

Глава 1

Ортогональные преобразования

1.1. Понятие ортогональных преобразований

Пусть задан некоторый линейный оператор $Q : \mathbb{R}^n \rightarrow \mathbb{R}^n$, определенный матрицей Q .

Определение 1.1.1 Оператор Q называется ортогональным, если осуществляемое им отображение сохраняет нормы векторов, то есть

$$\forall x \in \mathbb{R}^n : \|Qx\| = \|x\|.$$

Матрица Q , определяющая такой оператор, также называется ортогональной.

Для ортогональных операторов на основании этого определения можно записать

$$\|Qx\|_2^2 = \|x\|_2^2 \implies (Qx, Qx) = (x, x). \quad (1.1)$$

Далее рассмотрим действие ортогонального оператора на сумму двух векторов:

$$(Q(a+b), Q(a+b)) = (Qa, Qa) + 2(Qa, Qb) + (Qb, Qb).$$

С другой стороны,

$$(Q(a+b), Q(a+b)) = (a+b, a+b) = (a, a) + 2(a, b) + (b, b).$$

Вычитая из одного равенства другое и пользуясь тождеством (1.1), получим

$$(Qa, Qb) = (Q^T Qa, b) = (a, b).$$

В силу произвольности выбора векторов a и b отсюда вытекает, что $Q^T Q = I$. Это свойство часто используют в качестве альтернативного определения:

Определение 1.1.2 Линейный оператор Q называется ортогональным, если определяющая его матрица Q удовлетворяет соотношению

$$Q^T Q = I, \quad (1.2)$$

то есть ее обратная матрица совпадает с транспонированной.

Формула (1.2) объясняет происхождение термина «ортогональное преобразование». Действительно, выполнение этого условия означает, что как строки, так и столбцы матрицы Q образуют ортогональные (и, более того, ортонормированные) наборы векторов.

На основании определения 1.1.2 нетрудно показать, что

- произведение ортогональных матриц также является ортогональной матрицей;
- матрица, обратная к ортогональной, также ортогональна.

Сохранение нормы векторов часто оказывается полезным, когда в процессе преобразований необходимо обеспечить выполнение некоторых условий, например, гарантированное неувеличение нормы невязки и т.д.

1.2. Вращения Гивенса

Пусть i и j — некоторые целые индексы, такие что $1 \leq i < j \leq n$, а θ — некоторое вещественное число. Рассмотрим следующую матрицу G_{ij}^θ размерности n :

$$G_{ij}^\theta = \begin{pmatrix} I_{i-1} & \vdots & 0 & \vdots & 0 \\ \cdots & [\cos \theta]_{ii} & \cdots & [\sin \theta]_{ij} & \cdots \\ 0 & \vdots & I_{j-i-1} & \vdots & 0 \\ \cdots & [-\sin \theta]_{ji} & \cdots & [\cos \theta]_{jj} & \cdots \\ 0 & \vdots & 0 & \vdots & I_{n-j} \end{pmatrix} \quad (1.3)$$

(такая матрица отличается от единичной лишь тем, что на пересечениях ее (i, j) строк и столбцов расставлены величины $\cos \theta$ и $\sin \theta$).

Непосредственной проверкой как определения 1.1.1 так и определения 1.1.2 легко убедиться, что определяемый матрицей (1.3) оператор $\mathfrak{G}_{ij}^\theta$ является ортогональным.

Пусть имеется некоторый вектор $x \in \mathbb{R}^n$ и его образ $x' = \mathfrak{G}_{ij}^\theta x$. Очевидно, что они отличаются лишь i -ым и j -ым компонентами, для которых справедливо соотношение

$$\begin{pmatrix} x'_i \\ x'_j \end{pmatrix} = \begin{pmatrix} \cos \theta & \sin \theta \\ -\sin \theta & \cos \theta \end{pmatrix} \begin{pmatrix} x_i \\ x_j \end{pmatrix} = \begin{pmatrix} x_i \cos \theta + x_j \sin \theta \\ x_j \cos \theta - x_i \sin \theta \end{pmatrix}. \quad (1.4)$$

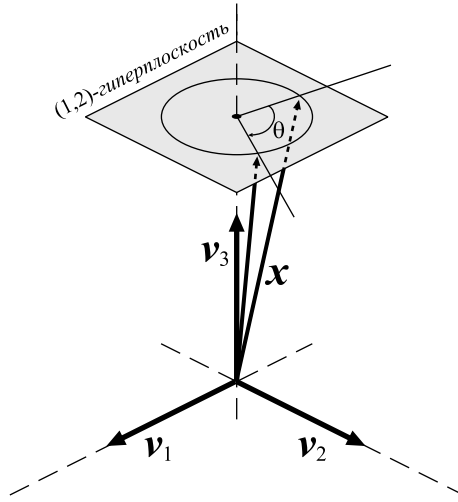


Рис. 1.1. Вращение Гивенса $\mathfrak{G}_{12}^\theta$ в пространстве $\mathbb{R}^3$

Как известно из аналитической геометрии, такое преобразование означает поворот вектора $(x_i, x_j)^T$ на угол θ . Соответственно, для n -мерного вектора x вращение будет происходить в (i, j) -гиперплоскости $\mathcal{M}_{ij}$, образованной всевозможными линейными комбинациями базисных векторов v_i и v_j пространства $\mathbb{R}^n$ (рис. 1.1):

$$\mathcal{M}_{ij} = \{z \mid z = \alpha v_i + \beta v_j, \alpha \in \mathbb{R}, \beta \in \mathbb{R}\}.$$

Определение 1.2.1 Ортогональный оператор $\mathfrak{G}_{ij}^\theta$, определяемый матрицей (1.3), называется вращением Гивенса на угол θ в (i, j) -гиперплоскости $\mathcal{M}_{ij}$ пространства $\mathbb{R}^n$.

На практике обычно возникает задача нахождения такого вращения, которое обнуляло бы один из компонентов x_i или x_j ¹. Обозначим для краткости $c = \cos \theta$ и $s = \sin \theta$. Если требуется обнулить x_j , то согласно (1.4) нужно найти такие c и s , чтобы выполнялось $cx_j - sx_i = 0$. Добавив к этому уравнению связывающее c и s основное тригонометрическое тождество, получим систему из двух уравнений с двумя неизвестными:

$$\begin{cases} cx_j - sx_i = 0 \\ c^2 + s^2 = 1, \end{cases}$$

имеющую решение

$$s = \sin \theta = \frac{x_j}{\sqrt{x_i^2 + x_j^2}}, \quad c = \cos \theta = \frac{x_i}{\sqrt{x_i^2 + x_j^2}}.$$

¹Так называемая задача поляризации вектора [?]. Пример применения см. §??.

Для хранения матриц вращения вида (1.3) достаточно двух целых (i, j) и двух вещественных (s, c) чисел; процедура матрично-векторного умножения при этом выполняется согласно (1.4).

1.3. Отражения Хаусхолдера

Пусть задан некоторый ненулевой вектор $\mathbf{p} \in \mathbb{R}^n$. Будем говорить, что он определяет $(n - 1)$ -мерную гиперплоскость $\mathcal{N}_{\mathbf{p}}$, образованную всеми ортогональными к $\mathbf{p}$ векторами пространства $\mathbb{R}^n$:

$$\mathcal{N}_{\mathbf{p}} = \{z \mid (z, \mathbf{p}) = 0\}.$$

При этом очевидно, что норма вектора $\mathbf{p}$ для задания гиперплоскости не имеет значения, поэтому для простоты дальнейших рассуждений будем полагать его нормированным: $\|\mathbf{p}\|_2 = 1$.

Рассмотрим произвольный вектор $x \in \mathbb{R}^n$ и представим его в виде суммы двух компонент:

$$x = x_{\mathbf{p}} + x_N, \quad \text{где} \quad x_{\mathbf{p}} = (\mathbf{p}, x)\mathbf{p};$$

$$x_N = x - (\mathbf{p}, x)\mathbf{p}.$$

Векторы $\mathbf{p}$ и $x_{\mathbf{p}}$ коллинеарны по построению. Покажем что $x_N \in \mathcal{N}_{\mathbf{p}}$. Действительно,

$$(\mathbf{p}, x_N) = (\mathbf{p}, x) - (\mathbf{p}, x)(\mathbf{p}, \mathbf{p}) = (\mathbf{p}, x) - (\mathbf{p}, x)\|\mathbf{p}\|_2^2 = 0.$$

Таким образом, $x_{\mathbf{p}} \perp x_N$.

Из приведенных рассуждений следует, что вектор $x - x_{\mathbf{p}}$ является проекцией x на гиперплоскость $\mathcal{M}_{\mathbf{p}}$, а вектор $x' = x - 2x_{\mathbf{p}}$ симметричен к вектору x относительно той же гиперплоскости (рис. 1.2).

Запишем вектор $x_{\mathbf{p}}$ в виде $x_{\mathbf{p}} = \mathbf{p}(\mathbf{p}^T x) = (\mathbf{p}\mathbf{p}^T)x$, тогда вектор x' можно записать как

$$x' = (\mathbf{I} - 2\mathbf{p}\mathbf{p}^T)x.$$

Введем обозначение

$$\mathbf{S}_p = \mathbf{I} - 2\mathbf{p}\mathbf{p}^T \tag{1.5}$$

и проверим произведение $\mathbf{S}_p^T \mathbf{S}_p$:

$$\begin{aligned} \mathbf{S}_p^T \mathbf{S}_p &= (\mathbf{I} - 2\mathbf{p}\mathbf{p}^T)^T (\mathbf{I} - 2\mathbf{p}\mathbf{p}^T) = \mathbf{I} - 4\mathbf{p}\mathbf{p}^T + 4\mathbf{p}\mathbf{p}^T \mathbf{p}\mathbf{p}^T \\ &= \mathbf{I} - 4\mathbf{p}\mathbf{p}^T + 4\|\mathbf{p}\|_2^2 \mathbf{p}\mathbf{p}^T = \mathbf{I}. \end{aligned}$$

Таким образом, линейный оператор $\mathfrak{H}_p$, заданный матрицей $\mathbf{S}_p$ (1.5) по определению 1.1.2 является ортогональным.

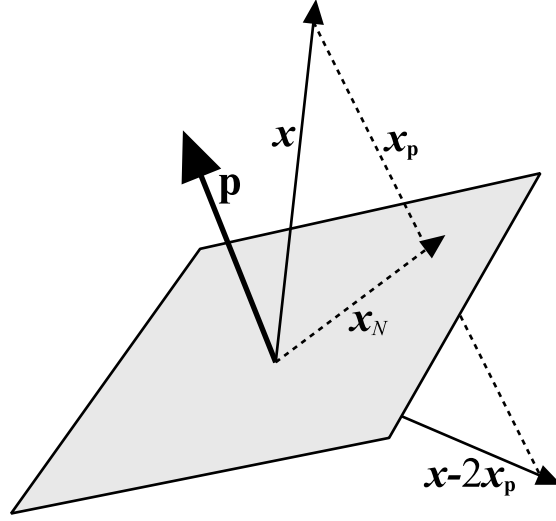


Рис. 1.2. Симметрия относительно гиперплоскости $\mathcal{N}_p$

Определение 1.3.1 Оператор $\mathfrak{H}_p$ с матрицей S_p , определяемой соотношением (1.5) для вектора p , такого что $\|p\|_2 = 1$, называется отражением Хаусхолдера относительно гиперплоскости $\mathcal{N}_p$ в пространстве $\mathbb{R}^n$.

Заметим, что матрица S_p по построению симметрична, а следовательно

$$S_p^T S_p = S_p S_p = (S_p)^2 = I.$$

Действительно, геометрически двукратное отражение вектора относительно одной и той же гиперплоскости означает его перевод сам в себя.

Основной задачей при использовании отражений является подбор такого вектора p , чтобы оператор $\mathfrak{H}_p$ переводил заранее заданный вектор $x \in \mathbb{R}^n$ в вектор, у которого заданное число последних компонент равно нулю:

$$\mathfrak{H}_p : (x_1, x_2, \dots, x_n)^T \rightarrow (x'_1, x'_2, \dots, x'_k, 0, \dots, 0). \quad (1.6)$$

Рассмотрим сначала вспомогательную задачу о подборе отражения, переводящего x в вектор вида αe_1 , $\alpha \in \mathbb{R}$. Необходимо для x подобрать такой p , чтобы

$$(I - 2pp^T)x = \alpha e_1. \quad (1.7)$$

Раскрыв скобки в левой части и перенеся x в правую, получим

$$2p(p^T x) = x - \alpha e_1. \quad (1.8)$$

Поскольку величина $2\mathbf{r}^T x$ является скалярной, равенство (1.8) означает что векторы $\mathbf{r}$ и $(x - \alpha \mathbf{e}_1)$ коллинеарны. С учетом условия $\|\mathbf{r}\|_2 = 1$ можно записать:

$$\mathbf{r} = \pm \frac{x - \alpha \mathbf{e}_1}{\|x - \alpha \mathbf{e}_1\|_2}. \quad (1.9)$$

Выбор знака перед дробью не играет роли²; будем для определенности брать знак «+». Остается определить коэффициент α .

Сделать это можно из условия ортогональности оператора $\mathfrak{H}_p$:

$$\|\mathfrak{H}_p x\|_2 = \|x\|_2 = \|\alpha \mathbf{e}_1\|_2 = |\alpha|,$$

откуда $\alpha = \pm \|x\|_2$.

Из соображений численной устойчивости, чтобы сделать знаменатель (1.9) возможно больше, обычно выбирают

$$\alpha = -\operatorname{sgn}(x_1)\|x\|_2. \quad (1.10)$$

Тогда формулы (1.9) и (1.10) вместе дают искомое выражение для $\mathbf{r}$:

$$\mathbf{r} = \frac{x + \operatorname{sgn}(x_1)\|x\|_2 \mathbf{e}_1}{\|x + \operatorname{sgn}(x_1)\|x\|_2 \mathbf{e}_1\|_2}.$$

Для решения общей задачи (1.6) применим блочный подход. Будем искать вектор $\mathbf{r}$ в виде

$$\mathbf{r} = (0_1, \dots, 0_{k-1}, p_k, p_{k+1}, \dots, p_n)^T.$$

Тогда матрица $\mathbf{S}_p$ будет иметь вид

$$\mathbf{S}_p = \begin{pmatrix} \mathbf{I}_{k-1} & \mathbf{0} \\ \mathbf{0} & \tilde{\mathbf{S}}_p^{(n-k+1)} \end{pmatrix}$$

и (1.6) сведется к задаче (1.7) уменьшенной размерности для вектора $(x_k, x_{k+1}, \dots, x_n)^T$.

С учетом сказанного, вектор $\mathbf{r}$, решающий задачу (1.6), находится по следующим формулам:

$$\mathbf{r} = z / \|z\|_2; \quad (1.11)$$

$$z_i = \begin{cases} 0, & i < k \\ \alpha + x_k, & i = k \\ x_i, & i > k \end{cases} \quad (1.12)$$

$$\alpha = \operatorname{sgn}(x_k) \left(\sum_{j=k}^n x_j^2 \right)^{1/2}. \quad (1.13)$$

Для хранения матриц отражений достаточно задать вектор $\mathbf{r}$, тогда элементы $\mathbf{S}_p$ могут быть легко вычислены по формуле

$$[\mathbf{S}_p]_{ij} = \delta_{ij} - \mathbf{r}_i \mathbf{r}_j,$$

где δ_{ij} — символ Кронекера-Капелли.

²Так как векторы p и $(-p)$ определяют одну и ту же гиперплоскость $\mathcal{N}_p$.

1.4. Ортогонализация Грама-Шмидта

В §1.1 отмечалось, что строки и столбцы ортогональной матрицы являются ортонормированными наборами векторов. Сформулируем следующую задачу.

Пусть задан произвольный набор векторов $\{x_i\}_{i=1}^m$ из пространства $\mathbb{R}^n$. требуется построить такой ортонормальный набор векторов $\{q_i\}_{i=1}^m$, чтобы линейные оболочки, натянутые на эти два набора, совпадали:

$$\text{span}\{x_1, \dots, x_m\} = \text{span}\{q_1, \dots, q_m\}, \quad (q_i, q_j) = \delta_{ij}.$$

Если ввести обозначение $\mathbf{Q} = [q_1 | \dots | q_m]$, то требование ортонормированности векторов q_i означает ортогональность в смысле определения 1.1.2 матрицы $\mathbf{Q}$. Очевидно, добиться этого можно лишь тогда, когда исходные векторы q_i линейно независимы.

Для построения векторов q_i применим индуктивный подход. Положим

$$q_1 = x_1 / \varrho_{11}, \quad \varrho_{11} = \|x_1\|_2. \quad (1.14)$$

Далее, полагая векторы $q_1, \dots, q_{k-1}$ уже построенными, найдем вектор q_k в виде

$$\varrho_{kk} q_k = x_k - \sum_{i=1}^{k-1} \varrho_{ik} q_i, \quad (1.15)$$

где ϱ_{kk} выбирается из условия $\|q_k\|_2 = 1$ (ситуация $\varrho_{kk} = 0$ означает линейную зависимость векторов и прекращение процесса), а коэффициенты ϱ_{ik} — из условия ортогональности векторов.

Умножим (1.15) скалярно на q_j для $j = \overline{1, k-1}$:

$$\varrho_{kk}(q_k, q_j) = 0 = (x_k, q_j) - \sum_{i=1}^{k-1} \varrho_{ik}(q_i, q_j) = (x_k, q_j) - \varrho_{jk}(q_j, q_j),$$

откуда с учетом $(q_j, q_j) = \|q_j\|_2^2 = 1$

$$\varrho_{jk} = (x_k, q_j). \quad (1.16)$$

Определение 1.4.1 Процесс перехода от векторов x_i к векторам q_i согласно формулам (1.14)–(1.16) называется ортогонализацией Грама-Шмидта.

С точки зрения программной реализации ортогонализация Грама-Шмидта обладает существенным недостатком: коэффициенты ϱ в ней вычисляются до операций линейного комбинирования, без учета возникающих при этом погрешностей округления, что ведет к

Обычная	Модифицированная
Для j от 1 до m Для i от 1 до $j - 1$ $\varrho_{ij} := (x_j, q_i)$ увеличить i $\tilde{q}_j := x_j - \sum_{i=1}^{j-1} \varrho_{ij} q_i$ $\varrho_{jj} := \ \tilde{q}_j\ _2$ Если $\varrho_{jj} = 0$ то КОНЕЦ $q_j := \tilde{q}_j / \varrho_{jj}$ увеличить j	Для j от 1 до m $\tilde{q}_j = x_j$ Для i от 1 до $j - 1$ $\varrho_{ij} := (\tilde{q}_j, q_i)$ $\tilde{q}_j := \tilde{q}_j - \varrho_{ij} q_i$ увеличить i $\varrho_{jj} := \ \tilde{q}_j\ _2$ Если $\varrho_{jj} = 0$ то КОНЕЦ $q_j := \tilde{q}_j / \varrho_{jj}$ увеличить j

Таблица 1.1. Ортогонализация Грама-Шмидта

потере ортогональности векторов q . Избежать этого можно, совместив вычисление ϱ с линейным комбинированием. В таблице 1.1 приведены оба варианта алгоритма.

Если записать (1.15) в виде

$$x_k = \sum_{i=1}^k \varrho_{ik} q_i$$

и обозначить $\mathbf{X} = [x_1 | \dots | x_m]$, то ортогонализацию Грама-Шмидта можно записать в матричной форме

$$\mathbf{X} = \mathbf{Q}\mathbf{R}, \quad (1.17)$$

где $\mathbf{R}$ — верхнетреугольная матрица, составленная из коэффициентов ϱ_{ij} :

$$\mathbf{R} = \begin{pmatrix} \varrho_{11} & \varrho_{12} & \cdots & \varrho_{1m} \\ & \varrho_{22} & \cdots & \varrho_{2m} \\ & & \ddots & \vdots \\ 0 & & & \varrho_{mm} \end{pmatrix}.$$

Определение 1.4.2 Представление произвольной матрицы $\mathbf{X}$ в виде (1.17), где матрицы $\mathbf{Q}$ и $\mathbf{R}$ являются ортогональной и треугольной соответственно, называется QR-факторизацией.

QR-факторизация играет важную роль не только при решении СЛАУ, но и в численных методах решения проблемы собственных значений.

1.5. Ортогонализация Хаусхолдера

Наряду с рассмотренной в предыдущем параграфе ортогонализацией Грама-Шмидта можно построить еще один алгоритм QR-факторизации, в котором ортогональность матрицы Q обеспечивается автоматически. Воспользуемся для этого отражениями Хаусхолдера.

Матрица X имеет размерность $n \times m$, причем $m \leq n$. В §1.3 был описан способ построения отражения, позволяющего обнулить заданное число последних компонент любого вектора (формулы (1.11)–(1.13)). Построим набор отражений $\mathfrak{H}_1, \mathfrak{H}_2, \dots, \mathfrak{H}_m$ следующим образом: $\mathfrak{H}_1$ обнуляет все компоненты вектора x_1 за исключением самой первой, $\mathfrak{H}_2$ обнуляет все компоненты вектора x_2 за исключением первых двух и т.д. Последний оператор $\mathfrak{H}_m$ обнуляет все компоненты вектора x_m за исключением m начальных.

Результатом применения построенных операторов к матрице X будет

$$\mathfrak{H}_m \dots \mathfrak{H}_2 \mathfrak{H}_1 X = \begin{pmatrix} R_{m \times m} \\ Z_{(n-m) \times m} \end{pmatrix}, \quad (1.18)$$

где R — верхнетреугольная матрица, а Z — нулевой блок.

Введем обозначения

$$P = \mathfrak{H}_m \dots \mathfrak{H}_2 \mathfrak{H}_1, \quad E_m = \begin{pmatrix} I_m \\ Z_{(n-m) \times m} \end{pmatrix},$$

тогда (1.18) можно записать в виде

$$PX = E_m R.$$

Будучи произведением ортогональных матриц, P сама является ортогональной, и для нее $P^{-1} = P^T$. Следовательно,

$$X = (P^T E_m) R. \quad (1.19)$$

Рассмотрим участвующую в этом соотношении матрицу $P^T E_m$.

$$(P^T E_m)^T (P^T E_m) = E_m^T (P^T P) E_m = E_m^T E_m = I.$$

В силу определения 1.1.2 данная матрица ортогональна. Таким образом, если ввести для нее обозначение $Q = P^T E_m$, то (1.19) есть искомая QR-факторизация матрицы X .

Описанную схему можно вкратце сформулировать следующим образом:

1. Последовательно рассматривать векторы $x_1, x_2, \dots, x_m$.

2. Для очередного вектора x_k найти $\tilde{r}_k = \mathfrak{H}_{k-1}\mathfrak{H}_{k-2}\dots\mathfrak{H}_1x_k$.
3. По формулам (1.11)–(1.13) найти вектор $\mathbf{p}_k$, обнуляющий последние $(n - k)$ компонент вектора x_k .
4. Положить $\mathfrak{H}_k = \mathbf{I} - \mathbf{p}_k\mathbf{p}_k^T$.
5. Найти $r_k = \mathfrak{H}_k\tilde{r}_k$.
6. Найти $q_k = \mathfrak{H}_1\mathfrak{H}_2\dots\mathfrak{H}_k\mathbf{e}_k$.
7. Перейти к следующему вектору x_k .

Такой алгоритм называется ортогонализацией Хаусхолдера; здесь векторы r_k являются столбцами матрицы $\mathbf{R}$.

Замечание. Если $n = m$, то $\mathbf{E}_m = \mathbf{I}_m$ и $\mathbf{Q} = \mathbf{P}^T = \mathbf{P}$. В этом случае одновременное применение операторов $\mathfrak{H}_k$ к $\mathbf{X}$ и вектору b означает решение СЛАУ $\mathbf{X}y = b$ путем приведения ее к верхнетреугольному виду (метод отражений).