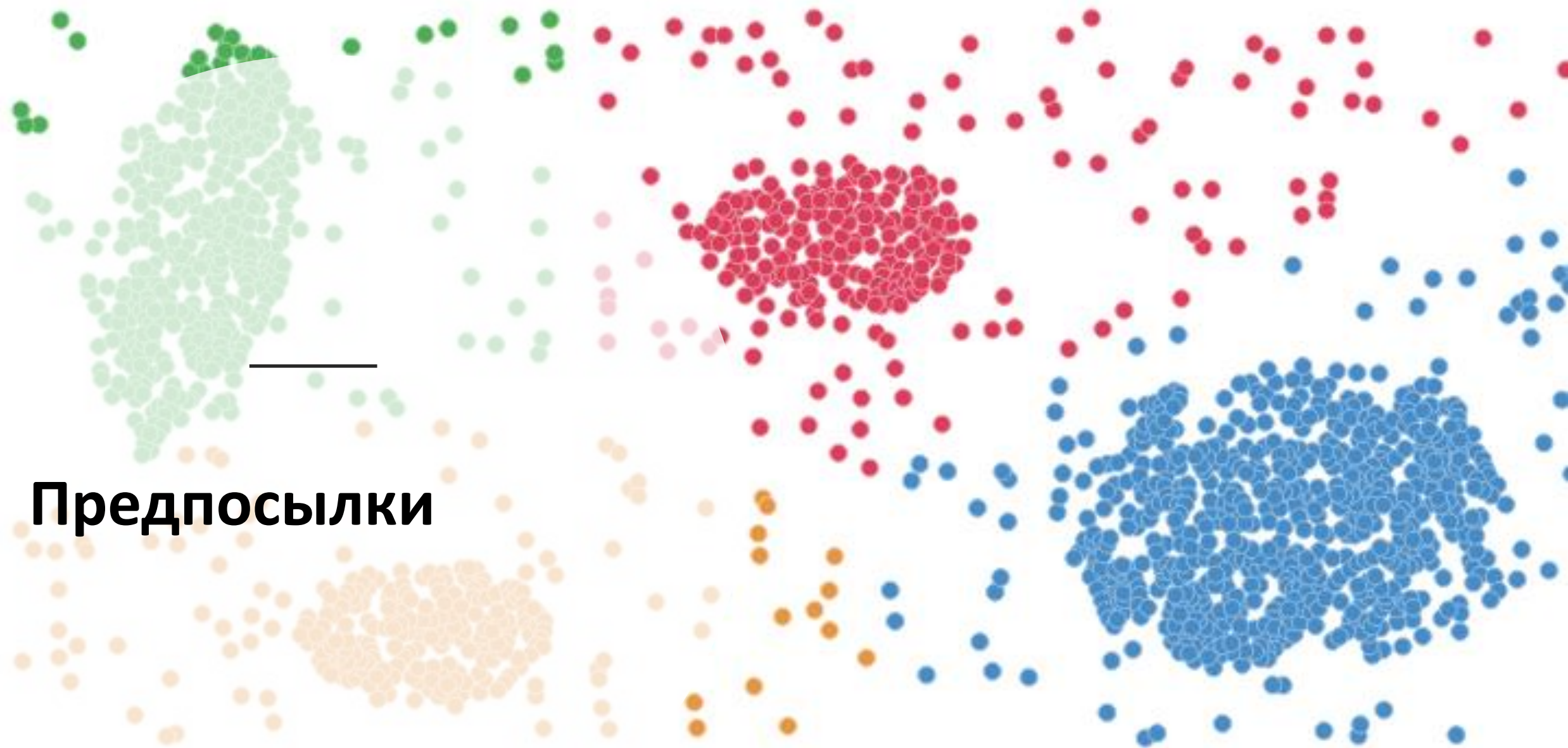




Machine Learning

Лекция #3

Снижение размерности



Предпосылки

Пусть данные — это матрица $\mathbf{D}$ размера

$n \times d$ (n точек с d признаками):

$$\mathbf{D} = \left(\begin{array}{c|ccccc} X & X_1 & X_2 & \cdots & X_d \\ \mathbf{x}_1 & x_{11} & x_{12} & \cdots & x_{1d} \\ \mathbf{x}_2 & x_{21} & x_{22} & \cdots & x_{2d} \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ \mathbf{x}_n & x_{n1} & x_{n2} & \cdots & x_{nd} \end{array} \right)$$

Каждая точка $\mathbf{x}_i = (x_{i1}, x_{i2}, \dots, x_{id})^T$ — 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

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 $\mathbf{u}_i^T \mathbf{u}_i = 1$), возможно представление каждой точки $\mathbf{x}$ как линейной комбинации:

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

где вектор $\mathbf{a} = (a_1, a_2, \dots, a_d)^T$ — это координаты $\mathbf{x}$ в новом базисе. Или в матричной форме: $\mathbf{x} = \mathbf{U}\mathbf{a}$, (7.2)

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

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

$\mathbf{U}$ = 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:

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

Пусть $\mathbf{x}$ — произвольная точка и пусть базисные векторы отсортированы в порядке уменьшения значимости — тогда выражение 7.1 может быть усечено до r членов:

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

Здесь $\mathbf{x}'$ — это проекция $\mathbf{x}$ на первые r базисных векторов, или в терминах матричных операций:

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

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 $\mathbf{P}_r^2 = \mathbf{P}_r$. Поскольку: $\mathbf{P}_r^T = (\mathbf{U}_r \mathbf{U}_r^T)^T = \mathbf{U}_r \mathbf{U}_r^T = \mathbf{P}_r$ и $\mathbf{P}_r^2 = (\mathbf{U}_r \mathbf{U}_r^T)(\mathbf{U}_r \mathbf{U}_r^T) = \mathbf{U}_r \mathbf{U}_r^T = \mathbf{P}_r$. Используя 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 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 следует, 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, возможно сделать 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.



Метод главных компонент
Наилучшая линейная аппроксимация

Метод главных компонент (PCA) – 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.

Наилучшая линейная аппроксимация

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 PCA для наилучшего $1 \leq r \leq d$ размерного базиса для $\mathbf{D}$.

Предполагаем, что $\mathbf{u}$ имеет величину $\|\mathbf{u}\|^2 = \mathbf{u}^T \mathbf{u} = 1$; в ина́че можно продолжать увеличивать проецируемую дисперсию, просто увеличивая величину $\mathbf{u}$. Также предполагаем, что данные центрированы и 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 $\mathbf{u}$. Поскольку средняя точка $\boldsymbol{\mu} = \mathbf{0}$, то ее координата вдоль $\mathbf{u}$ равна $\mu_{\mathbf{u}} = 0$. Необходимо выбрать направление $\mathbf{u}$ так, чтобы дисперсия проецируемых точек была максимальной. Проецируемая дисперсия вдоль $\mathbf{u}$: где $\boldsymbol{\Sigma}$ — это ковариационная матрица для центрированных данных $\mathbf{D}$.

Задача ограниченной оптимизации (максимизировать $\sigma_{\mathbf{u}}^2$ с учетом ограничения $\mathbf{u}^T \mathbf{u} = 1$). Для решения вводим множитель Лагранжа α для ограничения, чтобы получить задачу безусловной максимизации.

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

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

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

$$\frac{\partial}{\partial \mathbf{u}} J(\mathbf{u}) = \mathbf{0}$$

$$\frac{\partial}{\partial \mathbf{u}} (\mathbf{u}^T \boldsymbol{\Sigma} \mathbf{u} - \alpha(\mathbf{u}^T \mathbf{u} - 1)) = \mathbf{0}$$

$$2\boldsymbol{\Sigma} \mathbf{u} - 2\alpha \mathbf{u} = \mathbf{0}$$

$$\boldsymbol{\Sigma} \mathbf{u} = \alpha \mathbf{u} \quad (7.11)$$

Следовательно α - собственное значение 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:

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

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.

Предположим, что набор данных $\mathbf{D}$, 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

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

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

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

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

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

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

Далее, из уравнения (2.28) имеем (*см. основной учебник курса*)

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

Таким образом, можно переписать уравнение (7.15) как

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

Поскольку первый член $\text{var}(\mathbf{D})$ является константой для данного набора данных $\mathbf{D}$, вектор $\mathbf{u}$, который минимизирует $MSE(\mathbf{u})$, максимизирует второй член, проецируемую дисперсию $\mathbf{u}^T \mathbf{\Sigma} \mathbf{u}$. Поскольку мы знаем, что $\mathbf{u}_1$, доминирующий собственный вектор матрицы $\mathbf{\Sigma}$, максимизирует проецируемую дисперсию, имеем

$$MSE(\mathbf{u}_1) = \text{var}(\mathbf{D}) - \mathbf{u}_1^T \mathbf{\Sigma} \mathbf{u}_1 = \text{var}(\mathbf{D}) - \mathbf{u}_1^T \lambda_1 \mathbf{u}_1 = \text{var}(\mathbf{D}) - \lambda_1 \quad (7.16)$$

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



Метод главных компонент

Наилучшая двумерная аппроксимация

Найдем наилучшее двумерное приближение к $\mathbf{D}$. Требуется найти направление $\mathbf{v}$, которое также максимизирует проецируемую дисперсию, но ортогонально $\mathbf{u}_1$. Согласно (7.9) проецируемая на $\mathbf{v}$ дисперсия определяется как

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

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

$$\begin{aligned} \mathbf{v}^T \mathbf{u}_1 &= 0 \\ \mathbf{v}^T \mathbf{v} &= 1 \end{aligned}$$

Тогда условие оптимизации определяется как

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

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

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

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

$$\begin{aligned} 2\mathbf{u}_1^T \Sigma \mathbf{v} - 2\alpha \mathbf{u}_1^T \mathbf{v} - \beta \mathbf{u}_1^T \mathbf{u}_1 &= 0 \\ 2\mathbf{v}^T \Sigma \mathbf{u}_1 - \beta &= 0, \text{ откуда следует, что} \\ \beta &= 2\mathbf{v}^T \lambda_1 \mathbf{u}_1 = 2\lambda_1 \mathbf{v}^T \mathbf{u}_1 = 0 \end{aligned}$$

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

$$\begin{aligned} 2\Sigma \mathbf{v} - 2\alpha \mathbf{v} &= \mathbf{0} \\ \Sigma \mathbf{v} &= \alpha \mathbf{v} \end{aligned}$$

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 $\mathbf{v}$, выбираем $\alpha = \lambda_2$, второе по величине собственное значение Σ , 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 $\mathbf{A}$. Так как 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

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

$$= \frac{1}{n} \sum_{i=1}^n \mathbf{x}_i^T (\mathbf{U}_2 \mathbf{U}_2^T) \mathbf{x}_i = \frac{1}{n} \sum_{i=1}^n \mathbf{x}_i^T \mathbf{P}_2 \mathbf{x}_i \quad (7.19)$$

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

$$var(\mathbf{A}) = \frac{1}{n} \sum_{i=1}^n \mathbf{x}_i^T \mathbf{P}_2 \mathbf{x}_i \quad (7.20)$$

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

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

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

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

Среднеквадратичная ошибка

Покажем, что первые два главных компонента минимизируют среднеквадратичную ошибку. Среднеквадратичная ошибка задается как

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

Таким образом, MSE минимизируется, когда общая проецируемая дисперсия $\text{var}(\mathbf{A})$ максимальна. Из (7.22) имеем $MSE = \text{var}(\mathbf{D}) - \lambda_1 - \lambda_2$



Метод главных компонент
Наилучшая r -мерная аппроксимация

Найдем наилучшую r -мерную аппроксимацию 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 $\mathbf{\Sigma}$, 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

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

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

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

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

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

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

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

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

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

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

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

Чтобы максимизировать дисперсию по $\mathbf{v}$, 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 D , вычисляем собственные значения матрицы Σ . Т.к. Σ является положительно полуопределенной, все ее собственные числа должны быть неотрицательными, и мы можем отсортировать их в порядке убывания следующим образом:

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

Затем 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

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

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 $\mathbf{A}$ - 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

Из (7.23), среднеквадратичная ошибка в r -измерениях может быть записана как

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

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

Общая дисперсия

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

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

Алгоритм 7.1 демонстрирует псевдокод для алгоритма метода главных компонент PCA.

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. $\boldsymbol{\Sigma} = \frac{1}{n} (\mathbf{Z}^T \mathbf{Z})$ // вычисление ковариационной матрицы
4. $(\lambda_1, \lambda_2, \dots, \lambda_d) = \text{eigenvalues}(\boldsymbol{\Sigma})$ // вычисление собственных значений
5. $\mathbf{U} = (\mathbf{u}_1 \quad \mathbf{u}_2 \quad \dots \quad \mathbf{u}_d) = \text{eigenvectors}(\boldsymbol{\Sigma})$ // 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



Ядерный анализ главных компонент

Анализ главных компонент может быть расширен, чтобы искать нелинейные “направления” в данных с использованием ядерных методов. То есть, вместо поиска линейных комбинаций во входном пространстве, алгоритм ищет линейные комбинации в пространстве признаков, полученном из нелинейного преобразования входного пространства.

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:

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.

$$\begin{aligned} \left(\frac{1}{n} \sum_{i=1}^n \phi(\mathbf{x}_i) \phi(\mathbf{x}_i)^T \right) \mathbf{u}_1 &= \lambda_1 \mathbf{u}_1 \\ \frac{1}{n} \sum_{i=1}^n \phi(\mathbf{x}_i) (\phi(\mathbf{x}_i)^T \mathbf{u}_1) &= \lambda_1 \mathbf{u}_1 \quad (7.31) \\ \sum_{i=1}^n \left(\frac{\phi(\mathbf{x}_i)^T \mathbf{u}_1}{n\lambda_1} \right) \phi(\mathbf{x}_i) &= \mathbf{u}_1 \\ \sum_{i=1}^n c_i \phi(\mathbf{x}_i) &= \mathbf{u}_1 \quad (7.32) \end{aligned}$$

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

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

Заменяем 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

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

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

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

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

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

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:

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

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:

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

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

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

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

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. $(\mathbf{c}_1 \quad \mathbf{c}_2 \quad \cdots \quad \mathbf{c}_n) = \text{eigenvectors}(\mathbf{K})$ // 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. $\mathbf{C}_r = (\mathbf{c}_1 \quad \mathbf{c}_2 \quad \cdots \quad \mathbf{c}_r)$ // 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



Сингулярное разложение

Анализ главных компонент является частным случаем более общего метода разложения матриц, называемого сингулярным разложением (Singular Value Decomposition - SVD).

SVD обобщает приведенную выше факторизацию для любой матрицы. В частности, для матрицы данных $\mathbf{D}$ размера $n \times d$ с n точками и d столбцами SVD факторизует $\mathbf{D}$ следующим образом:

$$\mathbf{D} = \mathbf{L}\mathbf{\Delta}\mathbf{R}^T \quad 7.38$$

где $\mathbf{L}$ - ортогональная матрица размера $n \times n$, $\mathbf{R}$ - ортогональная матрица размера $d \times d$, а $\mathbf{\Delta}$ - «диагональная» матрица $n \times d$. Столбцы $\mathbf{L}$ называются левыми сингулярными векторами, а столбцы $\mathbf{R}$ (или строки $\mathbf{R}^T$) называются правыми сингулярными векторами. Матрица $\mathbf{\Delta}$ 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 $\mathbf{\Delta}$ называются сингулярными числами $\mathbf{D}$, и все они неотрицательны.

Если ранг $\mathbf{D}$ 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:

$$\mathbf{D} = \mathbf{L}_r \mathbf{\Delta}_r \mathbf{R}_r^T \quad 7.39$$

где $\mathbf{L}_r$ - матрица левых сингулярных векторов размера $n \times r$, $\mathbf{R}_r$ - матрица правых сингулярных векторов размера $d \times r$, а $\mathbf{\Delta}_r$ - диагональная матрица размера $r \times r$, содержащая положительные сингулярные вектора. Уменьшенный SVD приводит непосредственно к спектральному разложению $\mathbf{D}$, заданному как

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

Спектральное разложение $\mathbf{D}$ 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 $\mathbf{D}$.

То есть, если $\mathbf{D}_q$ - матрица, определенная как

$$\mathbf{D}_q = \sum_{i=1}^r \delta_i \mathbf{l}_i \mathbf{r}_i^T$$

то можно показать, что $\mathbf{D}_q$ - это матрица ранга q , которая минимизирует следующее выражение $\|\mathbf{D} - \mathbf{D}_q\|_F$

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

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

Связь между SVD и PCA

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

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

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

$$\mathbf{D}^T\mathbf{D} = n\Sigma = n\mathbf{U}\Lambda\mathbf{U}^T = \mathbf{U}(n\Lambda)\mathbf{U}^T \quad 7.41$$

Приравнивая уравнение (7.40) и уравнение (7.41), мы получаем, что правые сингулярные вектора $\mathbf{R}$ совпадают с собственными векторами Σ . Кроме того, соответствующие сингулярные значения $\mathbf{D}$ 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 $\mathbf{D}\mathbf{D}^T$. Имеем

$$\mathbf{D}\mathbf{D}^T = (\mathbf{L}\Delta\mathbf{R}^T)(\mathbf{L}\Delta\mathbf{R}^T)^T = \mathbf{L}\Delta\mathbf{R}^T\mathbf{R}\Delta^T\mathbf{L}^T = \mathbf{L}(\Delta\Delta^T)\mathbf{L}^T = \mathbf{L}\Delta_n^2\mathbf{L}^T$$

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 $\mathbf{L}$ являются собственными векторами матрицы $n \times n$ $\mathbf{D}\mathbf{D}^T$, а соответствующие собственные числа задаются как δ_i^2 .



Вопросы?