

В. Босс

ЛЕКЦИИ *по* МАТЕМАТИКЕ

ТОМ

3

Линейная алгебра

МОСКВА



URSS

Босс В.

Лекции по математике: линейная алгебра. Т. 3. — М.: КомКнига, 2005.
224 с.

ISBN 5–484–00046–7

Книга отличается краткостью и прозрачностью изложения. Объяснения даются «человеческим языком» — лаконично и доходчиво. Значительное внимание уделяется мотивации результатов и прикладным аспектам. Даже в устоявшихся темах ощущается свежий взгляд, в связи с чем преподаватели найдут для себя немало интересного. Книга легко читается.

Аналитическая геометрия рассматривается как вспомогательный предмет, способствующий освоению понятий векторного пространства. Охват линейной алгебры достаточно широкий, но изложение построено так, что можно ограничиться любым желаемым срезом содержания.

Для студентов, преподавателей, инженеров и научных работников.

Издательство «КомКнига». 117312, г. Москва, пр-т 60-летия Октября, 9
Подписано к печати 24.02.2005 г. Формат 60х90/16. Печ. л. 14. Заказ № 1912

Отпечатано с готовых диапозитивов во ФГУП ИПК «Ульяновский Дом печати»
432980, г. Ульяновск, ул. Гончарова, 14

ISBN 5–484–00046–7

© КомКнига, 2005

НАУЧНАЯ И УЧЕБНАЯ ЛИТЕРАТУРА	
	E-mail: URSS@URSS.ru
	Каталог изданий в Интернете:
	http://URSS.ru
	Тел./факс: 7 (095) 135–42–16
URSS	Тел./факс: 7 (095) 135–42–46

3143 ID 26865



Оглавление

Предисловие к «Лекциям»	7
Предисловие к тому	9
Глава 1. Аналитическая геометрия	10
1.1. Координаты и векторы	10
1.2. Описание геометрических объектов	15
1.3. Векторное произведение	19
1.4. Определители	22
1.5. Матрицы и преобразования	23
1.6. Прямые и плоскости	29
1.7. Геометрические задачи	32
1.8. Кривые и поверхности второго порядка	35
Глава 2. Векторы и матрицы	38
2.1. Примеры линейных задач	38
2.2. Векторы	39
2.3. Распознавание образов	43
2.4. Линейные отображения и матрицы	45
2.5. Прямоугольные и клеточные матрицы	49
2.6. Два примера	51
2.7. Элементарные преобразования	52
2.8. Теория определителей	57
2.9. Системы уравнений	62
2.10. Задачи и дополнения	65
Глава 3. Линейные преобразования	66
3.1. Замена координат	66
3.2. Собственные значения и комплексные пространства	68
3.3. Собственные векторы	72

3.4. Эскиз спектральной теории	74
3.5. Линейные пространства	76
3.6. Манипуляции с подпространствами	78
3.7. Задачи и дополнения	80
Глава 4. Квадратичные формы	81
4.1. Квадратичные формы	81
4.2. Положительная определенность	86
4.3. Инерция и сигнатура	89
4.4. Условный экстремум	90
4.5. Сингулярные числа	91
4.6. Биортогональные базисы	92
4.7. Сопряженное пространство	94
4.8. Преобразования и тензоры	98
4.9. Задачи и дополнения	100
Глава 5. Канонические представления	103
5.1. Унитарные матрицы	103
5.2. Триангуляция Шура	105
5.3. Жордановы формы	108
5.4. Аннулирующий многочлен	112
5.5. Корневые подпространства	113
5.6. Теорема Гамильтона—Кэли	117
5.7. λ -матрицы	118
5.8. Задачи и дополнения	120
Глава 6. Функции от матриц	123
6.1. Матричные ряды	123
6.2. Нормы векторов и матриц	125
6.3. Спектральный радиус	130
6.4. Сходимость итераций	131
6.5. Функции как ряды	132
6.6. Матричная экспонента	133
6.7. Конечные алгоритмы	135
6.8. Задачи и дополнения	138

Глава 7. Матричные уравнения	140
7.1. Типичные задачи	140
7.2. Кронекерово произведение	141
7.3. Уравнения	143
Глава 8. Неравенства	147
8.1. Теоремы об альтернативах	147
8.2. Выпуклые множества и конусы	149
8.3. Теоремы о пересечениях	152
8.4. P -матрицы	153
8.5. Линейное программирование	156
8.6. Задачи и дополнения	161
Глава 9. Положительные матрицы	162
9.1. Полуупорядоченность и монотонность	162
9.2. Теорема Перрона	163
9.3. Неразложимость	168
9.4. Положительная обратимость	170
9.5. Оператор сдвига и устойчивость	172
9.6. Импримитивность	176
9.7. Стохастические матрицы	177
9.8. Конус положительно определенных матриц	179
9.9. Задачи и дополнения	180
Глава 10. Численные методы	182
10.1. Предмет изучения	182
10.2. Ошибки счета и обусловленность	184
10.3. Оценки сверху и по вероятности	187
10.4. Возмущения спектра	188
10.5. Итерационные методы	191
10.6. Вычисление собственных значений	194
Глава 11. Сводка основных определений и результатов	196
11.1. Аналитическая геометрия	196
11.2. Векторы и матрицы	200
11.3. Линейные преобразования	205

11.4. Квадратичные формы	208
11.5. Канонические представления	210
11.6. Функции от матриц	211
11.7. Неравенства	213
11.8. Положительные матрицы	214
Обозначения	216
Литература	218
Предметный указатель	219

Предисловие к «Лекциям»

Самолеты позволяют летать, но добираться до аэропорта приходится самому.

Для нормального изучения любого математического предмета необходимы по крайней мере четыре ингредиента:

- 1) *живой учитель;*
- 2) *обыкновенный подробный учебник;*
- 3) *рядовой задачник;*
- 4) *учебник, освобожденный от рутины, но дающий общую картину, мотивы, связи, «что зачем».*

До четвертого пункта у системы образования руки не доходили. Конечно, подобная задача иногда ставилась и решалась, но в большинстве случаев — при параллельном исполнении функций обыкновенного учебника. Акценты из-за перегрузки менялись, и намерения со второй-третьей главы начинали дрейфовать, не достигая результата. В виртуальном пространстве так бывает. Аналог объединения гантели с теннисной ракеткой перестает решать обе задачи, хотя это не сразу бросается в глаза.

«Лекции» ставят 4-й пункт своей главной целью. Сопутствующая идея — экономия слов и средств. Правда, на фоне деклараций о краткости и ясности изложения предполагаемое издание около 20 томов может показаться тяжеловесным, но это связано с обширностью математики, а не с перегрузкой деталями.

Необходимо сказать, на кого рассчитано. Ответ «на всех» выглядит наивно, но он в какой-то мере отражает суть дела. Обозримый вид, обнаженные конструкции доказательств — такого сорта

книги удобно иметь под рукой. Не секрет, что специалисты самой высокой категории тратят массу сил и времени на освоение математических секторов, лежащих за рамками собственной специализации. Здесь же ко многим проблемам предлагается короткая дорога, позволяющая быстро освоить новые области и освежить старые. Для начинающих «короткие дороги» тем более полезны, поскольку облегчают движение любыми другими путями.

В вопросе «на кого рассчитано» — есть и другой аспект. На сильных или слабых? На средний вуз или физтех? Опять-таки выходит «на всех». Звучит странно, но речь не идет о регламентации кругозора. Простым языком, коротко и прозрачно описывается предмет. Из этого каждый извлечет свое и двинется дальше.

Наконец, последнее. В условиях информационного наводнения инструменты вчерашнего дня перестают работать. Не потому, что изучаемые дисциплины чересчур разрослись, а потому, что новых секторов жизни стало слишком много. И в этих условиях мало кто готов уделять много времени чему-то одному. Поэтому учить всему — надо как-то иначе. «Лекции» дают пример. Плохой ли, хороший — покажет время. Но в любом случае, это продукт нового поколения. Те же «колеса», тот же «руль», та же математическая суть, — но по-другому.

Предисловие к тому

Дом строят месяц, сдают под ключ — год.

Если что и дает ясное представление о высшей математике, так это линейная алгебра. Барьер повседневности здесь преодолевается легко и просто. При этом оказывается, что удивительные вещи находятся не в туманной дали, а совсем рядом.

Для освоения, разумеется, нужна еще определенная составляющая у вектора жизненных интересов. Некоторая готовность к трудностям и, конечно, время. Вы представьте, что переехали жить в другой город. День-другой побродили по главным улицам. Взглянули на панораму с какой-нибудь вышки. Разве этого достаточно — для знакомства? Чтобы узнать город, надо исходить его пешком. Много раз, вдоль и поперек, и в дождь, и в снег. И в хорошем, и в плохом настроении. Сжиться с людьми, побегать по магазинам, поездить на трамвае — и по делу, и просто так. И тогда лишь, в тысячный раз выходя из дома, вы увидите вдруг — *знакомый* город.

Глава 1

Аналитическая геометрия

Аналитическая геометрия представляет собой удобную ступеньку к пониманию линейной алгебры. Рене Декарт создавал ее, правда, с другой целью — чтобы геометрические задачи можно было решать алгебраически. Но результат часто не совпадает с намерением. Декартовы координаты оказались полезнее в противоположной ситуации, когда числам приписывается роль координат выдуманного пространства. При этом, конечно, получается виртуальная фикция, но она работает. Будит воображение там, где изначально никакой геометрии нет.

Когда дело было сделано, об аналитической геометрии стали забывать, начиная сразу с n измерений. У подсознания ушла почва из под ног, — ибо ему нужна другая последовательность. Геометрия, потом алгебра. Наглядность, затем абстракция.

1.1. Координаты и векторы

Координатами называют числа, определяющие положение точки на плоскости R^2 или в пространстве R^3 . Обозначения R^2 , R^3 —

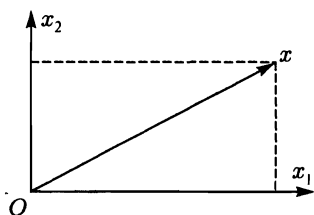


Рис. 1.1

пока для удобства. Прямоугольные (*декартовы* [9]) координаты точки на плоскости — суть снабженные знаками *плюс* или *минус* расстояния от точки x до двух взаимно перпендикулярных прямых Ox_1 и Ox_2 — *осей координат*, — точка пересечения которых считается началом координат. На рис. 1.1 точка $x = \{x_1, x_2\}$

описывается двумя координатами x_1 и x_2 . Горизонтальная x_1 — называется *абсциссой*, вертикальная x_2 — *ординатой*.

Систему декартовых координат в пространстве задают три взаимно перпендикулярные плоскости, относительно которых положение точки определяется тремя числами аналогично предыдущему, $x = \{x_1, x_2, x_3\}$.

Точку $x = \{x_1, x_2, x_3\}$ называют также *вектором* либо *радиусом-вектором*. Изначально, правда, вектор определяют как направленный отрезок прямой, изображая его со стрелочкой на конце. Однако все векторы, которые одинаково направлены и равны по длине, считаются равными¹⁾, — что, собственно, и позволяет ограничиться рассмотрением векторов, исходящих из начала координат.

Сумма $x = \{x_1, x_2, x_3\}$ и $y = \{y_1, y_2, y_3\}$ определяется как

$$x + y = \{x_1 + y_1, x_2 + y_2, x_3 + y_3\},$$

что равносильно сложению векторов по правилу параллелограмма, эквивалентом которого является *правило треугольника*. Преимущества последнего выявляются при сложении нескольких векторов (рис. 1.2): каждый следующий слагаемый вектор приставляется началом к концу предыдущего — замыкающий вектор дает сумму. Вычитание выводится из сложения: $b - a$ определяется как вектор, который в сумме с a дает b . Этому соответствует простой геометрический трюк: начала a и b совмещаются, а концы соединяются отрезком, направленным к b , что и дает разность $b - a$ (рис. 1.2).

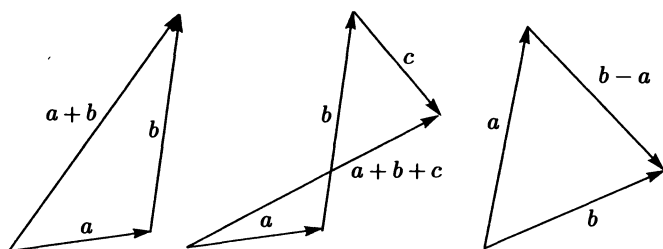


Рис. 1.2

Умножение на скаляр λ ,

$$\lambda x = \{\lambda x_1, \lambda x_2, \lambda x_3\},$$

растягивает ($|\lambda| > 1$) или сжимает ($|\lambda| < 1$) вектор x , не меняя направления при $\lambda > 0$, и меняет его на противоположное при $\lambda < 0$.

¹⁾ Это характерно для идеологии свободных векторов, которая удобна в геометрических задачах.

Вектор иногда определяют как направленный отрезок, но тогда в пограничных ситуациях возникают проблемы. Скажем, вращение около некоторой оси на угол φ — вектор или не вектор? С одной стороны, ему можно сопоставить направленный по оси отрезок прямой длины φ , но это не решает проблему. Неприятность заключается в том, что вращения около разных осей не складываются по правилу параллелограмма²⁾, — и это приводит к отрицательному ответу на исходный вопрос. Поэтому в подобное определение вектора необходимо добавить требование, чтобы векторы складывались по правилу параллелограмма.

В данном контексте речь идет о математических объектах вида $\mathbf{x} = \{x_1, x_2, x_3\}$, которые складываются «как надо» по определению. Что касается правил сложения сил, моментов, потоков и других физических величин, — это задача другой епархии.

«Второе дно». Идея декартовых координат проста до гениальности, но громоздка, — и поначалу возникает впечатление, что векторные понятия нужны для краткости. Это лишь половина правды. Геометрия имеет «некоординатный» характер, и ее координатное описание часто уводит мысль в ложном направлении. Векторный язык точнее отражает суть геометрических свойств, что обнаруживается на каждом шагу.

Векторы, лежащие на одной прямой (одинаково или противоположно направленные), называют *коллинеарными*; лежащие в одной плоскости, — *компланарными*.

Говорят, что векторы³⁾ $\{x_1, \dots, x_k\}$ *линейно зависимы*, если существуют такие коэффициенты $\lambda_1, \dots, \lambda_k$, не все равные нулю, что

$$\lambda_1 x_1 + \dots + \lambda_k x_k = 0.$$

В противном случае говорят о *линейной независимости* векторов $\{x_1, \dots, x_k\}$.

Коллинеарные векторы всегда линейно зависимы. Компланарные — линейно зависимы, если их больше двух. (?)

Линейно независимое множество $\{e_1, e_2, e_3\}$ в пространстве считается *базисом*, если любой вектор $\mathbf{x}$ можно представить в виде

²⁾ Складываются по правилу параллелограмма бесконечно малые вращения, что влечет за собой векторную природу угловой скорости [3].

³⁾ Как правило, векторы с индексом выделяются жирным шрифтом, чтобы отличить их от координат.

линейной комбинации

$$x = x_1 e_1 + x_2 e_2 + x_3 e_3.$$

Величины x_i называют *координатами* точки $x \in R^3$.

Стандартный базис R^3 (единичные векторы, орты, $\{e_1, e_2, e_3\}$ направлены по осям декартовых координат):

$$e_1 = \begin{bmatrix} 1 \\ 0 \\ 0 \end{bmatrix}, \quad e_2 = \begin{bmatrix} 0 \\ 1 \\ 0 \end{bmatrix}, \quad e_3 = \begin{bmatrix} 0 \\ 0 \\ 1 \end{bmatrix}. \quad (1.1)$$

Подготавливая плацдарм для линейной алгебры, естественно нумеровать оси координат и орты $\{e_1, e_2, e_3\}$, что потом облегчает следующий шаг — тройка увеличивается до n , и обобщение готово. Но в обычном пространстве более удобны и привычны другие обозначения: оси координат x, y, z ; соответствующие орты i, j, k , — чему далее отдается предпочтение.

Упражнения

1. Из векторов a, b, c можно сложить треугольник, если $a + b + c = 0$.
2. Любые три вектора на плоскости либо четыре — в пространстве, — линейно зависимы.
3. Любые два не коллинеарных вектора a, b определяют плоскость⁴⁾, все точки которой могут быть записаны в виде

$$r = \lambda a + \mu b.$$

4. В случае $\lambda + \mu = 1$ точка $r = \lambda a + \mu b$ лежит на прямой, проходящей через концы векторов a и b . При дополнительном условии $\lambda, \mu \geq 0$ точка $r = \lambda a + \mu b$ лежит на отрезке, соединяющем a и b .
5. Если a, b, c — радиус-векторы вершин треугольника ABC , то

$$r = \frac{1}{3}(a + b + c)$$

является радиусом-вектором точки пересечения медиан.

Проекция a_b вектора a на вектор (направление) b определяется формулой

$$a_b = a \cos \varphi,$$

⁴⁾ Проходящую через три точки: нуль, a и b .

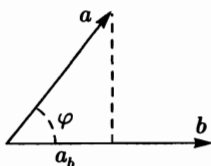


Рис. 1.3

где φ — угол между векторами a и b (рис. 1.3), не важно по малой или большой дуге измеренный — в силу $\cos \varphi = \cos (2\pi - \varphi)$.

Проекции a_x, a_y, a_z на декартовы оси x, y, z представляют собой координаты вектора a .

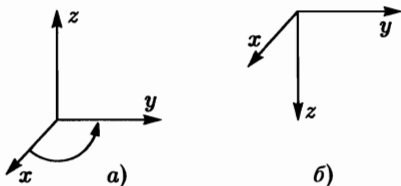


Рис. 1.4

При дальнейшем развитии теории важным оказывается деление систем координат на две категории. Прямоугольную систему *Oxyz* называют *правой*, если буравчик при вращении от x к y движется вдоль z — рис. 1.4 а, и *левой*, если ось z направлена противоположно — рис. 1.4 б.

Косоугольные координаты. Иногда декартовой называют косоугольную систему координат общего вида, в которой координаты определяются разложением вектора⁵⁾ — рис. 1.5 — по двум (на плоскости) или по трем (в пространстве) направлениям, которые не обязательно взаимно перпендикулярны.

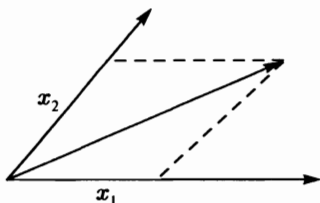


Рис. 1.5

Атрибутом «декартовости» при этом является наличие базиса. Примером не декартовых координат могут служить (нелинейные) полярные координаты (вектор характеризуется длиной r и углом φ с выделенным направлением).

*В общем случае не прямоугольных координат понятие ориентации определяется на следующее определение. Упорядоченная тройка некопланарных векторов a, b, c , исходящих из точки O , называется **правой**, если для наблюдателя, расположенного в нуле, обход концов a, b, c в указанном порядке происходит по часовой стрелке. В противном случае тройка a, b, c — **левая**. Соответственно классифицируются базисы.*

⁵⁾ Операция разложения вектора по заданным направлениям известна со школы (разложение сил, скоростей на несколько составляющих — две на плоскости, три в пространстве), и заключается в построении подходящего параллелограмма (либо параллелепипеда) с последующим представлением диагонали в виде векторной суммы сторон.

Если не оговорено противное, под декартовыми — подразумеваются далее прямоугольные координаты.

1.2. Описание геометрических объектов

На плоскости с декартовыми координатами x, y уравнение

$$x^2 + y^2 = r^2$$

описывает окружность радиуса r с центром в начале координат — рис. 1.6 а. Хорошо известно также, что графиком функции $y = kx + b$ служит прямая линия — рис. 1.6 б. Уравнение прямой обычно записывают в симметричном виде

$$\alpha x + \beta y + \gamma = 0.$$

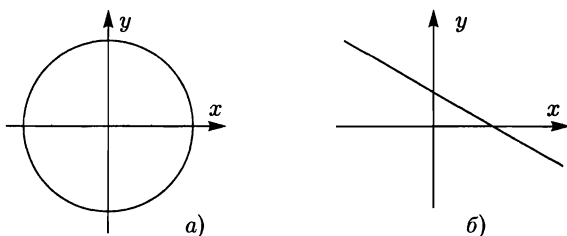


Рис. 1.6

Аналогичное уравнение в R^3

$$\alpha x + \beta y + \gamma z + \delta = 0$$

(1.2)

описывает плоскость.

«Школьный вариант» освоения этих уравнений довольно неуклюж, но он выполняет определенную задачу установления контакта с подсознанием. Проблема ведь заключается в том, чтобы сначала понять суть не умом, а нутром. Для этого строится прямая $y = x + 1$, затем $y = 2x - 1$, потом $y = -x + 3$. На каком-то этапе подсознание начинает злиться на однообразие — и тогда уже все примеры можно зачеркнуть, оставив одну запись $y = kx + b$.

Что касается плоскости, то $\alpha x + \beta y + \gamma z + \delta = 0$ еще не конец пути. Беда в том, что язык описания кое-что не ухватывает. Новую точку зрения дает понятие скалярного произведения векторов.

Скалярное произведение определяется как произведение длин векторов на косинус угла между ними⁶⁾ (рис. 1.7),

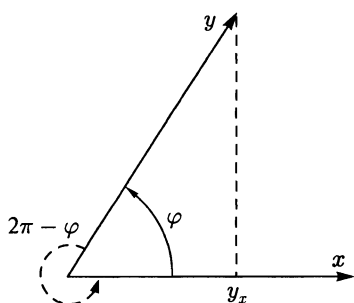


Рис. 1.7

$$x \cdot y = |x| \cdot |y| \cos \varphi. \quad (1.3)$$

Другими словами, xy (точка в обозначении скалярного произведения часто опускается) есть произведение длины x на проекцию вектора y на вектор x , т. е.

$$xy = |x|y_x.$$

Для обозначения скалярного произведения используется также более громоздкая, но иногда полезная запись (x, y) .

Скалярное произведение удовлетворяет обычным свойствам умножения:

$$xy = yx \quad (\text{коммутативность}),$$

$$x(y + z) = xy + xz \quad (\text{дистрибутивный закон}).$$

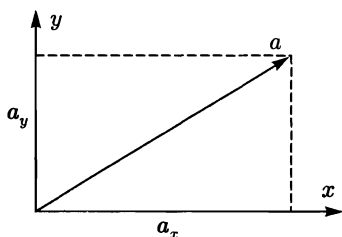


Рис. 1.8

Второе вытекает из равенства проекции суммы векторов — сумме проекций.

Отметим попутно, что длина $|a|$ (еще говорят — *норма*, и пишут $\|a\|$) вектора $a = \{a_x, a_y\}$ по теореме Пифагора равна (рис. 1.8)

$$\|a\| = \sqrt{a_x^2 + a_y^2},$$

$$\text{т. е. } a^2 = a \cdot a = a_x^2 + a_y^2.$$

В пространстве, соответственно,

$$\|a\|^2 = a^2 = a_x^2 + a_y^2 + a_z^2,$$

⁶⁾ Из-за четности и периодичности косинуса, $\cos(2\pi - \varphi) = \cos \varphi$, — не важно, как измеряется угол.

откуда

$$(a \pm b)^2 = (a_x \pm b_x)^2 + (a_y \pm b_y)^2 + (a_z \pm b_z)^2,$$

что приводит к

$$a \cdot b = \frac{(a+b)^2 - (a-b)^2}{4} = a_x b_x + a_y b_y + a_z b_z.$$

Последнюю формулу

$$a \cdot b = a_x b_x + a_y b_y + a_z b_z \quad (1.4)$$

иногда принимают за исходное определение скалярного произведения, выводя (1.3) в качестве следствия.

Произведение $a \cdot b = |a| \cdot |b| \cos \varphi$ отражает широко распространенный в природе способ взаимодействия векторов, характеризуемый умножением длины одного вектора на проекцию другого: $|a| b_a$. Например, работа силы F на перемещении s равна

$$F \cdot s = |F| \cdot |s| \cos \varphi,$$

поскольку «работает» только составляющая F , направленная вдоль s .

Поток жидкости через площадку площади ΔS равен скалярному произведению $v \Delta S$, где v — скорость течения (в районе ΔS), а вектор ΔS считается направленным по нормали к площадке. Физически опять-таки понятная ситуация: количество жидкости, протекающей через ΔS , определяется нормальной составляющей потока и не зависит от касательной составляющей.

Скалярное произведение — удобный инструмент при решении многих задач. Пусть, скажем, исходная декартова система координат имеет орты $\mathbf{i}, \mathbf{j}, \mathbf{k}$, а новая (тоже декартова) — орты $\mathbf{i}', \mathbf{j}', \mathbf{k}'$. Поскольку координаты вектора $a = \{a_x, a_y, a_z\}$ — есть проекции a на оси x, y, z , в новой системе координаты $a' = \{a'_{x'}, a'_{y'}, a'_{z'}\}$ — снова проекции, но уже на оси x', y', z' . Поэтому, например,

$$\begin{aligned} a'_{x'} &= a \cdot \mathbf{i}' = a_x (\mathbf{i} \cdot \mathbf{i}') + a_y (\mathbf{j} \cdot \mathbf{i}') + a_z (\mathbf{k} \cdot \mathbf{i}') = \\ &= a_x i'_{x'} + a_y i'_{y'} + a_z i'_{z'}, \end{aligned}$$

где коэффициенты дают запись орта $\mathbf{i}' = \{i'_{x'}, i'_{y'}, i'_{z'}\}$ в старой системе координат.

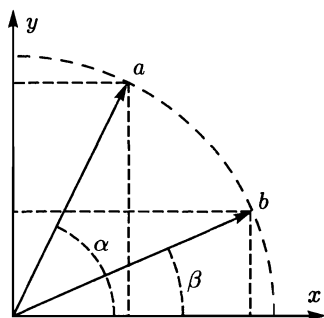


Рис. 1.9

- Произведение единичных векторов (рис. 1.9)

$$a = \{\cos \alpha, \sin \alpha\}, \quad b = \{\cos \beta, \sin \beta\}$$

сразу дает формулу

$$ab = \cos(\alpha - \beta) = \cos \alpha \cos \beta + \sin \alpha \sin \beta,$$

поскольку $|a| = |b| = 1$.

- Векторное описание треугольника ABC (рис. 1.10) легко приводит к теореме косинусов

$$(a + b)(a + b) = |c|^2 = |a|^2 + |b|^2 - 2|a| \cdot |b| \cos C,$$

так как $\cos(\pi - C) = -\cos C$.

- В косоугольных координатах скалярные произведения единичных направляющих (по осям) векторов — получаются ненулевыми:

$$i \cdot j = \cos(x, y), \quad i \cdot k = \cos(x, z), \quad j \cdot k = \cos(y, z).$$

Это, соответственно, усложняет формулы. Например,

$$a^2 = a_x^2 + a_y^2 + a_z^2 + 2a_x a_y \cos(x, y) + 2a_x a_z \cos(x, z) + 2a_y a_z \cos(y, z),$$

что получается возведением в квадрат $a = a_x i + a_y j + a_z k$.

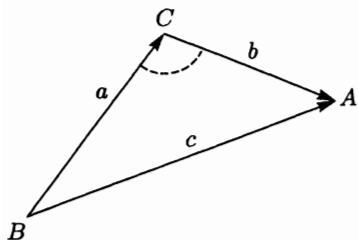


Рис. 1.10

Уравнение

Равенство нулю скалярного произведения ab означает — в силу $\cos(a, b) = 0$ — *ортogonalность* ненулевых векторов a и b . Поэтому $a r = 0$ (при фиксированном a и текущем $r = \{x, y, z\}$) представляет собой уравнение плоскости, *ортogonalной* вектору a и проходящей через начало координат.

$$a r = \delta,$$

(1.5)

как и (1.2), задает плоскость, не обязательно проходящую через начало координат, и означает, что проекция радиуса-вектора $r = \{x, y, z\}$ любой точки этой плоскости на направление a одна и та же. Если, дополнительно, $|a| = 1$, то эта проекция численно равна δ .

Понятно, что (1.5) лишь другая (векторная) форма записи уравнения (1.2), но она привносит дополнительную интерпретацию, которая позволяет наглядно мыслить.

Прямые и плоскости в евклидовой геометрии — одни из основных объектов. Соответствующую роль в аналитической геометрии играет уравнение (1.5), манипуляции с которым помогают описать многие задачи. Например, совокупности двух уравнений

$$ar = \delta, \quad br = \zeta$$

удовлетворяет прямая в R^3 (пересечение плоскостей).

Понятия точки, прямой и плоскости в геометрии Евклида *первичны* и потому — *неопределяемы*. Аналитическая геометрия дает определение этих понятий⁷⁾, но никакого чуда при этом не происходит — первичные понятия отодвигаются в другую область.

1.3. Векторное произведение

Векторное произведение,

$$c = a \times b,$$

определяется как вектор c , ортогональный векторам a и b , и составляющий тройку $\{a, b, c\}$, совпадающую по ориентации с тройкой базисных векторов, а длина c равна площади параллелограмма, построенного на векторах a, b (рис. 1.11),

$$|c| = |a| \cdot |b| \sin \varphi.$$

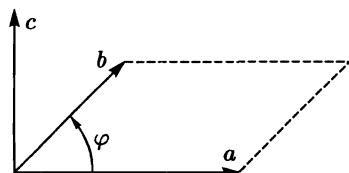


Рис. 1.11

Часто подразумевают, что система координат правая, — тогда проблему направления c решает «правило буравчика», и странность определения $c = a \times b$ забывается. «Странность» заключается в том, что векторное произведение $a \times b$ как бы ощущает ориентацию пространства (меняет направление при замене левой системы координат на правую). Такие векторы называют *аксиальными*, и даже — *псевдовекторами*. Обычные векторы (сила, скорость) — их называют *полярными* — на замену системы координат не реагируют⁸⁾.

⁷⁾ Как множество решений тех или иных уравнений.

⁸⁾ Реагирует их описание. При замене координат $\{x, y, z\}$ на $\{-x, -y, -z\}$ описание полярного вектора $\{a_x, a_y, a_z\}$ переходит в $\{-a_x, -a_y, -a_z\}$. Запись аксиального — не меняется.

Векторное произведение не ассоциативно,

$$(a \times b) \times c \neq a \times (b \times c),$$

антикоммутативно,

$$x \times y = -y \times x,$$

но справедлив дистрибутивный закон,

$$(x + y) \times u = x \times u + y \times u, \quad (1.6)$$

проверка которого требует некоторых усилий.

◀ Для обоснования (1.6) представим вектор x в виде суммы $x = x_{\perp} + x_{\parallel}$, где составляющая $x_{\perp}$ перпендикулярна вектору u , а $x_{\parallel}$ — параллельна. Очевидно,

$$x \times u = x_{\perp} \times u, \quad (x + y)_{\perp} = x_{\perp} + y_{\perp}.$$

Равенство (1.6) теперь следует из почти очевидного

$$(x_{\perp} + y_{\perp}) \times u = x_{\perp} \times u + y_{\perp} \times u. \quad \blacktriangleright$$

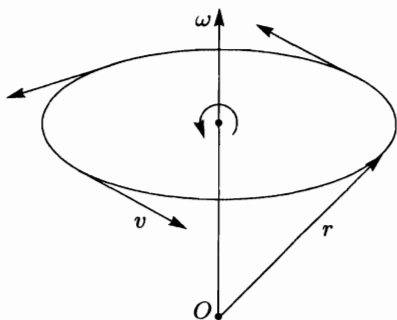


Рис. 1.12

Линейная скорость v конца радиуса-вектора r при вращении вокруг оси, проходящей через начало координат O , с угловой скоростью ω равна

$$v = \omega \times r, \quad (1.7)$$

где (пока формально) вектор ω направлен по оси вращения (в сторону, определяемую по «правилу буравчика», рис. 1.12).

Если тело участвует в двух вращениях ω^1 и ω^2 (с пересекающимися осями), то линейные скорости

$$v^1 = \omega^1 \times r \quad \text{и} \quad v^2 = \omega^2 \times r$$

складываются как векторы, и в силу (1.6) —

$$v = \omega^1 \times r + \omega^2 \times r = (\omega^1 + \omega^2) \times r.$$

Это означает, что результирующее движение происходит с угловой скоростью $\omega = \omega^1 + \omega^2$. Таким образом, угловые скорости складываются по правилу параллелограмма, и поэтому являются полноценными векторами⁹⁾.

⁹⁾ Аксиальными.

В правой декартовой системе координат:

$$\mathbf{i} \times \mathbf{j} = \mathbf{k}, \quad \mathbf{j} \times \mathbf{k} = \mathbf{i}, \quad \mathbf{k} \times \mathbf{i} = \mathbf{j},$$

что совместно с дистрибутивным законом (1.6) легко приводит¹⁰⁾ к формулам, определяющим координаты векторного произведения $\mathbf{a} \times \mathbf{b}$:

$$\begin{aligned} (a \times b)_x &= a_y b_z - a_z b_y, \\ (a \times b)_y &= a_z b_x - a_x b_z, \\ (a \times b)_z &= a_x b_y - a_y b_x. \end{aligned} \tag{1.8}$$

Векторное произведение оказывается удобным инструментом при описании многих физических явлений.

Моментом силы $\mathbf{F}$ относительно точки O называется вектор

$$\mathbf{M} = \mathbf{r} \times \mathbf{F}.$$

Аналогично определяется момент количества движения:

$$\mathbf{N} = \mathbf{r} \times m\mathbf{v}.$$

Дифференцируя $\mathbf{N}$, получаем

$$\dot{\mathbf{N}} = \dot{\mathbf{r}} \times m\mathbf{v} + \mathbf{r} \times m\dot{\mathbf{v}} = \mathbf{r} \times m\dot{\mathbf{v}} = \mathbf{r} \times \mathbf{F} = \mathbf{M},$$

поскольку $m\dot{\mathbf{v}} = \mathbf{F}$ и $\dot{\mathbf{r}} \times m\mathbf{v} = \mathbf{v} \times m\mathbf{v} = 0$.

Вот несколько примеров из геометрии.

- Раскрытие скобок в $(\mathbf{a} + \mathbf{b}) \times (\mathbf{a} - \mathbf{b})$ приводит к

$$(\mathbf{a} + \mathbf{b}) \times (\mathbf{a} - \mathbf{b}) = -2\mathbf{a} \times \mathbf{b},$$

откуда следует, что площадь параллелограмма, построенного на диагоналях, в два раза больше площади исходного параллелограмма.

- Перемножение векторов, изображенных на рис. 1.9, дает (в случае правой ориентации $\{\mathbf{i}, \mathbf{j}, \mathbf{k}\}$)

$$\mathbf{a} \times \mathbf{b} = -\sin(\alpha - \beta) \mathbf{k},$$

т. е. ненулевой является только координата

$$(a \times b)_z = a_x b_y - a_y b_x = \cos \alpha \sin \beta - \sin \alpha \cos \beta.$$

Поэтому

$$\sin(\alpha - \beta) = \sin \alpha \cos \beta - \cos \alpha \sin \beta.$$

- Уравнением прямой, параллельной вектору $\mathbf{a}$ и проходящей через точку $\mathbf{b}$ (конец вектора $\mathbf{b}$), служит

$$(\mathbf{r} - \mathbf{b}) \times \mathbf{a} = 0.$$

¹⁰⁾ В результате раскрытия скобок в $(a_x \mathbf{i} + a_y \mathbf{j} + a_z \mathbf{k}) \times (b_x \mathbf{i} + b_y \mathbf{j} + b_z \mathbf{k})$.

1.4. Определители

Объем V параллелепипеда, построенного на векторах a, b, c (как на ребрах), определяется смешанным произведением

$$a \cdot (b \times c) = \pm V,$$

где $b \times c$ дает площадь параллелограмма, лежащего в основании, а проекция $a_{b \times c}$ — высота параллелепипеда. Что касается знака, то плюс получается в ситуации, когда ориентация тройки $\{a, b, c\}$ такая же, как у базиса, минус — в противном случае (базис левый, тройка $\{a, b, c\}$ правая; либо наоборот).

В случае компланарных векторов $a \cdot (b \times c) = 0$, поскольку параллелепипед сплюснут, и его объем — нуль.

Рутинное перемножение векторов в координатной форме после исключения нулевых членов приводит к

$$a \cdot (b \times c) = a_x b_y c_z - a_x b_z c_y + a_y b_z c_x - a_y b_x c_z + a_z b_x c_y - a_z b_y c_x. \quad (1.9)$$

Тот же результат, что и (1.9), дает раскрытие *определителя*, или *детерминанта*

$$\begin{vmatrix} a_x & a_y & a_z \\ b_x & b_y & b_z \\ c_x & c_y & c_z \end{vmatrix}, \quad (1.10)$$

общая теория которых будет изложена в следующей главе, но частные случаи определителей 2-го и 3-го порядков — это те самые примеры, которые более поучительны, чем правила.

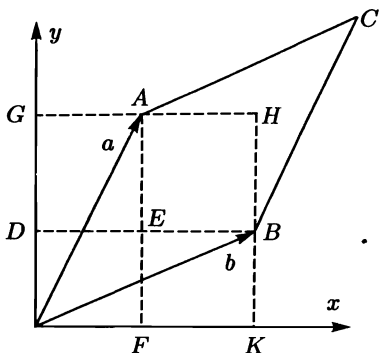


Рис. 1.13

Детерминант второго порядка по определению равен

$$\begin{vmatrix} a_x & a_y \\ b_x & b_y \end{vmatrix} = a_x b_y - a_y b_x. \quad (1.11)$$

Если ось z перпендикулярна плоскости векторов a, b , то $a_z = b_z = 0$ и из формул (1.8) для векторного произведения становится понятно, что детерминант (1.11) равен, по модулю, площади параллелограмма, построенного на векторах a и b .

С учетом (1.11) отсюда следует, что площадь параллелограмма $OACB$ (рис. 1.13) равна разности площадей прямоугольников $OGHK$ и $ODEF$. Вне векторной идеологии это не такая уж легкая геометрическая задача.

Сопоставление (1.11) с (1.9) приводит к равенству

$$\begin{vmatrix} a_x & a_y & a_z \\ b_x & b_y & b_z \\ c_x & c_y & c_z \end{vmatrix} = a_x \begin{vmatrix} b_y & b_z \\ c_y & c_z \end{vmatrix} - a_y \begin{vmatrix} b_x & b_z \\ c_x & c_z \end{vmatrix} + a_z \begin{vmatrix} b_x & b_y \\ c_x & c_y \end{vmatrix}, \quad (1.12)$$

которое сводит вычисление детерминанта 3-го порядка к вычислению детерминантов 2-го порядка.

Из этой формулы следует, что векторное произведение a на b можно записать в форме определителя

$$a \times b = \begin{vmatrix} \mathbf{i} & \mathbf{j} & \mathbf{k} \\ a_x & a_y & a_z \\ b_x & b_y & b_z \end{vmatrix},$$

где $\mathbf{i}, \mathbf{j}, \mathbf{k}$ — единичные орты.

1.5. Матрицы и преобразования

При манипулировании векторами часто встречается квадратная таблица чисел вида

$$\begin{bmatrix} a_x & a_y & a_z \\ b_x & b_y & b_z \\ c_x & c_y & c_z \end{bmatrix},$$

которая фигурировала в записи определителя (1.10). На такую таблицу можно смотреть как на координатную запись друг под другом трех векторов (вектор-строк) a, b, c ; либо как на координатную запись в строчку трех вектор-столбцов

$$\begin{bmatrix} a_x \\ b_x \\ c_x \end{bmatrix}, \quad \begin{bmatrix} a_y \\ b_y \\ c_y \end{bmatrix}, \quad \begin{bmatrix} a_z \\ b_z \\ c_z \end{bmatrix}.$$

Более удобной оказывается запись таких таблиц (*матриц*) в виде

$$A = \begin{bmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \\ a_{31} & a_{32} & a_{33} \end{bmatrix}, \quad (1.13)$$

где первый индекс у элемента a_{ij} обозначает номер строки, второй — столбца. При этом через a_i будем обозначать i -ю строку A ,

через $a_{\cdot j}$ — j -й столбец, т. е.

$$a_{i\cdot} = [a_{i1}, a_{i2}, a_{i3}], \quad a_{\cdot j} = \begin{bmatrix} a_{1j} \\ a_{2j} \\ a_{3j} \end{bmatrix}.$$

Матрицы как операторы. Матрица (1.13) чаще всего возникает как таблица коэффициентов линейного преобразования

$$\begin{cases} a_{11}x + a_{12}y + a_{13}z = \alpha, \\ a_{21}x + a_{22}y + a_{23}z = \beta, \\ a_{31}x + a_{32}y + a_{33}z = \gamma, \end{cases} \quad (1.14)$$

переводящего вектор $r = \{x, y, z\}$ в вектор $d = \{\alpha, \beta, \gamma\}$.

Преобразование (1.14) записывают в виде $Ar = d$, что определяет по сути **правило умножения матрицы на вектор** и позволяет смотреть на матрицу как на линейный оператор¹¹⁾, действующий в пространстве трех переменных и преобразующий векторы по указанному правилу.

Если на произведение $C = AB$ матриц $A = [a_{ij}]$ и $B = [b_{ij}]$ смотреть как на последовательное применение линейного оператора B , потом A , то **правило умножения матриц**

$$c_{ij} = \sum_{k=1}^3 a_{ik}b_{kj} \quad (1.15)$$

получается обыкновенным приведением подобных в равенстве

$$Cx = ABx = A(Bx).$$

Из (1.15) видно, что элемент c_{ij} матрицы C получается скалярным умножением вектор-строки $a_{i\cdot}$ на вектор-столбец $b_{\cdot j}$, т. е.

$$c_{ij} = a_{i\cdot} \cdot b_{\cdot j}.$$

¹¹⁾ Отображение, функцию.

Обратной к A называют матрицу A^{-1} , такую что

$$A^{-1}A = AA^{-1} = I,$$

где I — *единичная матрица*,

$$I = \begin{bmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{bmatrix},$$

определяющая тождественное преобразование $Ix \equiv x$.

Наконец, матрица A^* с элементами $a_{ij}^* = a_{ji}$ (строки становятся столбцами, столбцы — строками) называется *транспонированной* к A . Вместо A^* используются также равноценные обозначения A^T либо A' .

Легко проверяется, что $|A| = |A^*|$.

Линейные уравнения. На (1.14) можно смотреть также как на уравнение относительно $r = \{x, y, z\}$ при заданном $d = \{\alpha, \beta, \gamma\}$.

Здесь полезно рассмотрение уравнения (1.14) с разных точек зрения. С одной стороны, (1.14) есть

$$\begin{bmatrix} \alpha \\ \beta \\ \gamma \end{bmatrix} = \begin{bmatrix} a_{11} \\ a_{21} \\ a_{31} \end{bmatrix} x + \begin{bmatrix} a_{12} \\ a_{22} \\ a_{32} \end{bmatrix} y + \begin{bmatrix} a_{13} \\ a_{23} \\ a_{33} \end{bmatrix} z,$$

т. е. решение $\{x, y, z\}$ представляет собой координаты

$$\text{вектора } d = \begin{bmatrix} \alpha \\ \beta \\ \gamma \end{bmatrix} \text{ в базисе } \left\{ \begin{bmatrix} a_{11} \\ a_{21} \\ a_{31} \end{bmatrix}, \begin{bmatrix} a_{12} \\ a_{22} \\ a_{32} \end{bmatrix}, \begin{bmatrix} a_{13} \\ a_{23} \\ a_{33} \end{bmatrix} \right\}.$$

При этом понятно, что вектор-столбцы могут составлять базис, если они некопланарны (не лежат в одной плоскости). Для этого необходимо и достаточно, чтобы определитель $|A|$ был не равен нулю.

С другой стороны, (1.14) можно рассматривать как

$$\begin{cases} a \cdot r = \alpha, \\ b \cdot r = \beta, \\ c \cdot r = \gamma, \end{cases} \quad (1.16)$$

где

$$a = \{a_{11}, a_{12}, a_{13}\}, \quad b = \{a_{21}, a_{22}, a_{23}\}, \quad c = \{a_{31}, a_{32}, a_{33}\}.$$

Каждое из трех уравнений (1.16) описывает плоскость — см. (1.5), — причем векторы a, b, c перпендикулярны к этим плоскостям. Понятно, что плоскости пересекаются в единственной точке, если нормали a, b, c некопланарны, для чего опять-таки необходимо и достаточно $|A| \neq 0$. Остальные возможности геометрически также очевидны¹²⁾.

Хотя решить линейную систему уравнений численно всегда легко последовательным исключением переменных, интерес представляют единообразные формулы и возможные интерпретации.

С целью получения общего решения (1.16) рассмотрим сначала систему

$$\begin{cases} a \cdot r = 1, \\ b \cdot r = 0, \\ c \cdot r = 0. \end{cases} \quad (1.17)$$

Поскольку в (1.17) скалярное произведение r на b и на c равно нулю, то r перпендикулярен b и c , т. е. параллелен $b \times c$, а значит $r = \lambda(b \times c)$ при некотором λ , которое определяется подстановкой $r = \lambda(b \times c)$ в первое уравнение (1.17), откуда

$$\lambda[a \cdot (b \times c)] = 1 \Rightarrow \lambda = \frac{1}{a \cdot (b \times c)} \Rightarrow r = \frac{(b \times c)}{a \cdot (b \times c)}.$$

Обозначая полученное решение r через a° и два аналогичных через b° и c° ,

$$a^\circ = \frac{(b \times c)}{a \cdot (b \times c)}, \quad b^\circ = \frac{(c \times a)}{a \cdot (b \times c)}, \quad c^\circ = \frac{(a \times b)}{a \cdot (b \times c)},$$

имеем

$$\begin{cases} a \cdot a^\circ = 1, \\ b \cdot a^\circ = 0, \\ c \cdot a^\circ = 0, \end{cases} \quad \begin{cases} a \cdot b^\circ = 0, \\ b \cdot b^\circ = 1, \\ c \cdot b^\circ = 0, \end{cases} \quad \begin{cases} a \cdot c^\circ = 0, \\ b \cdot c^\circ = 0, \\ c \cdot c^\circ = 1. \end{cases} \quad (1.18)$$

Умножая теперь системы (1.18), соответственно, на α, β, γ и складывая, получаем

$$\begin{cases} a \cdot (\alpha a^\circ + \beta b^\circ + \gamma c^\circ) = \alpha, \\ b \cdot (\alpha a^\circ + \beta b^\circ + \gamma c^\circ) = \beta, \\ c \cdot (\alpha a^\circ + \beta b^\circ + \gamma c^\circ) = \gamma, \end{cases}$$

откуда видно, что

$$\boxed{r = \alpha a^\circ + \beta b^\circ + \gamma c^\circ} \quad (1.19)$$

является решением системы (1.16).

¹²⁾ У плоскостей может не быть общей точки пересечения (решение системы (1.16) невозможно). Либо пересечением плоскостей является прямая — решений системы (1.16) бесконечно много.

Таким образом, знание векторов $a^\circ, b^\circ, c^\circ$ позволяет практически мгновенно решать исходную систему уравнений (1.16) при любой правой части.

Говорят, что векторы $a^\circ, b^\circ, c^\circ$ взаимны по отношению к a, b, c . Из (1.18) видно, что точно так же a, b, c взаимны по отношению к $a^\circ, b^\circ, c^\circ$. Очевидно,

$$\begin{bmatrix} a_x & a_y & a_z \\ b_x & b_y & b_z \\ c_x & c_y & c_z \end{bmatrix} \cdot \begin{bmatrix} a_x^\circ & b_x^\circ & c_x^\circ \\ a_y^\circ & b_y^\circ & c_y^\circ \\ a_z^\circ & b_z^\circ & c_z^\circ \end{bmatrix} = \begin{bmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{bmatrix},$$

т. е. матрица, составленная из вектор-столбцов $a^\circ, b^\circ, c^\circ$, является обратной по отношению к исходной. Именно поэтому система (1.16) сразу решается, как только определены $a^\circ, b^\circ, c^\circ$. На матричном языке это выглядит тривиально. Уравнение $A\mathbf{r} = \mathbf{d}$ после умножения слева на A^{-1} переходит в $\mathbf{r} = A^{-1}\mathbf{d}$, что сразу дает решение, т. е. выражает $\mathbf{r}$ через $\mathbf{d}$.

Правило Крамера. Многое из сказанного может показаться переливанием из пустого в порожнее, поскольку мы толчемся вокруг да около. Как решить систему — ясно, а смена позиций меняет форму, но ничего не прибавляет по существу. По крайней мере, так кажется. Тем не менее разные взгляды на один и тот же факт позволяют математике расти и достигать новых областей, очень далеких от первоначальных постановок задач. Разные взгляды — разные ассоциации, связи, возможные обобщения.

Остановимся еще на одной интерпретации, для чего перенумеруем переменные, считая $\mathbf{r} = \{x_1, x_2, x_3\}$, и перепишем систему уравнений (1.14) в виде

$$\begin{cases} a_{11}x_1 + a_{12}x_2 + a_{13}x_3 = \alpha, \\ a_{21}x_1 + a_{22}x_2 + a_{23}x_3 = \beta, \\ a_{31}x_1 + a_{32}x_2 + a_{33}x_3 = \gamma, \end{cases} \quad (1.20)$$

т. е. по-прежнему $A\mathbf{r} = \mathbf{d}$.

Пусть $D = |A|$, а D_i обозначает определитель матрицы, которая получается из A заменой i -го столбца на вектор-столбец $\mathbf{d}$. Тогда решение (1.19) с учетом формулы разложения определителя (1.12) можно записать как

$$\mathbf{r} = \{x_1, x_2, x_3\}, \quad \boxed{x_i = \frac{D_i}{D}},$$

что называют *правилом Крамера*, которое точно так же формулируется в случае двух переменных¹³⁾.

В качестве упражнения полезно вернуться к представлению решения как координат вектора $\mathbf{d}$ в базисе, составленном из вектор-столбцов матрицы A . В двумерном случае

$$\begin{cases} a_1x_1 + b_1x_2 = \alpha, \\ a_2x_1 + b_2x_2 = \beta, \end{cases}$$

¹³⁾ Равно как и n переменных — см. следующую главу.

результат получается совсем легко. Векторная запись системы уравнений имеет вид $a\mathbf{x}_1 + b\mathbf{x}_2 = \mathbf{d}$ и после умножения (векторно) на b дает $(a \times b)\mathbf{x}_1 = (d \times b)$ (поскольку $b \times b = 0$), откуда $\mathbf{x}_1 = (d \times b)/(a \times b)$, что есть $\mathbf{x}_1 = D_1/D$. Аналогично получается $\mathbf{x}_2 = D_2/D$. Деление в данном случае вектора на вектор возможно в силу коллинеарности векторов.

Переход к другим координатам. Пусть имеется два базиса,

$$\{\mathbf{e}_1, \mathbf{e}_2\} \quad \text{и} \quad \{\mathbf{e}'_1, \mathbf{e}'_2\},$$

и

$$\begin{cases} \mathbf{e}'_1 = s_{11}\mathbf{e}_1 + s_{12}\mathbf{e}_2, \\ \mathbf{e}'_2 = s_{21}\mathbf{e}_1 + s_{22}\mathbf{e}_2 \end{cases} \quad (1.21)$$

определяет их взаимосвязь.

Допустим теперь, что x_i обозначают координаты вектора $\mathbf{x}$ в базисе $\{\mathbf{e}\}$, а x'_i — в базисе $\{\mathbf{e}'\}$. Подставляя $\mathbf{e}'_i$ из (1.21) в

$$\mathbf{x} = x_1\mathbf{e}_1 + x_2\mathbf{e}_2 = x'_1\mathbf{e}'_1 + x'_2\mathbf{e}'_2$$

и приравнивая коэффициенты при $\mathbf{e}_i$, — после раскрытия скобок и приведения подобных, получим

$$\begin{cases} x_1 = s_{11}x'_1 + s_{21}x'_2, \\ x_2 = s_{12}x'_1 + s_{22}x'_2, \end{cases}$$

т. е. $\mathbf{x} = S^*\mathbf{x}'$ и, соответственно, $\mathbf{x}' = (S^*)^{-1}\mathbf{x}$. Другими словами, если связь между базисами определяется матрицей S — в смысле (1.21), — то связь между координатами $\mathbf{x}' = T\mathbf{x}$ осуществляет матрица $T = (S^*)^{-1}$.

Неудобство, связанное с тем, что взаимоотношение базисов и координат выражается разными матрицами, не играет роли, поскольку рассматривать то и другое одновременно обычно не требуется. Пусть, говорят, новые координаты через старые (либо наоборот) выражаются с помощью T , — и о базисе не вспоминают. Главное, что T линейно.

В R^3 ситуация аналогична. Выражение одних координат через другие — это всегда определение проекций вектора

$$\mathbf{r} = \{r_x, r_y, r_z\}$$

на новые направления, задаваемые векторами другого базиса. При этом решает простое правило — проекция суммы векторов равна сумме проекций:

$$r_e = r_x \cos(\widehat{e, x}) + r_y \cos(\widehat{e, y}) + r_z \cos(\widehat{e, z}).$$

Если при переходе к новым координатам растяжения осей не происходит, — остается только поворот, который описывается матрицей

$$S(\varphi) = \begin{bmatrix} \cos \varphi & -\sin \varphi \\ \sin \varphi & \cos \varphi \end{bmatrix}.$$

Транспонированию S здесь соответствует замена φ на $-\varphi$ (замена вращения на угол φ противоположным вращением),

$$S^*(\varphi) = S(-\varphi), \quad S^{-1}(\varphi) = S^*(\varphi).$$

Упражнение

$$\begin{bmatrix} \cos(\alpha + \beta) & -\sin(\alpha + \beta) \\ \sin(\alpha + \beta) & \cos(\alpha + \beta) \end{bmatrix} = \begin{bmatrix} \cos \alpha & -\sin \alpha \\ \sin \alpha & \cos \alpha \end{bmatrix} \cdot \begin{bmatrix} \cos \beta & -\sin \beta \\ \sin \beta & \cos \beta \end{bmatrix}. \quad (?)$$

1.6. Прямые и плоскости

Различные описания плоскости. Как уравнение

$$\alpha x + \beta y + \gamma z + \delta = 0,$$

т. е. $a \cdot r + \delta = 0$, так и

$$r = \lambda a + \mu b + c,$$

задают плоскость. Первое описание служит уравнением плоскости с нормалью $a = \{\alpha, \beta, \gamma\}$ и позволяет легко проверить, принадлежит ли предъявляемая точка $r = \{x, y, z\}$ этой плоскости (проверкой равенства $a \cdot r + \delta = 0$), второе — параметрически задает плоскость, проходящую через точку c и ортогональную $a \times b$, — с его помощью удобно генерировать точки плоскости, выбирая произвольно параметры λ, μ .

Подобная двойственность характерна для математики вообще, и между такими описаниями иногда лежит пропасть¹⁴⁾. В данном случае различие не так велико. Из одного описания легко получить другое. (?)

¹⁴⁾ См. например, в [3] главу о вычислимости и доказуемости.

Приведенные варианты являются опорными, но не исчерпывают возможностей. Необходимо, скажем, описать плоскость, проходящую через три заданные точки

$$\mathbf{r}^1 = \{x_1, y_1, z_1\}, \quad \mathbf{r}^2 = \{x_2, y_2, z_2\}, \quad \mathbf{r}^3 = \{x_3, y_3, z_3\}.$$

Задача выглядит по-другому, но легко сводится к предыдущему. Полагая, например,

$$\mathbf{a} = \mathbf{r}^1 - \mathbf{r}^2, \quad \mathbf{b} = \mathbf{r}^1 - \mathbf{r}^3$$

и выбирая в качестве $\mathbf{c}$ любую из точек $\mathbf{r}^1, \mathbf{r}^2, \mathbf{r}^3$, получаем параметрическую запись плоскости¹⁵⁾ $\mathbf{r} = \lambda \mathbf{a} + \mu \mathbf{b} + \mathbf{c}$. Описанием первого типа будет

$$(\mathbf{a} \times \mathbf{b}) \cdot \mathbf{r} + \delta = 0,$$

где δ определяется подстановкой вместо $\mathbf{r}$ любой из точек $\mathbf{r}^1, \mathbf{r}^2, \mathbf{r}^3$. В иной форме это можно записать как равенство нулю определителя

$$\begin{vmatrix} x - x_1 & y - y_1 & z - z_1 \\ x_2 - x_1 & y_2 - y_1 & z_2 - z_1 \\ x_3 - x_1 & y_3 - y_1 & z_3 - z_1 \end{vmatrix} = 0,$$

что означает линейную зависимость векторов

$$\mathbf{r} - \mathbf{r}^1, \quad \mathbf{r}^1 - \mathbf{r}^2, \quad \mathbf{r}^1 - \mathbf{r}^3.$$

Кстати, вместо уравнения $\alpha x + \beta y + \gamma z + \delta = 0$ часто целесообразно использование его эквивалента

$$\alpha(x - x_0) + \beta(y - y_0) + \gamma(z - z_0) = 0,$$

где вместо не всегда удобного параметра δ фигурирует точка $\{x_0, y_0, z_0\}$, через которую проходит плоскость.

Прямая. Прямая на плоскости играет роль аналогичную той, которую играет плоскость в пространстве трех переменных. Поэтому

$$\alpha(x - x_0) + \beta(y - y_0) = 0$$

¹⁵⁾ Почему в качестве $\mathbf{c}$ можно выбирать любую из точек $\mathbf{r}^1, \mathbf{r}^2, \mathbf{r}^3$?

есть уравнение прямой, ортогональной вектору $\{\alpha, \beta\}$ и проходящей через точку $\{x_0, y_0\}$. Но для описания прямой в пространстве требуются уже два линейных уравнения

$$\boxed{a \cdot r + \delta = 0, \quad b \cdot r + \zeta = 0.} \quad (1.22)$$

Если r удовлетворяет обоим уравнениям, это означает, что r лежит на пересечении плоскостей, т. е. на прямой.

Разумеется, можно ограничиться одним нелинейным уравнением

$$(a \cdot r + \delta)^2 + (b \cdot r + \zeta)^2 = 0,$$

но это не меняет сути. Суть в каком-то смысле меняет запись (1.22) в виде одного линейного уравнения

$$\boxed{r \times (a \times b) = \delta b - \zeta a,} \quad (1.23)$$

однако в координатной записи (1.22) — это два уравнения, тогда как (1.23) приводит к трем уравнениям.

Поскольку о происхождении вектора $(a \times b)$ можно забыть, каноническое уравнение прямой можно записать так

$$\boxed{r \times c = d.}$$

Вектор c здесь направлен вдоль прямой, а $d = r^0 \times c$, где r^0 — любая точка, через которую проходит прямая.

Другой вариант описания прямой, проходящей через точку r^0 , — параметрическое задание, $r - r^0 = a\tau$, т. е.

$$x - x^0 = a_x \tau, \quad y - y^0 = a_y \tau, \quad z - z^0 = a_z \tau,$$

где вектор a задает направление прямой.

При необходимости описать прямую, проходящую через две точки

$$r^1 = \{x_1, y_1, z_1\}, \quad r^2 = \{x_2, y_2, z_2\},$$

достаточно в качестве a взять вектор $r^1 - r^2$. После исключения параметра τ , получается уравнение прямой ¹⁶⁾

$$\frac{x - x_0}{x_2 - x_1} = \frac{y - y_0}{y_2 - y_1} = \frac{z - z_0}{z_2 - z_1},$$

где точка $r^0 = \{x_0, y_0, z_0\}$, понятно, не может быть произвольной. Обычно полагают $r^0 = r^1$ либо $r^0 = r^2$.

Касательная плоскость. Пусть речь идет о функции двух переменных $z = f(x, y)$. Как известно, градиент

$$\text{grad } f = \left\{ \frac{\partial f}{\partial x}, \frac{\partial f}{\partial y} \right\}$$

перпендикулярен к линии постоянного уровня $f(x, y) = \text{const}$ в любой точке $\{x_0, y_0\}$. Поэтому касательная плоскость — проходящая через точку $\{x_0, y_0\}$ — к линии постоянного уровня функции $z = f(x, y)$ описывается уравнением (на рис. 1.14 это пунктирная прямая)

$$\frac{\partial f}{\partial x}(x - x_0) + \frac{\partial f}{\partial y}(y - y_0) = 0.$$

Если интерес представляет касательная плоскость к поверхности графика функции $z = f(x, y)$, то это сводится к предыдущему случаю рассмотрением функции $u = z - f(x, y)$ трех переменных. Ее градиент

$$\left\{ 1, -\frac{\partial f}{\partial x}, -\frac{\partial f}{\partial y} \right\}.$$

Поэтому касательная плоскость к поверхности $z = f(x, y)$ в точке $\{z_0, x_0, y_0\}$ определяется уравнением

$$z - z_0 = \frac{\partial f}{\partial x}(x - x_0) + \frac{\partial f}{\partial y}(y - y_0).$$

1.7. Геометрические задачи

Расстояние до плоскости. Пусть плоскость задает уравнение $r \cdot a = \delta$. Если r разложить по взаимно перпендикулярным направлениям

¹⁶⁾ На самом деле — два уравнения.

$\mathbf{r} = \mathbf{r}_a + \mathbf{r}_{\perp a}$, становится ясно, что минимальная длина у $\mathbf{r}$ будет в случае $\mathbf{r}_{\perp a} = 0$. Таким образом, расстояние от начала координат до плоскости равно

$$\|\mathbf{r}\|_{\min} = |\mathbf{r}_a| = \frac{|\delta|}{\|\mathbf{a}\|}.$$

Расстояние от произвольной точки $\mathbf{b}$ определяется так же — после переноса в $\mathbf{b}$ начала координат, что достигается заменой $\mathbf{r} = \mathbf{r}' + \mathbf{b}$, и в итоге даёт величину $|\delta - \mathbf{a} \cdot \mathbf{b}|/\|\mathbf{a}\|$.

Если требуется не само минимальное расстояние, а ближайшая к началу координат точка (на плоскости), то для ее определения надо решить систему из двух векторных уравнений: $\mathbf{r} \cdot \mathbf{a} = \delta$ и $\mathbf{r} \times \mathbf{a} = 0$, что равносильно $\mathbf{r}_{\perp a} = 0$.

При нежелании иметь дело с векторами — можно рассматривать координатную запись задачи:

$$x^2 + y^2 + z^2 \rightarrow \min, \quad \alpha x + \beta y + \gamma z = \delta,$$

которая легко решается методом множителей Лагранжа (см. [4, т. 1]).

Расстояние до прямой. Разнообразие возможностей здесь определяется вариантами задания прямой, которые легко сводятся друг к другу. В варианте (1.23), т. е. $\mathbf{r} \times \mathbf{a} = \mathbf{c}$, ближайшая к началу координат точка на прямой определяется системой уравнений

$$\mathbf{r} \times \mathbf{a} = \mathbf{c}, \quad \mathbf{r} \cdot \mathbf{a} = 0.$$

Находя решение и вычисляя длину соответствующего вектора $\mathbf{r}$, получаем расстояние до прямой (от начала координат).

Если сама ближайшая точка не нужна, то расстояние определяется совсем просто. Раскладывая $\mathbf{r}$ по взаимно перпендикулярным направлениям $\mathbf{r} = \mathbf{r}_{\parallel a} + \mathbf{r}_{\perp a}$, в ближайшей к началу координат точке имеем $|\mathbf{r}_{\perp a}| \cdot \|\mathbf{a}\| = \|\mathbf{c}\|$, откуда искомое расстояние: $\boxed{\|\mathbf{c}\|/\|\mathbf{a}\|}$.

Скрещивающиеся прямые. Пусть имеются две прямые, описываемые уравнениями

$$\mathbf{r} \times \mathbf{a} = \mathbf{A}, \quad \mathbf{s} \times \mathbf{b} = \mathbf{B},$$

где векторы $\mathbf{r}$ и $\mathbf{s}$ лежат, соответственно, на первой и второй прямой.

Умножая первое уравнение скалярно на $\mathbf{b}$, второе — на $\mathbf{a}$, имеем

$$(\mathbf{r} \times \mathbf{a}) \cdot \mathbf{b} = \mathbf{A} \cdot \mathbf{b}, \quad (\mathbf{s} \times \mathbf{b}) \cdot \mathbf{a} = \mathbf{B} \cdot \mathbf{a},$$

или равносильно ¹⁷⁾

$$\mathbf{r} \cdot (\mathbf{a} \times \mathbf{b}) = A \cdot b, \quad \mathbf{s} \cdot (\mathbf{b} \times \mathbf{a}) = B \cdot a.$$

Последние два уравнения при сложении дают

$$(\mathbf{r} - \mathbf{s}) \cdot (\mathbf{a} \times \mathbf{b}) = A \cdot b + B \cdot a.$$

Это означает, что проекция вектора $\mathbf{r} - \mathbf{s}$ на направление $(\mathbf{a} \times \mathbf{b})$ постоянна. Поэтому длина (норма) $\mathbf{r} - \mathbf{s}$ будет минимальна, когда $(\mathbf{r} - \mathbf{s}) \parallel (\mathbf{a} \times \mathbf{b})$. В этом случае $\|\mathbf{r} - \mathbf{s}\| \cdot \|\mathbf{a} \times \mathbf{b}\| = \|A \cdot b + B \cdot a\|$, откуда (минимальное) расстояние между прямыми равно

$$\frac{\|A \cdot b + B \cdot a\|}{\|\mathbf{a} \times \mathbf{b}\|}.$$

При $A \cdot b + B \cdot a = 0$ прямые пересекаются.

Чтобы найти сами точки $\mathbf{r}$ и $\mathbf{s}$, на которых достигается минимум расстояния, надо решить систему трех векторных уравнений

$$(\mathbf{r} - \mathbf{s}) \times (\mathbf{a} \times \mathbf{b}) = 0, \quad \mathbf{r} \times \mathbf{a} = A, \quad \mathbf{s} \times \mathbf{b} = B.$$

Частные задачи.

- В ряде случаев полезна следующая формула для двойного векторного произведения ¹⁸⁾

$$\mathbf{a} \times (\mathbf{b} \times \mathbf{c}) = \mathbf{b}(\mathbf{a} \cdot \mathbf{c}) - \mathbf{c}(\mathbf{a} \cdot \mathbf{b}).$$

Вот одно из полезных следствий:

$$\mathbf{n} \times (\mathbf{b} \times \mathbf{n}) = \mathbf{b} - \mathbf{n}(\mathbf{n} \cdot \mathbf{b}), \quad \text{если} \quad \mathbf{n} \cdot \mathbf{n} = 1,$$

что приводит к формуле

$$\mathbf{b} = \mathbf{n}(\mathbf{n} \cdot \mathbf{b}) + \mathbf{n} \times (\mathbf{b} \times \mathbf{n}), \quad \|\mathbf{n}\| = 1,$$

которая дает разложение $\mathbf{b}$ на две составляющие — параллельную и перпендикулярную единичному вектору $\mathbf{n}$.

- Координатно-векторное мышление «переворачивает» и упрощает почти любую геометрическую задачу. Вот, например, как выглядит векторное доказательство известного факта: *если диагонали четырехугольника $ABCD$ делят друг друга пополам, то $ABCD$ — параллелограмм.*

¹⁷⁾ Поскольку смешанное произведение $\mathbf{a} \cdot (\mathbf{b} \times \mathbf{c})$ не меняется при циклической перестановке векторов (и меняет знак при — не циклической) — объем параллелепипеда тот же самый. знак определяет ориентация.

¹⁸⁾ Мнемоническое правило: «абэцэ равно бац минус цаб».

Пусть a, b, c, d — радиус-векторы, соответственно, вершин A, B, C, D . По условию

$$\frac{1}{2}(a + c) = \frac{1}{2}(b + d),$$

откуда следует $b - a = c - d$, т.е. стороны AB и CD равны и параллельны. ►

Упражнения

- Пусть a, b, c — радиус-векторы вершин треугольника ABC . Тогда вектор

$$s = a \times b + a \times c + b \times c$$

равен удвоенной площади $\triangle ABC$ и перпендикулярен плоскости, в которой лежит треугольник.

- Уравнение

$$x_0x + y_0y + z_0z = R^2$$

определяет плоскость, касающуюся сферы $x^2 + y^2 + z^2 = R^2$ в точке $\{x_0, y_0, z_0\}$.

- Пусть a, b, c — радиус-векторы вершин треугольника ABC . Тогда радиус-вектор точки пересечения биссектрис равен

$$r = \frac{aBC + bAC + cAB}{BC + AC + AB},$$

где BC, AC, AB — длины соответствующих сторон,

$$r = \frac{a + b + c}{3} \quad \text{— точка пересечения медиан,}$$

$$r = \frac{a \operatorname{tg} A + b \operatorname{tg} B + c \operatorname{tg} C}{\operatorname{tg} A + \operatorname{tg} B + \operatorname{tg} C} \quad \text{— точка пересечения высот.}$$

1.8. Кривые и поверхности второго порядка

Уравнение на плоскости

$$A\xi^2 + 2B\xi\eta + C\eta^2 = D$$

поворотом осей координат

$$\begin{bmatrix} \xi \\ \eta \end{bmatrix} = \begin{bmatrix} \cos \varphi & -\sin \varphi \\ \sin \varphi & \cos \varphi \end{bmatrix} \cdot \begin{bmatrix} x \\ y \end{bmatrix}$$

приводится к виду

$$A'x^2 + 2B'xy + C'y^2 = D.$$

При этом выбором угла φ всегда можно обеспечить¹⁹⁾ $B' = 0$. Варианты в результате считаются на пальцах одной руки. Вырожденные случаи, когда одна из констант A' , C' обращается в нуль²⁰⁾, малоинтересны. В остающихся вариантах получаются эллипс и гипербола.

Эллипс. Каноническое уравнение эллипса:

$$\frac{x^2}{a^2} + \frac{y^2}{b^2} = 1, \quad a \geq b > 0.$$

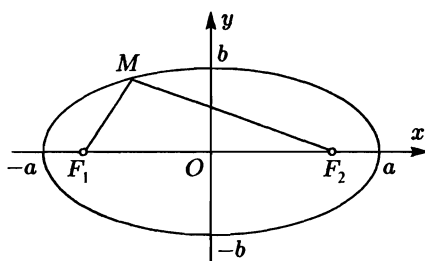


Рис. 1.15. Эллипс

На рис. 1.15 изображена соответствующая геометрическая фигура. Величины a и b равны длинам большой и малой полуосей. В случае $a = b$ получается окружность.

Эллипс обладает рядом замечательных свойств и может быть охарактеризован как:

- Множество точек плоскости, для каждой из которых сумма расстояний до фокусов F_1 , F_2 постоянна и равна $2a$. Расстояние между фокусами — $2c$, где $c^2 = a^2 - b^2$.
- Кривая, получаемая в пересечении кругового конуса с плоскостью.

Оптическое свойство: световой луч, исходящий из одного фокуса, после отражения от эллипса проходит через другой фокус.

Отношение $e = c/a$ называют *эксцентриситетом* эллипса. В полярных координатах ρ , φ эллипс описывается уравнением

$$\rho = \frac{p}{1 + e \cos \varphi},$$

где p — *фокальный параметр*, равный половине длины вертикальной хорды, проходящей через фокус.

¹⁹⁾ Принять на веру или проверить — каждый решает сам. Хорошо (плохо) и то, и другое. Проверки изошряют навыки, но замедляют ход. Вера, как воздушный шар, быстро возносит в облака, но отрывает от Земли.

²⁰⁾ Это соответствует ситуации

$$\det \begin{bmatrix} A & B \\ B & C \end{bmatrix} = \det \begin{bmatrix} A' & B' \\ B' & C' \end{bmatrix} = 0.$$

Пространственное обобщение — *эллипсоид*. Уравнение:

$$\frac{x^2}{a^2} + \frac{y^2}{b^2} + \frac{z^2}{c^2} = 1.$$

Гипербола. Гипербола — плоская кривая,

$$\frac{x^2}{a^2} - \frac{y^2}{b^2} = 1,$$

как и эллипс, представляет собой сечение конуса плоскостью.

На рис. 1.16 изображена соответствующая геометрическая фигура. Имеет две асимптоты $y = \pm \frac{a}{b}x$. Расстояние между фокусами F_1 , $F_2 - 2c$, где $c^2 = a^2 + b^2$.

Характеризуется как геометрическое множество точек плоскости, модуль разности расстояний которых до F_1 и F_2 постоянен и равен $2a$.

При $a = b$ асимптоты взаимно перпендикулярны. Если их принять за координатные оси, — уравнение гиперболы трансформируется в $y = k/x$.

Пространственное обобщение — *гиперболоид*. Уравнение:

$$\frac{x^2}{a^2} + \frac{y^2}{b^2} - \frac{z^2}{c^2} = 1.$$

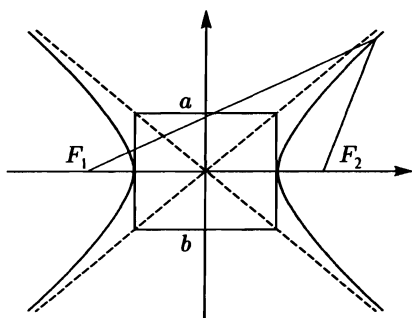


Рис. 1.16. Гипербола

Глава 2

Векторы и матрицы

2.1. Примеры линейных задач

Изучать абстрактную науку удобнее, имея в голове несколько содержательных задач. Конечно, преодолевать отвращение к конкретике нелегко, но это потом окупается.

Модель производства. Пусть x_j обозначает интенсивность j -го технологического процесса, a_{ij} — количество i -го продукта, производимого ($a_{ij} > 0$) или потребляемого ($a_{ij} < 0$) при единичной интенсивности j -го технологического процесса. Суммарное производство (потребление) i -го продукта определяется суммой

$$y_i = \sum_{j=1}^n a_{ij} x_j, \quad i = 1, \dots, m. \quad (2.1)$$

На этом фоне могут решаться задачи оптимизации типа $\sum_i c_i y_i \rightarrow \max$.

Межотраслевой баланс. Имеется n отраслей, i -я отрасль выпускает i -й продукт в количестве x_i . На выпуск единицы i -го продукта в системе затрачивается $a_{ij} \geq 0$ единиц j -го продукта. Понятно, что при выпуске набора

$$x = \{x_1, \dots, x_n\}$$

чистый выпуск i -го продукта равен

$$y_i = x_i - \sum_j a_{ij} x_j.$$

Вопрос о продуктивности модели, таким образом, сводится к положительной разрешимости системы уравнений¹⁾

$$x_i - \sum_j a_{ij} x_j = y_i, \quad i = 1, \dots, n.$$

¹⁾ То ли при любом, то ли при некотором наборе $y > 0$. Как оказывается, существование продуктивного плана $x > 0$ при некотором $y > 0$ достаточно для существования решения $x \geq 0$ — при любом $y \geq 0$.

Транспортная задача. Допустим, имеется n пунктов, в которых производится и потребляется некий вид продукта (хлеб, например). Из i -го пункта в j -й продукт перевозится в количестве x_{ij} . Естественно, что из i -го пункта нельзя вывезти больше продукта, чем там производится: $\sum_j x_{ij} \leq P_i$, а в j -й пункт надо

завозить не меньше, чем там потребляется $\sum_i x_{ij} \geq C_j$.

Если цена перевозки единицы продукта из i -го пункта в j -й — равна λ_{ij} , общая задача минимизации транспортных расходов выглядит так:

$$\sum_{ij} \lambda_{ij} x_{ij} \rightarrow \min, \quad \sum_j x_{ij} \leq P_i, \quad \sum_i x_{ij} \geq C_j.$$

Пассажиропотоки. В городе имеется n районов, P_i — число жителей i -го района, W_j — число работающих в j -м районе, наконец, x_{ij} — число живущих в i -м районе и работающих в j -м. Очевидно,

$$\sum_j x_{ij} = P_i, \quad \sum_i x_{ij} = W_j.$$

К этим линейным взаимосвязям переменных могут добавляться другие факторы и соотношения.

2.2. Векторы

При изучении вектор-функций $f(x_1, \dots, x_n)$ набор $x = \{x_1, \dots, x_n\}$ называют *вектором*, мысленно откладывая значения $x_1, \dots, x_n$ по координатным осям «выдуманного» пространства.

Аналогия с обычным представлением о векторах становится осмысленной при подходящем определении операций:

- *умножение на скаляр* λ ,

$$\lambda x = \{\lambda x_1, \dots, \lambda x_n\}, \quad (2.2)$$

- *сложение*,

$$x + y = \{x_1 + y_1, \dots, x_n + y_n\}, \quad (2.3)$$

- *скалярное произведение*,

$$x \cdot y = x_1 y_1 + \dots + x_n y_n, \quad (2.4)$$

для которого используются эквивалентные обозначения²⁾ (x, y) и xy .

²⁾ Эквивалентные обозначения не менее ценная вещь, чем синонимы в обычном языке.

С помощью скалярного произведения вводится норма (длина вектора),

$$\|x\| = \sqrt{x^2} = \sqrt{x_1^2 + \dots + x_n^2},$$

которую принято считать *евклидовой*.

Множество векторов, на которых введены перечисленные операции, называют n -мерным евклидовым пространством и обозначают R^n .

Операции (2.2)—(2.4) представляют собой обобщения аналогичных понятий в R^3 . Покоординатное сложение (2.3) в R^3 равносильно определению суммы векторов по правилу параллелограмма. Скалярное произведение в R^3 изначально иногда определяется как произведение длин векторов на косинус угла между ними, но потом выясняется, что формула (2.4) при $n = 3$ дает то же самое. В общем случае такая интерпретация может быть сохранена благодаря известному *неравенству Коши—Буняковского*³⁾

$$x \cdot y \leq \|x\| \cdot \|y\|, \quad (2.5)$$

которое дает возможность ввести понятие косинуса угла между векторами при любом n ,

$$\cos \varphi_{xy} = \frac{x \cdot y}{\|x\| \cdot \|y\|}. \quad (2.6)$$

◀ Неравенство (2.5) доказывается совсем просто. Раскрывая скобки в очевидном неравенстве $(x - \lambda y)^2 \geq 0$, имеем

$$\lambda^2 \|y\|^2 - 2\lambda xy + \|x\|^2 \geq 0$$

для любого λ , — что возможно лишь при неположительности дискриминанта квадратного многочлена — а это и есть (2.5). ▶

³⁾ Альтернативное название: *неравенство Коши—Шварца*.

Линейная независимость и базис. Говорят, что множество векторов⁴⁾ $\{x_1, \dots, x_k\}$ *линейно зависимо*, если существуют такие коэффициенты $\lambda_1, \dots, \lambda_k$, не все равные нулю, что

$$\lambda_1 x_1 + \dots + \lambda_k x_k = 0.$$

В противном случае говорят о *линейной независимости* элементов⁵⁾ $\{x_1, \dots, x_k\}$. *Коллинеарные* векторы⁶⁾, например, — линейно зависимы.

Если $u = \lambda_1 x_1 + \dots + \lambda_k x_k$, то говорят, что вектор u — *линейная комбинация* векторов $\{x_1, \dots, x_k\}$. Совокупность всевозможных

$$u = \lambda_1 x_1 + \dots + \lambda_k x_k$$

считают *линейной оболочкой* множества $\{x_1, \dots, x_k\}$. Линейная оболочка представляет собой *линейное пространство* P , *натянутое* на $\{x_1, \dots, x_k\}$. Характеристическое свойство P :

$$x, y \in P \Rightarrow \alpha x + \beta y \in P.$$

Линейно независимое множество $\{e_1, \dots, e_n\}$ называют *базисом*, если любой вектор $x \in R^n$ можно представить в виде линейной комбинации⁷⁾

$$x = x_1 e_1 + \dots + x_n e_n.$$

Величины $x_1, \dots, x_n$ именуются *координатами* точки $x \in R^n$.

Стандартный базис R^n :

$$e_1 = \begin{bmatrix} 1 \\ 0 \\ \vdots \\ 0 \end{bmatrix}, \quad e_2 = \begin{bmatrix} 0 \\ 1 \\ \vdots \\ 0 \end{bmatrix}, \quad \dots, \quad e_n = \begin{bmatrix} 0 \\ 0 \\ \vdots \\ 1 \end{bmatrix}. \quad (2.7)$$

Число n векторов, составляющих базис (не зависящее от выбора последнего), называют *размерностью* пространства.

Любые n линейно независимых векторов в R^n могут служить базисом.

⁴⁾ Напомним, что векторы с индексом выделяются жирным шрифтом, чтобы отличить их от координат.

⁵⁾ Элемент и точка — употребляются как синонимы вектора.

⁶⁾ Коллинеарными называют векторы, лежащие на одной прямой, т. е. векторы, которые одинаково или противоположно направлены.

⁷⁾ То есть R^n натянуто на свой базис.

Упражнения

- Если множество векторов $\{f_1, \dots, f_m\}$ составляет базис R^n , то $m = n$.
- Множество

$$f_1 = \begin{bmatrix} 1 \\ 0 \\ \vdots \\ 0 \end{bmatrix}, \quad f_2 = \begin{bmatrix} 1 \\ 1 \\ 0 \\ \vdots \end{bmatrix}, \quad \dots, \quad f_n = \begin{bmatrix} 1 \\ 1 \\ \vdots \\ 1 \end{bmatrix}$$

является базисом в множестве векторов вида $\begin{bmatrix} x_1 \\ \vdots \\ x_n \end{bmatrix}$.

Ортогональность. Векторы x, y определяют как *ортогональные*, если их скалярное произведение равно нулю, $x \cdot y = 0$.

Важная роль ортогональности заключена в том, что с ее помощью определяется понятие *плоскости* P_a , как множества элементов x , ортогональных некоторому вектору a ,

$$a \cdot x = 0, \quad \text{т. е.} \quad a_1 x_1 + \dots + a_n x_n = 0.$$

Понятно, что $x, y \in P_a$ влечет за собой $\alpha x + \beta y \in P_a$.

Описанная стартовая позиция представляет собой упрощенный взгляд на предмет — достаточный для понимания основных результатов линейной алгебры при некотором поэтапном расширении угла зрения. Важно лишь отдавать себе отчет в ограниченности избранного курса, чтобы для необходимых дополнений были открыты «форточки».

Вот некоторые бреши в изложении, которые будут заделаны позже. Неравенство Коши—Буняковского гарантирует $\cos \varphi \leq 1$, но этого мало для определения угла.

Правда, сама формула (2.6) далее не требуется, но она полезна в идеологическом отношении — ибо легализация n -мерного пространства в сфере подсознания нуждается в том, чтобы R^n имело те же атрибуты, что и R^3 . Обоснование (2.6) тривиально. Два вектора x, y определяют плоскость P (проходящую через три точки $0, x, y$), представляющую собой двумерное пространство R^2 , — а скалярное произведение $x \cdot y$ в R^2 как раз и дает косинус. Таким образом, при любой размерности R^n угол остается явлением двумерным, что льет воду на мельницу «легализации».

Однако легко сказка сказывается. На указанном пути возникает проблема с толкованием $x \cdot y$ в плоскости P . Беда в том, что в основе определения (2.4)

лежит координатное представление⁸⁾ векторов. При переходе к другой системе координат скалярное произведение как-то меняется (по форме записи) — пока не ясно как. В плоскости P система координат вообще не задана. Как ее задать, и нужно ли? Благодаря двумерности P , можно просто ограничиться произведением длин векторов на косинус угла между ними. Но совпадет ли это с результатом перемножения по формуле (2.4)?

Имеется также некоторое ощущение дискомфорта в связи с понятиями ортогональности, базиса и размерности. Пока есть пробелы в определении угла, ортогональность слегка расплывается. Как понимать $\mathbf{x} \cdot \mathbf{y} = 0$? Лучше всего, как $\cos \varphi_{xy} = 0$, но это уже сказка про белого бычка.

На упомянутые вопросы легко ответить, особенно, если начинать с более точных определений абстрактного характера (см. раздел 3.5). В данном случае избран не самый экономный путь для того, чтобы лучше почувствовать мотивировку некоторых абстракций, освоение которых помогает геометрическому стилю мышления. Линейную алгебру, строго говоря, можно изучать вообще без геометрических представлений, как науку о прямоугольных таблицах чисел, — но тогда она больше походит на китайскую грамоту. Здесь предпочтение отдается ее трактовке как теории линейных операторов, действующих в евклидовом пространстве. Это требует некоторых дополнительных усилий на освоение исходных понятий, что окупается впоследствии наглядностью и другими удобствами.

2.3. Распознавание образов

Абстрактные высоты хороши при обеспеченных тылах. Поэтому время от времени имеет смысл возвращаться к практическим задачам, которые напоминают о секторе прицела.

В одной из возможных постановок задачи распознавания образов объекты характеризуются наборами параметров $\mathbf{x} = \{x_1, \dots, x_n\}$ и группируются в классы (образы), которым в R^n соответствуют некоторые множества точек.

Допустим, есть всего два класса объектов⁹⁾ и, соответственно, два множества $P_1, P_2 \in R^n$, которые для простоты будем считать *конечными*. Предположим, что существует плоскость, разделяющая P_1 и P_2 . Это означает существование такого вектора $\mathbf{c}$, что $\mathbf{c}\mathbf{x} > 0$, если $\mathbf{x} \in P_1$, и $\mathbf{c}\mathbf{x} < 0$, если $\mathbf{x} \in P_2$.

В процессе обучения линейной распознающей машине предъявляется бесконечная обучающая последовательность

$$\mathbf{x}^1, \dots, \mathbf{x}^k, \dots,$$

в которой любой объект из P_1 или P_2 встречается бесконечное число раз. Машина в качестве $\mathbf{c}$ принимает сначала произвольный вектор $\mathbf{c}^0$, и потом меняет

⁸⁾ Уступки в общности определений обычно выигрывают в наглядности, но проигрывают в широте охвата.

⁹⁾ Скажем, жук и бегемот, либо A и B , либо свой самолет и самолет противника.

его следующим образом. Если с помощью c^k точка x^k распознается правильно, то $c^{k+1} = c^k$. Если же $x^k \in P_1$, но $c^k x^k \leq 0$, то $c^{k+1} = c^k + x^k$. И $c^{k+1} = c^k - x^k$, если $x^k \in P_2$, но $c^k x^k \geq 0$. Вопрос заключается в том, научится ли машина за конечное число шагов безошибочно распознавать образы.

Многословие описания здесь связано с содержательной интерпретацией. Грамотный математический эквивалент выглядит проще. Существование разделяющей плоскости равносильно наличию такого вектора c , что

$$cx > 0 \quad \text{для любого} \quad x \in P = P_1 \cap (-P_2).$$

Процедура подстройки c^k :

$$c^{k+1} = \begin{cases} c^k, & \text{если } c^k x^k > 0, \\ c^k + x^k, & \text{если } c^k x^k \leq 0, \end{cases} \quad \text{все } x^k \in P. \quad (2.8)$$

◀ Покажем сходимость (2.8) за конечное число шагов. Пусть x^k обозначает сокращенную обучающую последовательность, из которой выброшены правильно распознаваемые «показы». Тогда

$$c^{k+1} = x^1 + \dots + x^k, \quad (2.9)$$

если для простоты положить $c^0 = 0$. Умножая (2.9) на c , получаем

$$c^{k+1} \cdot c = x^1 \cdot c + \dots + x^k \cdot c \geq k\theta,$$

где $\theta > 0$ — минимальное значение $x \cdot c$ при условии $x \in P$.

Опираясь далее на неравенство Коши—Буняковского, приходим к оценке

$$\|c^{k+1}\| \geq \frac{\theta}{\|c\|} k. \quad (2.10)$$

С другой стороны, возводя $c^{j+1} = c^j + x^j$ в квадрат, имеем

$$\|c^{j+1}\|^2 = \|c^j\|^2 + \|x^j\|^2 + 2c^j \cdot x^j,$$

что, в силу $c^j \cdot x^j \leq 0$, дает $\|c^{j+1}\|^2 - \|c^j\|^2 \leq \|x^j\|^2$. Суммируя эти неравенства по j , приходим к

$$\|c^{k+1}\|^2 \leq \sum_j^k \|x^j\|^2 \leq k\Theta^2 \Rightarrow \|c^{k+1}\| \leq \Theta\sqrt{k}, \quad (2.11)$$

где Θ — максимум $\|x^j\|$ при условии $x \in P$.

Сопоставление неравенств (2.11) и (2.10) показывает, что k не может расти неограниченно. ►

2.4. Линейные отображения и матрицы

Линейная функция

$$y = a \cdot x = a_1 x_1 + \dots + a_n x_n$$

принимает постоянные значения $a \cdot x = \beta$ на «плоскостях» параллельных плоскости $a \cdot x = 0$. Действительно, из $a \cdot u = \beta$, $a \cdot v = \beta$ следует $a \cdot (u - v) = 0$, т. е. вектор $u - v$ параллелен (лежит в) плоскости $a \cdot x = 0$.

Множества $a \cdot x = \beta$ называют *гиперплоскостями*¹⁰⁾, чтобы подчеркнуть, что они не являются линейными пространствами, которым положено вместе с векторами содержать их сумму. Другими словами, *гиперплоскость — это плоскость, не проходящая через начало координат.*

Важную роль в анализе играет понятие *линейного оператора*, сопоставляющего вектору $x = \{x_1, \dots, x_n\}$ вектор $y = \{y_1, \dots, y_n\}$ по правилу:

[illegible]

Таблицу коэффициентов $A = [a_{ij}]$,

$$A = \begin{bmatrix} a_{11} & a_{12} & \dots & a_{1n} \\ a_{21} & a_{22} & \dots & a_{2n} \\ \dots & \dots & \dots & \dots \\ a_{n1} & a_{n2} & \dots & a_{nn} \end{bmatrix},$$

называют *матрицей*, и коротко пишут $y = Ax$, определяя тем самым умножение матрицы на вектор.

Из приведенной записи видно, что первый индекс у a_{ij} обозначает номер строки, второй — столбца. Совокупность элементов $a_{11}, \dots, a_{nn}$ называют *главной диагональю* матрицы.

¹⁰⁾ А иногда — и просто плоскостями, допуская вольность терминологии.

Через $a_{i\cdot}$ обозначается i -я строка матрицы A , через $a_{\cdot j}$ — j -й столбец, т. е.

$$a_{i\cdot} = [a_{i1}, \dots, a_{in}], \quad a_{\cdot j} = \begin{bmatrix} a_{1j} \\ \vdots \\ a_{nj} \end{bmatrix}.$$

Двоеточие как бы заменяет весь список индексов.

Матрица называется *невыврожденной*, если ее вектор-строки линейно независимы.

Забегая вперед, отметим, что из линейной независимости вектор-строк вытекает линейная независимость вектор-столбцов, которую с равным успехом можно было бы взять за определение невырожденной матрицы.

Умножение матрицы на скаляр γ и сложение матриц $A = [a_{ij}]$ и $B = [b_{ij}]$ определяются так:

$$\gamma A = [\gamma a_{ij}], \quad A + B = [a_{ij} + b_{ij}].$$

При этом, очевидно,

$$A(\lambda x + \mu y) = \lambda Ax + \mu Ay.$$

Умножение A на B дает матрицу $C = AB$ с элементами

$$c_{ij} = \sum_{k=1}^n a_{ik} b_{kj}. \quad (2.13)$$

Иными словами, c_{ij} оказывается равным скалярному произведению i -й вектор-строки $a_{i\cdot}$ на j -й вектор-столбец $b_{\cdot j}$, т. е.

$$c_{ij} = a_{i\cdot} \cdot b_{\cdot j}.$$

Полезно также представление

$$c_{\cdot j} = A \cdot b_{\cdot j}, \quad (2.14)$$

которое подчеркивает, что j -й столбец матрицы $C = AB$ получается умножением j -го столбца матрицы B на матрицу A . Соотношение (2.14) означает

$$\begin{bmatrix} c_{1j} \\ \vdots \\ c_{nj} \end{bmatrix} = \begin{bmatrix} a_{11} \\ \vdots \\ a_{n1} \end{bmatrix} b_{1j} + \dots + \begin{bmatrix} a_{1n} \\ \vdots \\ a_{nn} \end{bmatrix} b_{nj},$$

т.е. каждый столбец матрицы C представляет собой линейную комбинацию столбцов A . Все это, конечно, тавтология, не прибавляющая содержания. Но такая перезапись (2.13) меняет форму восприятия и помогает укрупненно мыслить.

Правило умножения матриц возникает автоматически, если на матрицы смотреть как на линейные операторы. Вектор x под действием оператора B переходит в $z = Bx$, а вектор z под действием оператора A переходит в $y = Az$. Результирующее преобразование x в y определяется матрицей $C = AB$ с элементами c_{ij} , формула вычисления которых (2.13) определяется обыкновенным приведением подобных,

$$y_i = \sum_k a_{ik} \sum_j b_{kj} x_j = \sum_j \sum_k a_{ik} b_{kj} x_j = \sum_j c_{ij} x_j.$$

Произведение матриц в общем случае не коммутативно

$$AB \neq BA, \text{ но ассоциативно } A(BC) = (AB)C$$

и дистрибутивно $A(B + C) = AB + AC$, что легко проверяется.

В случае $AB = BA$ *говорят, что матрицы* A *и* B *коммутируют.*

Матрица A^* с элементами $a_{ij}^* = a_{ji}$ называется *транспонированной* к A , и получается как бы пространственным поворотом A вокруг главной диагонали на 180° , в результате чего столбцы и строки меняются местами. Матрицу называют *симметричной* (*симметрической*), если $A = A^*$, и — *кососимметричной*, если $A = -A^*$.

Для транспонированной матрицы широко распространены обозначения A' либо A^T , тогда как через A^* обозначают сопряженную матрицу, которая получается из A транспонированием с последующим комплексным сопряжением всех элементов a_{ij} . Например,

$$\begin{bmatrix} i & 3 \\ 2-i & 1+i \end{bmatrix}^* = \begin{bmatrix} -i & 2+i \\ 3 & 1-i \end{bmatrix}.$$

Таким образом, в случае матриц с вещественными элементами, какие пока имеются в виду, сопряженная матрица A^* — есть транспонированная. Поэтому использование обозначения A^* экономит символы «штрих» и «Т» для других целей.

При транспонировании произведения часто используется формула

$$(AB)^* = B^* A^*. \quad (2.15)$$

◀ Доказательство занимает одну строчку,

$$(AB)^* = \left[\sum_k a_{ik} b_{kj} \right]^* = \left[\sum_k a_{jk} b_{ki} \right] = \left[\sum_k b_{ik}^* a_{kj}^* \right] = B^* A^*. \quad \blacktriangleright$$

Обратной к A называют матрицу A^{-1} , такую что

$$A^{-1} A = A A^{-1} = I,$$

где I — единичная матрица¹¹⁾,

$$I = \begin{bmatrix} 1 & 0 & \dots & 0 \\ 0 & 1 & \dots & 0 \\ \dots & \dots & \dots & \dots \\ 0 & 0 & \dots & 1 \end{bmatrix},$$

определяющая тождественное преобразование $Ix \equiv x$. Очевидно,

$$IA = AI = A$$

для любой матрицы A .

¹¹⁾ Иногда используют обозначение $I = \text{diag} \{1, \dots, 1\}$.

Вопрос существования обратной матрицы рассматривается далее. Пока речь идет об определении, но уже из определения вытекают некоторые следствия. В частности, формула, аналогичная (2.15),

$$(AB)^{-1} = B^{-1}A^{-1}, \quad (2.16)$$

при условии, что матрицы A и B обратимы.

◀ После умножения (2.16) слева на B получается $B(AB)^{-1} = A^{-1}$, что после умножения слева на A переходит в очевидное $AB(AB)^{-1} = I$. Разумеется, эта цепочка обладает доказательностью в обратном направлении. Очевидное $AB(AB)^{-1} = I$ после умножения слева на A^{-1} , потом на B^{-1} , — дает (2.16)¹². ▶

Транспонируя $A^{-1}A = I$, получаем $A^*(A^{-1})^* = I$, откуда

$$(A^*)^{-1} = (A^{-1})^*,$$

что означает *перестановочность* операций транспонирования и обращения матрицы.

2.5. Прямоугольные и клеточные матрицы

Матрицу A , имеющую m строк и n столбцов, называют $m \times n$ матрицей. При переходе от квадратных матриц к прямоугольным кое-что теряет смысл, но многое — сохраняется. С некоторыми уточнениями, конечно. Например, произведение $C = AB$ вычисляется по той же формуле (2.13), но теперь при условии, что *число столбцов A равно числу строк B* . В противном случае AB не имеет смысла.

Понятие прямоугольной матрицы позволяет считать векторы тоже матрицами. При этом надо следить, какую «матрицу» представляет собой вектор — вектор-столбец или вектор-строку. Чтобы матрицу A умножить на x *справа*, надо, чтобы x был обязательно вектор-столбцом. После транспонирования вектор-столбца возможно x^*A . Скалярное произведение (x, y) записывают тогда в виде x^*y .

¹²) Обычно мы стараемся не засорять текст подобными уточнениями. Концентрация на идейной стороне дела требует жертв.

Для наглядности:

$$[x_1, \dots, x_n] \cdot \begin{bmatrix} y_1 \\ \vdots \\ y_n \end{bmatrix} = \sum_i x_i y_i,$$

$$\begin{bmatrix} x_1 \\ \vdots \\ x_n \end{bmatrix} \cdot [y_1, \dots, y_n] = \begin{bmatrix} x_1 y_1 & \dots & x_1 y_n \\ \dots & \dots & \dots \\ x_n y_1 & \dots & x_n y_n \end{bmatrix}.$$

При этом важно не увлекаться общностью, чтобы не потерять различие между векторами и матрицами, которое позволяет считать векторы точками пространства, а матрицы — линейными операторами. Возможность подобной интерпретации составляет главную силу линейной алгебры.

Заметим, наконец, что $y = Ax$ в случае прямоугольной $m \times n$ матрицы A можно интерпретировать как линейное преобразование m -мерного вектора x в n -мерный вектор y . Понимать это можно по-разному. Векторы x , y могут принадлежать совершенно разным пространствам, не имеющим друг к другу никакого отношения, а могут — быть векторами одного и того же «объемлющего» пространства. Но это связано не с природой задачи, а с волей, как говорят, лица, принимающего решение¹³⁾. Придание того или иного смысла координатам

$$\{x_1, \dots, x_m\}, \quad \{y_1, \dots, y_n\}$$

будет порождать различные геометрические интерпретации, предоставляя или упраздняя те или иные возможности.

Ортогональна, например, система координат в записи вектора $x = \{x_1, \dots, x_m\}$ или нет? Годится любой ответ. На усмотрение «заказчика». До поры до времени можно обходиться даже без ответа. Но если возникает, скажем, необходимость в скалярном произведении, то без декларации свойств системы координат уже не обойтись¹⁴⁾. Ортогональная — тогда скалярное произведение определяется формулой (2.4), нет — тогда нужна другая формула (см. далее).

¹³⁾ Другое дело, что разумная воля опирается на понимание задачи и действует в интересах эффективного решения.

¹⁴⁾ На практике часто обходятся без явных деклараций. Какая система координат имеется в виду — вытекает из контекста.

Блочное умножение. Матрицы могут состоять из прямоугольных блоков,

$$A = \begin{bmatrix} A_{11} & A_{12} & \dots & A_{1k} \\ A_{21} & A_{22} & \dots & A_{2k} \\ \dots & \dots & \dots & \dots \\ A_{q1} & A_{q2} & \dots & A_{qk} \end{bmatrix}, \quad B = \begin{bmatrix} B_{11} & B_{12} & \dots & B_{1s} \\ B_{21} & B_{22} & \dots & B_{2s} \\ \dots & \dots & \dots & \dots \\ B_{41} & B_{42} & \dots & B_{ks} \end{bmatrix},$$

где каждая из клеток A_{ij} , B_{ij} представляет собой прямоугольную матрицу. Если размеры клеток согласованы, а именно: A_{ij} имеет размер $m_i \times n_j$, а B_{ij} — $n_i \times r_j$, то матрицы A и B можно умножать прямо в блочном виде по «обычному» правилу

$$C_{it} = \sum_j A_{ij} B_{jt},$$

где блоки C_{it} получаются размера $m_i \times r_t$, и они образуют матрицу $C = AB$.

Если A и B квадратные матрицы, то

$$C = \begin{bmatrix} A & B \\ B & -A \end{bmatrix} \Rightarrow C^2 = \begin{bmatrix} A^2 + B^2 & AB - BA \\ BA - AB & A^2 + B^2 \end{bmatrix}.$$

2.6. Два примера

Матрицей перестановки называют матрицу, у которой в каждом столбце и каждой строке — ровно один элемент равен 1, остальные нули. Умножение на такую матрицу слева — переставляет строки, справа — столбцы. Например,

$$P = \begin{bmatrix} 1 & 0 & 0 \\ 0 & 0 & 1 \\ 0 & 1 & 0 \end{bmatrix} \Rightarrow P \begin{bmatrix} a_1 \\ a_2 \\ a_3 \end{bmatrix} = \begin{bmatrix} a_1 \\ a_3 \\ a_2 \end{bmatrix}.$$

Матрицы перестановок, вообще говоря, не коммутируют, но их произведение остается матрицей перестановки.

У **треугольных** матриц ненулевые элементы стоят исключительно выше (ниже) главной диагонали:

$$\begin{bmatrix} t_{11} & t_{12} & \dots & t_{1n} \\ 0 & t_{22} & \dots & t_{2n} \\ \dots & \dots & \dots & \dots \\ 0 & \dots & 0 & t_{nn} \end{bmatrix}, \quad \begin{bmatrix} t_{11} & 0 & \dots & 0 \\ t_{21} & t_{22} & \dots & \dots \\ \dots & \dots & \dots & 0 \\ t_{n1} & t_{n2} & \dots & t_{nn} \end{bmatrix}.$$

Например,

$$\begin{bmatrix} 1 & 0 & 0 \\ 0 & \lambda & 0 \\ 0 & 0 & 1 \end{bmatrix} \cdot \begin{bmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \\ a_{31} & a_{32} & a_{33} \end{bmatrix} = \begin{bmatrix} a_{11} & a_{12} & a_{13} \\ \lambda a_{21} & \lambda a_{22} & \lambda a_{23} \\ a_{31} & a_{32} & a_{33} \end{bmatrix},$$

$$\begin{bmatrix} 1 & 0 & 1 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{bmatrix} \cdot \begin{bmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \\ a_{31} & a_{32} & a_{33} \end{bmatrix} = \begin{bmatrix} a_{11} + a_{31} & a_{12} + a_{32} & a_{13} + a_{33} \\ a_{21} & a_{22} & a_{23} \\ a_{31} & a_{32} & a_{33} \end{bmatrix}.$$

К элементарным преобразованиям (матрицам) относят также преобразования, получаемые последовательным применением операций (2.18). Другими словами, если все Λ_i описывают одну из двух операций (2.18), то матрица

$$\Lambda = \Lambda_1 \cdots \Lambda_N$$

считается элементарной. К элементарным — относятся, например, *любые перестановки строк*, которые всегда могут быть получены в результате последовательного применения перестановок двух строк (*транспозиций*). Транспозиции, в свою очередь, сводятся к последовательности операций вида (2.18) по схеме:

$$\begin{bmatrix} a \\ b \end{bmatrix} \Rightarrow \begin{bmatrix} a+b \\ b \end{bmatrix} \Rightarrow \begin{bmatrix} a+b \\ b-a-b \end{bmatrix} \Rightarrow \begin{bmatrix} a+b \\ -a \end{bmatrix} \Rightarrow \begin{bmatrix} b \\ -a \end{bmatrix} \Rightarrow \begin{bmatrix} b \\ a \end{bmatrix}. \quad (2.19)$$

Любое элементарное преобразование обратимо. Это важный, хотя и простой факт. Действительно, если первое из преобразований (2.18) умножает строку на λ , то обратное — получается умножением строки на $1/\lambda$. Обратным ко второму из преобразований (2.18) будет вычитание, т. е. прибавление строки после ее предварительного умножения на -1 .

Элементарные преобразования занимают в линейной алгебре скромное место, но в самом начале пути позволяют легко установить основополагающие факты.

Метод Гаусса. Принципиально следующее утверждение.

2.7.1. *Любая невырожденная матрица элементарным преобразованием, т. е. последовательностью операций (2.18), может быть преобразована в единичную.*

◀ Из невырожденности матрицы A следует, что у нее в первой строке, заведомо, есть ненулевой элемент¹⁶⁾ a_{1j_1} . Разделим строку a_1 на a_{1j_1} , после чего умножим a_1 на a_{1j_1} и вычтем из i -й строки. Прделаав такую операцию для всех

¹⁶⁾ Иначе вектор-строки были бы линейно зависимы.

$i = 2, \dots, n$, получим в итоге матрицу A^{new} , у которой в j_1 -м столбце стоят везде нули, за исключением $a_{1j_1} = 1$.

Матрица A^{new} невырождена¹⁷⁾, поэтому во второй строке у нее есть ненулевой элемент $a_{2j_2}^{new}$. Далее, по аналогии с предыдущим, делим строку a_2^{new} на $a_{2j_2}^{new}$, затем умножаем вновь полученную строку на $a_{ij_2}^{new}$ и вычитаем из i -й строки. Прделав это для всех $i \neq 2$, получим в итоге матрицу, у которой в j_2 -м столбце стоят везде нули, за исключением $a_{2j_2} = 1$. Заметим, что в результате второго этапа преобразования j_1 -й столбец остается прежним (везде нули, за исключением 1 на $\{1, j_1\}$ -м месте).

Продолжая далее в том же духе, в конце концов получим матрицу, у которой в каждой строке и в каждом столбце всего одна единица — остальные нули. Такая матрица перестановкой строк легко преобразуется в единичную¹⁸⁾. ►

Только что описанный процесс называют *методом Гаусса*. В применении к сопутствующей задаче решения системы уравнений $Ax = b$ метод дает решение x' ($Ax' = b$), при условии, что умножению и сложению строк A соответствует умножение и сложение уравнений $Ax = b$. В итоге вектор b преобразуется в x' .

Если матрица A вырождена, то для решения $Ax = b$ можно применять *модифицированный метод Гаусса*, в котором на каждом шаге переход осуществляется не обязательно к следующей строке, а к первой ненулевой. В итоге часть строк A в результате преобразований станут нулевыми. Если при этом соответствующие координаты вектора b тоже преобразуются в нулевые, то система $Ax = b$, очевидно, решается с запасом¹⁹⁾. Если нет²⁰⁾ — система $Ax = b$ не имеет решения.

Обратная матрица. Утверждение 2.7.1 означает наличие такой последовательности элементарных матриц $\Lambda_1, \dots, \Lambda_N$, что

$$\Lambda_N \cdots \Lambda_1 A = I,$$

откуда следует существование у невырожденной матрицы A обратной,

$$A^{-1} = \Lambda_N \cdots \Lambda_1. \quad (2.20)$$

¹⁷⁾ Поскольку элементарные операции переводят линейно независимые векторы в линейно независимые (проверьте).

¹⁸⁾ Если в j -м столбце $a_{ij} = 1$, то i_j -ю строку переставляем на j -е место.

¹⁹⁾ Часть переменных можно задать произвольно.

²⁰⁾ Т.е. среди соответствующих b_j есть ненулевые.

Точнее говоря, следует существование «левой обратной матрицы» Λ ($\Lambda A = I$). Но элементарная матрица $\Lambda = \Lambda_N \cdots \Lambda_1$ обратима (см. выше). Поэтому, умножая $\Lambda A = I$ слева на Λ^{-1} , получаем $A = \Lambda^{-1}$, что после умножения справа на Λ , дает $\Lambda A = I$, гарантируя тем самым (2.20). ►

Таким образом,

$$\Lambda A = I \Rightarrow A = \Lambda^{-1},$$

а поскольку Λ^{-1} элементарная матрица, то можно утверждать, что *любая невырожденная матрица A представима в виде произведения элементарных матриц* из категории (2.18). Равно как и наоборот, если A представима в виде такого произведения, то она невырождена, что ясно из представления $A = \Lambda^{-1}I$, поскольку здесь невырожденная матрица I элементарным преобразованием переводится в невырожденную.

Упражнения

- Элементарные преобразования переводят линейно зависимые строки в линейно зависимые.
- Вырожденная матрица не имеет обратной.

2.7.2. Лемма. При элементарных преобразованиях **строк** линейная независимость (зависимость) **столбцов** сохраняется.

◀ Столбцами матрицы ΛA являются векторы $\Lambda a_{:,j}$. В силу обратимости Λ , очевидно,

$$\alpha_1 a_{:,1} + \dots + \alpha_n a_{:,n} = (\neq) 0 \Leftrightarrow \alpha_1 \Lambda a_{:,1} + \dots + \alpha_n \Lambda a_{:,n} = (\neq) 0. \quad \blacktriangleright$$

Из леммы 2.7.2 сразу вытекает, что у вырожденной (невырожденной) матрицы вектор-столбцы линейно зависимы (независимы). Другими словами, A и транспонированная матрица A^* вырождены или невырождены одновременно.

Отметим, наконец, что в определении элементарных преобразований строки можно заменить столбцами, не меняя выводов по существу.

Ранг матрицы. Пусть теперь речь идет о прямоугольных матрицах $m \times n$ (m строк, n столбцов).

2.7.3. Определение. Максимальное число линейно независимых вектор-строк либо вектор-столбцов матрицы A называется ее **рангом** и обозначается как $\text{rank } A$.

В данном определении «заложена теорема», поскольку совпадение максимального числа линейно независимых строк и столбцов сразу неочевидно, но это так.

◀ Пусть r обозначает максимальное число линейно независимых строк матрицы A , s — максимальное число линейно независимых столбцов. Допустим, $r > s$. Оставим эти строки, уберем остальные, и применим к оставшейся матрице метод Гаусса. Выше он был описан для квадратной матрицы, но в «прямоугольном» случае почти ничего не меняется. Точно так же в итоге элементарных преобразований (2.18) получатся r столбцов вида

$$\{0, \dots, 0, 1, 0 \dots, 0\}^*,$$

у которых единицы стоят в разных строках²¹⁾.

Это и есть линейно независимые столбцы. Исходные линейно независимые столбцы имеют те же номера по лемме 2.7.2, которая справедлива и в «прямоугольном» случае. ►

Теперь понятно, что ранг матрицы можно было бы определить еще одним эквивалентным способом. *Ранг $m \times n$ матрицы — это максимальный размер ее квадратной невырожденной подматрицы*²²⁾.

Фактическое определение ранга возможно *модифицированным методом Гаусса*. На каждом шаге переход осуществляется не обязательно к следующей строке, а к первой ненулевой. По завершении процесса число ненулевых строк дает ранг.

Довольно часто требуется оценка ранга произведения AB . В большинстве случаев можно обойтись неравенством

$$\text{rank } AB \leq \min\{\text{rank } A, \text{rank } B\}. \quad (2.21)$$

◀ Поскольку столбцы матрицы AB являются линейными комбинациями столбцов A , то ранг матрицы $[AB, A]$ ²³⁾ не может быть больше ранга A . В свою очередь, ранг AB не больше ранга $[AB, A]$. Поэтому

$$\text{rank } AB \leq \text{rank } A.$$

Неравенство $\text{rank } AB \leq \text{rank } B$ устанавливается аналогично. ►

²¹⁾ В отличие от ситуации квадратной матрицы здесь будут «лишние» столбцы, но они не мешают, так как для противоречия нам достаточно показать $s \geq r$.

²²⁾ **Подматрицей** матрицы A называют матрицу, элементы которой стоят на пересечении выбранных строк и выбранных столбцов матрицы A . Если номера строк и столбцов одинаковы — подматрицу называют *главной*.

²³⁾ Полученной приписыванием справа к AB матрицы A .

Если A и B квадратные матрицы, то из (2.21) вытекает, что *произведение любой матрицы на вырожденную — дает вырожденную матрицу*.

Упражнение

Если P невырожденная матрица, то

$$\text{rank } PQ = \text{rank } Q, \quad \text{rank } SP = \text{rank } S.$$

2.8. Теория определителей

Определители, или детерминанты, иногда воспринимаются как «ложка дегтя», которая омрачает настроение при изучении линейной алгебры. С одной стороны, все достаточно просто в любом варианте *формального* изложения теории. С другой — не все «почему» находят простые и ясные ответы, и это оставляет чувство неудовлетворенности.

Вот один из формальных вариантов. *Определителем* (или *детерминантом*) матрицы A называется величина

$$|A| = \sum_{\sigma} (-1)^{N(\sigma)} a_{\sigma(1)1} \cdots a_{\sigma(n)n}, \quad (2.22)$$

где σ обозначает *перестановку* (отображение) множества индексов

$$\{1, \dots, n\} \text{ в } \{\sigma(1), \dots, \sigma(n)\},$$

$N(\sigma)$ — число *транспозиций*²⁴⁾ в σ , а суммирование в (2.22) идет по всем возможным перестановкам.

Иными словами, определитель представляет собой алгебраическую сумму *н* невозможных произведений $a_{\sigma(1)1} \cdots a_{\sigma(n)n}$, знак которых определяется четностью или нечетностью перестановки σ (четностью или нечетностью числа транспозиций в σ). Вычисление $N(\sigma)$ представляется обременительным, но его можно заменить рутинным определением числа *беспорядков*²⁵⁾ в σ .

Кстати, последовательность элементов $a_{\sigma(1)1}, \dots, a_{\sigma(n)n}$ называют *диагональю матрицы A* , соответствующей перестановке σ . В случае тождественной перестановки $\sigma(j) = j$ диагональ $a_{11}, \dots, a_{nn}$ называют *главной*. Если сказать иначе, то

²⁴⁾ Любая перестановка может быть представлена в виде последовательно выполненных транспозиций (перестановок двух элементов) — см. [11].

²⁵⁾ Число беспорядков в σ — это сумма беспорядков для каждого числа $\sigma(i)$, равных количеству чисел меньших $\sigma(i)$, но стоящих после $\sigma(i)$.

диагональ матрицы, это совокупность n ее элементов, таких что никакие два не стоят в одной строке или в одном столбце.

Из определения (2.22) все свойства детерминантов достаточно легко выводятся (см., например, [11]), и при определенном устройстве мозгов проблема оказывается исчерпанной. Но кому-то мешает ощущение фокуса, и тогда приходится заходить с другой стороны.

В основу определения детерминанта может быть положено понятие ориентированного объема (по аналогии с трехмерным вариантом). Геометрическая фантазия в этом направлении может быть такой. *Детерминант (определитель) $|A|$* — равносильное обозначение $\det A$ — это объем параллелепипеда, построенного на вектор-строках матрицы A , — взятый со знаком плюс (минус), если набор

вектор строк $\begin{bmatrix} a_{1:} \\ \vdots \\ a_{n:} \end{bmatrix}$ одинаково (противоположно) ориентирован с базисом.

Пока такое определение несостоятельно, поскольку в наличии нет определения объема параллелепипеда и ориентации упорядоченных наборов векторов в R^n . В подобных ситуациях часто спасает принцип «мы не знаем, что это такое, но „оно“ должно обладать такими-то свойствами». Самое удивительное, что обычно два-три очевидных свойства однозначно характеризуют объект.

2.8.1. Определение. *Детерминантом квадратной матрицы A назовем скалярную функцию $\det A$, обладающую тремя свойствами*

- (i) *при умножении строки матрицы на скаляр λ величина $\det A$ умножается на λ ;*
- (ii) *прибавление одной строки к другой не меняет детерминанта;*
- (iii) *детерминант единичной матрицы $\det I = 1$.*

Сразу, естественно, возникают вопросы. Существует ли такая функция? Если да, то единственна ли и как ее вычислять?

Все ответы благоприятные, но сначала удобнее из (i)–(iii) вывести некоторые следствия.

1. Если матрица A имеет нулевую строку, то $\det A = 0$, что сразу вытекает из (i).
2. Прибавление одной строки, умноженной на любой скаляр, к другой строке — не меняет детерминанта. Это чуть более общее свойство, чем (ii). Факт легко

выводится. ◀ Если записать $A = \begin{bmatrix} a \\ b \end{bmatrix}$, выделяя только строчки a и b , с которыми что-то происходит, то

$$\det \begin{bmatrix} a \\ b \end{bmatrix} = \frac{1}{\lambda} \det \begin{bmatrix} \lambda a \\ b \end{bmatrix} = \frac{1}{\lambda} \det \begin{bmatrix} \lambda a \\ b + \lambda a \end{bmatrix} = \det \begin{bmatrix} a \\ b + \lambda a \end{bmatrix}. \quad \blacktriangleright$$

3. Прибавление к строке линейной комбинации других — не меняет определителя. ◀ Операция сводится к последовательности операций из предыдущего пункта. ▶
4. Детерминант вырожденной матрицы равен нулю. ◀ При линейной зависимости строк прибавление к одной из строк некоторой линейной комбинации других — дает нулевую строку. ▶
5. При перестановке двух строк знак определителя меняется на противоположный. ◀ Преобразования по схеме (2.19) меняют значение определителя только на шаге умножения строки на -1 . ▶

Наконец, последнее и самое важное. ◀ Если A невырождена, то применение к ней метода Гаусса превращает ее в единичную — с детерминантом равным 1, в силу (iii). В процессе преобразований $\det A$ меняется n раз. При делении первой строки на a_{1j_1} , второй — на $a_{2j_2}^{new}$, и т. д. (см. раздел 2.7). Поэтому, в силу (i),

$$\det A = a_{1j_1} \cdot a_{2j_2}^{new} \cdots, \quad (2.23)$$

что не только устанавливает существование и единственность $\det A$, но и дает удобный алгоритм его вычисления. ▶

Итак, детерминант определен, причем *однозначно*. Ориентированный ли это объем — теперь, вообще говоря, не играет роли. Более того, первоначальный замысел определить $\det A$ как объем²⁶⁾ — можно вывернуть наизнанку, определив объем как детерминант. Для этого надо лишь убедиться в естественности свойств (i)–(iii) для объема ориентированного параллелепипеда.

Свойство (iii) есть равенство единице объема прямоугольного n -мерного куба со стороной 1. При увеличении одной из сторон в λ раз объем увеличивается тоже в λ раз — это (i). Обратим внимание, что при $\lambda < 0$ знак объема меняется на противоположный — и это как раз тот минимум требований к ориентации, который можно высказать, не понимая, что она (ориентация) есть такое. Наконец, (ii) — следствие линейности объема по ребру: если ребро a равно сумме $b + c$, то объем равен сумме объемов двух параллелепипедов, один из которых вместо a имеет ребро b , другой — c . Если при этом $c = a$, а b другое ребро, то «другой»

²⁶⁾ Понятно, замысел был чисто эмоционален, ибо не было ясно (и не ясно до сих пор), что такое объем в R^n .

параллелепипед будет иметь нулевой объем, поскольку у него два одинаковых (векторных) ребра. Это приводит к (ii).

Разумеется, все свойства (i)–(iii) выполняются при $n = 2, 3$, где для площади и объема есть другие идеологические платформы. Условие (ii), например, для площади параллелограмма $a \times b$ есть очевидное

$$(a + b) \times b = a \times b + b \times b = a \times b.$$

Для ориентированного объема параллелепипеда

$$a \cdot (b \times c) = \det \begin{bmatrix} a \\ b \\ c \end{bmatrix}$$

справедливость (ii) вытекает из очевидных равенств типа

$$a \cdot (b \times c) = (a + b) \cdot (b \times c), \quad a \cdot (b \times c) = a \cdot [(b + a) \times c].$$

В результате получено важное (за пределами линейной алгебры) определение ориентации n -мерного пространства. Базис из упорядоченного набора векторов $\{a, \dots\}$ ориентирован так же, как исходный, если

$$\det \begin{bmatrix} a \\ \vdots \\ \vdots \end{bmatrix} > 0, \text{ и — противоположно, если детерминант } < 0.$$

На $\det A$ полезно взглянуть еще как на функцию линейного преобразования A . Единичный куб, построенный на векторах (ребрах) стандартного базиса (2.7) под действием оператора $A = [a_{ij}]$ переходит в параллелепипед, построенный на вектор-столбцах $a_{\cdot 1}, \dots, a_{\cdot n}$. Объем этого параллелепипеда, равный по модулю $|\det A|$, естественно считать коэффициентом искажения объема линейного оператора A .

С тем же коэффициентом искажения объема оператор A преобразует любое тело V . Разбивая V на N мелких кубиков (со стороной ε и объемом ε^n), объем V можно сколь угодно точно приблизить величиной $N \cdot \varepsilon^n$. Каждый кубик под действием A переходит в параллелепипед объема $|\det A| \cdot \varepsilon^n$, что и дает необходимое заключение.

Теперь вернемся к альтернативному определению детерминанта (2.22). Совпадает ли оно по сути с определением 2.8.1? Совпадает.

Поскольку легко проверяется, что (2.22) удовлетворяет свойствам (i)–(iii), а в этом случае детерминант однозначно вычисляется методом Гаусса — см. (2.23).

Из (2.22), в свою очередь, сразу вытекают некоторые свойства детерминантов, которые проверять иначе несколько обременительно. Например,

$$\det A = \det A^*, \quad (2.24)$$

т. е. транспонирование не меняет определителя матрицы.

Еще одно легко выводимое следствие — *разложение определителя по элементам столбца (строки)*:

$$\det A = \sum_i (-1)^{i+j} a_{ij} M_{ij} \quad \left(\det A = \sum_j (-1)^{i+j} a_{ij} M_{ij} \right),$$

где M_{ij} обозначает *дополнительный минор*²⁷⁾ элемента a_{ij} .

Такое разложение в большей степени является теоретическим инструментом. В принципе оно сводит вычисление детерминанта n -го порядка к вычислению детерминантов $(n-1)$ -го порядка, но при больших n вычислять лучше по более простым схемам (методом Гаусса, например).

Еще одна часто используемая формула:

$$\det AB = \det A \det B, \quad (2.25)$$

определитель произведения матриц равен произведению определителей.

В частности,

$$\det I = \det AA^{-1} = \det A \det A^{-1},$$

откуда $\det A^{-1} = (\det A)^{-1}.$

◀ Формула (2.25) легко устанавливается разными способами. В курсе [11] это делается с помощью изящного фокуса²⁸⁾.

²⁷⁾ Дополнительный минор M_{ij} элемента a_{ij} есть определитель подматрицы, которая получается из A вычеркиванием i -й строки и j -го столбца.

²⁸⁾ Фокусы плохи тем, что непонятны, но хороши тем, что инициируют поиск.

Для конструктивного понимания полезна опора на понятие коэффициента искажения объема. Если A искажает объем в α раз, B — в β раз, то ясно, что последовательное применение A и B в любом порядке меняет объем в $\alpha\beta$ раз — и это справедливо с учетом изменения ориентации.

Простейший вариант доказательства основывается на представлении обеих матриц A и B в виде произведения элементарных матриц, для которых равенство

$$\det(\Lambda_1 \cdots \Lambda_N) = \det \Lambda_1 \cdots \det \Lambda_N$$

очевидно. Если же одна из матриц A , B вырождена, то (2.25) вытекает из (2.21). ►

В заключение обратим внимание, что при определении детерминанта вместо строк можно было бы взять столбцы, ничего не меняя по сути²⁹⁾. В итоге получились бы те же самые результаты, с заменой в формулировках строк на столбцы.

2.9. Системы уравнений

Системы $n \times n$. Система линейных уравнений $Ax = b$ с квадратной невырожденной матрицей A всегда имеет единственное решение $x = A^{-1}b$.

Результат $x = A^{-1}b$ можно интерпретировать и записывать по-разному. Перезапись $Ax = b$ в виде

$$\begin{bmatrix} a_{11} \\ \vdots \\ a_{n1} \end{bmatrix} x_1 + \dots + \begin{bmatrix} a_{1n} \\ \vdots \\ a_{nn} \end{bmatrix} x_n = b \quad (2.26)$$

показывает, что решение $x = \{x_1, \dots, x_n\}^*$ представляет собой координаты вектора b в базисе $\{a_{:1}, \dots, a_{:n}\}$ из вектор-столбцов матрицы A .

Определение таких координат относительно просто. Пусть $A^{j\downarrow b}$ обозначает матрицу, полученную из A заменой j -го столбца вектором b . С учетом представления (2.26) —

$$\det A^{j\downarrow b} = x_j \det A, \quad (2.27)$$

поскольку $A^{j\downarrow b}$ получается из A умножением j -го столбца A на x_j — с последующим прибавлением линейной комбинации

²⁹⁾ Что сразу следует из (2.24).

других столбцов A , что уже не играет роли (не меняет детерминанта)³⁰⁾. Из (2.27) вытекает *правило Крамера*:

$$\boxed{x_j = \frac{\det A^{j \downarrow b}}{\det A}}, \quad j = 1, \dots, n,$$

дающее в случае невырожденной матрицы A решение системы $Ax = b$.

Еще одна идея решения $Ax = b$ заключается в использовании *принципа суперпозиции*. Если векторы

$$b_{:,1}, \dots, b_{:,n} \quad (2.28)$$

линейно независимы и $Ax_{:,j} = b_{:,j}$, то разложение $b = \sum \lambda_j b_{:,j}$ вектора b по базису (2.28) сразу дает решение $x = \sum \lambda_j x_{:,j}$ системы $Ax = b$. Иными словами, знание n решений $x_{:,j}$ позволяет «легко» выписывать решение $Ax = b$ при любом b .

Конечно, легко сказка сказывается. Такой путь может оказаться длиннее, поскольку λ_j надо определять, решая систему $B\lambda = b$, где матрица B образована столбцами (2.28). Однако при выборе в качестве (2.28) базиса

$$e_1 = \begin{bmatrix} 1 \\ 0 \\ \vdots \\ 0 \end{bmatrix}, \quad e_2 = \begin{bmatrix} 0 \\ 1 \\ \vdots \\ 0 \end{bmatrix}, \quad \dots, \quad e_n = \begin{bmatrix} 0 \\ \vdots \\ 0 \\ 1 \end{bmatrix}$$

матрица $X = [x_{:,1}, \dots, x_{:,n}]$ оказывается обратной к A , т. е. $X = A^{-1}$.

Вырожденный случай. Если $n \times n$ матрица A вырождена, то применение к системе $Ax = b$ модифицированного метода Гаусса приводит ее к виду

$$\left[\begin{array}{c|c} I & 0 \\ \hline 0 & 0 \end{array} \right] x = b^{new},$$

где I — единичная матрица размерности $\leq n-1$, остальные блоки нулевые. При этом ясно, что система имеет решение, если у преобразованного вектора b^{new} координаты напротив нулевых строчек матрицы — нулевые. В противном случае исходная система не имеет решения.

³⁰⁾ См. раздел 2.8.

Упражнение

Если ранг $n \times n$ матрицы A меньше n , то однородное уравнение $Ax = 0$ имеет ненулевое решение x .

2.10. Задачи и дополнения

- В тех случаях, когда коэффициенты матрицы $A(t)$ зависят от времени (параметра), могут использоваться естественные операции поэлементного дифференцирования и интегрирования,

$$\dot{A}(t) = [\dot{a}_{ij}(t)], \quad \int A(t) dt = \left[\int a_{ij}(t) dt \right].$$

При этом очевидны стандартные свойства:

$$\frac{d}{dt}[A + B] = \dot{A} + \dot{B}, \quad \frac{d}{dt}[AB] = \dot{A}B + A\dot{B}.$$

Дифференцирование тождества $A(t)A^{-1} \equiv I$ приводит к

$$\frac{dA}{dt}A^{-1} + A\frac{dA^{-1}}{dt} \equiv 0,$$

что означает

$$\frac{dA^{-1}}{dt} = -A^{-1}\frac{dA}{dt}A^{-1}.$$

- Обратная матрица A^{-1} имеет элементами $(-1)^{i+j}M_{ij}/\det A$, где M_{ij} дополнительные миноры элементов a_{ij} матрицы A .
- $\det A = (-1)^{i+j}a_{ij}M_{ij} + \det \tilde{A}_{ij}$, где матрица $\tilde{A}_{ij}$ получается заменой в A элемента a_{ij} нулем.
- Матрица Вандермонда

$$A = [a_{ij}] = [\lambda_j^{i-1}]$$

имеет определитель

$$\det A = \prod_{j>i} (\lambda_j - \lambda_i).$$

- Если матрица A содержит нулевую подматрицу $p \times q$, и $p + q \geq n + 1$, то A вырождена. (?)
- Если $\text{rank} A = 1$, то найдутся такие векторы u и v , что $A = uv^*$.
- Матрица ранга k всегда есть сумма k матриц ранга 1.

Глава 3

Линейные преобразования

3.1. Замена координат

Изучение матриц в отрыве от геометрических представлений сильно меняет краски, и взгляд на матрицу как на линейный оператор начинает казаться если не откровением, то изобретательным ходом. На самом деле понятие линейного оператора «первичнее». Естественное изучение функций нескольких переменных, $y = f(x)$, в первую очередь обращает внимание на простейший тип линейных преобразований — характеризуемых свойством $f(\alpha x + \beta y) = \alpha f(x) + \beta f(y)$, — каковыми и оказываются¹⁾ $f(x) = Ax$.

При изучении же функций эффективным инструментом служит переход к другой (более выгодной) системе координат. Например, перенос начала координат преобразует $y = x^2 + \beta x + \gamma$ в $u = v^2$, а вращение осей координат сильно упрощает запись кривых второго порядка (глава 1). Аналогичные выгоды можно извлечь и в n -мерном случае.

При переходе к другой (штрихованной) системе координат с помощью невырожденной матрицы T :

$$x = Tx',$$

соотношение $u = Av$ после подстановки $u = Tu'$, $v = Tv'$ превращается в $Tu' = ATv'$, что в результате умножения слева на T^{-1} переходит в

$$u' = T^{-1}ATv',$$

на основании чего можно утверждать, что в новой системе координат линейному оператору A соответствует матрица

$$\boxed{A' = T^{-1}AT.} \quad (3.1)$$

¹⁾ Без дополнительного требования непрерывности линейные функции могут быть разрывными — см., например, [3].

Пусть имеется два базиса,

и f через e выражается линейно:

$$\begin{cases} \mathbf{f}_1 = s_{11}\mathbf{e}_1 + \dots + s_{1n}\mathbf{e}_n, \\ \dots\dots\dots \\ \mathbf{f}_n = s_{n1}\mathbf{e}_1 + \dots + s_{nn}\mathbf{e}_n. \end{cases} \quad (3.2)$$

Допустим теперь, что x_i обозначает координату точки x в базисе $\{e\}$, а x'_i — в базисе $\{f\}$. Подставляя f_i из (3.2) в

и приравнявая коэффициенты при ϵ_i , — после раскрытия скобок и приведения подобных, получим

$$\begin{cases} x_1 = s_{11}x'_1 + \dots + s_{n1}x'_n, \\ \dots\dots\dots \\ x_n = s_{1n}x'_1 + \dots + s_{nn}x'_n, \end{cases}$$

Замечание. Далее будут рассматриваться *ортогональные матрицы*, которые осуществляют поворот векторов, не меняя их длины. Для таких матриц S транспонированная является обратной, $S^* = S^{-1}$. В этом случае получается

$$T = (S^*)^{-1} = (S^{-1})^{-1} = S.$$

Полиномы. Если $A' = T^{-1}AT$, то $B = A^2$ в новой системе координат:

$$B' = (A')^2 = T^{-1}AT \cdot T^{-1}AT = T^{-1}A^2T.$$

То же самое можно сказать по поводу любой другой степени A^k ($k = 2, 3, \dots$), а также по поводу любого матричного полинома

$$p(A) = A^k + \gamma_{k-1}A^{k-1} + \dots + \gamma_0I,$$

²⁾ Избегать трудностей — другая игра, которая лишь выглядит более разумной.

т. е.

$$p(T^{-1}AT) = T^{-1}p(A)T.$$

Наконец, $(T^{-1}AT)^{-1} = T^{-1}A^{-1}T$, т. е. $(A')^{-1} = (A^{-1})'$.

3.2. Собственные значения и комплексные пространства

Обнадеживающе выглядит идея, исходя из (3.1), преобразовать матрицу так, чтобы она приобрела в новой системе координат наиболее простой вид. Одна из ориентаций при этом может быть на *матрицы диагонального вида*:

$$\Lambda = \begin{bmatrix} \lambda_1 & & 0 \\ & \ddots & \\ 0 & & \lambda_n \end{bmatrix} = \text{diag} \{ \lambda_1, \dots, \lambda_n \}, \quad (3.3)$$

у которых ненулевые элементы могут стоять только на главной диагонали.

Другими словами, соблазнительно задаться целью найти такое преобразование T , чтобы матрица A в новой системе координат стала диагональной, т. е.

$$A' = T^{-1}AT = \Lambda. \quad (3.4)$$

Матрицы A' и A в соотношении $A' = T^{-1}AT$ называют подобными. Это различные матрицы, но они описывают один и тот же оператор (в разных системах координат). Здесь начинается ощущаться разница между матрицей и оператором.

Понятно, что в случае (3.4) оператор A на векторы нового базиса $\{f'_1, \dots, f'_n\}$ действует предельно просто:

$$A'f'_j = \lambda_j f'_j. \quad (3.5)$$

Умножая (3.5) слева на T , справа на T^{-1} , получаем

$$TA'T^{-1} \cdot Tf'_jT^{-1} = \lambda_j Tf'_jT^{-1},$$

т. е. $Af_j = \lambda_j f_j$, где $f_j = Tf'_jT^{-1}$.

Таким образом, для определения чисел λ_j и векторов $\mathbf{f}_j$ надо решить уравнение

$$\boxed{A\mathbf{x} = \lambda\mathbf{x}}, \quad (3.6)$$

которое в эквивалентной записи выглядит как $(A - \lambda I)\mathbf{x} = 0$ и может иметь ненулевое решение $\mathbf{x}$ лишь в том случае, если матрица $A - \lambda I$ вырождена, т. е.

$$\boxed{\det(A - \lambda I) = 0}. \quad (3.7)$$

Уравнение (3.7) называют *характеристическим уравнением* матрицы A , его решения λ_j — *собственными значениями*, а соответствующие ненулевые решения $\mathbf{x} = \mathbf{f}_j$ уравнения (3.6) — *собственными векторами*³⁾ матрицы A .

3.2.1. Если матрица A имеет собственные значения λ_j , которым отвечают собственные векторы $\mathbf{f}_j$, то матричный полином

$$p(A) = A^n + \gamma_{n-1}A^{n-1} + \dots + \gamma_0 I$$

имеет собственные значения $p(\lambda_j)$ и те же собственные векторы $\mathbf{f}_j$.

◀ Применяя к $A\mathbf{f}_j = \lambda_j\mathbf{f}_j$ оператор A , получаем

$$A^2\mathbf{f}_j = \lambda_j A\mathbf{f}_j = \lambda_j^2\mathbf{f}_j.$$

Продолжая аналогично, имеем

$$A^k\mathbf{f}_j = \lambda_j^k\mathbf{f}_j, \quad (\alpha A^p + \beta A^q)\mathbf{f}_j = (\alpha\lambda_j^p + \beta\lambda_j^q)\mathbf{f}_j,$$

что в итоге дает $p(A)\mathbf{f}_j = p(\lambda_j)\mathbf{f}_j$. ▶

Если матрица A невырождена (все $\lambda_j \neq 0$), то

$$A\mathbf{f}_j = \lambda_j\mathbf{f}_j \Rightarrow \mathbf{f}_j = \lambda_j A^{-1}\mathbf{f}_j,$$

откуда вытекает, что собственными значениями обратной матрицы A^{-1} являются λ_j^{-1} и те же собственные векторы.

³⁾ Собственные векторы определены с точностью до умножения на константу. Принято считать, что все $\mathbf{f}_j$ нормированы: $\|\mathbf{f}_j\| = 1$.

Из записи (3.7) в форме

$$\begin{vmatrix} a_{11} - \lambda & \dots & a_{1n} \\ \dots & \ddots & \dots \\ a_{n1} & \dots & a_{nn} - \lambda \end{vmatrix} = 0$$

видно, что $p_A(\lambda) = \det(A - \lambda I)$ представляет собой многочлен n -й степени,

$$p_A(\lambda) = \lambda^n + \gamma_{n-1}\lambda^{n-1} + \dots + \gamma_1\lambda + \gamma_0,$$

который называют *характеристическим*.

Из

$$|T^{-1}AT - \lambda I| = |T^{-1}(A - \lambda I)T| = |T^{-1}| \cdot |A - \lambda I| \cdot |T| = |A - \lambda I|$$

вытекает, что полином $p_A(\lambda)$ инвариантен при замене координат, и потому — характеризует как матрицу, так и сам линейный оператор ⁴⁾ (независимо от базиса).

В фокус внимания часто попадают два коэффициента: γ_{n-1} , равный, по теореме Виета, сумме корней $\lambda_1(A) + \dots + \lambda_n(A)$, и γ_0 , равный произведению $\lambda_1(A) \dots \lambda_n(A)$,

$$\gamma_0 = \lambda_1(A) \dots \lambda_n(A) = \det A,$$

$$\gamma_{n-1} = \lambda_1(A) + \dots + \lambda_n(A) = \operatorname{tr} A,$$

где $\operatorname{tr} A = a_{11} + \dots + a_{nn}$ называют *следом матрицы (оператора)*.

В этом месте линейный анализ подходит к *фундаментальному поворотному моменту* своего развития. Если уравнение

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

имеет n различных действительных корней λ_j , то все слишком хорошо. Каждому λ_j соответствует свой собственный вектор $\mathbf{f}_j$. Эти векторы составляют базис новой системы, в которой матрица $T^{-1}AT$ принимает диагональный вид, а матрица перехода T образуется столбцами $\mathbf{f}_j$.

Действительно, для $T = [\mathbf{f}_1, \dots, \mathbf{f}_n]$ получается

$$T^{-1}AT = T^{-1}[A\mathbf{f}_1, \dots, A\mathbf{f}_n] = T^{-1}[\lambda_1\mathbf{f}_1, \dots, \lambda_n\mathbf{f}_n] = T^{-1}[\mathbf{f}_1, \dots, \mathbf{f}_n]\Lambda = T^{-1}T\Lambda = \Lambda.$$

Было бы опрометчиво сразу отказываться от всех этих благ из-за того, что полином n -й степени редко имеет n действительных корней. Обстоятельства подталкивают к выходу в комплексную плоскость. Там характеристическое уравнение всегда (по основной

⁴⁾ Естественно, что коэффициенты $p_A(\lambda)$ не меняются при переходе к другому базису.

теореме алгебры) имеет n корней. Но лиха беда начало. Одна уступка влечет за собой необходимость другой. Если в (3.6) параметр λ комплексный, то к рассмотрению надо допускать комплексные векторы⁵⁾, иначе нет смысла говорить о ненулевых решениях x уравнения (3.6).

Далее возникает необходимость в сдаче следующей позиции. Столбцы f_j (теперь комплексные)⁶⁾ образуют в совокупности матрицу перехода T , которая приводит A к диагональной форме $T^{-1}AT$, причем у $T^{-1}AT$ на диагонали могут стоять в том числе комплексные λ_j . Получается, что надо допускать к рассмотрению комплексные матрицы. Поэтому проще всего — и целесообразнее — с самого начала все считать комплексным⁷⁾.

Но тогда необходимо вернуться назад и пересмотреть уже сделанное. Оказывается, в ранее изложенном ничего не меняется, исключая понятия скалярного произведения и транспонирования матрицы. Последние приходится переопределить.

Под скалярным произведением в общем случае понимается

$$\langle x, y \rangle = \sum_j x_j y_j^*,$$

что при действительных координатах переходит в стандартное определение (2.4).

Что касается «транспонирования» комплексной матрицы, то в билинейных формах $\langle Vx, y \rangle$ действие матрицы V приходится перебрасывать с первого вектора на второй, не меняя значения функции, т. е.

$$\langle Vx, y \rangle = \langle x, V^*y \rangle. \quad (3.8)$$

Для справедливости (3.8) в случае действительной матрицы V надо, чтобы V^* была транспонированной матрицей V .

Действительно,

$$\langle Vx, y \rangle = \sum_i \left(\sum_k v_{ik} x_k \right) y_i = \sum_k \left(\sum_i v_{ki}^* y_i \right) x_k = \langle x, V^*y \rangle.$$

⁵⁾ Имеющие комплексные координаты.

⁶⁾ Если, конечно, таковые существуют.

⁷⁾ Логика комплексификации математического знания на других примерах прослеживается в [3].

Аналогичная выкладка в случае комплексных векторов и матриц показывает, что V^* должна получаться из V транспонированием с последующей заменой всех элементов на комплексно сопряженные. Такая матрица V^* называется **сопряженной** V . Формулы $(AB)^* = B^*A^*$ и $(A^*)^{-1} = (A^{-1})^*$ сохраняются и при таком понимании операции «звездочка».

3.3. Собственные векторы

Выход в комплексную плоскость обеспечивает существование n собственных значений $\lambda_1, \dots, \lambda_n$, но не n собственных векторов. «Потери» возникают при совпадении некоторых λ_j между собой, но это не выглядит поначалу большой неприятностью. Так или иначе, комплексный вариант теории начинает казаться не злонамеренной выдумкой, а удачной находкой, которая во многом спасает теорию матриц от блуждания впотьмах.

3.3.1. Теорема. Если матрица A имеет попарно различные собственные значения $\lambda_1, \dots, \lambda_n$, то каждому λ_j отвечает свой собственный вектор f_j , множество которых $\{f_1, \dots, f_n\}$ линейно независимо.

◀ Существование f_j очевидно, поскольку уравнение $(A - \lambda_j I)x = 0$ заведомо имеет ненулевое решение в силу вырожденности матрицы $A - \lambda_j I$.

Допустим, $\{f_1, \dots, f_n\}$ линейно зависимы. Выберем из $\{f_1, \dots, f_n\}$ минимальное число линейно зависимых векторов. Можно считать ⁸⁾, что это первые k векторов, которые удовлетворяют соотношению

$$\gamma_1 f_1 + \dots + \gamma_k f_k = 0, \quad (3.9)$$

где все γ_j ненулевые, иначе k не было бы минимально. В силу $Af_j = \lambda_j f_j$,

$$A(\gamma_1 f_1 + \dots + \gamma_k f_k) = \gamma_1 \lambda_1 f_1 + \dots + \gamma_k \lambda_k f_k = 0. \quad (3.10)$$

Умножая (3.9) на λ_k и вычитая из (3.10), получаем

$$\gamma_1 (\lambda_1 - \lambda_k) f_1 + \dots + \gamma_{k-1} (\lambda_{k-1} - \lambda_k) f_{k-1} = 0,$$

что противоречит минимальности k . ►

При $k < n$ попарно различных собственных значениях — теорема 3.3.1 по сути остается в силе. Гарантируется существование не менее k линейно независимых собственных векторов.

⁸⁾ По крайней мере, после перенумерации переменных.

3.3.2. Теорема. Пусть $Ax = \lambda x$, $yA = \mu y$ и $\lambda \neq \mu$. Тогда $\langle x, y \rangle = 0$. Другими словами, левый и правый собственные векторы, отвечающие разным собственным значениям, всегда ортогональны друг другу.

Собственный вектор-столбец x в $Ax = \lambda x$ называют *правым*, вектор-строку y в $yA = \mu y$ — *левым собственным вектором*. Очевидно, $A^*y^* = \bar{\mu}y^*$.

◀ Вычисление yAx двумя способами,

$$y(Ax) = \langle y, \lambda x \rangle = \lambda \langle y, x \rangle, \quad (yA)x = \langle \mu y, x \rangle = \mu \langle y, x \rangle,$$

приводит к $\langle y, x \rangle = 0$ в силу $\lambda \neq \mu$. ▶

Ортогональность левых и правых собственных векторов в предположениях теоремы 3.3.2 служит обычно источником различного рода двойственных результатов — см. главу 8.

Пример. Замена переменных $x = Ty$ приводит задачу Коши

$$\dot{x} = Ax, \quad x(0) = x_0$$

к виду

$$\dot{y} = T^{-1}ATy, \quad y(0) = T^{-1}x_0.$$

Если

$$T^{-1}AT = \begin{bmatrix} \lambda_1 & \dots & 0 \\ \vdots & \ddots & \vdots \\ 0 & \dots & \lambda_n \end{bmatrix},$$

то это «развязывает» исходную динамическую систему. В новых координатах система $\dot{y} = T^{-1}ATy$ распадается на независимые друг от друга скалярные уравнения $\dot{y}_k = \lambda_k y_k$, интегрирование которых по отдельности дает

$$y_k(t) = e^{\lambda_k t} y_k(0),$$

т. е.

$$y(t) = \begin{bmatrix} e^{\lambda_1 t} & \dots & 0 \\ \vdots & \ddots & \vdots \\ 0 & \dots & e^{\lambda_n t} \end{bmatrix} y(0).$$

Возврат в исходное пространство, $y = T^{-1}x$, приводит к решению

$$x(t) = T \begin{bmatrix} e^{\lambda_1 t} & \dots & 0 \\ \vdots & \ddots & \vdots \\ 0 & \dots & e^{\lambda_n t} \end{bmatrix} T^{-1}x(0).$$

3.4. Эскиз спектральной теории

Акцент ниже делается на фактах — обоснования рассматриваются в следующих главах.

Множество $\sigma(A)$ собственных значений λ_j матрицы A называют *спектром*, а максимальное значение модуля $|\lambda_j|$ — *спектральным радиусом*, и обозначают $\rho(A)$.

Спектр в значительной мере характеризует свойства матрицы. Например, из $\operatorname{Re} \lambda_j < 0$ для всех j — следует асимптотическая устойчивость (см. [4, т. 2]). системы $\dot{x} = Ax + b$. Из $\rho(A) < 1$ вытекает сходимость последовательных приближений $x^{k+1} = Ax^k + b$ к решению системы уравнений $x = Ax + b$.

Спектр не меняется при переходе к другому базису, т. е. характеризует сам линейный оператор, а не его конкретную запись в той или иной системе координат.

Если у $n \times n$ матриц A и B спектры состоят из n *попарно различных* точек и совпадают, то матрицы *подобны* — переходят одна в другую при замене системы координат, $B = T^{-1}AT$, — и представляют собой запись одного и того же оператора в различных базисах.

В общем случае это не так. Погоду портит возможность равенства собственных значений. Кратность λ как корня характеристического многочлена $p_A(\lambda)$ называется *алгебраической кратностью* собственного значения λ . Число линейно независимых решений уравнения $Ax = \lambda x$ называют *геометрической кратностью* собственного значения λ .

Если бы геометрическая кратность всегда совпадала с алгебраической, — у любой матрицы A было бы n линейно независимых собственных векторов $\{f_1, \dots, f_n\}$ и преобразование $T = [f_1, \dots, f_n]$ приводило бы A к диагональному виду.

Матрица, у которой геометрическая кратность какого-либо собственного значения меньше алгебраической, называется *дефектной*. Таковой является, например, $\begin{bmatrix} 0 & 1 \\ 0 & 0 \end{bmatrix}$. Собственное значение $\lambda = 0$ здесь имеет кратность 2, но ему соответствует лишь один собственный вектор $\begin{bmatrix} 1 \\ 0 \end{bmatrix}$. Вот еще несколько примеров

дефектных матриц:

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

Идея предельного перехода. Тот факт, что наличие кратных собственных значений имеет принципиальное значение, может показаться странным, ибо совпадение корней многочлена, равно как и вообще двух чисел, — вещь экзотическая. Более того, кратные λ всегда можно ликвидировать сколь угодно мало «пошевелив» матрицу. Или скажем иначе. Любую матрицу A можно представить как предел

$$A = \lim_{\varepsilon \rightarrow 0} A(\varepsilon), \quad (3.11)$$

где матрица $A(\varepsilon)$ при любом $\varepsilon \neq 0$ имеет «хороший» спектр — попарно различные собственные значения и, соответственно, n линейно независимых собственных векторов $f_j(\varepsilon)$. При этом возникает впечатление, что предельная матрица должна наследовать свойства $A(\varepsilon)$. В интересующем нас случае этого как раз не происходит.

«Беда» заключается в том, что линейно независимые векторы $f_j(\varepsilon)$ в пределе при $\varepsilon = 0$ могут стать зависимыми — и тогда предельный переход оказывается в некотором смысле холостым выстрелом. Однако идея представления матриц в виде (3.11) иногда эффективно работает (даже в тех случаях, где речь идет о собственных векторах, см. раздел 4.1).

Очевидно, $A + \varepsilon I$ имеет собственные значения $\lambda_j(A) + \varepsilon$. Поэтому с помощью сколь угодно малой добавки εI вырожденную матрицу A можно превратить в невырожденную $A(\varepsilon) = A + \varepsilon I$, сдвигая все точки спектра $\sigma(A)$ на ε .

Выше, однако, речь шла о более сложной задаче — о возмущении матрицы, ликвидирующем кратные собственные значения. Либерально настроенной части населения это кажется само собой разумеющимся, ибо для «управления» n коэффициентами характеристического полинома $|A - \lambda I| = 0$ имеется n^2 степеней свободы (n^2 элементов матрицы A). Тем не менее строгое обоснование очевидных вещей иногда наталкивается на определенные трудности⁹⁾.

⁹⁾ Формальное доказательство дано в главе 5.

3.5. Линейные пространства

Ориентация на модель, в которой векторы изначально представляются в форме $x = \{x_1, \dots, x_n\}$, где x_i вещественные или комплексные числа, — освобождает от массы скучных подробностей и охватывает 99 % приложений. Но из-за недостачи одного процента — дело может стоять на месте.

Вот как выглядит более фундаментальный подход. *Линейное пространство* X определяется как множество элементов, на котором заданы две операции: *сумма* $x + y$ и *произведение* λx (элемента $x \in X$ на число λ из некоторого поля¹⁰⁾), причем

$$x, y \in X \Rightarrow \alpha x + \beta y \in X,$$

а сами операции удовлетворяют следующим аксиомам:

- $1 \cdot x = x, \quad \lambda(\mu x) = \lambda\mu x,$
- $(\lambda + \mu)x = \lambda x + \mu x, \quad \lambda(x + y) = \lambda x + \lambda y,$

а также:

1. $x + y = y + x$ (коммутативность сложения),
2. $(x + y) + z = x + (y + z)$ (ассоциативность сложения),
3. существует нулевой элемент $0 \in X$ такой, что $x + 0 = x$,
4. существует обратный элемент $-x$ такой, что $x + (-x) = 0$.

Ощущения среднестатистического студента здесь естественны. Понятие линейного пространства «наводит тень на плетень», чтобы обучающей стороне было легче надувать щеки. И в этом есть свой резон, если после общих многозначительных разговоров курс целиком опирается на единственный частный случай.

А при широком охвате цель не достигается по другой причине. На первом этапе знакомства с предметом, дай бог, уследить за одной линией изложения. Дополнительная информация вредит, особенно, если не подчеркивается, что ее можно игнорировать полностью или частично.

Поэтому данный раздел на первом витке изучения линейной алгебры можно пропустить, но лучше — бегло просмотреть. Второе предпочтительнее, поскольку традиционный стиль обучения, тяготеющий к черно-белым тонам, далеко не оптимален. Нет никаких оснований отказываться от обычного взгляда на мир, когда что-то в фокусе внимания, а что-то — на периферии.

Примеры. Главный пример в данном контексте это, конечно, множество X из элементов

$$x = \{x_1, \dots, x_n\}. \quad (3.12)$$

¹⁰⁾ Пространство называют комплексным, если λ принадлежит комплексной плоскости.

Линейным пространством является множество матриц $m \times n$, а также совокупность многочленов

$$P_k(t) = \alpha_0 + \alpha_1 t + \dots + \alpha_{k-1} t^{k-1}, \quad k \leq n,$$

с обычными операциями сложения и умножения многочленов на число.

Координатами здесь можно считать коэффициенты $x_i = \alpha_{i-1}$, что соответствует выбору в качестве базиса $\{e_1 = 1, e_2 = t, \dots, e_n = t^{n-1}\}$. Но можно взять и другой базис. Например, $\{e_i = (t-a)^{i-1}\}$, — и тогда координатами будут коэффициенты в разложении Тэйлора:

$$P_n(t) = P_n(a) + P'_n(a)t + \dots + \frac{1}{(n-1)!} P_n^{(n-1)}(a).$$

С точки зрения данного определения линейными пространствами оказываются множества бесконечных последовательностей

$$\alpha_1, \dots, \alpha_k, \dots$$

с естественными операциями сложения и умножения, а также различные функциональные пространства.

Последние примеры относятся к линейным пространствам бесконечной размерности и рассматриваются в других математических дисциплинах. Линейная алгебра изучает пространства *конечной размерности*. Последняя определяется как максимальное число линейно независимых векторов в X (векторов в базисе), и обозначается как $\dim X$.

Разговоры о большой общности абстрактного подхода заканчиваются, правда, простой *теоремой об изоморфизме n -мерных пространств* [2, 7], которая означает следующее. Между линейными пространствами X и X' одинаковой размерности всегда можно установить взаимно однозначное соответствие $x \leftrightarrow x'$, при котором

$$x + y \leftrightarrow x' + y', \quad \lambda x \leftrightarrow \lambda x'.$$

Поэтому в любом случае анализ сводится (выбором базиса, например) к изучению арифметического пространства с элементами (3.12). Однако теорема об изоморфизме вовсе не перечеркивает целесообразность общего подхода. Здесь все, как в жизни. Одни на мир смотрят просто, иные все усложняют, но те и другие оказываются в проигрыше попеременно.

Евклидово пространство из линейного получается введением скалярного произведения. На докоординатном этапе форма $x \cdot y = x_1 y_1 + \dots + x_n y_n$, разумеется, дает осечку.

Формальное определение таково. *Скалярное произведение* — это так или иначе определенная операция, которая паре, вообще говоря, комплексных векторов x, y ставит в соответствие комплексное число $\langle x, y \rangle$, и удовлетворяет следующим требованиям (аксиомам):

- (i)° $\langle x, x \rangle$ — обязательно вещественное число, строго большее нуля при ненулевом x и равное нулю только при $x = 0$.
- (ii)° $\langle x, y \rangle = \langle y, x \rangle^*$, где звездочка обозначает комплексное сопряжение.
- (iii)° $\langle x + z, y \rangle = \langle x, y \rangle + \langle z, y \rangle$.

$$(iii)^{\circ} \quad \langle \lambda x, y \rangle = \lambda \langle x, y \rangle.$$

Заметим, что $\langle x, \lambda y \rangle = \lambda^* \langle x, y \rangle$, поскольку

$$\langle x, \lambda y \rangle = \langle \lambda y, x \rangle^* = \lambda^* \langle y, x \rangle^* = \lambda^* \langle x, y \rangle.$$

Даже на «координатном» этапе, как выясняется, перечисленным требованиям удовлетворяет не только стандартное

$$\langle x, y \rangle = \sum_j x_j y_j^*,$$

но и $\langle Vx, y \rangle$ с положительно определенной матрицей V (см. след. главу).

При изучении многочленов в качестве скалярного произведения может использоваться, например,

$$\langle P, Q \rangle = \int_0^1 P(t)Q(t) dt. \quad (3.13)$$

Отсюда, кстати, многое следует. Дело в том, что все предшествующие результаты, опиравшиеся на скалярное произведение, были связаны не с конкретной формой

$$\langle x, y \rangle = \sum_j x_j y_j^*,$$

а только с наличием свойств $(i)^{\circ} - (iii)^{\circ}$. Ничто другое в доказательствах не использовалось¹¹⁾. Поэтому в ситуации (3.13) сразу можно утверждать справедливость неравенства Коши—Буняковского (интегральный аналог)

$$\int_0^1 P(t)Q(t) dt \leq \sqrt{\int_0^1 P^2(t) dt} \sqrt{\int_0^1 Q^2(t) dt}.$$

3.6. Манипуляции с подпространствами

Было бы ошибкой думать, что общее понятие линейного пространства нужно лишь для того, чтобы пустить пыль в глаза и охватить полиномы.

При изучении обычной матрицы в обычной ситуации возникает потребность говорить о множестве значений Ax (образе, обозначаемом $\text{im } A$) либо о множестве решений $Ax = 0$ (ядре $\ker A$), т. е. о множестве векторов x , которые оператором отображаются в нуль.

¹¹⁾ Точнее говоря, все доказательства могли быть получены только лишь с использованием свойств $(i)^{\circ} - (iii)^{\circ}$.

И то и другое является линейным пространством (*подпространством*¹²⁾ R^n) в смысле предыдущего раздела, — и без соответствующих определений кое-что на понятийном уровне ускользает¹³⁾.

Что такое, например, множество значений Ax в случае R^3 ? Это может быть все пространство R^3 (матрица A невырождена), плоскость ($\text{rank } A = 2$), прямая ($\text{rank } A = 1$), точка ($\text{rank } A = 0$). Но при исходной модели описания $x = \{x_1, \dots, x_n\}$ нет возможности говорить о таких объектах до тех пор, пока там не введены координаты. В координатах же исходного пространства R^3 это достаточно неудобно, чтобы воспрепятствовать даже самым простым рассуждениям.

Нередко возникает необходимость оперировать *инвариантными* подпространствами матрицы (оператора) A .

3.6.1. Определение. Подпространство $X \subset R^n$ называется *инвариантным относительно* A , если $x \in X \Rightarrow Ax \in X$, т. е. $AX \subset X$.

Оператор A , рассматриваемый на инвариантном подпространстве X , т. е. сужение $A|_X$ оператора A на X , является, очевидно, линейным оператором.

При работе с подпространствами часто возникает необходимость рассматривать их комбинации. Суммой $X + Y$ подпространств $X, Y \subset R^n$ называется множество линейных комбинаций $\alpha x + \beta y$ элементов $x \in X$ и $y \in Y$. Другими словами, — линейная оболочка X и Y . Пересечением подпространств называют $X \cap Y$, где подразумевается обычное теоретико-множественное пересечение, т. е. $u \in X \cap Y$, когда и $u \in X$, и $u \in Y$. Наконец, $X^\perp$ определяют как *ортогональное дополнение*, при условии что $X + X^\perp = R^n$ и

$$x \in X, \quad y \in X^\perp \quad \Rightarrow \quad (x, y) = 0.$$

Если X плоскость, а Y прямая, не лежащая в этой плоскости, $X + Y$ даст трехмерное пространство, а $X \cap Y$ — точку. Если же $Y \subset X$ (прямая лежит в плоскости), то $X + Y = X$, $X \cap Y = Y$. В случае

$$\dim(X + Y) = \dim X + \dim Y$$

сумму $X + Y$ называют *прямой*.

Упражнения

- $\dim(X + Y) = \dim X + \dim Y - \dim X \cap Y$.
- Если линейный оператор A действует в R^n , то

$$\dim(\text{im } A) + \dim(\ker A) = n.$$

¹²⁾ Непустое подмножество X' линейного пространства X называется *линейным подпространством*, если $x, y \in X' \Rightarrow \alpha x + \beta y \in X'$.

¹³⁾ Как ускользает все, что не имеет имени.

- $\dim(\operatorname{im} A) = \operatorname{rank} A$.
- $\dim X^\perp = n - \dim X$.
- При замене координат меняется матрица, но не ее инвариантное подпространство.

3.7. Задачи и дополнения

- Из $\boxed{AB = B^{-1}(BA)B}$ очевидно: если хотя бы одна из матриц A, B невырождена, то AB и BA подобны, и $\sigma(AB) = \sigma(BA)$. Совпадение спектров, $\sigma(AB) = \sigma(BA)$, имеет место и в общем случае, когда обе матрицы вырождены.
 ◀ Доказательство легко сводится к предыдущему. Для достаточно малых $\epsilon \neq 0$ матрица $B(\epsilon) = B + \epsilon I$ невырождена, поэтому спектры $AB(\epsilon)$ и $B(\epsilon)A$ совпадают. Необходимое заключение дает предельный переход¹⁴⁾ при $\epsilon \rightarrow 0$. ▶
- У любого линейного оператора A , действующего в R^n , всегда существуют n инвариантных подпространств X_j , таких что размерность X_j равна j и

$$X_1 \subset X_2 \subset \dots \subset X_n = R^n.$$

◀ Это весьма важный факт, который очень просто доказывается по индукции. Пусть R — некоторое одномерное инвариантное подпространство¹⁵⁾ A^* , т. е. $A^*x' = \lambda x'$. Ортогональное дополнение $R^\perp$ имеет размерность $n-1$ и инвариантно относительно A . Действительно, если $y \in R^\perp$, т. е. $\langle x', y \rangle = 0$, то¹⁶⁾

$$\langle x', Ay \rangle = \langle A^*x', y \rangle = \langle \lambda x', y \rangle = 0.$$

Сужение A на $X_{n-1} = R^\perp$ имеет (по индуктивному предположению) цепочку вложенных инвариантных подпространств. Остается присовокупить $X_n = R^n$. ▶

¹⁴⁾ Предельный переход такого типа, как всегда, хорошо работает, когда поведение собственных векторов не играет роли.

¹⁵⁾ Существование которого следует из существования у A^* хотя бы одного собственного вектора.

¹⁶⁾ См. «правило переброски» по предметному указателю.

Глава 4

Квадратичные формы

4.1. Квадратичные формы

Кинетическая и потенциальная энергия механической системы с n степенями свободы обычно выражаются через обобщенные координаты q_j и обобщенные скорости $\dot{q}_j$ следующим образом:

$$T = \frac{1}{2} \sum_{j,k} a_{jk} \dot{q}_j \dot{q}_k, \quad \Pi = \frac{1}{2} \sum_{j,k} c_{jk} q_j q_k. \quad (4.1)$$

Квадратичные формы (4.1) проходят красной нитью через всю теорию малых колебаний [5].

Это хороший содержательный пример, на котором видна принципиальная роль квадратичных форм и их преобразований. Дело в том, что в задачах механики при выборе обобщенных координат всегда имеется большая свобода, и попасть в десятку экспромтом не так легко. Поэтому q_j как-то задают, а потом уже пытаются заменой координат привести T и Π к более простому виду. Отсюда, кстати, возникает задача о приведении к диагональному виду двух форм одним преобразованием — но об этом речь впереди.

Другой источник повышенного интереса к квадратичным формам — их роль в изучении стационарных точек, характеризуемых условием $\nabla \varphi(x) = 0$. Потенциал $\varphi(x)$ в окрестности таких точек представим в виде

$$\varphi(x) - \varphi(x + \Delta x) = \sum_{i,j} \frac{\partial^2 \varphi}{\partial x_i \partial x_j} \Delta x_i \Delta x_j + o(\dots).$$

Главный источник, конечно, сама линейная алгебра. Но на этом список далеко не заканчивается.

*Квадратичная форма*¹⁾

$$(Vx, x) = \sum_{i,j} v_{ij} x_i x_j, \quad (4.2)$$

¹⁾ Пока имеются в виду действительные переменные.

задается $n \times n$ матрицей V , которая обычно предполагается симметричной, $v_{ij} = v_{ji}$, поскольку в (Vx, x) коэффициент при $x_i x_j$ равен $v_{ij} + v_{ji}$ — и матрицу можно симметризовать, заменив v_{ij} полусуммой $(v_{ij} + v_{ji})/2$. Однако при рассмотрении *билинейной формы* (Vx, y) симметризация не проходит.

Напомним *правило переброски*:

$$(Vx, y) = (x, V^*y),$$

где V^* — (в данном случае) транспонированная матрица V .

Необходимость использования формулы $(Vx, y) = (x, V^*y)$ возникает довольно часто. Например, производная функции Ляпунова²⁾ (Vx, x) вдоль траекторий линейной динамической системы $\dot{x} = Ax$ равна³⁾

$$(V\dot{x}, x) + (Vx, \dot{x}) = (VAx, x) + (Vx, Ax) = ((VA + A^*V)x, x).$$

Задача упрощения. При переходе к другой (штрихованной) системе координат с помощью невырожденной матрицы S :

$$x = Sx',$$

функция (Vx, x) после подстановки $x = Sx'$ переходит в

$$(Vx, x) = (VSx', Sx') = (S^*VSx', x') = (V'x', x'),$$

откуда ясно, что матрица V , задающая квадратичную форму, при линейной замене системы координат $x = Sx'$ преобразуется по правилу

$$V' = S^*VS, \quad (4.3)$$

которое отличается от (3.1)⁴⁾. Но для ортогональных преобразований формулы (3.1) и (4.3) совпадают.

Если взаимоотношения с линейной алгеброй складываются так, что формулы типа (4.3) приходится как-то специально запоминать, то эти взаимоотношения нуждаются в небольшой коррекции. Проще помнить не формулы, а их выводы. Точнее даже, не выводы, а идеи, — представляющие собой кнопки, нажатие которых разворачивает вывод сам собой, и формула всплывает как побочный эффект.

²⁾ См. [4, т. 2].

³⁾ По правилу дифференцирования произведения.

⁴⁾ $V' = S^{-1}VS$.

4.1.1. Определение. Преобразование S называют *ортогональным*, если оно не меняет длину векторов,

$$(x, x) = (Sx, Sx) = (x, S^* Sx) = (S^* Sx, x),$$

откуда $S^* S = S^* S = I$, что означает

$$S^* = S^{-1}.$$

Иными словами, если матрица ортогональна, для получения обратной — ее достаточно просто транспонировать.

4.1.2. Определение. Множество векторов $\{f_1, \dots, f_n\}$ называется *ортогональным*, если $(f_i, f_j) = 0$ при $i \neq j$. Если, дополнительно, все векторы нормированы, $(f_i, f_i) = 1$, то их множество называется *ортонормированным*.

Сопоставление определения 4.1.2 с правилами транспонирования и умножения матриц показывает, что *столбцы (равно как и строки) ортогональной матрицы являются ортонормированным множеством*. При этом очевидно, что *множество всех ортогональных матриц компактно*.

Хотелось бы обратить внимание на некоторую двойственность, присущую линейной алгебре. Геометрический стиль мышления и алгебраический стиль доказательств. Геометрически ясно, например, что взаимно ортогональные векторы линейно независимы, но почти любой учебник предпочитает алгебраический вывод. И проблема не так проста, как может показаться при поверхностном взгляде.

Если $n > 3$, то о чем идет речь? В межотраслевом балансе, например, — что там такое ортогональность? Поэтому результат проще получить, манипулируя формулами. Геометрия, конечно, во главе угла, — но на заднем плане. Как ориентир, прожектор, профессиональный секрет, — но не как абсолютно законный инструмент. Беда здесь заключается в том, что ни у кого не доходили руки до того, чтобы геометрический стиль в алгебре полностью узаконить и стандартизовать. И дело даже не в том, что не нашлось нового Евклида. Просто путь наименьшего сопротивления иной. Там, где возникает геометрическая ясность, дать алгебраический эквивалент, как правило, не вызывает труда.

Упражнения

- Произведение ортогональных матриц — ортогональная матрица.
- Транспонирование ортогональной матрицы дает — ортогональную.

- Решением системы уравнений $Cx = b$ с ортогональной матрицей C является $x = C^*b$.
- Определитель ортогональной матрицы по модулю равен единице.
- Матрицы

$$\begin{bmatrix} 1 & 0 \\ 0 & -1 \end{bmatrix}, \quad \begin{bmatrix} \cos \varphi & -\sin \varphi \\ \sin \varphi & \cos \varphi \end{bmatrix}, \quad \begin{bmatrix} \cos \varphi & 0 & \sin \varphi \\ 0 & 1 & 0 \\ -\sin \varphi & 0 & \cos \varphi \end{bmatrix}$$

ортогональны.

Симметричные матрицы ($a_{ij} = a_{ji}$) широко распространены в физических приложениях. Таковы, в частности, формы (4.1).

4.1.3. *Все собственные значения вещественной симметричной матрицы вещественны, и им соответствуют вещественные собственные векторы.*

◀ Из $Ax = \lambda x$ следует $A\bar{x} = \bar{\lambda}\bar{x}$, где черта означает обычное комплексное сопряжение без транспонирования. Умножая скалярно первое равенство на $\bar{x}$, второе — на x , получаем

$$(\bar{x}, Ax) = \lambda(\bar{x}, x), \quad (x, A\bar{x}) = \bar{\lambda}(x, \bar{x}).$$

А поскольку левые части равенств здесь равны, то равны и правые. Поэтому

$$(\lambda - \bar{\lambda})(x, \bar{x}) = 0 \quad \Rightarrow \quad \lambda = \bar{\lambda}.$$

Вещественная разрешимость $Ax = \lambda x$ относительно x при вещественном λ очевидна. ►

4.1.4. Лемма. *Собственные векторы вещественной симметричной матрицы, соответствующие различным собственным значениям, — ортогональны.*

◀ Умножая скалярно $Ax = \lambda x$, $Ay = \mu y$, соответственно, на y , x , получаем $(y, Ax) = \lambda(y, x)$, $(x, Ay) = \mu(y, x)$, откуда, в силу $\lambda \neq \mu$,

$$(\lambda - \mu)(x, y) = 0 \quad \Rightarrow \quad (x, y) = 0. \quad \blacktriangleright$$

Таким образом, для вещественной симметричной матрицы A с попарно различными собственными значениями всегда существует ортонормированный базис, в котором A принимает диагональный вид. Оказывается, что требование попарного различия собственных значений необязательно, но это надо еще доказать.

◀ Пусть произвольная симметричная матрица A является пределом вида (3.11), т. е. $A(\varepsilon) \rightarrow A$ при $\varepsilon \rightarrow 0$, где $A(\varepsilon)$ при любом $\varepsilon \neq 0$ симметрична и все ее собственные значения попарно различны. Поэтому $A(\varepsilon \neq 0)$ могут быть приведены **ортогональным** преобразованием $T(\varepsilon)$ к диагональному виду.

А поскольку множество всех ортогональных матриц компактно, то $T(\varepsilon)$ сходится к некоторой тоже ортогональной матрице T . При этом ортонормированное множество $\{t_{\cdot 1}(\varepsilon), \dots, t_{\cdot n}(\varepsilon)\}$ собственных векторов $A(\varepsilon)$ сходится к ортонормированному множеству $\{t_{\cdot 1}, \dots, t_{\cdot n}\}$ собственных векторов матрицы A . ▶

Успех доказательства определился тем, что столбцы $T(\varepsilon)$, оставаясь в процессе предельного перехода взаимно ортогональными, не могли стать линейно зависимыми, как это могло бы случиться в общей ситуации.

Лемма 4.1.4, таким образом, приводит к общему результату.

4.1.5. Теорема. *Вещественная симметричная $n \times n$ матрица A всегда имеет n различных собственных векторов $\{f_1, \dots, f_n\}$, которые образуют ортонормированное множество, а ортогональное преобразование $T = [f_1, \dots, f_n]$ приводит A к диагональному виду:*

$$T^{-1}AT = T^*AT.$$

В терминах квадратичных форм (4.2) это означает, что (x, Vx) всегда может быть приведена *ортогональной заменой переменных* $x = Tu$ к виду

$$(x, Vx) = \sum_j \lambda_j u_j^2, \quad (4.4)$$

где λ_j — собственные значения V , если матрица V симметрична, — либо собственные значения симметризованной матрицы $(V + V^*)/2$.

Совокупность приведенных результатов создает впечатление, что с симметричными матрицами «все ясно». Но такого рода ситуации бывают обманчивы, и самые простые вопросы могут ставить в тупик.

Симметрично ли произведение симметричных матриц? Интуитивно почему-то кажется «да». Простая выкладка показывает, что «нет»,

$$(AB)^* = B^*A^* = BA.$$

Поэтому для $(AB)^* = AB$ надо, чтобы A и B коммутировали.

Тот же факт, что любая матрица A подобна некоторой симметричной S , $A = T^{-1}ST$, и всегда может быть представлена в виде произведения двух симметричных⁵⁾, мягко говоря, вызывает удивление. Но по здравому размышлению

⁵⁾ См. например, [18].

становится ясно, конечно, что здесь речь не может идти о вещественных матрицах. Однако странным кажется, что симметричные комплексные матрицы (в отличие от вещественных) никакими специфическими свойствами не обладают. Тем не менее это так. А вот *эрмитовы* матрицы, или *самосопряженные*, определяемые условием $A = A^*$, где звездочка обозначает уже транспонирование *плюс* комплексное сопряжение всех элементов, служат естественным обобщением вещественных симметричных матриц, обладая основными свойствами, присущими последним.

Эрмитовы матрицы. Если матрица A эрмитова, то:

- Все собственные значения A действительны.
- Квадратичная форма (x, Ax) принимает действительные значения при любых комплексных x .
- Матрица S^*AS тоже эрмитова.
- Существует представление $A = U^*\Lambda U$, где U — унитарная матрица, а Λ — вещественная диагональная. Обратное тоже верно. Кроме того, $A = C^*C$, где $C = \Lambda^{1/2}U$.
- Любая матрица B представима в виде $B = S + iT$, где обе матрицы S и T эрмитовы. (?)

4.2. Положительная определенность

Следующим шагом преобразования (4.4) может быть линейная замена переменных $z_j = \sqrt{|\lambda_j|} u_j$, после которой квадратичная форма приобретает вид

$$(x, Vx) = \sum_j \pm z_j^2. \quad (4.5)$$

В (4.5) коэффициенты при z_j^2 равны плюс или минус единице, а некоторые могут быть равны нулю.

Результирующее преобразование переменных x в z имеет вид $x = TSz$, где T ортогональная, а S диагональная матрица, у которой на диагонали стоят $1/\sqrt{|\lambda_j|}$, если $\lambda_j \neq 0$, и — нули в противном случае. Матрица V трансформируется в

$$(TS)^*V(TS) = S^*T^*VTS.$$

Разумеется TS уже не обязана быть ортогональной, а потому $(TS)^*V(TS)$ — не тот же самый линейный оператор V , как было бы в случае $(TS)^{-1}V(TS)$. В этом, собственно, и заключается некоторая путаница при изучении преобразований квадратичных форм.

На определенной стадии возникает привычка к тому, что преобразование матрицы меняет матрицу, а не оператор. Но это в случае $T^{-1}AT$. Матрица же

квадратичной формы меняется совсем по другому правилу: T^*AT . Опираясь на результаты предыдущего раздела, можно было бы ограничиться ортогональными преобразованиями, для которых $T^{-1}AT = T^*AT$, — и вопрос приведения к диагональному виду решался бы не только положительно, но и однозначно. Конечно, квадратичную форму к диагональному виду приводит «миллион» других преобразований, неортогональных. Казалось бы, их можно исключить из рассмотрения, чтобы не усложнять. Велика ли выгода от замены коэффициентов λ_j в (4.4) на ± 1 ?

Приведение двух форм. Выгода есть, и она выявляется в задаче приведения к диагональному виду сразу двух квадратичных форм⁶⁾. Оказывается, *если одна из матриц A, B положительно определена, то обе формы (x, Ax) и (x, Bx) приводятся к диагональной форме одновременно (одним и тем же преобразованием).*

4.2.1. Определение. Матрицу V , равно как и квадратичную форму (x, Vx) , называют *положительно определенной*, если

$$(x, Vx) > 0$$

при любом ненулевом x .

Совершенно ясно, что ортогональное преобразование приводит форму (x, Vx) к виду (4.4) со строго положительными λ_j — в противном случае для каких-то x было бы $(x, Vx) \leq 0$. На следующем шаге:

$$(x, Vx) = (z, z) = \sum_j z_j^2, \quad (4.6)$$

т. е. при переходе к переменным z матрица V преобразуется в единичную I , которую потом можно «крутить» любым ортогональным преобразованием T , ничего не меняя, ибо $T^*IT = I$.

Теперь, если одна из матриц A, B положительно определена (например, A), то (x, Ax) можно привести к сумме квадратов типа (4.6). При этом A и B перейдут, соответственно, в $I = S^*AS$ и $B' = S^*BS$. После этого квадратичная форма

$$(x, Bx) = (z, B'z)$$

приводится к диагональной форме некоторым ортогональным преобразованием T , которое на первую форму (z, Iz) уже не влияет⁷⁾.

⁶⁾ Кинетической и потенциальной энергии, например.

⁷⁾ Если же $\Lambda \neq I$, то $T^*\Lambda T$ не обязана быть диагональной. В этом, собственно, и заключается роль предположения о положительной определенности одной из матриц, которое позволяет обеспечить $\Lambda = I$.

Преобразование ST приводит обе формы (x, Ax) и (x, Bx) к диагональному виду — одновременно.

Пример. Пусть связанные линейные осцилляторы описываются системой дифференциальных уравнений

$$\begin{cases} b_{11}\ddot{x}_1 + b_{12}\ddot{x}_2 + a_{11}x_1 + a_{12}x_2 = 0, \\ b_{12}\ddot{x}_1 + b_{22}\ddot{x}_2 + a_{12}x_1 + a_{22}x_2 = 0. \end{cases} \quad (4.7)$$

Это могут быть индуктивно связанные колебательные контуры, либо связанные пружиной маятники.

Квадратичные формы

$$\begin{cases} (x, Ax) = a_{11}x_1^2 + 2a_{12}x_1x_2 + a_{22}x_2^2, \\ (\dot{x}, B\dot{x}) = b_{11}\dot{x}_1^2 + 2b_{12}\dot{x}_1\dot{x}_2 + b_{22}\dot{x}_2^2 \end{cases}$$

представляют потенциальную и кинетическую энергию системы, и потому — положительно определены.

Поиск решения (4.7) в форме $x = \begin{bmatrix} c_1 \\ c_2 \end{bmatrix} \cos(\omega t + \alpha)$ приводит к системе уравнений

$$\begin{cases} (a_{11} - b_{11}\omega^2)c_1 + (a_{12} - b_{12}\omega^2)c_2 = 0, \\ (a_{12} - b_{12}\omega^2)c_1 + (a_{22} - b_{22}\omega^2)c_2 = 0, \end{cases}$$

совместной лишь в случае равенства нулю детерминанта,

$$\begin{vmatrix} a_{11} - b_{11}\omega^2 & a_{12} - b_{12}\omega^2 \\ a_{12} - b_{12}\omega^2 & a_{22} - b_{22}\omega^2 \end{vmatrix} = 0, \quad (4.8)$$

что называют *секулярным*, или *вековым уравнением*.

Корни (4.8) дают две так называемые *нормальные частоты*⁸⁾ ω_1 и ω_2 . В итоге решение (4.7) принимает вид

$$x(t) = \begin{bmatrix} c_1 \\ c_2 \end{bmatrix} \cos \omega_1(t) + \begin{bmatrix} c'_1 \\ c'_2 \end{bmatrix} \cos \omega_2(t).$$

Интересно, что при этом в системе присутствуют две «виртуальные» гармонические компоненты ξ_1 и ξ_2 , не взаимодействующие друг с другом. Переменные⁹⁾

$$\xi = \begin{bmatrix} \xi_1 \\ \xi_2 \end{bmatrix} \text{ линейно выражаются через } x = \begin{bmatrix} x_1 \\ x_2 \end{bmatrix},$$

$$\xi = Tx,$$

причем T «развязывает» систему (4.7), приводя ее к виду

$$\begin{cases} \ddot{\xi}_1 + \omega_1^2 \xi_1 = 0, \\ \ddot{\xi}_2 + \omega_2^2 \xi_2 = 0. \end{cases}$$

⁸⁾ Проверьте, что $\omega_1^2, \omega_2^2 > 0$.

⁹⁾ Их называют *нормальными координатами*.

Преобразование T — это то самое преобразование, которое приводит одновременно квадратичные формы (x, Ax) и (x, Bx) к диагональному виду. При определенных навыках владения лагранжевым формализмом с этого можно было бы начинать.

Роль положительно определенных матриц, конечно, не сводится к участию в управлении парами квадратичных форм. Есть другое важное направление, связанное с изучением экстремальных точек. Если градиент $\nabla\varphi(x)$ обращается в нуль в точке x^* , то минимум φ характеризуется условием

$$\Delta\varphi = \varphi(x) - \varphi(x^*) = \left(\Delta x, \left[\frac{\partial^2 \varphi}{\partial x_i \partial x_j} \right] \Delta x \right) + o(\Delta x^2) > 0$$

при достаточно малых по норме Δx . Таким образом, вопрос сводится к положительной определенности гессиана $\left[\frac{\partial^2 \varphi}{\partial x_i \partial x_j} \right]$.

4.3. Инерция и сигнатура

Квадратичная форма (4.5) после перенумерации переменных может быть приведена к виду

$$(x, Vx) = \sum_{j=1}^p z_j^2 - \sum_{j=p+1}^q z_j^2. \quad (4.9)$$

Число $q \leq n$ называют *рангом квадратичной формы*, p и $q - p$ — *индексами инерции* (положительным и отрицательным), а разность индексов — *сигнатурой*.

За пределами положительной определенности квадратичные формы не так интересны и популярны, но в отдельные моменты понимание их устройства в общем случае — играет принципиальную роль.

Естественные вопросы по отношению к (4.9) заключаются в правомочности данных выше определений. Не могут ли p и q поменяться, если от (x, Vx) к правой части (4.9) идти другим путем?

Не могут. Это называют *законом инерции*. Разумеется, необходимо обоснование.

◀ В предположении противного возможны два представления с различными $\{p, q\}$ и $\{r, s\}$:

$$(x, Vx)_I = z_1^2 + \dots + z_p^2 - z_{p+1}^2 - \dots - z_q^2,$$

$$(x, Vx)_{II} = \xi_1^2 + \dots + \xi_r^2 - \xi_{r+1}^2 - \dots - \xi_s^2.$$

Сразу заметим, что $q = s = \text{rank } V$, поскольку при невырожденном преобразовании

$$\text{rank } T^*VT = \text{rank } V,$$

что узаконивает понятие ранга квадратичной формы.

Допустим теперь, что $p > r$, и переменным z и ξ соответствуют базисы $\{e_1, \dots, e_p\}$ и $\{f_1, \dots, f_n\}$. Пусть X обозначает линейную оболочку векторов $\{e_1, \dots, e_p\}$, а Y — линейную оболочку $\{f_{r+1}, \dots, f_n\}$.

В силу $p > r$ размерность $\dim X \cap Y \geq 1$, и для ненулевых $x \in X \cap Y$ очевидно, $(x, Vx)_I > 0$, а $(x, Vx)_{II} \leq 0$, что дает противоречие. ►

4.4. Условный экстремум

Использование метода множителей Лагранжа для решения задачи на условный экстремум¹⁰⁾

$$(x, Vx) \rightarrow \max, \quad (x, x) = 1 \quad (4.10)$$

приводит к необходимости решения параметрической системы уравнений

$$Vx = \lambda x,$$

проистекающей из максимизации лагранжиана

$$L(x, \lambda) = (x, Vx) - \lambda(x, x)$$

по x .

Если упорядочить собственные значения матрицы V в порядке убывания

$$\lambda_1 \geq \dots \geq \lambda_n,$$

то

$$\lambda_1 = \max_x \frac{(x, Vx)}{(x, x)}, \quad \lambda_n = \min_x \frac{(x, Vx)}{(x, x)},$$

что сразу вытекает из представления (4.4), поскольку в

$$(x, Vx) = \lambda_1 u_1^2 + \dots + \lambda_n u_n^2$$

¹⁰⁾ Матрица V предполагается симметричной.

имеет место $(x, x) = (u, u)$, так как переход от x к u осуществляется ортогональным преобразованием, которое не меняет длину векторов.

Промежуточные собственные значения могут быть охарактеризованы в терминах минимаксов [1, 18], что называют вариационным описанием собственных значений симметричных матриц.

4.5. Сингулярные числа

Далее речь идет о вещественной, вообще говоря *несимметричной* матрице A . Другими словами, о матрице общего вида. Оказывается, что всегда существует *ортонормированная* система векторов¹¹⁾ $\{f_1, \dots, f_n\}$, переводимая матрицей A в *ортогональную*¹²⁾ $\{Af_1, \dots, Af_n\}$. Длины и направления векторов f_i под воздействием A меняются, но взаимное расположение (в смысле попарной ортогональности) сохраняется.

Это полезный факт, который легко устанавливается, и заодно выводит на понятие сингулярных чисел матрицы.

◀ У симметричного¹³⁾ отображения A^*A всегда существует ортонормированный базис $\{f_1, \dots, f_n\}$ из его собственных векторов (теорема 4.1.5).

Преобразование A (не A^*A !) ортогональные векторы этого базиса преобразует в ортогональные, поскольку

$$(Af_i, Af_j) = (A^*Af_i, f_j) = \lambda_i(f_i, f_j),$$

откуда ясно, что $(Af_i, Af_j) = 0$, если $i \neq j$, и $(Af_i)^2 = \lambda_i$, где λ_i — собственные значения оператора A^*A . ▶

Величины $\alpha_i = \sqrt{\lambda_i}$ называют *сингулярными числами матрицы* A . Максимальное α_i представляет собой евклидову норму A (раздел 6.2).

4.5.1. Любая матрица A может быть разложена в произведение
 $A = TS$ *ортогональной матрицы T и симметричной неотрицательно определенной матрицы S .*

¹¹⁾ Единственная, если A невырождена.

¹²⁾ Точнее говоря, часть векторов Af_i могут быть нулевыми, если A вырождена.

¹³⁾ Поскольку $(A^*A)^* = A^*A$.

◀ Множество ортонормированных векторов $\mathbf{e}_i = A\mathbf{f}_i/\sqrt{\lambda_i}$, $\lambda_i \neq 0$, дополним (если A вырождена) произвольным образом до ортонормированного базиса, и пусть T ортогональная матрица, переводящая $\{\mathbf{f}_i\}$ в $\{\mathbf{e}_i\}$,

$$T\{\mathbf{f}_1, \dots, \mathbf{f}_n\} = \{\mathbf{e}_1, \dots, \mathbf{e}_n\}.$$

Преобразование $S = T^{-1}A$ в ортонормированном базисе $\mathbf{e}$ имеет диагональный вид $\Lambda = \text{diag}\{\sqrt{\lambda_1}, \dots, \sqrt{\lambda_n}\}$. Поэтому в исходном, ортонормированном базисе $S = U^* \Lambda U$ — симметрична. Остается заметить

$$S = T^{-1}A \Rightarrow A = TS. \quad \blacktriangleright$$

Из приведенного рассуждения следует также возможность *сингулярного разложения* любой матрицы A в виде произведения

$$A = U \Lambda W,$$

где U и W — ортогональные матрицы, а Λ — диагональная матрица с сингулярными числами α_i на диагонали¹⁴⁾.

Разложения подобного рода возможны и для прямоугольных матриц [18].

4.6. Биортогональные базисы

У матрицы A общего вида, имеющей попарно различные собственные значения $\{\lambda_1, \dots, \lambda_n\}$, существует базис из собственных векторов

$$\{\mathbf{e}_1, \dots, \mathbf{e}_n\},$$

но он не обязан быть ортогональным, что с инструментальной точки зрения — серьезная потеря. «Дух ортогональности», тем не менее, сохраняется в общем случае, хотя и в замаскированном виде.

Транспонированная матрица A^* имеет те же собственные значения, но другие собственные векторы

$$\{\mathbf{f}_1, \dots, \mathbf{f}_n\},$$

отвечающие сопряженным собственным значениям $\bar{\lambda}_i$.

4.6.1. Теорема. Базисы $\{\mathbf{e}_i\}$ и $\{\mathbf{f}_i\}$ биортогональны:

$$\langle \mathbf{e}_i, \mathbf{f}_j \rangle = 0 \quad \text{при} \quad i \neq j.$$

¹⁴⁾ Отсюда сразу следует $|\det A| = \alpha_1 \cdots \alpha_n$.

◀ С одной стороны,

$$\langle Ae_i, f_j \rangle = \langle \lambda_i e_i, f_j \rangle = \lambda_i \langle e_i, f_j \rangle.$$

С другой —

$$\langle Ae_i, f_j \rangle = \langle e_i, A^* f_j \rangle = \langle e_i, \bar{\lambda}_j f_j \rangle = \lambda_j \langle e_i, f_j \rangle,$$

что при $\lambda_i \neq \lambda_j$ возможно лишь в случае $\langle e_i, f_j \rangle = 0$. ▶

Базисы $\{e_i\}$ и $\{f_i\}$ можно так нормировать, что

$$\langle e_i, f_i \rangle = 1, \quad i = 1, \dots, n.$$

◀ Действительно, разложение e_i по базису $\{f_1, \dots, f_n\}$:

$$e_i = \xi_1 f_1 + \dots + \xi_n f_n,$$

после умножения на e_i дает

$$\langle e_i, e_i \rangle = \xi_1 \langle f_1, e_i \rangle + \dots + \xi_n \langle f_n, e_i \rangle = \xi_i \langle f_i, e_i \rangle.$$

Если базис $\{e_i\}$ ортонормирован, то векторы $\{f_i\}$ достаточно заменить на $\{f_i/\xi_i\}$. ▶

Результаты типа теоремы 4.6.1 выглядят несколько схоластично, когда излагаются в отрыве от мотивировок. На вопрос «зачем это нужно?» — здесь ответить легко, но приходится отвлекаться.

Вот одна из возможных сфер приложения.

В численных методах (глава 10) при изучении спектральных свойств матриц требуется, например, переходить от оператора A к его сужениям на собственные подпространства. Скажем, A имеет базис из собственных векторов $\{e_1, \dots, e_n\}$. Возникает нужда в конкретных рецептах. Как, скажем, записать матрицу, действующую в собственном подпространстве $\{e_2, \dots, e_n\}$?

Если $Ae_i = \lambda_i e_i$, $A^* f_i = \lambda_i^* f_i$ и

$$e_i = \begin{bmatrix} e_{1i} \\ \vdots \\ e_{ni} \end{bmatrix}, \quad f_i = \begin{bmatrix} f_{1i} \\ \vdots \\ f_{ni} \end{bmatrix},$$

то матрица

$$A_i = A - \lambda_i e_i f_i^* = A - \lambda_i \begin{bmatrix} e_{1i} \\ \vdots \\ e_{ni} \end{bmatrix} [f_{1i} \dots f_{ni}]$$

имеет те же собственные значения, что и A , за исключением λ_i , которое заменяется нулем:

$$A_i e_i = A e_i - \lambda_i (e_i f_i^*) e_i = \lambda_i e_i - \lambda_i e_i (f_i^* e_i) = 0,$$

поскольку¹⁵⁾ $f_i^* e_i = 1$.

Что касается остальных собственных значений и векторов, то при $i \neq j$:

$$A_i e_j = A e_j - \lambda_i (e_i f_i^*) e_j = \lambda_j e_j - \lambda_i e_i (f_i^* e_j) = \lambda_j e_j,$$

так как $f_i^* e_j = 0$

Билинейное разложение. Приведенный факт представляет собой частный случай более общего трюка.

4.6.2. Теорема. С помощью биортогональных базисов $\{e_i\}$ и $\{f_i\}$ — при условии нормировки $\langle e_i, f_j \rangle = \delta_{ij}$ ¹⁶⁾ — матрица A с попарно различными собственными значениями $\{\lambda_1, \dots, \lambda_n\}$ может быть представлена как

$$A = \lambda_1 e_1 f_1^* + \dots + \lambda_n e_n f_n^*. \quad (4.11)$$

◀ Матрица $H_k = e_k f_k^*$ имеет элементы $h_{ij}^k = e_{ik} f_{jk}^*$, а матрица $\sum_k H_k$ — элементы

$$\sum_k h_{ij}^k = \sum_k e_{ik} f_{jk}^* = \langle e_i, f_j^* \rangle = \delta_{ij}.$$

Поэтому $\sum_k H_k = I$, т. е.

$$I = e_1 f_1^* + \dots + e_n f_n^*.$$

Умножая последнее матричное равенство слева на A и пользуясь ассоциативностью матричного умножения, получаем (4.11). ▶

4.7. Сопряженное пространство

Наличие в евклидовом пространстве $E = R^n$ скалярного произведения мешает оценить целесообразность понятия *сопряженного пространства* E^* . Последнее естественно воспринимается там, где без него не обойтись. В линейной же алгебре многозначительные разговоры о E^* кончаются почти безрезультатно, порождая разочарование.

Поэтому лучше заранее отказаться от особых ожиданий, воспринимая сопряженное пространство как инструмент, необходимый

¹⁵⁾ Изменение порядка сомножителей превращает матрицу $e_i f_i^*$ в скаляр $f_i^* e_i$.

¹⁶⁾ Где $\delta_{ij} = \begin{cases} 1, & \text{если } i = j, \\ 0, & \text{если } i \neq j \end{cases}$ — символ Кронекера.

в смежных областях (функциональный анализ, тензорное исчисление). Что касается матриц, то здесь его можно рассматривать как дополнительный способ мышления, подготавливающий заодно к изучению других территорий.

В то же время нельзя сказать, что в линейной алгебре E^* совсем не имеет реальной почвы. Даже наоборот. Другое дело, что в процессе исследования появляются возможности обойтись без E^* . Но соответствующие «следы» остаются. В частности, матрицы, задающие линейное преобразование и квадратичную форму, при замене координат преобразуются по разным законам. Это как раз из-за E^* .

Вообще использование в линейной алгебре разных пространств изначально выглядит более разумно, чем «приведение к одному знаменателю» R^n . Пусть $y = Ax$. Почему x и y надо считать векторами одного пространства? Обычно они даже измеряются в разных единицах (x , скажем, в трудоднях, а y в копейках), и надо сильно постараться, чтобы сгладить это различие¹⁷⁾. А уж если речь идет о билинейной форме (Ax, y) , то это функция двух переменных, и от рассмотрения двух пространств, одного — для x , другого — для y , уже не уклониться. На диагонали $x = y$, правда, можно обойтись одним пространством, но происхождение (Ax, x) от (Ax, y) дает о себе знать.

Короче говоря, *сопряженное пространство* E^* определяется как пространство линейных функционалов¹⁸⁾ $f(x)$, заданных на E (в частности, $E = R^n$).

В базисе $\{e_1, \dots, e_n\}$

$$f(x) = f(\xi^1 e_1 + \dots + \xi^n e_n) = f_1 \xi^1 + \dots + f_n \xi^n,$$

где коэффициенты $f_j = f(e_j)$.

Сумма функционалов и умножение на число определяются естественным образом, откуда становится ясно, что E^* — линейное пространство, причем n -мерное, если n -мерно E , поскольку любой функционал f определяется n числами (координатами) $f_j = f(e_j)$.

Дальше иногда мешают двигаться возгласы. Дескать, $f(x)$ это скалярное произведение, и огород можно больше не городить.

¹⁷⁾ Это делается фиксацией системы единиц, но тогда замена координат порождает новую головную боль.

¹⁸⁾ Функционалом принято называть скалярную функцию. Под линейностью $f(x)$ понимается обычное $f(\alpha x + \beta y) = \alpha f(x) + \beta f(y)$.

В этом есть свой резон, но резон есть и в продолжении начатого дела. Более того, намек насчет скалярного произведения стоит учесть, вводя обозначение $f(x) = \langle f, x \rangle$, которое вызывает полезные ассоциации, подчеркивая некое равноправие x и f . Число $\langle f, x \rangle$, как функция двух переменных, линейно по каждому аргументу в отдельности. Как x , так и f — векторы, но природа их различна. Векторы x из E называют *контравариантными*, f из E^* — *ковариантными*. Различие, конечно, не только в названии.

Главное различие фундаментально и заключено в том, что x и f подчиняются совсем разным законам преобразования координат (см. далее). Векторы x — это аргументы, а наборы координат f — это коэффициенты преобразований. Там и там, конечно, числа. Но за числами, как известно, прячутся разные вещи. Мы уже едва не попали впросак, согласившись, что линейное преобразование и квадратичную форму определяет матрица. В обоих случаях, правда, суть дела заслоняют похожие таблицы чисел, — но разве это основание игнорировать остальное? Шахматные фигуры тоже можно характеризовать размером и весом, но мы их классифицируем все же по роли в игре. Пространства E и E^ нуждаются в таком же разумном подходе. Попытка игнорировать их различие мешает правду с ложью и порождает «таинственные» явления. Скорость — вектор, и угловая скорость — вектор. Почему же при определенных обстоятельствах они ведут себя различно? Линейное отображение — «матрица», и квадратичная форма — «матрица». Почему опять возникают аномалии? Да потому, что все это тензоры различной природы (см. следующий раздел), имеющие из-за стечения обстоятельств одинаковый внешний вид.*

Чтобы лучше почувствовать проблему, полезно рассмотреть нелинейное преобразование

$$y_j = \max_k a_{jk} x_k, \quad j, k = 1, \dots, n,$$

однозначно характеризуемое матрицей $[a_{jk}]$. Пытаясь произвести линейную замену переменных, легко убедиться, что от матрицы здесь — одно название. Объект совсем из другой оперы.

Вернемся, однако, к взаимосвязи $E \Leftrightarrow E^*$.

Векторы x и f считаются ортогональными, если $\langle f, x \rangle = 0$. Определение выглядит странно (ибо x и f находятся в разных пространствах), но это лишь первое впечатление.

4.7.1. Базисы

$$\{e_1, \dots, e_n\} \subset E, \quad \{f^1, \dots, f^n\} \subset E^*$$

называются взаимными (биортогональными), если $\langle f^j, e_k \rangle = \delta_k^j$.

Символ Кронекера $\delta_k^j = \begin{cases} 1, & \text{если } j = k, \\ 0, & \text{если } j \neq k \end{cases}$ здесь в отличие от обычного варианта имеет один индекс верхний, другой — нижний. Манипулирование индексами вообще характерно для тензорного исчисления.

При заданном базисе $\{e_1, \dots, e_n\} \subset E$ всегда существует взаимный ему базис $\{f^1, \dots, f^n\} \subset E^*$, поскольку система уравнений (при каждом j)

$$\langle f^j, e_k \rangle = \delta_k^j, \quad k = 1, \dots, n$$

однозначно определяет вектор (линейный функционал) f^j .

4.7.2. Взаимные базисы характеризуются тем, что координаты любых векторов

$$x = \xi^1 e_1 + \dots + \xi^n e_n, \quad f = \eta_1 f^1 + \dots + \eta_n f^n,$$

удовлетворяют соотношениям

$$\xi^j = \langle f^j, x \rangle,$$

$$\eta_j = \langle f, e_j \rangle,$$

$$\langle f, x \rangle = \sum_k \eta_k \xi^k.$$

◀ Очевидно,

$$\langle f^j, x \rangle = \left\langle f^j, \sum_k \xi^k e_k \right\rangle = \sum_k \xi^k \langle f^j, e_k \rangle = \xi^j.$$

Остальные соотношения устанавливаются аналогично. ▶

Сопряженный оператор. Если A — линейный оператор, действующий в E , то $f(Ax) = \langle f, Ax \rangle$ — есть некоторый линейный функционал $g(x) = \langle g, x \rangle$. Сопряженный оператор A^* определяется условием $g = A^* f$, т. е.

$$\langle f, Ax \rangle = \langle A^* f, x \rangle.$$

Таким образом, если $f(x)$ меняется в связи с воздействием на аргумент линейного преобразования A , то A^* позволяет достичь того же результата (изменения), не трогая x , а воздействуя на функционал, $f \Rightarrow A^* f$.

Сопряженный к сопряженному — снова исходный оператор, $(A^*)^* = A$. Так же, как $(E^*)^* = E$. (?)

Двойственность. Взаимоотношения между векторами и операторами сопряженных пространств лежат в основе различного рода двойственных результатов. К этой категории, например, может быть отнесено содержание раздела 4.6, где изучение «объекта» Ax ведется по наблюдению его «реакций». Очень естественный принцип исследования, кстати. Заключение о свойствах Ax делается на основании того, как Ax ведет себя под воздействием линейных функционалов (f, Ax) . Успех дела в данном случае обеспечивается «сопряженной равноправностью» x и f .

Так сложилось, что о двойственности часто говорят полувосторженно-полузагадочно. Начинаешь искать, где бы почитать, — нигде. Встречаются одни частности. Но где общая теория?

Секрет в том, что никакой общей теории нет. Двойственность — это идея, работающая в разных областях в разных формах. Между различными объектами устанавливаются некие взаимоотношения (сходства, противоположности, дополненности), благодаря которым из свойств одних объектов делаются заключения о поведении — других. В подобном ключе естественно интерпретируются многие результаты о линейных неравенствах. См. также разделы 6.2 — о двойственной норме, и 8.5 — о двойственности в линейном программировании¹⁹⁾.

4.8. Преобразования и тензоры

Пусть

$$\hat{e}_j = \sum_k t_j^k e_k$$

задает переход от базиса

$$\{e_1, \dots, e_n\} \quad \text{к} \quad \{\hat{e}_1, \dots, \hat{e}_n\},$$

а $\{f^1, \dots, f^n\}$ и $\{\hat{f}^1, \dots, \hat{f}^n\}$ — соответствующие биортогональные базисы.

¹⁹⁾ Двойственность широко используется вообще: в оптимизации, геометрии, топологии, теории аналитических функций и т. д.

Взаимосвязь между последними определяется соотношениями:

$$\boxed{\mathbf{r}^k = \sum_j t_j^k \widehat{\mathbf{f}}^j}, \quad (4.12)$$

где, обратим внимание, суммирование идет по нижнему индексу t_j^k , что соответствует транспонированию $[t_j^k]$. При необходимости выразить $\widehat{\mathbf{f}}$ через $\mathbf{f}$ матрицу $[t_j^k]^*$ надо обратить,

$$\widehat{\mathbf{f}} = \{[t_j^k]^*\}^{-1} \mathbf{f}.$$

Если элементы матрицы $\{[t_j^k]^*\}^{-1}$ обозначить через s_j^k , то $\widehat{\mathbf{f}}^k = \sum_j s_j^k \mathbf{f}^j$.

Для тензорного исчисления характерно жонглирование индексами. Базисы E помечаются нижними индексами, координаты — верхними. В E^* правило обратное. Одна из причин — местоположение индекса показывает, о чем идет речь. Другая причина имеет стенографическую природу. Если у сомножителей один и тот же индекс стоит один раз внизу, другой вверху, то по этому индексу, считается, идет суммирование. Это позволяет опускать знак $\sum$, что при избытии сумм расчищает игровое поле.

Поскольку тензоры в данном контексте играют третьестепенную роль, выгоды стенографии не используются (чтобы не привыкать).

◀ Проверим теперь (4.12). С одной стороны, в силу взаимности базисов,

$$\langle \mathbf{r}^k, \widehat{\mathbf{e}}_j \rangle = \left\langle \mathbf{r}^k, \sum_q t_j^q \mathbf{e}_q \right\rangle = t_j^k,$$

с другой, — если связь $\mathbf{r}^k = \sum_p \tau_p^k \widehat{\mathbf{f}}^p$ считать пока неизвестной,

$$\langle \mathbf{r}^k, \widehat{\mathbf{e}}_j \rangle = \left\langle \sum_p \tau_p^k \widehat{\mathbf{f}}^p, \widehat{\mathbf{e}}_j \right\rangle = \tau_j^k,$$

откуда $\tau_j^k = t_j^k$. ▶

Что касается формул преобразования соответствующих координат, то

$$\widehat{\xi}^j = \sum_k s_k^j \xi^k, \quad \widehat{\eta}_j = \sum_k t_j^k \eta_k,$$

т. е.

$$\widehat{\xi} = \{[t_j^k]^*\}^{-1}\xi, \quad \widehat{\eta} = [t_j^k]\eta,$$

где ξ — координаты в E , а η — в E^* .

Различие формул преобразования координат у контравариантных и ковариантных векторов представляет собой крайне важное обстоятельство, игнорирование которого ведет к принципиальным ошибкам.

Если матрица $[t_j^k]$ ортогональна, то $\{[t_j^k]^\}^{-1} = [t_j^k]$, — и различие исчезает. Но при этом нельзя забывать, что не исчезает явление как таковое.*

Для многоиндексных объектов, каковые называются *тензорами*, такое различие играет аналогичную роль, о чем напоминает рассмотрение квадратичных форм. Мы не останавливаемся подробно на определении, да это и не может увеличить ясность²⁰⁾. Вектор — это тензор первого ранга, матрица — второго. В общем случае тензора n -го ранга стандартной моделью служит полилинейная функция — линейная по каждому вектор-аргументу при фиксации остальных.

Главной характеристикой тензора является количество верхних и нижних индексов, что определяет правила его преобразования при переходе к другому базису²¹⁾, — см., например, [2].

4.9. Задачи и дополнения

- Если $\sum_{i,j} v_{ij}x_iy_j$ представляет собой запись *билинейной формы* (Vx, y) в базисе $\{e_1, \dots, e_n\}$, то $(Ve_i, e_j) = v_{ij}$.

²⁰⁾ Ясности способствуют примеры, но это задача другого курса.

²¹⁾ При этом говорят, что тензор p раз контравариантный и q раз ковариантный, если p и q , соответственно, число верхних и нижних индексов.

- Пусть у симметричной матрицы V все главные миноры:

$$D_1 = v_{11}, \quad D_2 = \begin{vmatrix} v_{11} & v_{12} \\ v_{21} & v_{22} \end{vmatrix}, \quad \dots, \quad D_n = \begin{vmatrix} v_{11} & v_{12} & \dots & v_{1n} \\ v_{21} & v_{22} & \dots & v_{2n} \\ \dots & \dots & \dots & \dots \\ v_{n1} & v_{n2} & \dots & v_{nn} \end{vmatrix}$$

не равны нулю. Тогда существует базис $\{f_1, \dots, f_n\}$, в котором

$$(Vx, x) = \frac{D_0}{D_1} \xi_1^2 + \frac{D_1}{D_2} \xi_2^2 + \dots + \frac{D_{n-1}}{D_n} \xi_n^2, \quad (4.13)$$

где для единообразия записи принято $D_0 = 1$.

◀ Базис $\{f_1, \dots, f_n\}$ через исходный $\{e_1, \dots, e_n\}$ выражается следующим «треугольным» образом:

$$\begin{aligned} f_1 &= e_1, \\ f_2 &= a_{21}e_1 + e_2, \\ f_3 &= a_{31}e_1 + a_{32}e_2 + e_3, \\ &\dots \dots \dots \\ f_n &= a_{n1}e_1 + a_{n2}e_2 + \dots + e_n. \end{aligned}$$

Коэффициенты a_{ij} определяются из условий диагональной записи, $(V'f_i, f_j) = 0$ при $i \neq j$, квадратичной формы в новом базисе. ▶

- Из записи (4.13) следует, что для положительной определенности квадратичной формы достаточна положительность всех главных миноров D_k . Необходимость всех $D_k > 0$ доказывается совсем просто²²⁾, и в совокупности с достаточностью называется *критерием Сильвестра*.
- Скалярное произведение (x, x) можно рассматривать как квадратичную форму (x, Ix) , которая при замене координат $x = Ty$ переходит в (y, T^*Ty) , откуда следует положительная определенность матрицы T^*T , если T невырождена.

Из представления T в виде объединения вектор-строк t_i видно, что $\{ij\}$ -м элементом матрицы T^*T является скалярное произведение (t_i, t_j) . Главные миноры этой матрицы называются *определителями Грама*, и они все строго положительны, когда векторы $\{t_1, \dots, t_n\}$ линейно независимы²³⁾, — поскольку T^*T положительно определена.

- $\text{tr } A^*A = \sum_{ij} |a_{ij}|^2$.
- Теорема Шура о произведении.** Если A и B симметричные положительно определенные матрицы, то матрица²⁴⁾

$$C = [a_{ij}b_{ij}]$$

тоже положительно определена.

²²⁾ Если $D_j = 0$, то существует такой ненулевой $x = \{x_1, \dots, x_j, 0, \dots, 0\}$, что $(Vx, x) = 0$.

²³⁾ Что равносильно невырожденности T .

²⁴⁾ «Наивное» произведение матриц, сводящееся к поэлементному умножению, $A \circ B = [a_{ij}b_{ij}]$, называют *произведением по Адамару*.

◀ Поскольку A приводится к диагональному виду $T^*AT = \Lambda$, то $A = T\Lambda T^*$, в силу чего элементы A можно представить в виде

$$a_{ij} = \sum_k \lambda_k t_{ik} t_{jk}.$$

Поэтому

$$\sum_{i,j} a_{ij} b_{ij} x_i x_j = \sum_{i,j} \left(\sum_k \lambda_k t_{ik} t_{jk} \right) b_{ij} x_i x_j = \sum_k \lambda_k \left(\sum_{i,j} b_{ij} (t_{ik} x_i) (t_{jk} x_j) \right).$$

Окончательный вывод следует из $\lambda_k > 0$, положительной определенности B и невырожденности T . ▶

- Возможность полярного разложения сохраняется в случае прямоугольной матрицы. Любая $m \times n$ ($m \leq n$) матрица A представима в виде $A = ST$, где $S = (AA^*)^{1/2}$ неотрицательно определена, а T имеет ортонормированные строки ($TT^* = I$).
- Пусть случайные величины x_i имеют математические ожидания m_i . Ковариационная матрица R с элементами

$$R_{ij} = M\{(x_i - m_i)(x_j - m_j)\},$$

где M обозначает оператор математического ожидания, — неотрицательно определена:

$$\sum_{ij} R_{ij} \xi_i \xi_j = M \left\{ \sum_{ij} (x_i - m_i) \xi_i \right\}^2 \geq 0.$$

- Пусть m_i из предыдущего пункта равны нулю. При необходимости прогнозирования случайной величины y с помощью линейной комбинации $\sum_i c_i x_i$ часто ориентируются на минимизацию среднеквадратической ошибки

$$M \left\{ \left(y - \sum_i c_i x_i \right)^2 \right\} \rightarrow \min_c.$$

Производные по c_j должны быть равны нулю,

$$\frac{\partial}{\partial c_j} M \left\{ \left(y - \sum_i c_i x_i \right)^2 \right\} = M \left\{ \left(y - \sum_i c_i x_i \right) x_j \right\} = 0.$$

Оптимальный вектор c , таким образом, определяется решением системы

$$Rc = R_{xy},$$

где R — ковариационная матрица, а вектор ковариации R_{xy} имеет координаты $M\{x_i y\}$.

Описанный рецепт называют *методом наименьших квадратов*.

Глава 5

Канонические представления

В первых двух разделах главы рассматриваются простые и важные результаты, широко применяемые в различных ситуациях.

Затем излагаются факты, к элементарной линейной алгебре не относящиеся, и им не имеет смысла уделять много внимания, если математика не планируется в качестве специальности. Речь идет о жордановых формах, традиционно изучаемых чаще, чем это целесообразно.

Сами по себе жордановы формы не так важны, как методы, с помощью которых они изучаются. Поэтому аннулирующие многочлены, корневые подпространства и прочие «фокусы» — это не столько инструменты, необходимые для изучения жордановых форм, сколько эффективные категории мышления при изучении более высоких этажей линейной алгебры.

5.1. Унитарные матрицы

Унитарные матрицы — это комплексный вариант ортогональных.

Матрица U унитарна, если $U^*U = I$. Форма определения такая же, но теперь U^* — не просто транспонированная, а сопряженная¹⁾ матрица U , т. е. $u_{ij}^* = \overline{u_{ji}}$.

Источник интереса к унитарным матрицам — тот же самый, что и к ортогональным. Те и другие возникают обычно как инструмент исследования. Ортогональные матрицы, как матрицы преобразований, удобны благодаря ортогональности столбцов (строк). Оказалось, что их как раз «хватает» при изучении симметричных матриц.

В общем случае ортогональных преобразований уже недостаточно, и унитарные матрицы — следующий шаг. Выгоды, по сути, те же:

- $\langle u_i, u_j \rangle = 0$ при $i \neq j$, и $u_i^2 = 1$ (аналогично для строк u_i);

¹⁾ Эквивалентный термин: *эрмитово сопряженная* матрица.

- евклидова длина преобразованного вектора $\mathbf{x} = U\mathbf{y}$ ($\mathbf{y} = U^*\mathbf{x}$) сохраняется²⁾, $\langle \mathbf{x}, \mathbf{x} \rangle = \langle \mathbf{y}, \mathbf{y} \rangle$;
- $U^* = U^{-1}$;
- все собственные значения по модулю равны единице, $\lambda_k = e^{i\varphi_k}$.

Ортогонализация Грама—Шмидта. При манипуляциях с ортогональными и унитарными матрицами полезной часто оказывается замена некоторого множества линейно независимых векторов $\{\mathbf{e}_1, \dots, \mathbf{e}_n\}$ ортонормированной системой $\{\mathbf{f}_1, \dots, \mathbf{f}_n\}$, которая порождает то же самое пространство³⁾. Это всегда возможно, причем — разными способами. Удобный алгоритм дает стандартный процесс ортонормирования Грама—Шмидта.

◀ На первом шаге полагается $\mathbf{h}_1 = \mathbf{e}_1$ и

$$\mathbf{f}_1 = \frac{\mathbf{h}_1}{\langle \mathbf{h}_1, \mathbf{h}_1 \rangle^{1/2}}.$$

На втором — строится вектор $\mathbf{h}_2 = \mathbf{e}_2 - \langle \mathbf{e}_2, \mathbf{f}_1 \rangle \mathbf{f}_1$, ортогональный $\mathbf{f}_1$, и опять происходит нормировка: $\mathbf{f}_2 = \frac{\mathbf{h}_2}{\langle \mathbf{h}_2, \mathbf{h}_2 \rangle^{1/2}}$.

На k -м шаге строится вектор

$$\mathbf{h}_k = \mathbf{e}_k - \langle \mathbf{e}_k, \mathbf{f}_{k-1} \rangle \mathbf{f}_{k-1} - \dots - \langle \mathbf{e}_k, \mathbf{f}_1 \rangle \mathbf{f}_1,$$

ортогональный ко всем $\{\mathbf{f}_1, \dots, \mathbf{f}_{k-1}\}$ (проверьте), после чего нормируется,

$\mathbf{f}_k = \frac{\mathbf{h}_k}{\langle \mathbf{h}_k, \mathbf{h}_k \rangle^{1/2}}$. На n -м шаге процесс заканчивается. ▶

Обратим внимание, что в случае исходной системы векторов с действительными координатами процесс не выходит в комплексную плоскость, т. е. векторы $\{\mathbf{f}_1, \dots, \mathbf{f}_n\}$ получаются тоже действительными.

Заметим, что на k -м шаге процесса Грама—Шмидта векторы $\{\mathbf{f}_1, \dots, \mathbf{f}_k\}$ линейно выражаются через k первых векторов множества $\{\mathbf{e}_1, \dots, \mathbf{e}_n\}$. По-другому это может быть сформулировано так. Если матрицы X и Y составлены из вектор-столбцов $\{\mathbf{e}\}$ и $\{\mathbf{f}\}$, т. е.

$$X = [\mathbf{e}_1, \dots, \mathbf{e}_n], \quad Y = [\mathbf{f}_1, \dots, \mathbf{f}_n],$$

²⁾ Под скалярным произведением в комплексном пространстве понимается

$$\langle \mathbf{x}, \mathbf{y} \rangle = \sum_i x_i y_i^*.$$

³⁾ В том смысле, что линейные оболочки обеих систем, $\{\mathbf{e}\}$ и $\{\mathbf{f}\}$, совпадают.

то $Y = XQ$, $X = YR$, где Q и R — верхние треугольные матрицы⁴⁾. Этот факт обычно выделяют в самостоятельное утверждение:

5.1.1. Теорема. *Любая действительная невырожденная матрица A может быть представлена в виде произведения $A = QR$ ортогональной матрицы Q на верхнюю треугольную⁵⁾ R .*

Упражнения

- Для унитарности U необходимо и достаточно, чтобы $U^{-1} = U^*$. Отсюда легко следует, что *произведение унитарных матриц унитарно*.
- Детерминант унитарной матрицы равен по модулю единице.
- Матрицы A и B называют *унитарно эквивалентными*, если $A = U^*BU$, где U — унитарная матрица. Если A и B унитарно эквивалентны, то

$$\sum_{ij} |a_{ij}|^2 = \sum_{ij} |b_{ij}|^2.$$

5.2. Триангуляция Шура

5.2.1. Теорема об унитарной триангуляции. *Любая матрица A может быть приведена унитарным преобразованием к треугольному виду:*

$$U^*AU = T, \quad (5.1)$$

где T — *верхняя треугольная матрица*⁶⁾. Если матрица A и все ее собственные значения вещественны, то U может быть выбрана вещественной.

Таким образом, в комплексном пространстве всегда существует базис, в котором A имеет треугольный вид. На диагонали T , понятно, стоят собственные значения матрицы A . Доказательство совсем просто.

◀ Матрица U строится поэтапно. Пусть $\mathbf{f}_1$ обозначает собственный вектор A , отвечающий собственному значению⁷⁾ λ_1 . Дополнив (неважно как) $\mathbf{f}_1$ до базиса

$$\{\mathbf{f}_1, \mathbf{e}_2, \dots, \mathbf{e}_n\}$$

⁴⁾ $Q = R^{-1}$; $q_{ij}, r_{ij} = 0$ при $i > j$.

⁵⁾ Либо наоборот, верхней треугольной на ортогональную. Либо, наконец, в том и другом случае верхнюю треугольную матрицу можно заменить на — нижнюю.

⁶⁾ С тем же успехом — нижняя треугольная матрица.

⁷⁾ Любому собственному значению λ всегда отвечает хотя бы один собственный вектор.

и применив к последнему процесс ортонормирования Грама—Шмидта, придем к ортонормированному базису

$$\{f_1, g_2, \dots, g_n\},$$

в котором матрица A примет вид⁸⁾

$$U_1^* A U_1 = \begin{bmatrix} \lambda_1 & * \\ 0 & A_1 \end{bmatrix},$$

где 0 обозначает нулевой столбец, а * — некую вектор-строку. Наконец, A_1 — это $(n-1) \times (n-1)$ матрица.

Далее аналогичная процедура применяется к матрице A_1 , в результате чего находится такая матрица U_2' , что $U_2'^* A_1 U_2' = \begin{bmatrix} \lambda_2 & * \\ 0 & A_2 \end{bmatrix}$. Тогда

$$U_2^* U_1^* A U_1 U_2 = \begin{bmatrix} \lambda_1 & * & * \\ 0 & \lambda_2 & * \\ 0 & 0 & A_2 \end{bmatrix},$$

где $U_2 = \begin{bmatrix} 1 & 0 \\ 0 & U_2' \end{bmatrix}$. В итоге $U = U_1 \cdots U_{n-1}$ дает необходимое унитарное преобразование.

Если матрица A и все ее собственные значения вещественны, то им отвечают вещественные собственные векторы, и описанные шаги не выводят за пределы вещественных чисел. ►

В случае $U^* A U = B$ говорят об *унитарной эквивалентности* матриц A и B . Понятно, что унитарная эквивалентность является частным случаем подобия. Может быть, например, так: A приводится к диагональному виду, но с помощью унитарных преобразований — только к треугольному.

Зачем это нужно. Триангуляция Шура (теорема 5.2.1) обычно используется двояко: как теоретический инструмент и как алгоритм решения практических задач.

- Вернемся, например, к задаче Коши (раздел 3.2)

$$\dot{x} = Ax, \quad x(0) = x_0.$$

Пусть замена переменных $x = Uy$, с унитарной матрицей U , приводит ее к треугольному виду

$$\dot{y} = U^* A U y, \quad y(0) = U^* x_0,$$

⁸⁾ Потому что в этом базисе $f_1 = \{1, 0, \dots, 0\}^*$.

где

$$U^*AU = \begin{bmatrix} \lambda_1 & \dots & * \\ \vdots & \ddots & \vdots \\ 0 & \dots & \lambda_n \end{bmatrix}.$$

Хотя это и не «развязывает» полностью исходную систему (в отличие от диагонализации), но сводит ее к легкому последовательному интегрированию скалярных уравнений. Сначала интегрируется $\dot{y}_n = \lambda_n y_n$, затем $\dot{y}_{n-1} = \lambda_{n-1} y_{n-1} + \tilde{a}_{(n-1)n} y_n$ и т. д.

- Рассмотрим теперь задачу о малом возмущении матрицы A с кратными собственными значениями, приводящем к матрице $A(\varepsilon)$ с попарно различными собственными значениями (см. раздел 3.4). Пусть U приводит A к треугольному виду, $T = U^*AU$. Добавление к T диагональной матрицы⁹⁾

$$\Lambda(\varepsilon) = \begin{bmatrix} \varepsilon^1 & \dots & 0 \\ \vdots & \ddots & \vdots \\ 0 & \dots & \varepsilon^n \end{bmatrix}$$

дает матрицу $T + \Lambda(\varepsilon)$, имеющую при достаточно малых $\varepsilon \neq 0$ попарно различные собственные значения. Те же собственные значения имеет $A(\varepsilon) = U(T + \Lambda(\varepsilon))U^*$, и $A(\varepsilon) \rightarrow A$ при $\varepsilon \rightarrow 0$.

- Из треугольного представления матрицы сразу следует:

$$\det A = \lambda_1(A) \cdots \lambda_n(A), \quad \operatorname{tr} A = \lambda_1(A) + \dots + \lambda_n(A).$$

Рассмотренные примеры демонстрируют два вида преимуществ, которые дает треугольное представление. На диагонали у T стоят собственные значения матрицы — и это позволяет работать с T , почти как с диагональной матрицей. В частности, во втором примере возмущение диагонали было в точности возмущением спектра, а внедиагональные элементы никакой роли не играли. В первом примере главное было в другом — в нулях под диагональю, что позволило с комфортом интегрировать дифференциальные уравнения одно за другим.

Другие варианты. При желании обойтись вещественными преобразованиями можно ориентироваться на следующий результат.

5.2.2. В случае вещественной матрицы A всегда существует вещественная ортогональная матрица Q , преобразующая A к «почти

⁹⁾ Здесь верхние индексы являются показателями степени.

треугольному» виду

$$\begin{bmatrix} A_1 & & * \\ & A_2 & \\ & & \ddots \\ 0 & & & A_p \end{bmatrix},$$

где каждый диагональный блок A_j отвечает либо вещественному собственному значению — и тогда это просто скаляр, либо — паре комплексно сопряженных — и тогда $A_j = \begin{bmatrix} \alpha_j & -\beta_j \\ \beta_j & \alpha_j \end{bmatrix}$.

Результат поначалу воодушевляет, но наличие ненулевой «поддиагонали» все же портит малину, обесценивая достигнутое упрощение.

Существенно более полезен другой вариант триангуляции, но опять комплексный и, плюс к тому, уже не унитарный.

5.2.3. Пусть $\lambda_1, \dots, \lambda_k$ обозначают различные собственные значения матрицы A кратности n_j . Тогда A подобна¹⁰⁾ блочно-диагональной матрице

$$\begin{bmatrix} A_1 & & 0 \\ & A_2 & \\ & & \ddots \\ 0 & & & A_p \end{bmatrix},$$

где A_j — верхние треугольные матрицы размера $n_j \times n_j$ с диагональными элементами λ_j .

Результат будет, по существу, доказан в разделе 5.5. Техническое обоснование, не выходящее за пределы элементарных категорий линейной алгебры, можно найти в [18].

5.3. Жордановы формы

При переходе к треугольной форме матрицы возможен, образно говоря, миллион вариантов. Слишком легкая задача. Ее можно решать «и так и этак». Поэтому возникает желание распорядиться свободой выбора с целью извлечения дополнительных выгод.

¹⁰⁾ Но не обязательно «унитарно подобна».

На этом пути оказывается, что любая матрица A заменой базиса, $A = S^{-1}JS$, может быть приведена к *жордановой форме*:

$$J = \begin{bmatrix} J_{k_1}(\lambda_1) & & 0 \\ & \ddots & \\ 0 & & J_{k_p}(\lambda_p) \end{bmatrix}, \quad (5.2)$$

где у *жордановой матрицы* J на диагонали стоят *жордановы клетки* $J_k(\lambda)$,

$$J_k(\lambda) = \begin{bmatrix} \lambda & 1 & & 0 \\ & \ddots & \ddots & \\ & & \ddots & 1 \\ 0 & & & \lambda \end{bmatrix},$$

в которых на главной диагонали повторяется собственное значение λ , над главной диагональю — диагональ из единиц, все остальные элементы — нули¹¹⁾.

В (5.2) $k_1 + \dots + k_p = n$ и собственные значения λ_j в разных клетках не обязательно различны. С точностью до перестановки блоков $J_k(\cdot)$ — форма (5.2) определяется однозначно. Конкретный вид J для A связан со следующими факторами:

- Число жордановых клеток J равно максимальному числу линейно независимых собственных векторов¹²⁾ A .
- Сумма порядков жордановых клеток, отвечающих одному и тому же собственному значению λ_j , равна алгебраической кратности λ_j , а число этих клеток равно геометрической кратности λ_j .

Перечисленных факторов, однако, не хватает для полной определенности. Проблема будет «дожата» в следующих разделах.

Никаких особых выгод жорданова форма (по сравнению с треугольной) не дает, а по «литрам испорченной крови» занимает лидирующее место. К тому же, малейшее изменение матрицы способно изменить ее жорданову форму. В таких случаях обычно говорят о *структурной неустойчивости* изучаемого объекта. Ситуация, понятно, не из приятных. Сколь угодно малое изменение условий задачи — кардинально меняет ответ. Какой смысл тогда решать, казалось бы.

Проблема, однако, не так проста. Вычисление жордановой формы для конкретной матрицы практического смысла, конечно, не имеет, если речь не идет

¹¹⁾ У жордановой клетки размера 1×1 кроме λ нигде, разумеется, ничего не стоит, потому что нигде.

¹²⁾ При $p = n$ матрица J диагональна.

об упражнении учебного характера. Смысл имеет изучение самой задачи с целью выявления природы аномальных явлений, когда до них доходит дело.

Подводная часть айсберга. Тот факт, что не всякая матрица приводится к диагональному виду, становится очевиден из самых простых примеров. Матрица $A = \begin{bmatrix} 1 & 1 \\ 0 & 1 \end{bmatrix}$ имеет два равных собственных значения

$\lambda_{1,2} = 1$, но она не может быть подобна $I = \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix}$, иначе

$$S^{-1}AS = I \Rightarrow A = SIS^{-1} = I$$

приводило бы к противоречию¹³⁾.

Причина, как уже отмечалось (раздел 3.2), в пресловутой кратности собственных значений. Механизм возникновения «неприятностей» наглядно демонстрирует модель перехода к A с помощью

$A(\epsilon) \rightarrow A$ при $\epsilon \rightarrow 0$, где у матриц $A(\epsilon)$ все собственные значения различны ($\epsilon \neq 0$). По теореме 3.3.1 $A(\epsilon)$ имеет n линейно независимых собственных векторов

$$f(\epsilon) = \{f_1(\epsilon), \dots, f_n(\epsilon)\}, \quad (5.3)$$

которые могут становиться линейно зависимыми — в случае появления в спектре $\sigma(A)$ (при $\epsilon \rightarrow 0$) кратных λ .

Дело в том, что ранг матрицы $A - \lambda I$ никак не связан с кратностью λ как корня характеристического уравнения $|A - \lambda I| = 0$. Единственное, что можно утверждать, если λ собственное значение, что $\text{rang}(A - \lambda I) < n$. Поэтому ненулевую разрешимость $(A - \lambda I)x = 0$ (существование собственного вектора) можно гарантировать, но о количестве линейно независимых решений — полином $|A - \lambda I| = 0$ информации не дает.

Разнообразие возможностей при $f(\epsilon) \rightarrow f$ достаточно велико. Векторы $\{f_1(\epsilon), \dots, f_n(\epsilon)\}$ распадаются на подгруппы, соответствующие подмножествам собственных значений, которые уравниваются при $\epsilon \rightarrow 0$. Внутри подгруппы — пусть для определенности

¹³⁾ Но $A = \begin{bmatrix} 1-\epsilon & 1 \\ 0 & 1 \end{bmatrix}$ подобна $I = \begin{bmatrix} 1-\epsilon & 0 \\ 0 & 1 \end{bmatrix}$ при сколь угодно малом $\epsilon \neq 0$ (теорема 3.3.1).

это будет $\{f_1(\varepsilon), \dots, f_k(\varepsilon)\}$ — собственные векторы в пределе могут остаться линейно независимыми, а могут стать линейно зависимыми. При этом размерность линейной оболочки $\{f_1, \dots, f_k\}$ может уменьшиться вплоть до единицы¹⁴⁾, становясь в общем случае равной некоторому m в пределах от 1 до k . Число¹⁵⁾ $k - m + 1$ показывает, сколько линейно независимых собственных векторов отвечает данному собственному значению, но этим разнообразие не исчерпывается, поскольку вариантов «вырождения» — укладывания k векторов в m -мерное подпространство — тоже может быть несколько — векторы могут распадаться на подгруппы следующего уровня, вырождаясь там на свой манер (становясь коллинеарными, компланарными и т. п.).

Варианты друг от друга отличаются структурой линейной «зависимости-независимости» собственных векторов. Матрицы A и B , имеющие одинаковую структуру, переводятся друг в друга невырожденным преобразованием, $A = T^{-1}BT$.

Все это в какой-то мере проясняет существо возможных явлений, но от понимания сути до формального алгоритма имеется некоторое расстояние. Проблема с более конструктивной точки зрения рассматривается в следующих разделах.

О приложениях. Жордановы формы оказываются полезны¹⁶⁾ каждый раз, когда задаются максималистские вопросы, требующие предельно точных ответов.

В каком случае матрицы A и B подобны? — Если их спектры не содержат кратных значений и совпадают. — А если есть кратные собственные значения? — Тогда возможен универсальный ответ. Если совпадают жордановы формы A и B , — разумеется, с точностью до порядка расположения жордановых клеток.

Триангуляция Шура здесь бы уже не помогла. Другое дело, что циник бы сказал, что максималистские вопросы никого не интересуют, кроме диссертантов. Однако циники, заслуживая всяческого уважения, все же не исчерпывают Вселенной.

¹⁴⁾ То есть векторы $\{f_1(\varepsilon), \dots, f_k(\varepsilon)\}$ могут стать в пределе коллинеарными. Либо компланарными — предельные $\{f_1, \dots, f_k\}$ попадают в одну плоскость. Либо уложиться в пределе в трехмерное пространство и т. д.

¹⁵⁾ Геометрическая кратность собственного значения.

¹⁶⁾ Это компенсация высказанных выше скептических замечаний.

Подобна ли матрица своей транспонированной? Наивный вроде бы вопрос. Но положительный ответ очевиден лишь в случае «хорошего» спектра. А если в спектре кратные значения?

Выручают жордановы формы. Клетка транспонируется легко:

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

Поэтому транспонируется и вся жорданова матрица, $J^* = P^{-1}JP$, а значит, и $A = T^{-1}JT$, если речь идет о вещественной матрице¹⁷⁾.

5.4. Аннулирующий многочлен

В силу конечномерности R^n множество

$$x, Ax, \dots, A^k x,$$

при некотором минимальном k , становится линейно зависимым, что означает существование таких констант ν_j , что¹⁸⁾

$$A^k x + \nu_{k-1} A^{k-1} x + \dots + \nu_1 x = 0.$$

Последнее равенство может быть записано в виде $\varphi(A)x = 0$, где $\varphi(\lambda)$ обозначает многочлен

$$\varphi(\lambda) = \lambda^k + \nu_{k-1} \lambda^{k-1} + \dots + \nu_1,$$

который называют *аннулирующим многочленом*¹⁹⁾ вектора x .

Аннулирующими в смысле $\psi(A)x = 0$ могут быть и другие многочлены $\psi(\lambda)$, но из построения ясно, что $\varphi(\lambda)$ имеет наименьшую степень, и потому — любой аннулирующий многочлен $\psi(\lambda)$ делится на $\varphi(\lambda)$ без остатка.

Действительно, из $\psi(\lambda) = \zeta(\lambda)\varphi(\lambda) + r(\lambda)$ следует, в силу $\psi(A)x = 0$,

$$\psi(A)x = \zeta(A)\varphi(A)x + r(A)x = r(A)x \Rightarrow r(A)x = 0.$$

Пусть теперь $\varphi_j(\lambda)$ обозначает минимальный аннулирующий многочлен базисного вектора e_j . Наименьшее общее кратное $\varphi_A(\lambda)$

¹⁷⁾ Комплексная матрица тоже подобна своей транспонированной (без сопряжения).

¹⁸⁾ В частности, все ν_j могут быть равны нулю, если $A^k x = 0$.

¹⁹⁾ Для рассматриваемого оператора A .

многочленов $\varphi_1(\lambda), \dots, \varphi_n(\lambda)$ — называется *минимальным многочленом оператора* ²⁰⁾ A (иногда говорят: *минимальный аннулирующий*). Очевидно, $\varphi_A(A)x = 0$ для любого вектора

$$x = x_1 e_1 + \dots + x_n e_n,$$

и потому $\varphi_A(A)$ представляет собой матрицу из одних нулей.

5.4.1. Теорема. Если минимальный многочлен оператора A представим в виде произведения двух **взаимно простых** многочленов,

$$\varphi_A(\lambda) = \varphi_1(\lambda)\varphi_2(\lambda),$$

то R^n представимо в виде прямой суммы (расщепляется на сумму) двух инвариантных ²¹⁾ для A подпространств X_1 и X_2 .

◀ Легко видеть, что $X_1 = \ker \varphi_1(A)$, $X_2 = \ker \varphi_2(A)$. Действительно, из $x \in X_j$ и перестановочности A с любым своим многочленом следует

$$\varphi_j(A)x = 0 \Rightarrow A\varphi_j(A)x = 0 \Rightarrow \varphi_j(A)Ax = 0 \Rightarrow Ax \in X_j,$$

т. е. $AX_j \subset X_j$.

Остается показать, что $X_1 \cap X_2 = 0$. В противном случае существовал бы ненулевой вектор x такой, что $\varphi_1(A)x = 0$ и $\varphi_2(A)x = 0$, но это противоречило бы взаимной простоте многочленов $\varphi_1(\lambda)$ и $\varphi_2(\lambda)$, ибо тогда $\varphi_1(\lambda)$ и $\varphi_2(\lambda)$ были бы представимы в виде $\varphi_j(\lambda) = \varphi(\lambda)\psi_j(\lambda)$ с общим множителем, равным минимальному аннулирующему многочлену $\varphi(\lambda)$ вектора x . ▶

Разумеется, теорема 5.4.1 влечет за собой более общий факт: если минимальный многочлен оператора $\varphi_A(\lambda)$ представим в виде произведения p взаимно простых многочленов, то R^n расщепляется на сумму p инвариантных для A подпространств $X_1, \dots, X_p$.

5.5. Корневые подпространства

Поскольку речь идет о комплексных λ , то минимальный многочлен $\varphi_A(\lambda)$ разлагается в произведение ²²⁾

$$\varphi_A(\lambda) = (\lambda - \lambda_1)^{s_1} \dots (\lambda - \lambda_p)^{s_p}, \quad (5.4)$$

²⁰⁾ Предполагается, что все рассматриваемые многочлены имеют старший коэффициент, равный единице.

²¹⁾ Напомним, что подпространство $X \subset R^n$ называется *инвариантным для оператора A* , если $AX \subset X$.

²²⁾ Все λ_j попарно различны.

и по теореме 5.4.1 множителям $(\lambda - \lambda_j)^{g_j}$ отвечают инвариантные подпространства X_j , называемые **корневыми**, на которые расщепляется R^n .

5.5.1. *Корнями минимального многочлена (5.4) служат собственные значения матрицы A , и только они.*

◀ Аннулирующим для собственного вектора f_j ($Af_j = \lambda_j f_j$) является многочлен $\lambda - \lambda_j = 0$. Поэтому все собственные значения A являются корнями $\varphi_A(\lambda)$.

Теперь в обратную сторону. Пусть среди λ_j в (5.4) есть корень, не являющийся собственным значением A . Тогда сужение A на соответствующее корневое подпространство — было бы линейным оператором, не имеющим собственных значений. ►

Выбор в качестве базиса R^n объединения базисов корневых подпространств X_j приводит к блочно-диагональной записи матрицы A

$$\begin{bmatrix} A_1 & & & 0 \\ & A_2 & & \\ & & \ddots & \\ 0 & & & A_p \end{bmatrix},$$

где блоки A_j имеют собственное значение λ_j , кратность которого совпадает с размерностью X_j .

Далее каждым блоком A_j можно заниматься независимо от других блоков²³⁾, меняя базис X_j и не затрагивая базисы других подпространств X_k , $k \neq j$. Поэтому, чтобы не загромождать изложение, индекс j далее опускается.

Итак, можно считать, что речь идет о матрице A с совпадающими собственными значениями, равными λ_0 , где λ_0 обозначает одно из λ_j , а $\{e, f, \dots, h\}$ — собственные векторы, число k которых определяется условием

$$\text{rang}(A - \lambda_0 I) = r - k,$$

если A имеет размер $r \times r$.

Перейдем к построению базиса. В качестве e_1 возьмем собственный вектор e . Очевидно, $(A - \lambda_0 I)e_1 = 0$. Затем в качестве e_2

²³⁾ Результирующее преобразование A определится матрицей $T = \text{diag}\{T_1, \dots, T_p\}$, где блоки T_j влияют только на A_j .

возьмем решение уравнения $(A - \lambda_0 I)e_2 = e_1$. Далее продолжаем в том же духе, пока решаются уравнения

$$(A - \lambda_0 I)e_q = e_{q-1}. \quad (5.5)$$

Почему и до каких пор решаются эти уравнения, и почему в результате получается базис, мы обсудим позже, а пока остановимся на том, как будет выглядеть матрица в этом базисе. Запись (5.5) в виде

$$Ae_q = \lambda_0 e_q + e_{q-1}$$

сразу показывает, что в этом базисе A приобретает вид

$$\begin{bmatrix} J_m(\lambda_0) & 0 \\ 0 & \tilde{A} \end{bmatrix}.$$

Затем берется следующий собственный вектор f , проводится такая же манипуляция, и т. д. В итоге матрица (блок) A приобретает жорданову форму. Процесс не может остановиться в ситуации, когда в правом нижнем углу стоит не жорданова клетка $\tilde{A}'$, а собственные векторы уже кончились. Тогда бы матрица $\tilde{A}'$ не имела собственных векторов, чего не бывает.

◀ Теперь подготовим почву для обоснования. Начнем с рассмотрения последовательности подпространств $K_{\lambda_0}^{(q)} = \ker(A - \lambda_0 I)^q$. Очевидно, $K_{\lambda_0}^{(1)}$ — есть так называемое *собственное подпространство*, базис которого составляют собственные векторы $\{e, f, \dots, h\}$ (и любой вектор которого тоже является собственным).

Если $(A - \lambda_0 I)x = 0$, то тем более $(A - \lambda_0 I)^2 x = 0$, поэтому $K_{\lambda_0}^{(1)} \subset K_{\lambda_0}^{(2)}$. Действуя аналогично, получаем

$$\boxed{K_{\lambda_0}^{(1)} \subset K_{\lambda_0}^{(2)} \subset \dots \subset K_{\lambda_0}^{(p)},} \quad (5.6)$$

где включения *строгие* до некоторого $p \geq 1$, начиная с которого включения превращаются в равенства.

Содержательный вариант $p > 1$ имеет место в том случае, когда у минимального многочлена $(\lambda - \lambda_0)^{s_0}$ матрицы²⁴⁾ A показатель степени $s_0 > 1$. Причина заключается в том, что минимальный многочлен аннулирует матрицу, т. е. $(A - \lambda_0 I)^{s_0} = 0$, — и поэтому $K_{\lambda_0}^{(p)}$ при $p = s_0$ совпадает со всем пространством, в котором действует A . Следовательно, цепочка (5.6) должна быть

²⁴⁾ То есть у соответствующего множителя $(\lambda - \lambda_j)^{s_j}$ в (5.4).

возрастающей (по размерности) от $K_{\lambda_0}^{(1)}$ до $K_{\lambda_0}^{(s_0)}$. Причем затормозить в середине, $K_{\lambda_0}^{(q-1)} = K_{\lambda_0}^{(q)}$, она не может, поскольку в этом случае дальнейший рост размерности остановился бы. (?)

Наконец, заметим, что подпространства $K_{\lambda_0}^{(q)}$ инвариантны относительно A , что вытекает из очевидной импликации

$$(A - \lambda_0 I)^q x = 0 \Rightarrow A(A - \lambda_0 I)^q x = (A - \lambda_0 I)^q Ax = 0,$$

поскольку A коммутирует с $(A - \lambda_0 I)^q$. ►

Перейдем к обоснованию описанного выше алгоритма приведения матрицы к жордановой форме. Что касается уравнений (5.5), то их разрешимость очевидна, в силу

$$(A - \lambda_0 I)K_{\lambda_0}^{(q)} = K_{\lambda_0}^{(q-1)}, \quad (5.7)$$

причем ясно, что $e_q \in K_{\lambda_0}^{(q)} \setminus K_{\lambda_0}^{(q-1)}$, — и потому векторы $\{e_1, \dots, e_p\}$ линейно независимы.

Как только $K_{\lambda_0}^{(p)} = K_{\lambda_0}^{(p+1)}$, процесс останавливается. Линейная оболочка построенного базиса представляет инвариантное подпространство относительно A . Действительно, $x = \sum_q x_q e_q$ влечет за собой

$$Ax = \sum_q x_q A e_q = \sum_q (\lambda x_q + x_{q+1}) A e_q,$$

где $x_{p+1} = 0$. Далее процесс продолжается указанным выше образом. Берется следующий собственный вектор f и т. д.

Нильпотентные матрицы. Матрица A называется *нильпотентной*, если $A^k = 0$ при некотором²⁵⁾ k . Жорданова клетка может быть записана в виде $J_k = \lambda I + N_k$, где N_k нильпотентна, $N_k^k = 0$,

$$J_k(\lambda) = \begin{bmatrix} \lambda & 1 & & 0 \\ & \ddots & \ddots & \\ & & \ddots & 1 \\ 0 & & & \lambda \end{bmatrix} = \begin{bmatrix} \lambda & & & 0 \\ & \ddots & & \\ & & \ddots & \\ 0 & & & \lambda \end{bmatrix} + \begin{bmatrix} 0 & 1 & & 0 \\ & \ddots & \ddots & \\ & & \ddots & 1 \\ 0 & & & 0 \end{bmatrix}.$$

Отсюда ясно, что жорданова матрица в общем случае представима в виде суммы, $J = \Lambda + N$, где Λ диагональная матрица, а N — нильпотентная, причем Λ и N коммутируют²⁶⁾. В случае произвольной матрицы,

$$A = T^{-1} J T = T^{-1} \Lambda T + T^{-1} N T = A_\Lambda + A_N,$$

²⁵⁾ Показатель нильпотентности k не превышает порядка матрицы. (?)

²⁶⁾ Поскольку $J = \lambda I + N$, а единичная матрица I коммутирует с любой.

A_Δ диагонализируема, A_N нильпотентна, и они опять-таки коммутируют.

Из представления $A = A_\Delta + A_N$ ясно, что для сходимости A^m к нулевой матрице при $m \rightarrow \infty$ необходимо и достаточно условие $\rho(A) < 1$, т. е.

$$\max_j |\lambda_j| < 1.$$

Упражнения

- Матрицу A называют *простой*, если геометрическая кратность любого ее собственного значения равна единице, т. е. $\text{rank}(A - \lambda_j) = n - 1$. Число жордановых клеток простой матрицы равно числу различных собственных значений.
- Матрица A диагонализируема в том и только том случае, когда все корни минимального многочлена имеют кратность 1, т. е.

$$\varphi_A(\lambda) = (\lambda - \lambda_1) \cdots (\lambda - \lambda_p),$$

где все λ_j попарно различны. *Идемпотентная* матрица, определяемая условием $A^2 = A$, имеет минимальный многочлен $\lambda(\lambda - 1)$, — поэтому диагонализируема.

- Характеристический многочлен нильпотентной матрицы имеет вид

$$p_A(\lambda) = \lambda^n.$$

5.6. Теорема Гамильтона—Кэли

Излагаемые ниже результаты примыкают к рассматриваемой тематике, но имеют в большей степени эстетическое значение²⁷⁾.

У пространства $n \times n$ матриц — размерность n^2 , поэтому множество

$$A, A^2, \dots, A^N$$

при $N > n^2$ линейно зависимо, что гарантирует для A существование аннулирующего полинома степени N .

Теорема Гамильтона—Кэли дает намного более экономный результат.

²⁷⁾ Имеется в виду, что теорема Гамильтона—Кэли применяется довольно редко.

5.6.1. Теорема. Любая матрица удовлетворяет своему характеристическому уравнению.

В учебниках популярны предостережения от ошибочных доказательств этого результата²⁸⁾, что нередко инициирует слишком подробные рассуждения. На самом деле возможно совсем простое обоснование.

◀ Пусть

$$p_A(\lambda) = \lambda^n + \gamma_{n-1}\lambda^{n-1} + \dots + \gamma_0$$

обозначает характеристический полином, все собственные значения λ_j матрицы A попарно различны и T приводит A к диагональному виду. Тогда

$$p_A(A) = p_A(TAT^{-1}) = Tp_A(A)T^{-1} = [0],$$

ибо $p_A(A)$ — диагональная матрица с диагональными элементами $p_A(\lambda_j) = 0$.

Прямолинейная попытка достичь общего результата за счет предельного перехода, $A(\epsilon) \rightarrow A$, не проходит²⁹⁾, но это надо делать просто чуть иначе, избегая участия в рассуждениях матрицы перехода T_ϵ .

Попарно различные собственные значения при малых $\epsilon \neq 0$ гарантированно имеет матрица $A + \epsilon\Delta A$ (см. раздел 5.2). Поэтому из доказанного выше следует

$$p_{A+\epsilon\Delta A}(A + \epsilon\Delta A) = (A + \epsilon\Delta A)^n + \gamma_{n-1}(\epsilon)(A + \epsilon\Delta A)^{n-1} + \dots + \gamma_0(\epsilon)I = [0].$$

Но, очевидно, $p_{A+\epsilon\Delta A}(A + \epsilon\Delta A) = p_A(A) + o(\epsilon)$, что оправдывает предельный переход при $\epsilon \rightarrow 0$. ▶

Возможно и чисто алгебраическое доказательство без обращения к топологическим соображениям непрерывности. Подстановка $A = U^*SU$, где S — треугольная матрица, в $p_A(\lambda) = (\lambda - \lambda_1) \dots (\lambda - \lambda_n)$ вместо λ — приводит к

$$p_A(A) = U^*(S - \lambda_1 I) \dots (S - \lambda_n I)U.$$

У треугольной матрицы $S - \lambda_j I$ на диагонали j -й элемент равен нулю (поскольку $s_{jj} = \lambda_j$). Произведение n таких матриц, $(S - \lambda_1 I) \dots (S - \lambda_n I)$, дает чисто нулевую матрицу, но в этом надо убедиться, что можно использовать в качестве упражнения.

5.7. λ -матрицы

Матрицу

$$A(\lambda) = \begin{bmatrix} a_{11}(\lambda) & \dots & a_{1n}(\lambda) \\ \dots & \dots & \dots \\ a_{n1}(\lambda) & \dots & a_{nn}(\lambda) \end{bmatrix},$$

²⁸⁾ От наивного «Подставим A в $|A - \lambda I| = 0$ вместо λ — получим $|A - A| = 0$ » — до изощренных, но несостоятельных манипуляций.

²⁹⁾ При $\epsilon \rightarrow 0$ матрица T_ϵ в $T_\epsilon A_\epsilon T_\epsilon^{-1}$ вырождается, в результате чего некоторые элементы T_ϵ^{-1} стремятся к бесконечности.

элементы которой являются многочленами от λ , называют λ -матрицей. На λ -матрицы можно смотреть как на матричные многочлены,

$$A(\lambda) = B_0\lambda^k + B_1\lambda^{k-1} + \dots + B_{k-1}\lambda^1 + B_k.$$

Основной источник интереса к данной области — матрица $A - \lambda I$. Естественное ожидание, что канонизация $A - \lambda I$ связана с канонизацией самой матрицы A , оправдывается, и путь к жордановым формам оказывается даже короче и проще. Но этот путь, алгебраический по природе, менее нагляден. Ясно и просто теория λ -матриц изложена у Куроша [11]. Остановимся здесь лишь на стержневых моментах.

Элементарным преобразованием $A(\lambda)$ называется любая конечная последовательность четырех типов операций:

- *умножение строки (столбца) матрицы на ненулевой скаляр;*
- *прибавление строки (столбца) к другой (другому).*

По сравнению с элементарными преобразованиями обычных матриц — здесь добавлена возможность «то же самое» делать со столбцами. Результаты получаются аналогичные (см. раздел 2.7).

Центральный результат: *любая λ -матрица элементарным преобразованием приводится к каноническому виду*

$$E(\lambda) = \begin{bmatrix} e_1(\lambda) & & & 0 \\ & e_2(\lambda) & & \\ & & \ddots & \\ 0 & & & e_n(\lambda) \end{bmatrix},$$

где у диагональной матрицы $E(\lambda)$ многочлен $e_j(\lambda)$ ($j = 2, \dots, n$) делится (без остатка) на $e_{j-1}(\lambda)$, и старшие коэффициенты всех $e_j(\lambda)$, в том числе $e_1(\lambda)$, равны единице, если многочлен не нулевой.

Две λ -матрицы называют эквивалентными, если они приводятся к одному и тому же каноническому виду. В итоге все λ -матрицы распадаются на классы эквивалентных между собой (это, разумеется, теорема). Классы эквивалентности однозначно характеризуются последовательностью многочленов $e_1(\lambda), \dots, e_n(\lambda)$, называемых *инвариантными множителями*, которые, в свою очередь, определяются как

$$e_j(\lambda) = \frac{d_j(\lambda)}{d_{j-1}(\lambda)},$$

где $d_j(\lambda)$ — наибольший общий делитель всех миноров j -го порядка³⁰⁾ матрицы $A(\lambda)$.

³⁰⁾ Понятно, что $d_j(\lambda)$, таким образом, определяется однозначно и конструктивно матрицей $A(\lambda)$. Факт неизменности $d_j(\lambda)$ при элементарных преобразованиях $A(\lambda)$ очевиден. (?)

Теорию «заземляет» следующий принципиальный факт.

Числовые матрицы A и B подобны в том и только том случае, когда их характеристические матрицы $A - \lambda I$ и $B - \lambda I$ эквивалентны, т. е. приводятся к одному и тому же каноническому виду (имеют одни и те же инвариантные множители).

5.8. Задачи и дополнения

- Если матрицы A и B коммутируют, т. е. $AB = BA$, то существует унитарное преобразование U , одновременно приводящее A и B к треугольному виду. (?) *Подсказка.* Поскольку у коммутирующих матриц всегда есть общий собственный вектор (?), то такой вектор можно использовать на каждом шаге процедуры, описанной в доказательстве теоремы 5.2.1.
- Из предыдущего утверждения легко вытекает, что собственные значения суммы $A + B$ коммутирующих матриц равны $\lambda_j + \mu_j$, где λ_j, μ_j — определенным образом упорядоченные собственные значения матриц A и B .
- Треугольную форму матрицы всегда можно выбрать такой, что ее наддиагональные элементы будут сколь угодно малы.
 ◀ Действительно, у треугольной матрицы

$$T = S^{-1}AS$$

после дополнительного преобразования $\Lambda^{-1}T\Lambda$, где

$$\Lambda = \text{diag} \{N, N^2, \dots, N^n\},$$

наддиагональные элементы уменьшаются в N^{j-i} раз, и выбором достаточно большого N можно добиться желаемого результата. ►

Надо лишь отдавать себе отчет, что успех этот — иллюзорный³¹⁾, хотя и не всегда. Примерно такой же, как увеличение зарплаты за счет измерения ее в копейках.

- Как правило, указание базиса, в котором матрица имеет тот или иной вид, несущественно. Важно лишь, что базис и матрица перехода к нему существуют. С этой точки зрения эффективен следующий алгоритм. Ищутся все собственные значения, после чего анализируются ранги

$$\text{rank} (A - \lambda_j I)^k, \quad k = 1, 2, \dots,$$

что позволяет определить число и размеры жордановых клеток для каждого собственного значения λ_j . Для указания жордановой формы больше ничего не нужно.

- Невырожденные жордановы клетки у AB и BA одинаковы. (?)
- Подобные матрицы имеют один и тот же минимальный многочлен. Обратное неверно. (?)

³¹⁾ Из-за того, что $\|\Lambda\| \rightarrow \infty$ при $\varepsilon \rightarrow 0$.

- Оператор дифференцирования D , действующий в пространстве многочленов третьей степени в базисе $\{1, t, t^2, t^3\}$ записывается как

$$D = \begin{bmatrix} 0 & 1 & 0 & 0 \\ 0 & 0 & 2 & 0 \\ 0 & 0 & 0 & 3 \\ 0 & 0 & 0 & 0 \end{bmatrix}.$$

Жорданова форма J_D и матрица перехода T к соответствующему базису здесь:

$$J_D = \begin{bmatrix} 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \\ 0 & 0 & 0 & 0 \end{bmatrix}, \quad T^{-1} = \begin{bmatrix} 6 & 0 & 0 & 0 \\ 0 & 6 & 0 & 0 \\ 0 & 0 & 3 & 0 \\ 0 & 0 & 0 & 1 \end{bmatrix}.$$

- При возведении в степень жорданова клетка эволюционирует следующим образом

$$\begin{bmatrix} 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \\ 0 & 0 & 0 & 0 \end{bmatrix}^2 = \begin{bmatrix} 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \end{bmatrix}, \quad \begin{bmatrix} 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \\ 0 & 0 & 0 & 0 \end{bmatrix}^3 = \begin{bmatrix} 0 & 0 & 0 & 1 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \end{bmatrix},$$

четвертая степень оказывается уже чисто нулевой матрицей.

При замене единиц другими числами характер поведения степеней не меняется,

$$\begin{bmatrix} 0 & 7 & 0 & 0 \\ 0 & 0 & 5 & 0 \\ 0 & 0 & 0 & 2 \\ 0 & 0 & 0 & 0 \end{bmatrix}^2 = \begin{bmatrix} 0 & 0 & 35 & 0 \\ 0 & 0 & 0 & 10 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \end{bmatrix}, \quad \begin{bmatrix} 0 & 7 & 0 & 0 \\ 0 & 0 & 5 & 0 \\ 0 & 0 & 0 & 2 \\ 0 & 0 & 0 & 0 \end{bmatrix}^3 = \begin{bmatrix} 0 & 0 & 0 & 70 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \end{bmatrix}.$$

- Матрица

$$A = \begin{bmatrix} 0 & \dots & 0 & a_0 \\ 1 & \dots & \dots & \dots \\ 0 & \dots & \dots & \dots \\ \dots & \dots & \dots & 0 \\ \dots & \dots & \dots & a_{n-2} \\ 0 & \dots & 0 & 1 & a_{n-1} \end{bmatrix}$$

называется *сопровождающей матрицей многочлена*

$$q(\lambda) = \lambda^n - a_{n-1}\lambda^{n-1} - \dots - a_1\lambda - a_0,$$

который, со своей стороны, служит минимальным многочленом A .

Например, для

$$q(\lambda) = \lambda^3 - 5\lambda^2 + 8\lambda - 4 = (\lambda - 2)^2(\lambda - 1)$$

матрица A имеет вид $\begin{bmatrix} 0 & 0 & 4 \\ 1 & 0 & -8 \\ 0 & 1 & 5 \end{bmatrix}$. Жорданова форма A и матрица перехода T к соответствующему базису здесь:

$$J = \begin{bmatrix} 1 & 0 & 0 \\ 0 & 2 & 1 \\ 0 & 0 & 2 \end{bmatrix}, \quad T^{-1} = \begin{bmatrix} 4 & 2 & -3 \\ -4 & -3 & 4 \\ 1 & 1 & -1 \end{bmatrix}.$$

- Матрица подобна сопровождающей матрице своего характеристического многочлена в том и только том случае, когда характеристический многочлен является ее минимальным многочленом.
- Матрицу A называют *нормальной*, если она коммутирует со своей сопряженной, т. е. $AA^* = A^*A$. Интерес к нормальным матрицам объясняется тем, что они, обладая уникальными свойствами, объединяют «под одной крышей» унитарные матрицы,

$$U^*U = I = UU^*,$$

эрмитовы (определяемые условием $A = A^*$), *косозермитовы* ($A = -A^*$) и некоторые другие.

Нормальная матрица A унитарно диагонализуема, имеет ортонормированное множество собственных векторов,

$$\sum_{ij} |a_{ij}|^2 = \sum_j |\lambda_j|^2,$$

наконец, всегда существует ортогональная матрица Q , приводящая вещественную нормальную матрицу к блочно-диагональному виду

$$\begin{bmatrix} A_1 & & 0 \\ & \ddots & \\ 0 & & A_p \end{bmatrix}, \quad A_j = \begin{bmatrix} \alpha_j & -\beta_j \\ \beta_j & \alpha_j \end{bmatrix}.$$

- Если геометрическая кратность собственного значения λ матрицы A порядка n больше или равна k , то λ является собственным значением любой главной подматрицы A порядка $m > n - k$. (?)

Глава 6

Функции от матриц

6.1. Матричные ряды

Полином от матрицы $\gamma_0 I + \gamma_1 A + \dots + \gamma_l A^l$ дает простейший пример матричной функции. Следующий шаг — ряд

$$f(A) = \sum_{k=0}^{\infty} \gamma_k A^k, \quad A^0 = I. \quad (6.1)$$

Рассмотрение бесконечных сумм (6.1) не является метафизической выдумкой. Ряды вида (6.1) возникают в самых естественных ситуациях. Например, по аналогии с $(1-x)^{-1} = 1 + x + x^2 + \dots$ хотелось бы иметь возможность так же вычислять обратные матрицы,

$$(I - A)^{-1} = I + A + A^2 + \dots,$$

но здесь, понятно, нужны условия, в которых это работает.

Другой источник — дифференциальное уравнение $\dot{X} = AX$. Естественный поиск решения в виде

$$X(t) = \sum_{k=0}^{\infty} C_k t^k$$

с неопределенными матричными коэффициентами C_k — после подстановки в $\dot{X} = AX$ дает

$$\sum_{k=1}^{\infty} k C_k t^{k-1} = \sum_{k=0}^{\infty} A C_k t^k.$$

Приравнивая коэффициенты при равных степенях t , получаем $k C_k = A C_{k-1}$, что в итоге (при условии $X(0) = I$) приводит к ряду

$$e^{At} = I + At + \frac{(At)^2}{2!} + \frac{(At)^3}{3!} + \dots,$$

который было бы естественно назвать матричной экспонентой, но пока не на что опереться — в смысле идеологической базы.

Ситуация отнюдь не безнадежна, как может показаться, а при отсутствии у A кратных собственных значений — даже тривиальна. Если $A = T^{-1}\Lambda T$, где матрица Λ диагональна, то частичные суммы

$$f_k(A) = \sum_{k=0}^q \gamma_k A^k$$

диагонализуются тем же преобразованием¹⁾ T ,

$$f_k(T^{-1}\Lambda T) = T^{-1}f_k(\Lambda)T.$$

Поэтому стоящие по краям T^{-1} и T оказываются как бы внешними наблюдателями того, что происходит на диагонали,

где элементами являются обычные суммы $\sum_{k=0}^q \gamma_k \lambda_j^k$, переходящие в обыкновенные числовые ряды

$$\sum_{k=0}^{\infty} \gamma_k \lambda_j^k \quad \text{при} \quad k \rightarrow \infty. \quad (6.2)$$

Для сходимости же (6.2), как известно (см. ниже), достаточно, чтобы модуль λ_j был строго меньше радиуса сходимости r , и необходимо $|\lambda_j| \leq r$. Таким образом, если спектральный радиус $\rho(A)$, т. е. максимум $|\lambda_j|$ по j , строго меньше r , то все числовые ряды (6.2) на диагонали сходятся.

В случае $I + A + A^2/2! + \dots$, например, диагональные элементы (6.2) имеют вид $1 + \lambda_j + \lambda_j^2/2! + \dots = e^{\lambda_j}$. В результате

$$A = T \begin{bmatrix} \lambda_1 & \dots & 0 \\ \vdots & \ddots & \vdots \\ 0 & \dots & \lambda_n \end{bmatrix} T^{-1} \Rightarrow e^A = T \begin{bmatrix} e^{\lambda_1} & \dots & 0 \\ \vdots & \ddots & \vdots \\ 0 & \dots & e^{\lambda_n} \end{bmatrix} T^{-1}.$$

Однако идиллическая картина такого типа разрушается при столкновении с недиагонализуемыми матрицами. Предельный переход здесь не работает. Приходится либо залезать в дебри жордановых канонических форм, либо изобретать простые обходные пути — см. раздел 6.5.

О сходимости числовых рядов. Напомним некоторые определения и результаты о числовых рядах (см. например [4, т. 1]).

¹⁾ В силу $(T^{-1}\Lambda T)^p = T^{-1}\Lambda T T^{-1}\Lambda T \dots T^{-1}\Lambda T = T^{-1}\Lambda^p T$.

Конечный или бесконечный предел α частичной суммы $\alpha_n = a_1 + \dots + a_n$ определяют как сумму ряда $\sum_{k=1}^{\infty} a_k$. Ряд, имеющий конечную (бесконечную) сумму, называют *сходящимся* (*расходящимся*).

- Если ряд $\sum_{k=1}^{\infty} a_k$ сходится, то $a_k \rightarrow 0$. Обратное неверно.
- Ряд называют *абсолютно сходящимся*, если сходится ряд из абсолютных величин, $|a_1| + |a_2| + \dots$. Любой абсолютно сходящийся ряд сходится.
- Если степенной ряд $\sum_{k=0}^{\infty} \gamma_k x^k$ сходится при каком-то x , $|x| = r$, то при условии $|x| < r$ он сходится абсолютно. Отсюда ясно, что область сходимости степенного ряда — всегда круг некоторого радиуса r (возможно бесконечного), который называют *радиусом сходимости*.

Ряды

$$e^x = 1 + \frac{x}{1!} + \frac{x^2}{2!} + \dots + \frac{x^n}{n!} + \dots, \quad \sin x = x - \frac{x^3}{3!} + \dots + (-1)^n \frac{x^{2n+1}}{(2n+1)!} + \dots$$

сходятся при любом x , — радиус сходимости бесконечен.

Ряд

$$\ln(1+x) = x - \frac{x^2}{2} + \frac{x^3}{3} - \frac{x^5}{5} + \dots,$$

сходится при $|x| < 1$, $r = 1$.

6.2. Нормы векторов и матриц

Разговор о нормах здесь более подробен, чем того требует проблематика сходимости матричных рядов. Но тема важна сама по себе. Нормы относятся к тем инструментам, которые работают каждый день, а не к тем, которые «висят на стене», а все ждут «когда выстрелит».

До сих пор рассматривался вариант $\|x\| = \sqrt{(x, x)}$, — но есть и другие способы. Их контроль осуществляет следующее определение.

Нормой вектора $x \in R^n$ называют положительное число $\|x\|$, удовлетворяющее следующим требованиям:

- (I) $\|x\| = 0 \Leftrightarrow x = 0$;
- (II) $\|x + y\| \leq \|x\| + \|y\|$ (*неравенство треугольника*);
- (III) $\|\alpha x\| = |\alpha| \cdot \|x\|$.

Множество точек, удовлетворяющих условию $\|x\| \leq r$ ($\|x\| = r$) называется *шаром* (сферой) радиуса r . В силу (iii) любой шар центрально симметричен.

Норма представляет собой обобщенное расстояние от точки x до начала координат. Расстояние между точками определяется как $\rho(x, y) = \|x - y\|$. Функцию $\rho(x, y)$ называют *метрикой*.

Помимо *евклидовой* $\|x\| = \sqrt{(x, x)}$ иногда используется норма

$$\|x\|_V = \sqrt{(x, Vx)},$$

где V — положительно определенная матрица.

Подставляя $x = Ay$ в (x, x) , получаем

$$(x, x) = (Ay, Ay) = (y, A^*Ay),$$

откуда видно, что (y, Vy) с матрицей $V = A^*A$ — есть то же самое (x, x) в новых координатах, переход к которым осуществляет матрица A .

Отсюда, кстати, вытекает, что A^*A — всегда положительно определена, разумеется для невырожденной матрицы A .

Широкое распространение имеют также

$$\|x\|_m = \max_i |x_i| \quad (\text{кубическая норма})$$

и

$$\|x\|_l = \sum_{i=1}^n |x_i| \quad (\text{октаэдрическая норма}).$$

В случае комплексного пространства под $|x_i|$ здесь понимается модуль комплексного числа x_i .

Норма однозначно определяется указанием единичного шара, в качестве которого может быть «назначено» любое выпуклое центрально-симметричное тело.

Разнообразие норм объясняется двумя причинами. Любопытством и свободой маневра, которую дает большой ассортимент вариантов. Основу свободного манипулирования нормами обеспечивает следующее утверждение.

6.2.1. В R^n все нормы эквивалентны (с точки зрения сходимости), т. е. для любых двух норм $\|\cdot\|_1, \|\cdot\|_2$ можно указать такие константы α и β , что

$$\|x\|_1 \leq \alpha \|x\|_2, \quad \|x\|_2 \leq \beta \|x\|_1$$

при любом $x \in R^n$.

◀ Функция $\varphi(x) = \|x\|_1$ на множестве $\|x\|_2 = 1$ достигает своего минимального ($\gamma > 0$) и максимального ($\delta > 0$) значения. Поэтому

$$\gamma \|x\|_2 \leq \|x\|_1 \leq \delta \|x\|_2$$

при любом $x \in R^n$, что завершает доказательство. ▶

Роль эквивалентности норм конечномерного пространства выявляется при изучении вопросов сходимости. В R^n всегда справедлива импликация

$$\|x_k - x_*\|_1 \rightarrow 0 \quad \Rightarrow \quad \|x_k - x_*\|_2 \rightarrow 0.$$

Это позволяет рассматривать предельные переходы в той норме, которая в наибольшей степени подходит данной конкретной задаче²⁾.

Простейший вопрос. Означает ли сходимость по норме ($\|x_k - \tilde{x}\| \rightarrow 0$) покоординатную сходимость? — Означает, и судить об этом легче всего в кубической норме, после чего общий положительный ответ вытекает из эквивалентности норм.

Нормой матрицы A называют положительное число $\|A\|$, удовлетворяющее тем же самым условиям (i)–(iii), если в них векторы заменить на матрицы.

Поскольку матрицы — это одновременно и линейные операторы в R^n , необходимо позаботиться об оценке образов Ax . Норма матрицы $\|A\|$ называется *согласованной* с нормой вектора $\|x\|$, если для любых A и x

$$\|Ax\| \leq \|A\| \cdot \|x\|$$

и для любых A и B выполняется неравенство $\|AB\| \leq \|A\| \cdot \|B\|$.

Чтобы неравенство $\|Ax\| \leq \|A\| \cdot \|x\|$ давало хорошую оценку, оно должно быть неулучшаемо. Это мотивирует выделение в самостоятельное понятие *подчиненной нормы* матрицы:

$$\|A\| = \max_{\|x\|=1} \|Ax\|.$$

²⁾ В бесконечномерных пространствах нормы не эквивалентны, и там необходимы меры предосторожности.

Матричные нормы играют важную роль как в самой линейной алгебре, так и за ее пределами. Для нелинейного оператора $f(x)$, например,

$$\begin{aligned}\|f(x+y) - f(x)\| &= \left\| \int_0^1 f'(x+\tau y)y \, d\tau \right\| \leq \\ &\leq \int_0^1 \|f'(x+\tau y)\| \cdot \|y\| \, d\tau = \|f'(x+\theta y)\| \cdot \|y\|\end{aligned}$$

при некотором $\theta \in (0, 1)$.

Поэтому, если $\|f'(z)\| \leq \nu < 1$, то f сжимает в метрике $\rho(x, y) = \|x - y\|$ со всеми вытекающими последствиями [4].

В случае евклидовой нормы $\|x\|_E = \sqrt{(x, x)}$ норма матрицы равна максимуму

$$\sqrt{(Ax, Ax)} = \sqrt{(x, A^*Ax)}$$

на единичной сфере. Это приводит к результату³⁾ $\|A\|_E = \sqrt{\rho(A^*A)}$.

Нормам $\|\cdot\|_m$ и $\|\cdot\|_l$ подчинены, соответственно⁴⁾,

$$\|A\|_m = \max_i \sum_{j=1}^n |a_{ij}| \quad (\text{строчная норма}),$$

и

$$\|A\|_l = \max_j \sum_{i=1}^n |a_{ij}| \quad (\text{столбцовая норма}).$$

Если $\|x\|$ — некоторая норма, а M — невырожденная матрица, то $\|x\|_M = \|M^{-1}x\|$ — также норма. Норма матрицы, подчиненная норме $\|x\|_M$, определяется как

$$\|A\|_M = \|M^{-1}AM\|.$$

В случае $M = \text{diag}\{\mu_1, \dots, \mu_n\}$, $\mu_i > 0$,

$$\|x\|_{mM} = \max_i \frac{1}{\mu_i} |x_i|, \quad \|A\|_{mM} = \max_i \frac{1}{\mu_i} \sum_{j=1}^n \mu_j |a_{ij}|,$$

$$\|x\|_{lM} = \sum_{i=1}^n \frac{1}{\mu_i} |x_i|, \quad \|A\|_{lM} = \max_j \mu_j \sum_{i=1}^n \frac{1}{\mu_i} |a_{ij}|.$$

³⁾ См. раздел 4.4.

⁴⁾ Модуль $|a_{ij}|$ обозначает модуль, вообще говоря, комплексного числа a_{ij} .

Двойственная норма. Если $\|A\| \leq \lambda$, — что можно сказать о норме транспонированной (сопряженной) матрицы? Существует ли норма $\|\cdot\|'$, в которой $\|A^*\|' \leq \lambda$? Универсальный положительный ответ дает понятие *двойственной нормы*, определяемой как

$$\|x\|^d = \max_{\|y\|=1} \{x, y\}.$$

◀ Импликация $\|A\| \leq \lambda \Rightarrow \|A^*\|^d \leq \lambda$ вытекает из

$$\|A^*x\|^d = \max_{\|y\|=1} \{A^*x, y\} = \max_{\|y\|=1} \{x, Ay\} \leq \max_{\|z\|=\lambda} \{x, z\} = \lambda \|x\|^d. \blacktriangleright$$

Двойственная норма к двойственной — опять исходная, $\|x\|^{dd} = \|x\|$. (?) Евклидова норма двойственна сама себе. Двойственная к октаэдрической норме — кубическая, и наоборот. (?) Отметим также

$$\|x\|^d = \frac{\max_{y \neq 1} \{x, y\}}{\|y\|} \Rightarrow (x, y) \leq \|x\|^d \cdot \|y\|.$$

О комплексификации пространства. Одно из неудобств при изучении линейной алгебры заключается в том, что объекты внимания все время «двоятся». С одной стороны, как уже отмечалось, переменные проще всего считать комплексными. С другой — хочется обойтись действительными. В частности, что надо иметь в виду при определении норм? Максимум в

$$\|A\| = \max_{\|x\|=1} \|Ax\| \quad \text{либо} \quad \|x\|^d = \max_{\|y\|=1} \{x, y\}$$

надо считать по комплексным переменным или действительным? Оказывается, все равно. Для понимания этого обстоятельства полезна следующая точка зрения на комплексные пространства.

Комплексификацией любого нормированного пространства E считается комплексное пространство $\hat{E}$, элементами которого являются формальные пары $z = x + iy$ ($x, y \in E$) с определением линейных операций на базе сложения и умножения комплексных чисел,

$$(\alpha + i\beta)(x + iy) = (\alpha x - \beta y) + i(\beta x + \alpha y),$$

и продолжением линейных операторов на $\hat{E}$ по правилу $A(x + iy) = Ax + iAy$.

Норма в $\hat{E}$ определяется формулой

$$\|x + iy\| = \max_{\varphi} \|x \cos \varphi + y \sin \varphi\|,$$

откуда ясно, что в $\|A\| = \max_{\|x\|=1} \|Ax\|$, например, нет разницы, какие переменные имеются в виду.

6.3. Спектральный радиус

Напомним, что спектральным радиусом $\rho(A)$ называется максимальный модуль собственного числа матрицы A . Иными словами, $\rho(A)$ — это радиус минимального круга, содержащего весь спектр $\sigma(A)$.

Из $\lambda x = Ax$ следует

$$|\lambda| \cdot \|x\| = \|Ax\| \leq \|A\| \cdot \|x\| \Rightarrow |\lambda| \leq \|A\|,$$

т. е. любая норма матрицы $\geq$ любого собственного значения. Отсюда

$$\rho(A) \leq \|A\|.$$

В разделе 5.5 уже отмечалось, что для сходимости A^m к нулевой матрице при $m \rightarrow \infty$ необходимо и достаточно, чтобы $\rho(A) < 1$.

Для диагоналируемой матрицы это очевидно. В общем случае результат вытекает из всегда возможного представления $A = \Lambda_A + N_A$, где Λ_A — диагональная матрица, а N_A — нильпотентная, причем Λ_A и N_A коммутируют. Поскольку $N_A^k = 0$, начиная с некоторого $k = k_0$, то поведение A^m при больших m определяется поведением Λ_A^m .

6.3.1. Лемма. Если $\rho(A) < 1$, то существует норма $\|\cdot\|_*$, в которой $\|A\|_* < 1$.

◀ В качестве $\|\cdot\|_*$ можно взять

$$\|x\|_* = \max_{k=0,1,\dots} \frac{\|A^k x\|}{\lambda^k}, \quad (6.3)$$

где $\rho(A) < \lambda < 1$, а $\|\cdot\|$ — произвольная норма.

Максимум существует, поскольку $\rho(A/\lambda) < 1 \Rightarrow (A/\lambda)^k \rightarrow 0$ при $k \rightarrow \infty$. Аксиомы нормы для (6.3) проверяются без труда. Окончательный вывод следует из неравенства

$$\|Ax\|_* = \max_{k=0,1,\dots} \frac{\|A^k Ax\|}{\lambda^k} = \lambda \max_{k=1,\dots} \frac{\|A^k x\|}{\lambda^k} \leq \lambda \|x\|_*. \quad \blacktriangleright$$

Лемма 6.3.1 — важный и полезный результат, облегчающий многие рассуждения. Разумеется, всегда существует норма $\|\cdot\|_*$, в которой при любом $\epsilon > 0$

$$\boxed{\|A\|_* < \rho(A) + \epsilon,} \quad (6.4)$$

что сразу вытекает из леммы 6.3.1, поскольку оператор

$$(\rho(A) + \epsilon)^{-1} A,$$

независимо от A , имеет спектральный радиус строго меньший 1.

6.3.2. Какова бы ни была норма,

$$\rho(A) = \lim_{k \rightarrow \infty} \sqrt[k]{\|A^k\|}.$$

◀ С одной стороны, $\rho(A)^k = \rho(A^k) \leq \|A^k\|$ влечет за собой

$$\rho(A) \leq \sqrt[k]{\|A^k\|}.$$

С другой — спектральный радиус матрицы $\tilde{A} = [\rho(A) + \varepsilon]^{-1} A$ при $\varepsilon > 0$ строго меньше 1. Поэтому $\|\tilde{A}^k\| \rightarrow 0$ при $k \rightarrow \infty$, что означает

$$\|A^k\| \leq [\rho(A) + \varepsilon]^k,$$

при достаточно больших k . Таким образом, $\sqrt[k]{\|A^k\|} - \varepsilon \leq \rho(A) \leq \sqrt[k]{\|A^k\|}$. ▶

Наконец, $\rho(A) = \rho(A^*)$.

6.4. Сходимость итераций

Для вычислений часто используется итерационная процедура

$$x_{k+1} = Ax_k + b. \quad (6.5)$$

Процесс, очевидно, разворачивается следующим образом,

$$\begin{aligned} x_1 &= Ax_0 + b, \\ x_2 &= A^2x_0 + Ab + b, \\ &\dots \\ x_k &= A^kx_0 + A^{k-1}b + \dots + Ab + b, \\ &\dots \end{aligned}$$

При условии $\|A\| < 1$ слагаемое A^kx_0 стремится к нулю, поскольку $\|A^k\| \leq \|A\|^k \rightarrow 0$ при $k \rightarrow \infty$. Поэтому динамика процесса (6.5) определяется сходимостью ряда $\sum_{k=0}^{\infty} A^k$, который сходится в силу

$$\|I + A + A^2 + \dots + A^k + \dots\| \leq 1 + \|A\| + \|A\|^2 + \dots = \frac{1}{1 - \|A\|},$$

что влечет за собой сходимость x_k к решению уравнения

$$x = Ax + b, \quad x = (I - A)^{-1}b,$$

соответственно, $I + A + A^2 + \dots + A^k + \dots \rightarrow (I - A)^{-1}$.

6.5. Функции как ряды

Сходимость матричного ряда (6.1) обычно определяют через сходимость частичных сумм

$$f_q(A) = \sum_{k=0}^q \gamma_k A^k \quad (6.6)$$

к $f(A)$, что оформляется стандартным образом. Ряд (6.1) сходится к $f(A)$, если по любому $\varepsilon > 0$ можно указать такое N , что $\|f(A) - f_q(A)\| < \varepsilon$ как только $q > N$.

Чаше всего это «стреляет в молоко», поскольку $f(A)$ обычно определяется как сумма ряда, — и выходит, что сходимость выводится сама из себя. Разорвать порочный круг позволяет понятие *фундаментальной последовательности (последовательности Коши)*. В данном случае эффективный аналог звучит так.

Матричная последовательность $f_k(A)$ фундаментальна, если по любому $\varepsilon > 0$ можно указать такое N , что

$$\|f_p(A) - f_q(A)\| < \varepsilon,$$

как только $p, q > N$.

В конечномерном пространстве фундаментальные последовательности сходятся, а сходящиеся — фундаментальны (см. [4, т. 1]). Что касается (6.1), то здесь добавляется специфика степенного ряда.

◀ Пусть $\sum_{k=0}^{\infty} \gamma_k \lambda^k$ сходится, что влечет за собой $\gamma_k \lambda^k \rightarrow 0$, а значит, и ограниченность: $|\gamma_k \lambda^k| < M$. Пусть $\|A\| < |\lambda|$. Тогда

$$\sum_{k=0}^{\infty} \|\gamma_k A^k\| = \sum_{k=0}^{\infty} |\gamma_k \lambda^k| \left| \frac{\|A\|}{\lambda} \right|^k < M \sum_{k=0}^{\infty} \left| \frac{\|A\|}{\lambda} \right|^k < \infty.$$

Следовательно, сходится ряд $\sum_{k=0}^{\infty} \gamma_k \|A\|^k$. Это означает, что частичные суммы $\sum_{k=0}^q \gamma_k A^k$ представляют собой фундаментальную последовательность, поскольку

$$\left\| \sum_{k=p}^q \gamma_k A^k \right\| \leq \sum_{k=p}^q \gamma_k \|A\|^k.$$

В итоге гарантирована сходимость ряда $\sum_{k=0}^{\infty} \gamma_k A^k$. ▶

Подключение к проведенному рассуждению леммы 6.3.1 вместе с неравенством (6.4) приводит к следующему результату.

6.5.1. Если спектральный радиус $\rho(A) < r$, где r радиус сходимости ряда $\sum_{k=0}^{\infty} \gamma_k \lambda^k$, то матричный ряд $\sum_{k=0}^{\infty} \gamma_k A^k$ сходится, определяя некоторую функцию $f(A)$.

Никаких дополнительных требований на матрицу A здесь не накладывается.

6.6. Матричная экспонента

Главным мотивом изучения функций от матриц, конечно, является экспонента e^A , которая теперь строго определена:

$$e^A = I + A + \frac{A^2}{2!} + \frac{A^3}{3!} + \dots, \quad (6.7)$$

но, как всегда, по достижении цели становится ясно, что хотелось-то не совсем того, что получилось.

Формула (6.7) позволяет просто и с большой точностью вычислять e^A . Но что можно сказать, например, о спектре e^A ? Конечно, если матрица A диагонализуема, то (см. раздел 6.1)

$$e^A = T^{-1} \begin{bmatrix} e^{\lambda_1} & \dots & 0 \\ \vdots & \ddots & \vdots \\ 0 & \dots & e^{\lambda_n} \end{bmatrix} T, \quad (6.8)$$

где λ_j собственные значения A , а T диаболизирующее преобразование. Спектр e^A , получается, состоит из точек e^{λ_j} , да еще и собственные векторы — те же самые⁵⁾.

А если — не диаболизируема? Готового ответа нет. Определение (6.7), правда, пока выручает. Можно использовать стандартную идею предельного перехода $A(\epsilon) \rightarrow A$ с диаболизируемыми при $\epsilon \neq 0$ матрицами $A(\epsilon)$. Коэффициенты характеристического полинома непрерывно зависят от ϵ , и в спектре оказываются те же

⁵⁾ Вектор-столбцы матрицы T являются собственными векторами как самой матрицы A , так и экспоненты e^A .

точки⁶⁾ e^{λ_j} . Но конечной записи типа (6.8) уже не получается, и повисают вопросы, связанные с собственными векторами и каноническим представлением.

На базе определения функций с помощью матричных рядов, конечно, можно продвинуться достаточно далеко. Например, если T приводит A к блочно-диагональному виду

$$A = T^{-1} \begin{bmatrix} A_1 & & 0 \\ & \ddots & \\ 0 & & A_p \end{bmatrix} T,$$

то, очевидно,

$$f(A) = T^{-1} \begin{bmatrix} f(A_1) & & 0 \\ & \ddots & \\ 0 & & f(A_p) \end{bmatrix} T,$$

где $f(A_j)$ определяются рядом с теми же самыми коэффициентами, что и $\sum_{k=0}^{\infty} \gamma_k A^k$.

Если T приводит A к треугольному виду $A = T^{-1}ST$, то и $f(A) = T^{-1}f(S)T$, где $f(S)$ имеет тот же самый вид, что и S . Это происходит потому, что полином от треугольной матрицы — тоже треугольная матрица, и предельный переход не в состоянии разрушить эту импликацию.

Получается как бы, что «хороший» базис для A хорош и для $f(A)$, независимо от f . То есть при подборе базиса можно ориентироваться не на функцию, а на аргумент. Матрицы $f(A)$ и A выглядят родственной парой. Они всегда коммутируют, у них общие инвариантные подпространства, — и все это из-за того, что $f(A)$ является пределом полиномов от A .

Отметим популярную формулу

$$\boxed{\det e^A = e^{\text{tr} A}}, \quad (6.9)$$

которая в случае диагоналируемой матрицы A очевидна, поскольку при линейном преобразовании $T^{-1}AT$ ни детерминант, ни след не меняются⁷⁾, а в диагональном случае (6.9) есть простое

$$e^{\lambda_1} \dots e^{\lambda_n} = e^{\lambda_1 + \dots + \lambda_n}.$$

Окончательный вывод о справедливости (6.9), без предположений о свойствах A , получается стандартным предельным переходом $A(\varepsilon) \rightarrow A$.

⁶⁾ И аналогичный вывод получается при рассмотрении других функций типа $\sqrt{A}$ или $\ln A$.

⁷⁾ Поскольку детерминант равен произведению, а след — сумме собственных значений матрицы. Последние же инвариантны при переходе к другому базису.

Обратим, наконец, внимание на импликацию

$$AB = BA \Rightarrow \boxed{e^{A+B} = e^A e^B}.$$

Но $\boxed{e^{A+B} \neq e^A e^B}$, если $AB \neq BA$. Это становится очевидным при перемножении рядов e^A и e^B .

Упражнения

- Если матрица A гурвицева (все $\operatorname{Re} \lambda_j < 0$), то спектральный радиус $\rho(e^A) < 1$.
- $e^{A(t+s)} = e^{At} e^{As}$.
- Независимо от A экспонента всегда невырождена, $\det e^A \neq 0$.

6.7. Конечные алгоритмы

За ширмой бесконечных рядов прячутся все-таки какие-то другие механизмы вычисления функций. Точные результаты типа

$$A = \begin{bmatrix} 0 & 1 \\ -1 & 0 \end{bmatrix} \Rightarrow e^{At} = \begin{bmatrix} \cos t & \sin t \\ -\sin t & \cos t \end{bmatrix},$$

$$A = \begin{bmatrix} 0 & 1 & 0 \\ 0 & 0 & 1 \\ 0 & 0 & 0 \end{bmatrix} \Rightarrow e^{At} = \begin{bmatrix} 1 & t & \frac{t^2}{2} \\ 0 & 1 & t \\ 0 & 0 & 1 \end{bmatrix}$$

получаются, конечно, с помощью (6.7), но возникает впечатление, что за кадром остаются нераскрытые возможности.

Рассмотрим сначала диагоналируемую матрицу A . В этом случае

$$f(A) = T^{-1} \operatorname{diag} \{f(\lambda_1), \dots, f(\lambda_n)\} T, \quad (6.10)$$

что уже позволяет обойтись без (6.7). Но здесь интересно продвинуться чуть дальше. Из (6.10) видно, что $f(A)$ определяется значениями $f(\lambda)$ всего в n точках $\{\lambda_1, \dots, \lambda_n\}$. Это позволяет заменить f любой другой функцией, которая в этих точках принимает те же значения, — и результат не изменится (если, конечно, речь идет о вычислении $f(A)$ для данной конкретной матрицы A).

В частности, $f(\lambda)$ можно положить равной интерполяционному полиному Лагранжа ⁸⁾

$$f(\lambda) = P(\lambda) = a_{n-1}\lambda^{(n-1)} + \dots + a_1\lambda + a_0.$$

Тогда

$$\begin{aligned} f(A) &= T^{-1} \sum_{k=0}^{n-1} a_k \operatorname{diag}\{\lambda_1^k, \dots, \lambda_n^k\} T = \sum_{k=0}^{n-1} a_k T^{-1} \operatorname{diag}\{\lambda_1, \dots, \lambda_n\}^k T = \\ &= \sum_{k=0}^{n-1} a_k (T^{-1} \operatorname{diag}\{\lambda_1, \dots, \lambda_n\} T)^k = \sum_{k=0}^{n-1} a_k A^k. \end{aligned}$$

Выходит, что $f(A)$, какова бы ни была функция f , всегда можно вычислить как полином $(n-1)$ -й степени от A . Отличие от (6.1) состоит в том, что коэффициенты этого полинома зависят от A . Но это никак не мешает, если речь идет об одной конкретной матрице A . Преимущество более громоздкого рецепта (6.1) заключается в универсальности — алгоритм вычисления одинаков для любой матрицы. Тем не менее представление отдельно взятой матрицы $f(A)$ в виде конечного полинома от A нередко играет принципиальную роль.

Например, в динамических задачах об управляемости возникают системы уравнений типа $f(A)Bx = b$, для однозначной разрешимости которых требуется $\operatorname{rank} f(A)B = n$. При этом один лишь факт представимости $f(A)$ в виде $a_0I + \dots + a_k A^k$ позволяет заменить $\operatorname{rank} f(A)B = n$ конечно проверяемым условием $\operatorname{rank} U = n$ для матрицы

$$U = [B, AB, A^2B, \dots, A^{n-1}B].$$

Все это справедливо и в общем (недиагонализируемом) случае, но для обоснования приходится углубляться в дебри корневых подпространств.

Из жорданова представления $A = S^{-1}JS$,

$$J = \operatorname{diag}\{J_{k_1}(\lambda_1), \dots, J_{k_p}(\lambda_p)\},$$

⁸⁾ Где коэффициенты a_j определяются после приведения подобных в

$$P(\lambda) = \sum_k f(\lambda_k) \prod_{j \neq k} \frac{\lambda - \lambda_j}{\lambda_k - \lambda_j}.$$

следует

$$f(A) = S^{-1} \operatorname{diag} \{f(J_{k_1}(\lambda_1)), \dots, f(J_{k_p}(\lambda_p))\} S,$$

откуда становится ясно, что ключом к $f(A)$ является умение вычислять $f(J_k(\lambda))$.

Поскольку $J_k(\lambda) = \lambda I + D$, где наддиагональная матрица D нильпотентна, $D^j = 0$ при $j \geq k$, — то в разложении⁹⁾

$$J_k^q(\lambda) = (\lambda I + D)^q = \sum_j^q C_q^j D^j \lambda^{q-j}$$

отличаются от нуля только члены, содержащие D в степени $j < k$.

Приведение подобных в $f(J_k) = \sum_{q=0}^{\infty} \gamma_q J_k^q$ показывает, что частичные суммы этого ряда имеют вид

$$\sum_{q=0}^N \gamma_q J_k^q = \sum_{j=0}^{k-1} F_{Nj} D^j, \quad \text{где} \quad F_{Nj} = \sum_{s=0}^N \gamma_s C_s^j \lambda^{s-j}.$$

С другой стороны, если λ лежит в круге сходимости ряда $f(\lambda) = \sum_{q=0}^{\infty} \gamma_q \lambda^q$, то этот ряд можно почленно дифференцировать, не нарушая сходимости и получая

$$f^{(j)}(\lambda) = \sum_{q=j}^{\infty} \gamma_q \frac{C_q^j}{j!} \lambda^{q-j}.$$

Поэтому

$$F_{Nj} \rightarrow \frac{f^{(j)}(\lambda)}{j!} \quad \text{при} \quad N \rightarrow \infty.$$

В итоге (с учетом $D = J_k - \lambda I$) получается

$$f(J_k) = \sum_{j=0}^{k-1} \frac{f^{(j)}(\lambda)}{j!} (J_k - \lambda I)^j.$$

Например,

$$J_4 = \begin{bmatrix} \lambda & 1 & 0 & 0 \\ 0 & \lambda & 1 & 0 \\ 0 & 0 & \lambda & 1 \\ 0 & 0 & 0 & \lambda \end{bmatrix} \Rightarrow f(J_4) = \begin{bmatrix} f(\lambda) & f'(\lambda) & \frac{1}{2}f^{(2)}(\lambda) & \frac{1}{6}f^{(3)}(\lambda) \\ 0 & f(\lambda) & f'(\lambda) & \frac{1}{2}f^{(2)}(\lambda) \\ 0 & 0 & f(\lambda) & f'(\lambda) \\ 0 & 0 & 0 & f(\lambda) \end{bmatrix},$$

какова бы ни была функция f .

⁹⁾ Здесь C_q^j биномиальные коэффициенты.

Опять значение $f(A)$ для данной конкретной матрицы A оказывается зависящим только от *конечного числа* значений функции f и ее производных на точках спектра $\sigma(A)$. Поэтому $f(A)$ можно заменить полиномом $P(A)$. По сравнению с диагонализируемым вариантом добавились значения производных, но убавилось число точек спектра. Легко сообразить (разобравшись предварительно в корневых подпространствах), что степень $P(A)$ на единицу меньше степени минимального полинома A .

6.8. Задачи и дополнения

- Скалярные тождества с участием одной переменной¹⁰⁾ переходят — в матричные,

$$\sin^2 \varphi + \cos^2 \varphi = 1 \quad \Rightarrow \quad \sin^2 A + \cos^2 A = I.$$

- $e^{iA} = \cos A + i \sin A$.
- Унитарная матрица U всегда имеет представление $U = e^{iH}$, где H — самосопряженная (эрмитова) матрица. Если U вещественна и ортогональна, причем $\det U = 1$, то $U = e^S$, где S — кососимметричная матрица.
- Преобразование Лапласа¹¹⁾ дифференциального уравнения $\dot{x} = Ax$ приводит к

$$p\hat{x}(p) - x(0) = A\hat{x}(p),$$

откуда $\hat{x}(p) = (Ip - A)^{-1}x(0)$, т. е. $x(t) \doteq (Ip - A)^{-1}x(0)$. Теперь сравнение с известным фактом $x(t) = e^{At}x(0)$ дает

$$e^{At} \doteq (Ip - A)^{-1}.$$

Используя формулу обращения

$$f(t) = \frac{1}{2\pi i} \int_{\sigma - i\infty}^{\sigma + i\infty} e^{pt} \hat{f}(p) dp,$$

это можно записать как

$$e^{At} = \frac{1}{2\pi i} \int_{\sigma - i\infty}^{\sigma + i\infty} e^{pt} (Ip - A)^{-1} dp. \quad (6.11)$$

¹⁰⁾ При двух переменных мешает некоммутативность матриц.

¹¹⁾ Преобразованием Лапласа называют интегральное преобразование

$$\hat{f}(p) = \int_0^{\infty} e^{-pt} f(t) dt,$$

где $p = \sigma + i\omega$. Функцию $f(t)$ при этом называют *оригиналом*, а $\hat{f}(p)$ — *изображением*, что записывают обычно как $f(t) \doteq \hat{f}(p)$.

Если матрица A устойчива (гурвицева), то в (6.11) можно положить $\sigma = 0$, поскольку единственное назначение σ в преобразовании Лапласа — это подавление быстро растущих функций. Если матрица гурвицева, все $\|x(t)\| \rightarrow 0$, подавлять нечего. Задача попадает в рамки преобразования Фурье.

- Матричную функцию $(Ip - A)^{-1}$ комплексной переменной p называют *резольвентой*. В предыдущем пункте резольвента сыграла принципиальную роль в интегральном представлении экспоненты. Это обстоятельство не случайно. Имеет место общая формула

$$f(A) = \frac{1}{2\pi i} \oint_C f(p)(Ip - A)^{-1} dp, \quad (6.12)$$

обобщающая обычную интегральную формулу Коши из ТФКП (см. [4, т. 1]). Для справедливости (6.12) предполагается, что замкнутый контур C содержит внутри себя спектр A , а функция $f(p)$ аналитична внутри C .

Различные способы описания матричных функций — это различные инструменты, позволяющие решать разнородные задачи. Конечно, можно ограничиться чем-то одним, но это ограничивает возможности. Так же как приготовление пищи с помощью только сковороды.

Что касается представления (6.12), то его более естественно изучать вместе с аналитическими функциями, где базовая идеология способствует проникновению в суть дела.

Глава 7

Матричные уравнения

7.1. Типичные задачи

Матричные уравнения естественно возникают в задачах, которые изначально выглядят, как «векторные».

Скажем, поиск собственных векторов $\mathbf{x}_j$ ($A\mathbf{x}_j = \lambda_j\mathbf{x}_j$) объединяется одним уравнением

$$\boxed{AX = X\Lambda}, \quad (7.1)$$

где $\Lambda = \text{diag} \{\lambda_1, \dots, \lambda_n\}$, а X — искомая матрица со столбцами $\mathbf{x}_j$ в качестве собственных векторов.

Другой характерный пример — линейное дифференциальное уравнение $\dot{\mathbf{x}} = A\mathbf{x}$. Как известно, общее решение имеет вид

$$\mathbf{x}(t) = c_1\mathbf{x}_1(t) + \dots + c_n\mathbf{x}_n(t),$$

где линейно независимые решения $\mathbf{x}_j(t)$, как вектор-столбцы, составляют матрицу фундаментальных решений

$$X(t) = \{\mathbf{x}_1(t), \dots, \mathbf{x}_n(t)\},$$

удовлетворяющую матричному дифференциальному уравнению

$$\boxed{\dot{X} = AX},$$

которое выгоднее рассматривать с самого начала вместо $\dot{\mathbf{x}} = A\mathbf{x}$.

Поиск преобразования X , обеспечивающего подобие матриц A и B , порождает уравнение $X^{-1}AX = B$. После умножения слева

на X оно переходит в эквивалентное

$$AX = XB, \quad (7.2)$$

в предположении невырожденности X .

Очевидно, (7.1) при заданной матрице Λ представляет собой частный случай (7.2).

Наконец, поиск функции Ляпунова¹⁾ для линейной системы $\dot{x} = Ax$ приводит к уравнению

$$A^*V + VA = -I \quad (7.3)$$

относительно матрицы V .

Обозначая неизвестную матрицу V через X и обобщая (7.3), приходим к уравнению

$$AX + XB = C, \quad (7.4)$$

которое охватывает в качестве частных случаев все рассмотренные выше случаи, — разумеется, кроме дифференциального.

Уравнение (7.3) возникает следующим образом. Ищется положительно определенная квадратичная форма $\varphi(x) = (x, Vx)$, убывающая на траекториях $\dot{x} = Ax$. Очевидно,

$$\dot{\varphi} = (Ax, Vx) + (x, VAx) = (x, [A^*V + VA]x).$$

С точки зрения убывания $\varphi(x)$ устроило бы любое решение $A^*V + VA = W$ с отрицательно определенной матрицей W . Оказывается, однако, что выбор W не играет роли (что само по себе представляет интересный факт), — и потому берут максимально простое $W = -I$.

7.2. Кронекерово произведение

Уравнение (7.4) линейно относительно элементов x_{ij} неизвестной матрицы X , и этим замечанием, казалось бы, можно закончить исследование, сославшись на предыдущее изучение линейных уравнений. Но проблема заключается в том, что уравнение (7.4), как линейное, имеет нестандартную форму, опираясь на двухиндексное описание переменных. В принципе, нет никакой трудности

¹⁾ См. [4, т. 2].

в том, чтобы перенумеровать переменные x_{ij} , вытянув их в строчку. Но при бесхитростной перенумерации содержательная информация о матрицах A , B , C может разрушиться, что будет означать отсутствие смысловой связи между получаемыми линейными системами в R^{n^2} и исходными операторами A , B , C , действующими в R^n .

Для решения таких задач имеется специальный инструмент — *кронекерово произведение*²⁾ матриц, $A \otimes B$. Если A и B — прямоугольные матрицы размера, соответственно, $m \times n$ и $p \times q$, то

$$A \otimes B = \begin{bmatrix} a_{11}B & a_{12}B & \dots & a_{1n}B \\ a_{21}B & a_{22}B & \dots & a_{2n}B \\ \dots & \dots & \dots & \dots \\ a_{m1}B & a_{m2}B & \dots & a_{mn}B \end{bmatrix}.$$

Размер $A \otimes B$, очевидно, $mp \times nq$.

Свойства

$$(A + B) \otimes C = A \otimes C + B \otimes C,$$

$$A \otimes (B + C) = A \otimes B + A \otimes C,$$

$$(A \otimes B) \otimes C = A \otimes (B \otimes C) = A \otimes B \otimes C$$

легко проверяются.

Важную роль играет формула

$$(A \otimes B)(C \otimes D) = AC \otimes BD. \quad (7.5)$$

Разумеется, все операции в (7.5) предполагаются выполнимыми, т.е. размерности матриц должны быть согласованы. Часто рассматриваются только квадратные матрицы одного размера, тогда о проблеме согласованности можно не думать. В частности, для множества квадратных матриц A_k , B_k справедливо

$$(A_1 \otimes B_1) \cdots (A_p \otimes B_p) = (A_1 \cdots A_p) \otimes (B_1 \cdots B_p), \quad (7.6)$$

что легко выводится из (7.5) по индукции.

²⁾ Кронекерово произведение называют также *прямым* или *тензорным* произведением.

◀ Равенство (7.5) вытекает из

$$\sum_k (a_{ik}B)(c_{kj}D) = \left(\sum_k a_{ik}c_{kj} \right) BD,$$

поскольку $\sum_k a_{ik}c_{kj}$ есть (ij) -й элемент матрицы AC . ▶

Легко проверяется

$$(A \otimes B)^* = A^* \otimes B^*.$$

Менее очевидно, что в случае невырожденности квадратных матриц A и B произведение $A \otimes B$ тоже невырожденно, и ³⁾

$$(A \otimes B)^{-1} = A^{-1} \otimes B^{-1}.$$

Рассмотрим две квадратные матрицы A и B размера $n \times n$, каждая из которых имеет n попарно различных собственных значений, которые обозначим через λ_j , μ_j , и пусть e_j , f_j — соответствующие собственные векторы. В силу (7.5),

$$(A \otimes B)(e_i \otimes f_j) = Ae_i \otimes Bf_j = \lambda_i e_i \otimes \mu_j f_j = \lambda_i \mu_j e_i \otimes f_j.$$

Это влечет за собой справедливость следующего факта.

7.2.1. *Спектр матрицы $A \otimes B$, имеющей размер $n^2 \times n^2$, состоит из n^2 попарных произведений $\lambda_i \mu_j$, которым отвечают собственные векторы $e_i \otimes f_j$.*

Если теперь допустить кратные собственные значения у A и B , то идея предельного перехода $A(\varepsilon) \rightarrow A$, $B(\varepsilon) \rightarrow B$ (см. раздел 3.4) здесь работает без проблем. Собственные векторы в пределе могут становиться линейно зависимыми, но это в данном случае ничему не мешает. Поэтому утверждение 7.2.1 справедливо без каких бы то ни было предположений о матрицах A и B .

Сразу становится ясной отмечавшаяся выше невырожденность $A \otimes B$ в случае невырожденности A и B .

7.3. Уравнения

Обычное соотношение $AX = C$ может быть записано в виде

$$(I \otimes A)x = c,$$

³⁾ Из (7.5) следует $(A \otimes B)(A^{-1} \otimes B^{-1}) = AA^{-1} \otimes BB^{-1}$.

где вектор x это вытянутая в столбик матрица X ,

$$x = [x_{11}, \dots, x_{n1}; \dots, x_{1n}, \dots, x_{nn}]^*,$$

вектор c из C получен аналогично.

Соотношение $XB = C$ записывается иначе,

$$(B^* \otimes I)x = c.$$

Поэтому уравнение (7.4), т. е.

$$AX + XB = C,$$

с помощью кронекерова произведения можно переписать так

$$(I \otimes A + B^* \otimes I)x = c. \quad (7.7)$$

В этой перезаписи уравнения не было бы большого смысла, если бы она не позволяла делать выводы в терминах исходных матриц A и B . Но специфика кронекерова произведения как раз такова, что она дает возможность судить о спектральных свойствах « $\otimes$ -матриц» во многих практических ситуациях. Причиной является следующий факт.

7.3.1. Лемма. Если

$$A' = T^{-1}AT, \quad B' = S^{-1}BS,$$

то

$$A' \otimes B' = (T \otimes S)^{-1}(A \otimes B)(T \otimes S).$$

◀ Результат сразу вытекает из (7.6),

$$(T^{-1}AT) \otimes (S^{-1}BS) = (T^{-1} \otimes S^{-1})(A \otimes B)(T \otimes S). \quad \blacktriangleright$$

Лемма 7.3.1 означает, что матрицы A и B в $A \otimes B$ могут быть приведены к желаемому виду (диагональному, треугольному, жордановому) независимо друг от друга. Пусть, например,

$$A = T^{-1}J_AT, \quad B = S^{-1}J_BS,$$

где J_A, J_B — соответствующие жордановы формы. Тогда ясно, что спектры $A \otimes B$ и $J_A \otimes J_B$ совпадают⁴⁾, что еще раз доказывает утверждение 7.2.1.

⁴⁾ Матрица $J_A \otimes J_B$ верхняя треугольная с собственными значениями на диагонали.

Точно так же A и B могут быть приведены к своим жордановым формам⁵⁾ в

$$I \otimes A + B^* \otimes I, \quad (7.8)$$

после чего становится очевидным, что матрица (7.8) имеет собственные значения $\lambda_i + \mu_j$ при всевозможных комбинациях i и j . Под λ_i , μ_j подразумеваются собственные значения A и B .

Отсюда ясно, что для невырожденности (7.8) необходимо и достаточно, чтобы не нашлось противоположных собственных значений, $\lambda_i + \mu_j \neq 0$. Это и является условием однозначной разрешимости уравнения (7.7), т. е. (7.4).

Как теоретический инструмент иногда полезна формула

$$X = - \int_0^{\infty} e^{At} C e^{Bt} dt, \quad (7.9)$$

дающая решение уравнения $AX + XB = C$ в случае, когда матрицы A и B гурвицевы, т. е. действительные части их собственных значений строго отрицательны. Устанавливается это совсем легко. Решением задачи Коши

$$\dot{Y} = AY + YB, \quad Y(0) = C, \quad (7.10)$$

является $Y(t) = e^{At} C e^{Bt}$, что проверяется подстановкой. Из гурвицевости A и B следует⁶⁾ экспоненциально быстрое убывание до нуля e^{At} и e^{Bt} при $t \rightarrow \infty$. Это позволяет проинтегрировать (7.10) от 0 до ∞ , что сразу дает (7.9).

В частности, решение уравнения $A^*V + VA = -I$ в случае гурвицевой матрицы A приводит к положительно определенной функции Ляпунова

$$(x, Vx) = \left(x, \int_0^{\infty} e^{A^*t} I e^{At} dt x \right) = \int_0^{\infty} (e^{A^*t} x, e^{At} x) dt > 0.$$

⁵⁾ Достаточно говорить о диагонализации, но тогда надо иметь в виду предельный переход $A(\varepsilon) \rightarrow A$ (см. выше).

⁶⁾ См. [4, т. 2].

Стоит обратить внимание на то обстоятельство, что кронекерово произведение выручает в вопросах разрешимости матричных уравнений, но пасует при изучении свойств самого решения. Что будет, если «длинный» вектор x сложить в матрицу X ? Будет ли матрица X вырождена, положительно определена, диагонализируема? Здесь уже лучше анализировать само исходное уравнение.

Скажем, из проведенного выше анализа ясно, что однородное уравнение $AX - XB = 0$ имеет ненулевое решение, если найдутся $\lambda_i + \mu_j \neq 0$, дающие в сумме нуль. Но каковы условия невырожденности X ? Это уже довольно тонкий и громоздкий вопрос. Ответ положителен, когда жордановы формы A и B одинаковы.

Упражнения

- Уравнение $AXB = C$ эквивалентно

$$(B^* \otimes A)x = c,$$

где x и c — вытянутые в столбик матрицы X и C .

- Для однозначной разрешимости $AXB - X = C$ необходимо и достаточно отсутствие у A и B таких собственных значений, что $\lambda_i \mu_j = 1$.
- $\operatorname{tr} A \otimes B = \operatorname{tr} A \cdot \operatorname{tr} B$.
- Если матрицы A и B обе унитарны, симметричны и положительно определены, нормальны, — то матрица $A \otimes B$ такая же.

Глава 8

Неравенства

8.1. Теоремы об альтернативах

Матрицы в этой главе *не обязательно квадратные*, что естественно при рассмотрении систем линейных неравенств

$$Ax = \begin{bmatrix} a_{11} & a_{12} & \dots & a_{1n} \\ a_{21} & a_{22} & \dots & a_{2n} \\ \dots & \dots & \dots & \dots \\ a_{m1} & a_{m2} & \dots & a_{mn} \end{bmatrix} x > 0,$$

где « $>$ » означает совокупность покомпонентных неравенств,

$$\{c_1, \dots, c_n\} > 0 \Rightarrow \text{все } c_i > 0.$$

Аналогично определяется « $\geq$ » между векторами.

Разрешимость неравенства $Ax > 0$ означает существование такого вектора x , скалярное произведение которого на любую вектор-строку a_i положительно,

$$(a_i, x) = \sum_j a_{ij} x_j > 0.$$

Геометрически ясно, что подходящий элемент x существует, если все векторы a_i расположены строго по одну сторону некоторой $(n-1)$ -мерной плоскости.

Для этого, в свою очередь, требуется, чтобы множество $K(A)$ векторов

$$w = \sum_j y_j a_j$$

при всевозможных наборах положительных коэффициентов

$$y = \{y_1, \dots, y_m\} \geq 0$$

было заостренным конусом, не содержащим противоположно направленных векторов. Если же

$$u = \sum_j y_j a_j, \quad v = \sum_j z_j a_j$$

и $u = -v$, то $u + v = \sum_j (y_j + z_j) a_j = 0$.

Проведенное рассуждение указывает на справедливость следующего результата.

8.1.1. Теорема. *Либо неравенство $Ax > 0$ разрешимо, либо $yA = 0$ имеет ненулевое положительное решение.*

Если требуется строго положительное решение $x > 0$ неравенства $Ax > 0$, то это сводится к предыдущей задаче о разрешимости $Bz > 0$ с матрицей $B = \begin{bmatrix} A \\ I \end{bmatrix}$.

Задача о разрешимости неравенства $Ax > b$ тоже сводится к предыдущей — с матрицей

$$\begin{bmatrix} a_{11} & \dots & a_{1n} & -b_1 \\ a_{21} & \dots & a_{2n} & -b_2 \\ \dots & \dots & \dots & \dots \\ a_{m1} & \dots & a_{mn} & -b_m \\ 0 & \dots & 0 & 1 \end{bmatrix} x > 0.$$

◀ Если последнее неравенство имеет решение $\{x_1, \dots, x_{n+1}\}$, то решением является и

$$\left\{ \frac{x_1}{x_{n+1}}, \dots, \frac{x_n}{x_{n+1}}, 1 \right\}. \quad \blacktriangleright$$

8.1.2. Теорема. *Либо неравенство $Ax \geq 0$ разрешимо ($x \neq 0$), либо $yA = 0$ имеет строго положительное решение $y > 0$.*

◀ Для доказательства достаточно заметить, что $Ax \geq 0$ разрешимо, если $K(A)$ помещается в некотором полупространстве. В противном случае $K(A)$ совпадает со всем пространством R^n , и тогда для любого $u = \sum_j y_j a_j$, где все $y_j > 0$, существует противоположный вектор $v = \sum_j z_j a_j$, причем $y + z > 0$, и $(y + z)A = 0$. ▶

По существу теоремы 8.1.1, 8.1.2 в совокупности со стандартными приемами перехода к новым матрицам (неравенствам) охватывают многие результаты в области линейных неравенств (альтернативного характера «либо — либо»). Разнообразие вариантов в этой области, тем не менее, достаточно велико и формулировки результатов выглядят весьма непохожими друг на друга (см. раздел 8.6).

8.2. Выпуклые множества и конусы

Предыдущий раздел отражает упрощенческую позицию, каковая хороша в определенных обстоятельствах. Если смотреть глубже, то теория линейных неравенств довольно обширна, содержит тонкие результаты и уходит корнями в различные смежные дисциплины.

Наиболее естественно и продуктивно линейные неравенства трактуются в рамках выпуклого анализа [16]. Подробное обсуждение здесь неуместно. В то же время основные категории мышления выпуклого анализа становятся постепенно всеобщим достоянием, и их трудно обойти стороной.

Аффинные множества. Множество $A \subset R^n$ называется *аффинным*, если вместе с любыми двумя различными точками оно содержит проходящую через них прямую. В случае $0 \in A$ множество A является линейным пространством. Всякое аффинное множество — есть пересечение *гиперплоскостей*, описываемых уравнениями вида $ax = \beta$. Совокупность решений $Ax = b$, таким образом, это аффинное множество.

Отображение S из R^m в R^n , удовлетворяющее условию

$$S(\alpha x + \beta y) = \alpha Sx + \beta Sy,$$

называется *аффинным*. Аффинное отображение — это всегда преобразование вида $S(x) = Ax + b$, где A — линейное отображение.

Выпуклость. Множество $X \subset R^n$ называется *выпуклым*, если вместе с любыми двумя точками $x, y \in X$ оно содержит отрезок, их соединяющий,

$$\lambda x + (1 - \lambda)y \in X, \quad \lambda \in [0, 1],$$

а скалярная функция φ векторного аргумента считается *выпуклой*, если

$$\varphi[\lambda x + (1 - \lambda)y] \leq \lambda \varphi(x) + (1 - \lambda)\varphi(y), \quad \lambda \in [0, 1].$$

Линейную комбинацию $\sum_{j=1}^N \lambda_j f_j$ точек $\{f_1, \dots, f_N\}$, при условии $\lambda_j \geq 0$, $\sum_{j=1}^N \lambda_j = 1$, называют *выпуклой*.

Объединение всех выпуклых комбинаций конечных наборов точек из $X \subset R^n$ определяется как *выпуклая оболочка* множества X .

Теорема Каратеодори: выпуклая оболочка всегда совпадает с объединением всевозможных выпуклых комбинаций конечных подмножеств $X \subset R^n$, содержащих не более $n + 1$ точек.

Отделимость. Выпуклые множества $X, Y \subset R^n$ называются *отделимыми* (строغو отделимыми), если существует *разделяющая гиперплоскость* $H \subset R^n$ ($ax = \beta$), относительно которой множества располагаются в разных замкнутых (открытых) полупространствах H^+ , H^- :

$$x \in X \Rightarrow ax \leq \beta (< \beta), \quad x \in Y \Rightarrow ax \geq \beta (> \beta).$$

Теорема об отделимости. Непересекающиеся выпуклые множества X, Y всегда отделимы, и — строго отделимы, если хотя бы одно из этих множеств ограничено.

Опорной гиперплоскостью к X называют гиперплоскость H , которая имеет хотя бы одну общую точку с X и либо $X \subset H^+$, либо $X \subset H^-$.

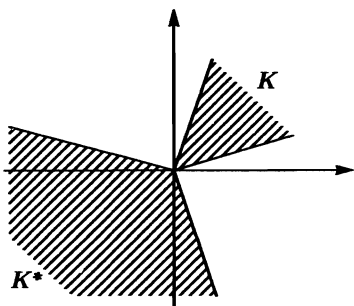


Рис. 8.1. Двойственный конус

Конусы. Выпуклое множество $K \subset R^n$, содержащее вместе с любой точкой $x \in K$ луч λx , проходящий через эту точку, называется **конусом**. Конус, не содержащий противоположных точек,

$$x \in K, x \neq 0 \Rightarrow -x \notin K,$$

называют **острым**, или **заостренным**.

Множество K^* , состоящее из точек y , для которых скалярное произведение $yx \leq 0$ при любом $x \in K$, называется **двойственным конусом** по отношению к K . На рис. 8.1 изображен пример на плоскости.

При изучении линейных неравенств важную роль играют следующие легко проверяемые свойства¹⁾:

- $(K^*)^* = K$.
- $K + K^* = R^n$.
- $K_1 \subset K_2 \Rightarrow K_1^* \supset K_2^*$.
- $(K_1 \cap K_2)^* = K_1^* + K_2^*$.
- $(K_1 + K_2)^* = K_1^* \cap K_2^*$.

Многогранники. Выпуклую оболочку конечного множества точек называют **многогранником**²⁾.

Пересечение F многогранника P с любой своей опорной гиперплоскостью называют **гранью многогранника**; k -гранью, если $\dim F = k$; 0-грани — **вершины**, 1-грани — **ребра**.

Каждая грань многогранника, в свою очередь, является многогранником. Всякий многогранник совпадает с выпуклой оболочкой своих вершин.

Пересечение конечного числа замкнутых полупространств называют **полиэдром**. Ограниченный полиэдр — **многогранник**.

Конусная техника. Докажем, для иллюстрации, следующее утверждение. *Либо существует ненулевой вектор $x \geq 0$ — такой, что $Ax \leq 0$, либо $yA > 0$ при некотором $y \geq 0$.*

◀ Множество решений системы $Ax \leq 0$ — есть конус $K = \{x : Ax \leq 0\}$, двойственный к которому

$$K^* = \{x : x = yA, y \geq 0\}.$$

¹⁾ Под суммой $X + Y$ подразумевается множество точек $z = x + y$, где $x \in X, y \in Y$.

²⁾ Или **выпуклым многогранником**, если рассматриваются еще и невыпуклые.

Совокупность положительных решений $Ax \leq 0$ является пересечением K с отрицательным ортантом R_+ . Если ненулевых положительных решений нет, то $K \cap R_+ = 0$, что влечет за собой $(K \cap R_+)^* = R^n$. Остается заметить,

$$(K \cap R_+)^* = K^* + R_+^* = R^n,$$

где, очевидно, R_+^* представляет собой отрицательный ортант R_- . Если теперь взять $u < 0$ ($u \in R_-$), то условие $K^* + R_- = R^n$ влечет за собой (в силу $0 \in R^n$)

$$yA = -u > 0, \quad y \geq 0. \quad \blacktriangleright$$

Инструментарий, таким образом, несколько иной. Он достаточно хорошо формализован и удобен, что становится ясно, правда, по приобретению необходимых навыков. Преимущества, однако, заключаются не только в технике, но и в результатах. Концентрация внимания на линейных неравенствах часто маскирует более фундаментальные причины. Во многих ситуациях линейность на самом деле ни при чем. Либо она ответственна лишь за второстепенную аранжировку результатов. Вот, например, один из общих нелинейных результатов [16], из которого легко выводится серия теорем о линейных альтернативах.

8.2.1. Пусть $\varphi_1, \dots, \varphi_m$ — выпуклые функции на выпуклом множестве C . Возможна лишь одна из двух альтернатив:

- существует такой вектор $x \in C$, что

$$\varphi_1(x) < 0, \quad \dots, \quad \varphi_m(x) < 0;$$

- существует ненулевой вектор $y = \{y_1, \dots, y_m\} \geq 0$ такой, что

$$y_1\varphi_1(x) + \dots + y_m\varphi_m(x) \geq 0 \quad \text{при любом } x \in C.$$

Линейные следствия можно получать, выбирая $\varphi_j(x) = \sum_k a_{jk}x_k + b_j$ и различные выпуклые множества в качестве C . Например, в случае $C = R_+$ можно утверждать: либо $Ax < 0$ положительно разрешимо, либо $yAx \geq 0$ при некотором ненулевом $y \geq 0$ и любом $x \geq 0$. Последнее означает существование такого ненулевого $y \geq 0$, что $yA \geq 0$.

На этом пути в качестве C могут фигурировать различные конусы,

$$K = \{x : Bx \geq 0\} \quad \text{либо} \quad K = \{x : x = By, Dy \geq 0\},$$

или — многогранники, $M = \{x : Bx \leq b\}$, — что дает обозримо трактуемые результаты в отдельных случаях.

8.3. Теоремы о пересечениях

Один из естественных подходов к разрешимости неравенств

$$f(x) = \{f_1(x), \dots, f_n(x)\} > 0$$

дают теоремы о непустоте пересечения³⁾.

8.3.1. Лемма⁴⁾. Пусть X — произвольное множество в R^n , $A(x)$ — точно-множественное отображение (из X в R^n) с компактными образами. Пусть для всякого конечного подмножества $\{x_1, \dots, x_k\} \subset X$ выпуклая оболочка точек $\{x_1, \dots, x_k\}$ содержится в объединении

$\bigcup_{i=1}^k A(x_i)$. Тогда

$$\bigcap_{x \in X} A(x) \neq \emptyset.$$

◀ В силу компактности образов достаточно показать, что не пусто пересечение любого конечного множества образов. Предположим противное, т. е.

$\bigcap_{i=1}^k A(x_i) = \emptyset$ для некоторого множества точек $\{x_1, \dots, x_k\}$. Тогда $\bigcup_{i=1}^k Y_i = R^n$, где Y_i обозначает дополнение к $A(x_i)$.

Пусть теперь $\rho_i(x)$, как говорят, разбиение единицы в R^n , согласованное с покрытием $\{Y_i\}$, т. е. такой набор функций, что⁵⁾

$$\sum_{i=1}^k \rho_i(x) \equiv 1, \quad \text{supp } \rho_i \subset Y_i.$$

Отображение $\varphi(x) = \sum_{i=1}^k \rho_i(x)x_i$, очевидно, преобразует в себя выпуклую

оболочку $\text{co}\{x_1, \dots, x_k\}$ и, по теореме Брауэра, имеет неподвижную точку x° . При этом $\rho_i(x^\circ) > 0$ для i из некоторого подмножества индексов $M \subset \{1, \dots, k\}$, и $\rho_i(x^\circ) = 0$ для $i \notin M$. Но тогда

$$x^\circ = \sum_{i \in M} \rho_i(x^\circ)x_i \in \text{co}\{x_i : i \in M\} \subset \bigcup_{i \in M} A(x_i).$$

Следовательно, $x^\circ \in A(x_i)$ для некоторого $i \in M$, т. е. $x^\circ \notin Y_i$, а значит, $\rho_i(x^\circ) = 0$, что приводит к противоречию. ►

³⁾ Если S_i обозначает множество точек x , для которых $f_i(x) > 0$, то разрешимость $f(x) > 0$ равносильна условию $\bigcap_i S_i \neq \emptyset$.

⁴⁾ Результат известен как лемма Кнастера—Кураковского—Мазуркевича (см., например, Ниренберг Л. Лекции по нелинейному функциональному анализу. М.: Мир, 1977).

⁵⁾ $\text{supp } \rho$ обозначает носитель функции $\rho(x)$, т. е. множество тех x , для которых $\rho(x) > 0$.

Другой тип утверждений о пересечениях дает классическая *теорема Хелли*.

8.3.2. Пусть S — семейство из не менее чем $n + 1$ выпуклых множеств в R^n , причем либо S конечно, либо каждое множество из S компактно. Тогда, если каждые $n + 1$ множеств из S имеют непустое пересечение, то пересечение всех множеств из S не пусто⁶⁾. ►

Проиллюстрируем возможности теоремы 8.3.2 на примере линейной модели производства (раздел 2.1), описываемой в целом системой уравнений

$$y = Ax \quad \left(y_i = \sum_{j=1}^n a_{ij}x_j \right)$$

с матрицей размера $m \times n$, $m \geq n$.

Здесь естественно возникает вопрос о продуктивности модели. Можно ли так подобрать интенсивности x_i технологических процессов, чтобы результирующее производство всех видов продуктов было строго положительным? Другими словами, существует ли положительное решение неравенства $Ax > 0$?

Теорема Хелли позволяет дать следующий ответ. *Если можно обеспечить строго положительное производство любых n видов продуктов, то можно обеспечить строго положительное производство всех видов продуктов.* ◀ Для доказательства достаточно рассмотреть на $(n - 1)$ -мерном симплексе $\{x : \|x\| = 1, x \geq 0\}$ семейство множеств

$$S_i = \left\{ x : \sum_{j=1}^n a_{ij}x_j > 0 \right\}, \quad i = 1, \dots, m. \quad \blacktriangleright$$

Упражнение

Пусть семейство S состоит из тысячи 100-мерных параллелепипедов с ребрами, параллельными осям координат. Все множества из S имеют общую точку, если общую точку имеют любые два из них.

8.4. *P*-матрицы

Квадратная $n \times n$ матрица A называется *P-матрицей*, если для любого ненулевого вектора x можно указать такой индекс j , что

$$x_j \cdot y_j > 0, \quad \text{где} \quad y = Ax.$$

Другими словами, *P*-матрица не обращает знак ни одного ненулевого вектора.

⁶⁾ Многочисленные доказательства и приложения имеются в [8].

8.4.1. Если A является P -матрицей, то неравенство $Ax > 0$ всегда имеет строго положительное решение $x > 0$.

Линейность здесь, вообще говоря, излишня. Результат сохраняет силу при естественном нелинейном обобщении. Оператор

$$f(x) = \{f_1(x), \dots, f_n(x)\}$$

назовем P -отображением, если для любого ненулевого вектора $x \geq 0$ можно указать такой индекс j , что $x_j \cdot f_j(x) > 0$.

8.4.2. Теорема⁷⁾. Если f — непрерывное P -отображение, то неравенство $f(x) > 0$ положительно разрешимо.

◀ Применим лемму 8.3.1, которая, заметим, остается справедливой при замене предположения о компактности образов $A(x)$ требованием их открытости и ограниченности. В качестве X возьмем n вершин симплекса $\Delta = \{x : \|x\| = 1, x \geq 0\}$

$$x_i = \{0, \dots, 0, 1, 0, \dots, 0\}$$

и положим

$$A(x_i) = \{x : f_i(x) > 0, x \in \Delta\}.$$

Из определения P -отображения следует, что выпуклая оболочка любого подмножества вершин симплекса Δ покрывается соответствующим объединением $\bigcup A(x_i)$, что позволяет (лемма 8.3.1) гарантировать непустоту пересечения всех $A(x_i)$. ▶

Упражнения

- Одновременная и одинаковая перестановка строк и столбцов преобразует P -матрицу в P -матрицу.
- Если A — P -матрица, $\Gamma = \text{diag} \{\gamma_1, \dots, \gamma_n\}$ и все $\gamma_j \neq 0$, то $\Gamma^{-1}A\Gamma$ — тоже P -матрица.

Остановимся на возможной содержательной интерпретации. Назовем переменные x_i действиями, y_i — результатами⁸⁾. Одинаковые номера за действиями как бы закрепляют «свои» результаты.

Понятно, что при наличии перекрестных взаимосвязей (побочных влияний) совместные стереотипные действия могут приводить к отрицательным последствиям.

⁷⁾ См. Karamardian S. Existence of solutions of certain systems of non-linear inequalities // Nimer. Math. 1968. 12. № 4. P. 327–334.

⁸⁾ Например, x_i — количество i -го лекарства, y_i — мера i -го симптома (температура, например).

миноры матрицы

$$\begin{bmatrix} \hat{a}_{22}(t) & \dots & \hat{a}_{2n}(t) \\ \dots & \dots & \dots \\ \hat{a}_{n2}(t) & \dots & \hat{a}_{nn}(t) \end{bmatrix}.$$

Далее срабатывает индуктивное предположение. ►

Утверждение 8.4.3 позволяет переопределить P -матрицы, положив в основу положительность главных миноров. Именно так P -матрицы были определены изначально¹²⁾.

8.5. Линейное программирование

Добавление к изучению неравенств оптимизационного аспекта проливает дополнительный свет, но уводит в другую область. Поэтому здесь рассматривается лишь эскиз одного ракурса.

Задача линейного программирования

$$(c, x) = \max, \quad Ax \leq b, \quad x \geq 0, \quad (8.1)$$

характеризуется линейностью *целевой функции* (c, x) и ограничений $\sum_j a_{ij}x_j \leq b_i$. Матрица A имеет размер $m \times n$ (m ограничений, n переменных).

Стандартная модификация

$$(c, u) = \max, \quad Bu = b, \quad u \geq 0 \quad (8.2)$$

получается из (8.1) добавлением фиктивных параметров $z_i \geq 0$ таких, что $\sum_j a_{ij}x_j + z_i = b_i$. В переменных $u = \begin{bmatrix} x \\ z \end{bmatrix}$ исходная задача (8.1) принимает вид (8.2).

Использование фиктивных переменных позволяет менять форму задачи до неузнаваемости. Например, вместо (c, x) максимизировать единственный скалярный параметр $x_{n+1} = \max$, добавляя к ограничениям условие

$$(c, x) - x_{n+1} = 0.$$

¹²⁾ Gale D., Nikaido H. The Jacobian matrix and global univalence of mappings // Math. Ann. 1965. 159. № 2.

Заметим, наконец, что варианты с минимизацией (c, x) ничего нового в постановку не привносят, ибо сводятся к максимизации после замены (c, x) на $(-c, x)$.

Линейная модель производства. Остановимся более подробно на модели производства (2.1). Ситуацию можно представлять себе так. На вход модели подается набор ресурсов $x = \{x_1, \dots, x_n\}$, на выходе образуется набор различных видов производимой продукции $y = \{y_1, \dots, y_m\}$. Внутреннее состояние производства характеризуется вектором $z = \{z_1, \dots, z_k\}$, компоненты которого принято называть интенсивностями технологических процессов, но терминология условна. В частности, z_i может обозначать количество рабочих, занятых в том или ином технологическом процессе; время работы некоторого агрегата; количество основного вида изготавливаемой продукции и т. д.

Как правило, имеются нормативы, определяющие взаимосвязи $x = Az$, $y = Bz$, а также ограничения по ресурсам $x \leq b$ и ассортименту продукции $y \geq d$. Возможные программы работы предприятия в свою очередь подвержены ограничениям. Агрегаты не могут работать больше чем 24 часа в сутки, рабочих нельзя задействовать в большем количестве, чем есть в наличии и т. п. Все это записывается в виде некоторой совокупности неравенств $Cz \leq e$. Сведение ограничений воедино дает $Gz \leq h$, $z \geq 0$.

Если цель производства заключается в максимизации объема реализации, $(c, y) = \max$, где c_i — цена i -го вида продукции, то после подстановки

$$(c, y) = (c, Bz) = (B^*c, z) = (g, z) = \max$$

возникает задача линейного программирования, приведенная к внутренним переменным z . При этом $g = B^*c$ также интерпретируется как вектор «внутренних» цен (g_i показывает, сколько в ценовом измерении продукции дает i -й технологический процесс, работающий с единичной интенсивностью).

Задача о диете. Имеется n продуктов $G_1, \dots, G_n$, каждый из которых содержит m питательных веществ (белки, витамины и прочее). Пусть x_i обозначает количество i -го продукта, c_i — цену G_i , a_{ij} — количество i -го питательного вещества в пищевом рационе. Задача минимизации стоимости пищевого рациона при ограничениях на содержание питательных веществ выглядит следующим образом

$$(c, x) = \max, \quad Ax \geq b, \quad x \geq 0,$$

где нормативы b определяются диетологами¹³⁾.

Геометрическая интерпретация. Каждое из неравенств

$$\sum_j a_{ij} x_j \leq b_i$$

¹³⁾ Либо традициями, либо (что чаще всего) «берутся с потолка».

определяет в R^n сдвинутое полупространство. Пересечение таких полупространств дает выпуклый многогранник¹⁴⁾, ограниченный или неограниченный (возможно, пустой). Линейная целевая функция (c, x) растет при движении точки x вдоль c . Поверхностями постоянного уровня, на которых функция (c, x) принимает одно и то же значение, являются гиперповерхности $(c, x) = \gamma$, ортогональные вектору c (рис. 8.2).

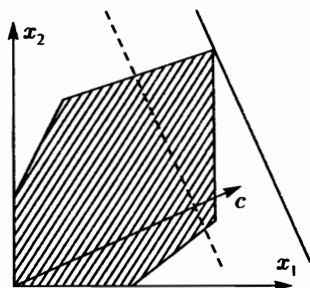


Рис. 8.2

Легко понять, что решение задачи линейного программирования, если существует, обязательно лежит в некоторой вершине допустимого многогранника¹⁵⁾. Решения может не быть, когда допустимый многогранник либо пуст, либо неограничен в направлении c .

Обратим внимание на определенную нечувствительность решения к данным задачи. Достаточно очевидно, что при малых изменениях вектора c оптимальное решение будет достигаться в той же самой вершине. Пределы нечувствительности решения к изменению критерия могут быть достаточно велики. Например, в ситуации, изображенной на рис. 8.2, решение одно и то же при любом положительном векторе c .

Двойственность. Вернемся к исходной задаче

$$(c, x) = \max, \quad Ax \leq b, \quad x \geq 0.$$

Двойственной к ней называется задача

$$(b, y) = \min, \quad A^*y \geq c, \quad y \geq 0.$$

Легко проверить, что двойственная к двойственной — снова исходная задача. Двойственные задачи тесно связаны между собой, и эта связь имеет принципиальное значение.

¹⁴⁾ Который называют *допустимым многогранником*, или *многогранником допустимых решений*.

¹⁵⁾ Решение не обязательно единственно. Когда одна из «дальних» граней допустимого многогранника ортогональна вектору c , — все точки этой грани (в том числе вершины) будут решениями.

8.5.1. Пусть x — допустимый вектор задачи (8.1), а y — допустимый вектор двойственной задачи. Тогда

$$(c, x) \leq (b, y).$$

◀ Таким образом, максимизируемая целевая функция в прямой (исходной) задаче всегда не больше, чем минимизируемая целевая функция двойственной задачи, если только эти функции вычисляются в точках, удовлетворяющих ограничениям соответствующих задач. Обоснование совсем просто. Умножая $Ax \leq b$ скалярно на y , а неравенство $A^*y \geq c$ — скалярно на x , получаем

$$(Ax, y) \leq (b, y), \quad (A^*y, x) \geq (c, x).$$

Откуда, в силу $(Ax, y) = (A^*y, x)$,

$$(c, x) \leq (Ax, y) \leq (b, y). \quad \blacktriangleright$$

Из утверждения 8.5.1 легко выводится следующий критерий оптимальности.

8.5.2. Пусть x° и y° — допустимые векторы прямой и двойственной задач, причем $(c, x^\circ) = (b, y^\circ)$. Тогда векторы x° , y° являются решениями соответствующих задач.

Утверждение 8.5.1 дает также ключ к пониманию основной теоремы двойственности.

8.5.3. Теорема. Если прямая и двойственная задача имеют допустимые векторы, то обе они имеют решения.

◀ При наличии допустимых векторов задача (8.1) неразрешима, если на допустимом многограннике функцию (c, x) выбором x можно сделать сколь угодно большой. Но по условию функция на допустимом многограннике ограничена: $(c, x) \leq (b, y)$.

Такие же рассуждения — с учетом того, что ищется минимум, а не максимум, — годятся для двойственной задачи. ▶

Активные и пассивные ограничения. Вернемся к рассмотрению пары двойственных задач:

$$(c, x) = \max, \quad Ax \leq b, \quad x \geq 0,$$

$$(b, y) = \min, \quad A^*y \geq c, \quad y \geq 0.$$

Пусть x°, y° — их оптимальные решения. Через $c(b)$ обозначим максимально возможное значение (c, x) . В силу утверждения 8.5.2,

$$c(b) = (c, x^\circ) = (b, y^\circ).$$

Приращение $\Delta c(b) = c(b + \Delta b) - c(b)$ при малом изменении вектора b , очевидно, равно

$$\Delta c(b) = (b + \Delta b, y^\circ + \Delta y) - (b, y^\circ) = (\Delta b, y^\circ) + (b, \Delta y) + (\Delta b, \Delta y),$$

где $y^\circ + \Delta y$ обозначает решение двойственной задачи с новой целевой функцией $(b + \Delta b, y)$ при старых ограничениях $A^*y \geq c$.

Как уже отмечалось, при малых изменениях b решение двойственной задачи не меняется, т.е. $\Delta y = 0$, что влечет за собой $\Delta c(b) = (\Delta b, y^\circ)$. В частном случае $\Delta b = \{0, \dots, \Delta b_i, 0, \dots, 0\}$ имеем $\Delta c(b) = y_i^\circ \Delta b_i$, т.е.

$$y_i^\circ = \frac{\Delta c(b)}{\Delta b_i}. \quad (8.3)$$

Таким образом, i -я компонента решения двойственной задачи показывает, во сколько раз приращение целевой функции больше, чем приращение Δb_i . Другим словами, y_i° это цена i -го ограничения.

При этом ясно, что если ограничение $\sum_j a_{ij} x_j^\circ \leq b_i$ не активно, т.е.

$\sum_j a_{ij} x_j^\circ < b_i$, то малые изменения величины b_i никакого влияния на решение x° не оказывают. Поэтому соответствующее $\Delta c(b) = 0$, а значит, и $y_i^\circ = 0$, в силу (8.3). На этом пути в итоге получается следующий результат.

8.5.4. Теорема. Для того чтобы допустимые векторы прямой и двойственной задач были решениями этих задач, необходимо и достаточно выполнение следующих условий:

$$\begin{cases} y_i^\circ = 0, & \text{если } \sum_j a_{ij} x_j^\circ < b_i; \\ x_j^\circ = 0, & \text{если } \sum_i a_{ij} y_i^\circ > c_j. \end{cases}$$

Понятно, что в решении исходной задачи часть неравенств обращается в равенства — эти ограничения называют *насыщающими* или *активными*, в противоположность *ненасыщающим*, или *пассивным*. Решение задачи (8.1) было бы предельно простым, если бы было известно, какие ограничения активные. В этом случае было бы достаточно неравенства, соответствующие активным ограничениям, заменить на равенства — и решить получившуюся систему линейных уравнений. Но решение двойственной задачи как раз однозначно указывает, какие ограничения активные. Поэтому, если известно двойственное решение, исходную задачу можно считать «почти решенной».

8.6. Задачи и дополнения

- Система линейных уравнений и неравенств,

$$Ax = 0, \quad x \geq 0, \quad yA \geq 0,$$

всегда имеет такое решение $\{x, y\}$, что $yA + x > 0$.

- Либо $Ax = b$ имеет решение $x \geq 0$, либо $yA \geq 0$ влечет за собой $yb \geq 0$ (лемма Минковского—Фаркаша)¹⁶⁾.
- Либо $Ax = b$ имеет решение $x \geq 0$, либо разрешима система неравенств $yA \geq 0$, $yb < 0$ (переформулировка предыдущего утверждения).
- Либо $Ax \leq b$ имеет решение $x \geq 0$, либо система неравенств $yA \geq 0$, $yb < 0$ имеет решение $y \geq 0$.
- Либо $Ax \geq b$ разрешимо, либо система $yA = 0$, $yb = 1$ имеет решение $y \geq 0$.
- Система неравенств

$$Ax \leq 0, \quad x \geq 0, \quad yA \geq 0, \quad y \geq 0$$

всегда имеет такое решение, что

$$x + yA > 0, \quad y - Ax > 0.$$

- Если в прямоугольной области X матрица Якоби $\partial f_i / \partial x_j$ всюду является P -матрицей, то отображение $f(x)$ взаимно однозначно в X [14].
- Если в прямоугольной области X матрица Якоби $\partial f_i / \partial x_j$ всюду является P -матрицей, то отображение $f(x)$ в области X является универсальным P -отображением.
- Введем обозначения¹⁷⁾

$$h(A, B) = \{C : C = tA + (1 - t)B, \quad t \in [0, 1]\},$$

$$r(A, B) = \{C : C = TA + (I - T)B, \quad T = \text{diag} \{t_1, \dots, t_n\} \quad t_j \in [0, 1]\},$$

$$i(A, B) = \{C = [c_{ij}] : c_{ij} = t_{ij}a_{ij} + (1 - t_{ij})b_{ij}, \quad t_{ij} \in [0, 1]\}.$$

- Все матрицы из $h(A, B)$ невырождены в том и только том случае, когда BA^{-1} не имеет отрицательных собственных значений.
- Если все матрицы из $i(A, B)$ невырождены, то AB^{-1} и BA^{-1} являются P -матрицами.
- Любая матрица из $r(A, B)$ невырождена в том и только том случае, когда BA^{-1} — P -матрица.

¹⁶⁾ Дополнительную информацию о линейных неравенствах можно найти в сборнике [12].

¹⁷⁾ По поводу результатов данного пункта см. Johnson Ch., Tsatsomeros M. Convex sets of nonsingular and P-matrices // Linear and Multilinear Algebra. 1995. 38. P. 233–239.

Глава 9

Положительные матрицы

Уникальные свойства положительных матриц выглядят тем удивительнее, чем дальше объяснение от истинных причин. Основные результаты в рассматриваемой области базируются на явлениях полуупорядоченности и монотонности. С этого ракурса предмет выглядит не так удивительно, зато более понятно.

9.1. Полуупорядоченность и монотонность

Заостренный конус $K \subset R^n$ позволяет ввести в R^n *полуупорядоченность*: $x \geq y$, если $x - y \in K$. Если $x - y$ лежит внутри K (для чего требуется телесность конуса), пишут $x > y$.

Основная линия изложения опирается далее на ситуацию $K = R_+$, где R_+ обозначает неотрицательный ортант¹⁾. В этом случае $x > y$ означает покоординатные неравенства $x_j > y_j$ для всех j .

Монотонность отображения²⁾ $f(x)$ означает:

$$x \geq y \Rightarrow f(x) \geq f(y).$$

Оператор $f(x)$ называют *положительным*, если

$$f(K) \subset K,$$

т. е. $x \geq 0$ влечет за собой $f(x) \geq 0$.

В случае

$$x > 0 \Rightarrow f(x) > 0$$

говорят, что оператор *сильно положителен*.

¹⁾ Множество векторов x с неотрицательными координатами.

²⁾ Имеется в виду отображение R^n в R^n .

Итерационный процесс

$$x_{k+1} = f(x_k)$$

с монотонным оператором f обладает следующим характерным свойством. Если точка x_0 под действием оператора f идет *вперед* (*назад*), т. е. $f(x_0) \geq x_0$ ($\leq x_0$), то вся последовательность x_k монотонно возрастает (убывает),

$$x_0 \leq \dots \leq x_k \leq \dots \quad (x_0 \geq \dots \geq x_k \geq \dots).$$

◀ Это сразу вытекает из применения к неравенству $x_k \geq x_{k-1}$ оператора f , что дает $x_{k+1} \geq x_k$ и т. д. ▶

Если теперь есть две точки $v_0 \leq w_0$, «меньшая» из которых идет вперед, а «большая» — назад, то для

$$v_{k+1} = f(v_k), \quad w_{k+1} = f(w_k)$$

имеем

$$v_0 \leq \dots \leq v_k \leq \dots \leq w_k \leq \dots \leq w_0.$$

Обе последовательности v_k, w_k оказываются ограниченными и в случае непрерывного f сходятся к неподвижным точкам $v^* = f(v^*), w^* = f(w^*)$. А если f имеет лишь одну неподвижную точку x^* , то любая последовательность x_k процесса $x_{k+1} = f(x_k)$ при условии $v_0 \leq x_0 \leq w_0$ оказывается зажатая в тиски,

$$v_k \leq x_k \leq w_k,$$

и потому тоже сходится, $x_k \rightarrow x^*$.

Заметим, что множество точек x , удовлетворяющих неравенствам

$$v \leq x \leq w,$$

называется *конусным отрезком* и обозначается $\langle v, w \rangle$.

9.2. Теорема Перрона

Матрицу A с неотрицательными элементами, $a_{ij} \geq 0$, назовем *положительной*, с положительными, $a_{ij} > 0$, — *строго положительной*.

Под $|A|$ подразумевается матрица с элементами $|a_{ij}|$, где, вообще говоря, имеется в виду модуль комплексного числа³⁾. Аналогично,

$$|x| = \{|x_1|, \dots, |x_n|\}.$$

9.2.1. Лемма. Пусть $A \geq 0$, $x_0 > 0$ и

$$\alpha x_0 \leq A x_0 \leq \beta x_0 \quad (<).$$

Тогда $\alpha \leq \rho(A) \leq \beta$ ($<$)⁴⁾.

◀ Из $Ax_0 \leq \beta x_0$ следует, что оператор $B = \frac{1}{\beta}A$ отображает в себя конусный отрезок $\langle -x_0, x_0 \rangle$. Отсюда вытекает ограниченность всех итераций B^k , что влечет за собой $\|(\gamma B)^k\| \rightarrow 0$ при $k \rightarrow \infty$ и любом $\gamma \in (0, 1)$. Следовательно (см. 6.3.2), $\rho(A) \leq \gamma \Rightarrow \rho(A) \leq \beta$.

Далее,

$$\alpha x_0 \leq A x_0 \Rightarrow \alpha \|x_0\| \leq \|A\| \cdot \|x_0\|,$$

т. е. $\alpha \leq \|A\|$, какова бы ни была норма. Но для любой матрицы всегда существует норма, сколь угодно близкая к спектральному радиусу (см. лемму 6.3.1). Поэтому $\alpha \leq \rho(A)$. При замене неравенств на строгие — рассуждения обрастают незначительными оговорками. ▶

Замечание. При рассмотрении векторных неравенств часто возникают неудобные ситуации. В том смысле, что вопрос, как говорится, не стоит выведенного яйца, но плодит многословие.

Рассмотрим в качестве примера лемму 9.2.1. Пусть $A > 0$ и существует $x > 0$, при котором $Ax \leq x$, но $Ax \neq x$. Лемма утверждает $\rho(A) \leq 1$. На самом деле $\rho(A) < 1$. Обоснование просто, но несколько громоздко.

◀ Без ограничения общности можно считать, что в $Ax \leq x$ первые несколько неравенств строгие, остальные — равенства. В блочном виде это записывается так,

$$\begin{cases} A_{11}x_1 + A_{12}x_2 < x_1, \\ A_{21}x_1 + A_{22}x_2 = x_2. \end{cases} \quad (9.1)$$

³⁾ $|A|$ обозначает также детерминант матрицы, но в данной главе используется только в смысле $|A| = [a_{ij}]$.

⁴⁾ Знак ($<$) говорит о том, что утверждение сохраняет силу при замене нестрогих неравенств на строгие в указанных местах.

Поскольку $A_{21}x_1 > 0$, то $A_{22}x_2 < x_2$. Прибавляя тогда $A_{22}\varepsilon x_2 < \varepsilon x_2$ ($\varepsilon > 0$) к равенству в (9.1) и увеличивая в неравенстве (9.1) x_2 на εx_2 , приходим к

$$\begin{cases} A_{11}x_1 + A_{12}(1 + \varepsilon)x_2 < x_1, \\ A_{21}x_1 + A_{22}(1 + \varepsilon)x_2 < (1 + \varepsilon)x_2, \end{cases}$$

если только ε достаточно мало⁵⁾. Таким образом, небольшое «шевеление» вектора $x = \{x_1, x_2\}$, связанное с переходом к вектору $\{x_1, (1 + \varepsilon)x_2\}$, позволяет заменить нестрогое неравенство $Ax \leq x$ строгим $Ax' < x'$. ►

Подобного рода «шевеление» неравенств, векторов и матриц дает эффект во многих ситуациях. Форма, правда, бывает разная, хотя идея одна и та же. Так или иначе, но ссылка на «метод шевеления» (оговариваемая или подразумеваемая) иногда спасает от тяжеловесных пояснений.

(!) В общем случае произвольной матрицы, неравенство $Ax \leq x$ не обязано шевелением икса приводиться к строгому.

9.2.2. Пусть $A \geq 0$ и

$$-A \leq B \leq A \quad (<).$$

Тогда $\rho(B) \leq \rho(A)$ ($<$).

◀ Очевидно, $|B| \leq A$. Далее

$$|B| \leq A \Rightarrow |B^k| \leq |B|^k \leq A^k \Rightarrow \|B^k\|_m \leq \|A^k\|_m.$$

Теперь остается сослаться на формулу $\rho(A) = \lim_{k \rightarrow \infty} \sqrt[k]{\|A^k\|}$, для которой годится любая норма (см. 6.3.2). ► Опробование метода шевеления предоставляется в качестве упражнения.

9.2.3. Если B главная подматрица⁶⁾ матрицы $A \geq 0$, то

$$\rho(B) \leq \rho(A).$$

Если $A > 0$, то $\rho(B) < \rho(A)$.

◀ Пусть матрица C получается из A обнулением всех элементов, не входящих в B . Тогда $\rho(B) = \rho(C)$, после чего решает ссылка на предыдущее утверждение. ►

⁵⁾ Чтобы не нарушить первого неравенства в (9.1).

⁶⁾ Подматрица называется главной, если ее элементы стоят на пересечении строк и столбцов исходной матрицы — с номерами из одного и того же подмножества индексов $\{1, \dots, n\}$.

Дальнейшее рассмотрение предполагает строгую положительность матрицы A .

1°. Понятно, что оператор

$$f(x) = \frac{Ax}{\|Ax\|_l},$$

где $\|x\|_l = \sum_i |x_i|$, отображает в себя симплекс

$$\sum_i |x_i| = 1, \quad \text{все } x_i \geq 0,$$

и по теореме Брауэра⁷⁾ имеет неподвижную точку, $\frac{Ax^0}{\|Ax^0\|_l} = x^0$.

Иными словами, у A существует положительное собственное значение и положительный собственный вектор⁸⁾,

$$Ax^0 = \lambda x^0, \quad \lambda = \|Ax^0\|_l.$$

2°. Вектор x^0 , ясно, строго положителен, поэтому⁹⁾

$$\|A\|_{mM} = \sum_{j=1}^n |a_{ij}| \frac{x_j^0}{x_i^0} = \lambda, \quad (9.2)$$

где подразумевается $M = \text{diag}\{x_1^0, \dots, x_n^0\}$. Теперь, с одной стороны (раздел 6.3), $\|A\|_{mM} \geq \rho(A)$, с другой — $\lambda \leq \rho(A)$. Поэтому найденное позитивное собственное значение равно спектральному радиусу,

$$\boxed{\lambda = \rho(A)}. \quad (9.3)$$

Собственное значение (9.3) называют *ведущим* и обозначают как λ_A . Очевидно, $\lambda_A = \lambda_{A^*}$.

3°. Допустим, что у A в R_+ собственному значению $\rho(A)$ отвечают два разных собственных вектора x_1 и x_2 , которые, силу $A > 0$, строго положительны. Пусть

⁷⁾ Теорема Брауэра (см. [4, т. 1]) гарантирует существование неподвижной точки у любого непрерывного оператора, отображающего в себя выпуклое, замкнутое и ограниченное множество.

⁸⁾ То же самое можно сказать о любом нелинейном строго положительном операторе.

⁹⁾ См. раздел 6.2.

$\alpha > 0$ обозначает число, максимальное в неравенстве $x_2 - \alpha x_1 \geq 0$. Поскольку $x_2 - \alpha x_1 \neq 0$ и $A > 0$, то $A(x_2 - \alpha x_1) > 0$, а значит $A(x_2 - \alpha x_1) \geq \beta x_1$ при некотором $\beta > 0$, т. е.

$$x_2 - \alpha x_1 \geq \frac{\beta}{\rho(A)} x_1,$$

что противоречит максимальности α и исключает возможность существования у A линейно независимых собственных векторов¹⁰⁾ $x_1, x_2 > 0$.

Возможность $x_2 \not\geq 0$ также исключена, поскольку тогда при достаточно малых $\tau > 0$ у A все равно будет другой положительный собственный вектор

$$(1 - \tau)x_1 + \tau x_2.$$

4°. Возможно ли существование у $A > 0$ других собственных значений λ , равных по модулю $\rho(A)$? Нет.

Если $Ax = \lambda x$, то либо $x = |x|e^{i\varphi}$, и тогда $A|x| = \lambda|x|$, что влечет за собой $\lambda > 0$, либо комплексные координаты x неодинаково направлены, и тогда

$$|\lambda| \cdot |x| = |\lambda x| = |Ax| < A|x|,$$

что приводит к противоречию $\rho(A) > \rho(A)$ (лемма 9.2.1).

5°. Выше уже было показано, что геометрическая кратность собственного значения $\lambda_A = \rho(A)$ равна единице. То же самое имеет место и для алгебраической кратности.

◀ Пусть A_k обозначает подматрицу, получаемую из A вычеркиванием k -й строки и k -го столбца. Производная

$$\frac{d|A - \lambda I|}{d\lambda} = -|A_1 - \lambda I| - \dots - |A_n - \lambda I|$$

в точке $\lambda = \lambda_A$ не равна нулю, поскольку¹¹⁾ $\lambda_A > \lambda_{A_k}$, — в результате чего все слагаемые $|A_k - \lambda_{A_k} I|$ имеют один знак. ►

По совокупности сказанного можно сделать следующие выводы.

¹⁰⁾ Возможность существования $x_2 > 0$, отвечающего другому положительному $\lambda \neq \rho(A)$, исключена рассуждениями предыдущего пункта, где установлена импликация

$$Ax = \lambda x, x > 0, \lambda > 0 \Rightarrow \lambda = \rho(A).$$

¹¹⁾ См. п. 9.2.3.

9.2.4. Теорема Перрона. Пусть $A > 0$. Тогда:

- Спектральный радиус $\rho(A) > 0$ является собственным значением матрицы A алгебраической кратности единица, и ему отвечает строго положительный собственный вектор.
- Других положительных собственных значений и векторов матрица A не имеет.
- Если $Ax = \lambda x$, $\lambda \neq \rho(A)$, $x \neq 0$, то $|\lambda| < \rho(A)$.

Обратим внимание, что $\rho(A) = 0$ для $A \geq 0$ возможно лишь в том случае, когда $A^p = 0$ при некотором p .

9.3. Неразложимость

Теорема Перрона опирается на предположение $A > 0$. В случае $A \geq 0$ картина усложняется: $\rho(A)$ остается собственным значением, которому отвечает собственный вектор $x \geq 0$ (уже не обязательно $x > 0$). Остальная часть выводов рушится.

Все восстанавливается, если $A^k > 0$ при некотором k , — поскольку $\lambda(A^k) = \lambda^k(A)$. Но это слишком просто. Поиск трудностей разворачивается в других направлениях.

Матрица A называется *разложимой*¹²⁾ (*неразложимой*), если одинаковой перестановкой строк и столбцов она приводится (не может быть приведена) к виду

$$\begin{bmatrix} A_{11} & A_{12} \\ 0 & A_{22} \end{bmatrix},$$

где A_{11} и A_{22} — квадратные матрицы.

Иными словами, A неразложима, если не существует такого подмножества индексов J , что $a_{ik} = 0$ для всех $i \in J$, $k \notin J$.

Система уравнений $Ax = b$ с разложимой матрицей, по сути, имеет вид

$$\begin{bmatrix} A_{11} & A_{12} \\ 0 & A_{22} \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \end{bmatrix} = \begin{bmatrix} b_1 \\ b_2 \end{bmatrix},$$

т. е.

$$A_{11}x_1 + A_{12}x_2 = b_1,$$

$$A_{22}x_2 = b_2.$$

¹²⁾ Разложимую матрицу называют также *приводимой*.

Наличие автономной подсистемы $A_{22}x_2 = b_2$, которую можно решать независимо, — характеристическое свойство разложимой матрицы.

Понятие неразложимости иногда используется при изучении матриц общего вида, но в основном все же — это атрибут теории положительных матриц, что далее подразумевается. *Расположение ненулевых элементов в $A \geq 0$ однозначно определяет расположение ненулевых элементов в степенях A^k .*

Неразложимость матрицы $A = [a_{ij}]$ равносильна существованию *дорожки невырожденности*, каковой называют последовательность ненулевых элементов

$$a_{i_1 i_2}, a_{i_2 i_3}, \dots, a_{i_k i_1},$$

где среди индексов $i_1, i_2, \dots, i_k$ имеются все числа $1, 2, \dots, n$.

Вот еще два эквивалента. Неразложимость $A \geq 0$ равносильна либо неравенству $(I + A)^{n-1} > 0$, либо существованию для любой пары индексов i, j такого k , что $a_{ij}^{(k)} > 0$, где $a_{ij}^{(k)}$ обозначает (ij) -й элемент матрицы A^k . Заметим, что отсюда не вытекает существование k , при котором $A^k > 0$.

Неразложимая матрица *положительно невырождена*: $Ax \neq 0$ при любом ненулевом $x \geq 0$.

Если главная диагональ неразложимой матрицы A строго положительна, то $A^{n-1} > 0$.

9.3.1. Если матрица $A \geq 0$ неразложима, то $\rho(A) > 0$ является ведущим собственным значением A алгебраической кратности 1, которому отвечает строго положительный собственный вектор. Других положительных собственных значений и векторов у A нет.

Из сравнения со сказанным в начале раздела видно, что неразложимость добавляет к предположению $A \geq 0$ не так много¹³⁾. Превосходство λ_A над остальными собственными значениями (по модулю) не достигается.

¹³⁾ Впрочем, как смотреть.

9.4. Положительная обратимость

Если $B = I - A$, $A \geq 0$ и $\rho(A) < 1$, то матрица B положительно обратима, что сразу следует из сходимости ряда ¹⁴⁾

$$B^{-1} = (I - A)^{-1} = I + A + A^2 + \dots$$

Условие $\rho(A) < 1$ в данном случае можно заменить эквивалентным: существуют такие положительные константы $\mu_1, \dots, \mu_n > 0$, что (см. разделы 6.2, 9.2) либо

$$\|A\|_{mM} = \max_i \frac{1}{\mu_i} \sum_{j=1}^n \mu_j |a_{ij}| < 1, \quad (9.4)$$

либо

$$\|A\|_{lM} = \max_j \mu_j \sum_{i=1}^n \frac{1}{\mu_i} |a_{ij}| < 1.$$

◀ Первое неравенство означает $A\mu < \mu$, где $\mu = \{\mu_1, \dots, \mu_n\}'$, второе — $A^*\mu < \mu$. Далее решает ссылка на лемму 9.2.1 ▶

Если $\lambda = \rho(A)$, то $(1 - \lambda)^{-1} = 1 + \lambda + \lambda^2 + \dots$ — ведущее собственное значение матрицы B^{-1} , а $(1 - \lambda)$ — наименьшее по модулю собственное значение матрицы B .

Положительная обратимость $I - A$ иногда проявляется в завуалированной форме, порождая эффект неожиданности. Следующий факт, например, нередко дает ощущение глубины.

Если система $x - Ax = y$ имеет решение $x > 0$ хотя бы при одном $y > 0$, то она имеет положительное решение $x \geq 0$ при любом ¹⁵⁾ $y \geq 0$.

Вся хитрость здесь заключается в том, что из существования $x > 0$, при котором $x - Ax = y > 0$, следует (лемма 9.2.1)

$$Ax < x \Rightarrow \rho(A) < 1.$$

Поэтому $x = (I - A)^{-1}y \geq 0$ при любом $y \geq 0$.

¹⁴⁾ Из-за очевидной положительности всех степеней $A^k \geq 0$.

¹⁵⁾ В применении к линейной модели межотраслевого баланса $x - Ax = y$ это звучит интригующе. Если возможен хотя бы один набор чистых выпусков $y > 0$ (хотя бы один продуктивный план), то возможен и любой другой.

Очевидным образом разрешается вопрос о положительной обратимости *внедиагонально отрицательной матрицы* $B = [b_{ij}]$ с положительной диагональю, характеризуемой неравенствами

$$\text{все } b_{ii} > 0, \quad b_{ij} \leq 0 \quad \text{при } i \neq j.$$

Понятно, что после деления каждой строки $b_{i\cdot}$ на b_{ii} задача сводится к предыдущей. При этом вопрос о справедливости $B^{-1} \geq 0$ заменяется эквивалентным вопросом о положительной обратимости матрицы $B' = I - A$, $A \geq 0$,

$$B' = DB, \quad D = \text{diag} \{b_{11}^{-1}, \dots, b_{nn}^{-1}\}.$$

Форма утвердительного ответа может быть, например, такой.

9.4.1. *Внедиагонально отрицательная матрица B с положительной диагональю положительно обратима, если*¹⁶⁾

$$\max_i \frac{1}{\gamma_i} \sum_{j \neq i} |b_{ij}| \gamma_j < |b_{ii}| \quad (9.5)$$

при некотором наборе констант $\gamma_1, \dots, \gamma_n > 0$.

Условие (9.5) в общем случае¹⁷⁾ характеризует матрицы, имеющие *доминирующую диагональ*. Если «все $b_{ii} > 0$ ($b_{ii} < 0$)», — говорят о положительной (отрицательной) доминирующей диагонали. Если в (9.5) все $\gamma_j = 1$, матрицу B называют *матрицей Адамара*. Интерес к таким матрицам объясняется тем, что они невырождены и, в случае *отрицательной* доминирующей диагонали, — устойчивы (см. след. раздел).

Ж-матрицы. Матрицу A назовем *Ж-матрицей*, если отображение A не переводит граничные точки неотрицательного ортанта R_+ внутрь R_+ .

Очевидно, образ неотрицательного ортанта AR_+ представляет собой некоторый конус в R^n . Если граничные точки R_+ не переходят внутрь R_+ , имеются две возможности: или пересечение $R_+ \cap AR_+$ не содержит внутренних точек R_+ , или конус AR_+ охватывает R_+ , т. е.

$$R_+ \subset AR_+. \quad (9.6)$$

¹⁶⁾ Результат естественным образом перекликается с (9.4).

¹⁷⁾ Без предположений о внедиагональной отрицательности или положительности B .

Легко видеть, что в случае (9.6) матрица A невырождена и положительно обратима, поскольку уравнение $Ax = y$ при любом положительном y имеет положительное решение. Ясно, что альтернатива (9.6) реализуется лишь в том случае, когда уравнение $Ax = y$ имеет положительное решение хотя бы при одном $y > 0$.

9.4.2. Для положительной обратимости матрицы A необходимо и достаточно выполнение двух условий: A является K -матрицей, уравнение $Ax = y$ имеет положительное решение хотя бы при одном $y > 0$.

Элементарные следствия:

- Пусть $C - B \geq 0$, матрица C положительно обратима и $Bu > 0$ для некоторого $u \geq 0$. Тогда $B^{-1} \geq 0$.
- Пусть $C \geq B \geq D$, причем матрицы C и D положительно обратимы. Тогда $B^{-1} \geq 0$.

9.5. Оператор сдвига и устойчивость

Рассмотрим задачу Коши

$$\dot{x} = A(t)x, \quad x(0) = x_0,$$

решением которой является $x(t) = X(t)x(0)$, где $X(t)$ — так называемый матрицант (см. [4, т. 2]), представляющий собой решение задачи

$$\dot{X} = A(t)X, \quad X(0) = I.$$

Другими словами, матрица $X(t)$ — это оператор сдвига U_{t0} по траекториям $\dot{x} = A(t)x$.

9.5.1. Линейная система $\dot{x} = Ax$ асимптотически устойчива¹⁸⁾, если (и только если) действительные части всех собственных значений матрицы A строго отрицательны.

◀ Сходимость всех решений $x(t) \rightarrow 0$ при $t \rightarrow \infty$ очевидна из представления $x(t) = \sum c_k e^{\lambda_k t}$. Если все λ_k различны, то c_k константы. При равных λ_k коэффициенты c_k могут полиномиально расти, но это не может перевесить экспоненциальное убывание множителей $e^{\lambda_k t}$. ▶

Матрица A , у которой все собственные значения λ_k находятся строго в левой полуплоскости, $\operatorname{Re} \lambda_k < 0$, — называется *устойчивой*, или *гурвицевой*.

¹⁸⁾ См. [4, т. 2].

Утверждение 9.5.1 гарантирует асимптотическую устойчивость нулевого равновесия линейной системы $\dot{x} = Ax$ с устойчивой матрицей.

Полиномы $P(\lambda)$, все корни λ_k которых удовлетворяют условию $\operatorname{Re} \lambda_k < 0$, также называют *устойчивыми*, либо *гурвицевыми*, а задачу об устойчивости $P(\lambda)$ — *проблемой Рауса—Гурвица* (см. [4, 15]).

При подстановке $\lambda = i\omega$ в

$$P(\lambda) = \lambda^n + a_1 \lambda^{n-1} + \dots + a_{n-1} \lambda + a_n$$

многочлен P записывается в форме $|P(i\omega)|e^{i \arg P(i\omega)}$, либо

$$P(i\omega) = g(\omega) + ih(\omega), \quad (9.7)$$

где g и h — многочлены от ω . Например, в случае $P(\lambda) = \lambda^2 + a_1 \lambda + a_2$,

$$g(\omega) = -\omega^2 + a_2, \quad h(\omega) = a_1 \omega.$$

Кривая (9.7) в комплексной плоскости, задаваемая параметром ω , называется *амплитудно-фазовой характеристикой* полинома, либо *годографом Михайлова*.

9.5.2. Если аргумент $P(i\omega)$ при изменении ω от 0 до $+\infty$ меняется

на величину $\boxed{\pi \frac{n}{2}}$, то многочлен $P(\lambda)$ устойчив.

Это так называемый *критерий Михайлова* [4, 15] — один из простых и эффективных методов проверки устойчивости полинома. Теория устойчивых многочленов достаточно популярна, ибо в миниатюре востребована практикой, а в общем виде служит удобным источником задач. В силу общих причин, однако, — см. главу 10 — об устойчивости надежнее судить непосредственно по матрице A , минуя запись характеристического полинома (конечно, если это возможно).

9.5.3. Лемма. Если матрица $A(t)$ внедиагонально положительна при любом t , то $X(t) \geq 0$.

◀ Траектории вспомогательной системы

$$\dot{x} = A(t)x + \varepsilon \quad (9.8)$$

со строго положительным вектором $\varepsilon > 0$ — не могут выйти из неотрицательного ортанта, так как в граничных точках направлены строго внутрь R_+ . Поэтому

решение $x(t)$ остается в R_+ при условии, что $x(0) \in R_+$, т. е. $X_\epsilon(t)x(0) \geq 0$, если $x(0) \geq 0$. При $\epsilon \rightarrow 0$ траектории (9.8) переходят в траектории $\dot{x} = A(t)x$. ►

Ограничимся далее рассмотрением автономного случая. Решением $\dot{x} = Ax$ служит $x(t) = e^{At}x(0)$. Положительность экспоненты e^{At} , при условии $A \geq 0$, очевидна из обыкновенного разложения в ряд (6.7). Лемма 9.5.3 показывает, что для $e^{At} \geq 0$ достаточно внедиагональной положительности матрицы A .

9.5.4. Лемма. Если траектория линейной системы $\dot{x} = Ax$ с внедиагонально положительной матрицей A начинает двигаться вперед (назад), то она все время будет двигаться вперед (назад).

◄ Умножая неравенство $\dot{x}(0) = Ax(0) \leq 0$ слева на положительную матрицу e^{At} , получаем $e^{At}Ax(0) \leq 0$, и далее, поскольку A и e^{At} коммутируют,

$$Ae^{At}x(0) \leq 0 \Rightarrow \dot{x}(t) = Ax(t) \leq 0,$$

т. е. из $\dot{x}(0) \leq 0$ следует $\dot{x}(t) \leq 0$ при любом $t \geq 0$. Аналогично,

$$\dot{x}(0) \geq 0 \Rightarrow \dot{x}(t) \geq 0 \quad (t \geq 0). \quad \blacktriangleright$$

9.5.5. Лемма. Траектории $u(t)$, $v(t)$ системы $\dot{x} = Ax$ с внедиагонально положительной матрицей A , находящиеся в начальный момент в положении $u(0) \leq v(0)$, остаются в том же отношении $u(t) \leq v(t)$ при любом $t \geq 0$.

◄ В силу положительности (а значит — монотонности) e^{At} ,

$$u(0) \leq v(0) \Rightarrow e^{At}u(0) \leq e^{At}v(0). \quad \blacktriangleright$$

Следующие шаги. Если A внедиагонально положительна, причем $A = B - I$, $B \geq 0$ и $\rho(B) < 1$, то матрица A гурвицева, поскольку ее собственные значения $\lambda_j(A) = \lambda_j(B) - 1$. Поэтому все $\operatorname{Re} \lambda_j(A) < 0$.

В более общем случае внедиагонально положительной матрицы A с отрицательной доминирующей диагональю вывод о гурвицевости A сделать уже не так легко. Запишем A в виде $A = B - \Gamma$, где $B \geq 0$ — внедиагональная часть A (на диагонали у B стоят нули), а

$$\Gamma = \operatorname{diag} \{\gamma_{11}, \dots, \gamma_{nn}\}, \quad \text{все } \gamma_{ii} > 0. \quad (9.9)$$

Умножение A слева на Γ^{-1} вроде бы сводит задачу к предыдущей, поскольку в $\Gamma^{-1}A = \Gamma^{-1}B - I$ исходное предположение

о доминирующей диагонали дает $\rho(\Gamma^{-1}B) < 1$, что обеспечивает устойчивость $A' = \Gamma^{-1}A$. Но из устойчивости A' , вообще говоря, не следует устойчивость $A = \Gamma A'$, ибо о спектре произведения некоммутирующих матриц трудно судить в общем случае. Тем не менее в данной ситуации матрица $A = \Gamma A'$ гурвицева.

◀ В силу $\rho(\Gamma^{-1}B) < 1$ матрица $\Gamma^{-1}B$ имеет собственный вектор $x_0 \geq 0$,

$$\Gamma^{-1}Bx_0 = \rho(\Gamma^{-1}B)x_0.$$

«Метод шевеления» позволяет утверждать существование такого $x_0 > 0$ (для простоты обозначение оставляем тем же), что $\Gamma^{-1}Bx_0 < x_0$, т. е.

$$Bx_0 < \Gamma x_0 \Rightarrow Ax_0 < 0.$$

Таким образом, траектория $x(t)$ уравнения $\dot{x} = Ax$, начинающаяся в точке x_0 , все время убывает (лемма 9.5.4) при сохранении неравенства $x(t) \geq 0$ (лемма 9.5.3), что влечет за собой $x(t) \rightarrow 0$ ($t \rightarrow \infty$). Наконец, лемма 9.5.5 гарантирует, что любая траектория $u(t)$, удовлетворяющая в начальный момент неравенству $-x(0) \leq u(0) \leq x(0)$, остается при любом $t \geq 0$ в пределах $-x(t) \leq u(t) \leq x(t)$, и потому $u(t) \rightarrow 0$. А поскольку система линейна, то сходимость всех траекторий к нулю влечет за собой асимптотическую устойчивость, что и означает гурвицевость матрицы A . ▶

Устойчивость при доминировании. Рассмотрим две матрицы с одинаковой отрицательной диагональю,

$$A = B - \Gamma, \quad G = H - \Gamma,$$

где B и H имеют нулевые диагонали, Γ определяется как (9.9). При этом $B \geq 0$, а знаковая структура H произвольна.

9.5.6. Лемма. Пусть $|H| \leq B$, $\dot{u} = Au$, $\dot{v} = Gv$. Тогда

$$u(0) \geq v(0) \Rightarrow u(t) \geq v(t)$$

при любом $t \geq 0$.

Это весьма примечательный и полезный результат из разряда дифференциальных неравенств.

◀ Векторное неравенство $u(t) \geq v(t)$ не может нарушиться. Как только $u_j(t_0) = v_j(t_0)$, так производная $\dot{u}_j(t_0) - \dot{v}_j(t_0) \geq 0$, в силу $B \geq |H|$. ▶

Лемма 9.5.6 легко позволяет сделать вывод об устойчивости матрицы $G = H - \Gamma$ с отрицательной доминирующей диагональю. Достаточно взять $B = |H|$ и применить лемму 9.5.6 в совокупности с установленной выше гурвицевостью $B - \Gamma$.

9.6. Импримитивность

Неразложимая матрица $A \geq 0$ может иметь несколько *периферических* собственных значений $\{\lambda_1, \dots, \lambda_k\}$, равных по модулю $\rho(A)$. Такие матрицы называют *импримитивными*¹⁹⁾.

Здесь возникает в некотором роде неожиданность. Оказывается, что все периферические λ_j обязаны иметь вид $\omega_j \rho(A)$, где ω_j — комплексные корни k -й степени из 1. Вопрос подробно излагается в [14, 18].

Примеры. У матрицы циклической перестановки $\begin{bmatrix} 0 & 1 & 0 \\ 0 & 0 & 1 \\ 1 & 0 & 0 \end{bmatrix}$ все три собственных значения равны корням кубическим из единицы. Ничего по существу не меняется в ситуации $\begin{bmatrix} 0 & a & 0 \\ 0 & 0 & b \\ c & 0 & 0 \end{bmatrix}$ с произвольными $a, b, c > 0$ ²⁰⁾. Замена любого нуля единицей превращает матрицу в примитивную.

Рассмотрим аналогичный вопрос для других n . Видению результатов здесь помогает следующая модель. Напомним предварительно, что если у $A \geq 0$ и $B \geq 0$ ненулевые элементы стоят на одинаковых местах, то

$$A > 0 \Leftrightarrow B > 0, \quad \text{равно как и} \quad A \neq 0 \Leftrightarrow B \neq 0.$$

Поэтому, изучая вопрос о строгой положительности A^k (положительной матрицы), можно всегда рассматривать матрицу из нулей и единиц как бы определяющей тип изучаемой матрицы.

Представим далее модель, в которой имеется n ячеек, расположенных по кругу, и заполненных, скажем, шариками в соответствии с расположением единиц в векторе $\{1, 0, 0, 1, \dots, 0\}$. Действие матрицы «0-1» (из нулей и единиц) на «0-1»-векторы можно ассоциировать с движением шариков, включая их размножение. Например, циклическая матрица один шарик будет просто гонять по кругу.

Но если у циклической перестановки, скажем $\begin{bmatrix} 0 & 1 & 0 \\ 0 & 0 & 1 \\ 1 & 0 & 0 \end{bmatrix}$, какой-либо нуль заменить на 1, например $\begin{bmatrix} 0 & 1 & 1 \\ 0 & 0 & 1 \\ 1 & 0 & 0 \end{bmatrix}$, то количество шариков со временем будет возрастать. Если в какой-то момент $x_3 = 1$, то в следующий момент окажутся заполнены 2-я и 3-я ячейки. В итоге заполнятся все ячейки, что будет означать $A^k > 0$.

¹⁹⁾ В отличие от — *примитивных*, имеющих только одно собственное значение с $|\lambda| = \rho(A)$. Необходимое и достаточное условие примитивности — строгая положительность некоторой степени A^k .

²⁰⁾ Собственные значения здесь равны трем корням $\sqrt[3]{abc}$.

В четырехмерном пространстве ситуация меняется. Матрица

$$A = \begin{bmatrix} 0 & 1 & 0 & 1 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \\ 1 & 0 & 0 & 0 \end{bmatrix} \text{ имеет собственные значения}$$

$$\lambda_{1,2} = \pm 1,272, \quad \lambda_{3,4} = \pm 0,786i.$$

Направления двух периферических $\lambda_{1,2}$ определяются корнями $\sqrt{1}$, а $\lambda_{3,4}$ — корнями $\sqrt{-1}$. Процесс увеличения числа строго положительных координат у положительного вектора при воздействии матрицы A здесь закликивается. Но в случае, когда к той же циклической перестановке одна единица добавляется иначе, $a_{13} = 1$, — снова $A^k > 0$ при некотором k .

Для $n = 5$ ситуация вновь упрощается. Достаточным оказывается добавление единицы вместо любого нуля. Например,

$$A = \begin{bmatrix} 0 & 1 & 0 & 0 & 1 \\ 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 1 \\ 1 & 0 & 0 & 0 & 0 \end{bmatrix} \Rightarrow A^{11} = \begin{bmatrix} 4 & 2 & 3 & 1 & 4 \\ 2 & 1 & 1 & 1 & 1 \\ 1 & 2 & 1 & 1 & 3 \\ 3 & 1 & 2 & 1 & 2 \\ 2 & 3 & 1 & 2 & 4 \end{bmatrix},$$

A^{10} еще имеет один нулевой элемент.

В другом варианте первая строго положительная степень только 17,

$$A = \begin{bmatrix} 0 & 1 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 1 \\ 1 & 0 & 0 & 0 & 0 \end{bmatrix} \Rightarrow A^{17} = \begin{bmatrix} 4 & 1 & 2 & 4 & 6 \\ 3 & 1 & 1 & 1 & 3 \\ 3 & 3 & 4 & 1 & 1 \\ 1 & 3 & 6 & 4 & 1 \\ 1 & 1 & 4 & 6 & 4 \end{bmatrix}.$$

Естественная гипотеза: если n простое, то для $A^k > 0$ при некотором k достаточно (в указанном выше смысле) у циклической матрицы перестановки заменить один нуль (любой) на единицу. (?)

9.7. Стохастические матрицы

Положительная матрица $A \geq 0$, у которой все столбцовые суммы $\sum_i a_{ij} = 1$, называется *стохастической*.

Термин побуждает думать, что речь идет о случайных матрицах, но все детерминированно. Просто источником интереса к таким матрицам первоначально была вероятностная задача. Имеется n состояний, и «нечто» находится в k -й момент времени в j -м состоянии с вероятностью p_j^k . Вероятность попадания

этого «нечто» в $(k+1)$ -й момент в i -е состояние определяется суммой²¹⁾

$$p_i^{k+1} = \sum_j a_{ij} p_j^k,$$

т. е. $\mathbf{p}^{k+1} = A\mathbf{p}^k$. Чтобы сумма вероятностей $\sum_i p_i^{k+1}$ была равна единице, необходимо как раз $\sum_i a_{ij} = 1$.

Естественный вопрос, конечно, предельное поведение итераций $\mathbf{p}^k$. Относительно прост случай примитивной матрицы A , в котором следующее по модулю собственное значение строго меньше λ_A , т. е. $|\lambda_2| < \lambda_A = 1$. Тогда итерации $\mathbf{p}^k$ сходятся к собственному вектору $\hat{\mathbf{p}}$ матрицы A , причем

$$\|\mathbf{p}^k - \hat{\mathbf{p}}\| \sim |\lambda_2|^k,$$

а итерации $A^k \rightarrow A_\infty$, где у A_∞ все столбцы одинаковы и равны $\hat{\mathbf{p}}$.

Если матрица не стохастическая, форма аналогичного результата обрастает деталями, но суть та же.

9.7.1. Пусть $A > 0$ и

$$A\mathbf{x} = \rho(A)\mathbf{x}, \quad A^*\mathbf{y} = \rho(A)\mathbf{y}, \quad \mathbf{x} > 0, \quad \mathbf{y} > 0, \quad \mathbf{x}^*\mathbf{y} = 1.$$

Тогда при $k \rightarrow \infty$

$$[\rho^{-1}(A)A]^k \rightarrow A_\infty = \mathbf{x}\mathbf{y}^*. \quad \blacktriangleright$$

В случае импримитивной неотрицательной матрицы предел у $[\rho^{-1}(A)A]^k$ может не существовать. Но предел имеют средневзвешенные суммы

$$\frac{1}{N} \sum_{k=1}^N [\rho^{-1}(A)A]^k \rightarrow \mathbf{x}\mathbf{y}^*.$$

Если обе матрицы, A и транспонированная A^* , стохастические, то матрицу A называют *двойкостохастической*. Классическая *теорема Биркгофа* утверждает, что любая двойкостохастическая матрица есть выпуклая комбинация конечного числа матриц перестановок.

²¹⁾ Величина a_{ij} интерпретируется как вероятность попадания из j -го состояния в i -е.

9.8. Конус положительно определенных матриц

Легко проверить, что множество K неотрицательно определенных матриц — конус. Полуупорядочим с помощью K пространство E симметричных матриц. Неравенства $A \succcurlyeq 0$, $A \succcurlyeq B$, таким образом, означают, что матрицы A , $A - B$ неотрицательно определены. Рассмотрим оператор

$$H(x) = xx^*,$$

действующий из R^n в E . Его производная вдоль траекторий линейной системы дифференциальных уравнений $\dot{x} = Ax$ равна

$$\frac{d}{dt}(xx^*) = xx^*A^* + Ax x^*,$$

т. е.

$$\dot{H} = HA^* + AH. \quad (9.10)$$

Уравнение (9.10) — частный случай (7.10). Его решением служит

$$H(t) = e^{At}H(0)e^{A^*t}.$$

Из $H_1(0) \succcurlyeq H_2(0)$ следует

$$H_1(t) - H_2(t) = e^{At}[H_1(0) - H_2(0)]e^{A^*t} \succcurlyeq 0,$$

что означает монотонность оператора сдвига по траекториям матричного дифференциального уравнения (9.10). Неравенство $e^{At}Be^{A^*t} \succcurlyeq 0$ для любой неотрицательно определенной матрицы B вытекает из

$$(e^{At}Be^{A^*t}x, x) = (Be^{A^*t}x, e^{A^*t}x) \geq 0$$

при любом x .

Для асимптотической устойчивости нулевого равновесия системы с монотонным оператором сдвига достаточно, чтобы нашлась точка $X \in \text{int } K$, идущая под действием оператора сдвига назад, и точка $Y \in -\text{int } K$, идущая вперед²²⁾. Наличие таких точек в данном случае гарантирует разрешимость в $\text{int } K$ уравнения

$$H_0A^* + AH_0 = -G_0 \quad (9.11)$$

при $G_0 \in \text{int } K$.

²²⁾ Обоснование этого факта почти не отличается от схемы рассуждений в ситуации $K = R_+^n$.

В случае гурвицевой матрицы A интеграл

$$H_0 = \int_0^{\infty} e^{At} G_0 e^{A^*t} dt$$

сходится и является решением (9.11).

В результате оказалось, что самый общий случай устойчивости линейной системы укладывается в рамки идеологии полуупорядоченности. Это в какой-то степени свидетельствует, что последняя не так узка, как иногда кажется.

Упражнения (см. [18])

Матрицы A, B считаются симметричными (эрмитовыми).

- Если $A \succ 0$, то

$$A \succcurlyeq B \succcurlyeq 0 \Leftrightarrow \rho(BA^{-1}) \leq 1.$$

- Для любой матрицы T , в том числе прямоугольной,

$$A \succcurlyeq B \Rightarrow T^*AT \succcurlyeq T^*BT.$$

•

$$A \succcurlyeq B \succ 0 \Leftrightarrow B^{-1} \succcurlyeq A^{-1} \succ 0.$$

- Если $A \succcurlyeq B \succ 0$ и собственные значения A и B упорядочены по возрастанию, то $\lambda_j(A) \geq \lambda_j(B)$.

9.9. Задачи и дополнения

- Функционал

$$\rho(x, y) = \ln \min \left\{ \frac{\beta}{\alpha} : \alpha x \leq y \leq \beta x \right\}$$

представляет собой метрику на множестве лучей, лежащих во внутренности неотрицательного ортанта. (?) Легко проверяется, что

$$\rho(x, y) = \ln \min_{j,k} \frac{x_j y_k}{x_k y_j}.$$

Линейный оператор, описываемый строго положительной матрицей A , сжимает по этой метрике (?)

$$\rho(Ax, Ay) \leq \theta \rho(x, y), \quad \theta < 1.$$

См. о фокусирующих операторах в [10].

- Матрицу A называют *вполне неразложимой*, если не существует таких матриц перестановок P и Q , что

$$PAQ = \begin{bmatrix} A_{11} & A_{12} \\ 0 & A_{22} \end{bmatrix},$$

где A_{11} и A_{22} квадратные матрицы²³⁾.

Матрица A вполне неразложима тогда и только тогда, когда существует такая матрица перестановки P , что главная диагональ PA строго положительна и PA неразложима.

Если A вполне неразложима, то существует k , при котором $A^k > 0$.

- Неразложимая матрица $A \geq 0$ примитивна, если хотя бы один диагональный элемент $a_{ii} > 0$.
- Любая симметричная матрица с положительной доминирующей диагональю положительно определена. Все главные миноры любой внедиагонально отрицательной матрицы с положительной доминирующей диагональю — строго положительны (результат известен как *условие Хокинса—Саймона* [14]).
- Положительно обратимая внедиагонально отрицательная матрица является P -матрицей.
- В случае гурвицевой матрицы A спектральный радиус $\rho(e^A) < 1$. Это означает, что оператор сдвига e^{At} ($t = 1$) по траекториям $\dot{x} = Ax$ сжимающий (в некоторой норме).
- Спектральный радиус $\rho(A) < 1$ в том и только том случае, когда существует положительно определенная матрица V такая, что $A^*VA - V$ отрицательно определена.
- Если положительная матрица $[a_{ij}]$ неразложима, то система уравнений

$$\lambda x_j = \max_k a_{jk} x_k, \quad j, k = 1, \dots, n,$$

имеет единственное положительное решение ($x \geq 0$, $\lambda > 0$)²⁴⁾.

²³⁾ В отличие от обыкновенной неразложимости перестановки строк и столбцов не обязаны быть одинаковыми, т. е. Q может быть не равно P^* .

²⁴⁾ См. Varat R. B. A max version of the Perron—Frobenius theorem // Linear Algebra and its Applications. 1998. 275—276. P. 3—18.

Глава 10

Численные методы

Облик «вычислительной математики» в значительной мере определяется причинами, которые давно перестали действовать. Сведение к минимуму объема арифметических операций, требуемой памяти, времени счета — все эти факторы совсем недавно были определяющими. А уж если говорить о временах ручного счета, откуда численные методы берут начало, то многие загадочные ныне акценты становятся понятны.

Сегодня системы уравнений «тысяча-на-тысячу» легко решает обыкновенный ноутбук, а удобный интерфейс превращает вычисления в задачу, не требующую понимания внутренних механизмов. И тогда возникает естественный вопрос: «а кому эта внутренняя кухня нужна?»

Разумеется, — тем, кто будет заниматься «этим» профессионально. И тем, как ни странно, кто заниматься «этим» профессионально не будет. Неосведомленность о подводных течениях рано или поздно дает о себе знать. Поэтому необходимость понимания, что там за кадром, всегда «висит на стене» — как ружье Станиславского — и время от времени стреляет.

10.1. Предмет изучения

Интуитивные ожидания и мнения опираются, как правило, на примитивные примеры — вплоть до уравнений, которые приходилось решать в седьмом классе. От этой коллекции уже не избавиться, но к ней неплохо добавить что-нибудь еще. Например, при решении уравнений в частных производных пространство переменных разбивается на клетки (кубики), в узлах (вершинах кубиков) записываются приближенные связи переменных — получаются тысячи и даже миллионы линейных уравнений. Для решения таких систем, оказывается, нужна совсем другая философия.

Издавна иногда кажется, что практические задачи линейной алгебры сводятся к решению систем уравнений. А спектры, базисы и другие премудрости в большей степени проистекают из любопытства ¹⁾.

¹⁾ Здесь полезно вспомнить $x(t) = c_1 e^{\lambda_1 t} + \dots + c_n e^{\lambda_n t}$, дающее решение линейной системы $\dot{x} = Ax$, где λ_j — собственные значения, а c_j — собственные векторы матрицы A .

Доля правды здесь есть, но в этом нет ничего предвзятого. Математика, как и все остальное, развивается на основе определенной игры. Одни придумывают задачи, другие — решают, третьи — учат, четвертые — втирают очки и греют руки у того же огня. Самое удивительное, что процесс не замыкается внутри себя, питает другие игры (вплоть до космических перелетов), — и потому оказывается фундаментально связан с жизнью всей цивилизации, как бы высокопарно это ни звучало.

Численные методы в линейной алгебре несут на себе отпечаток общей ситуации. В сектор прицела попадает многое, что с вычислениями связано лишь опосредованно. В первую очередь, это изучение линейной алгебры под другим углом зрения — конструктивным. То есть изучение всего, что потенциально может пригодиться. КПД здесь очень низок, зато широта охвата решает проблему занятости²⁾, совершенствуется понимание предмета и развивается чутье. Вторая мишень утилитарная — решение практических задач. При этом важно понимать, что собой представляют реальные задачи, откуда они берутся и в каком виде являются.

Возьмем простой пример. Модель межотраслевого баланса³⁾:

$$x - Ax = y.$$

На выпуск единицы i -го продукта, выпускаемого i -й отраслью в количестве x_i , в системе затрачивается $a_{ij} \geq 0$ единиц j -го продукта.

При изучении жизни только по книгам фантазия в постановке задачи не простирается дальше решения системы уравнений $x - Ax = y$ относительно x . Реальность шире. С матрицей затрат A , описывающей модель, приходится иметь дело из года в год, — что означает необходимость многократного решения задачи. Сегодня надо обеспечить один набор чистых выпусков y , «завтра» — другой. Какие-то коэффициенты a_{ij} известны приближенно. Это порождает свой круг проблем. От оценок точности решения до попыток организации децентрализованного решения задачи на базе итерационного обмена данными,

$$x_{k+1} = Ax_k + y,$$

исходя из предположения, что «на местах» информация точнее. «Вниз» спускается вариант плана x_k , там прикидывают необходимые затраты, отсылают «наверх», далее вычисляется новый вариант x_{k+1} — и так несколько раз. Сходится ли процедура? Как быстро?

²⁾ Как бы издевательски это ни звучало, поддерживать творческий процесс без топлива невозможно.

³⁾ См. раздел 2.1.

Так или иначе, становится ясно, что для овладения ситуацией в целом необходимо близкое знакомство с матрицей A . Проникновение в ее внутреннее устройство, понимание, как различные коэффициенты влияют на интегральные свойства системы. Матрица затрат задается ведь не на небесах, а зависит от используемых технологических процессов, потерь от неорганизованности и т. п. Поэтому в окрестности решаемых задач всегда маячат проблемы другого уровня. Как перестроить систему? На чем сосредоточить внимание, ресурсы?

Так что вычислительная практика — это совсем другая епархия. Теоретические акценты очень сильно меняются. Простейшая задача решения системы $Ax = b$, например, сама по себе обычно не интересует, хотя иногда так кажется. Причина довольно проста. Когда порядок системы, скажем, всего 10×10 , но от записи матрицы уже рябит в глазах, голому решению едва ли кто обрадуется, поскольку голое решение — кот в мешке⁴⁾. Хотя бы потому, что матрица может быть близка к вырожденной. Тогда ошибка больше самого решения — и проще было не решать, взяв что-нибудь от фонаря, получилось бы не хуже.

Из всего сказанного нетрудно сделать вывод, что всякое решение практической задачи должно проходить под творческим надзором куратора, мыслящего шире, чем это оговорено в алгоритме. Куратора, способного «в случае чего» переосмыслить постановку задачи и направить усилия, может быть, совсем в другое русло.

Если говорить открытым текстом, то успех дела в первую очередь определяется «до-» и «надматематическим» этапом. В каждом проекте нужен человек, который вникает, разбирается, переосмысливает — и ставит задачу⁵⁾. Потом без математики не обойдешься, но важно, чтобы решались задачи «те, что нужно», а не взятые с потолка⁶⁾. На практике в 99 % случаев решается «лишь бы что». Но это нормально. Так устроена природа. Из миллионов сперматозоидов, движущихся «лишь бы куда», один попадает в цель — и для поддержания интереса к численным методам этого оказывается достаточно.

10.2. Ошибки счета и обусловленность

В принципе ясно, что практическое решение систем уравнений $Ax = b$ сталкивается с трудностями, если матрица A близка к вы-

⁴⁾ Особенно, если размер матрицы $10^4 \times 10^4$ и в глазах не рябит, потому что никому в голову не придет выводить такую матрицу на бумагу из памяти.

⁵⁾ Люди, которые могут разобраться и поставить задачу — товар штучный, большая редкость.

⁶⁾ Разумеется, только в сказке ситуация аккуратно делится на два этапа: постановка, потом решение. Реально все переплетено и перепутано. Поэтому нужна здравая направляющая рука.

рожденной. Но что значит «близка»? До сих пор фигурировали черно-белые тона «вырождена — не вырождена». Каким параметром можно характеризовать нейтральную полосу?

Матрица $\begin{bmatrix} 0,99 & 1 \\ 1 & 1,01 \end{bmatrix}$ «плоха», потому что изменение ее элементов на 1 % может увести решение $Ax = b$ в бесконечность. Ее детерминант близок к нулю, $\det A = 10^{-4}$, — но служит ли это индикатором? Не служит. Потому что у матрицы $\begin{bmatrix} 0,01 & 0 \\ 0 & 0,01 \end{bmatrix}$ детерминант такой же, но это «хорошая» матрица, потому что изменение ее элементов на 1 % дает в решении относительную ошибку того же порядка.

Скрытый механизм определяется *числом обусловленности*.

Априорная оценка. Пусть

$$Ax = b \quad \text{и} \quad A(x + \Delta x) = b + \Delta b,$$

где Δx — смещение решения $Ax = b$ при возмущении правой части на Δb . Под возмущением исходных данных может подразумеваться что угодно. Ошибки при физическом измерении параметров, неточный прогноз, ошибки счета и округления.

◀ Очевидно,

$$\|b\| = \|Ax\| \leq \|A\| \cdot \|x\|, \quad \|\Delta x\| = \|A^{-1} \Delta b\| \leq \|A^{-1}\| \cdot \|\Delta b\|,$$

откуда

$$\frac{\|\Delta x\|}{\|x\|} \leq \|A\| \cdot \|A^{-1}\| \frac{\|\Delta b\|}{\|b\|} = \sigma(A) \frac{\|\Delta b\|}{\|b\|}.$$

►

Пусть теперь

$$Ax = b \quad \text{и} \quad (A + \Delta A)(x + \Delta x) = b,$$

где Δx по-прежнему смещение решения $Ax = b$, но уже при возмущении матрицы на ΔA .

◀ Аналогично предыдущему

$$\|\Delta x\| = \|A^{-1} \Delta A(x + \Delta x)\| \leq \|A^{-1}\| \cdot \|\Delta A\| \cdot \|x + \Delta x\|,$$

откуда

$$\frac{\|\Delta x\|}{\|x + \Delta x\|} \leq \|A\| \cdot \|A^{-1}\| \frac{\|\Delta A\|}{\|A\|} = \sigma(A) \frac{\|\Delta A\|}{\|A\|}.$$

►

В том и другом случае определяющую роль играет величина

$$\sigma(A) = \|A\| \cdot \|A^{-1}\|,$$

называемая *числом обусловленности*⁷⁾. В случае вырожденной матрицы полагают $\sigma(A) = \infty$.

Для невырожденной матрицы и евклидовой нормы,

$$\sigma(A) = \frac{\alpha_1(A)}{\alpha_n(A)},$$

где $\alpha_1(A)$ — максимальное сингулярное число, $\alpha_n(A)$ — минимальное. Для симметричной матрицы

$$\sigma(A) = \frac{|\lambda_1(A)|}{|\lambda_n(A)|},$$

где собственные значения $\lambda_1, \dots, \lambda_n$ расположены в порядке убывания модулей.

Число обусловленности всегда ≥ 1 , поскольку

$$I = AA^{-1} \Rightarrow 1 = \|AA^{-1}\| \leq \|A\| \cdot \|A^{-1}\| = \sigma(A).$$

Чем $\sigma(A)$ больше, тем хуже оценки. В этом случае говорят о *плохой обусловленности* матрицы.

Роль понятия обусловленности матрицы рассмотренными задачами не исчерпывается.

Апостериорная оценка. Приведенные выше оценки решения опирались только на информацию о возмущениях ΔA , Δb . Подход может быть иной. Уравнение $Ax = b$ так или иначе приближенно решается, что дает $Ax' \approx b$. Далее вывод о близости x' к настоящему решению x делается на основе вектора невязки $\delta = Ax' - b$.

◀ В силу $\|b\| = \|Ax\| \leq \|A\| \cdot \|x\|$ и $A^{-1}\delta = x' - A^{-1}b = x - x'$, имеем

$$\|x\| \geq \frac{\|b\|}{\|A\|}, \quad \|x - x'\| \leq \|A^{-1}\| \cdot \|\delta\|.$$

⁷⁾ Термин введен Тьюрингом.

Деление второго неравенства на первое приводит к

$$\frac{\|x - x'\|}{\|x\|} \leq \sigma(A) \frac{\|\delta\|}{\|b\|}. \quad \blacktriangleright$$

10.3. Оценки сверху и по вероятности

Оценки предыдущего раздела — это оценки сверху, дающие максимально пессимистический прогноз. Реальные ошибки могут быть меньше. Не в нашей Вселенной, правда. У нас бутерброд падает маслом вниз, и вообще, всякая потенциальная неприятность случается. Подтверждением тому в данном контексте служит ряд работ, где устанавливается, что вероятные ошибки в естественных предположениях близки к наихудшим⁸⁾. Воронка неприятностей, получается, затягивает.

◀ Вот простые наводящие соображения из [18]. Если матрица A с нормой порядка единицы плохо обусловлена, то в $A^{-1} = [a'_{ij}]$ обязательно найдутся большие элементы⁹⁾. Дифференцируя $x = A^{-1}b = [a'_{ij}]b$ по b_j , получаем

$$\frac{\partial x_i}{\partial b_j} = a'_{ij},$$

откуда ясен источник происхождения больших ошибок,

$$\Delta x_i \sim a'_{ij} \Delta b_j.$$

Что касается возмущений самой матрицы A , то здесь «возни» чуть больше. Дифференцирование тождеств

$$\sum_i a'_{is} a_{si} = \delta_{ii}$$

приводит к соотношениям

$$\sum_i \left\{ \frac{\partial a'_{is}}{\partial a_{jk}} a_{si} + \delta_{sj} \delta_{ik} a'_{is} \right\} = \sum_i \frac{\partial a'_{is}}{\partial a_{jk}} a_{si} + \delta_{ik} a'_{ij} = 0,$$

⁸⁾ См. например, *Красносельский М. А., Крейн С. Г.* Замечание о распределении ошибок при решении системы линейных уравнений при помощи итерационного процесса // УМН. 1952. 7, вып. 4.

⁹⁾ Чтобы норма $\|A^{-1}\|$ была большая. Иначе $\sigma(A)$ будет малю.

учет которых при дифференцировании $x = A^{-1}b = [a'_{ij}]b$ дает

$$\begin{aligned}\frac{\partial x_i}{\partial a_{jk}} &= \sum_s \frac{\partial a'_{is}}{\partial a_{jk}} b_s = \sum_s \sum_t \frac{\partial a'_{is}}{\partial a_{jk}} a_{st} x_t = \\ &= \sum_t \left\{ \sum_s \frac{\partial a'_{is}}{\partial a_{jk}} a_{st} \right\} x_t = - \sum_t \delta_{tk} a'_{ij} x_t = -a'_{ij} x_k.\end{aligned}$$

Окончательно,

$$\frac{\partial x_i}{\partial a_{jk}} = -a'_{ij} \sum_s a'_{ks} b_s.$$

Прогнозы опять неутешительные. ►

10.4. Возмущения спектра

Принципиальную роль в вопросах численного анализа играет влияние возмущений матрицы на ее спектр. Здесь возможны удивительные явления, способные пустить насмарку многие численные расчеты.

Остановимся пока на вспомогательном результате.

10.4.1. Лемма. *Все собственные значения матрицы A лежат в объединении кругов Гершгорина*

$$|z - a_{ii}| \leq R_i \quad (i = 1, \dots, n),$$

где

$$R_i = \sum_{j \neq i} |a_{ij}|, \quad \text{либо} \quad R_i = \sum_{j \neq i} |a_{ji}|.$$

◀ Если λ_k — собственное значение A , то $\det(A - \lambda_k I) = 0$, что влечет за собой $|\lambda_k - a_{ii}| \leq R_i$ хотя бы для одного i . В противном случае (все $|\lambda_k - a_{ii}| > R_i$) матрица $A - \lambda_k I$ имеет доминирующую диагональ¹⁰⁾, и потому — невырождена, что исключает возможность $\det(A - \lambda_k I) = 0$. ►

10.4.2. Теорема. *Если матрица A диагонализуема,*

$$A = T^{-1} \Lambda T, \quad \Lambda = \text{diag} \{ \lambda_1, \dots, \lambda_n \},$$

¹⁰⁾ См. утверждение 9.4.1 с последующим замечанием.

и λ' — собственное значение возмущенной матрицы $A + \Delta A$, то найдется такое λ_i , что

$$|\lambda' - \lambda_i| \leq \|T\| \cdot \|T^{-1}\| \cdot \|\Delta A\| = \sigma(A) \|\Delta A\|,$$

где $\|\cdot\|$ обозначает одну из двух норм — строчную или столбцовую.

◀ Поскольку матрицы

$$A + \Delta A \quad \text{и} \quad T(A + \Delta A)T^{-1} = \Lambda + T\Delta AT^{-1}$$

имеют одинаковые собственные значения, то проблема сводится к возмущению диагональной матрицы Λ . Пусть c_{ij} обозначают элементы матрицы $T\Delta AT^{-1}$.

Лемма 10.4.1 гарантирует, что все собственные значения $\Lambda + T\Delta AT^{-1}$ лежат в кругах

$$|z - \lambda_i - c_{ii}| \leq \sum_{j \neq i} |c_{ij}|,$$

которые, в свою очередь, содержатся в кругах

$$|z - \lambda_i| \leq \sum_j |c_{ij}|,$$

что влечет за собой оценку

$$|\lambda' - \lambda_i| \leq \|T\Delta AT^{-1}\|_m \leq \|T\|_m \cdot \|T^{-1}\|_m \cdot \|\Delta A\|_m.$$

Аналогично получается неравенство для нормы $\|\cdot\|_l$. ▶

Таким образом, при возмущении ΔA возмущением спектра A замаскировано управляет матрица T , приводящая A к диагональному виду. Задним числом это становится понятно. Но с самого начала догадаться не так легко, о чем свидетельствует история развития численных методов в докомпьютерную эпоху.

Многие разработанные тогда методы красиво выглядели на бумаге¹¹⁾, но при столкновении с вычислительной практикой на ЭВМ, сопровождаемой неизбежными ошибками округления, — были полностью вытеснены из алгоритмического арсенала. В первую очередь это коснулось методов определения спектра матрицы, опиравшихся на вычисление характеристического полинома $p_A(\lambda)$ с последующим определением его корней. Обе задачи, вычисления коэффициентов и корней $p_A(\lambda)$, таят в себе много ловушек и характеризуются очень сильной зависимостью результатов счета от исходных данных и точности манипуляций.

В [19] приводится такой пример. Если у полинома

$$p(\lambda) = (\lambda - 1)(\lambda - 2) \cdots (\lambda - 20),$$

¹¹⁾ И почти не проверялись фактически, поскольку решать вручную задачи хотя бы «10 на 10» никому не приходило в голову.

имеющего корни $\{1, 2, \dots, 20\}$, изменить коэффициент при x^{19} с -210 всего лишь на $-210 + 2^{-23}$, то у возмущенного полинома появляется, в частности, пара комплексно сопряженных корней с мнимыми частями $\sim \pm 3i$.

Теперь достаточно выписать сопровождающую матрицу многочлена $p(\lambda)$, — как получится пример матрицы 20×20 , у которой изменение одного элемента на величину 2^{-23} (приблизительно на $2^{-24}\%$) меняет спектр «революционным» образом.

Замечание. Формулировка теоремы 10.4.2 опирается на ограниченный список норм (строчную и столбцовую). Из эквивалентности всех норм в R^n можно сделать вывод, что с той или иной долей натяжки (с добавлением некоторого множителя) норму в неравенстве

$$|\lambda' - \lambda_i| \leq \|T\| \cdot \|T^{-1}\| \cdot \|\Delta A\|$$

можно считать любой. Соответствующий множитель будет не очень большим, если норма не очень вычурна, без больших перекосов¹²⁾. В естественных предположениях такой множитель вообще не требуется (равен единице).

Наличие скрытого механизма, управляющего спектром матрицы при ее возмущении, показывает, что первоочередной интерес может представлять вовсе не фасад задачи. Возмущение спектра совсем редко фигурирует в прикладных постановках, но почти всегда присутствует в подводной части айсберга. Спектральные свойства очень сильно влияют на решение задач, где речь идет вроде бы совсем о других вещах. Диагонализующее преобразование T , как «внутренний орган» A , напрямую редко запрашивается, но именно T , как показывает теорема 10.4.2, отделяет норму от аномалии.

Если возмущение в $2^{-24}\%$ «ломает спектр» (см. выше), то это катастрофическим образом сказывается на решении массы сопутствующих задач. Скудную, но комфортную ситуацию определяет равенство

$$\|T\| \approx \|T^{-1}\| \Rightarrow \sigma(A) \approx 1.$$

В идеальном варианте ортогональной (унитарной) матрицы T все ее собственные значения равны по модулю 1, что в евклидовой норме дает $\sigma(A) = \|T\| \cdot \|T^{-1}\| = 1$. Это одна из главных причин, по которой приятнее всего иметь дело с симметричными матрица-

¹²⁾ «Перекосы» возникают, например, в случае $\|x\| = (\gamma_1 x_1^2 + \dots + \gamma_n x_n^2)^{1/2}$ при большом разбросе значений γ_j .

ми. Возмущение симметричной матрицы приводит к возмущению ее спектра того же порядка.

Упражнения

- Если круги Гершгорина на комплексной плоскости не пересекаются, то все собственные значения матрицы — действительны. (?)
- Попытка улучшить обусловленность матрицы A в задаче $Ax = b$ с помощью умножения $Ax = b$ на A^* не проходит. Обусловленность A^*A в естественных предположениях оказывается не лучше, чем у матрицы A .

10.5. Итерационные методы

К численным методам можно отнести все, что можно использовать для счета. Значительной популярностью пользуются различного рода итерационные процедуры типа

$$x_{k+1} = Ax_k + b. \quad (10.1)$$

Процесс (10.1) при условии $\|A\| < 1$, а значит¹³⁾ и $\rho(A) < 1$, сходится к решению уравнения $x = Ax + b$ (раздел 6.4), но симпатии вычислительной математики к таким процедурам выглядят по меньшей мере странно. Странно — при наличии простых и надежных методов исключения переменных, например.

Простой и убедительный ответ здесь дать трудно. Кто-то будет настаивать, что иначе тем для диссертаций не хватит, — и будет, по-своему, прав, потому что жизнь устроена многопланово. Но существуют также другие причины, менее дискуссионные, роль которых выявляется постепенно в процессе вычислительной практики. Один из главных аргументов, конечно, специфика алгоритмизации, требующая однотипности действий.

Вот пример из другой оперы, где вопрос обострен до крайности. Вычислять значение полинома

$$P(\xi) = a_0\xi^n + a_1\xi^{n-1} + \dots + a_{n-1}\xi + a_n,$$

оказывается, выгоднее на основе представления

$$P(\xi) = (\dots ((a_0\xi + a_1)\xi + a_2)\xi + a_3)\xi + \dots + a_{n-1})\xi + a_n,$$

¹³⁾ Поскольку $\rho(A) < 1$ влечет за собой существование нормы, в которой $\|A\| < 1$ (лемма 6.3.1).

которое приводит к возможности последовательного выполнения однотипных действий: $b_0 = a_0$, и далее $b_{k+1} = a_{k+1} + b_k \xi$ вплоть до

$$b_n = a_n + b_{n-1} \xi = P(\xi).$$

Это так называемая *схема Горнера*. Любой здравомыслящий человек поначалу скажет: «Что в лоб, что по лбу». Второй вариант выглядит даже хуже — из-за противоестественности. Потом лишь выясняется, что в этом есть определенный смысл. Не потому, что схема Горнера приносит дополнительные дивиденды¹⁴⁾, а потому, что итерационный характер действий оказывается краеугольным камнем компьютерных алгоритмов.

Это, конечно, половина правды. При малых размерностях действительно «что в лоб, что по лбу», и всякая попытка агитации только усугубляет непонимание. Роль итерационных методов выявляется сама собой при столкновении с матрицами больших размеров. Там прямые, по идее точные, методы из-за ошибок округления оказываются хуже итерационных даже по точности. Не говоря о том, что всякое преобразование большой матрицы обычно разрушает структуру, порождая хаос. Достаточно представить себе матрицу A «тысяча на тысячу» с элементами

$$a_{jk} = \frac{j}{10^3} \sin \frac{k\pi}{10^3},$$

которую не надо хранить в памяти — ее можно вычислять по мере надобности. Перед тем как преобразовывать такую изюминку — тысячу раз подумаешь¹⁵⁾.

Метод простой итерации. Пусть $\hat{x}$ обозначает решение системы

$$\hat{x} = A\hat{x} + b. \quad (10.2)$$

Вычитая (10.2) из (10.1), получаем процесс

$$\Delta x_{k+1} = A\Delta x_k, \quad (10.3)$$

описывающий эволюцию ошибки $\Delta x_k = x_k - \hat{x}$.

¹⁴⁾ Коэффициенты $b_0, b_1, \dots, b_{n-1}$ определяют полином

$$Q(x) = b_0 x^{n-1} + b_1 x^{n-2} + \dots + b_{n-1}$$

в тождестве $P(x) = Q(x)(x - \xi) + b_n$.

¹⁵⁾ Покажите, что $\rho(A) < 1$.

Динамика вектора невязки

$$\delta_k = Ax_k + b - x_k = x_{k+1} - x_k,$$

очевидно, подчиняется тому же закону (10.3),

$$\delta_{k+1} = A\delta_k.$$

Процессы (10.1), (10.3) сходятся при условии $\rho(A) < 1$. Если критерием остановки счета служит неравенство $\|\delta_k\| \leq \varepsilon$, то ошибка Δx_k , в силу $\Delta x_k = (I - A)^{-1}\delta_k$, может быть значительно больше ε по норме¹⁶⁾. Формулируется это обычно так. Шару $\|\delta_k\| \leq \varepsilon$ ошибок по невязке отвечает эллипсоид ошибок Δx_k с полуосями, равными $\varepsilon/(1 - \alpha_k)$, где α_k — характеристические числа матрицы A (раздел 4.6).

Приведение к удобному виду. Уравнение $Ax = b$ при условии невырожденности A всегда может быть приведено к виду, пригодному для итерационного счета. После умножения $Ax = b$ на $(A^{-1} - \varepsilon I)$ получается эквивалентное уравнение

$$x = \varepsilon Ax + c, \quad c = (A^{-1} - \varepsilon I)b.$$

При этом $\rho(\varepsilon A) < 1$, если модуль ε достаточно мал.

Разумеется, действовать именно так необязательно. Свободой выбора имеет смысл распоряжаться более рационально, преследуя те или иные конкретные цели.

Довольно часто предпочтение отдается процессам с самосопряженными матрицами. Умножая $Ax = b$ на A^* , приходим к уравнению $A^*Ax = A^*b$ с симметричной матрицей.

Процедура Зейделя. Итерация (10.1) реализуется в компьютере последовательно. Сначала вычисляется первая координата x_{k+1} , затем — вторая и т. д. Естественно приходит мысль при вычислении второй координаты использовать вычисленное перед этим новое значение — первой. Затем при вычислении третьей — использовать уже вычисленные значения первой и второй. И так далее. Это и есть *метод Зейделя*.

¹⁶⁾ В том случае, когда у матрицы есть собственные значения, близкие по модулю к единице.

В более обозримой форме он записывается так. Матрица $A = [a_{ij}]$ представляется в виде суммы $A = B + C$ двух матриц

$$B = \begin{bmatrix} 0 & 0 & \dots & 0 & 0 \\ a_{21} & 0 & \dots & 0 & 0 \\ \dots & \dots & \dots & \dots & \dots \\ a_{n1} & a_{n2} & \dots & a_{n(n-1)} & 0 \end{bmatrix} \quad \text{и} \quad C = \begin{bmatrix} a_{11} & a_{12} & \dots & a_{1n} \\ 0 & a_{22} & \dots & a_{2n} \\ \dots & \dots & \dots & \dots \\ 0 & 0 & \dots & a_{nn} \end{bmatrix},$$

что позволяет процесс Зейделя записать в форме

$$x_{k+1} = Bx_{k+1} + Cx_k + b,$$

которая равносильна обычному итерационному процессу

$$x_{k+1} = (I - B)^{-1}Cx_k + (I - B)^{-1}b.$$

Вопрос о сходимости, таким образом, сводится к выяснению вопроса: $\rho[(I - B)^{-1}C] < 1$?

Условия сходимости процессов Зейделя и обычной итерации (10.1) пересекаются, но не совпадают. Любой из них может расходиться, когда другой сходится. Подробности можно найти в любом учебнике по вычислительной математике (см. [17]).

10.6. Вычисление собственных значений

Сначала о том, как действовать нецелесообразно.

Прямые методы определения спектра матрицы опираются на предварительное развертывание определителя

$$\det(A - \lambda I) = \lambda^n + p_1\lambda^{n-1} + \dots + p_n = p(\lambda)$$

с последующим решением характеристического уравнения $p(\lambda) = 0$.

Коэффициенты $p(\lambda)$ определяются, конечно, не в лоб, а с помощью тех или иных уловок. Например, умножение тождества Гамильтона—Кэли

$$A^n + p_1A^{n-1} + \dots + Ip_n = 0$$

на некоторый вектор x_0 дает систему линейных уравнений¹⁷⁾

$$x_n + p_1x_{n-1} + \dots + p_nx_0 = 0$$

относительно коэффициентов $p(\lambda)$. Так $p_1, \dots, p_n$ определяются в *методе Крылова*. При небольших размерностях, о которых только и могла идти речь в старое время, это хорошо работает (иногда). Но в случае «приличных» размерностей это перестает работать совсем, поскольку матрица из вектор-столбцов x_k получается плохо обусловленной, мягко говоря.

Причина заключается в следующем. Пусть, для простоты, A имеет собственное значение $\lambda_1 = 1$, а остальной спектр лежит в круге $|\lambda| \leq \lambda_2 < 1$. Тогда ясно,

¹⁷⁾ $x_k = A^k x_0$.

что итерации $A^k x_0$ при $k \rightarrow \infty$ сходятся по направлению к собственному вектору $x' = Ax'$. В результате столбцы x_k при больших k оказываются почти параллельными.

Хуже того, если минимальный многочлен A имеет степень $< n$, то векторы

$$x_0, \dots, x_n$$

вообще обязаны быть линейно зависимыми.

Поэтому коэффициенты $p(\lambda)$ определяются неточно, если вообще определяются. А присовокупив сюда соображения раздела 10.4 о патологической вычислительной неустойчивости корней многочленов, легко сделать вывод о необходимости искать другие инструменты.

С точки зрения изучения предмета описанная история может показаться холостым выстрелом. Но ошибки и неудачи обладают важным потенциалом, ригиорицивающим идущего в правильном направлении.

Используемые ныне методы вычисления собственных значений имеют итерационный характер. Наиболее популярен *QR-алгоритм*, опирающийся на возможность представления любой невырожденной матрицы A в виде произведения $A = QR$ ортогональной матрицы Q на верхнюю треугольную R (теорема 5.1.1).

Схема итераций *QR*-алгоритма такова.

Определяется *QR*-разложение матрицы $A_k = Q_k R_k$, после чего Q_k и R_k перемножаются в обратном порядке, давая $A_{k+1} = R_k Q_k$. Далее снова определяется ортогонально-треугольное разложение $A_{k+1} = Q_{k+1} R_{k+1}$ и т. д. Процедура начинается с исходной матрицы $A_0 = A$.

Если обозначить

$$\hat{Q}_k = Q_0 Q_k, \quad \hat{R}_k = R_0 \cdots R_k,$$

то $\hat{Q}_k \hat{R}_k = A^k$ и

$$A_{k+1} = \hat{Q}_k^* A \hat{Q}_k.$$

Алгоритм, таким образом, порождает последовательность матриц A_{k+1} , ортогонально подобных исходной A , сходящуюся, при определенных условиях, к верхней треугольной матрице (с собственными значениями A на диагонали).

Глава 11

Сводка основных определений и результатов

11.1. Аналитическая геометрия

✓ Прямоугольные (декартовы) координаты точки на плоскости — суть снабженные знаками *плюс* или *минус* расстояния от точки x до двух взаимно перпендикулярных прямых Ox_1 и Ox_2 — *осей координат*. Декартовы координаты в пространстве задают три взаимно перпендикулярные плоскости, относительно которых положение точки определяется тремя числами, $x = \{x_1, x_2, x_3\}$.

✓ Точку $x = \{x_1, x_2, x_3\}$ называют *вектором* либо *радиусом-вектором*.

✓ Сумма $x = \{x_1, x_2, x_3\}$ и $y = \{y_1, y_2, y_3\}$ определяется по правилу параллелограмма:

$$x + y = \{x_1 + y_1, x_2 + y_2, x_3 + y_3\}.$$

✓ Векторы, лежащие на одной прямой (одинаково или противоположно направленные), *коллинеарны*; в одной плоскости — *компланарны*.

✓ Векторы $\{x_1, \dots, x_k\}$ *линейно зависимы (независимы)*, если существуют (не существуют) такие $\lambda_1, \dots, \lambda_k$, не все равные нулю, что

$$\lambda_1 x_1 + \dots + \lambda_k x_k = 0.$$

✓ Стандартный базис R^3 составляют орты $\{e_1, e_2, e_3\}$ направленные по осям декартовых координат,

$$e_1 = \begin{bmatrix} 1 \\ 0 \\ 0 \end{bmatrix}, \quad e_2 = \begin{bmatrix} 0 \\ 1 \\ 0 \end{bmatrix}, \quad e_3 = \begin{bmatrix} 0 \\ 0 \\ 1 \end{bmatrix}.$$

✓ Проекция a_b вектора a на b определяется формулой $a_b = a \cos \varphi$, где φ — угол между векторами a и b . Проекции a_x, a_y, a_z на декартовы оси x, y, z представляют собой координаты вектора a .

✓ Скалярное произведение равно произведению длин векторов на косинус угла между ними,

$$a \cdot b = |a| \cdot |b| \cos \varphi = a_x b_x + a_y b_y + a_z b_z.$$

✓ **Векторное произведение** $c = a \times b$ определяется как вектор c , ортогональный векторам a и b и составляющий тройку $\{a, b, c\}$, совпадающую по ориентации с тройкой базисных векторов, а длина c равна площади параллелограмма, построенного на векторах a, b ,

$$|c| = |a| \cdot |b| \sin \varphi.$$

✓ Векторное произведение не ассоциативно, $(a \times b) \times c \neq a \times (b \times c)$, антикоммутитивно, $x \times y = -y \times x$, но справедлив дистрибутивный закон,

$$(x + y) \times u = x \times u + y \times u.$$

✓ **Определители.** Объем V параллелепипеда, построенного на векторах a, b, c (как на ребрах), определяется смешанным произведением

$$a \cdot (b \times c) = \pm V,$$

где $b \times c$ дает площадь параллелограмма, лежащего в основании, а проекция $a_{b \times c}$ — высота параллелепипеда. Что касается знака, то плюс получается в ситуации, когда ориентация тройки $\{a, b, c\}$ такая же как у базиса, минус — в противном случае (базис левый, тройка $\{a, b, c\}$ правая; либо наоборот).

✓ Детерминант второго порядка по определению равен

$$\begin{vmatrix} a_x & a_y \\ b_x & b_y \end{vmatrix} = a_x b_y - a_y b_x.$$

✓ Вычисление детерминанта 3-го порядка сводится к вычислению детерминантов 2-го порядка:

$$\begin{vmatrix} a_x & a_y & a_z \\ b_x & b_y & b_z \\ c_x & c_y & c_z \end{vmatrix} = a_x \begin{vmatrix} b_y & b_z \\ c_y & c_z \end{vmatrix} - a_y \begin{vmatrix} b_x & b_z \\ c_x & c_z \end{vmatrix} + a_z \begin{vmatrix} b_x & b_y \\ c_x & c_y \end{vmatrix}.$$

✓ **Прямые и плоскости.** Уравнение прямой в плоскости:

$$\alpha x + \beta y + \delta = 0,$$

плоскости в пространстве:

$$\alpha x + \beta y + \gamma z + \delta = 0.$$

✓ Оба уравнения в векторном виде записываются как $ar + \delta = 0$, что интерпретируется как фиксация проекции текущего вектора $r = \{x, y\}$ либо $r = \{x, y, z\}$ на вектор (направление) a .

✓ Совокупности двух уравнений в пространстве

$$ar = \delta, \quad br = \zeta$$

удовлетворяет прямая (пересечение плоскостей).

✓ Оба описания,

$$\alpha x + \beta y + \gamma z + \delta = 0 \quad \text{и} \quad \mathbf{r} = \lambda \mathbf{a} + \mu \mathbf{b} + \mathbf{c},$$

задают плоскость. Первое описание служит уравнением плоскости с нормалью $\mathbf{a} = \{\alpha, \beta, \gamma\}$ и позволяет легко проверить, принадлежит ли $\mathbf{r} = \{x, y, z\}$ этой плоскости (проверкой равенства $\mathbf{a} \cdot \mathbf{r} + \delta = 0$), второе — параметрически задает плоскость, проходящую через точку $\mathbf{c}$ и ортогональную $\mathbf{a} \times \mathbf{b}$. С его помощью удобно генерировать точки плоскости, выбирая произвольно параметры λ, μ .

✓ Через три заданные точки $\mathbf{r}^1, \mathbf{r}^2, \mathbf{r}^3$ проходит плоскость

$$(\mathbf{a} \times \mathbf{b}) \cdot \mathbf{r} + \delta = 0,$$

где $\mathbf{a} = \mathbf{r}^1 - \mathbf{r}^2$, $\mathbf{b} = \mathbf{r}^1 - \mathbf{r}^3$, величина δ определяется подстановкой вместо $\mathbf{r}$ любой из точек $\mathbf{r}^1, \mathbf{r}^2, \mathbf{r}^3$. В иной форме это можно записать как равенство нулю определителя

$$\begin{vmatrix} x - x_1 & y - y_1 & z - z_1 \\ x_2 - x_1 & y_2 - y_1 & z_2 - z_1 \\ x_3 - x_1 & y_3 - y_1 & z_3 - z_1 \end{vmatrix} = 0.$$

✓ Вместо $\alpha x + \beta y + \gamma z + \delta = 0$ целесообразно использование эквивалента:

$$\alpha(x - x_0) + \beta(y - y_0) + \gamma(z - z_0) = 0,$$

где вместо не всегда удобного параметра δ фигурирует точка $\{x_0, y_0, z_0\}$, через которую проходит плоскость.

✓ Аналогично,

$$\alpha(x - x_0) + \beta(y - y_0) = 0$$

есть уравнение прямой в плоскости, ортогональной вектору $\{\alpha, \beta\}$ и проходящей через точку $\{x_0, y_0\}$.

✓ Каноническим уравнением прямой удобно считать

$$\mathbf{r} \times \mathbf{c} = \mathbf{d}.$$

Вектор $\mathbf{c}$ здесь направлен вдоль прямой, а $\mathbf{d} = \mathbf{r}^0 \times \mathbf{c}$, где $\mathbf{r}^0$ — любая точка, через которую проходит прямая.

✓ Другой вариант описания прямой, проходящей через точку $\mathbf{r}^0$, — параметрическое задание, $\mathbf{r} - \mathbf{r}^0 = \alpha \mathbf{t}$, где вектор $\mathbf{a}$ задает направление прямой.

При необходимости описать прямую, проходящую через две точки, достаточно в качестве $\mathbf{a}$ взять вектор $\mathbf{r}^1 - \mathbf{r}^2$. После исключения параметра τ получается

уравнение

$$\frac{x - x_0}{x_2 - x_1} = \frac{y - y_0}{y_2 - y_1} = \frac{z - z_0}{z_2 - z_1},$$

где точку r^0 полагают обычно равной либо r^1 , либо r^2 .

✓ Касательная плоскость, проходящая через точку $\{x_0, y_0\}$, к линии постоянного уровня функции $z = f(x, y)$ описывается уравнением

$$\frac{\partial f}{\partial x}(x - x_0) + \frac{\partial f}{\partial y}(y - y_0) = 0.$$

✓ Касательная плоскость к поверхности графика функции $z = f(x, y)$ в точке $\{z_0, x_0, y_0\}$ определяется уравнением

$$z - z_0 = \frac{\partial f}{\partial x}(x - x_0) + \frac{\partial f}{\partial y}(y - y_0).$$

✓ Расстояние от начала координат до плоскости $ra = \delta$ равно

$$\|r\|_{\min} = |r_a| = \frac{|\delta|}{\|a\|}.$$

✓ Расстояние до прямой $r \times a = c$. Раскладываем r по взаимно перпендикулярным направлениям $r = r_{\parallel a} + r_{\perp a}$. В ближайшей к началу координат точке имеем $|r_{\perp a}| \cdot \|a\| = \|c\|$, откуда искомое расстояние:

$$\frac{\|c\|}{\|a\|}.$$

✓ Минимальное расстояние между скрещивающимися прямыми,

$$r \times a = A, \quad r \times b = B,$$

равно

$$\frac{\|A \cdot b + B \cdot a\|}{\|a \times b\|}.$$

✓ Эллипс. Каноническое уравнение эллипса:

$$\frac{x^2}{a^2} + \frac{y^2}{b^2} = 1, \quad a \geq b > 0.$$

Эллипс получается в пересечении кругового конуса с плоскостью и может быть охарактеризован как множество точек плоскости, для каждой из которых

сумма расстояний до фокусов F_1, F_2 постоянна и равна $2a$. Расстояние между фокусами — $2c$, где $c^2 = a^2 - b^2$.

Оптическое свойство: световой луч, исходящий из одного фокуса, после отражения от эллипса проходит через другой фокус.

✓ Уравнение эллипсоида:

$$\frac{x^2}{a^2} + \frac{y^2}{b^2} + \frac{z^2}{c^2} = 1.$$

✓ Гипербола. Гипербола — плоская кривая,

$$\frac{x^2}{a^2} - \frac{y^2}{b^2} = 1,$$

как и эллипс, представляет собой сечение конуса плоскостью.

Характеризуется как геометрическое множество точек плоскости, модуль разности расстояний которых до F_1 и F_2 постоянен и равен $2a$.

При $a = b$ асимптоты взаимно перпендикулярны. Если их принять за координатные оси, — уравнение гиперболы трансформируется в $y = k/x$.

✓ Пространственное обобщение — *гиперboloид*. Уравнение:

$$\frac{x^2}{a^2} + \frac{y^2}{b^2} - \frac{z^2}{c^2} = 1.$$

11.2. Векторы и матрицы

✓ Операции над векторами $x = \{x_1, \dots, x_n\}$:

- умножение на скаляр: $\lambda x = \{\lambda x_1, \dots, \lambda x_n\}$,
- сложение: $x + y = \{x_1 + y_1, \dots, x_n + y_n\}$,
- скалярное произведение: $x \cdot y = x_1 y_1 + \dots + x_n y_n$. Эквивалентные обозначения (x, y) , xy .

✓ Евклидова норма:

$$\|x\| = \sqrt{x^2} = \sqrt{x_1^2 + \dots + x_n^2}.$$

✓ Множество векторов, на котором введены перечисленные операции, называют n -мерным *евклидовым* пространством и обозначают R^n .

✓ Неравенство Коши—Буняковского (Коши—Шварца)

$$x \cdot y \leq \|x\| \cdot \|y\|.$$

✓ Векторы $\{x_1, \dots, x_k\}$ *линейно зависимы (независимы)*, если существуют (не существуют) такие $\lambda_1, \dots, \lambda_k$, не все равные нулю, что

$$\lambda_1 x_1 + \dots + \lambda_k x_k = 0.$$

✓ Вектор $u = \lambda_1 x_1 + \dots + \lambda_k x_k$ называют *линейной комбинацией* векторов $\{x_1, \dots, x_k\}$, а совокупность всевозможных u — *линейной оболочкой* множества $\{x_1, \dots, x_k\}$.

✓ **Базисом** называют линейно независимое множество $\{e_1, \dots, e_n\}$ при условии, что любой вектор $x \in R^n$ представим в виде линейной комбинации $x = x_1 e_1 + \dots + x_n e_n$. Величины $x_1, \dots, x_n$ — координаты точки $x \in R^n$. Число n векторов, составляющих базис, — *размерность* пространства.

✓ **Ортогональность.** Векторы x, y *ортогональны*, если $x \cdot y = 0$. Плоскость (линейное подпространство) определяется как множество элементов x , ортогональных некоторому вектору a , т. е. $a \cdot x = 0$.

✓ **Линейная функция** $y = a \cdot x = a_1 x_1 + \dots + a_n x_n$ принимает постоянные значения $a \cdot x = \beta$ на «плоскостях», параллельных плоскости $a \cdot x = 0$. Множества $a \cdot x = \beta$ называют *гиперплоскостями*.

✓ Таблицу коэффициентов $A = [a_{ij}]$,

$$A = \begin{bmatrix} a_{11} & \dots & a_{1n} \\ \dots & \dots & \dots \\ a_{n1} & \dots & a_{nn} \end{bmatrix},$$

называют *матрицей*, которая определяет *линейный оператор* $y = Ax$, сопоставляющий вектору x вектор y по правилу

$$y_i = \sum_j a_{ij} x_j, \quad i = 1, \dots, n,$$

что задает тем самым умножение матрицы на вектор.

✓ Совокупность элементов $a_{11}, \dots, a_{nn}$ называют *главной диагональю* матрицы.

✓ Через a_i обозначается i -я строка матрицы A , через $a_{\cdot j}$ — j -й столбец, т. е.

$$a_i = [a_{i1}, \dots, a_{in}], \quad a_{\cdot j} = \begin{bmatrix} a_{1j} \\ \vdots \\ a_{nj} \end{bmatrix}.$$

✓ Матрица называется *невыврожденной*, если ее вектор-строки линейно независимы.

✓ По определению: $\gamma A = [\gamma a_{ij}]$, $A + B = [a_{ij} + b_{ij}]$. Умножение A и B дает матрицу $C = AB$ с элементами $c_{ij} = \sum_k a_{ik} b_{kj}$. Иными словами, c_{ij} равно скалярному произведению i -й вектор-строки на j -й вектор-столбец, т. е. $c_{ij} = a_{i \cdot} \cdot b_{\cdot j}$.

✓ Произведение матриц в общем случае не коммутативно $AB \neq BA$, но ассоциативно, $A(BC) = (AB)C$, и дистрибутивно, $A(B + C) = AB + AC$. В случае $AB = BA$ говорят, что матрицы A и B коммутируют.

✓ Матрица A^* с элементами $a_{ij}^* = a_{ji}$ называется *транспонированной*. Симметричная матрица характеризуется равенством $A = A^*$.

$$(AB)^* = B^* A^*.$$

✓ Обратной к A называют матрицу A^{-1} такую, что

$$A^{-1}A = AA^{-1} = I,$$

где *единичная матрица* I определяет тождественный оператор $Ix \equiv x$.

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

$$(A^*)^{-1} = (A^{-1})^*.$$

✓ Если матрица A невырождена, то:

1. $\text{rank} A = n$.
2. Существует A^{-1} .
3. $\det A \neq 0$.
4. Система уравнений $Ax = b$ имеет единственное решение при любом b .

✓ **Прямоугольные и клеточные матрицы.** Матрицу A , имеющую m строк и n столбцов, называют $m \times n$ матрицей. Произведение $C = AB$ вычисляется по обычной формуле, но при условии, что *число столбцов A равно числу строк B* .

✓ Понятие прямоугольной матрицы позволяет считать векторы тоже матрицами. Чтобы матрицу A умножить на x *справа*, надо, чтобы элемент x был обязательно вектор-столбцом. После транспонирования вектор-столбца возможно x^*A . Скалярное произведение (x, y) записывают тогда в виде x^*y .

✓ Пример:

$$[x_1, \dots, x_n] \cdot \begin{bmatrix} y_1 \\ \vdots \\ y_n \end{bmatrix} = \sum_i x_i y_i, \quad \begin{bmatrix} x_1 \\ \vdots \\ x_n \end{bmatrix} \cdot [y_1, \dots, y_n] = \begin{bmatrix} x_1 y_1 & \dots & x_1 y_n \\ \dots & \dots & \dots \\ x_n y_1 & \dots & x_n y_n \end{bmatrix}.$$

✓ **Элементарные преобразования** матрицы представляют собой последовательность операций двух типов:

*умножение строки матрицы на ненулевой скаляр,
прибавление одной строки к другой.*

✓ Любая невырожденная матрица элементарным преобразованием может быть преобразована в единичную. Соответствующий процесс называют *методом Гаусса*.

✓ При элементарных преобразованиях линейно зависимые (независимые) строки и столбцы переходят в линейно зависимые (независимые).

✓ У вырожденной (невырожденной) матрицы вектор-столбцы линейно зависимы (независимы), т. е. A и A^* вырождены или невырождены одновременно.

✓ Максимальное число линейно независимых вектор-строк либо вектор-столбцов (то и другое совпадает) прямоугольной матрицы A называется ее **рангом** и обозначается как $\text{rank } A$. Эквивалент: ранг $m \times n$ матрицы — это максимальный размер ее квадратной *невырожденной* подматрицы.

$$\text{rank } AB \leq \min \{ \text{rank } A, \text{rank } B \}.$$

✓ Если P невырожденная матрица, то

$$\text{rank } PQ = \text{rank } Q, \quad \text{rank } SP = \text{rank } S.$$

✓ **Определители.** *Детерминантом, или определителем*, квадратной матрицы A называется скалярная функция $d(A)$, обладающая тремя свойствами:

- (i) *при умножении строки матрицы на скаляр λ величина $d(A)$ умножается на λ ;*
- (ii) *прибавление одной строки к другой не меняет детерминанта;*
- (iii) *детерминант единичной матрицы $d(I) = 1$.*

✓ Альтернативное определение:

$$d(A) = \det A = |A| = \sum_{\sigma} (-1)^{N(\sigma)} a_{\sigma(1)1} \cdots a_{\sigma(n)n},$$

где σ обозначает *перестановку* индексов $\{1, \dots, n\}$ в $\{\sigma(1), \dots, \sigma(n)\}$, а $N(\sigma)$ — число *транспозиций* (перестановок двух элементов) в σ . Суммирование идет по всем возможным перестановкам.

✓ Детерминант представляет собой алгебраическую сумму всевозможных произведений $a_{\sigma(1)1} \cdots a_{\sigma(n)n}$, знак которых определяется четностью или нечетностью перестановки σ (четностью или нечетностью числа транспозиций в σ).

11.3. Линейные преобразования

✓ При переходе к новой системе координат с помощью невырожденной матрицы, $x = Tx'$, матрица A преобразуется по правилу:

$$A' = T^{-1}AT.$$

✓ Аналогично преобразуется любой полином от матрицы

$$p(T^{-1}AT) = T^{-1}p(A)T.$$

✓ Операции взятия обратной матрицы и перехода к новой системе координат — перестановочны: $(A')^{-1} = (A^{-1})'$.

✓ У диагональной матрицы

$$\Lambda = \begin{bmatrix} \lambda_1 & & 0 \\ & \ddots & \\ 0 & & \lambda_n \end{bmatrix} = \text{diag} \{ \lambda_1, \dots, \lambda_n \}$$

ненулевые элементы могут стоять только на главной диагонали.

✓ Матрица A диагонализуема, если существует такая невырожденная матрица T , что $T^{-1}AT = \Lambda$.

✓ Матрицы A' и A в соотношении $A' = T^{-1}AT$ называют *подобными*. Это различные матрицы, но они описывают один и тот же оператор (в разных системах координат).

✓ Уравнение $Ax = \lambda x$ может иметь ненулевое решение x лишь в случае, когда матрица $A - \lambda I$ вырождена, т. е.

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

что называют *характеристическим уравнением* матрицы A , его решения λ_j — *собственными значениями*, а соответствующие ненулевые решения $x = f_j$ уравнения $(A - \lambda I)x = 0$ — *собственными векторами* матрицы A .

✓ Полином

$$p_A(\lambda) = \det(A - \lambda I) = \lambda^n + \gamma_{n-1}\lambda^{n-1} + \dots + \gamma_1\lambda + \gamma_0$$

называют *характеристическим*.

✓ Если матрица A имеет собственные значения λ_j , которым отвечают собственные векторы f_j , то матричный полином

$$p(A) = A^n + \gamma_{n-1}A^{n-1} + \dots + \gamma_0I$$

имеет собственные значения $p(\lambda_j)$ и те же собственные векторы f_j .

✓ Собственными значениями обратной матрицы A^{-1} являются λ_j^{-1} и те же собственные векторы.

✓ В фокус внимания часто попадают два коэффициента: γ_{n-1} , равный сумме корней (собственных значений) $\lambda_1(A) + \dots + \lambda_n(A)$, и γ_0 , равный произведению $\lambda_1(A) \dots \lambda_n(A)$. По теореме Виета

$$\alpha_0 = \lambda_1(A) \dots \lambda_n(A) = \det A,$$

$$\alpha_{n-1} = \lambda_1(A) + \dots + \lambda_n(A) = \operatorname{tr} A,$$

где $\operatorname{tr} A = a_{11} + \dots + a_{nn}$ — след матрицы (оператора)

✓ Под скалярным произведением в случае комплексных векторов понимается

$$\langle x, y \rangle = \sum_j x_j y_j^*,$$

что при действительных координатах переходит в стандартное определение.

✓ Аналогом «транспонирования» в комплексном случае служит операция сопряжения матрицы: V^* получается из V транспонированием с последующей заменой всех элементов на комплексно сопряженные — и называется сопряженной V .

✓ В билинейных формах $\langle Vx, y \rangle$ переброска матрицы на другой сомножитель происходит по правилу:

$$\langle Vx, y \rangle = \langle x, V^*y \rangle.$$

✓ Формулы $(AB)^* = B^*A^*$ и $(A^*)^{-1} = (A^{-1})^*$ сохраняются в комплексном случае.

Если матрица A имеет попарно различные собственные значения $\lambda_1, \dots, \lambda_n$, то каждому λ_j отвечает свой собственный вектор f_j , и множество собственных векторов $\{f_1, \dots, f_n\}$ линейно независимо.

Пусть $Ax = \lambda x$, $yA = \mu y$ и $\lambda \neq \mu$. Тогда $\langle x, y \rangle = 0$. Другими словами, левый и правый собственные векторы, отвечающие разным собственным значениям, всегда ортогональны друг другу.

✓ Множество $\sigma(A)$ собственных значений λ_j матрицы A называют *спектром*, а максимальное значение модуля $|\lambda_j|$ — *спектральным радиусом*, и обозначают $\rho(A)$.

✓ Спектр в значительной мере характеризует свойства матрицы. Например, из $\operatorname{Re} \lambda_j < 0$ для всех j — следует асимптотическая устойчивость равновесия системы $\dot{x} = Ax + b$. Из $\rho(A) < 1$ вытекает сходимость последовательных приближений $x^{k+1} = Ax^k + b$ к решению $x = Ax + b$.

✓ Спектр не меняется при переходе к другому базису, т. е. характеризует сам линейный оператор, а не его конкретную запись в той или иной системе координат.

✓ Если у $n \times n$ матриц A и B спектры состоят из n попарно различных точек и совпадают, то матрицы *подобны* — переходят одна в другую при замене системы координат, $B = T^{-1}AT$, — и представляют собой запись одного и того же оператора в различных базисах.

✓ В общем случае это не так. Мешают кратные собственные значения. Кратность λ как корня характеристического многочлена $p_A(\lambda)$ называется *алгебраической кратностью* собственного значения λ . Число линейно независимых решений x уравнения $Ax = \lambda x$ называют *геометрической кратностью* собственного значения λ .

✓ Если бы геометрическая кратность всегда совпадала с алгебраической, — у любой матрицы A было бы n линейно независимых собственных векторов $\{f_1, \dots, f_n\}$ и преобразование $T = [f_1, \dots, f_n]$ приводило бы A к диагональному виду.

✓ Матрица, у которой геометрическая кратность какого-либо собственного значения меньше алгебраической, называется *дефектной*. Таковой является, например, $\begin{bmatrix} 0 & 1 \\ 0 & 0 \end{bmatrix}$. Собственное значение $\lambda = 0$ здесь имеет кратность 2, но ему соответствует лишь один собственный вектор $\begin{bmatrix} 1 \\ 0 \end{bmatrix}$.

✓ Множество значений Ax называют *образом* и обозначают как $\operatorname{im} A$, а множество решений $Ax = 0$ называют *ядром* и пишут $\ker A$, имея в виду множество векторов x , которые оператором отображаются в нуль. И то и другое является линейным пространством.

✓ Подпространство $X \subset R^n$ называется *инвариантным* относительно A , если $x \in X \Rightarrow Ax \in X$, т. е. $AX \subset X$.

✓ Оператор A , рассматриваемый на инвариантном подпространстве X , т. е. *сужение* $A|_X$ оператора A на X , является, очевидно, линейным оператором.

✓ При работе с подпространствами возникает необходимость рассматривать их комбинации. Суммой $X + Y$ подпространств $X, Y \subset R^n$ называется множество линейных комбинаций $\alpha x + \beta y$ элементов $x \in X$ и $y \in Y$. Пересечением подпространств называют $X \cap Y$, где подразумевается обычное теоретическое пересечение, т. е. $u \in X \cap Y$, когда и $u \in X$, и $u \in Y$. Наконец, $X^\perp$ определяют как ортогональное дополнение, при условии что $X + X^\perp \neq R^n$ и

$$x \in X, \quad x^\perp \in X^\perp \Rightarrow (x, x^\perp) = 0.$$

✓ Если X плоскость, а Y прямая, не лежащая в этой плоскости, $X + Y$ даст трехмерное пространство, а $X \cap Y$ — точку. Если же $Y \subset X$ (прямая лежит в плоскости), то $X + Y = X$, $X \cap Y = Y$. В случае

$$\dim(X + Y) = \dim X + \dim Y$$

сумму $X + Y$ называют прямой.

11.4. Квадратичные формы

✓ Квадратичная форма

$$(Vx, x) = \sum_{i,j} v_{ij} x_i x_j$$

задается $n \times n$ матрицей V , которая обычно предполагается симметричной.

✓ При переходе к другой (штрихованной) системе координат с помощью невырожденной матрицы, $x = Sx'$, матрица V преобразуется по правилу

$$V' = S^* V S.$$

✓ Преобразование S называют ортогональным, если оно не меняет длину векторов, что влечет за собой

$$S^* = S^{-1}.$$

✓ Множество векторов $\{f_1, \dots, f_n\}$ ортогонально, если $(f_i, f_j) = 0$ при $i \neq j$. Если, дополнительно, все векторы нормированы, $(f_i, f_i) = 1$, то их множество считается ортонормированным.

✓ Все собственные значения вещественной симметричной матрицы вещественны, и им соответствуют вещественные собственные векторы.

Теорема. Вещественная симметричная $n \times n$ матрица A всегда имеет n различных собственных векторов $\{f_1, \dots, f_n\}$, которые образуют ортонормированное множество, а ортогональное преобразование $T = [f_1, \dots, f_n]$ приводит A к диагональному виду $T^* A T$.

✓ (x, Vx) всегда может быть приведена *ортогональной заменой переменных* $x = Tu$ к виду

$$(x, Vx) = \sum_j \lambda_j u_j^2$$

где λ_j — собственные значения V (при условии $V = V^*$), и — к виду

$$(x, Vx) = \sum_{j=1}^p z_j^2 - \sum_{j=p+1}^q z_j^2,$$

если требование ортогональности на T не накладывается.

✓ Число $q \leq n$ называют *рангом квадратичной формы*, p и $q-p$ — *индексами инерции* (положительным и отрицательным), а разность индексов — *сигнатурой*, которая не зависит от способа преобразований, ведущего к указанной форме (закон инерции).

✓ Матрицу V , равно как и квадратичную форму (x, Vx) , называют *положительно определенной*, если $(x, Vx) > 0$ при любом ненулевом x .

✓ Если одна из матриц A, B *положительно определена*, то обе формы (x, Ax) и (x, Bx) приводятся к диагональной форме *одновременно* (одним и тем же преобразованием).

✓ Если упорядочить собственные значения матрицы V в порядке убывания $\lambda_1 \geq \dots \geq \lambda_n$, то

$$\lambda_1 = \max_x \frac{(x, Vx)}{(x, x)}, \quad \lambda_n = \min_x \frac{(x, Vx)}{(x, x)}.$$

✓ *Эрмитова матрица*, или *самосопряженная* (комплексный вариант симметричной), определяется условием $A = A^*$, где звездочка обозначает уже транспонирование *плюс* комплексное сопряжение всех элементов.

✓ Если матрица A эрмитова, то:

- Все собственные значения A действительны.
- Квадратичная форма (x, Ax) принимает действительные значения при любых комплексных x .
- Матрица S^*AS тоже эрмитова.
- Существует представление $A = U^* \Lambda U$, где U — унитарная матрица, а Λ — вещественная диагональная. Обратное тоже верно.

✓ **Сингулярные числа.** Ортонормированная система собственных векторов $\{f_1, \dots, f_n\}$ преобразования A^*A переводится матрицей A в *ортогональную*:

$$\{Af_1, \dots, Af_n\}.$$

Величины $\alpha_i = \sqrt{\lambda_i}$, где λ_i — собственные значения A^*A , называют *сингулярными числами матрицы A* . Максимальное α_i представляет собой евклидову норму A .

✓ Любая матрица A может быть разложена в произведение

$$A = TS$$

ортогональной матрицы T и симметричной неотрицательно определенной матрицы S .

✓ **Биортогональные базисы.** У матрицы A общего вида, имеющей попарно различные собственные значения $\{\lambda_1, \dots, \lambda_n\}$, существует базис из собственных векторов $\{e_1, \dots, e_n\}$, не обязательно ортогональный.

Транспонированная матрица A^* имеет те же собственные значения, но другие собственные векторы $\{f_1, \dots, f_n\}$.

Базисы $\{e_i\}$ и $\{f_i\}$ *биортогональны*:

$$(e_i, f_j) = 0 \quad \text{при} \quad i \neq j.$$

✓ С помощью биортогональных базисов $\{e_i\}$ и $\{f_i\}$ — при условии нормировки $(e_i, f_i) = 1$ — матрица A с попарно различными собственными значениями $\{\lambda_1, \dots, \lambda_n\}$ может быть представлена как

$$A = \lambda_1 e_1 f_1^* + \dots + \lambda_n e_n f_n^*.$$

11.5. Канонические представления

✓ **Унитарные матрицы** — это комплексный вариант ортогональных. Матрица U *унитарна*, если $U^*U = I$. Форма определения такая же, но теперь U^* — не просто транспонированная, а сопряженная.

✓ Источник интереса к унитарным матрицам — тот же самый, что и к ортогональным. Те и другие возникают обычно как инструмент исследования. Ортогональные матрицы, как матрицы преобразований, удобны благодаря ортогональности столбцов. В общем случае ортогональных преобразований уже недостаточно, и унитарные матрицы — следующий шаг. Выгоды, по сути, те же:

- $(u_i, u_j) = 0$ при $i \neq j$, и $u_i^2 = 1$ (аналогично для строк u_i);
- евклидова длина преобразованного вектора $x = Uy$ ($y = U^*x$) сохраняется, $\langle x, x \rangle = \langle y, y \rangle$;
- $U^* = U^{-1}$;
- все собственные значения по модулю равны единице, $\lambda_k = e^{i\varphi_k}$.

Теорема об унитарной триангуляции. Любая матрица A может быть приведена унитарным преобразованием к треугольному виду:

$$U^*AU = T,$$

где T — верхняя треугольная матрица. Если матрица A и все ее собственные значения вещественны, то U может быть выбрана вещественной.

✓ Любая действительная невырожденная матрица A может быть представлена в виде произведения $A = QR$ ортогональной матрицы Q на верхнюю треугольную R .

✓ **Жордановы формы.** Матрица A заменой базиса, $A = S^{-1}JS$, может быть всегда приведена к жордановой форме:

$$J = \begin{bmatrix} J_{k_1}(\lambda_1) & & 0 \\ & \ddots & \\ 0 & & J_{k_p}(\lambda_p) \end{bmatrix},$$

где на диагонали стоят жордановы клетки $J_k(\lambda)$,

$$J_k(\lambda) = \begin{bmatrix} \lambda & 1 & & 0 \\ & \ddots & \ddots & \\ & & \ddots & 1 \\ 0 & & & \lambda \end{bmatrix}.$$

✓ Конкретный вид J для A связан со следующими факторами:

- Число жордановых клеток J равно максимальному числу линейно независимых собственных векторов A .
- Сумма порядков жордановых клеток, отвечающих одному и тому же собственному значению λ_j , равна алгебраической кратности λ_j , а число этих клеток равно геометрической кратности λ_j .

11.6. Функции от матриц

✓ Матричный ряд

$$f(A) = \sum_{k=0}^{\infty} \gamma_k A^k, \quad A^0 = I,$$

сходится, если $\rho(A) < r$, где r — радиус сходимости ряда $\sum \gamma_k x^k$.

$$\checkmark \quad (I - A)^{-1} = I + A + A^2 + \dots, \quad \rho(A) < 1.$$

$$\checkmark \quad e^A = I + A + \frac{A^2}{2!} + \frac{A^3}{3!} + \dots \quad (\text{матричная экспонента}).$$

$$\checkmark \quad AB = BA \Rightarrow \boxed{e^{A+B} = e^A e^B}, \text{ но } \boxed{e^{A+B} \neq e^A e^B}, \text{ если } AB \neq BA.$$

$$\checkmark \quad \det e^A = e^{\operatorname{tr} A}, \quad e^{iA} = \cos A + i \sin A.$$

✓ Если матрица A диагонализируема преобразованием T , то

$$e^A = T^{-1} \begin{bmatrix} e^{\lambda_1} & \dots & 0 \\ \vdots & \ddots & \vdots \\ 0 & \dots & e^{\lambda_n} \end{bmatrix} T.$$

По любому $\varepsilon > 0$ можно указать норму $\|\cdot\|_*$, в которой

$$\|A\|_* < \rho(A) + \varepsilon.$$

$$\checkmark \quad \text{Какова бы ни была норма, } \rho(A) = \lim_{k \rightarrow \infty} \sqrt[k]{\|A^k\|}.$$

✓ **Нормы.** Помимо евклидовой $\|x\| = \sqrt{(x, x)}$ широкое распространение имеют также

$$\|x\|_m = \max_i |x_i| \quad (\text{кубическая норма})$$

и

$$\|x\|_l = \sum_{i=1}^n |x_i| \quad (\text{октаэдрическая норма}).$$

$$\checkmark \quad \text{Двойственная норма: } \|x\|^* = \max_{\|y\|=1} \{x, y\}.$$

✓ В R^n все нормы эквивалентны, т. е. для любых двух норм $\|\cdot\|_1, \|\cdot\|_2$ можно указать такие константы α и β , что

$$\|x\|_1 \leq \alpha \|x\|_2, \quad \|x\|_2 \leq \beta \|x\|_1 \quad (x \in R^n).$$

✓ Норма матрицы определяется как

$$\|A\| = \max_{\|x\|=1} \|Ax\|.$$

✓ Нормам $\|\cdot\|_m$ и $\|\cdot\|_l$ соответствуют:

$$\|A\|_m = \max_i \sum_{j=1}^n |a_{ij}| \quad (\text{строчная норма}),$$

и

$$\|A\|_l = \max_j \sum_{i=1}^n |a_{ij}| \quad (\text{столбцовая норма}).$$

11.7. Неравенства

✓ Теоремы об альтернативах.

- Либо неравенство $Ax > 0$ разрешимо, либо $yA = 0$ имеет ненулевое положительное решение.
- Либо неравенство $Ax \geq 0$ разрешимо ($x \neq 0$), либо $yA = 0$ имеет строго положительное решение $y > 0$.
- Система линейных уравнений и неравенств,

$$Ax = 0, \quad x \geq 0, \quad yA \geq 0,$$

всегда имеет такое решение $\{x, y\}$, что $yA + x > 0$.

- Либо $Ax = b$ имеет решение $x \geq 0$, либо $yA \geq 0$ влечет за собой $yb \geq 0$ (лемма Минковского—Фаркаша).
- Либо $Ax = b$ имеет решение $x \geq 0$, либо разрешима система неравенств $yA \geq 0$, $yb < 0$ (переформулировка предыдущего утверждения).
- Либо $Ax \leq b$ имеет решение $x \geq 0$, либо система неравенств $yA \geq 0$, $yb < 0$ имеет решение $y \geq 0$.
- Либо $Ax \geq b$ разрешимо, либо система $yA = 0$, $yb = 1$ имеет решение $y \geq 0$.
- Система неравенств

$$Ax \leq 0, \quad x \geq 0, \quad yA \geq 0, \quad y \geq 0$$

всегда имеет такое решение, что

$$x + yA > 0, \quad y - Ax > 0.$$

✓ **P-матрицы.** Квадратная матрица A называется *P-матрицей*, если она не обращает знак ни одного ненулевого вектора, т. е. для любого $x \neq 0$ можно указать такой индекс j , что $x_j \cdot y_j > 0$, где $y = Ax$.

✓ Если A является *P-матрицей*, то неравенство $Ax > 0$ всегда имеет строго положительное решение $x > 0$.

✓ Матрица является *P-матрицей* в том и только том случае, когда все ее главные миноры строго положительны.

11.8. Положительные матрицы

✓ Матрица A с неотрицательными элементами, $a_{ij} \geq 0$, называется *положительной*, с положительными, $a_{ij} > 0$, — *строго положительной*.

✓ Пусть $A \geq 0$, $x_0 > 0$ и

$$\alpha x_0 \leq Ax_0 \leq \beta x_0.$$

Тогда $\alpha \leq \rho(A) \leq \beta$.

✓ Пусть $A \geq 0$ и $-A \leq B \leq A$. Тогда $\rho(B) \leq \rho(A)$.

✓ Если B главная подматрица матрицы $A \geq 0$, то

$$\rho(B) \leq \rho(A).$$

Если $A > 0$, то $\rho(B) < \rho(A)$.

✓ **Теорема Перрона.** Пусть $A > 0$. Тогда:

- Спектральный радиус $\rho(A) > 0$ является собственным значением матрицы A алгебраической кратности единица, и ему отвечает строго положительный собственный вектор.
- Других положительных собственных значений и векторов матрица A не имеет.
- Если $Ax = \lambda x$, $\lambda \neq \rho(A)$, $x \neq 0$, то $|\lambda| < \rho(A)$.

✓ Матрица A называется *разложимой* (неразложимой), если одинаковой перестановкой строк и столбцов она приводится (не может быть приведена) к виду

$$\begin{bmatrix} A_{11} & A_{12} \\ 0 & A_{22} \end{bmatrix},$$

где A_{11} и A_{22} — квадратные матрицы.

✓ Если матрица $A \geq 0$ неразложима, то $\rho(A) > 0$ является ведущим собственным значением A алгебраической кратности 1, которому отвечает строго положительный собственный вектор. Других положительных собственных значений и векторов у A нет.

✓ Внидиагонально отрицательная матрица B с положительной диагональю положительно обратима, если

$$\max_i \frac{1}{\gamma_i} \sum_{j \neq i} |b_{ij}| \gamma_j < |b_{ii}|$$

при некотором наборе констант $\gamma_1, \dots, \gamma_n > 0$.

✓ Пусть $C - B \geq 0$, матрица C положительно обратима и $Bu > 0$ для некоторого $u \geq 0$. Тогда $B^{-1} \geq 0$.

✓ Пусть $C \geq B \geq D$, причем матрицы C и D положительно обратимы. Тогда $B^{-1} \geq 0$.

✓ Неразложимая матрица $A \geq 0$ может иметь несколько *периферических* собственных значений $\{\lambda_1, \dots, \lambda_k\}$, равных по модулю $\rho(A)$. Такие матрицы называют *импримитивными*. Оказывается, что все периферические λ_j обязаны иметь вид $\omega_j \rho(A)$, где ω_j — комплексные корни k -й степени из 1.

✓ Положительная матрица $A \geq 0$, у которой все столбцовые суммы

$$\sum_i a_{ij} = 1,$$

называется *стохастической*.

✓ Пусть $A > 0$ и

$$Ax = \rho(A)x, \quad A^*y = \rho(A)y,$$

$$x > 0, \quad y > 0, \quad x^*y = 1.$$

Тогда при $k \rightarrow \infty$

$$[\rho^{-1}(A)A]^k \rightarrow A_\infty = xy^*.$$

Обозначения

◀ и ▶ — начало и конец рассуждения, темы, доказательства

(?) — предлагает проверить или доказать утверждение в качестве упражнения, либо довести рассуждение до «логической точки»

$A \Rightarrow B$ — из A следует B

$x \in X$ — x принадлежит X

$X \cup Y$, $X \cap Y$, $X \setminus Y$ — объединение, пересечение и разность множеств X и Y

$X \subset Y$ — X подмножество Y , в том числе имеется в виду возможность $X \subseteq Y$, т. е. между $X \subset Y$ и $X \subseteq Y$ различия не делается

$\emptyset$ — пустое множество

$\text{int } \Omega$ — внутренность Ω

R^2 — плоскость, R^3 — трехмерное, R^n — n -мерное пространство

$\dim X$ — размерность линейного пространства X

$\delta_{ij} = \delta_i^j = \begin{cases} 1, & \text{если } i = j, \\ 0, & \text{если } i \neq j, \end{cases}$ — символ Кронекера

i — мнимая единица, $i^2 = -1$

$z = x + iy$ — комплексное число, $z = r(\cos \varphi + i \sin \varphi)$ — его тригонометрическая запись, $x = \text{Re } z$ — действительная часть, $y = \text{Im } z$ — мнимая; $\bar{z} = z^* = x - iy$ — комплексно сопряженное число.

$\mathbf{x} = \begin{bmatrix} x_1 \\ \vdots \\ x_n \end{bmatrix}$ — вектор, x_i — его координаты

$$\bar{x} = \begin{bmatrix} \bar{x}_1 \\ \vdots \\ \bar{x}_n \end{bmatrix}$$

$$x^* = [\bar{x}_1, \dots, \bar{x}_n]$$

A^* — обозначает транспонированную (сопряженную) матрицу, если A имеет действительные (комплексные) элементы

$\ker A$ — ядро линейного оператора A

(x, y) либо $\langle x, y \rangle$ — скалярное произведение векторов x и y ; в общем случае комплексных векторов

$$\langle x, y \rangle = x_1 y_1^* + \dots + x_n y_n^*.$$

Для скалярного произведения используются также эквивалентные обозначения $x \cdot y$ и xy

$|A| = \det A$ — определитель (детерминант), в главе 9 $|A|$ обозначает матрицу с элементами $|a_{ij}|$

$\operatorname{tr} A$ — след матрицы A (сумма диагональных элементов)

$\rho(A)$ — спектральный радиус матрицы A

$$\frac{df(t)}{dt} = f'(t) \text{ — производная } f(t)$$

$\frac{\partial u}{\partial x}$ — частная производная функции u по переменной x . Эквивалентное обозначение u'_x

$\nabla f(x)$ — градиент функции $f(x)$

Литература

1. Беккенбах Э., Беллман Р. Неравенства. М., 2004.
2. Беклемишев Д. В. Курс аналитической геометрии и линейной алгебры. М., 2004.
3. Босс В. Интуиция и математика. М., 2003.
4. Босс В. Лекции по математике. М.: УРСС, 2004.
5. Гантмахер Ф. Р. Лекции по аналитической механике. М., 1966.
6. Гантмахер Ф. Р. Теория матриц. М., 1988.
7. Гельфанд И. М. Лекции по линейной алгебре. М., 1998.
8. Данцер Л., Грюнбаум Б., Кли В. Теорема Хелли и ее применения. М.: Мир, 1968.
9. Декарт Р. Геометрия. М., 1938.
10. Красносельский М. А., Лифшиц Е. А., Соболев А. В. Позитивные линейные системы. Метод положительных операторов. М., 1985.
11. Курош А. Г. Курс высшей алгебры. М., 2004.
12. Линейные неравенства и смежные вопросы: Сб. статей // Под ред. Г. У. Куна и А. У. Таккера. М., 1959.
13. Маркус М., Минк Х. Обзор по теории матриц и матричных неравенств. М.: УРСС, 2004.
14. Никайдо Х. Выпуклые структуры и математическая экономика. М.: Мир, 1972.
15. Постников М. М. Устойчивые многочлены. М.: УРСС, 2004.
16. Рокафеллар Р. Выпуклый анализ. М.: Мир, 1973.
17. Фаддеев Д. К., Фаддеева И. Н. Вычислительные методы линейной алгебры. М.: Физматгиз, 1963.
18. Хорн Р., Джонсон Ч. Матричный анализ. М.: Мир, 1989.
19. Wilkinson J. H. Rounding errors in algebraic processes. Englewood Cliffs, N.J., 1963.

Предметный указатель

Абсцисса 10

алгебраическая кратность 74, 109, 211

амплитудно-фазовая
характеристика 173

Базис 12, 41

беспорядок 57

билинейная форма 71, 100

билинейное разложение 94

биортогональность 92, 210

Ведущее собственное значение
166

вековое уравнение 88

вектор 11, 39, 196

— аксиальный 19

— ковариантный 96

— контравариантный 96

— полярный 19

векторное произведение 19

векторы коллинеарные 12, 196

— компланарные 12, 196

— ортогональные 96

выпуклая комбинация 149

— оболочка 149

Гаусса метод 54

геометрическая кратность 74, 111

гессиан 89

гиперплоскость 45, 149

главная диагональ 45, 201

— подматрица 56

главный минор 101

годограф Михайлова 173

грань многогранника 150

Декартовы координаты 10,
196

детерминант 22, 57, 58

диагональ главная 57

— матрицы 57

доминирующая диагональ 171

допустимый многогранник 158

дорожка невырожденности 169

Евклидово пространство 40

Жорданова клетка 109, 211

— форма 109, 211

Задача двойственная 158

закон инерции 89

Индексы инерции 89

Касательная плоскость 32

квадратичная форма 81

коллинеарные векторы 41

конус 150

— двойственный 150

— заостренный 150

— острый 150

конусный отрезок 163

координаты 10, 13, 41

корневое подпространство 114

критерий Михайлова 173

— Сильвестра 101

кронекерово произведение 142
круги Гершгорина 188

Лемма Кнастера—
Куратовского—Мазуркевича
152

— Минковского—Фаркаша 161,
213

линейная комбинация 41, 201

— независимость 12, 41

— оболочка 41, 201

линейно зависимые векторы 12,
41

линейное пространство 76

линейный оператор 45

Матрица 23, 45

— Адамара 171

— Вандермонда 65

— внедиагонально отрицательная
171

— вполне неразложимая 181

— гурвицева 145, 172

— двоякостochasticкая 178

— дефектная 74

— диагональная 68

— единичная 25, 48

— жорданова 109

— идемпотентная 117

— импримитивная 176, 215

— ковариации 102

— кососимметричная 47

— невырожденная 46, 201

— неразложимая 168, 214

— нильпотентная 116

— нормальная 122

— обратная 25, 48, 202

— ортогональная 67

— перестановки 51

— подобная 74

— положительная 163, 214

— простая 117

— прямоугольная 49

— разложимая 168, 214

— самосопряженная 86

— симметричная 47

— сопровождающая 121

— сопряженная 48, 72

— стохастическая 177, 215

— строго положительная 163, 214

— транспонированная 25, 47

— треугольная 51

— унитарная 103, 210

— устойчивая 172

— элементарная 52

— эрмитова 86, 122

— — сопряженная 103

матрицы элементарные
преобразования 52

метод Гаусса модифицированный
54, 56

— Зейделя 193

— наименьших квадратов 102

— шевеления 165

метрика 126

минор дополнительный 61

многогранник 150

многочлен аннулирующий 112

— минимальный 113

множество аффинное 149

— выпуклое 149

монотонность 162

Невязка 186, 193

неравенство

Коши—Буняковского 40, 78

— Коши—Шварца 40

норма вектора 16, 40, 125

— двойственная 129

— евклидова 40, 126, 212

— кубическая 126, 212
матрицы 127
— октаэдрическая 126, 212
подчиненная 127
нормальные координаты 88

Образ линейного оператора 78
оператор положительный 162
— сильно положительный 162
опорная гиперплоскость 150
определитель 22, 57
— Грама 101
ордината 10
ортогонализация Грама—Шмидта 104
ортогональное дополнение 79
множество 83
ортонормированное множество 83
отделимость 149
отображение аффинное 149

Перестановка 57
плоскость 42
подобные матрицы 68
полином гурвицев 173
— устойчивый 173
полиэдр 150
полуупорядоченность 162
правило Крамера 27, 63
— переброски 71
преобразование ортогональное 83
приведение двух форм 87
проекция на направление 13
произведение по Адамару 101
пространство инвариантное 79
— натянутое 41
прямая сумма 79
прямое произведение 142
псевдовектор 19

Радиус-вектор 11, 196
разделяющая гиперплоскость 149
разложение определителя 61
размерность 41, 77
ранг квадратичной формы 89
— матрицы 55, 203
резольвента 139

Секулярное уравнение 88
сигнатура 89
сингулярное разложение 92
сингулярные числа 91, 210
система координат правая, левая 14
скалярное произведение 16, 39, 200
след матрицы 70
собственное значение 69
— подпространство 115
собственный вектор 69
— — левый 73
— — правый 73
сопряженное пространство 95
спектр матрицы 74
спектральный радиус 74
столбцовая норма 128, 213
строчная норма 128, 213
сужение A на X 79
сумма подпространств 79

Тензор 100
тензорное произведение 142
теорема Биркгофа 178
— Гамильтона—Кэли 117
— Каратеодори 149
— Кронекера—Капелли 64
— об изоморфизме 77
— об отделимости 150
— Перрона 168, 214

— Хелли 153

транспозиция 53, 57

Унитарная эквивалентность 105,
106

Функционал 95

функция выпуклая 149

Характеристический полином 70
характеристическое уравнение 69

Число обусловленности 186

Эквивалентность норм 127, 212

Ядро линейного оператора 78

λ -матрица 119

K -матрица 171

P -матрица 153, 213

P -отображение 154

QR -алгоритм 195

Уважаемые читатели! Уважаемые авторы!

Наше издательство специализируется на выпуске научной и учебной литературы, в том числе монографий, журналов, трудов ученых Российской академии наук, научно-исследовательских институтов и учебных изданий. Мы предлагаем авторам свои услуги на выгодных экономических условиях. При этом мы берем на себя всю работу по подготовке и печати — от набора, редактирования и верстки до тиражирования и распространения.



URSS

Среди вышедших и готовящихся к изданию книг мы предлагаем Вам следующие:



В. Босс

Лекции по математике: анализ

Книга отличается краткостью и прозрачностью изложения, вплоть до объяснений «на пальцах». Значительное внимание уделяется мотивации результатов и укрупненному видению. В первой части дается обширный материал стандартных курсов математического анализа. Во второй, «необязательной», части излагаются — в стиле обзоров и очерков — примыкающие к анализу предметы: аналитические функции, топология и неподвижные точки, векторный анализ. «Высокие материи» рассматриваются на доступном уровне. Книга легко читается.

Для студентов, преподавателей, инженеров и научных работников.

В. Босс. Лекции по математике

Планируются к изданию следующие тома:

Анализ (вышел)
Дифференциальные уравнения (вышел)
Линейная алгебра (вышел)
Вероятность, информация, статистика
Геометрические методы нелинейного анализа
Дискретные задачи
ТФКП
Вычислимость и доказуемость
Оптимизация
Уравнения математической физики
Функциональный анализ
Алгебраические методы
Случайные процессы
Топология
Численные методы

По всем вопросам Вы можете обратиться к нам:
тел./факс (095) 135-42-16, 135-42-46
или **электронной почтой** URSS@URSS.ru
Полный каталог изданий представлен
в **Интернет-магазине**: <http://URSS.ru>

**Научная и учебная
литература**

Представляем Вам наши лучшие книги:



Бэр Р. Линейная алгебра и проективная геометрия.
Золотаревская Д. И. Сборник задач по линейной алгебре.
Чеботарев Н. Г. Основы теории Галуа. В 2 кн.
Чеботарев Н. Г. Введение в теорию алгебр.
Чеботарев Н. Г. Теория алгебраических функций.
Чеботарев Н. Г. Теория групп Ли.
Супруненко Д. А. Группы подстановок.
Супруненко Д. А., Тышкевич Р. И. Перестановочные матрицы.
Маркус М., Минк Х. Обзор по теории матриц и матричных неравенств.
Шевалле К. Введение в теорию алгебраических функций.
Яглом И. М. Необыкновенная алгебра.
Фейнман Р., Лейтон Р., Сэндс М. Фейнмановские лекции по физике.
Вайнберг С. Мечты об окончательной теории.
Грин Б. Элегантная Вселенная. Суперструны и поиски окончательной теории.
Пенроуз Р. НОВЫЙ УМ КОРОЛЯ. О компьютерах, мышлении и законах физики.

В. Босс. Наваждение

Смирнов Ю. М. Курс аналитической геометрии.
Жуков А. В. Вездесущее число «пи».
Вейль А. Основы теории чисел.
Вейль Г. Алгебраическая теория чисел.
Хинчин А. Я. Три жемчужины теории чисел.
Гнеденко Б. В. Курс теории вероятностей.
Боровков А. А. Теория вероятностей.
Гильберт Д., Кон-Фоссен С. Наглядная геометрия.
Клейн Ф. Неевклидова геометрия.
Клейн Ф. Высшая геометрия.
Рашиевский П. К. Курс дифференциальной геометрии.
Дубровин Б. А., Новиков С. П., Фоменко А. Т. Современная геометрия. Т. 1–3.
Боярчук А. К. и др. Справочное пособие по высшей математике (Антидемилович). Т. 1–5.
Краснов М. Л. и др. Вся высшая математика. Т. 1–6.
Краснов М. Л. и др. Сборники задач с подробными решениями.



Тел./факс:

(095) 135-42-46,
(095) 135-42-16,

E-mail:

URSS@URSS.ru

<http://URSS.ru>

Наши книги можно приобрести в магазинах:

«Библио-Глобус» (м. Лубянка, ул. Мясницкая, 6. Тел. (095) 925-2457)
«Московский дом книги» (м. Арбатская, ул. Новый Арбат, 8. Тел. (095) 203-8242)
«Москва» (м. Охотный ряд, ул. Тверская, 8. Тел. (095) 229-7355)
«Молодая гвардия» (м. Полянка, ул. Б. Полянка, 28. Тел. (095) 238-5083, 238-1144)
«Дом деловой книги» (м. Пролетарская, ул. Марсиская, 9. Тел. (095) 270-5421)
«Гнозис» (м. Университет, 1 г-м. корпус МГУ, комн. 141. Тел. (095) 939-4713)
«У Кентавра» (РГТУ) (м. Новослободская, ул. Чапаева, 15. Тел. (095) 973-4301)
«СПб. дом книги» (Невский пр., 28. Тел. (812) 311-3954)



Проект издания 20 томов «Лекций по математике» В. Босса обретает краски.

Автор наращивает обороты
и поднимает планку все выше.

Читательская аудитория ширится.



*В условиях
информационного
наводнения
инструменты
вчерашнего дня
перестают
работать.
Поэтому учить
надо как-то иначе.
«Лекции» дают
пример.
Плохой ли, хороший –
покажет время.
Но в любом случае,
это продукт нового
поколения.
Те же «колеса»,
тот же «руль», та же
математическая
суть, – но по-другому.*

В. Босс

НАУЧНАЯ И УЧЕБНАЯ ЛИТЕРАТУРА



E-mail: URSS@URSS.ru
Каталог изданий в Интернете:
<http://URSS.ru>

Тел./факс: 7 (095) 135-42-16
Тел./факс: 7 (095) 135-42-46

3143 ID 26865



9 785484 000463 >



Из отзывов читателей:

Чтобы усвоить предмет, надо освободить его от деталей, обнажить центральные конструкции, понять, как до теорем можно было додуматься. Это тяжелая работа, на которую не всегда хватает сил и времени. В «Лекциях» такая работа проделывается автором.

Популярность книг В. Босса среди преподавателей легко объяснима. Дается то, чего недостает. Общая картина, мотивация, взаимосвязи. И самое главное — легкость вхождения в любую тему.

Содержание продумано и хорошо увязано. Громоздкие доказательства ужаты до нескольких строчек. Виртуозное владение языком. Что касается замысла изложить всю математику в 20 томах, с трудом верится, что это по силам одному человеку.

Лекции В. Босса — замечательные математические книги. Как учебные пособия, они не всегда отвечают канонам преподавания, но студентам это почему-то нравится.

Отзывы о настоящем издании,
а также обнаруженные опечатки присылайте
по адресу URSS@URSS.ru.
Ваши замечания и предложения будут учтены
и отражены на web-странице этой книги
в нашем интернет-магазине <http://URSS.ru>

