

clustering $\leftarrow$ unsupervised method
dimensional reduction (입력만)

Preclass 05: K-Means Clustering

[SCS4049] Machine Learning and Data Science

Seongsik Park (s.park@dgu.edu)

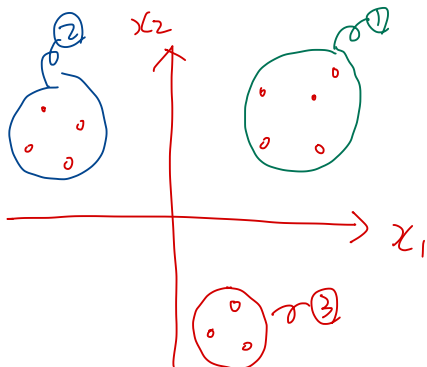
School of AI Convergence & Department of Artificial Intelligence, Dongguk University

중간고사

- 10월 26일 (수) 오전 10시 (5시간)
- P404호
- Clustering 가지, (이론, 실험, 과제)
- Cheat sheet 1장 (양면, 2쪽), A4용지.
- Closed book

$\therefore$
 input samples $\rightarrow$ clustering $\rightarrow$ cluster label

①
 ②
 ③



$\Rightarrow$ ① kmeans clustering
 ② mixture of Gaussian.

입력

$x_1, x_2, \dots, x_N$

→ K means clustering

→ one-of-K coding label $R_{N \times K}$

cluster의 개수 K

K 개의 중심 = 평균
centroid = mean

iterative method

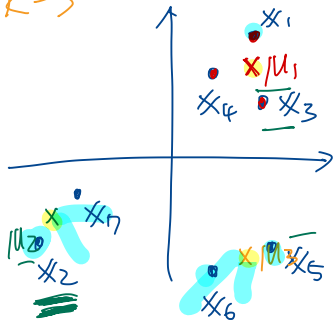
주어진 샘플

주어진 K 개의 중심

label을 계산

←
주어진 샘플

$K=3$



$$\mathbb{R} = \begin{matrix} N \times K \\ 7 \times 3 \end{matrix}$$

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

$$\begin{aligned} \mu_1 &\leftarrow \frac{1}{3} (x_1 + x_3 + x_4) \\ \mu_2 &\leftarrow \frac{1}{2} (x_2 + x_7) \\ \mu_3 &\leftarrow \frac{1}{2} (x_5 + x_6) \end{aligned}$$

3개의 sample
 $x_1 \dots x_7$
 3개의 label $\mathbb{R}$
 $\downarrow$
 중심 $\mu_1 \dots \mu_3$ 갱신 0/1

$$D_{N \times K} = \left\| \underline{x_n} - u_k \right\|_2$$

$$D_{5 \times 3} = \begin{bmatrix} & u_1 & u_2 & u_3 \\ 1 & 4 & 9 \\ 2 & 5 & 6 \\ 8 & 2 & 7 \\ 6 & 1 & 9 \\ 10 & 5 & 1 \end{bmatrix}$$

$$\rightarrow R_{5 \times 3} = \begin{bmatrix} 1 & 0 & 0 \\ 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{bmatrix}$$

K-means clustering

$K=2$
↓
10의 초기 위치
가 결정

초기화.

샘플 + centroid
→ label

샘플 + label
→ centroid

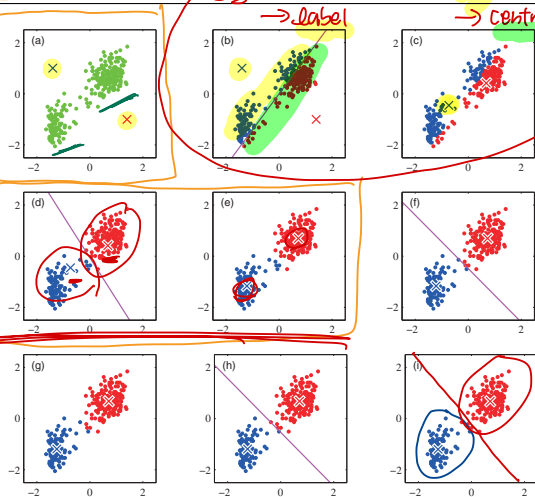


Figure 9.1 Illustration of the K -means algorithm using the re-scaled Old Faithful data set. (a) Green points denote the data set in a two-dimensional Euclidean space. The initial choices for centres μ_1 and μ_2 are shown by the red and blue crosses, respectively. (b) In the initial E step, each data point is assigned either to the red cluster or to the blue cluster, according to which cluster centre is nearer. This is equivalent to classifying the points according to which side of the perpendicular bisector of the two cluster centres, shown by the magenta line, they lie on. (c) In the subsequent M step, each cluster centre is re-computed to be the mean of the points assigned to the corresponding cluster. (d)–(i) show successive E and M steps through to final convergence of

K-means: method

We begin by considering the problem of identifying groups, or clusters, of data points in a multidimensional space. Suppose we have a data set $\{\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_N\}$ consisting of N observations of a random D -dimensional Euclidean variable $\mathbf{x}$. Our goal is to partition the data set into some number K of clusters, where we shall suppose for the moment that the value of K is given.

Intuitively, we might think of a cluster as comprising a group of data points whose inter-point distances are small compared with the distances to points outside of the cluster. We can formalize this notion by first introducing a set of D -dimensional vectors $\boldsymbol{\mu}_k$, where $k = 1, \dots, K$, in which $\boldsymbol{\mu}_k$ is a prototype associated with the k -th cluster.

K-means: method

As we shall see shortly, we can think of the μ_k as representing the centres of the clusters. Our goal is then to find an assignment of data points to clusters, as well as a set of vectors $\{\mu_k\}$, such that the sum of the squares of the distances of each data point to its closest vector μ_k , is a minimum.

It is convenient at this point to define some notation to describe the assignment of data points to clusters. For each data point $\mathbf{x}_n$, we introduce a corresponding set of binary indicator variables $r_{nk} \in \{0, 1\}$, where $k = 1, \dots, K$ describing which of the K clusters the data point $\mathbf{x}_n$ is assigned to, so that if data point $\mathbf{x}_n$ is assigned to cluster k then $r_{nk} = 1$, and $r_{nj} = 0$ for $j \neq k$. This is known as **the 1-of- K coding scheme**.

K-means: method

We can then define an objective function, sometimes called a *distortion measure*, given by

$$J = \sum_{n=1}^N \sum_{k=1}^K r_{nk} \|\mathbf{x}_n - \boldsymbol{\mu}_k\|^2 \quad (1)$$

which represents the sum of the squares of the distances of each data point to its assigned vector $\boldsymbol{\mu}_k$.

Our goal is to find values for the $\{r_{nk}\}$ and the $\{\boldsymbol{\mu}_k\}$ so as to minimize J .

$$R = \begin{bmatrix} 1 & 0 \\ 0 & 1 \\ 1 & 0 \end{bmatrix}$$

$$\sum_{n=1}^3 \sum_{k=1}^2 r_{nk} \|\mathbf{x}_n - \boldsymbol{\mu}_k\|^2$$
$$\sum_{k=1}^2 \left(\sum_{n=1}^3 r_{nk} \right) \|\mathbf{x}_1 - \boldsymbol{\mu}_k\|^2$$

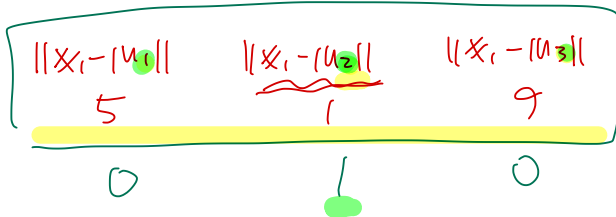
$$= 1 \cdot \|\mathbf{x}_1 - \boldsymbol{\mu}_1\|^2 + 0 \cdot \|\mathbf{x}_1 - \boldsymbol{\mu}_2\|^2$$

K-means: method

The terms involving different n are independent and so we can optimize for each n separately by choosing r_{nk} to be 1 for whichever value of k gives the minimum value of $\|\mathbf{x}_n - \boldsymbol{\mu}_k\|^2$. In other words, we simply assign the n -th data point to the closest cluster centre.

$$r_{nk} = \begin{cases} 1 & \text{if } k = \arg \min_j \|\mathbf{x}_n - \boldsymbol{\mu}_j\|^2 \\ 0 & \text{otherwise.} \end{cases} \quad (2)$$

R



K-means: method

Now consider the optimization of the μ_k with the r_{nk} held fixed. The objective function J is a quadratic function of μ_k , and it can be minimized by setting its derivative with respect to μ_k to zero giving

$$2 \sum_{n=1}^N r_{nk} (\mathbf{x}_n - \mu_k) = 0 \quad (3)$$

$$\mu_k = \frac{\sum_n r_{nk} \mathbf{x}_n}{\sum_n r_{nk}} \quad \begin{array}{l} \text{K번及n번이 속하는 샘플의 합} \\ \text{K번及n번이 속하는 샘플개수} \end{array} \quad (4)$$

The denominator in this expression is equal to the number of points assigned to cluster k , and so this result has a simple interpretation, namely set μ_k equal to the mean of all of the data points $\mathbf{x}_n$ assigned to cluster k .

The two phases of re-assigning data points to clusters and re-computing the cluster means are repeated in turn until there is no further change in the assignments (or until some maximum number of iterations is exceeded).

K-means: python example

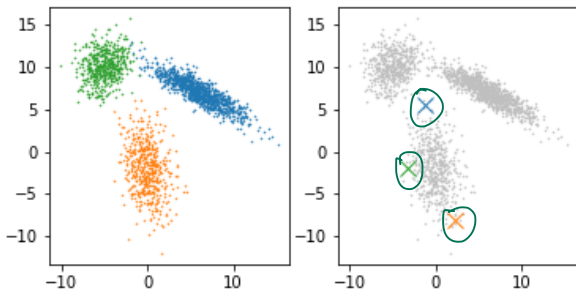


Figure 1: K-means algorithm python example: dataset (left) and the initialization (right).

K-means: python example

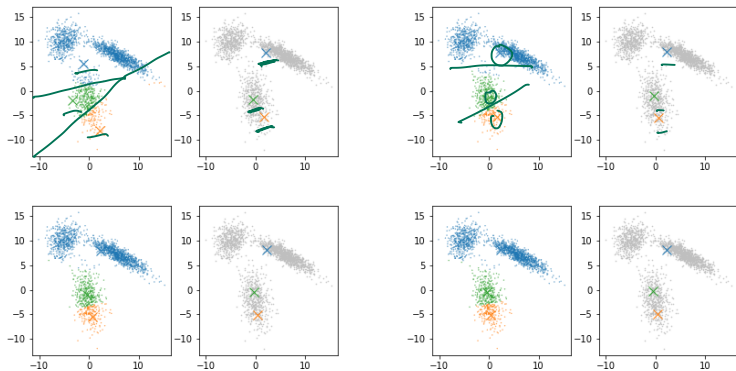


Figure 2: From the 1st to 4th iterations.

K-means: python example

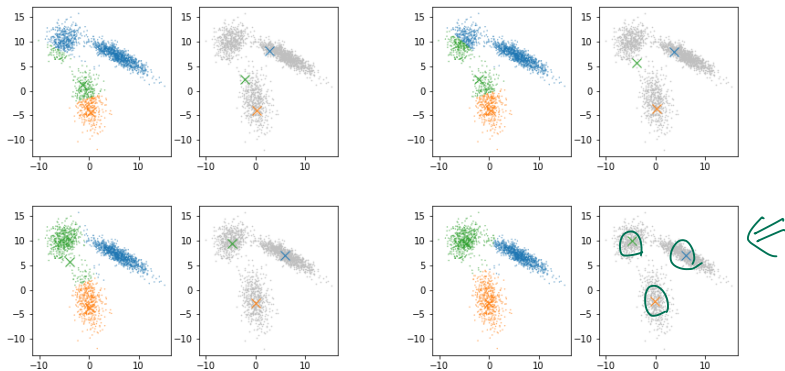


Figure 3: From the 9th to 12th iterations.

K-means: example

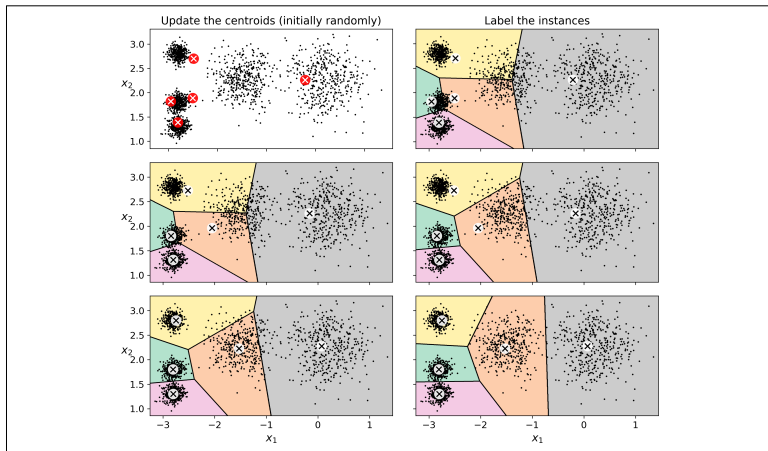
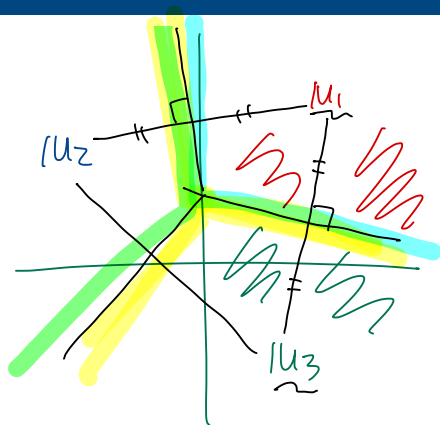


Figure 9-4. The K-Means algorithm



① convex

② single connect.



스킵이등분선

K-means: decision boundary

Decision boundary = Voronoi tessellation

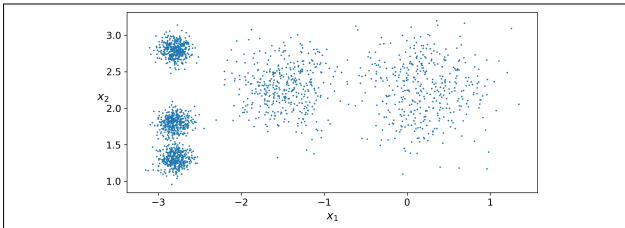


Figure 9-2. An unlabeled dataset composed of five blobs of instances

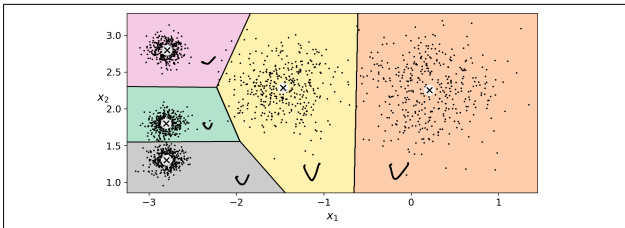


Figure 9-3. K-Means decision boundaries (Voronoi tessellation)

Appendix

Reference and further reading

- “Chap 9” of A. Geron, Hands-On Machine Learning with Scikit-Learn, Keras & TensorFlow
- “Chap 9” of C. Bishop, Pattern Recognition and Machine Learning
- Variational Bayesian mixtures of Gaussians