned to electrical engineering—the matrices came from networks and circuits. The exercises involve chemical engineering and I could go on! Economics and management and engineering design come later in this chapter (the key is optimization).

Engineering leads to linear algebra in two ways, directly and indirectly:

Direct way The physical problem has only a finite number of pieces. The laws connecting their position or velocity are *linear* (movement is not too big or too fast). The laws are expressed by *matrix equations*.

Indirect way The physical system is “continuous”. Instead of individual masses, the mass density and the forces and the velocities are functions of x or x, y or x, y, z . The laws are expressed by *differential equations*. **To find accurate solutions we approximate by finite difference equations or finite element equations.**

Both ways produce matrix equations and linear algebra. I really believe that you cannot do modern engineering without matrices.

Here we present equilibrium equations $K \mathbf{u} = \mathbf{f}$. With motion, $M d^2 \mathbf{u} / dt^2 + K \mathbf{u} = \mathbf{f}$ becomes dynamic. Then we would use eigenvalues from $K \mathbf{x} = \lambda M \mathbf{x}$, or finite differences.

Differential Equation to Matrix Equation

Differential equations are continuous. Our basic example will be $-d^2 u / dx^2 = f(x)$. Matrix equations are discrete. Our basic example will be $K_0 \mathbf{u} = \mathbf{f}$. By taking the step from second derivatives to second differences, you will see the big picture in a very short space. *Start with fixed boundary conditions at both ends $x = 0$ and $x = 1$:*

**Fixed-fixed
boundary value problem**

$$-\frac{d^2 u}{dx^2} = 1 \text{ with } u(0) = 0 \text{ and } u(1) = 0. \quad (1)$$

That differential equation is linear. A particular solution is $u_p = -\frac{1}{2}x^2$ (then $d^2 u / dx^2 = -1$). We can add any function “in the nullspace”. Instead of solving $A \mathbf{x} = \mathbf{0}$ for a vector $\mathbf{x}$, we solve $-d^2 u / dx^2 = 0$ for a function $u_n(x)$. (Main point: The right side is zero.)

The nullspace solutions are $u_n(x) = C + Dx$ (a 2-dimensional nullspace for a second order differential equation). The complete solution is $u_p + u_n$:

**Complete
solution to**

$$-\frac{d^2u}{dx^2} = 1$$

$$u(x) = -\frac{1}{2}x^2 + C + Dx.$$

(2)

Now find C and D from the two boundary conditions: Set $x = 0$ and then $x = 1$. At $x = 0, u(0) = 0$ forces $C = 0$. At $x = 1, u(1) = 0$ forces $-\frac{1}{2} + D = 0$. Then $D = \frac{1}{2}$:

$$u(x) = -\frac{1}{2}x^2 + \frac{1}{2}x = \frac{1}{2}(x - x^2) \text{ solves the fixed-fixed boundary value problem. (3)}$$

Differences Replace Derivatives

To get matrices instead of derivatives, we have three basic choices—*forward or backward or centered differences*. Start with first derivatives and first differences:

$$\frac{du}{dx} \approx \frac{u(x + \Delta x) - u(x)}{\Delta x} \quad \text{or} \quad \frac{u(x) - u(x - \Delta x)}{\Delta x} \quad \text{or} \quad \frac{u(x + \Delta x) - u(x - \Delta x)}{2\Delta x}.$$

Between $x = 0$ and $x = 1$, we divide the interval into $n + 1$ equal pieces. The pieces have width $\Delta x = 1/(n + 1)$. The values of u at the n breakpoints $\Delta x, 2\Delta x, \dots$ will be the unknowns u_1 to u_n in our matrix equation $Ku = f$:

Solution to compute: $u = (u_1, u_2, \dots, u_n) \approx (u(\Delta x), u(2\Delta x), \dots, u(n\Delta x))$.

Zero values $u_0 = u_{n+1} = 0$ come from the boundary conditions $u(0) = u(1) = 0$.

Replace the derivatives in $-\frac{d}{dx} \left(\frac{du}{dx} \right) = 1$ by forward and backward differences:

$$\frac{1}{(\Delta x)^2} \begin{bmatrix} 1 & -1 & 0 & 0 \\ 0 & 1 & -1 & 0 \\ 0 & 0 & 1 & -1 \end{bmatrix} \begin{bmatrix} 1 & 0 & 0 \\ -1 & 1 & 0 \\ 0 & -1 & -1 \\ 0 & 0 & -1 \end{bmatrix} \begin{bmatrix} u_1 \\ u_2 \\ u_3 \end{bmatrix} = \begin{bmatrix} 1 \\ 1 \\ 1 \end{bmatrix} \quad (4)$$

This is our matrix equation when $n = 3$ and $\Delta x = \frac{1}{4}$. The two first differences are transposes of each other! The equation is $A^T A u = (\Delta x)^2 f$. When we multiply $A^T A$, we get the positive definite second difference matrix K_0 :

$$K_0 u = \frac{1}{(\Delta x)^2} \begin{bmatrix} 2 & -1 & 0 \\ -1 & 2 & -1 \\ 0 & -1 & 2 \end{bmatrix} \begin{bmatrix} u_1 \\ u_2 \\ u_3 \end{bmatrix} = \frac{1}{16} \begin{bmatrix} 1 \\ 1 \\ 1 \end{bmatrix} \quad \text{gives} \quad \begin{bmatrix} u_1 \\ u_2 \\ u_3 \end{bmatrix} = \frac{1}{32} \begin{bmatrix} 3 \\ 4 \\ 3 \end{bmatrix}. \quad (5)$$

The wonderful fact in this example is that those numbers u_1, u_2, u_3 are exactly correct! They agree with the true solution $u = \frac{1}{2}(x - x^2)$ at the three meshpoints $x = \frac{1}{4}, \frac{2}{4}, \frac{3}{4}$. Figure 10.3 shows the true solution (continuous curve) and the approximations u_1, u_2, u_3 (lying exactly on the curve). This curve is a parabola.

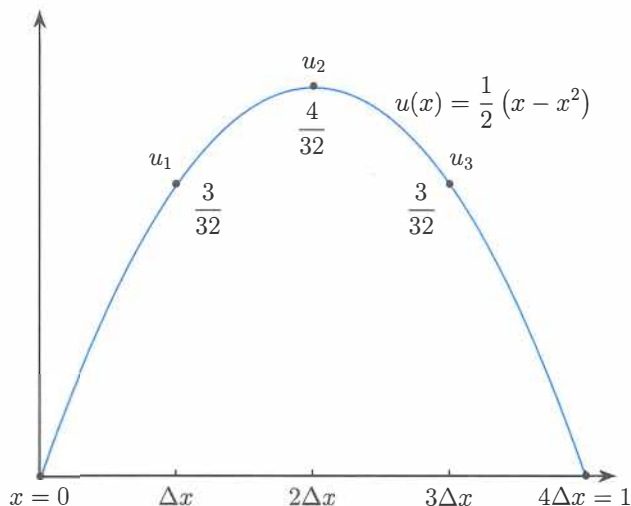


Figure 10.3: Solutions to $-\frac{d^2u}{dx^2} = 1$ and $K_0 u = (\Delta x)^2 f$ with fixed-fixed boundaries.

How to explain this perfect answer, lying right on the graph of $u(x)$? In the matrix equation, $K_0 = A^T A$ is a “second difference matrix.” It gives a centered approximation to $-d^2u/dx^2$. I included the minus sign because the first derivative is *antisymmetric*. The second derivative by itself is *negative*:

The “transpose” of $\frac{d}{dx}$ is $-\frac{d}{dx}$. Then $\left(-\frac{d}{dx}\right)\left(\frac{d}{dx}\right)$ is positive definite.

You can see that in the matrices A and A^T . The transpose of $A = \text{forward difference}$ is $A^T = -\text{backward difference}$. I don’t want to choose a centered $u(x + \Delta x) - u(x - \Delta x)$. Centered is the best for a first difference, but then the second difference $A^T A$ would stretch from $u(x + 2\Delta x)$ to $u(x - 2\Delta x)$: not good.

Now we can explain the perfect answers, exactly on the true curve $u(x) = \frac{1}{2}(x - x^2)$. Second differences $-1, 2, -1$ are exactly correct for straight lines $y = x$ and parabolas!

$$y = x \quad -\frac{d^2y}{dx^2} = 0 \quad -(x + \Delta x) + 2x - (x - \Delta x) = 0(\Delta x)^2$$

$$y = x^2 \quad -\frac{d^2y}{dx^2} = -2 \quad -(x + \Delta x)^2 + 2x^2 - (x - \Delta x)^2 = -2(\Delta x)^2$$

The miracle continues to $y = x^3$. The correct $-d^2y/dx^2 = -6x$ is produced by second differences. But for $y = x^4$ we return to earth. Second differences don’t exactly match $-y'' = -12x^2$. The approximations u_1, u_2, u_3 won’t fall on the graph of $u(x)$.

Fixed End and Free End and Variable Coefficient $c(x)$

To see two new possibilities, I will change the equation and also one boundary condition:

$$\boxed{-\frac{d}{dx} \left((1+x) \frac{du}{dx} \right) = f(x) \quad \text{with} \quad u(0) = 0 \quad \text{and} \quad \frac{du}{dx}(1) = 0.} \quad (6)$$

The end $x = 1$ is now **free**. There is no support at that end. “A hanging bar is fixed only at the top.” There is no force at the free end $x = 1$. That translates to $du/dx = 0$ instead of the fixed condition $u = 0$ at $x = 1$.

The other change is in the coefficient $c(x) = 1 + x$. The stiffness of the bar is varying as you go from $x = 0$ to $x = 1$. Maybe its width is changing, or the material changes. This coefficient $1 + x$ will bring a new matrix C into the difference equation.

Since u_4 is no longer fixed at 0, it becomes a new unknown. The backward difference A is 4 by 4. And the multiplication by $c(x) = 1 + x$ becomes a diagonal matrix C —which multiplies by $1 + \Delta x, \dots, 1 + 4\Delta x$ at the meshpoints. Here are A^T , C , and A :

$$A^T C A = \begin{bmatrix} 1 & -1 & 0 & 0 \\ 0 & 1 & -1 & 0 \\ 0 & 0 & 1 & -1 \\ 0 & 0 & 0 & 1 \end{bmatrix} \begin{bmatrix} 1.25 & & & \\ & 1.5 & & \\ & & 1.75 & \\ & & & 2.0 \end{bmatrix} \begin{bmatrix} 1 & 0 & 0 & 0 \\ -1 & 1 & 0 & 0 \\ 0 & -1 & 1 & 0 \\ 0 & 0 & -1 & 1 \end{bmatrix}. \quad (7)$$

This matrix $K = A^T C A$ will be symmetric and positive definite! Symmetric because $(A^T C A)^T = A^T C^T A^{TT} = A^T C A$. Positive definite because it passes the energy test: A has independent columns, so $Ax \neq 0$ when $x \neq 0$.

$$\text{Energy} = x^T A^T C A x = (Ax)^T C (Ax) > 0 \text{ for every } x \neq 0, \text{ because } Ax \neq 0.$$

When you multiply the matrices $A^T A$ and $A^T C A$ for this fixed-free combination, watch how 1 replaces 2 in the last corner of $A^T A$. That fourth equation has $u_4 - u_3$, a first (not second) difference coming from the free boundary condition $du/dx = 0$.

Notice in $A^T C A$ how c_1, c_2, c_3, c_4 come from $c(x) = 1 + x$ in equation (7). Previously the c 's were simply 1, 1, 1, 1. Here are the **fixed-free** matrices:

$$A^T A = \begin{bmatrix} 2 & -1 & & \\ -1 & 2 & -1 & \\ & -1 & 2 & -1 \\ & & -1 & 1 \end{bmatrix} \quad A^T C A = \begin{bmatrix} c_1 + c_2 & -c_2 & & \\ -c_2 & c_2 + c_3 & -c_3 & \\ & -c_3 & c_3 + c_4 & -c_4 \\ & & -c_4 & c_4 \end{bmatrix}. \quad (8)$$

Free-free Boundary Conditions

Suppose both ends of the bar are free. Now $du/dx = 0$ at both $x = 0$ and $x = 1$. Nothing is holding the bar in place! Physically it is unstable—it can move with no force. Mathematically all constant functions like $u = 1$ satisfy these free conditions. **Algebraically our matrices $A^T A$ and $A^T C A$ will not be invertible:**

Free-free examples
Unknown u_0, u_1, u_2
 $\Delta x = 0.5$

$$A^T A = \begin{bmatrix} 1 & -1 & 0 \\ -1 & 2 & -1 \\ 0 & -1 & 1 \end{bmatrix} \quad A^T C A = \begin{bmatrix} c_0 & -c_0 & \\ -c_0 & c_0 + c_1 & -c_1 \\ & -c_1 & c_1 \end{bmatrix}.$$

The vector $(1, 1, 1)$ is in both nullspaces. This matches $u(x) = 1$ in the continuous problem. Free-free $A^T A \mathbf{u} = \mathbf{f}$ and $A^T C A \mathbf{u} = \mathbf{f}$ are generally unsolvable.

Before explaining more physical examples, may I write down six of the matrices? The tridiagonal K_0 appears many times in this textbook. Now we are seeing its applications. These matrices are all symmetric, and the first four are positive definite:

$$K_0 = A_0^T A_0 = \begin{bmatrix} 2 & -1 & \\ -1 & 2 & -1 \\ & -1 & 2 \end{bmatrix} \quad A_0^T C_0 A_0 = \begin{bmatrix} c_1 + c_2 & -c_2 & \\ -c_2 & c_2 + c_3 & -c_3 \\ & -c_3 & c_3 + c_4 \end{bmatrix}$$

Fixed-fixed

Spring constants included

$$K_1 = A_1^T A_1 = \begin{bmatrix} 2 & -1 & \\ -1 & 2 & -1 \\ & -1 & 1 \end{bmatrix} \quad A_1^T C_1 A_1 = \begin{bmatrix} c_1 + c_2 & -c_2 & \\ -c_2 & c_2 + c_3 & -c_3 \\ & -c_3 & c_3 \end{bmatrix}$$

Fixed-free

Spring constants included

$$K_{\text{singular}} = \begin{bmatrix} 1 & -1 & \\ -1 & 2 & -1 \\ & -1 & 1 \end{bmatrix} \quad K_{\text{circular}} = \begin{bmatrix} 2 & -1 & -1 \\ -1 & 2 & -1 \\ -1 & -1 & 2 \end{bmatrix}$$

Free-free

Periodic $u(0) = u(1)$

The matrices $K_0, K_1, K_{\text{singular}}$, and K_{circular} have $C = I$ for simplicity. This means that all the “spring constants” are $c_i = 1$. We included $A_0^T C_0 A_0$ and $A_1^T C_1 A_1$ to show how the spring constants enter the matrix (without changing its positive definiteness). Our next goal is to see these same stiffness matrices in other engineering problems.

A Line of Springs and Masses

Figure 10.4 shows three masses m_1, m_2, m_3 connected by a line of springs. The fixed-fixed case has four springs, with top and bottom fixed. That leads to K_0 and $A_0^T C_0 A_0$. The fixed-free case has only three springs; the lowest mass hangs freely. That will lead to K_1 and $A_1^T C_1 A_1$. A **free-free** problem produces K_{singular} .

We want equations for the mass movements $\mathbf{u}$ and the spring tensions $\mathbf{y}$:

$$\begin{aligned}\mathbf{u} &= (u_1, u_2, u_3) = \text{movements of the masses (down is positive)} \\ \mathbf{y} &= (y_1, y_2, y_3, y_4) \text{ or } (y_1, y_2, y_3) = \text{tensions in the springs}\end{aligned}$$

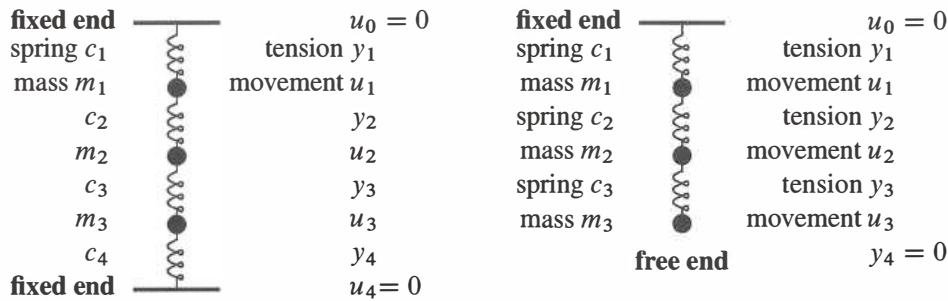


Figure 10.4: Lines of springs and masses: **fixed-fixed** and **fixed-free** ends.

When a mass moves downward, its displacement is positive ($u_j > 0$). For the springs, tension is positive and compression is negative ($y_i < 0$). In tension, the spring is stretched so it pulls the masses inward. Each spring is controlled by its own Hooke's Law $y = ce$: (stretching force $\mathbf{y}$) = (spring constant c) times (stretching distance e).

Our job is to link these one-spring equations $y = ce$ into a vector equation $K\mathbf{u} = \mathbf{f}$ for the whole system. The force vector $\mathbf{f}$ comes from gravity. The gravitational constant g will multiply each mass to produce downward forces $\mathbf{f} = (m_1g, m_2g, m_3g)$.

The real problem is to find the stiffness matrix (**fixed-fixed** and **fixed-free**). The best way to create K is in three steps, not one. Instead of connecting the movements $\mathbf{u}_j$ directly to the forces f_i , it is much better to connect each vector to the next in this list:

$$\begin{aligned}\mathbf{u} &= \text{Movements of } n \text{ masses} &= (u_1, \dots, u_n) \\ \mathbf{e} &= \text{Elongations of } m \text{ springs} &= (e_1, \dots, e_m) \\ \mathbf{y} &= \text{Internal forces in } m \text{ springs} &= (y_1, \dots, y_m) \\ \mathbf{f} &= \text{External forces on } n \text{ masses} &= (f_1, \dots, f_n)\end{aligned}$$

A great framework for applied mathematics connects $\mathbf{u}$ to $\mathbf{e}$ to $\mathbf{y}$ to $\mathbf{f}$. Then $A^T C A \mathbf{u} = \mathbf{f}$:

$$\begin{array}{ccc} \boxed{\mathbf{u}} & \boxed{\mathbf{f}} & \mathbf{e} = A\mathbf{u} \quad A \text{ is } m \text{ by } n \\ \downarrow A & \uparrow A^T & \mathbf{y} = C\mathbf{e} \quad C \text{ is } m \text{ by } m \\ \boxed{\mathbf{e}} & \xrightarrow{C} \boxed{\mathbf{y}} & \mathbf{f} = A^T \mathbf{y} \quad A^T \text{ is } n \text{ by } m \end{array}$$

We will write down the matrices A and C and A^T for the two examples, first with fixed ends and then with the lower end free. Forgive the simplicity of these matrices, it is their form that is so important. Especially the appearance of A together with A^T .

The elongation e is the stretching distance—how far the springs are extended. Originally there is no stretching—the system is lying on a table. When it becomes vertical and upright, gravity acts. The masses move down by distances u_1, u_2, u_3 . Each spring is stretched or compressed by $e_i = u_i - u_{i-1}$, the difference in displacements of its ends:

$$\begin{array}{ll} \text{First spring:} & e_1 = u_1 \quad (\text{the top is fixed so } u_0 = 0) \\ \text{Second spring:} & e_2 = u_2 - u_1 \\ \text{Third spring:} & e_3 = u_3 - u_2 \\ \text{Fourth spring:} & e_4 = -u_3 \quad (\text{the bottom is fixed so } u_4 = 0) \end{array}$$

If both ends move the same distance, that spring is not stretched: $u_j = u_{j-1}$ and $e_j = 0$. The matrix in those four equations is a 4 by 3 difference matrix A , and $e = Au$:

$$\begin{array}{l} \text{Stretching} \\ \text{distances} \\ \text{(elongations)} \end{array} \quad e = Au \quad \text{is} \quad \begin{bmatrix} e_1 \\ e_2 \\ e_3 \\ e_4 \end{bmatrix} = \begin{bmatrix} 1 & 0 & 0 \\ -1 & 1 & 0 \\ 0 & -1 & 1 \\ 0 & 0 & -1 \end{bmatrix} \begin{bmatrix} u_1 \\ u_2 \\ u_3 \end{bmatrix}. \quad (9)$$

The next equation $y = Ce$ connects spring elongation e with spring tension y . This is Hooke's Law $y_i = c_i e_i$ for each separate spring. It is the "constitutive law" that depends on the material in the spring. A soft spring has small c , so a moderate force y can produce a large stretching e . Hooke's linear law is nearly exact for real springs, before they are overstretched and the material becomes plastic.

Since each spring has its own law, the matrix in $y = Ce$ is a diagonal matrix C :

$$\begin{array}{l} \text{Hooke's} \\ \text{Law} \\ y = Ce \end{array} \quad \begin{array}{l} y_1 = c_1 e_1 \\ y_2 = c_2 e_2 \\ y_3 = c_3 e_3 \\ y_4 = c_4 e_4 \end{array} \quad \text{is} \quad \begin{bmatrix} y_1 \\ y_2 \\ y_3 \\ y_4 \end{bmatrix} = \begin{bmatrix} c_1 & & & \\ & c_2 & & \\ & & c_3 & \\ & & & c_4 \end{bmatrix} \begin{bmatrix} e_1 \\ e_2 \\ e_3 \\ e_4 \end{bmatrix} \quad (10)$$

Combining $e = Au$ with $y = Ce$, the spring forces (tension forces) are $y = CAu$.

Finally comes the balance equation, the most fundamental law of applied mathematics. The internal forces from the springs balance the external forces on the masses. Each mass is pulled or pushed by the spring force y_j above it. From below it feels the spring force y_{j+1} plus f_j from gravity. Thus $y_j = y_{j+1} + f_j$ or $f_j = y_j - y_{j+1}$:

$$\begin{array}{l} \text{Force} \\ \text{balance} \\ f = A^T y \end{array} \quad \begin{array}{l} f_1 = y_1 - y_2 \\ f_2 = y_2 - y_3 \\ f_3 = y_3 - y_4 \end{array} \quad \text{is} \quad \begin{bmatrix} f_1 \\ f_2 \\ f_3 \end{bmatrix} = \begin{bmatrix} 1 & -1 & 0 & 0 \\ 0 & 1 & -1 & 0 \\ 0 & 0 & 1 & -1 \end{bmatrix} \begin{bmatrix} y_1 \\ y_2 \\ y_3 \\ y_4 \end{bmatrix} \quad (11)$$

That matrix is A^T ! The equation for balance of forces is $f = A^T y$. Nature transposes the rows and columns of the $e - u$ matrix to produce the $f - y$ matrix. This is the beauty of the framework, that A^T appears along with A . The three equations combine into $Ku = f$.

$$\left\{ \begin{array}{l} \mathbf{e} = \mathbf{A}\mathbf{u} \\ \mathbf{y} = \mathbf{C}\mathbf{e} \\ \mathbf{f} = \mathbf{A}^T\mathbf{y} \end{array} \right\} \quad \begin{array}{l} \text{combine into } \mathbf{A}^T\mathbf{C}\mathbf{A}\mathbf{u} = \mathbf{f} \text{ or } \mathbf{K}\mathbf{u} = \mathbf{f} \\ \mathbf{K} = \mathbf{A}^T\mathbf{C}\mathbf{A} \text{ is the } \mathbf{\text{stiffness matrix}} \text{ (mechanics)} \\ \mathbf{K} = \mathbf{A}^T\mathbf{C}\mathbf{A} \text{ is the } \mathbf{\text{conductance matrix}} \text{ (networks)} \end{array}$$

Finite element programs spend major effort on assembling $\mathbf{K} = \mathbf{A}^T\mathbf{C}\mathbf{A}$ from thousands of smaller pieces. We find $\mathbf{K}$ for four springs (**fixed-fixed**) by multiplying $\mathbf{A}^T$ times $\mathbf{C}\mathbf{A}$:

$$\begin{bmatrix} 1 & -1 & 0 & 0 \\ 0 & 1 & -1 & 0 \\ 0 & 0 & 1 & -1 \end{bmatrix} \begin{bmatrix} c_1 & 0 & 0 \\ -c_2 & c_2 & 0 \\ 0 & -c_3 & c_3 \\ 0 & 0 & -c_4 \end{bmatrix} = \begin{bmatrix} c_1 + c_2 & -c_2 & 0 \\ -c_2 & c_2 + c_3 & -c_3 \\ 0 & -c_3 & c_3 + c_4 \end{bmatrix}$$

If all springs are identical, with $c_1 = c_2 = c_3 = c_4 = 1$, then $\mathbf{C} = \mathbf{I}$. The stiffness matrix reduces to $\mathbf{A}^T\mathbf{A}$. It becomes the special $-1, 2, -1$ matrix $\mathbf{K}_0$.

Note the difference between $\mathbf{A}^T\mathbf{A}$ from engineering and $\mathbf{LU}$ from linear algebra. The matrix $\mathbf{A}$ from four springs is 4 by 3. The triangular matrices from elimination are square. The stiffness matrix $\mathbf{K}$ is assembled from $\mathbf{A}^T\mathbf{A}$, and then broken up into $\mathbf{LU}$. One step is applied mathematics, the other is computational mathematics. Each $\mathbf{K}$ is built from rectangular matrices and factored into square matrices.

May I list some properties of $\mathbf{K} = \mathbf{A}^T\mathbf{C}\mathbf{A}$? You know almost all of them:

1. $\mathbf{K}$ is **tridiagonal**, because mass 3 is not connected to mass 1.
2. $\mathbf{K}$ is **symmetric**, because $\mathbf{C}$ is symmetric and $\mathbf{A}^T$ comes with $\mathbf{A}$.
3. $\mathbf{K}$ is **positive definite**, because $c_i > 0$ and $\mathbf{A}$ has **independent columns**.
4. $\mathbf{K}^{-1}$ is a **full matrix** (not sparse) with **all positive entries**.

Property 4 leads to an important fact about $\mathbf{u} = \mathbf{K}^{-1}\mathbf{f}$: *If all forces act downwards ($f_j > 0$) then all movements are downwards ($u_j > 0$).* Notice that “positive” is different from “positive definite”. $\mathbf{K}^{-1}$ is positive ($\mathbf{K}$ is not). Both are positive definite.

Example 1 Suppose all $c_i = c$ and $m_j = m$. Find the movements $\mathbf{u}$ and tensions $\mathbf{y}$.

All springs are the same and all masses are the same. But all movements and elongations and tensions will *not* be the same. $\mathbf{K}^{-1}$ includes $\frac{1}{c}$ because $\mathbf{A}^T\mathbf{C}\mathbf{A}$ includes c :

$$\text{Movements} \quad \mathbf{u} = \mathbf{K}^{-1}\mathbf{f} = \frac{1}{4c} \begin{bmatrix} 3 & 2 & 1 \\ 2 & 4 & 2 \\ 1 & 2 & 3 \end{bmatrix} \begin{bmatrix} mg \\ mg \\ mg \end{bmatrix} = \frac{mg}{c} \begin{bmatrix} 3/2 \\ 2 \\ 3/2 \end{bmatrix}$$

The displacement u_2 , for the mass in the middle, is greater than u_1 and u_3 . The units are correct: the force mg divided by force per unit length c gives a length u . Then

$$\text{Elongations} \quad \mathbf{e} = \mathbf{A}\mathbf{u} = \begin{bmatrix} 1 & 0 & 0 \\ -1 & 1 & 0 \\ 0 & -1 & 1 \\ 0 & 0 & -1 \end{bmatrix} \frac{mg}{c} \begin{bmatrix} 3/2 \\ 2 \\ 3/2 \end{bmatrix} = \frac{mg}{c} \begin{bmatrix} 3/2 \\ 1/2 \\ -1/2 \\ -3/2 \end{bmatrix}.$$

Warning: Normally you cannot write $K^{-1} = A^{-1}C^{-1}(A^T)^{-1}$.

The three matrices are mixed together by A^TCA , and they cannot easily be untangled. In general, $A^T\mathbf{y} = \mathbf{f}$ has many solutions. And four equations $A\mathbf{u} = \mathbf{e}$ would usually have no solution with three unknowns. But A^TCA gives the correct solution to all three equations in the framework. Only when $m = n$ and the matrices are square can we go from $\mathbf{y} = (A^T)^{-1}\mathbf{f}$ to $\mathbf{e} = C^{-1}\mathbf{y}$ to $\mathbf{u} = A^{-1}\mathbf{e}$. We will see that now.

Fixed End and Free End

Remove the fourth spring. All matrices become 3 by 3. The pattern does not change! The matrix A loses its fourth row and (of course) A^T loses its fourth column. The new stiffness matrix K_1 becomes a product of square matrices:

$$A_1^T C_1 A_1 = \begin{bmatrix} 1 & -1 & 0 \\ 0 & 1 & -1 \\ 0 & 0 & 1 \end{bmatrix} \begin{bmatrix} c_1 & & \\ & c_2 & \\ & & c_3 \end{bmatrix} \begin{bmatrix} 1 & 0 & 0 \\ -1 & 1 & 0 \\ 0 & -1 & 1 \end{bmatrix}.$$

The missing column of A^T and row of A multiplied the missing c_4 . So the quickest way to find the new A^TCA is to set $c_4 = 0$ in the old one:

$$\begin{array}{l} \text{FIXED} \\ \text{FREE} \end{array} \quad A_1^T C_1 A_1 = \begin{bmatrix} c_1 + c_2 & -c_2 & 0 \\ -c_2 & c_2 + c_3 & -c_3 \\ 0 & -c_3 & c_3 \end{bmatrix}. \quad (12)$$

Example 2 If $c_1 = c_2 = c_3 = 1$ and $C = I$, this is the $-1, 2, -1$ tridiagonal matrix K_1 . The last entry of K_1 is 1 instead of 2 because the spring at the bottom is free. Suppose all $m_j = m$:

$$\text{Fixed-free} \quad \mathbf{u} = K_1^{-1}\mathbf{f} = \frac{1}{c} \begin{bmatrix} 1 & 1 & 1 \\ 1 & 2 & 2 \\ 1 & 2 & 3 \end{bmatrix} \begin{bmatrix} mg \\ mg \\ mg \end{bmatrix} = \frac{mg}{c} \begin{bmatrix} 3 \\ 5 \\ 6 \end{bmatrix}.$$

Those movements are greater than the free-free case. The number 3 appears in u_1 because all three masses are pulling the first spring down. The next mass moves by that 3 plus an additional 2 from the masses below it. The third mass drops even more ($3 + 2 + 1 = 6$). The elongations $\mathbf{e} = A\mathbf{u}$ in the springs display those numbers 3, 2, 1:

$$\mathbf{e} = \begin{bmatrix} 1 & 0 & 0 \\ -1 & 1 & 0 \\ 0 & -1 & 1 \end{bmatrix} \frac{mg}{c} \begin{bmatrix} 3 \\ 5 \\ 6 \end{bmatrix} = \frac{mg}{c} \begin{bmatrix} 3 \\ 2 \\ 1 \end{bmatrix}.$$

Two Free Ends: K is Singular

Freedom at *both ends* means trouble. The whole line can move. A is 2 by 3:

$$\begin{array}{l} \text{FREE-FREE} \\ e = Au \end{array} \quad \begin{bmatrix} e_1 \\ e_2 \end{bmatrix} = \begin{bmatrix} u_2 - u_1 \\ u_3 - u_2 \end{bmatrix} = \begin{bmatrix} -1 & 1 & 0 \\ 0 & -1 & 1 \end{bmatrix} \begin{bmatrix} u_1 \\ u_2 \\ u_3 \end{bmatrix}. \quad (13)$$

Now there is a nonzero solution to $Au = 0$. *The masses can move with no stretching of the springs.* The whole line can shift by $u = (1, 1, 1)$ and this leaves $e = (0, 0)$:

$$Au = \begin{bmatrix} -1 & 1 & 0 \\ 0 & -1 & 1 \end{bmatrix} \begin{bmatrix} 1 \\ 1 \\ 1 \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \end{bmatrix} = \text{no stretching}. \quad (14)$$

$Au = 0$ certainly leads to $A^T CAu = 0$. Then $A^T CA$ is only *positive semidefinite*, without c_1 and c_4 . The pivots will be c_2 and c_3 and *no third pivot*. The rank is only 2:

$$\begin{bmatrix} -1 & 0 \\ 1 & -1 \\ 0 & 1 \end{bmatrix} \begin{bmatrix} c_2 & & \\ & c_3 & \end{bmatrix} \begin{bmatrix} -1 & 1 & 0 \\ 0 & -1 & 1 \end{bmatrix} = \begin{bmatrix} c_2 & -c_2 & 0 \\ -c_2 & c_2 + c_3 & -c_3 \\ 0 & -c_3 & c_3 \end{bmatrix} \quad (15)$$

Two eigenvalues will be positive but $x = (1, 1, 1)$ is an eigenvector for $\lambda = 0$. We can solve $A^T CAu = f$ only for special vectors f . The forces have to add to $f_1 + f_2 + f_3 = 0$, or the whole line of springs (with both ends free) will take off like a rocket.

Circle of Springs

A third spring will complete the circle from mass 3 back to mass 1. This doesn't make K invertible—the stiffness matrix K_{circular} matrix is still singular:

$$A_{\text{circular}}^T A_{\text{circular}} = \begin{bmatrix} 1 & -1 & 0 \\ 0 & 1 & -1 \\ -1 & 0 & 1 \end{bmatrix} \begin{bmatrix} 1 & 0 & -1 \\ -1 & 1 & 0 \\ 0 & -1 & 1 \end{bmatrix} = \begin{bmatrix} 2 & -1 & -1 \\ -1 & 2 & -1 \\ -1 & -1 & 2 \end{bmatrix}. \quad (16)$$

The only pivots are 2 and $\frac{3}{2}$. The eigenvalues are 3 and 3 and 0. The determinant is zero. The nullspace still contains $x = (1, 1, 1)$, when all the masses move together. This movement vector $(1, 1, 1)$ is in the nullspace of A_{circular} and $K_{\text{circular}} = A_{\text{circular}}^T CA_{\text{circular}}$.

May I summarize this section? I hope the example will help you connect calculus with linear algebra, replacing differential equations by difference equations. If your step Δx is small enough, you will have a totally satisfactory solution.

$$\text{The equation is } -\frac{d}{dx} \left(c(x) \frac{du}{dx} \right) = f(x) \text{ with } u(0) = 0 \text{ and } \left[u(1) \text{ or } \frac{du}{dx}(1) \right] = 0$$

Divide the bar into N pieces of length Δx . Replace du/dx by Au and $-dy/dx$ by $A^T y$. Now A and A^T include $1/\Delta x$. The end conditions are $u_0 = 0$ and $[u_N = 0 \text{ or } y_N = 0]$.

The three steps $-d/dx$ and $c(x)$ and d/dx correspond to A^T and C and A :

$$\mathbf{f} = A^T \mathbf{y} \text{ and } \mathbf{y} = C\mathbf{e} \text{ and } \mathbf{e} = A\mathbf{u} \text{ give } A^T C A \mathbf{u} = \mathbf{f}.$$

This is a fundamental example in computational science and engineering.

1. Model the problem by a differential equation
2. Discretize the differential equation to a difference equation
3. Understand and solve the difference equation (and boundary conditions!)
4. Interpret the solution; visualize it; redesign if needed.

Numerical simulation has become a third branch of science, beside experiment and deduction. Computer design of the Boeing 777 was much less expensive than a wind tunnel.

The two texts *Introduction to Applied Mathematics* and *Computational Science and Engineering* (Wellesley-Cambridge Press) develop this whole subject further—see the course page math.mit.edu/18085 with video lectures (The lectures are also on ocw.mit.edu and **YouTube**). I hope this book helps you to see the framework behind the computations.

Problem Set 10.2

- 1 Show that $\det A_0^T C_0 A_0 = c_1 c_2 c_3 + c_1 c_3 c_4 + c_1 c_2 c_4 + c_2 c_3 c_4$. Find also $\det A_1^T C_1 A_1$ in the fixed-free example.
- 2 Invert $A_1^T C_1 A_1$ in the fixed-free example by multiplying $A_1^{-1} C_1^{-1} (A_1^T)^{-1}$.
- 3 In the free-free case when $A^T C A$ in equation (15) is singular, add the three equations $A^T C A \mathbf{u} = \mathbf{f}$ to show that we need $f_1 + f_2 + f_3 = 0$. Find a solution to $A^T C A \mathbf{u} = \mathbf{f}$ when the forces $\mathbf{f} = (-1, 0, 1)$ balance themselves. Find all solutions!
- 4 Both end conditions for the free-free differential equation are $du/dx = 0$:

$$-\frac{d}{dx} \left(c(x) \frac{du}{dx} \right) = f(x) \quad \text{with} \quad \frac{du}{dx} = 0 \quad \text{at both ends.}$$

Integrate both sides to show that the force $f(x)$ must balance itself, $\int f(x) dx = 0$, or there is no solution. The complete solution is one particular solution $u(x)$ plus any constant. The constant corresponds to $\mathbf{u} = (1, 1, 1)$ in the nullspace of $A^T C A$.

- 5 In the fixed-free problem, the matrix A is square and invertible. We can solve $A^T \mathbf{y} = \mathbf{f}$ separately from $A\mathbf{u} = \mathbf{e}$. Do the same for the differential equation:

$$\text{Solve } -\frac{dy}{dx} = f(x) \text{ with } y(1) = 0. \text{ Graph } y(x) \text{ if } f(x) = 1.$$

- 6 The 3 by 3 matrix $K_1 = A_1^T C_1 A_1$ in equation (6) splits into three “element matrices” $c_1 E_1 + c_2 E_2 + c_3 E_3$. Write down those pieces, one for each c . Show how they come from *column times row* multiplication of $A_1^T C_1 A_1$. This is how finite element stiffness matrices are actually assembled.
- 7 For five springs and four masses with both ends fixed, what are the matrices A and C and K ? With $C = I$ solve $K\mathbf{u} = \text{ones}(4)$.
- 8 Compare the solution $\mathbf{u} = (u_1, u_2, u_3, u_4)$ in Problem 7 to the solution of the continuous problem $-u'' = 1$ with $u(0) = 0$ and $u(1) = 0$. The parabola $u(x)$ should correspond at $x = \frac{1}{5}, \frac{2}{5}, \frac{3}{5}, \frac{4}{5}$ to $\mathbf{u}$ —is there a $(\Delta x)^2$ factor to account for?
- 9 Solve the fixed-free problem $-u'' = mg$ with $u(0) = 0$ and $u'(1) = 0$. Compare $u(x)$ at $x = \frac{1}{3}, \frac{2}{3}, \frac{3}{3}$ with the vector $\mathbf{u} = (3mg, 5mg, 6mg)$ in Example 2.
- 10 Suppose $c_1 = c_2 = c_3 = c_4 = 1$, $m_1 = 2$ and $m_2 = m_3 = 1$. Solve $A^T C A \mathbf{u} = (2, 1, 1)$ for this fixed-fixed line of springs. Which mass moves the most (largest u)?
- 11 (MATLAB) Find the displacements $u(1), \dots, u(100)$ of 100 masses connected by springs all with $c = 1$. Each force is $f(i) = .01$. Print graphs of $\mathbf{u}$ with **fixed-fixed** and **fixed-free** ends. Note that $\text{diag}(\text{ones}(n, 1), d)$ is a matrix with n ones along diagonal d . This print command will graph a vector u :

```
plot(u, '+' );    xlabel('mass number');    ylabel('movement');    print
```

- 12 (MATLAB) Chemical engineering has a first derivative du/dx from fluid velocity as well as d^2u/dx^2 from diffusion. Replace du/dx by a *forward* difference, then a *centered* difference, then a *backward* difference, with $\Delta x = \frac{1}{8}$. Graph your three numerical solutions of

$$-\frac{d^2u}{dx^2} + 10 \frac{du}{dx} = 1 \quad \text{with } u(0) = u(1) = 0.$$

This **convection-diffusion equation** appears everywhere. It transforms to the Black-Scholes equation for option prices in mathematical finance.

Problem 12 is developed into the first MATLAB homework in my 18.085 course on Computational Science and Engineering at MIT. Videos on ocw.mit.edu.

10.3 Markov Matrices, Population, and Economics

This section is about **positive matrices**: every $a_{ij} > 0$. The key fact is quick to state: **The largest eigenvalue is real and positive and so is its eigenvector.** In economics and ecology and population dynamics and random walks, that fact leads a long way:

Markov $\lambda_{\max} = 1$ **Population** $\lambda_{\max} > 1$ **Consumption** $\lambda_{\max} < 1$

$\lambda_{\max}$ controls the powers of A . We will see this first for $\lambda_{\max} = 1$.

Markov Matrices

Multiply a positive vector $\mathbf{u}_0$ again and again by this matrix A :

$$\begin{array}{l} \text{Markov} \\ \text{matrix} \end{array} \quad A = \begin{bmatrix} .8 & .3 \\ .2 & .7 \end{bmatrix} \quad \mathbf{u}_1 = A\mathbf{u}_0 \quad \mathbf{u}_2 = A\mathbf{u}_1 = A^2\mathbf{u}_0$$

After k steps we have $A^k\mathbf{u}_0$. The vectors $\mathbf{u}_1, \mathbf{u}_2, \mathbf{u}_3, \dots$ will approach a “steady state” $\mathbf{u}_\infty = (.6, .4)$. This final outcome does not depend on the starting vector $\mathbf{u}_0$. **For every** $\mathbf{u}_0 = (a, 1-a)$ **we converge to the same** $\mathbf{u}_{\infty(.6,.4)}$. The question is why.

The steady state equation $A\mathbf{u}_\infty = \mathbf{u}_\infty$ makes $\mathbf{u}_\infty$ **an eigenvector with eigenvalue 1**:

$$\begin{array}{l} \text{Steady state} \end{array} \quad \begin{bmatrix} .8 & .3 \\ .2 & .7 \end{bmatrix} \begin{bmatrix} .6 \\ .4 \end{bmatrix} = \begin{bmatrix} .6 \\ .4 \end{bmatrix} = \mathbf{u}_\infty.$$

Multiplying by A does not change $\mathbf{u}_\infty$. But this does not explain why so many vectors $\mathbf{u}_0$ lead to $\mathbf{u}_\infty$. Other examples might have a steady state, but it is not necessarily attractive:

$$\text{Not Markov} \quad B = \begin{bmatrix} 1 & 0 \\ 0 & 2 \end{bmatrix} \quad \text{has the unattractive steady state} \quad B \begin{bmatrix} 1 \\ 0 \end{bmatrix} = \begin{bmatrix} 1 \\ 0 \end{bmatrix}.$$

In this case, the starting vector $\mathbf{u}_0 = (0, 1)$ will give $\mathbf{u}_1 = (0, 2)$ and $\mathbf{u}_2 = (0, 4)$. The second components are doubled. In the language of eigenvalues, B has $\lambda = 1$ but also $\lambda = 2$ —this produces instability. The component of $\mathbf{u}$ along that unstable eigenvector is multiplied by λ , and $|\lambda| > 1$ means blowup.

This section is about two special properties of A that guarantee a *stable steady state*. These properties define a positive **Markov matrix**, and A above is one particular example:

Markov matrix

1. Every entry of A is positive: $a_{ij} > 0$.
2. Every column of A adds to 1.

Column 2 of B adds to 2, not 1. When A is a Markov matrix, two facts are immediate: Because of 1: Multiplying $\mathbf{u}_0 \geq 0$ by A produces a nonnegative $\mathbf{u}_1 = A\mathbf{u}_0 \geq 0$. Because of 2: If the components of $\mathbf{u}_0$ add to 1, so do the components of $\mathbf{u}_1 = A\mathbf{u}_0$.

Reason: The components of $\mathbf{u}_0$ add to 1 when $\begin{bmatrix} 1 & \dots & 1 \end{bmatrix} \mathbf{u}_0 = 1$. This is true for each column of A by Property 2. Then by matrix multiplication $\begin{bmatrix} 1 & \dots & 1 \end{bmatrix} A = \begin{bmatrix} 1 & \dots & 1 \end{bmatrix}$:

$$\text{Components of } A\mathbf{u}_0 \text{ add to 1} \quad \begin{bmatrix} 1 & \dots & 1 \end{bmatrix} A\mathbf{u}_0 = \begin{bmatrix} 1 & \dots & 1 \end{bmatrix} \mathbf{u}_0 = 1.$$

The same facts apply to $\mathbf{u}_2 = A\mathbf{u}_1$ and $\mathbf{u}_3 = A\mathbf{u}_2$. **Every vector $A^k\mathbf{u}_0$ is nonnegative with components adding to 1.** These are “*probability vectors*.” The limit $\mathbf{u}_\infty$ is also a probability vector—but we have to prove that there is a limit. We will show that $\lambda_{\max} = 1$ for a positive Markov matrix.

Example 1 The fraction of rental cars in Denver starts at $\frac{1}{50} = .02$. The fraction outside Denver is .98. Every month, 80% of the Denver cars stay in Denver (and 20% leave). Also 5% of the outside cars come in (95% stay outside). This means that the fractions $\mathbf{u}_0 = (.02, .98)$ are multiplied by A :

$$\text{First month} \quad A = \begin{bmatrix} .80 & .05 \\ .20 & .95 \end{bmatrix} \quad \text{leads to} \quad \mathbf{u}_1 = A\mathbf{u}_0 = A \begin{bmatrix} .02 \\ .98 \end{bmatrix} = \begin{bmatrix} .065 \\ .935 \end{bmatrix}.$$

Notice that $.065 + .935 = 1$. All cars are accounted for. Each step multiplies by A :

$$\text{Next month} \quad \mathbf{u}_2 = A\mathbf{u}_1 = (.09875, .90125). \text{ This is } A^2\mathbf{u}_0.$$

All these vectors are positive because A is positive. Each vector $\mathbf{u}_k$ will have its components adding to 1. The first component has grown from .02 and cars are moving toward Denver. What happens in the long run?

This section involves powers of matrices. The understanding of A^k was our first and best application of diagonalization. Where A^k can be complicated, the diagonal matrix Λ^k is simple. The eigenvector matrix X connects them: A^k equals $X\Lambda^kX^{-1}$. The new application to Markov matrices uses the eigenvalues (in Λ) and the eigenvectors (in X). We will show that $\mathbf{u}_\infty$ is an eigenvector of A corresponding to $\lambda = 1$.

Since every column of A adds to 1, nothing is lost or gained. We are moving rental cars or populations, and no cars or people suddenly appear (or disappear). The fractions add to 1 and the matrix A keeps them that way. The question is how they are distributed after k time periods—which leads us to A^k .

Solution $A^k\mathbf{u}_0$ gives the fractions in and out of Denver after k steps. We diagonalize A to understand A^k . The eigenvalues are $\lambda = 1$ and $.75$ (the trace is 1.75).

$$A\mathbf{x} = \lambda\mathbf{x} \quad A \begin{bmatrix} .2 \\ .8 \end{bmatrix} = 1 \begin{bmatrix} .2 \\ .8 \end{bmatrix} \quad \text{and} \quad A \begin{bmatrix} -1 \\ 1 \end{bmatrix} = .75 \begin{bmatrix} -1 \\ 1 \end{bmatrix}.$$

The starting vector $\mathbf{u}_0$ combines $\mathbf{x}_1$ and $\mathbf{x}_2$, in this case with coefficients 1 and .18:

$$\text{Combination of eigenvectors} \quad \mathbf{u}_0 = \begin{bmatrix} .02 \\ .98 \end{bmatrix} = \begin{bmatrix} .2 \\ .8 \end{bmatrix} + .18 \begin{bmatrix} -1 \\ 1 \end{bmatrix}.$$

Now multiply by A to find $\mathbf{u}_1$. The eigenvectors are multiplied by $\lambda_1 = 1$ and $\lambda_2 = .75$:

$$\text{Each } \mathbf{x} \text{ is multiplied by } \lambda \quad \mathbf{u}_1 = 1 \begin{bmatrix} .2 \\ .8 \end{bmatrix} + (.75)(.18) \begin{bmatrix} -1 \\ 1 \end{bmatrix}.$$

Every month, another $\lambda = .75$ multiplies the vector x_2 . The eigenvector x_1 is unchanged:

After k steps
$$u_k = A^k u_0 = 1^k \begin{bmatrix} .2 \\ .8 \end{bmatrix} + (.75)^k (.18) \begin{bmatrix} -1 \\ 1 \end{bmatrix}.$$

This equation reveals what happens. **The eigenvector x_1 with $\lambda = 1$ is the steady state.** The other eigenvector x_2 disappears because $|\lambda| < 1$. The more steps we take, the closer we come to $u_\infty = (.2, .8)$. In the limit, $\frac{2}{10}$ of the cars are in Denver and $\frac{8}{10}$ are outside. This is the pattern for Markov chains, even starting from $u_0 = (0, 1)$:

If A is a *positive* Markov matrix (entries $a_{ij} > 0$, each column adds to 1), then $\lambda_1 = 1$ is larger than any other eigenvalue. The eigenvector x_1 is the **steady state**:

$$u_k = x_1 + c_2(\lambda_2)^k x_2 + \cdots + c_n(\lambda_n)^k x_n \quad \text{always approaches} \quad u_\infty = x_1.$$

The first point is to see that $\lambda = 1$ is an eigenvalue of A . *Reason:* Every column of $A - I$ adds to $1 - 1 = 0$. The rows of $A - I$ add up to the zero row. Those rows are linearly dependent, so $A - I$ is singular. Its determinant is zero and $\lambda = 1$ is an eigenvalue.

The second point is that no eigenvalue can have $|\lambda| > 1$. With such an eigenvalue, the powers A^k would grow. But A^k is also a Markov matrix! A^k has positive entries still adding to 1—and that leaves no room to get large.

A lot of attention is paid to the possibility that another eigenvalue has $|\lambda| = 1$.

Example 2 $A = \begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix}$ has no steady state because $\lambda_2 = -1$.

This matrix sends all cars from inside Denver to outside, and vice versa. The powers A^k alternate between A and I . The second eigenvector $x_2 = (-1, 1)$ will be multiplied by $\lambda_2 = -1$ at every step—and does not become smaller: No steady state.

Suppose the entries of A or any power of A are all *positive*—zero is not allowed. In this “regular” or “primitive” case, $\lambda = 1$ is strictly larger than any other eigenvalue. The powers A^k approach the rank one matrix that has the steady state in every column.

Example 3 (“Everybody moves”) Start with three groups. At each time step, half of group 1 goes to group 2 and the other half goes to group 3. The other groups also *split in half and move*. Take one step from the starting populations p_1, p_2, p_3 :

New populations
$$u_1 = Au_0 = \begin{bmatrix} 0 & \frac{1}{2} & \frac{1}{2} \\ \frac{1}{2} & 0 & \frac{1}{2} \\ \frac{1}{2} & \frac{1}{2} & 0 \end{bmatrix} \begin{bmatrix} p_1 \\ p_2 \\ p_3 \end{bmatrix} = \begin{bmatrix} \frac{1}{2}p_2 + \frac{1}{2}p_3 \\ \frac{1}{2}p_1 + \frac{1}{2}p_3 \\ \frac{1}{2}p_1 + \frac{1}{2}p_2 \end{bmatrix}.$$

A is a Markov matrix. Nobody is born or lost. A contains zeros, which gave trouble in Example 2. But after two steps in this new example, the zeros disappear from A^2 :

Two-step matrix
$$u_2 = A^2 u_0 = \begin{bmatrix} \frac{1}{2} & \frac{1}{4} & \frac{1}{4} \\ \frac{1}{4} & \frac{1}{2} & \frac{1}{4} \\ \frac{1}{4} & \frac{1}{4} & \frac{1}{2} \end{bmatrix} \begin{bmatrix} p_1 \\ p_2 \\ p_3 \end{bmatrix}.$$

The eigenvalues of A are $\lambda_1 = 1$ (because A is Markov) and $\lambda_2 = \lambda_3 = -\frac{1}{2}$. For $\lambda = 1$, **the eigenvector $x_1 = (\frac{1}{3}, \frac{1}{3}, \frac{1}{3})$ will be the steady state.** When three equal populations split in half and move, the populations are again equal. Starting from $u_0 = (8, 16, 32)$, the Markov chain approaches its steady state:

$$u_0 = \begin{bmatrix} 8 \\ 16 \\ 32 \end{bmatrix} \quad u_1 = \begin{bmatrix} 24 \\ 20 \\ 12 \end{bmatrix} \quad u_2 = \begin{bmatrix} 16 \\ 18 \\ 22 \end{bmatrix} \quad u_3 = \begin{bmatrix} 20 \\ 19 \\ 17 \end{bmatrix}.$$

The step to u_4 will split some people in half. This cannot be helped. The total population is $8 + 16 + 32 = 56$ at every step. The steady state is 56 times $(\frac{1}{3}, \frac{1}{3}, \frac{1}{3})$. You can see the three populations approaching, but never reaching, their final limits $56/3$.

Challenge Problem 6.7.16 created a Markov matrix A from the number of links between websites. The steady state u will give the Google rankings. **Google finds u_∞ by a random walk that follows links** (*random surfing*). That eigenvector comes from counting the fraction of visits to each website—a quick way to compute the steady state.

The size $|\lambda_2|$ of the second eigenvalue controls the speed of convergence to steady state.

Perron-Frobenius Theorem

One matrix theorem dominates this subject. The Perron-Frobenius Theorem applies when all $a_{ij} \geq 0$. There is no requirement that columns add to 1. We prove the neatest form, when all $a_{ij} > 0$: any positive matrix A (not necessarily positive definite!).

Perron-Frobenius for $A > 0$

All numbers in $Ax = \lambda_{\max}x$ are strictly positive.

Proof The key idea is to look at all numbers t such that $Ax \geq tx$ for some nonnegative vector x (other than $x = 0$). We are allowing inequality in $Ax \geq tx$ in order to have many small positive candidates t . For the largest value $t_{\max}$ (which is attained), we will show that **equality holds**: $Ax = t_{\max}x$.

Otherwise, if $Ax \geq t_{\max}x$ is not an equality, multiply by A . Because A is positive that produces a strict inequality $A^2x > t_{\max}Ax$. Therefore the positive vector $y = Ax$ satisfies $Ay > t_{\max}y$, and $t_{\max}$ could be increased. This contradiction forces the equality $Ax = t_{\max}x$, and **we have an eigenvalue**. Its eigenvector x is positive because on the left side of that equality, Ax is sure to be positive.

To see that no eigenvalue can be larger than $t_{\max}$, suppose $Az = \lambda z$. Since λ and z may involve negative or complex numbers, we take absolute values: $|\lambda||z| = |Az| \leq A|z|$ by the “triangle inequality.” This $|z|$ is a nonnegative vector, so this $|\lambda|$ is one of the possible candidates t . Therefore $|\lambda|$ cannot exceed $t_{\max}$ —which must be $\lambda_{\max}$.

Population Growth

Divide the population into three age groups: age < 20 , age 20 to 39, and age 40 to 59. At year T the sizes of those groups are n_1, n_2, n_3 . Twenty years later, the sizes have changed for three reasons: births, deaths, and getting older.

1. **Reproduction** $n_1^{\text{new}} = F_1 n_1 + F_2 n_2 + F_3 n_3$ gives a new generation
2. **Survival** $n_2^{\text{new}} = P_1 n_1$ and $n_3^{\text{new}} = P_2 n_2$ gives the older generations

The fertility rates are F_1, F_2, F_3 (F_2 largest). The *Leslie matrix* A might look like this:

$$\begin{bmatrix} n_1 \\ n_2 \\ n_3 \end{bmatrix}^{\text{new}} = \begin{bmatrix} F_1 & F_2 & F_3 \\ P_1 & 0 & 0 \\ 0 & P_2 & 0 \end{bmatrix} \begin{bmatrix} n_1 \\ n_2 \\ n_3 \end{bmatrix} = \begin{bmatrix} .04 & 1.1 & .01 \\ .98 & 0 & 0 \\ 0 & .92 & 0 \end{bmatrix} \begin{bmatrix} n_1 \\ n_2 \\ n_3 \end{bmatrix}.$$

This is population projection in its simplest form, the same matrix A at every step. In a realistic model, A will change with time (from the environment or internal factors). Professors may want to include a fourth group, age ≥ 60 , but we don't allow it.

The matrix has $A \geq 0$ but not $A > 0$. The Perron-Frobenius theorem still applies because $A^3 > 0$. The largest eigenvalue is $\lambda_{\max} \approx 1.06$. You can watch the generations move, starting from $n_2 = 1$ in the middle generation:

$$\text{eig}(A) = \begin{matrix} 1.06 \\ -1.01 \\ -0.01 \end{matrix} \quad A^2 = \begin{bmatrix} 1.08 & 0.05 & .00 \\ 0.04 & 1.08 & .01 \\ 0.90 & 0 & 0 \end{bmatrix} \quad A^3 = \begin{bmatrix} 0.10 & 1.19 & .01 \\ 0.06 & 0.05 & .00 \\ 0.04 & 0.99 & .01 \end{bmatrix}.$$

A fast start would come from $u_0 = (0, 1, 0)$. That middle group will reproduce 1.1 and also survive .92. The newest and oldest generations are in $u_1 = (1.1, 0, .92) = \text{column 2 of } A$. Then $u_2 = Au_1 = A^2 u_0$ is the second column of A^2 . The early numbers (transients) depend a lot on u_0 , but **the asymptotic growth rate $\lambda_{\max}$ is the same from every start.** Its eigenvector $x = (.63, .58, .51)$ shows all three groups growing steadily together.

Caswell's book on *Matrix Population Models* emphasizes sensitivity analysis. The model is never exactly right. If the F 's or P 's in the matrix change by 10%, does $\lambda_{\max}$ go below 1 (which means extinction)? Problem 19 will show that a matrix change ΔA produces an eigenvalue change $\Delta\lambda = y^T(\Delta A)x$. Here x and y^T are the right and left eigenvectors of A , with $Ax = \lambda x$ and $A^T y = \lambda y$.

Linear Algebra in Economics: The Consumption Matrix

A long essay about linear algebra in economics would be out of place here. A short note about one matrix seems reasonable. The **consumption matrix** tells how much of each input goes into a unit of output. This describes the manufacturing side of the economy.

Consumption matrix We have n industries like chemicals, food, and oil. To produce a unit of chemicals may require .2 units of chemicals, .3 units of food, and .4 units of oil. Those numbers go into row 1 of the consumption matrix A :

$$\begin{bmatrix} \text{chemical output} \\ \text{food output} \\ \text{oil output} \end{bmatrix} = \begin{bmatrix} .2 & .3 & .4 \\ .4 & .4 & .1 \\ .5 & .1 & .3 \end{bmatrix} \begin{bmatrix} \text{chemical input} \\ \text{food input} \\ \text{oil input} \end{bmatrix}.$$

Row 2 shows the inputs to produce food—a heavy use of chemicals and food, not so much oil. Row 3 of A shows the inputs consumed to refine a unit of oil. The real consumption matrix for the United States in 1958 contained 83 industries. The models in the 1990's are much larger and more precise. We chose a consumption matrix that has a convenient eigenvector.

Now comes the question: Can this economy meet demands y_1, y_2, y_3 for chemicals, food, and oil? To do that, the inputs p_1, p_2, p_3 will have to be higher—because part of p is consumed in producing y . The input is p and the consumption is Ap , which leaves the output $p - Ap$. This net production is what meets the demand y :

Problem Find a vector p such that $p - Ap = y$ or $p = (I - A)^{-1}y$.

Apparently the linear algebra question is whether $I - A$ is invertible. But there is more to the problem. The vector y of required outputs is nonnegative, and so is A . *The production levels in $p = (I - A)^{-1}y$ must also be nonnegative.* The real question is:

When is $(I - A)^{-1}$ a nonnegative matrix?

This is the test on $(I - A)^{-1}$ for a productive economy, which can meet any demand. If A is small compared to I , then Ap is small compared to p . There is plenty of output. If A is too large, then production consumes too much and the demand y cannot be met.

“Small” or “large” is decided by the largest eigenvalue λ_1 of A (which is positive):

- If $\lambda_1 > 1$ then $(I - A)^{-1}$ has negative entries
- If $\lambda_1 = 1$ then $(I - A)^{-1}$ fails to exist
- If $\lambda_1 < 1$ then $(I - A)^{-1}$ is nonnegative as desired.

The main point is that last one. The reasoning uses a nice formula for $(I - A)^{-1}$, which we give now. The most important infinite series in mathematics is the **geometric series** $1 + x + x^2 + \cdots$. This series adds up to $1/(1 - x)$ provided x lies between -1 and 1 . When $x = 1$ the series is $1 + 1 + 1 + \cdots = \infty$. When $|x| \geq 1$ the terms x^n don't go to zero and the series has no chance to converge.

The nice formula for $(I - A)^{-1}$ is the **geometric series of matrices**:

Geometric series

$$(I - A)^{-1} = I + A + A^2 + A^3 + \cdots.$$

If you multiply the series $S = I + A + A^2 + \cdots$ by A , you get the same series except for I . Therefore $S - AS = I$, which is $(I - A)S = I$. The series adds to $S = (I - A)^{-1}$ if it converges. **And it converges if all eigenvalues of A have $|\lambda| < 1$.**

In our case $A \geq 0$. All terms of the series are nonnegative. Its sum is $(I - A)^{-1} \geq 0$.

Example 4 $A = \begin{bmatrix} .2 & .3 & .4 \\ .4 & .4 & .1 \\ .5 & .1 & .3 \end{bmatrix}$ has $\lambda_{\max} = .9$ and $(I - A)^{-1} = \frac{1}{93} \begin{bmatrix} 41 & 25 & 27 \\ 33 & 36 & 24 \\ 34 & 23 & 36 \end{bmatrix}$.

This economy is productive. A is small compared to I , because $\lambda_{\max}$ is .9. To meet the demand y , start from $p = (I - A)^{-1}y$. Then Ap is consumed in production, leaving $p - Ap$. This is $(I - A)p = y$, and the demand is met.

Example 5 $A = \begin{bmatrix} 0 & 4 \\ 1 & 0 \end{bmatrix}$ has $\lambda_{\max} = 2$ and $(I - A)^{-1} = -\frac{1}{3} \begin{bmatrix} 1 & 4 \\ 1 & 1 \end{bmatrix}$.

This consumption matrix A is too large. Demands can't be met, because production consumes more than it yields. The series $I + A + A^2 + \cdots$ does not converge to $(I - A)^{-1}$ because $\lambda_{\max} > 1$. The series is growing while $(I - A)^{-1}$ is actually negative.

In the same way $1 + 2 + 4 + \cdots$ is not really $1/(1 - 2) = -1$. But not entirely false!

Problem Set 10.3

Questions 1–12 are about Markov matrices and their eigenvalues and powers.

- 1 Find the eigenvalues of this Markov matrix (their sum is the trace):

$$A = \begin{bmatrix} .90 & .15 \\ .10 & .85 \end{bmatrix}.$$

What is the steady state eigenvector for the eigenvalue $\lambda_1 = 1$?

- 2 Diagonalize the Markov matrix in Problem 1 to $A = X\Lambda X^{-1}$ by finding its other eigenvector:

$$A = \begin{bmatrix} & \\ & \end{bmatrix} \begin{bmatrix} 1 & \\ & .75 \end{bmatrix} \begin{bmatrix} & \\ & \end{bmatrix}.$$

What is the limit of $A^k = X\Lambda^k X^{-1}$ when $\Lambda^k = \begin{bmatrix} 1 & 0 \\ 0 & .75^k \end{bmatrix}$ approaches $\begin{bmatrix} 1 & 0 \\ 0 & 0 \end{bmatrix}$?

- 3 What are the eigenvalues and steady state eigenvectors for these Markov matrices?

$$A = \begin{bmatrix} 1 & .2 \\ 0 & .8 \end{bmatrix} \quad A = \begin{bmatrix} .2 & 1 \\ .8 & 0 \end{bmatrix} \quad A = \begin{bmatrix} \frac{1}{2} & \frac{1}{4} & \frac{1}{4} \\ \frac{1}{4} & \frac{1}{2} & \frac{1}{4} \\ \frac{1}{4} & \frac{1}{4} & \frac{1}{2} \end{bmatrix}.$$

- 4 For every 4 by 4 Markov matrix, what eigenvector of A^T corresponds to the (known) eigenvalue $\lambda = 1$?

- 5 Every year 2% of young people become old and 3% of old people become dead. (No births.) Find the steady state for

$$\begin{bmatrix} \text{young} \\ \text{old} \\ \text{dead} \end{bmatrix}_{k+1} = \begin{bmatrix} .98 & .00 & 0 \\ .02 & .97 & 0 \\ .00 & .03 & 1 \end{bmatrix} \begin{bmatrix} \text{young} \\ \text{old} \\ \text{dead} \end{bmatrix}_k.$$

- 6 For a Markov matrix, the sum of the components of x equals the sum of the components of Ax . If $Ax = \lambda x$ with $\lambda \neq 1$, prove that the components of this non-steady eigenvector x add to zero.
- 7 Find the eigenvalues and eigenvectors of A . Explain why A^k approaches A^∞ :

$$A = \begin{bmatrix} .8 & .3 \\ .2 & .7 \end{bmatrix} \quad A^\infty = \begin{bmatrix} .6 & .6 \\ .4 & .4 \end{bmatrix}.$$

Challenge problem: Which Markov matrices produce that steady state $(.6, .4)$?

- 8 The steady state eigenvector of a permutation matrix is $(\frac{1}{4}, \frac{1}{4}, \frac{1}{4}, \frac{1}{4})$. This is *not* approached when $u_0 = (0, 0, 0, 1)$. What are u_1 and u_2 and u_3 and u_4 ? What are the four eigenvalues of P , which solve $\lambda^4 = 1$?

Permutation matrix = Markov matrix

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

- 9 Prove that the square of a Markov matrix is also a Markov matrix.
- 10 If $A = \begin{bmatrix} a & b \\ c & d \end{bmatrix}$ is a Markov matrix, its eigenvalues are 1 and _____. The steady state eigenvector is $x_1 =$ _____.
- 11 Complete A to a Markov matrix and find the steady state eigenvector. When A is a symmetric Markov matrix, why is $x_1 = (1, \dots, 1)$ its steady state?

$$A = \begin{bmatrix} .7 & .1 & .2 \\ .1 & .6 & .3 \\ - & - & - \end{bmatrix}.$$

- 12 A Markov differential equation is not $du/dt = Au$ but $du/dt = (A - I)u$. The diagonal is negative, the rest of $A - I$ is positive. The columns add to zero, not 1.

Find λ_1 and λ_2 for $B = A - I = \begin{bmatrix} -.2 & .3 \\ .2 & -.3 \end{bmatrix}$. Why does $A - I$ have $\lambda_1 = 0$?

When $e^{\lambda_1 t}$ and $e^{\lambda_2 t}$ multiply x_1 and x_2 , what is the steady state as $t \rightarrow \infty$?

Questions 13–15 are about linear algebra in economics.

- 13** Each row of the consumption matrix in Example 4 adds to .9. Why does that make $\lambda = .9$ an eigenvalue, and what is the eigenvector?
- 14** Multiply $I + A + A^2 + A^3 + \cdots$ by $I - A$ to get I . The series adds to $(I - A)^{-1}$. For $A = \begin{bmatrix} 0 & \frac{1}{2} \\ 1 & 0 \end{bmatrix}$, find A^2 and A^3 and use the pattern to add up the series.
- 15** For which of these matrices does $I + A + A^2 + \cdots$ yield a nonnegative matrix $(I - A)^{-1}$? Then the economy can meet any demand:

$$A = \begin{bmatrix} 0 & 1 \\ 0 & 0 \end{bmatrix} \quad A = \begin{bmatrix} 0 & 4 \\ .2 & 0 \end{bmatrix} \quad A = \begin{bmatrix} .5 & 1 \\ .5 & 0 \end{bmatrix}.$$

If the demands are $\mathbf{y} = (2, 6)$, what are the vectors $\mathbf{p} = (I - A)^{-1}\mathbf{y}$?

- 16** (Markov again) This matrix has zero determinant. What are its eigenvalues?

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

Find the limits of $A^k \mathbf{u}_0$ starting from $\mathbf{u}_0 = (1, 0, 0)$ and then $\mathbf{u}_0 = (100, 0, 0)$.

- 17** If A is a Markov matrix, why doesn't $I + A + A^2 + \cdots$ add up to $(I - A)^{-1}$?
- 18** For the Leslie matrix show that $\det(A - \lambda I) = 0$ gives $F_1 \lambda^2 + F_2 P_1 \lambda + F_3 P_1 P_2 = \lambda^3$. The right side λ^3 is larger as $\lambda \rightarrow \infty$. The left side is larger at $\lambda = 1$ if $F_1 + F_2 P_1 + F_3 P_1 P_2 > 1$. In that case the two sides are equal at an eigenvalue $\lambda_{\max} > 1$: *growth*.
- 19** **Sensitivity of eigenvalues:** A matrix change ΔA produces eigenvalue changes $\Delta \lambda$. Those changes $\Delta \lambda_1, \dots, \Delta \lambda_n$ are on the diagonal of $(X^{-1} \Delta A X)$. **Challenge:** Start from $AX = X\Lambda$. The eigenvectors and eigenvalues change by ΔX and $\Delta \Lambda$:

$$(A + \Delta A)(X + \Delta X) = (X + \Delta X)(\Lambda + \Delta \Lambda) \text{ becomes } A(\Delta X) + (\Delta A)X = X(\Delta \Lambda) + (\Delta X)\Lambda.$$

Small terms $(\Delta A)(\Delta X)$ and $(\Delta X)(\Delta \Lambda)$ are ignored. Multiply the last equation by X^{-1} . From the inner terms, the diagonal part of $X^{-1}(\Delta A)X$ gives $\Delta \Lambda$ as we want. Why do the outer terms $X^{-1} \Delta A \Delta X$ and $X^{-1} \Delta X \Delta \Lambda$ cancel on the diagonal?

$$\text{Explain } X^{-1}A = \Lambda X^{-1} \text{ and then } \mathbf{diag}(\Lambda X^{-1} \Delta X) = \mathbf{diag}(X^{-1} \Delta X \Lambda).$$

- 20** Suppose $B > A > 0$, meaning that each $b_{ij} > a_{ij} > 0$. How does the Perron-Frobenius discussion show that $\lambda_{\max}(B) > \lambda_{\max}(A)$?

10.4 Linear Programming

Linear programming is linear algebra plus two new ideas: *inequalities* and *minimization*. The starting point is still a matrix equation $Ax = b$. But the only acceptable solutions are *nonnegative*. We require $x \geq 0$ (meaning that no component of x can be negative). The matrix has $n > m$, more unknowns than equations. If there are any solutions $x \geq 0$ to $Ax = b$, there are probably a lot. Linear programming picks the solution $x^* \geq 0$ that minimizes the cost:

The cost is $c_1x_1 + \cdots + c_nx_n$. The winning vector x^ is the nonnegative solution of $Ax = b$ that has smallest cost.*

Thus a linear programming problem starts with a matrix A and two vectors b and c :

- i) A has $n > m$: for example $A = \begin{bmatrix} 1 & 1 & 2 \end{bmatrix}$ (one equation, three unknowns)
- ii) b has m components for m equations $Ax = b$: for example $b = \begin{bmatrix} 4 \end{bmatrix}$
- iii) The **cost vector** c has n components: for example $c = \begin{bmatrix} 5 & 3 & 8 \end{bmatrix}$.

Then the problem is to minimize $c \cdot x$ subject to the requirements $Ax = b$ and $x \geq 0$:

Minimize $5x_1 + 3x_2 + 8x_3$ **subject to** $x_1 + x_2 + 2x_3 = 4$ **and** $x_1, x_2, x_3 \geq 0$.

We jumped right into the problem, without explaining where it comes from. Linear programming is actually the most important application of mathematics to management. Development of the fastest algorithm and fastest code is highly competitive. You will see that finding x^* is harder than solving $Ax = b$, because of the extra requirements: $x^* \geq 0$ and minimum cost $c^T x^*$. We will explain the background, and the famous *simplex method*, and *interior point methods*, after solving the example.

Look first at the “constraints”: $Ax = b$ and $x \geq 0$. The equation $x_1 + x_2 + 2x_3 = 4$ gives a plane in three dimensions. The nonnegativity $x_1 \geq 0, x_2 \geq 0, x_3 \geq 0$ chops the plane down to a triangle. The solution x^* must lie in the triangle PQR in Figure 8.6.

Inside that triangle, all components of x are positive. On the edges of PQR , one component is zero. At the corners P and Q and R , two components are zero. **The optimal solution x^* will be one of those corners!** We will now show why.

The triangle contains all vectors x that satisfy $Ax = b$ and $x \geq 0$. Those x 's are called *feasible points*, and the triangle is the **feasible set**. These points are the allowed candidates in the minimization of $c \cdot x$, which is the final step:

Find x^* in the triangle PQR to minimize the cost $5x_1 + 3x_2 + 8x_3$.

The vectors that have *zero* cost lie on the plane $5x_1 + 3x_2 + 8x_3 = 0$. That plane does not meet the triangle. We cannot achieve zero cost, while meeting the requirements on x . So increase the cost C until the plane $5x_1 + 3x_2 + 8x_3 = C$ does meet the triangle. As C increases, we have *parallel planes moving toward the triangle*.

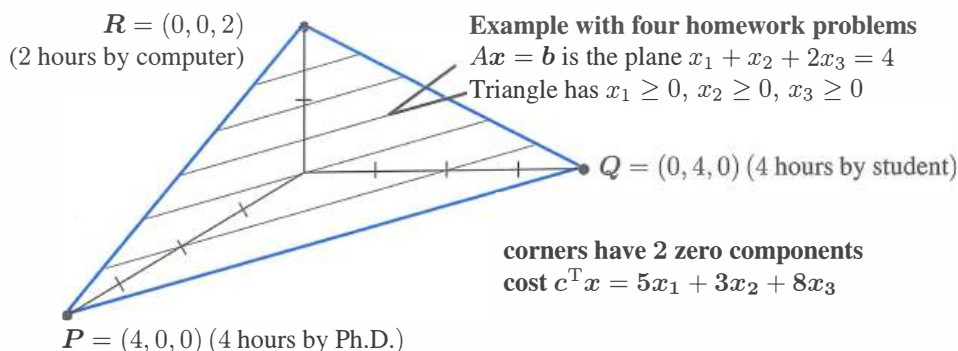


Figure 10.5: The triangle contains all nonnegative solutions: $Ax = b$ and $x \geq 0$. The lowest cost solution x^* is a corner P , Q , or R of this feasible set.

The first plane $5x_1 + 3x_2 + 8x_3 = C$ to touch the triangle has minimum cost C . **The point where it touches is the solution x^* .** This touching point must be one of the corners P or Q or R . A moving plane could not reach the inside of the triangle before it touches a corner! So check the cost $5x_1 + 3x_2 + 8x_3$ at each corner:

$P = (4, 0, 0)$ costs 20	$Q = (0, 4, 0)$ costs 12	$R = (0, 0, 2)$ costs 16.
--------------------------	--------------------------	---------------------------

The winner is Q . Then $x^* = (0, 4, 0)$ solves the linear programming problem.

If the cost vector c is changed, the parallel planes are tilted. For small changes, Q is still the winner. For the cost $c \cdot x = 5x_1 + 4x_2 + 7x_3$, the optimum x^* moves to $R = (0, 0, 2)$. The minimum cost is now $7 \cdot 2 = 14$.

Note 1 Some linear programs *maximize profit* instead of minimizing cost. The mathematics is almost the same. The parallel planes start with a large value of C , instead of a small value. They move toward the origin (instead of away), as C gets smaller. *The first touching point is still a corner.*

Note 2 The requirements $Ax = b$ and $x \geq 0$ could be impossible to satisfy. The equation $x_1 + x_2 + x_3 = -1$ cannot be solved with $x \geq 0$. *That feasible set is empty.*

Note 3 It could also happen that the feasible set is *unbounded*. If the requirement is $x_1 + x_2 - 2x_3 = 4$, the large positive vector $(100, 100, 98)$ is now a candidate. So is the larger vector $(1000, 1000, 998)$. The plane $Ax = b$ is no longer chopped off to a triangle. The two corners P and Q are still candidates for x^* , but R moved to infinity.

Note 4 With an unbounded feasible set, the minimum cost could be $-\infty$ (minus infinity). Suppose the cost is $-x_1 - x_2 + x_3$. Then the vector $(100, 100, 98)$ costs $C = -102$. The vector $(1000, 1000, 998)$ costs $C = -1002$. We are being paid to include x_1 and x_2 , instead of paying a cost. In realistic applications this will not happen. But it is theoretically possible that A , b , and c can produce unexpected triangles and costs.

The Primal and Dual Problems

This first problem will fit A, b, c in that example. The unknowns x_1, x_2, x_3 represent hours of work by a Ph.D. and a student and a machine. The costs per hour are \$5, \$3, and \$8. (*I apologize for such low pay.*) The number of hours cannot be negative: $x_1 \geq 0, x_2 \geq 0, x_3 \geq 0$. The Ph.D. and the student get through one homework problem per hour. *The machine solves two problems in one hour.* In principle they can share out the homework, which has four problems to be solved: $x_1 + x_2 + 2x_3 = 4$.

The problem is to finish the four problems at minimum cost $c^T x$.

If all three are working, the job takes one hour: $x_1 = x_2 = x_3 = 1$. The cost is $5 + 3 + 8 = 16$. But certainly the Ph.D. should be put out of work by the student (who is just as fast and costs less—this problem is getting realistic). When the student works two hours and the machine works one, the cost is $6 + 8$ and all four problems get solved. We are on the edge QR of the triangle because the Ph.D. is not working: $x_1 = 0$. But the best point is all work by student (at Q) or all work by machine (at R). In this example the student solves four problems in four hours for \$12—the minimum cost.

With only one equation in $Ax = b$, the corner $(0, 4, 0)$ has only one nonzero component. **When $Ax = b$ has m equations, corners have m nonzeros.** We solve $Ax = b$ for those m variables, with $n - m$ free variables set to zero. But unlike Chapter 3, **we don't know which m variables to choose.**

The number of possible corners is the number of ways to choose m components out of n . This number “ n choose m ” is heavily involved in gambling and probability. With $n = 20$ unknowns and $m = 8$ equations (still small numbers), the “feasible set” can have $20!/8!12!$ corners. That number is $(20)(19) \cdots (13) = 5,079,110,400$.

Checking three corners for the minimum cost was fine. Checking five billion corners is not the way to go. The simplex method described below is much faster.

The Dual Problem In linear programming, problems come in pairs. There is a minimum problem and a maximum problem—the original and its “dual.” The original problem was specified by a matrix A and two vectors b and c . The dual problem transposes A and switches b and c : **Maximize $b \cdot y$.** Here is the dual to our example:

A cheater offers to solve homework problems by selling the answers.

The charge is y dollars per problem, or $4y$ altogether. (Note how $b = 4$ has gone into the cost.) The cheater must be as cheap as the Ph.D. or student or machine: $y \leq 5$ and $y \leq 3$ and $2y \leq 8$. (Note how $c = (5, 3, 8)$ has gone into inequality constraints). The cheater maximizes the income $4y$.

Dual Problem **Maximize $b \cdot y$ subject to $A^T y \leq c$.**

The maximum occurs when $y = 3$. The income is $4y = 12$. The maximum in the dual problem (\$12) equals the minimum in the original (\$12). *Max = min* is duality.

If either problem has a best vector (x^ or y^*) then so does the other.*

Minimum cost $c \cdot x^*$ equals maximum income $b \cdot y^*$

This book started with a row picture and a column picture. The first “duality theorem” was about rank: The number of independent rows equals the number of independent columns. That theorem, like this one, was easy for small matrices. Minimum cost = maximum income is proved in our text *Linear Algebra and Its Applications*. One line will establish the easy half of the theorem: ***The cheater’s income $b^T y$ cannot exceed the honest cost:***

$$\text{If } Ax = b, x \geq 0, A^T y \leq c \text{ then } b^T y = (Ax)^T y = x^T (A^T y) \leq x^T c. \quad (1)$$

The full duality theorem says that when $b^T y$ reaches its maximum and $x^T c$ reaches its minimum, they are equal: $b \cdot y^* = c \cdot x^*$. Look at the last step in (1), with $\leq$ sign:

The dot product of $x \geq 0$ and $s = c - A^T y \geq 0$ gave $x^T s \geq 0$. This is $x^T A^T y \leq x^T c$.

Equality needs $x^T s = 0$ So the optimal solution has $x_j^* = 0$ or $s_j^* = 0$ for each j .

The Simplex Method

Elimination is the workhorse for linear equations. The simplex method is the workhorse for linear inequalities. We cannot give the simplex method as much space as elimination, but the idea can be clear. *The simplex method goes from one corner to a neighboring corner of lower cost.* Eventually (and quite soon in practice) it reaches the corner of minimum cost.

A **corner** is a vector $x \geq 0$ that satisfies the m equations $Ax = b$ with at most m positive components. *The other $n - m$ components are zero.* (Those are the free variables. Back substitution gives the m basic variables. All variables must be nonnegative or x is a false corner.) For a *neighboring corner*, one zero component of x becomes positive and one positive component becomes zero.

The simplex method must decide which component “enters” by becoming positive, and which component “leaves” by becoming zero. That exchange is chosen so as to lower the total cost. This is one step of the simplex method, moving toward x^* .

Here is the overall plan. Look at each zero component at the current corner. If it changes from 0 to 1, the other nonzeros have to adjust to keep $Ax = b$. Find the new x by back substitution and compute the change in the total cost $c \cdot x$. This change is the “reduced cost” r of the new component. The **entering variable** is the one that gives the **most negative r** . This is the greatest cost reduction for a single unit of a new variable.

Example 1 Suppose the current corner is $P = (4, 0, 0)$, with the Ph.D. doing all the work (the cost is \$20). If the student works one hour, the cost of $x = (3, 1, 0)$ is down to \$18. The reduced cost is $r = -2$. If the machine works one hour, then $x = (2, 0, 1)$ also

costs \$18. The reduced cost is also $r = -2$. In this case the simplex method can choose either the student or the machine as the entering variable.

Even in this small example, the first step may not go immediately to the best x^* . The method chooses the entering variable before it knows how much of that variable to include. We computed r when the entering variable changes from 0 to 1, but one unit may be too much or too little. The method now chooses the leaving variable (the Ph.D.). It moves to corner Q or R in the figure.

The more of the entering variable we include, the lower the cost. This has to stop when one of the positive components (which are adjusting to keep $Ax = b$) hits zero. *The leaving variable is the first positive x_i to reach zero.* When that happens, a neighboring corner has been found. Then start again (from the new corner) to find the next variables to enter and leave.

When all reduced costs are positive, the current corner is the optimal x^* . No zero component can become positive without increasing $c \cdot x$. No new variable should enter. The problem is solved (and we can show that y^* is found too).

Note Generally x^* is reached in αn steps, where α is not large. But examples have been invented which use an exponential number of simplex steps. Eventually a different approach was developed, which is guaranteed to reach x^* in fewer (but more difficult) steps. The new methods travel through the *interior* of the feasible set.

Example 2 Minimize the cost $c \cdot x = 3x_1 + x_2 + 9x_3 + x_4$. The constraints are $x \geq 0$ and two equations $Ax = b$:

$$\begin{array}{rcl} x_1 + 2x_3 + x_4 = 4 & m = 2 & \text{equations} \\ x_2 + x_3 - x_4 = 2 & n = 4 & \text{unknowns.} \end{array}$$

A starting corner is $x = (4, 2, 0, 0)$ which costs $c \cdot x = 14$. It has $m = 2$ nonzeros and $n - m = 2$ zeros. The zeros are x_3 and x_4 . The question is whether x_3 or x_4 should enter (become nonzero). Try one unit of each of them:

If $x_3 = 1$ and $x_4 = 0$, then $x = (2, 1, 1, 0)$ costs 16.

If $x_4 = 1$ and $x_3 = 0$, then $x = (3, 3, 0, 1)$ costs 13.

Compare those costs with 14. The reduced cost of x_3 is $r = 2$, positive and useless. The reduced cost of x_4 is $r = -1$, negative and helpful. *The entering variable is x_4 .*

How much of x_4 can enter? One unit of x_4 made x_1 drop from 4 to 3. Four units will make x_1 drop from 4 to zero (while x_2 increases all the way to 6). *The leaving variable is x_1 .* The new corner is $x = (0, 6, 0, 4)$, which costs only $c \cdot x = 10$. This is the optimal x^* , but to know that we have to try another simplex step from $(0, 6, 0, 4)$. Suppose x_1 or x_3 tries to enter:

Start from the corner $(0, 6, 0, 4)$	If $x_1 = 1$ and $x_3 = 0$,	then $x = (1, 5, 0, 3)$ costs 11.
	If $x_3 = 1$ and $x_1 = 0$,	then $x = (0, 3, 1, 2)$ costs 14.

Those costs are higher than 10. Both r 's are positive—it does not pay to move. The current corner $(0, 6, 0, 4)$ is the solution $\mathbf{x}^*$.

These calculations can be streamlined. Each simplex step solves three linear systems with the same matrix B . (This is the m by m matrix that keeps the m basic columns of A .) When a column enters and an old column leaves, there is a quick way to update B^{-1} . That is how most codes organize the simplex method.

Our text on *Computational Science and Engineering* includes a short code with comments. (The code is also on math.mit.edu/cse) The best $\mathbf{y}^*$ solves m equations $A^T \mathbf{y}^* = \mathbf{c}$ in the m components that are nonzero in $\mathbf{x}^*$. Then we have optimality $\mathbf{x}^T \mathbf{s} = 0$ and this is duality: *Either $x_j^* = 0$ or the “slack” in $\mathbf{s}^* = \mathbf{c} - A^T \mathbf{y}^*$ has $s_j^* = 0$.*

When $\mathbf{x}^* = (0, 4, 0)$ was the optimal corner $\mathbf{Q}$, the cheater's price was set by $y^* = 3$.

Interior Point Methods

The simplex method moves along the edges of the feasible set, eventually reaching the optimal corner $\mathbf{x}^*$. **Interior point methods move inside the feasible set** (where $\mathbf{x} > 0$). These methods hope to go more directly to $\mathbf{x}^*$. They work well.

One way to stay inside is to put a barrier at the boundary. Add extra cost as a *logarithm that blows up* when any variable x_j touches zero. The best vector has $\mathbf{x} > 0$. The number θ is a small parameter that we move toward zero.

Barrier problem	Minimize $\mathbf{c}^T \mathbf{x} - \theta (\log x_1 + \cdots + \log x_n)$ with $A\mathbf{x} = \mathbf{b}$	(2)
------------------------	--	-----

This cost is nonlinear (but linear programming is already nonlinear from inequalities). The constraints $x_j \geq 0$ are not needed because $\log x_j$ becomes infinite at $x_j = 0$.

The barrier gives an *approximate problem* for each θ . The m constraints $A\mathbf{x} = \mathbf{b}$ have Lagrange multipliers $y_1, \dots, y_m$. This is the good way to deal with constraints.

$\mathbf{y}$ from Lagrange $L(\mathbf{x}, \mathbf{y}, \theta) = \mathbf{c}^T \mathbf{x} - \theta (\sum \log x_i) - \mathbf{y}^T (A\mathbf{x} - \mathbf{b})$ (3)

$\partial L / \partial \mathbf{y} = 0$ brings back $A\mathbf{x} = \mathbf{b}$. The derivatives $\partial L / \partial x_j$ are interesting !

Optimality in barrier pbm	$\frac{\partial L}{\partial x_j} = c_j - \frac{\theta}{x_j} - (A^T \mathbf{y})_j = 0$	which is $x_j s_j = \theta$.	(4)
----------------------------------	---	-------------------------------	-----

The true problem has $x_j s_j = 0$. The barrier problem has $x_j s_j = \theta$. The solutions $\mathbf{x}^*(\theta)$ lie on the **central path** to $\mathbf{x}^*(0)$. Those n optimality equations $x_j s_j = \theta$ are nonlinear, and we solve them iteratively by Newton's method.

The current $\mathbf{x}, \mathbf{y}, \mathbf{s}$ will satisfy $A\mathbf{x} = \mathbf{b}$, $\mathbf{x} \geq 0$ and $A^T \mathbf{y} + \mathbf{s} = \mathbf{c}$, but not $x_j s_j = \theta$. Newton's method takes a step $\Delta \mathbf{x}, \Delta \mathbf{y}, \Delta \mathbf{s}$. By ignoring the second-order term $\Delta \mathbf{x} \Delta \mathbf{s}$

in $(x + \Delta x)(s + \Delta s) = \theta$, the corrections in x, y, s come from linear equations:

$$\begin{aligned} A \Delta x &= 0 \\ \text{Newton step} \quad A^T \Delta y + \Delta s &= 0 \\ s_j \Delta x_j + x_j \Delta s_j &= \theta - x_j s_j \end{aligned} \quad (5)$$

Newton iteration has quadratic convergence for each θ , and then θ approaches zero. The duality gap $x^T s$ generally goes below 10^{-8} after 20 to 60 steps. The explanation in my *Computational Science and Engineering* textbook takes one Newton step in detail, for the example with four homework problems. I didn't intend that the student should end up doing all the work, but x^* turned out that way.

This interior point method is used almost “as is” in commercial software, for a large class of linear and nonlinear optimization problems.

Problem Set 10.4

- 1 Draw the region in the xy plane where $x + 2y = 6$ and $x \geq 0$ and $y \geq 0$. Which point in this “feasible set” minimizes the cost $c = x + 3y$? Which point gives maximum cost? Those points are at corners.
- 2 Draw the region in the xy plane where $x + 2y \leq 6$, $2x + y \leq 6$, $x \geq 0$, $y \geq 0$. It has four corners. Which corner minimizes the cost $c = 2x - y$?
- 3 What are the corners of the set $x_1 + 2x_2 - x_3 = 4$ with x_1, x_2, x_3 all ≥ 0 ? Show that the cost $x_1 + 2x_3$ can be very negative in this feasible set. This is an example of unbounded cost: no minimum.
- 4 Start at $x = (0, 0, 2)$ where the machine solves all four problems for \$16. Move to $x = (0, 1, \quad)$ to find the reduced cost r (the savings per hour) for work by the student. Find r for the Ph.D. by moving to $x = (1, 0, \quad)$ with 1 hour of Ph.D. work.
- 5 Start Example 1 from the Ph.D. corner $(4, 0, 0)$ with c changed to $[5 \ 3 \ 7]$. Show that r is better for the machine even when the total cost is lower for the student. The simplex method takes two steps, first to the machine and then to the student for x^* .
- 6 Choose a different cost vector c so the Ph.D. gets the job. Rewrite the dual problem (maximum income to the cheater).
- 7 A six-problem homework on which the Ph.D. is fastest gives a second constraint $2x_1 + x_2 + x_3 = 6$. Then $x = (2, 2, 0)$ shows two hours of work by Ph.D. and student on each homework. Does this x minimize the cost $c^T x$ with $c = (5, 3, 8)$?
- 8 These two problems are also dual. Prove weak duality, that always $y^T b \leq c^T x$:

Primal problem Minimize $c^T x$ with $Ax \geq b$ and $x \geq 0$.

Dual problem Maximize $y^T b$ with $A^T y \leq c$ and $y \geq 0$.

10.5 Fourier Series: Linear Algebra for Functions

This section goes from finite dimensions to *infinite* dimensions. I want to explain linear algebra in infinite-dimensional space, and to show that it still works. First step: look back. This book began with vectors and dot products and linear combinations. We begin by converting those basic ideas to the infinite case—then the rest will follow.

What does it mean for a vector to have infinitely many components? There are two different answers, both good:

1. The vector is infinitely long: $\mathbf{v} = (v_1, v_2, v_3, \dots)$. It could be $(1, \frac{1}{2}, \frac{1}{4}, \dots)$.
2. The vector is a function $f(x)$. It could be $\mathbf{v} = \sin x$.

We will go both ways. Then the idea of a Fourier series will connect them.

After vectors come *dot products*. The natural dot product of two infinite vectors $(v_1, v_2, \dots)$ and $(w_1, w_2, \dots)$ is an infinite series:

$$\text{Dot product} \quad \mathbf{v} \cdot \mathbf{w} = v_1 w_1 + v_2 w_2 + \dots \quad (1)$$

This brings a new question, which never occurred to us for vectors in $\mathbf{R}^n$. Does this infinite sum add up to a finite number? Does the series converge? Here is the first and biggest difference between finite and infinite.

When $\mathbf{v} = \mathbf{w} = (1, 1, 1, \dots)$, the sum certainly does not converge. In that case $\mathbf{v} \cdot \mathbf{w} = 1 + 1 + 1 + \dots$ is infinite. Since $\mathbf{v}$ equals $\mathbf{w}$, we are really computing $\mathbf{v} \cdot \mathbf{v} = \|\mathbf{v}\|^2$, the length squared. The vector $(1, 1, 1, \dots)$ has infinite length. *We don't want that vector.* Since we are making the rules, we don't have to include it. The only vectors to be allowed are those with finite length:

DEFINITION The vector $\mathbf{v} = (v_1, v_2, \dots)$ and the function $f(x)$ are in our infinite-dimensional “*Hilbert spaces*” if and only if their lengths $\|\mathbf{v}\|$ and $\|f\|$ are finite:

$$\|\mathbf{v}\|^2 = \mathbf{v} \cdot \mathbf{v} = v_1^2 + v_2^2 + v_3^2 + \dots \quad \text{must add to a finite number.}$$

$$\|f\|^2 = (f, f) = \int_0^{2\pi} |f(x)|^2 dx \quad \text{must be a finite integral.}$$

Example 1 The vector $\mathbf{v} = (1, \frac{1}{2}, \frac{1}{4}, \dots)$ is included in Hilbert space, because its length is $2/\sqrt{3}$. We have a geometric series that adds to $4/3$. The length of $\mathbf{v}$ is the square root:

$$\text{Length squared} \quad \mathbf{v} \cdot \mathbf{v} = 1 + \frac{1}{4} + \frac{1}{16} + \dots = \frac{1}{1 - \frac{1}{4}} = \frac{4}{3}.$$

Question If $\mathbf{v}$ and $\mathbf{w}$ have finite length, how large can their dot product be?

Answer The sum $\mathbf{v} \cdot \mathbf{w} = v_1 w_1 + v_2 w_2 + \dots$ also adds to a finite number. We can safely take dot products. The Schwarz inequality is still true:

$$\text{Schwarz inequality} \quad |\mathbf{v} \cdot \mathbf{w}| \leq \|\mathbf{v}\| \|\mathbf{w}\|. \quad (2)$$

The ratio of $\mathbf{v} \cdot \mathbf{w}$ to $\|\mathbf{v}\| \|\mathbf{w}\|$ is still the cosine of θ (the angle between $\mathbf{v}$ and $\mathbf{w}$). Even in infinite-dimensional space, $|\cos \theta|$ is not greater than 1.

Now change over to functions. Those are the “vectors.” The space of functions $f(x)$, $g(x)$, $h(x)$, . . . defined for $0 \leq x \leq 2\pi$ must be somehow bigger than $\mathbf{R}^n$. **What is the dot product of $f(x)$ and $g(x)$? What is the length of $f(x)$?**

Key point in the continuous case: *Sums are replaced by integrals.* Instead of a sum of v_j times w_j , the dot product is an integral of $f(x)$ times $g(x)$. Change the “dot” to parentheses with a comma, and change the words “dot product” to *inner product*:

DEFINITION The *inner product* of $f(x)$ and $g(x)$, and the *length squared* of $f(x)$, are

$$(f, g) = \int_0^{2\pi} f(x)g(x) dx \quad \text{and} \quad \|f\|^2 = \int_0^{2\pi} (f(x))^2 dx. \quad (3)$$

The interval $[0, 2\pi]$ where the functions are defined could change to a different interval like $[0, 1]$ or $(-\infty, \infty)$. We chose 2π because our first examples are $\sin x$ and $\cos x$.

Example 2 The length of $f(x) = \sin x$ comes from its inner product with itself:

$$(f, f) = \int_0^{2\pi} (\sin x)^2 dx = \pi. \quad \text{The length of } \sin x \text{ is } \sqrt{\pi}.$$

That is a standard integral in calculus—not part of linear algebra. By writing $\sin^2 x$ as $\frac{1}{2} - \frac{1}{2} \cos 2x$, we see it go above and below its average value $\frac{1}{2}$. Multiply that average by the interval length 2π to get the answer π .

More important: $\sin x$ **and** $\cos x$ **are orthogonal in function space**: $(f, g) = 0$

$$\text{Inner product is zero} \quad \int_0^{2\pi} \sin x \cos x dx = \int_0^{2\pi} \frac{1}{2} \sin 2x dx = \left[-\frac{1}{4} \cos 2x\right]_0^{2\pi} = 0. \quad (4)$$

This zero is no accident. It is highly important to science. The orthogonality goes beyond the two functions $\sin x$ and $\cos x$, to an infinite list of sines and cosines. The list contains $\cos 0x$ (which is 1), $\sin x$, $\cos x$, $\sin 2x$, $\cos 2x$, $\sin 3x$, $\cos 3x$, . . .

Every function in that list is orthogonal to every other function in the list.

Fourier Series

The Fourier series of a function $f(x)$ is its expansion into sines and cosines:

$$f(x) = a_0 + a_1 \cos x + b_1 \sin x + a_2 \cos 2x + b_2 \sin 2x + \cdots. \quad (5)$$

We have an orthogonal basis! The vectors in “function space” are combinations of the sines and cosines. On the interval from $x = 2\pi$ to $x = 4\pi$, all our functions repeat what they did from 0 to 2π . They are “*periodic*.” The distance between repetitions is the period 2π .

Remember: The list is infinite. The Fourier series is an infinite series. We avoided the vector $\mathbf{v} = (1, 1, 1, \dots)$ because its length is infinite, now we avoid a function like $\frac{1}{2} + \cos x + \cos 2x + \cos 3x + \dots$. (Note: This is π times the famous **delta function** $\delta(x)$. It is an infinite “spike” above a single point. At $x = 0$ its height $\frac{1}{2} + 1 + 1 + \dots$ is infinite. At all points inside $0 < x < 2\pi$ the series adds in some average way to zero.) The integral of $\delta(x)$ is 1. But $\int \delta^2(x) = \infty$, so delta functions are not allowed into Hilbert space.

Compute the length of a typical sum $f(x)$:

$$\begin{aligned}(f, f) &= \int_0^{2\pi} (a_0 + a_1 \cos x + b_1 \sin x + a_2 \cos 2x + \dots)^2 dx \\ &= \int_0^{2\pi} (a_0^2 + a_1^2 \cos^2 x + b_1^2 \sin^2 x + a_2^2 \cos^2 2x + \dots) dx \\ \|f\|^2 &= 2\pi a_0^2 + \pi(a_1^2 + b_1^2 + a_2^2 + \dots).\end{aligned}\quad (6)$$

The step from line 1 to line 2 used orthogonality. All products like $\cos x \cos 2x$ integrate to give zero. Line 2 contains what is left—the integrals of each sine and cosine squared. Line 3 evaluates those integrals. (The integral of 1^2 is 2π , when all other integrals give π .) If we divide by their lengths, our functions become *orthonormal*:

$$\frac{1}{\sqrt{2\pi}}, \frac{\cos x}{\sqrt{\pi}}, \frac{\sin x}{\sqrt{\pi}}, \frac{\cos 2x}{\sqrt{\pi}}, \dots \text{ is an orthonormal basis for our function space.}$$

These are unit vectors. We could combine them with coefficients $A_0, A_1, B_1, A_2, \dots$ to yield a function $F(x)$. Then the 2π and the π 's drop out of the formula for length.

$$\text{Function length} = \text{vector length} \quad \|F\|^2 = (F, F) = A_0^2 + A_1^2 + B_1^2 + A_2^2 + \dots \quad (7)$$

Here is the important point, for $f(x)$ as well as $F(x)$. *The function has finite length exactly when the vector of coefficients has finite length.* Fourier series gives us a perfect match between the Hilbert spaces for functions and for vectors. The function is in L^2 , its Fourier coefficients are in ℓ^2 .

The function space contains $f(x)$ exactly when the Hilbert space contains the vector $\mathbf{v} = (a_0, a_1, b_1, \dots)$ of Fourier coefficients of $f(x)$. Both must have finite length.

Example 3 Suppose $f(x)$ is a “square wave,” equal to 1 for $0 \leq x < \pi$. Then $f(x)$ drops to -1 for $\pi \leq x < 2\pi$. The $+1$ and -1 repeat forever. This $f(x)$ is an odd function like the sines, and all its cosine coefficients are zero. We will find its Fourier series, containing only sines:

$$\text{Square wave} \quad f(x) = \frac{4}{\pi} \left[\frac{\sin x}{1} + \frac{\sin 3x}{3} + \frac{\sin 5x}{5} + \dots \right]. \quad (8)$$

The length of this function is $\sqrt{2\pi}$, because at every point $(f(x))^2$ is $(-1)^2$ or $(+1)^2$:

$$\|f\|^2 = \int_0^{2\pi} (f(x))^2 dx = \int_0^{2\pi} 1 dx = 2\pi.$$

At $x = 0$ the sines are zero and the Fourier series gives zero. This is half way up the jump from -1 to $+1$. The Fourier series is also interesting when $x = \frac{\pi}{2}$. At this point the square wave equals 1, and the sines in (8) alternate between $+1$ and -1 :

$$\text{Formula for } \pi \quad 1 = \frac{4}{\pi} \left(1 - \frac{1}{3} + \frac{1}{5} - \frac{1}{7} + \cdots \right). \quad (9)$$

Multiply by π to find a magical formula $4(1 - \frac{1}{3} + \frac{1}{5} - \frac{1}{7} + \cdots)$ for that famous number.

The Fourier Coefficients

How do we find the a 's and b 's which multiply the cosines and sines? For a given function $f(x)$, we are asking for its Fourier coefficients a_k and b_k :

$$\text{Fourier series} \quad f(x) = a_0 + a_1 \cos x + b_1 \sin x + a_2 \cos 2x + \cdots$$

Here is the way to find a_1 . *Multiply both sides by $\cos x$. Then integrate from 0 to 2π .* The key is orthogonality! All integrals on the right side are zero, except for $\cos^2 x$:

$$\text{For coefficient } a_1 \quad \int_0^{2\pi} f(x) \cos x \, dx = \int_0^{2\pi} a_1 \cos^2 x \, dx = \pi a_1. \quad (10)$$

Divide by π and you have a_1 . To find any other a_k , multiply the Fourier series by $\cos kx$. Integrate from 0 to 2π . Use orthogonality, so only the integral of $a_k \cos^2 kx$ is left. That integral is πa_k , and divide by π :

$$a_k = \frac{1}{\pi} \int_0^{2\pi} f(x) \cos kx \, dx \quad \text{and similarly} \quad b_k = \frac{1}{\pi} \int_0^{2\pi} f(x) \sin kx \, dx. \quad (11)$$

The exception is a_0 . This time we multiply by $\cos 0x = 1$. The integral of 1 is 2π :

$$\text{Constant term} \quad a_0 = \frac{1}{2\pi} \int_0^{2\pi} f(x) \cdot 1 \, dx = \text{average value of } f(x). \quad (12)$$

I used those formulas to find the Fourier coefficients for the square wave in equation (8). The integral of $f(x) \cos kx$ was zero. The integral of $f(x) \sin kx$ was $4/k$ for odd k .

Compare Linear Algebra in $\mathbf{R}^n$

Infinite-dimensional Hilbert space is very much like the n -dimensional space $\mathbf{R}^n$. Suppose the nonzero vectors $v_1, \dots, v_n$ are orthogonal in $\mathbf{R}^n$. We want to write the vector b (instead of the function $f(x)$) as a combination of those v 's:

$$\text{Finite orthogonal series} \quad b = c_1 v_1 + c_2 v_2 + \cdots + c_n v_n. \quad (13)$$

Multiply both sides by v_1^T . Use orthogonality, so $v_1^T v_2 = 0$. Only the c_1 term is left:

$$\text{Coefficient } c_1 \quad v_1^T b = c_1 v_1^T v_1 + 0 + \cdots + 0. \quad \text{Therefore } c_1 = v_1^T b / v_1^T v_1. \quad (14)$$

The denominator $v_1^T v_1$ is the length squared, like π in equation (11). The numerator $v_1^T b$ is the inner product like $\int f(x) \cos kx \, dx$. **Coefficients are easy to find when the**

basis vectors are orthogonal. We are just doing one-dimensional projections, to find the components along each basis vector.

The formulas are even better when the vectors are orthonormal. Then we have unit vectors in Q . The denominators $\mathbf{v}_k^T \mathbf{v}_k$ are all 1. You know $c_k = \mathbf{v}_k^T \mathbf{b}$ in another form:

$$\text{Equation for } c\text{'s} \quad c_1 \mathbf{v}_1 + \cdots + c_n \mathbf{v}_n = \mathbf{b} \quad \text{or} \quad \begin{bmatrix} \mathbf{v}_1 & \cdots & \mathbf{v}_n \end{bmatrix} \begin{bmatrix} c_1 \\ \vdots \\ c_n \end{bmatrix} = \mathbf{b}.$$

$$Q\mathbf{c} = \mathbf{b} \quad \text{yields} \quad \mathbf{c} = Q^T \mathbf{b}. \quad \text{Row by row this is } c_k = \mathbf{q}_k^T \mathbf{b}.$$

Fourier series is like having a matrix with infinitely many orthogonal columns. Those columns are the basis functions $1, \cos x, \sin x, \dots$. After dividing by their lengths we have an “infinite orthogonal matrix.” Its inverse is its transpose, Q^T . Orthogonality is what reduces a series of terms to one single term, when we integrate.

Problem Set 10.5

- 1 Integrate the trig identity $2 \cos jx \cos kx = \cos(j+k)x + \cos(j-k)x$ to show that $\cos jx$ is orthogonal to $\cos kx$, provided $j \neq k$. What is the result when $j = k$?
- 2 Show that $1, x$, and $x^2 - \frac{1}{3}$ are orthogonal, when the integration is from $x = -1$ to $x = 1$. Write $f(x) = 2x^2$ as a combination of those orthogonal functions.
- 3 Find a vector $(w_1, w_2, w_3, \dots)$ that is orthogonal to $\mathbf{v} = (1, \frac{1}{2}, \frac{1}{4}, \dots)$. Compute its length $\|\mathbf{w}\|$.
- 4 The first three *Legendre polynomials* are $1, x$, and $x^2 - \frac{1}{3}$. Choose c so that the fourth polynomial $x^3 - cx$ is orthogonal to the first three. All integrals go from -1 to 1 .
- 5 For the square wave $f(x)$ in Example 3 jumping from 1 to -1 , show that

$$\int_0^{2\pi} f(x) \cos x \, dx = 0 \quad \int_0^{2\pi} f(x) \sin x \, dx = 4 \quad \int_0^{2\pi} f(x) \sin 2x \, dx = 0.$$

Which three Fourier coefficients come from those integrals?

- 6 The square wave has $\|f\|^2 = 2\pi$. Then (6) gives what remarkable sum for π^2 ?
- 7 Graph the square wave. Then graph by hand the sum of two sine terms in its series, or graph by machine the sum of 2, 3, and 10 terms. The famous **Gibbs phenomenon** is the oscillation that overshoots the jump (this doesn't die down with more terms).
- 8 Find the lengths of these vectors in Hilbert space:

$$(a) \quad \mathbf{v} = \left(\frac{1}{\sqrt{1}}, \frac{1}{\sqrt{2}}, \frac{1}{\sqrt{4}}, \frac{1}{\sqrt{8}}, \dots \right)$$

- (b) $v = (1, a, a^2, \dots)$
 (c) $f(x) = 1 + \sin x$.
- 9 Compute the Fourier coefficients a_k and b_k for $f(x)$ defined from 0 to 2π :
- (a) $f(x) = 1$ for $0 \leq x \leq \pi$, $f(x) = 0$ for $\pi < x < 2\pi$
 (b) $f(x) = x$.
- 10 When $f(x)$ has period 2π , why is its integral from $-\pi$ to π the same as from 0 to 2π ? If $f(x)$ is an *odd* function, $f(-x) = -f(x)$, show that $\int_{-\pi}^{\pi} f(x) dx$ is zero. Odd functions only have sine terms, even functions only have cosines.
- 11 Using trigonometric identities find the two terms in the Fourier series for $f(x)$:
- (a) $f(x) = \cos^2 x$ (b) $f(x) = \cos(x + \frac{\pi}{3})$ (c) $f(x) = \sin^3 x$
- 12 The functions $1, \cos x, \sin x, \cos 2x, \sin 2x, \dots$ are a basis for Hilbert space. Write the derivatives of those first five functions as combinations of the same five functions. What is the 5 by 5 “differentiation matrix” for these functions?
- 13 Find the Fourier coefficients a_k and b_k of the square pulse $F(x)$ centered at $x = 0$: $F(x) = 1/h$ for $|x| \leq h/2$ and $F(x) = 0$ for $h/2 < |x| \leq \pi$.
 As $h \rightarrow 0$, this $F(x)$ approaches a delta function. Find the limits of a_k and b_k .

Section 4.1 of *Computational Science and Engineering* explains the sine series, cosine series, complete series, and complex series $\sum c_k e^{ikx}$ on math.mit.edu/cse.

Section 9.3 of this book explains the *Discrete Fourier Transform*. This is “Fourier series for vectors” and it is computed by the **Fast Fourier Transform**. That fast algorithm comes quickly from special complex numbers $z = e^{i\theta} = \cos \theta + i \sin \theta$ when the angle is $\theta = 2\pi k/n$.

10.6 Computer Graphics

Computer graphics deals with images. The images are moved around. Their scale is changed. Three dimensions are projected onto two dimensions. All the main operations are done by matrices—but the shape of these matrices is surprising.

The transformations of three-dimensional space are done with 4 by 4 matrices. You would expect 3 by 3. The reason for the change is that one of the four key operations cannot be done with a 3 by 3 matrix multiplication. Here are the four operations:

Translation (shift the origin to another point $P_0 = (x_0, y_0, z_0)$)

Rescaling (by c in all directions or by different factors c_1, c_2, c_3)

Rotation (around an axis through the origin or an axis through P_0)

Projection (onto a plane through the origin or a plane through P_0).

Translation is the easiest—just add (x_0, y_0, z_0) to every point. But this is not linear! No 3 by 3 matrix can move the origin. So we change the coordinates of the origin to $(0, 0, 0, 1)$. This is why the matrices are 4 by 4. The “homogeneous coordinates” of the point (x, y, z) are $(x, y, z, 1)$ and we now show how they work.

1. Translation Shift the whole three-dimensional space along the vector $\mathbf{v}_0$. The origin moves to (x_0, y_0, z_0) . This vector $\mathbf{v}_0$ is added to every point $\mathbf{v}$ in $\mathbf{R}^3$. Using homogeneous coordinates, the 4 by 4 matrix T shifts the whole space by $\mathbf{v}_0$:

$$\text{Translation matrix} \quad T = \begin{bmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ x_0 & y_0 & z_0 & 1 \end{bmatrix}.$$

Important: *Computer graphics works with row vectors.* We have row times matrix instead of matrix times column. You can quickly check that $[0 \ 0 \ 0 \ 1]T = [x_0 \ y_0 \ z_0 \ 1]$.

To move the points $(0, 0, 0)$ and (x, y, z) by $\mathbf{v}_0$, change to homogeneous coordinates $(0, 0, 0, 1)$ and $(x, y, z, 1)$. Then multiply by T . A row vector times T gives a row vector. **Every $\mathbf{v}$ moves to $\mathbf{v} + \mathbf{v}_0$:** $[x \ y \ z \ 1]T = [x + x_0 \ y + y_0 \ z + z_0 \ 1]$.

The output tells where any $\mathbf{v}$ will move. (It goes to $\mathbf{v} + \mathbf{v}_0$.) Translation is now achieved by a matrix, which was impossible in $\mathbf{R}^3$.

2. Scaling To make a picture fit a page, we change its width and height. A copier will rescale a figure by 90%. In linear algebra, we multiply by .9 times the identity matrix. That matrix is normally 2 by 2 for a plane and 3 by 3 for a solid. In computer graphics, with homogeneous coordinates, the matrix is *one size larger*:

$$\text{Rescale the plane:} \quad S = \begin{bmatrix} .9 & & \\ & .9 & \\ & & 1 \end{bmatrix}$$

$$\text{Rescale a solid:} \quad S = \begin{bmatrix} c & 0 & 0 & 0 \\ 0 & c & 0 & 0 \\ 0 & 0 & c & 0 \\ 0 & 0 & 0 & 1 \end{bmatrix}.$$

Important: S is not cI . We keep the “1” in the lower corner. Then $[x, y, 1]$ times S is the correct answer in homogeneous coordinates. The origin stays in its normal position because $[0\ 0\ 1]S = [0\ 0\ 1]$.

If we change that 1 to c , the result is strange. **The point (cx, cy, cz, c) is the same as $(x, y, z, 1)$.** The special property of homogeneous coordinates is that *multiplying by cI does not move the point*. The origin in $\mathbf{R}^3$ has homogeneous coordinates $(0, 0, 0, 1)$ and $(0, 0, 0, c)$ for every nonzero c . This is the idea behind the word “homogeneous.”

Scaling can be different in different directions. To fit a full-page picture onto a half-page, scale the y direction by $\frac{1}{2}$. To create a margin, scale the x direction by $\frac{3}{4}$. The graphics matrix is diagonal but not 2 by 2. It is 3 by 3 to rescale a plane and 4 by 4 to rescale a space:

$$\text{Scaling matrices} \quad S = \begin{bmatrix} \frac{3}{4} & & \\ & \frac{1}{2} & \\ & & 1 \end{bmatrix} \quad \text{and} \quad S = \begin{bmatrix} c_1 & & \\ & c_2 & \\ & & c_3 \\ & & & 1 \end{bmatrix}.$$

That last matrix S rescales the x, y, z directions by positive numbers c_1, c_2, c_3 . The extra column in all these matrices leaves the extra 1 at the end of every vector.

Summary The scaling matrix S is the same size as the translation matrix T . They can be multiplied. To translate and then rescale, multiply vTS . To rescale and then translate, multiply vST . Are those different? *Yes*.

The point (x, y, z) in $\mathbf{R}^3$ has homogeneous coordinates $(x, y, z, 1)$ in $\mathbf{P}^3$. This “projective space” is not the same as $\mathbf{R}^4$. It is still three-dimensional. To achieve such a thing, (cx, cy, cz, c) is the same point as $(x, y, z, 1)$. Those points of projective space $\mathbf{P}^3$ are really lines through the origin in $\mathbf{R}^4$.

Computer graphics uses **affine** transformations, *linear plus shift*. An affine transformation T is executed on $\mathbf{P}^3$ by a 4 by 4 matrix with a special fourth column:

$$A = \begin{bmatrix} a_{11} & a_{12} & a_{13} & 0 \\ a_{21} & a_{22} & a_{23} & 0 \\ a_{31} & a_{32} & a_{33} & 0 \\ a_{41} & a_{42} & a_{43} & 1 \end{bmatrix} = \begin{bmatrix} T(1, 0, 0) & 0 \\ T(0, 1, 0) & 0 \\ T(0, 0, 1) & 0 \\ T(0, 0, 0) & 1 \end{bmatrix}.$$

The usual 3 by 3 matrix tells us three outputs, this tells four. The usual outputs come from the inputs $(1, 0, 0)$ and $(0, 1, 0)$ and $(0, 0, 1)$. When the transformation is linear, three outputs reveal everything. When the transformation is affine, the matrix also contains the output from $(0, 0, 0)$. Then we know the shift.

3. Rotation A rotation in $\mathbf{R}^2$ or $\mathbf{R}^3$ is achieved by an orthogonal matrix Q . The determinant is $+1$. (With determinant -1 we get an extra reflection through a mirror.) Include the extra column when you use homogeneous coordinates!

Plane rotation $Q = \begin{bmatrix} \cos \theta & -\sin \theta \\ \sin \theta & \cos \theta \end{bmatrix}$ becomes

$$R = \begin{bmatrix} \cos \theta & -\sin \theta & 0 \\ \sin \theta & \cos \theta & 0 \\ 0 & 0 & 1 \end{bmatrix}.$$

This matrix rotates the plane around the origin. *How would we rotate around a different point* $(4, 5)$? The answer brings out the beauty of homogeneous coordinates. *Translate* $(4, 5)$ *to* $(0, 0)$, *then rotate by* θ , *then translate* $(0, 0)$ *back to* $(4, 5)$:

$$vT_-RT_+ = \begin{bmatrix} x & y & 1 \end{bmatrix} \begin{bmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ -4 & -5 & 1 \end{bmatrix} \begin{bmatrix} \cos \theta & -\sin \theta & 0 \\ \sin \theta & \cos \theta & 0 \\ 0 & 0 & 1 \end{bmatrix} \begin{bmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 4 & 5 & 1 \end{bmatrix}.$$

I won't multiply. The point is to apply the matrices one at a time: v translates to vT_- , then rotates to vT_-R , and translates back to vT_-RT_+ . Because each point $\begin{bmatrix} x & y & 1 \end{bmatrix}$ is a row vector, T_- acts first. The center of rotation $(4, 5)$ —otherwise known as $(4, 5, 1)$ —moves first to $(0, 0, 1)$. Rotation doesn't change it. Then T_+ moves it back to $(4, 5, 1)$. All as it should be. The point $(4, 6, 1)$ moves to $(0, 1, 1)$, then turns by θ and moves back.

In three dimensions, every rotation Q turns around an axis. The axis doesn't move—it is a line of eigenvectors with $\lambda = 1$. Suppose the axis is in the z direction. The 1 in Q is to leave the z axis alone, the extra 1 in R is to leave the origin alone:

$$Q = \begin{bmatrix} \cos \theta & -\sin \theta & 0 \\ \sin \theta & \cos \theta & 0 \\ 0 & 0 & 1 \end{bmatrix} \quad \text{and} \quad R = \begin{bmatrix} & & 0 \\ & Q & \\ 0 & 0 & 0 & 1 \end{bmatrix}.$$

Now suppose the rotation is around the unit vector $\mathbf{a} = (a_1, a_2, a_3)$. With this axis $\mathbf{a}$, the rotation matrix Q which fits into R has three parts:

$$Q = (\cos \theta)I + (1 - \cos \theta) \begin{bmatrix} a_1^2 & a_1a_2 & a_1a_3 \\ a_1a_2 & a_2^2 & a_2a_3 \\ a_1a_3 & a_2a_3 & a_3^2 \end{bmatrix} - \sin \theta \begin{bmatrix} 0 & a_3 & -a_2 \\ -a_3 & 0 & a_1 \\ a_2 & -a_1 & 0 \end{bmatrix}. \quad (1)$$

The axis doesn't move because $\mathbf{a}Q = \mathbf{a}$. When $\mathbf{a} = (0, 0, 1)$ is in the z direction, this Q becomes the previous Q —for rotation around the z axis.

The linear transformation Q always goes in the upper left block of R . Below it we see zeros, because rotation leaves the origin in place. When those are not zeros, the transformation is affine and the origin moves.

4. Projection In a linear algebra course, most planes go through the origin. In real life, most don't. A plane through the origin is a vector space. The other planes are affine spaces, sometimes called “flats.” An affine space is what comes from translating a vector space.

We want to project three-dimensional vectors onto planes. Start with a plane through the origin, whose unit normal vector is $\mathbf{n}$. (We will keep $\mathbf{n}$ as a column vector.) The vectors in the plane satisfy $\mathbf{n}^T \mathbf{v} = 0$. *The usual projection onto the plane is the matrix* $I - \mathbf{n}\mathbf{n}^T$. To project a vector, multiply by this matrix. The vector $\mathbf{n}$ is projected to zero, and the in-plane vectors $\mathbf{v}$ are projected onto themselves:

$$(I - \mathbf{n}\mathbf{n}^T)\mathbf{n} = \mathbf{n} - \mathbf{n}(\mathbf{n}^T \mathbf{n}) = \mathbf{0} \quad \text{and} \quad (I - \mathbf{n}\mathbf{n}^T)\mathbf{v} = \mathbf{v} - \mathbf{n}(\mathbf{n}^T \mathbf{v}) = \mathbf{v}.$$

In homogeneous coordinates the projection matrix becomes 4 by 4 (but the origin doesn't move):

$$\text{Projection onto the plane } \mathbf{n}^T \mathbf{v} = 0 \quad P = \begin{bmatrix} I - \mathbf{n}\mathbf{n}^T & 0 \\ 0 & 0 & 0 & 1 \end{bmatrix}.$$

Now project onto a plane $\mathbf{n}^T(\mathbf{v} - \mathbf{v}_0) = 0$ that does *not* go through the origin. One point on the plane is $\mathbf{v}_0$. This is an affine space (or a *flat*). It is like the solutions to $A\mathbf{v} = \mathbf{b}$ when the right side is not zero. One particular solution $\mathbf{v}_0$ is added to the nullspace—to produce a flat.

The projection onto the flat has three steps. Translate $\mathbf{v}_0$ to the origin by T_- . Project along the $\mathbf{n}$ direction, and translate back along the row vector $\mathbf{v}_0$:

$$\text{Projection onto a flat} \quad T_- P T_+ = \begin{bmatrix} I & 0 \\ -\mathbf{v}_0 & 1 \end{bmatrix} \begin{bmatrix} I - \mathbf{n}\mathbf{n}^T & 0 \\ 0 & 1 \end{bmatrix} \begin{bmatrix} I & 0 \\ \mathbf{v}_0 & 1 \end{bmatrix}.$$

I can't help noticing that T_- and T_+ are inverse matrices: translate and translate back. They are like the elementary matrices of Chapter 2.

The exercises will include reflection matrices, also known as *mirror matrices*. These are the fifth type needed in computer graphics. A reflection moves each point twice as far as a projection—the reflection goes through the plane and out the other side. So change the projection $I - \mathbf{n}\mathbf{n}^T$ to $I - 2\mathbf{n}\mathbf{n}^T$ for a mirror matrix.

The matrix P gave a “parallel” projection. All points move parallel to $\mathbf{n}$, until they reach the plane. The other choice in computer graphics is a “perspective” projection. This is more popular because it includes foreshortening. With perspective, an object looks larger as it moves closer. Instead of staying parallel to $\mathbf{n}$ (and parallel to each other), the lines of projection come *toward the eye*—the center of projection. This is how we perceive depth in a two-dimensional photograph.

The basic problem of computer graphics starts with a scene and a viewing position. Ideally, the image on the screen is what the viewer would see. The simplest image assigns just one bit to every small picture element—called a *pixel*. It is light or dark. This gives a black and white picture with no shading. You would not approve. In practice, we assign shading levels between 0 and 2^8 for three colors like red, green, and blue. That means $8 \times 3 = 24$ bits for each pixel. Multiply by the number of pixels, and a lot of memory is needed!

Physically, a *raster frame buffer* directs the electron beam. It scans like a television set. The quality is controlled by the number of pixels and the number of bits per pixel. In this area, the standard text is *Computer Graphics: Principles and Practice* by Hughes, Van Dam, McGuire, Skylar, Foley, Feiner, and Akeley (3rd edition, Addison-Wesley, 2014). Notes by Ronald Goldman and by Tony DeRose were excellent references.

■ REVIEW OF THE KEY IDEAS ■

1. Computer graphics needs shift operations $T(\mathbf{v}) = \mathbf{v} + \mathbf{v}_0$ as well as linear operations $T(\mathbf{v}) = A\mathbf{v}$.
2. A shift in $\mathbf{R}^n$ can be executed by a matrix of order $n + 1$, using homogeneous coordinates.
3. The extra component 1 in $[x \ y \ z \ 1]$ is preserved when all matrices have the numbers 0, 0, 0, 1 as last column.

Problem Set 10.6

- 1 A typical point in $\mathbf{R}^3$ is $x\mathbf{i} + y\mathbf{j} + z\mathbf{k}$. The coordinate vectors $\mathbf{i}$, $\mathbf{j}$, and $\mathbf{k}$ are $(1, 0, 0)$, $(0, 1, 0)$, $(0, 0, 1)$. The coordinates of the point are (x, y, z) .
This point in computer graphics is $x\mathbf{i} + y\mathbf{j} + z\mathbf{k} + \mathbf{origin}$. Its homogeneous coordinates are $(\quad, \quad, \quad, \quad)$. Other coordinates for the same point are $(\quad, \quad, \quad, \quad)$.
- 2 A linear transformation T is determined when we know $T(\mathbf{i})$, $T(\mathbf{j})$, $T(\mathbf{k})$. For an affine transformation we also need $T(\quad)$. The input point $(x, y, z, 1)$ is transformed to $xT(\mathbf{i}) + yT(\mathbf{j}) + zT(\mathbf{k}) + \quad$.
- 3 Multiply the 4 by 4 matrix T for translation along $(1, 4, 3)$ and the matrix T_1 for translation along $(0, 2, 5)$. The product TT_1 is translation along $\quad$.
- 4 Write down the 4 by 4 matrix S that scales by a constant c . Multiply ST and also TS , where T is translation by $(1, 4, 3)$. To blow up the picture around the center point $(1, 4, 3)$, would you use $\mathbf{v}ST$ or $\mathbf{v}TS$?
- 5 What scaling matrix S (in homogeneous coordinates, so 3 by 3) would produce a 1 by 1 square page from a standard 8.5 by 11 page?
- 6 What 4 by 4 matrix would move a corner of a cube to the origin and then multiply all lengths by 2? The corner of the cube is originally at $(1, 1, 2)$.
- 7 When the three matrices in equation 1 multiply the unit vector $\mathbf{a}$, show that they give $(\cos \theta)\mathbf{a}$ and $(1 - \cos \theta)\mathbf{a}$ and $\mathbf{0}$. Addition gives $\mathbf{a}Q = \mathbf{a}$ and the rotation axis is not moved.
- 8 If $\mathbf{b}$ is perpendicular to $\mathbf{a}$, multiply by the three matrices in 1 to get $(\cos \theta)\mathbf{b}$ and $\mathbf{0}$ and a vector perpendicular to $\mathbf{b}$. So $Q\mathbf{b}$ makes an angle θ with $\mathbf{b}$. **This is rotation.**
- 9 What is the 3 by 3 projection matrix $I - \mathbf{nn}^T$ onto the plane $\frac{2}{3}x + \frac{2}{3}y + \frac{1}{3}z = 0$? In homogeneous coordinates add 0, 0, 0, 1 as an extra row and column in P .

- 10 With the same 4 by 4 matrix P , multiply T_-PT_+ to find the projection matrix onto the plane $\frac{2}{3}x + \frac{2}{3}y + \frac{1}{3}z = 1$. The translation T_- moves a point on that plane (choose one) to $(0, 0, 0, 1)$. The inverse matrix T_+ moves it back.
- 11 Project $(3, 3, 3)$ onto those planes. Use P in Problem 9 and T_-PT_+ in Problem 10.
- 12 If you project a square onto a plane, what shape do you get?
- 13 If you project a cube onto a plane, what is the outline of the projection? Make the projection plane perpendicular to a diagonal of the cube.
- 14 The 3 by 3 mirror matrix that reflects through the plane $\mathbf{n}^T\mathbf{v} = 0$ is $M = I - 2\mathbf{n}\mathbf{n}^T$. Find the reflection of the point $(3, 3, 3)$ in the plane $\frac{2}{3}x + \frac{2}{3}y + \frac{1}{3}z = 0$.
- 15 Find the reflection of $(3, 3, 3)$ in the plane $\frac{2}{3}x + \frac{2}{3}y + \frac{1}{3}z = 1$. Take three steps T_-MT_+ using 4 by 4 matrices: translate by T_- so the plane goes through the origin, reflect the translated point $(3, 3, 3, 1)T_-$ in that plane, then translate back by T_+ .
- 16 The vector between the origin $(0, 0, 0, 1)$ and the point $(x, y, z, 1)$ is the difference $\mathbf{v} = \underline{\hspace{1cm}}$. In homogeneous coordinates, vectors end in $\underline{\hspace{1cm}}$. So we add a $\underline{\hspace{1cm}}$ to a point, not a point to a point.
- 17 If you multiply only the *last* coordinate of each point to get (x, y, z, c) , you rescale the whole space by the number $\underline{\hspace{1cm}}$. This is because the point (x, y, z, c) is the same as $(\underline{\hspace{1cm}}, \underline{\hspace{1cm}}, \underline{\hspace{1cm}}, 1)$.

10.7 Linear Algebra for Cryptography

- 1 Codes can use finite fields as alphabets: letters in the message become numbers $0, 1, \dots, p-1$.
- 2 The numbers are added and multiplied ($\text{mod } p$). Divide by p , keep the remainder.
- 3 A Hill Cipher multiplies blocks of the message by a secret matrix $E \pmod{p}$.
- 4 To decode, multiply each block by the inverse matrix $D \pmod{p}$. Not a very secure cipher!

Cryptography is about encoding and decoding messages. Banks do this all the time with financial information. Amazingly, modern algorithms can involve extremely deep mathematics. “Elliptic curves” play a part in cryptography, as they did in the sensational proof by Andrew Wiles of Fermat’s Last Theorem.

This section will not go that far! But it will be our first experience with *finite fields* and *finite vector spaces*. The field for $\mathbf{R}^n$ contains all real numbers. The field for “modular arithmetic” contains only p integers $0, 1, \dots, p-1$. There were infinitely many vectors in $\mathbf{R}^n$ —now there will only be p^n messages of length n in message space. The alphabet from A to Z is finite (as in $p = 26$).

The codes in this section will be easily breakable—they are much too simple for practical security. The power of computers demands more complex cryptography, because that power would quickly detect a small encoding matrix. But a matrix code (the Hill Cipher) will allow us to see linear algebra at work in a new way.

All our calculations in encoding and decoding will be “**mod p** ”. But the central concepts of linear independence and bases and inverse matrices and determinants survive this change. We will be doing “linear algebra with finite fields”. Here is the meaning of $\text{mod } p$:

$$27 \equiv 2 \pmod{5} \quad \text{means that } 27 - 2 \text{ is divisible by } 5$$

$$y \equiv x \pmod{p} \quad \text{means that } y - x \text{ is divisible by } p$$

Dividing y by 5 produces one of the five possible remainders $x = 0, 1, 2, 3, 4$. All the numbers $5, -5, 10, -10, \dots$ with no remainder are congruent to zero ($\text{mod } 5$). The numbers $y = 6, -4, 11, -9, \dots$ are all congruent to $x = 1 \pmod{5}$.

We use the word **congruent** for the symbol $\equiv$ and we call this “modular arithmetic”. Every integer y produces one of the values $x = 0, 1, 2, \dots, p-1$.

The theory is best if p is a prime number. With $p = 26$ letters from A to Z, we unfortunately don’t start with a prime p . Cryptography can deal with this problem.

Modular Arithmetic

Linear algebra is based on linear combinations of vectors. Now our vectors $(x_1, \dots, x_n)$ are strings of integers limited to $x = 0, 1, \dots, p-1$. All calculations produce these integers when we work “mod p ”. This means: *Every integer y outside that range is divided by p and x is the remainder:*

$$y = qp + x \quad y \equiv x \pmod{p} \quad y \text{ divided by } p \text{ has remainder } x$$

Addition mod 3 $10 \equiv 1 \pmod{3}$ and $16 \equiv 1 \pmod{3}$ and $10 + 16 \equiv 1 + 1 \pmod{3}$

I could add $10 + 16$ and divide 26 by 3 to get the remainder 2.

Or I can just add remainders $1 + 1$ to reach the same answer 2.

Addition mod 2 $11 \equiv 1 \pmod{2}$ and $17 \equiv 1 \pmod{2}$ and $11 + 17 = 28 \equiv 0 \pmod{2}$

The remainders added to $1 + 1$ *but this is not 2*. The final step was $2 \equiv 0 \pmod{2}$.

Addition mod p is completely reasonable. So is **multiplication mod p** . Here $p = 3$:

$$10 \equiv 1 \pmod{3} \text{ times } 16 \equiv 1 \pmod{3} \text{ gives } 1 \text{ times } 1 \equiv 1 \quad 160 \equiv 1 \pmod{3}$$

$$5 \equiv 2 \pmod{3} \text{ times } 8 \equiv 2 \pmod{3} \text{ gives } 2 \text{ times } 2 \equiv 1 \quad 40 \equiv 1 \pmod{3}$$

Conclusion: We can safely add and multiply modulo p . So we can take linear combinations. This is the key operation in linear algebra. **But can we divide?**

In the real number field, the inverse is $1/y$ (for any number except $y = 0$). This means: We found another real number z so that $yz = 1$. Invertibility is a requirement for a field. **Is inversion always possible mod p ?** For every number $y = 1, \dots, p-1$ can we find another number $z = 1, \dots, p-1$ so that $yz \equiv 1 \pmod{p}$?

The examples $3^{-1} \equiv 4 \pmod{11}$ and $2^{-1} \equiv 6 \pmod{11}$ and $5^{-1} \equiv 9 \pmod{11}$ all succeed. Can you solve $7z \equiv 1 \pmod{11}$? Inverting numbers will be the key to inverting matrices.

Let me show that inversion mod p has a problem when p is not a prime number. The example $p = 26$ factors into 2 times 13. **Then $y = 2$ cannot have an inverse $z \pmod{26}$.** The requirement $2z \equiv 1 \pmod{26}$ is impossible to satisfy because $2z$ and 26 are even.

Similarly 5 has no inverse z when p is 25. We can't solve $5z \equiv 1 \pmod{25}$. The number $5z - 1$ is never going to be a multiple of 5, so it can't be a multiple of 25.

Inversion of every y ($0 < y < p$) will be possible if and only if p is prime.

Inversion needs $y, 2y, 3y, \dots, py$ to have different remainders when divided by p .

If my and ny had the same remainder x then $(m - n)y$ would be divisible by p .

The prime number p would have to divide either $m - n$ or y . Both are impossible.

So $y, \dots, py$ have different remainders: **One of those remainders must be $x = 1$.**

The Enigma Machine and the Hill Cipher

Lester Hill published his cipher (his system for encoding and decoding) in the *American Mathematical Monthly* (1929). The idea was simple, but in some way it started the transition of cryptography from linguistics to mathematics. Codes up to that time mainly mixed up alphabets and rearranged messages. The **Enigma code** used by the German Navy in World War II was a giant advance—using machines that look to us like primitive computers. The English set up Bletchley Park to break Enigma. They hired puzzle solvers and language majors. And by good luck they also happened to get Alan Turing.

I don't know if you have seen the movie about him: *The Imitation Game*. A lot of it is unrealistic (like *Good Will Hunting* and *A Beautiful Mind* at MIT). But the core idea of breaking the Enigma code was correct, using human weaknesses in the encoding and broadcasting. The German naval command openly sent out their coded orders—knowing that the codes were too complicated to break (if it hadn't been for those weaknesses). The codebreaking required English electronics to undo the German electronics. It also required genius.

Alan Turing was surely a genius—England's most exceptional mathematician. His life was ultimately tragic and he ended it in 1954. The biography by Andrew Hodges is excellent. Turing arrived at Bletchley Park the day after Poland was invaded. It is to Winston Churchill's credit that he gave fast and full support when his support was needed.

The Enigma Machine had gears and wheels. The Hill Cipher only needs a matrix. That is the code to be explained now, using linear algebra. You will see how decoding involved inverse matrices. All steps use modular arithmetic, multiplying and inverting $\text{mod } p$.

I will follow the neat exposition of Professor Spickler of Salisbury State University, which he made available on the Web: facultyfp.salisbury.edu/despickler/personal/index.asp

Modular Arithmetic with Matrices

Addition, subtraction, and multiplication are all we need for Ax (matrix times vector). To multiply $\text{mod } p$ we can multiply the integers in A times the integers in x as usual—and then replace every entry of Ax by its value $\text{mod } p$.

Key questions: When can we solve $Ax \equiv b \pmod{p}$? Do we still have the four subspaces $C(A)$, $N(A)$, $C(A^T)$, $N(A^T)$? Are they still orthogonal in pairs? Is there still an inverse matrix $\text{mod } p$ whenever the determinant of A is nonzero $\text{mod } p$? I am happy to say that the last three answers are *yes* (but the inverse question requires p to be a prime number).

We can find $A^{-1} \pmod{p}$ by Gauss-Jordan elimination, reducing $[A \ I]$ to $[I \ A^{-1}]$ as in Section 2.5. Or we can use determinants and the cofactor matrix C in the formula $A^{-1} = (\det A)^{-1} C^T$. I will work $\text{mod } 3$ with a 2 by 2 integer matrix A :

$$[A \ I] = \begin{bmatrix} 2 & 0 & 1 & 0 \\ 2 & 1 & 0 & 1 \end{bmatrix} \rightarrow \begin{bmatrix} 2 & 0 & 1 & 0 \\ 0 & 1 & 2 & 1 \end{bmatrix} \xrightarrow[\text{by } 2^{-1} \equiv 2]{\text{multiply row 1}} \begin{bmatrix} 1 & 0 & 2 & 0 \\ 0 & 1 & 2 & 1 \end{bmatrix} = [I \ A^{-1}]$$

By pure chance $A^{-1} \equiv A$! Multiplying A times $A \pmod 3$ does give the identity matrix:

$$A^2 = AA^{-1} = \begin{bmatrix} 2 & 0 \\ 2 & 1 \end{bmatrix} \begin{bmatrix} 2 & 0 \\ 2 & 1 \end{bmatrix} = \begin{bmatrix} 4 & 0 \\ 6 & 1 \end{bmatrix} \equiv \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix} \pmod 3.$$

The determinant of A is 2, and the cofactor formula from Section 5.3 also gives $A^{-1} \equiv A$:

$$\begin{bmatrix} 2 & 0 \\ 2 & 1 \end{bmatrix}^{-1} = 2^{-1} \begin{bmatrix} 1 & -0 \\ -2 & 2 \end{bmatrix} \equiv 2 \begin{bmatrix} 1 & -0 \\ -2 & 2 \end{bmatrix} \equiv \begin{bmatrix} 2 & 0 \\ 2 & 1 \end{bmatrix} \pmod 3.$$

Theorem. A^{-1} exists $\pmod p$ if and only if $(\det A)^{-1}$ exists $\pmod p$.

The requirement is: $\det A$ and p have no common factors.

Encryption with the Hill Cipher

The original cipher used the letters A to Z with $p = 26$. Hill chose an n by n encryption matrix E so that $\det E$ is not divisible by 2 or 13. Then the number $\det E$ has an inverse $\pmod 26$ and so does the matrix E . The inverse matrix $E^{-1} \equiv D \pmod 26$ will be the decryption matrix that decodes the message.

Now convert each letter of the message into a number from 0 to 25. The obvious choice from $A = 0$ to $Z = 25$ is acceptable because the matrix will make this cipher stronger.

Ignore spaces and divide the message into blocks $v_1, v_2, \dots$ of size n .

Then multiply each message block $\pmod p$ by the encryption matrix E .

The coded message is $Ev_1, Ev_2, \dots$ and you know what the decoder will do.

$$\begin{aligned} \text{Spikler's example has } D = E^{-1} &= \begin{bmatrix} 2 & 3 & 15 \\ 5 & 8 & 12 \\ 1 & 13 & 4 \end{bmatrix}^{-1} \equiv \begin{bmatrix} 10 & 19 & 16 \\ 4 & 23 & 7 \\ 17 & 5 & 19 \end{bmatrix} \pmod{26}. \\ \det E = 583 \equiv 11 \pmod{26} \end{aligned}$$

Of course a codebreaker will not know E or D . And the block size n is generally unknown too. For the matrices Hill had in mind n would not be very large and a computer could quickly discover E and D .

I am not sure if Hill's Cipher could become seriously difficult to break by choosing very large matrices and a large prime number p . And by encoding the coded message a second time, using a different block size n_2 and large matrix E_2 and large prime p_2 .

Finite Fields and Finite Vector Spaces

In algebra, a field $\mathbf{F}$ is a set of scalars that can be added and multiplied and inverted (except 0 can't be inverted). Familiar examples are the real numbers $\mathbf{R}$ and the complex numbers $\mathbf{C}$ and the rational numbers $\mathbf{Q}$ (containing every ratio p/q of integers). From a field you build vectors $v = (f_1, f_2, \dots, f_n)$. From linear combinations of vectors you build vector spaces. *So linear algebra begins with a field $\mathbf{F}$.*

I taught for ten years from a textbook that started with fields. On the way to $\mathbf{R}^n$, we lost a lot of students. That was a signal—the emphasis was misplaced if we wanted the

course to be useful. I believe the right way is to understand $\mathbf{R}^n$ and its subspaces first, as you do. Then you can look at other fields and vector spaces with a natural question in mind: *What is new when the field is not $\mathbf{R}$?*

These pages are asking that question for **finite fields**. The possibilities become more limited but also highly interesting. The starting point (and not quite the ending point) is the finite field $\mathbf{F}_p$. It contains only the numbers $0, 1, \dots, p-1$ and p is a prime number. I will focus first on the field $\mathbf{F}_2$ with only 2 members “0” and “1”. You could think of 0 and 1 as “even” and “odd” because the rules to add and multiply are obeyed by the even numbers and odd numbers: even $+$ odd = *odd* and even $\times$ odd = *even*.

		0	1			0	1								
Addition	0	<table><tr><td>0</td><td>1</td></tr><tr><td>1</td><td>0</td></tr></table>		0	1	1	0	Multiplication	0	<table><tr><td>0</td><td>0</td></tr><tr><td>0</td><td>1</td></tr></table>		0	0	0	1
0	1														
1	0														
0	0														
0	1														
table	1			table	1										

This is addition and multiplication “*mod 2*”.

From this field $\mathbf{F}_2$ we can build vectors like $\mathbf{v} = (0, 0, 1)$ and $\mathbf{w} = (1, 0, 1)$. There are three components with two choices each: a total of $2^3 = 8$ different vectors in the vector space $(\mathbf{F}_2)^3$. You know the requirements on a subspace and the possibilities it opens up:

- The zero-dimensional subspace containing only $\mathbf{0} = (0, 0, 0)$.
- One-dimensional subspaces containing $\mathbf{0}$ and a vector like $\mathbf{v}$. Notice $\mathbf{v} + \mathbf{v} = \mathbf{0}$!
- Two-dimensional subspaces with a basis like $\mathbf{v}$ and $\mathbf{w}$ and 4 vectors $\mathbf{0}, \mathbf{v}, \mathbf{w}, \mathbf{v} + \mathbf{w}$.
- The full three-dimensional subspace $(\mathbf{F}_2)^3$ with 8 vectors.

What are the possible bases for $(\mathbf{F}_2)^3$? The standard basis contains $(1, 0, 0)$ and $(0, 1, 0)$ and $(0, 0, 1)$. Those vectors are linearly independent and they span $(\mathbf{F}_2)^3$. Their eight combinations with coefficients 0 and 1 fill all of $(\mathbf{F}_2)^3$.

What about matrices that multiply those vectors? The matrices will be 1 by 3, or 2 by 3, or 3 by 3. When they are 3 by 3 we can ask if they are invertible. Their determinants can only be 0 (singular matrix) or 1 (invertible matrix). Let me leave you the pleasure of deciding whether these matrices are invertible. *And how would you find the inverse?*

$$A = \begin{bmatrix} 1 & 0 & 0 \\ 1 & 1 & 0 \\ 1 & 1 & 1 \end{bmatrix} \quad B = \begin{bmatrix} 1 & 1 & 0 \\ 0 & 1 & 1 \\ 1 & 0 & 1 \end{bmatrix} \quad C = \begin{bmatrix} 1 & 1 & 1 \\ 0 & 0 & 1 \\ 1 & 0 & 0 \end{bmatrix}$$

Out of 2^9 possible matrices over $\mathbf{F}_2$, I will guess that most are singular.

To conclude this discussion of $\mathbf{F}_2$, I mention a field with $2^2 = 4$ members. It will not come from multiplication (*mod 4*), because 4 is not prime. The multiplication 2 times 2 will give 0 (and 2 has no inverse): *not a field*. But we can start with the numbers 0 and 1 in $\mathbf{F}_2$ and invent two more numbers a and $1 + a$ —provided they follow these two rules: $(a + a = 0)$ and $(a \times a = 1 + a)$. Then a and $1 + a$ are inverses. Not obvious!

Add	0	1	a	$1 + a$
0	0	1	a	$1 + a$
1	1	0	$1 + a$	a
a	a	$1 + a$	0	1
$1 + a$	$1 + a$	a	1	0

Multiply	0	1	a	$1 + a$
0	0	0	0	0
1	0	1	a	$1 + a$
a	0	a	$1 + a$	1
$1 + a$	0	$1 + a$	1	a

Beyond $p = 2$, we have the fields $\mathbf{F}_p$ for all prime numbers p . They use addition and multiplication $\text{mod } p$. They are alphabets for codes. They provide the components for vectors $\mathbf{v} = (f_1, \dots, f_n)$ in the space $(\mathbf{F}_p)^n$. They provide the entries for matrices that multiply those vectors. These fields $\mathbf{F}_p$ are the most frequently used finite fields.

The only other finite fields have p^k members. The example above of $0, 1, a, 1 + a$ had $2^2 = 4$ members. We will leave it there and get back safely to $\mathbf{R}$.

Problem Set 10.7

- 1 If you multiply n whole numbers (even or odd) when is the answer odd? Translate into multiplication ($\text{mod } 2$): If you multiply 0's and 1's when is the answer 1?
- 2 If you add n whole numbers (even or odd) when is the sum of the numbers odd? Translate into adding 0's and 1's ($\text{mod } 2$). When do they add to 1?
- 3 (a) If $y_1 \equiv x_1$ and $y_2 \equiv x_2$, why is $y_1 + y_2 \equiv x_1 + x_2$? All are $\text{mod } p$.
Suggestion: $y_1 = p q_1 + x_1$ and $y_2 = p q_2 + x_2$. Now add $y_1 + y_2$.
(b) Can you be sure that $x_1 + x_2$ is smaller than p ? No. Give an example where there is a smaller x with $(y_1 + y_2) = x \pmod{p}$.
- 4 $p = 39$ is not prime. Find a number a that has no inverse $z \pmod{39}$. This means that $az \equiv 1 \pmod{39}$ has no solution. Then find a 2 by 2 matrix A that has no inverse matrix $Z \pmod{39}$. This means that $AZ \equiv I \pmod{39}$ has no solution.
- 5 Show that $y \equiv x \pmod{p}$ leads to $-y \equiv -x \pmod{p}$.
- 6 Find a matrix that has independent columns in $\mathbf{R}^2$ but dependent columns ($\text{mod } 5$).
- 7 What are all the 2 by 2 matrices of 0's and 1's that are invertible ($\text{mod } 2$)?
- 8 Is the row space of A still orthogonal to the nullspace in modular arithmetic ($\text{mod } 11$)? Are bases for those subspaces still bases ($\text{mod } 11$)?
- 9 (Hill's Cipher) Separate the message THISWHOLEBOOKISINCODE into blocks of 3 letters. Replace each letter by a number from 1 to 26 (normal order). Multiply each block by the 3 by 3 matrix L with 1's on and below the diagonal. What is the coded message (in numbers) and how would you decode it?
- 10 Suppose you know the original message (the plaintext). Suppose you also see the coded message. How would you start to discover the matrix in Hill's Cipher? For a very long message do you expect success?