

Глава 7 Снижение размерности

В главе 6 мы рассмотрели некоторые характеристики многомерных данных, которые подчас противоречат здравому смыслу. Например, в пространствах высокой размерности точки чаще рассредоточены по поверхности пространства или его углам, а центр практически пуст. Можно наблюдать явное увеличение ортогональных осей. Как следствие, данные высокой размерности создают определённые проблемы для анализа, хотя в некоторых применениях (например, в задачах классификации) высокие размерности могут быть полезны. Тем не менее, важно проверить, возможно ли понижение размерности без потери основных свойств исходной матрицы. Это может помочь как в интеллектуальном анализе данных, так и при их визуализации. В этой главе мы изучим методы, которые позволяют получить оптимальные проекции данных на меньшие размерности.

7.1 Предпосылки

Пусть данные представляют из себя матрицу $\mathbf{D}$ размера $n \times d$ (n точек с d признаками):

$$\mathbf{D} = \begin{pmatrix} \mathbf{x}_1 & \begin{array}{c|ccc} X_1 & X_2 & \cdots & X_d \\ \hline x_{11} & x_{12} & \cdots & x_{1d} \end{array} \\ \mathbf{x}_2 & \begin{array}{c|ccc} x_{21} & x_{22} & \cdots & x_{2d} \end{array} \\ \vdots & \begin{array}{c|ccc} \vdots & \vdots & \ddots & \vdots \end{array} \\ \mathbf{x}_n & \begin{array}{c|ccc} x_{n1} & x_{n2} & \cdots & x_{nd} \end{array} \end{pmatrix}.$$

Каждая точка $\mathbf{x}_i = (x_{i1}, x_{i2}, \dots, x_{id})^T$ — это вектор в d -мерном пространстве с базисом $\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_d$, в котором компонент $\mathbf{e}_i$ соотносится с i -тым признаком X_i . Напомним, что стандартным базисом называется ортонормированный базис для пространства данных, то есть, базисные векторы попарно ортогональны: $\mathbf{e}_i^T \mathbf{e}_j = 0$, и имеют единичную длину: $\|\mathbf{e}_i\| = 1$.

Таким образом, для любого набора d ортонормированных векторов $\mathbf{u}_1, \mathbf{u}_2, \dots, \mathbf{u}_d$ такого, что $\mathbf{u}_i^T \mathbf{u}_j = 0$ и $\|\mathbf{u}_i\| = 1$ (или $\mathbf{u}_i^T \mathbf{u}_i = 1$), мы можем представить каждую точку $\mathbf{x}$ как линейную комбинацию:

$$\mathbf{x} = a_1 \mathbf{u}_1 + a_2 \mathbf{u}_2 + \cdots + a_d \mathbf{u}_d, \quad (7.1)$$

где вектор $\mathbf{a} = (a_1, a_2, \dots, a_d)^T$ — это координаты $\mathbf{x}$ в новом базисе. Линейную комбинацию выше можно записать в виде умножения матриц:

$$\mathbf{x} = \mathbf{U} \mathbf{a}, \quad (7.2)$$

где $\mathbf{U}$ — матрица размера $d \times d$, i -ый столбец которой содержит i -тый базисный вектор:

$$\mathbf{U} = \begin{pmatrix} | & | & \cdots & | \\ \mathbf{u}_1 & \mathbf{u}_2 & \cdots & \mathbf{u}_d \\ | & | & & | \end{pmatrix}.$$

Матрица $\mathbf{U}$ является *ортогональной*, все её столбцы, будучи базисными векторами, *ортогональны*, то есть, они попарно ортогональны и имеют единичную длину:

$$\mathbf{u}_i^T \mathbf{u}_j = \begin{cases} 1, & \text{если } i = j \\ 0, & \text{если } i \neq j \end{cases}$$

Поскольку $\mathbf{U}$ ортогональна, её обратная матрица равна её транспонированной:

$$\mathbf{U}^{-1} = \mathbf{U}^T,$$

следовательно: $\mathbf{U}^T \mathbf{U} = \mathbf{I}$, где $\mathbf{I}$ — единичная матрица размера $d \times d$.

Умножив ур. (7.2) слева на $\mathbf{U}^T$, получим выражение для вычисления координат $\mathbf{x}$ в новом базисе:

$$\begin{aligned} \mathbf{U}^T \mathbf{x} &= \mathbf{U}^T \mathbf{U} \mathbf{a} \\ \mathbf{a} &= \mathbf{U}^T \mathbf{x}. \end{aligned} \quad (7.3)$$

Поскольку теоретически существует бесконечное число наборов ортонормальных базисных векторов, возникает вполне естественный вопрос: существует ли оптимальный базис для конкретного толкования оптимальности? Кроме того, часто входная размерность d достаточно велика, что может вызвать различные проблемы из-за проклятия размерности (см. главу 6). Вполне естественно поставить вопрос о поиске пространства меньшей размерности, которое сохраняло бы основные характеристики данных. То есть, нам интересно найти r -мерное представление $\mathbf{D}$ такое, чтобы $r \ll d$, или, иными словами, рассматривая произвольную точку $\mathbf{x}$ и полагая, что базисные векторы отсортированы в порядке уменьшения значимости, мы могли бы усечь выражение 7.1 до r членов и получить следующее:

$$\mathbf{x}' = a_1 \mathbf{u}_1 + a_2 \mathbf{u}_2 + \dots + a_r \mathbf{u}_r = \sum_{i=1}^r a_i \mathbf{u}_i. \quad (7.4)$$

Здесь $\mathbf{x}'$ — это проекция $\mathbf{x}$ на первые r базисных векторов, что в терминах матричных операций можно записать следующим образом:

$$\mathbf{x}' = \begin{pmatrix} | & | & \dots & | \\ \mathbf{u}_1 & \mathbf{u}_2 & \dots & \mathbf{u}_r \\ | & | & & | \end{pmatrix} \begin{pmatrix} a_1 \\ a_2 \\ \vdots \\ a_r \end{pmatrix} = \mathbf{U}_r \mathbf{a}_r, \quad (7.5)$$

где $\mathbf{U}_r$ — матрица, состоящая из r первых базисных векторов, а $\mathbf{a}_r$ — вектор, содержащий первые r координат. Далее, ограничив ур-е 7.3 первыми r членами, получим:

$$\mathbf{a}_r = \mathbf{U}_r^T \mathbf{x}. \quad (7.6)$$

Подставив 7.6 в 7.5, можно получить окончательное выражение для проекции $\mathbf{x}$ на первые r базисных векторов:

$$\mathbf{x}' = \mathbf{U}_r \mathbf{U}_r^T \mathbf{x} = \mathbf{P}_r \mathbf{x}, \quad (7.7)$$

где $\mathbf{P}_r = \mathbf{U}_r \mathbf{U}_r^T$ — *ортогональная проекционная матрица* для подпространства над первыми r базисными векторами. То есть, $\mathbf{P}_r$ симметрична, и $\mathbf{P}_r^2 = \mathbf{P}_r$. Это легко

проверить, поскольку: $\mathbf{P}_r^T = (\mathbf{U}_r \mathbf{U}_r^T)^T = \mathbf{U}_r \mathbf{U}_r^T = \mathbf{P}_r$ и $\mathbf{P}_r^2 = (\mathbf{U}_r \mathbf{U}_r^T)(\mathbf{U}_r \mathbf{U}_r^T) = \mathbf{U}_r \mathbf{U}_r^T = \mathbf{P}_r$.

Здесь мы использовали тот факт, что $\mathbf{U}_r^T \mathbf{U}_r = \mathbf{I}_{r \times r}$, единичная матрица размера $r \times r$.

Проекционную матрицу $\mathbf{P}_r$ можно легко записать как декомпозицию:

$$\mathbf{P}_r = \mathbf{U}_r \mathbf{U}_r^T = \sum_{i=1}^r \mathbf{u}_i \mathbf{u}_i^T. \quad (7.8)$$

Из выражений 7.1 и 7.4 понятно, что проекция $\mathbf{x}$ на оставшиеся размерности составляет *вектор ошибки*:

$$\epsilon = \sum_{i=r+1}^d a_i \mathbf{u}_i = \mathbf{x} - \mathbf{x}'.$$

Следует заметить, что $\mathbf{x}'$ и ϵ — ортогональные векторы:

$$\mathbf{x}'^T \epsilon = \sum_{i=1}^r \sum_{j=r+1}^d a_i a_j \mathbf{u}_i^T \mathbf{u}_j = 0,$$

что следует из ортонормированности базиса. Более того, мы можем сделать ещё более сильное утверждение: подпространство над первыми r базисными векторами:

$$S_r = \text{span}(\mathbf{u}_1, \dots, \mathbf{u}_r)$$

и подпространство над оставшимися базисными векторами:

$$S_{d-r} = \text{span}(\mathbf{u}_{r+1}, \dots, \mathbf{u}_d)$$

являются *ортогональными подпространствами*, то есть, любая пара векторов $\mathbf{x} \in S_r$ и $\mathbf{y} \in S_{d-r}$ должна быть ортогональной. Подпространство S_{d-r} называют также *ортогональным дополнением* S_r .

Цель понижения размерности — найти такой r -мерный базис, который даст наилучшее приближение $\mathbf{x}'_i$ для всех точек $\mathbf{x}_i \in \mathbf{D}$, либо уменьшить ошибку $\epsilon = \mathbf{x} - \mathbf{x}'$ для всех точек.

7.2 Метод главных компонент

Метод главных компонент (РСА) – это метод, который ищет r -мерный базис, наилучшим образом описывающий дисперсию данных. Направление с наибольшей проецируемой дисперсией называется первым главным компонентом. Ортогональное направление, которое захватывает вторую по величине проецируемую дисперсию, называется вторым главным компонентом и т.д. Как мы увидим, направление, которое максимизирует дисперсию, также минимизирует среднеквадратичную ошибку.

7.2.1 Наилучшая линейная аппроксимация

Мы начнем с $r = 1$, т.е. с одномерного подпространства или прямой $\mathbf{u}$, которая наилучшим образом аппроксимирует $\mathbf{D}$, с точки зрения дисперсии проецируемых точек. Это приведет к общей методике РСА для наилучшего $1 \leq r \leq d$ размерного базиса для $\mathbf{D}$.

Без потери общности, мы предполагаем, что $\mathbf{u}$ имеет величину $\|\mathbf{u}\|^2 = \mathbf{u}^T \mathbf{u} = 1$; в противном случае можно продолжать увеличивать проецируемую дисперсию, просто увеличивая величину $\mathbf{u}$. Мы также предполагаем, что данные центрированы так, что среднее значение $\boldsymbol{\mu} = \mathbf{0}$.

Проекция $\mathbf{x}_i$ на вектор $\mathbf{u}$ берется как

$$\mathbf{x}'_i = \left(\frac{\mathbf{u}^T \mathbf{x}_i}{\mathbf{u}^T \mathbf{u}} \right) \mathbf{u} = (\mathbf{u}^T \mathbf{x}_i) \mathbf{u} = a_i \mathbf{u}$$

где скаляр

$$a_i = \mathbf{u}^T \mathbf{x}_i$$

дает координату $\mathbf{x}'_i$ вдоль прямой $\mathbf{u}$. Заметим, что поскольку средняя точка $\boldsymbol{\mu} = \mathbf{0}$, ее координата вдоль $\mathbf{u}$ равна $\mu_u = 0$.

Необходимо выбрать направление $\mathbf{u}$ так, чтобы дисперсия проецируемых точек была максимальной. Проецируемая дисперсия вдоль $\mathbf{u}$ берется как:

$$\begin{aligned} \sigma_u^2 &= \frac{1}{n} \sum_{i=1}^n (a_i - \mu_u)^2 \\ &= \frac{1}{n} \sum_{i=1}^n (\mathbf{u}^T \mathbf{x}_i)^2 \\ &= \frac{1}{n} \sum_{i=1}^n \mathbf{u}^T (\mathbf{x}_i \mathbf{x}_i^T) \mathbf{u} \\ &= \mathbf{u}^T \left(\frac{1}{n} \sum_{i=1}^n \mathbf{x}_i \mathbf{x}_i^T \right) \mathbf{u} \\ &= \mathbf{u}^T \boldsymbol{\Sigma} \mathbf{u} \end{aligned} \tag{7.9}$$

где $\boldsymbol{\Sigma}$ – это ковариационная матрица для центрированных данных $\mathbf{D}$.

Чтобы максимизировать проецируемую дисперсию, мы должны решить задачу ограниченной оптимизации, а именно, чтобы максимизировать σ_u^2 с учетом ограничения $\mathbf{u}^T \mathbf{u} = 1$. Это может быть решено путем введения множителя Лагранжа α для ограничения, чтобы получить задачу безусловной максимизации.

$$\max_{\mathbf{u}} J(\mathbf{u}) = \mathbf{u}^T \boldsymbol{\Sigma} \mathbf{u} - \alpha (\mathbf{u}^T \mathbf{u} - 1) \tag{7.10}$$

Приравнявая производную $J(\mathbf{u})$ по $\mathbf{u}$ к нулевому вектору, получаем

$$\begin{aligned}\frac{\partial}{\partial \mathbf{u}} J(\mathbf{u}) &= \mathbf{0} \\ \frac{\partial}{\partial \mathbf{u}} (\mathbf{u}^T \Sigma \mathbf{u} - \alpha (\mathbf{u}^T \mathbf{u} - 1)) &= \mathbf{0} \\ 2\Sigma \mathbf{u} - 2\alpha \mathbf{u} &= \mathbf{0} \\ \Sigma \mathbf{u} &= \alpha \mathbf{u}\end{aligned}\tag{7.11}$$

Это означает, что α является собственным значением ковариационной матрицы Σ , с соответствующим собственным вектором $\mathbf{u}$. Далее, взяв скалярное произведение с $\mathbf{u}$ по обе стороны от уравнения (7.11) получаем

$$\mathbf{u}^T \Sigma \mathbf{u} = \mathbf{u}^T \alpha \mathbf{u}$$

Из уравнения (7.9) имеем

$$\begin{aligned}\sigma_{\mathbf{u}}^2 &= \alpha \mathbf{u}^T \mathbf{u} \\ \text{или } \sigma_{\mathbf{u}}^2 &= \alpha\end{aligned}\tag{7.12}$$

Таким образом, чтобы максимизировать проецируемую дисперсию $\sigma_{\mathbf{u}}^2$, необходимо выбрать наибольшее собственное значение матрицы Σ . Другими словами, доминирующий собственный вектор $\mathbf{u}_1$ задает направление наибольшей дисперсии, также называемое первым главным компонентом, т.е. $\mathbf{u} = \mathbf{u}_1$. Кроме того, наибольшее собственное значение λ_1 определяет проецируемую дисперсию, т.е. $\sigma_{\mathbf{u}}^2 = \alpha = \lambda_1$.

Минимальная квадратичная ошибка

Теперь покажем, что направление, которое максимизирует проецируемую дисперсию, также минимизирует среднеквадратичную ошибку. Как и раньше, предположим, что набор данных $\mathbf{D}$, был центрирован путем вычитания среднего значения из каждой точки. Для точки $\mathbf{x}_i \in \mathbf{D}$ пусть $\mathbf{x}'_i$ обозначает ее проекцию вдоль направления $\mathbf{u}$, и пусть $\boldsymbol{\epsilon}_i = \mathbf{x}_i - \mathbf{x}'_i$ обозначает вектор ошибки. Среднеквадратичная ошибка (MSE) условия оптимизации определяется как

$$MSE(\mathbf{u}) = \frac{1}{n} \sum_{i=1}^n \|\boldsymbol{\epsilon}_i\|^2\tag{7.13}$$

$$\begin{aligned}&= \frac{1}{n} \sum_{i=1}^n \|\mathbf{x}_i - \mathbf{x}'_i\|^2 \\ &= \frac{1}{n} \sum_{i=1}^n (\mathbf{x}_i - \mathbf{x}'_i)^T (\mathbf{x}_i - \mathbf{x}'_i) \\ &= \frac{1}{n} \sum_{i=1}^n (\|\mathbf{x}_i\|^2 - 2\mathbf{x}_i^T \mathbf{x}'_i + (\mathbf{x}'_i)^T \mathbf{x}'_i)\end{aligned}\tag{7.14}$$

Заметив, что $\mathbf{x}'_i = (\mathbf{u}^T \mathbf{x}_i) \mathbf{u}$, имеем

$$= \frac{1}{n} \sum_{i=1}^n (\|\mathbf{x}_i\|^2 - 2\mathbf{x}_i^T (\mathbf{u}^T \mathbf{x}_i) \mathbf{u} + ((\mathbf{u}^T \mathbf{x}_i) \mathbf{u})^T ((\mathbf{u}^T \mathbf{x}_i) \mathbf{u}))$$

$$\begin{aligned}
&= \frac{1}{n} \sum_{i=1}^n (\|x_i\|^2 - 2(u^T x_i)(x_i^T u) + (u^T x_i)(x_i^T u)u^T u) \\
&= \frac{1}{n} \sum_{i=1}^n (\|x_i\|^2 - (u^T x_i)(x_i^T u)) \\
&= \frac{1}{n} \sum_{i=1}^n \|x_i\|^2 - \frac{1}{n} \sum_{i=1}^n u^T (x_i x_i^T) u \\
&= \frac{1}{n} \sum_{i=1}^n \|x_i\|^2 - u^T \left(\frac{1}{n} \sum_{i=1}^n x_i x_i^T \right) u \\
&= \sum_{i=1}^n \frac{\|x_i\|^2}{n} - u^T \Sigma u
\end{aligned} \tag{7.15}$$

Заметим, что из уравнения (1.4) общая дисперсия центрированных данных (т.е. $\mu = \mathbf{0}$) получается как

$$\text{var}(\mathbf{D}) = \frac{1}{n} \sum_{i=1}^n \|x_i - \mathbf{0}\|^2 = \frac{1}{n} \sum_{i=1}^n \|x_i\|^2$$

Далее, из уравнения (2.28) имеем

$$\text{var}(\mathbf{D}) = \text{tr}(\Sigma) = \sum_{i=1}^d \sigma_i^2$$

Таким образом, можно переписать уравнение (7.15) как

$$MSE(u) = \text{var}(\mathbf{D}) - u^T \Sigma u = \sum_{i=1}^d \sigma_i^2 - u^T \Sigma u$$

Поскольку первый член $\text{var}(\mathbf{D})$ является константой для данного набора данных $\mathbf{D}$, вектор u , который минимизирует $MSE(u)$, максимизирует второй член, проецируемую дисперсию $u^T \Sigma u$. Поскольку мы знаем, что u_1 , доминирующий собственный вектор матрицы Σ , максимизирует проецируемую дисперсию, имеем

$$MSE(u_1) = \text{var}(\mathbf{D}) - u_1^T \Sigma u_1 = \text{var}(\mathbf{D}) - u_1^T \lambda_1 u_1 = \text{var}(\mathbf{D}) - \lambda_1 \tag{7.16}$$

Таким образом, главный компонент u_1 , являющийся направлением, которое максимизирует проецируемую дисперсию, также является направлением, которое минимизирует среднеквадратичную ошибку MSE .

7.2.2 Наилучшая двумерная аппроксимация

Теперь нас интересует наилучшее двумерное приближение к $\mathbf{D}$. Как и раньше, предположим, что $\mathbf{D}$ уже центрирован, т.е. $\mu = \mathbf{0}$. Мы уже вычислили направление с наибольшей дисперсией, а именно u_1 , которое является собственным вектором с соответствующим наибольшим собственным значением λ_1 матрицы Σ . Теперь мы хотим найти другое направление v , которое также максимизирует проецируемую дисперсию, но ортогонально u_1 . Согласно формуле (7.9) проецируемая на v дисперсия определяется как

$$\sigma_v^2 = v^T \Sigma v$$

Далее мы потребуем, чтобы $\mathbf{v}$ был единичным вектором, ортогональным $\mathbf{u}_1$, т.е.

$$\mathbf{v}^T \mathbf{u}_1 = 0$$

$$\mathbf{v}^T \mathbf{v} = 1$$

Тогда условие оптимизации определяется как

$$\max_{\mathbf{v}} J(\mathbf{v}) = \mathbf{v}^T \Sigma \mathbf{v} - \alpha(\mathbf{v}^T \mathbf{v} - 1) - \beta(\mathbf{v}^T \mathbf{u}_1 - 0) \quad (7.17)$$

Взяв производную $J(\mathbf{v})$ по $\mathbf{v}$ и приравняв ее к нулевому вектору, получаем

$$2\Sigma\mathbf{v} - 2\alpha\mathbf{v} - \beta\mathbf{u}_1 = \mathbf{0} \quad (7.18)$$

Если умножить левую часть на $\mathbf{u}_1^T$, то получим

$$2\mathbf{u}_1^T \Sigma \mathbf{v} - 2\alpha \mathbf{u}_1^T \mathbf{v} - \beta \mathbf{u}_1^T \mathbf{u}_1 = 0$$

$$2\mathbf{v}^T \Sigma \mathbf{u}_1 - \beta = 0, \text{ откуда следует, что}$$

$$\beta = 2\mathbf{v}^T \Sigma \mathbf{u}_1 = 2\lambda_1 \mathbf{v}^T \mathbf{u}_1 = 0$$

При выводе выше мы использовали тот факт, что $\mathbf{u}_1^T \Sigma \mathbf{v} = \mathbf{v}^T \Sigma \mathbf{u}_1$, и что $\mathbf{v}$ ортогонален $\mathbf{u}_1$. Подставляя $\beta = 0$ в уравнение (7.18), получаем

$$2\Sigma\mathbf{v} - 2\alpha\mathbf{v} = \mathbf{0}$$

$$\Sigma\mathbf{v} = \alpha\mathbf{v}$$

Это означает, что $\mathbf{v}$ — это еще один собственный вектор Σ . Так же, как в уравнении (7.12), имеем $\sigma_{\mathbf{v}}^2 = \alpha$. Чтобы максимизировать дисперсию по $\mathbf{v}$, мы должны выбрать $\alpha = \lambda_2$, второе по величине собственное значение Σ , причем второй главный компонент задается соответствующим собственным вектором, т.е. $\mathbf{v} = \mathbf{u}_2$.

Общая проецируемая дисперсия

Пусть $\mathbf{U}_2$ — матрица, столбцы которой соответствуют двум главным компонентам

$$\mathbf{U}_2 = \begin{pmatrix} | & | \\ \mathbf{u}_1 & \mathbf{u}_2 \\ | & | \end{pmatrix}$$

Для точки $\mathbf{x}_i \in \mathbf{D}$ координаты в двумерном подпространстве, натянутом на $\mathbf{u}_1$ и $\mathbf{u}_2$, могут быть вычислены по формуле (7.6) следующим образом

$$\mathbf{a}_i = \mathbf{U}_2^T \mathbf{x}_i$$

Предположим, что каждая точка $\mathbf{x}_i \in \mathbb{R}^d$ в $\mathbf{D}$ была спроецирована для получения ее координат $\mathbf{a}_i \in \mathbb{R}^d$, что дало новый набор данных $\mathbf{A}$. Кроме того, поскольку предполагается, что $\mathbf{D}$ центрирован ($\boldsymbol{\mu} = \mathbf{0}$), координаты спроецированного среднего также равны нулю, т.к. $\mathbf{U}_2 \boldsymbol{\mu} = \mathbf{U}_2^T \mathbf{0} = \mathbf{0}$.

Общая дисперсия для $\mathbf{A}$ определяется как

$$\begin{aligned} \text{var}(\mathbf{A}) &= \frac{1}{n} \sum_{i=1}^n \|\mathbf{a}_i - \mathbf{0}\|^2 \\ &= \frac{1}{n} \sum_{i=1}^n (\mathbf{U}_2^T \mathbf{x}_i)^T (\mathbf{U}_2^T \mathbf{x}_i) \end{aligned}$$

$$\begin{aligned}
&= \frac{1}{n} \sum_{i=1}^n \mathbf{x}_i^T (\mathbf{U}_2 \mathbf{U}_2^T) \mathbf{x}_i \\
&= \frac{1}{n} \sum_{i=1}^n \mathbf{x}_i^T \mathbf{P}_2 \mathbf{x}_i
\end{aligned} \tag{7.19}$$

Где $\mathbf{P}_2$ – матрица ортогональной проекции [уравнение (7.8)], которая определяется как

$$\mathbf{P}_2 = \mathbf{U}_2 \mathbf{U}_2^T = \mathbf{u}_1 \mathbf{u}_1^T + \mathbf{u}_2 \mathbf{u}_2^T$$

Подставляя это в (7.19), проецируемая общая дисперсия определяется как

$$\begin{aligned}
\text{var}(\mathbf{A}) &= \frac{1}{n} \sum_{i=1}^n \mathbf{x}_i^T \mathbf{P}_2 \mathbf{x}_i \\
&= \frac{1}{n} \sum_{i=1}^n \mathbf{x}_i^T (\mathbf{u}_1 \mathbf{u}_1^T + \mathbf{u}_2 \mathbf{u}_2^T) \mathbf{x}_i \\
&= \frac{1}{n} \sum_{i=1}^n (\mathbf{u}_1^T \mathbf{x}_i)(\mathbf{x}_i^T \mathbf{u}_1) + \frac{1}{n} \sum_{i=1}^n (\mathbf{u}_2^T \mathbf{x}_i)(\mathbf{x}_i^T \mathbf{u}_2) \\
&= \mathbf{u}_1^T \mathbf{\Sigma} \mathbf{u}_1 + \mathbf{u}_2^T \mathbf{\Sigma} \mathbf{u}_2
\end{aligned} \tag{7.21}$$

Поскольку $\mathbf{u}_1$ и $\mathbf{u}_2$ – собственные вектора $\mathbf{\Sigma}$, имеем $\mathbf{\Sigma} \mathbf{u}_1 = \lambda_1 \mathbf{u}_1$ и $\mathbf{\Sigma} \mathbf{u}_2 = \lambda_2 \mathbf{u}_2$, т.е.

$$\text{var}(\mathbf{A}) = \mathbf{u}_1^T \mathbf{\Sigma} \mathbf{u}_1 + \mathbf{u}_2^T \mathbf{\Sigma} \mathbf{u}_2 = \mathbf{u}_1^T \lambda_1 \mathbf{u}_1 + \mathbf{u}_2^T \lambda_2 \mathbf{u}_2 = \lambda_1 + \lambda_2 \tag{7.22}$$

Таким образом, сумма собственных значений представляет собой общую дисперсию проецируемых точек, а первые два главных компонента максимизируют эту дисперсию.

Среднеквадратичная ошибка

Теперь мы продемонстрируем, что первые два главных компонента также минимизируют среднеквадратичную ошибку. Среднеквадратичная ошибка задается как

$$\begin{aligned}
\text{MSE} &= \frac{1}{n} \sum_{i=1}^n \|\mathbf{x}_i - \mathbf{x}'_i\|^2 \\
&= \frac{1}{n} \sum_{i=1}^n (\|\mathbf{x}_i\|^2 - 2\mathbf{x}_i^T \mathbf{x}'_i + (\mathbf{x}'_i)^T \mathbf{x}'_i), \text{ используя (7.14)} \\
&= \text{var}(\mathbf{D}) + \frac{1}{n} \sum_{i=1}^n (-2\mathbf{x}_i^T \mathbf{P}_2 \mathbf{x}_i + (\mathbf{P}_2 \mathbf{x}_i)^T \mathbf{P}_2 \mathbf{x}_i), \text{ используя (7.7), что } \mathbf{x}'_i = \mathbf{P}_2 \mathbf{x}_i \\
&= \text{var}(\mathbf{D}) - \frac{1}{n} \sum_{i=1}^n (\mathbf{x}_i^T \mathbf{P}_2 \mathbf{x}_i) \\
&= \text{var}(\mathbf{D}) - \text{var}(\mathbf{A}), \text{ используя (7.20)}
\end{aligned} \tag{7.23}$$

Таким образом, MSE минимизируется именно тогда, когда общая проецируемая дисперсия $\text{var}(\mathbf{A})$ максимальна. Из (7.22) имеем

$$\text{MSE} = \text{var}(\mathbf{D}) - \lambda_1 - \lambda_2$$

7.2.3 Наилучшая r-мерная аппроксимация

Теперь нас интересует наилучшая r-мерная аппроксимация к $\mathbf{D}$, где $2 < r \leq d$. Предположим, что мы уже вычислили первые $j - 1$ главных компонент или собственных векторов, $\mathbf{u}_1, \mathbf{u}_2, \dots, \mathbf{u}_{j-1}$,

соответствующих $j - 1$ наибольшим собственным значениям матрицы Σ , для $1 \leq j \leq r$. Чтобы вычислить j -й новый базисный вектор $\mathbf{v}$, мы должны убедиться, что он нормализован до единичной длины, т.е. $\mathbf{v}^T \mathbf{v} = 1$, и ортогонален всем предыдущим компонентам $\mathbf{u}_i$, т.е. $\mathbf{u}_i^T \mathbf{v} = 0$, для $1 \leq i < j$. Как и до этого, проецируемая дисперсия по $\mathbf{v}$ задается как

$$\sigma_{\mathbf{v}}^2 = \mathbf{v}^T \Sigma \mathbf{v}$$

В сочетании с ограничениями на $\mathbf{v}$ это приводит к следующей задаче максимизации с множителями Лагранжа:

$$\max_{\mathbf{v}} J(\mathbf{v}) = \mathbf{v}^T \Sigma \mathbf{v} - \alpha(\mathbf{v}^T \mathbf{v} - 1) - \sum_{i=1}^{j-1} \beta_i(\mathbf{u}_i^T \mathbf{v} - 0)$$

Взяв производную $J(\mathbf{v})$ по $\mathbf{v}$ и приравняв ее к нулевому вектору, получаем

$$2\Sigma\mathbf{v} - 2\alpha\mathbf{v} - \sum_{i=1}^{j-1} \beta_i \mathbf{u}_i = \mathbf{0} \quad (7.24)$$

Если умножить левую часть на $\mathbf{u}_k^T$ для $1 \leq k < j$, то получим

$$2\mathbf{u}_k^T \Sigma \mathbf{v} - 2\alpha \mathbf{u}_k^T \mathbf{v} - \beta_k \mathbf{u}_k^T \mathbf{u}_k - \sum_{\substack{i=1 \\ i \neq k}}^{j-1} \beta_i \mathbf{u}_k^T \mathbf{u}_i = 0$$

$$2\mathbf{v}^T \Sigma \mathbf{u}_k - \beta_k = 0$$

$$\beta_k = 2\mathbf{v}^T \lambda_k \mathbf{u}_k = 2\lambda_k \mathbf{v}^T \mathbf{u}_k = 0$$

где мы использовали тот факт, что $\Sigma \mathbf{u}_k = \lambda_k \mathbf{u}_k$, т.к. $\mathbf{u}_k$ является собственным вектором, которому соответствует k -е наибольшее собственное значение λ_k матрицы Σ . Таким образом, мы находим, что $\beta_i = 0$ для всех $i < j$ в уравнении (7.24), откуда следует

$$\Sigma \mathbf{v} = \alpha \mathbf{v}$$

Чтобы максимизировать дисперсию по $\mathbf{v}$, мы устанавливаем $\alpha = \lambda_j$, j -е наибольшее собственное значение матрицы Σ , где $\mathbf{v} = \mathbf{u}_j$ задает j -ю главную компоненту.

Таким образом, чтобы найти наилучшее r -мерное приближение к $\mathbf{D}$, мы вычисляем собственные значения матрицы Σ . Поскольку Σ является положительно полуопределенной, все ее собственные числа должны быть неотрицательными, и мы можем отсортировать их в порядке убывания следующим образом:

$$\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_r \geq \lambda_{r+1} \dots \geq \lambda_d \geq 0$$

Затем мы выбираем r наибольших собственных значений и соответствующие им собственные векторы, чтобы сформировать лучшее r -мерное приближение.

Общая проецируемая дисперсия

Пусть $\mathbf{U}_r$ – r -мерная матрица базисных векторов

$$\mathbf{U}_r = \begin{pmatrix} | & | & \dots & | \\ \mathbf{u}_1 & \mathbf{u}_2 & \dots & \mathbf{u}_r \\ | & | & \dots & | \end{pmatrix}$$

с матрицей проекции, заданной как

$$\mathbf{P}_r = \mathbf{U}_r \mathbf{U}_r^T = \sum_{i=1}^r \mathbf{u}_i \mathbf{u}_i^T$$

Пусть $\mathbf{A}$ обозначает набор данных, сформированный координатами проецируемых точек в r -мерном подпространстве, т.е. $\mathbf{a}_i = \mathbf{U}_r^T \mathbf{x}_i$, и пусть $\mathbf{x}'_i = \mathbf{P}_r \mathbf{x}_i$ обозначает спроецированную точку в исходном d -мерном пространстве. Из (7.19), (7.21) и (7.22) проецируемая дисперсия определяется как

$$\text{var}(\mathbf{A}) = \frac{1}{n} \sum_{i=1}^n \mathbf{x}_i^T \mathbf{P}_r \mathbf{x}_i = \sum_{i=1}^r \mathbf{u}_i^T \mathbf{\Sigma} \mathbf{u}_i = \sum_{i=1}^r \lambda_i$$

Таким образом, общая проецируемая дисперсия – это просто сумма r наибольших собственных значений матрицы $\mathbf{\Sigma}$.

Среднеквадратичная ошибка

На основании вывода для уравнения (7.23), среднеквадратичная ошибка в r -измерениях может быть записана как

$$\begin{aligned} \text{MSE} &= \frac{1}{n} \sum_{i=1}^n \|\mathbf{x}_i - \mathbf{x}'_i\|^2 \\ &= \text{var}(\mathbf{D}) - \text{var}(\mathbf{A}) \\ &= \text{var}(\mathbf{D}) - \sum_{i=1}^r \mathbf{u}_i^T \mathbf{\Sigma} \mathbf{u}_i \\ &= \text{var}(\mathbf{D}) - \sum_{i=1}^r \lambda_i \end{aligned}$$

Первые r главных компонент максимизируют проецируемую дисперсию $\text{var}(\mathbf{A})$ и, таким образом, они также минимизируют MSE.

Общая дисперсия

Обратите внимание, что общая дисперсия $\mathbf{D}$ инвариантна к изменению базисных векторов. Следовательно, имеем следующее тождество:

$$\text{var}(\mathbf{D}) = \sum_{i=1}^d \sigma_i^2 = \sum_{i=1}^d \lambda_i$$

Выбор размерности

Часто мы можем не знать, сколько измерений r использовать для хорошей аппроксимации. Одним из критериев выбора r является вычисление доли общей дисперсии, охваченной первыми r главными компонентами

$$f(r) = \frac{\lambda_1 + \lambda_2 + \dots + \lambda_r}{\lambda_1 + \lambda_2 + \dots + \lambda_d} = \frac{\sum_{i=1}^r \lambda_i}{\sum_{i=1}^d \lambda_i} = \frac{\sum_{i=1}^r \lambda_i}{\text{var}(\mathbf{D})} \quad (7.25)$$

Учитывая определенный желаемый порог дисперсии, скажем α , начиная с первого главного компонента, мы продолжаем добавлять дополнительные компоненты и останавливаемся на наименьшем значении r , для которого $f(r) \geq \alpha$. Другими словами, мы выбираем наименьшее количество измерений так, что подпространство, охватываемое этими r измерениями, захватывает

по крайней мере долю α общей дисперсии. На практике α обычно устанавливается равным 0.9 или выше, т.ч. уменьшенный набор данных фиксирует не менее 90% общей дисперсии.

Алгоритм 7.1 демонстрирует псевдокод для алгоритма метода главных компонент PCA. Центрируются входные данные $\mathbf{D} \in \mathbb{R}^{n \times d}$ вычитанием среднего значения из каждой точки. Затем вычисляются собственные векторы и собственные значения ковариационной матрицы $\mathbf{\Sigma}$. Учитывая желаемый порог дисперсии α , выбирается наименьший набор измерений r , который охватывает по меньшей мере долю α общей дисперсии. В заключении, вычисляются координаты каждой точки в новом r -мерном подпространстве главных компонент, чтобы получить новую матрицу данных $\mathbf{A} \in \mathbb{R}^{n \times r}$.

АЛГОРИТМ 7.1. Метод главных компонент PCA

PCA($\mathbf{D}, \alpha$):

1. $\boldsymbol{\mu} = \frac{1}{n} \sum_{i=1}^n \mathbf{x}_i$ // вычисление среднего
 2. $\mathbf{Z} = \mathbf{D} - \mathbf{1} \cdot \boldsymbol{\mu}^T$ // центрирование данных
 3. $\mathbf{\Sigma} = \frac{1}{n} (\mathbf{Z}^T \mathbf{Z})$ // вычисление ковариационной матрицы
 4. $(\lambda_1, \lambda_2, \dots, \lambda_d) = \text{eigenvalues}(\mathbf{\Sigma})$ // вычисление собственных значений
 5. $\mathbf{U} = (\mathbf{u}_1 \quad \mathbf{u}_2 \quad \dots \quad \mathbf{u}_d) = \text{eigenvectors}(\mathbf{\Sigma})$ // собственных векторов
 6. $f(r) = \frac{\sum_{i=1}^r \lambda_i}{\sum_{i=1}^d \lambda_i}$, for all $r = 1, 2, \dots, d$ // доля от общей дисперсии
 7. Choose smallest r so that $f(r) \geq \alpha$ // выбор размерности
 8. $\mathbf{U}_r = (\mathbf{u}_1 \quad \mathbf{u}_2 \quad \dots \quad \mathbf{u}_r)$ // уменьшенный базис
 9. $\mathbf{A} = \{\mathbf{a}_i | \mathbf{a}_i = \mathbf{U}_r^T \mathbf{x}_i, \text{ for } i = 1, \dots, n\}$ // данные уменьшенной размерности
-

7.2.4 Геометрия метода главных компонент

Геометрически, когда $r = d$, PCA соответствует ортогональному изменению базиса, так что общая дисперсия фиксируется суммой дисперсий по каждому из главных направлений $\mathbf{u}_1, \mathbf{u}_2, \dots, \mathbf{u}_d$, и далее, все ковариации равны нулю. Это можно увидеть, посмотрев на коллективное действие полного набора главных компонент, которые могут быть организованы в ортогональную матрицу $d \times d$.

$$\mathbf{U} = \begin{pmatrix} | & | & & | \\ \mathbf{u}_1 & \mathbf{u}_2 & \dots & \mathbf{u}_d \\ | & | & & | \end{pmatrix}$$

$$\mathbf{U}^{-1} = \mathbf{U}^T$$

Каждый главный компонент $\mathbf{u}_i$ соответствует собственному вектору ковариационной матрицы $\mathbf{\Sigma}$, т.е.

$$\mathbf{\Sigma} \mathbf{u}_i = \lambda_i \mathbf{u}_i \text{ for all } 1 \leq i \leq d$$

Который можно компактно записать в матричных обозначениях следующим образом:

$$\mathbf{\Sigma} \begin{pmatrix} | & | & & | \\ \mathbf{u}_1 & \mathbf{u}_2 & \dots & \mathbf{u}_d \\ | & | & & | \end{pmatrix} = \begin{pmatrix} | & | & & | \\ \lambda_1 \mathbf{u}_1 & \lambda_2 \mathbf{u}_2 & \dots & \lambda_d \mathbf{u}_d \\ | & | & & | \end{pmatrix}$$

$$\mathbf{\Sigma} \mathbf{U} = \mathbf{U} \begin{pmatrix} \lambda_1 & 0 & \dots & 0 \\ 0 & \lambda_2 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & \lambda_d \end{pmatrix}$$

$$\mathbf{\Sigma} \mathbf{U} = \mathbf{U} \mathbf{\Lambda} \tag{7.26}$$

Если мы умножим уравнение (7.26) слева на $\mathbf{U}^{-1} = \mathbf{U}^T$, то получим

$$\mathbf{U}^T \mathbf{\Sigma} \mathbf{U} = \mathbf{U}^T \mathbf{U} \mathbf{\Lambda} = \mathbf{\Lambda} = \begin{pmatrix} \lambda_1 & 0 & \cdots & 0 \\ 0 & \lambda_2 & \cdots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \cdots & \lambda_d \end{pmatrix}$$

Это означает, что если мы изменим базис на $\mathbf{U}$, то ковариационная матрица $\mathbf{\Sigma}$ изменится на аналогичную матрицу $\mathbf{\Lambda}$, которая фактически является ковариационной матрицей в новом базисе. Тот факт, что $\mathbf{\Lambda}$ является диагональной, подтверждает, что после изменения базиса все ковариации исчезают, и остается только дисперсия по каждой из главных компонент, причем дисперсия по каждому новому направлению $\mathbf{u}_i$ задается соответствующим собственным значением λ_i .

Стоит отметить, что в новом базисе уравнение

$$\mathbf{x}^T \mathbf{\Sigma}^{-1} \mathbf{x} = 1 \quad (7.27)$$

Определяет d-мерный эллипсоид (или гиперэллипс). Собственные векторы $\mathbf{u}_i$ матрицы $\mathbf{\Sigma}$, т.е. главные компоненты, являются направлениями главных осей эллипсоида. Квадратные корни собственных значений, т.е. $\sqrt{\lambda_i}$, являются длинами полуосей эллипса.

Умножая (7.26) справа на $\mathbf{U}^{-1} = \mathbf{U}^T$, имеем

$$\mathbf{\Sigma} = \mathbf{U} \mathbf{\Lambda} \mathbf{U}^T \quad (7.28)$$

Предполагая, что $\mathbf{\Sigma}$ обратимая или невырожденная, имеем

$$\mathbf{\Sigma}^{-1} = (\mathbf{U} \mathbf{\Lambda} \mathbf{U}^T)^{-1} = (\mathbf{U}^{-1})^T \mathbf{\Lambda}^{-1} \mathbf{U}^{-1} = \mathbf{U} \mathbf{\Lambda}^{-1} \mathbf{U}^T$$

Где

$$\mathbf{\Lambda}^{-1} = \begin{pmatrix} \frac{1}{\lambda_1} & 0 & \cdots & 0 \\ 0 & \frac{1}{\lambda_2} & \cdots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \cdots & \frac{1}{\lambda_d} \end{pmatrix}$$

Подставив $\mathbf{\Sigma}^{-1}$ в уравнение (7.27) и используя тот факт, что $\mathbf{x} = \mathbf{U} \mathbf{a}$ из (7.2), где $\mathbf{a} = (a_1, a_2, \dots, a_d)^T$ представляет координаты $\mathbf{x}$ в новом базисе, получаем

$$\mathbf{x}^T \mathbf{\Sigma}^{-1} \mathbf{x} = 1$$

$$(\mathbf{a}^T \mathbf{U}^T) \mathbf{U} \mathbf{\Lambda}^{-1} \mathbf{U}^T (\mathbf{U} \mathbf{a}) = 1$$

$$\mathbf{a}^T \mathbf{\Lambda}^{-1} \mathbf{a} = 1$$

$$\sum_{i=1}^d \frac{a_i^2}{\lambda_i} = 1$$

что и есть уравнение эллипса с центром в $\mathbf{0}$ и длиной полуосей $\sqrt{\lambda_i}$. Таким образом $\mathbf{x}^T \mathbf{\Sigma}^{-1} \mathbf{x} = 1$ или эквивалентно $\mathbf{a}^T \mathbf{\Lambda}^{-1} \mathbf{a} = 1$ в новом базисе главных компонент определяет d-мерный эллипсоид, где длины полуосей равны стандартным отклонениям (квадратный корень из дисперсии $\sqrt{\lambda_i}$) вдоль каждой оси. Точно так же уравнение $\mathbf{x}^T \mathbf{\Sigma}^{-1} \mathbf{x} = s$ или эквивалентное $\mathbf{a}^T \mathbf{\Lambda}^{-1} \mathbf{a} = s$, для различных значений скалярной величины s , представляет концентрические эллипсоиды.

7.3 Ядерный анализ главных компонент

Анализ главных компонент может быть расширен, чтобы искать нелинейные “направления” в данных с использованием ядерных методов. Ядерный PCA ищет направления наибольшей дисперсии в пространстве признаков вместо входного пространства. То есть, вместо поиска линейных комбинаций во входном пространстве, алгоритм ищет линейные комбинации в пространстве признаков, полученном из нелинейного преобразования входного пространства. Таким образом, линейные главные компоненты пространства признаков соответствуют нелинейным направлениям во входном пространстве. Как мы увидим, с использованием *ядерного трюка* все операции могут быть выполнены в терминах входного пространства, без необходимости преобразования данных в пространстве признаков.

Пусть ϕ – отображение из входного пространства в пространство признаков. Каждой точке в пространстве признаков соответствует образ $\phi(\mathbf{x}_i)$ точки $\mathbf{x}_i$ во входном пространстве. Во входном пространстве первая главная компонента показывает направление наибольшей дисперсии; её собственный вектор совпадает с наибольшим собственным значением в матрице ковариантности. Таким же образом, через поиск собственного вектора, соответствующий наибольшему собственному значению в матрице ковариантности в пространстве признаков, можно найти первую главную компоненту в пространстве признаков $\mathbf{u}_1$ (с $\mathbf{u}_1^T \mathbf{u}_1 = 1$):

$$\Sigma_\phi = \lambda_1 \mathbf{u}_1 \quad (7.29)$$

где Σ_ϕ – матрица ковариантности в пространстве признаков, определенная как

$$\Sigma_\phi = \frac{1}{n} \sum_{i=1}^n \phi(\mathbf{x}_i) \phi(\mathbf{x}_i)^T \quad (7.30)$$

Здесь предполагается, что точки центрированы, т.е. $\phi(\mathbf{x}_i) = \phi(\mathbf{x}_i) - \boldsymbol{\mu}_\phi$, где $\boldsymbol{\mu}_\phi$ – среднее в пространстве признаков.

Подставляя (7.30) в (7.29) получаем:

$$\begin{aligned} \left(\frac{1}{n} \sum_{i=1}^n \phi(\mathbf{x}_i) \phi(\mathbf{x}_i)^T \right) \mathbf{u}_1 &= \lambda_1 \mathbf{u}_1 \\ \frac{1}{n} \sum_{i=1}^n \phi(\mathbf{x}_i) (\phi(\mathbf{x}_i)^T \mathbf{u}_1) &= \lambda_1 \mathbf{u}_1 \\ \sum_{i=1}^n \left(\frac{\phi(\mathbf{x}_i)^T \mathbf{u}_1}{n \lambda_1} \right) \phi(\mathbf{x}_i) &= \mathbf{u}_1 \end{aligned} \quad (7.31)$$

$$\sum_{i=1}^n c_i \phi(\mathbf{x}_i) = \mathbf{u}_1 \quad (7.32)$$

где $c_i = \frac{\phi(\mathbf{x}_i)^T \mathbf{u}_1}{n \lambda_1}$ – скалярная величина. Из (7.32) видно, что лучшее направление в пространстве признаков, $\mathbf{u}_1$ есть линейная комбинация преобразованных точек, где скаляры c_i показывают важность каждой точки в направлении наибольшей дисперсии.

Теперь можно подставить (7.32) назад в (7.31) и получить

$$\begin{aligned}
\left(\frac{1}{n} \sum_{i=1}^n \phi(\mathbf{x}_i) \phi(\mathbf{x}_i)^T \right) \left(\sum_{i=1}^n c_i \phi(\mathbf{x}_i) \right) &= \lambda_1 \sum_{i=1}^n c_i \phi(\mathbf{x}_i) \\
\frac{1}{n} \sum_{i=1}^n \sum_{j=1}^n c_j \phi(\mathbf{x}_i) \phi(\mathbf{x}_i)^T \phi(\mathbf{x}_j) &= \lambda_1 \sum_{i=1}^n c_i \phi(\mathbf{x}_i) \\
\sum_{i=1}^n \left(\phi(\mathbf{x}_i) \sum_{j=1}^n c_j \phi(\mathbf{x}_i)^T \phi(\mathbf{x}_j) \right) &= n \lambda_1 \sum_{i=1}^n c_i \phi(\mathbf{x}_i)
\end{aligned}$$

В предыдущем уравнении можно заменить скалярное произведение в пространстве признаков, $(\mathbf{x}_i)^T \phi(\mathbf{x}_j)$, соответствующей ядерной функцией во входном пространстве, а именно $K(\mathbf{x}_i, \mathbf{x}_j)$

$$\sum_{i=1}^n \left(\phi(\mathbf{x}_i) \sum_{j=1}^n c_j K(\mathbf{x}_i, \mathbf{x}_j) \right) = n \lambda_1 \sum_{i=1}^n c_i \phi(\mathbf{x}_i) \quad (7.33)$$

Здесь предполагается, что точки в пространстве признаков уже центрированы, т.е. ядерная матрица была центрирована через уравнение (5.14):

$$\mathbf{K} = \left(\mathbf{I} - \frac{1}{n} \mathbf{1}_{n \times n} \right) \mathbf{K} \left(\mathbf{I} - \frac{1}{n} \mathbf{1}_{n \times n} \right)$$

где $\mathbf{I}$ – единичная матрица, у которой элементы главной диагонали 1, а остальные 0, и $\mathbf{1}_{n \times n}$ – матрица $n \times n$, у которой все элементы – 1.

Мы заменили одно скалярное произведение ядерной функцией. Чтобы убедиться, что все вычисления в пространстве признаков выполняются только в терминах скалярного произведения, возьмем любую точку $\phi(\mathbf{x}_k)$ и умножим (7.33) на неё с обеих сторон чтобы получить

$$\begin{aligned}
\sum_{i=1}^n \left(\phi(\mathbf{x}_k) \phi(\mathbf{x}_i) \sum_{j=1}^n c_j K(\mathbf{x}_i, \mathbf{x}_j) \right) &= n \lambda_1 \sum_{i=1}^n c_i \phi(\mathbf{x}_i) \phi(\mathbf{x}_k) \\
\sum_{i=1}^n \left(K(\mathbf{x}_k, \mathbf{x}_i) \sum_{j=1}^n c_j K(\mathbf{x}_i, \mathbf{x}_j) \right) &= n \lambda_1 \sum_{i=1}^n c_i K(\mathbf{x}_k, \mathbf{x}_i)
\end{aligned} \quad (7.34)$$

Затем, обозначим $\mathbf{K}_i$ строку i центрированной ядерной матрицы, записанную как вектор-столбец:

$$\mathbf{K}_i = (K(\mathbf{x}_i, \mathbf{x}_1) \quad K(\mathbf{x}_i, \mathbf{x}_2) \quad \dots \quad K(\mathbf{x}_i, \mathbf{x}_n))^T$$

Пусть $\mathbf{c}$ обозначает вектор-столбец весов

$$\mathbf{c} = (c_1 \quad c_2 \quad \dots \quad c_n)^T$$

Теперь можно записать (7.34) как:

$$\sum_{i=1}^n K(\mathbf{x}_k, \mathbf{x}_i) \mathbf{K}_i^T \mathbf{c} = n \lambda_1 \mathbf{K}_k^T \mathbf{c}$$

На самом деле, поскольку можно выбрать любые из n точек $\phi(\mathbf{x}_k)$ в пространстве признаков, чтобы получить (7.34), имеется набор уравнений:

$$\begin{aligned}
\sum_{i=1}^n K(\mathbf{x}_1, \mathbf{x}_i) \mathbf{K}_i^T \mathbf{c} &= n\lambda_1 \mathbf{K}_1^T \mathbf{c} \\
\sum_{i=1}^n K(\mathbf{x}_2, \mathbf{x}_i) \mathbf{K}_i^T \mathbf{c} &= n\lambda_1 \mathbf{K}_2^T \mathbf{c} \\
&\vdots \\
\sum_{i=1}^n K(\mathbf{x}_n, \mathbf{x}_i) \mathbf{K}_i^T \mathbf{c} &= \sum_{i=1}^n K(\mathbf{x}_k, \mathbf{x}_i) \mathbf{K}_n^T \mathbf{c}
\end{aligned}$$

Можно представить эти n уравнений как:

$$\mathbf{K}^2 \mathbf{c} = n\lambda_1 \mathbf{K} \mathbf{c}$$

где $\mathbf{K}$ – центрированная ядерная матрица. Умножением на $\mathbf{K}^{-1}$ получаем:

$$\begin{aligned}
\mathbf{K}^{-1} \mathbf{K}^2 \mathbf{c} &= n\lambda_1 \mathbf{K}^{-1} \mathbf{K} \mathbf{c} \\
\mathbf{K} \mathbf{c} &= n\lambda_1 \mathbf{c} \\
\mathbf{K} \mathbf{c} &= \eta_1 \mathbf{c}
\end{aligned} \tag{7.35}$$

где $\eta_1 = n\lambda_1$. Таким образом, вектор весов $\mathbf{c}$ есть собственный вектор, соответствующий наибольшему собственному значению η_1 ядерной матрицы $\mathbf{K}$.

Как только найдено $\mathbf{c}$, её можно подставить обратно в (7.32), чтобы получить первую ядерную основную компоненту $\mathbf{u}_1$. Единственное ограничение – надо нормализовать $\mathbf{u}_1$ до единичного вектора следующим образом:

$$\begin{aligned}
\mathbf{u}_1^T \mathbf{u}_1 &= 1 \\
\sum_{i=1}^n \sum_{j=1}^n c_i c_j \phi(\mathbf{x}_i)^T \phi(\mathbf{x}_j) &= 1 \\
\mathbf{c}^T \mathbf{K} \mathbf{c} &= 1
\end{aligned}$$

Из (7.35) имеем:

$$\begin{aligned}
\mathbf{c}^T (\eta_1 \mathbf{c}) &= 1 \\
\eta_1 \mathbf{c}^T \mathbf{c} &= 1 \\
\|\mathbf{c}\|^2 &= \frac{1}{\eta_1}
\end{aligned}$$

Однако, поскольку $\mathbf{c}$ – собственный вектор $\mathbf{K}$, он будет иметь единичную норму. То есть, чтобы удостовериться, что $\mathbf{u}_1$ – единичный вектор, надо масштабировать вектор весов $\mathbf{c}$ так, что его норма $\|\mathbf{c}\| = \sqrt{\frac{1}{\eta_1}}$, что можно получить, домножив $\mathbf{c}$ на $\sqrt{\frac{1}{\eta_1}}$.

В общем случае, поскольку мы не отображаем входные точки в пространство признаков через ϕ , невозможно непосредственно вычислить главное направление в терминах $\phi(\mathbf{x}_i)$, как показано в (7.32). Однако, здесь важно, что любую точку $\phi(\mathbf{x})$ можно отобразить на главное направление $\mathbf{u}_1$ следующим образом:

$$\mathbf{u}_1^T \phi(\mathbf{x}) = \sum_{i=1}^n c_i \phi(\mathbf{x}_i)^T \phi(\mathbf{x}) = \sum_{i=1}^n c_i K(\mathbf{x}_i, \mathbf{x})$$

что требует только ядерные операции. Когда $\mathbf{x} = \mathbf{x}_i$ – одна из входных точек, проекция $\phi(\mathbf{x}_i)$ на главную компоненту $\mathbf{u}_1$ может быть записана как скалярное произведение

$$\mathbf{a}_i = \mathbf{u}_1^T \phi(\mathbf{x}_i) = \mathbf{K}_i^T \mathbf{c} \tag{7.36}$$

где $\mathbf{K}_i$ – вектор-столбец, соответствующий i -й строке в ядерной матрице. Таким образом, мы показали, что все вычисления, или для решения главной компоненты, или для проекции точек, могут быть выполнены с использованием одной только ядерной функции. Затем, мы можем получить дополнительные главные компоненты путём решения 7.35 для других собственных значений и собственных векторов. Другими словами, отсортировав собственные значения $\mathbf{K}$ в порядке убывания $\eta_1 \geq \eta_2 \geq \dots \geq \eta_n \geq 0$, можно получить j -ю главную компоненту как соответствующий собственный вектор $\mathbf{c}_j$, который должен быть нормализован по норме $\|\mathbf{c}_j\| = \sqrt{\frac{1}{n}}$, при условии $\eta_j > 0$. И поскольку $\eta_j = n\lambda_j$, дисперсия по j -й главной компоненте определяется как $\lambda_j = \frac{\eta_j}{n}$. Алгоритм 7.2 приводит псевдокод для метода ядерного PCA.

Алгоритм 7.2. Ядерный анализ главных компонент

KernelPCA($\mathbf{D}, \mathbf{K}, \alpha$):

1. $\mathbf{K} = \{K(\mathbf{x}_i, \mathbf{x}_j)\}_{\{i,j=1\dots n\}}$ // вычислить ядерную матрицу $n \times n$
2. $\mathbf{K} = \left(\mathbf{I} - \frac{1}{n}\mathbf{1}_{n \times n}\right) \mathbf{K} \left(\mathbf{I} - \frac{1}{n}\mathbf{1}_{n \times n}\right)$ // центрировать ядерную матрицу
3. $(\eta_1, \eta_2, \dots, \eta_d) = \text{eigenvalues}(\mathbf{K})$ // собственные значения
4. $(\mathbf{c}_1 \ \mathbf{c}_2 \ \dots \ \mathbf{c}_n) = \text{eigenvectors}(\mathbf{K})$ // собственные вектора
5. $\lambda_i = \frac{\eta_i}{n}, \forall i = 1, \dots, n$ // вычислить дисперсию каждой компоненты
6. $\mathbf{c}_i = \sqrt{\frac{1}{\eta_i}} \cdot \mathbf{c}_i, \forall i = 1, \dots, n$ // убедиться, что $\mathbf{u}_i^T \mathbf{u}_i = 1$
7. $f(r) = \frac{\sum_{i=1}^r \lambda_i}{\sum_{i=1}^d \lambda_i}, \forall r = 1, 2, \dots, d$ // часть общей дисперсии
8. Выбрать наименьшую r такую, что $f(r) \geq \alpha$ // выбор размерности
9. $\mathbf{C}_r = (\mathbf{c}_1 \ \mathbf{c}_2 \ \dots \ \mathbf{c}_r)$ // Уменьшенный базис
10. $\mathbf{A} = \{\mathbf{a}_i | \mathbf{a}_i = \mathbf{C}_r^T \mathbf{K}_i, \forall i = 1, \dots, n\}$ // Данные меньшей размерности

7.4. Сингулярное разложение

Анализ главных компонент является частным случаем более общего метода разложения матриц, называемого сингулярным разложением (Singular Value Decomposition - SVD). Мы знаем, что по ур-ю 7.28 PCA дает следующее разложение ковариационной матрицы:

$$\Sigma = U \Lambda U^T \quad 7.37$$

где ковариационная матрица была представлена через ортогональную матрицу U , содержащую ее собственные вектор и диагональную матрицу Λ , содержащую ее собственные числа (отсортированные в порядке убывания). SVD обобщает приведенную выше факторизацию для любой матрицы. В частности, для матрицы данных D размера $n \times d$ с n точками и d столбцами SVD факторизует D следующим образом:

$$D = L \Delta R^T \quad 7.38$$

где L - ортогональная матрица размера $n \times n$, R - ортогональная матрица размера $d \times d$, а Δ - «диагональная» матрица $n \times d$. Столбцы L называются левыми сингулярными векторами, а столбцы R (или строки R^T) называются правыми сингулярными векторами. Матрица Δ определяется как

$$\Delta(i, j) = \begin{cases} \delta_i, & \text{если } i = j \\ 0, & \text{если } i \neq j \end{cases}$$

где $i = 1, \dots, n$ и $j = 1, \dots, d$. Элементы $\Delta(i, i) = \delta_i$ вдоль главной диагонали Δ называются сингулярными числами D , и все они неотрицательны. Если ранг D равен $r \leq \min(n, d)$, то существует только r ненулевых особых значений, которые, как мы предполагаем, упорядочены следующим образом:

$$\delta_1 \geq \delta_2 \geq \dots \geq \delta_r \geq 0$$

Можно отбросить эти левые и правые сингулярные вектора, которые соответствуют нулевым сингулярным значениям, чтобы получить сокращенный SVD в следующем виде:

$$D = L_r \Delta_r R_r^T \quad 7.39$$

где L_r - матрица левых сингулярных векторов размера $n \times r$, R_r - матрица правых сингулярных векторов размера $d \times r$, а Δ_r - диагональная матрица размера $r \times r$, содержащая положительные сингулярные вектора. Уменьшенный SVD приводит непосредственно к спектральному разложению D , заданному как

$$\begin{aligned} D = L_r \Delta_r R_r^T &= \begin{pmatrix} | & | & \dots & | \\ l_1 & l_2 & \dots & l_r \\ | & | & \dots & | \end{pmatrix} \begin{pmatrix} \delta_1 & 0 & \dots & 0 \\ 0 & \delta_2 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & \delta_r \end{pmatrix} \begin{pmatrix} - & r_1^T & - \\ - & r_2^T & - \\ - & \vdots & - \\ - & r_r^T & - \end{pmatrix} \\ &= \delta_1 l_1 r_1^T + \delta_2 l_2 r_2^T + \dots + \delta_r l_r r_r^T = \sum_{i=1}^r \delta_i l_i r_i^T \end{aligned}$$

Спектральное разложение D представляет собой сумму матриц первого ранга вида $\delta_i l_i r_i^T$. Выбрав q наибольших сингулярных значений $\delta_1, \delta_2, \dots, \delta_q$ и соответствующие левый и правый сингулярные вектора, мы получаем наилучшее приближение ранга q к исходной матрице D . То есть, если D_q - матрица, определенная как

$$D_q = \sum_{i=1}^r \delta_i l_i r_i^T$$

то можно показать, что D_q - это матрица ранга q , которая минимизирует следующее выражение

$$\|D - D_q\|_F$$

где $\|A\|_F$ - норма Фробениуса матрицы A размера $n \times d$, определяемая как

$$\|A\|_F = \sqrt{\sum_{i=1}^n \sum_{j=1}^d A(i, j)^2}$$

7.4.1. Геометрия сингулярного разложения

В общем, любая матрица D размера $n \times d$ представляет собой линейное преобразование $D: \mathbb{R}^d \rightarrow \mathbb{R}^n$ из пространства d -мерных векторов в пространство n -мерных векторов, поскольку для любого $x \in \mathbb{R}^d$ существует $y \in \mathbb{R}^n$ такое, что

$$Dx = y$$

Множество всех векторов $y \in \mathbb{R}^n$ таких, что $Dx = y$ по всем возможным $x \in \mathbb{R}^d$, называется пространством столбцов D , а множество всех векторов $x \in \mathbb{R}^d$, таких что $D^T y = x$ по всем $y \in \mathbb{R}^n$, называется пространством строк D , что эквивалентно пространству столбцов D^T . Другими словами, пространство столбцов D — это набор всех векторов, которые могут быть получены как линейные комбинации столбцов D , а пространство строк D - это набор всех векторов, которые могут быть получены как линейные комбинации строк D (или столбцов D^T). Также обратите внимание, что множество всех векторов $x \in \mathbb{R}^d$, таких что $Dx = 0$, называется нулевым пространством D , и, наконец, множество всех векторов $y \in \mathbb{R}^n$, таких что $D^T y = 0$, называется левым нулевым пространством D .

Одним из основных свойств сингулярного разложения является то, что оно дает основу для каждого из четырех фундаментальных пространств, связанных с матрицей D . Если D имеет ранг r , это означает, что у него есть только r независимых столбцов и r независимых строк. Таким образом, r левых сингулярных векторов $l_1, l_2, \dots, l_r$, соответствующие r ненулевым сингулярным значениям D в ур-и (7.38), представляют собой основу для пространства столбцов D . Остальные $n - r$ левых сингулярных векторов $l_{r+1}, l_{r+2}, \dots, l_n$, представляют собой основу для левого нулевого пространства D . Для пространства строк r правых сингулярных векторов $r_1, r_2, \dots, r_r$, соответствующие r ненулевым сингулярным значениям, представляют собой базис для пространства строк D , а оставшиеся $d - r$ правых сингулярных векторов r_j ($j = r + 1, \dots, d$) представляют собой базис для нулевого пространства D .

Рассмотрим сокращенное представление сингулярного разложения из ур-я (7.39). Умножив обе части уравнения справа на R_r и зная, что $R_r^T R_r = I_r$, где I_r - единичная матрица размера $r \times r$, имеем

$$DR_r = L_r \Delta_r R_r^T R_r$$

$$DR_r = L_r \Delta_r$$

$$DR_r = L_r \begin{pmatrix} \delta_1 & 0 & \dots & 0 \\ 0 & \delta_2 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & \delta_r \end{pmatrix}$$

$$D \begin{pmatrix} | & | & \dots & | \\ r_1 & r_2 & \dots & r_r \\ | & | & \dots & | \end{pmatrix} = \begin{pmatrix} | & | & \dots & | \\ \delta_1 l_1 & \delta_2 l_2 & \dots & \delta_r l_r \\ | & | & \dots & | \end{pmatrix}$$

Из вышесказанного получаем, что

$$DR_i = \delta_i l_i \text{ для всех } i = 1, \dots, r$$

Другими словами, SVD — это специальная факторизация матрицы D , так что любой базисный вектор r_i для пространства строк отображается в соответствующий базисный вектор l_i в пространстве столбцов, масштабируемый сингулярным значением δ_i . Таким образом, мы можем рассматривать SVD как отображение ортонормированного базиса $(r_1, r_2, \dots, r_r) \mathbb{R}^d$ (пространство строк) в ортонормированный базис $(l_1, l_2, \dots, l_r) \mathbb{R}^n$ (пространство столбцов) с соответствующими осями, масштабируемыми согласно сингулярным значениям $(\delta_1, \delta_2, \dots, \delta_r)$.

7.4.2. Связь между SVD и PCA

Предположим, что матрица D центрирована, и, что она разложена с помощью SVD [ур-е. (7.38)] как $D = L\Delta R^T$. Рассмотрим матрицу рассеяния для D , заданную как $D^T D$. Имеем

$$D^T D = (L\Delta R^T)^T (L\Delta R^T) = R\Delta^T L^T L\Delta R^T = R(\Delta^T \Delta)R^T = R(\Delta_d^2)R^T \quad 7.40$$

где Δ_d^2 - диагональная матрица $d \times d$, определенная как $\Delta_d^2(i, i) = \delta_i^2$ для $i = 1, \dots, d$. Причем только $r \leq \min(d, n)$ из этих собственных чисел положительны, а остальные - нули.

Поскольку ковариационная матрица центрированного D задается как $\Sigma = \frac{1}{n} D^T D$ и поскольку ее можно разложить как $\Sigma = U\Lambda U^T$ через PCA [ур-е. (7.37)], имеем

$$D^T D = n\Sigma = nU\Lambda U^T = U(n\Lambda)U^T \quad 7.41$$

Приравнявая уравнение (7.40) и уравнение (7.41), мы получаем, что правые сингулярные вектора R совпадают с собственными векторами Σ . Кроме того, соответствующие сингулярные значения D связаны с собственными числами Σ через выражение

$$n\lambda_i = \delta_i^2 \text{ или } \lambda_i = \frac{\delta_i^2}{n} \text{ для } i = 1, \dots, d \quad 7.42$$

Рассмотрим теперь матрицу DD^T . Имеем

$$DD^T = (L\Delta R^T)(L\Delta R^T)^T = L\Delta R^T R\Delta^T L^T = L(\Delta\Delta^T)L^T = L\Delta_n^2 L^T$$

где Δ_n^2 - диагональная матрица $d \times d$, определенная как $\Delta_n^2(i, i) = \delta_i^2$ для $i = 1, \dots, d$. Только r из этих сингулярных значений положительны, а все остальные - нули. Таким образом, левые сингулярные вектора L являются собственными векторами матрицы $n \times n$ DD^T , а соответствующие собственные числа задаются как δ_i^2 .