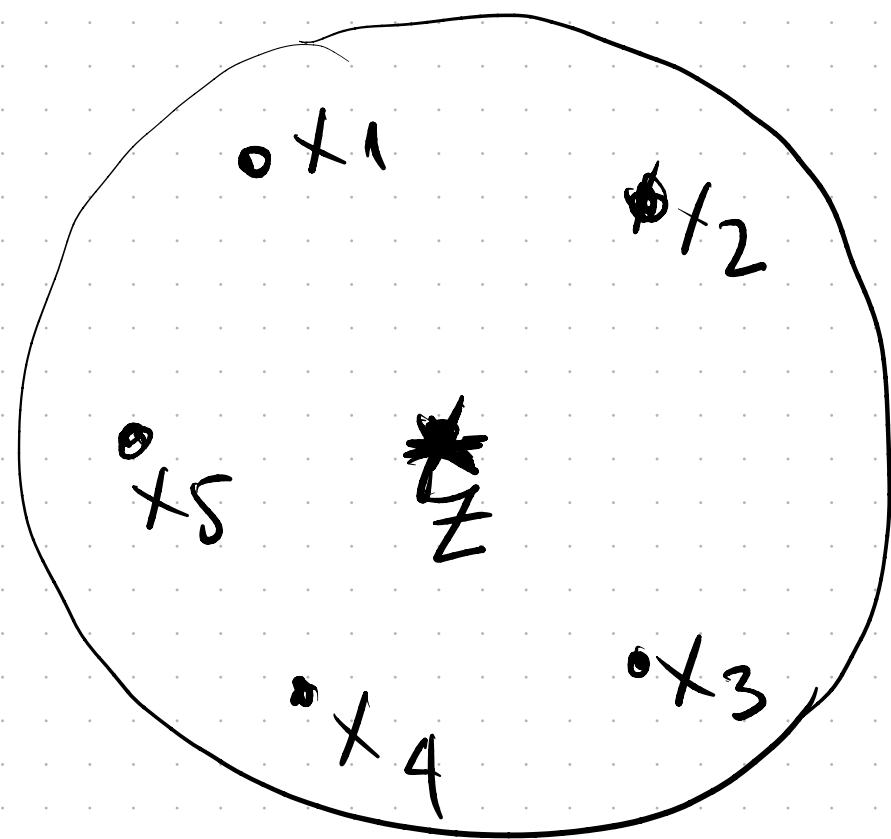


Lecture 7/18

- KNN
- optional: tsne
- Kernel PCA
- Hw5 demo

K Nearest Neighbor : For each point $z \in \text{Test}$
- define neighbourhood $N_z = \{ \text{set of closest points to } z \}$



- predict label/class/target by estimating from N_z

classification: predict class Y_z that is most present in N_z $Y_z = \text{mode}\{N_z\}$

quantity/regression: predict AVG $\{N_z\}$ label

ex $y = \text{house price} \Rightarrow \text{pred } Y_z = \text{AVG}_{x_i \in N_z} \{y_i\}$

not really training, just pred $\hat{y}_Z =$ estimate from closest neighbors to Z

● need dist/similarity (kernel) across datapoints

$$k_{ij} = \text{sim}(x_i, x_j)$$

$$k_{iz} = \text{sim}(x_i, z)$$

appropriate for data
and for task

ex $x_i = \text{patients} \Rightarrow k_{ij} = \text{similarity of patients w.r.t diabetes}$
 $y = \text{diabetes}$

$x_i = \text{images}$
 $y = \text{prod. price}$

$\Rightarrow k_{ij} = \text{similarity of image prod w.r.t. price (willingness to pay)}$

$x_i = \text{email}$
 $y = \text{spam/not}$

$\Rightarrow k_{ij} = \text{similarity of emails w.r.t. spam}$

$x_i = \text{movies}$
 $y = \text{rating}$

$\Rightarrow k_{ij} = \text{similarity of movies w.r.t user satisfaction}$

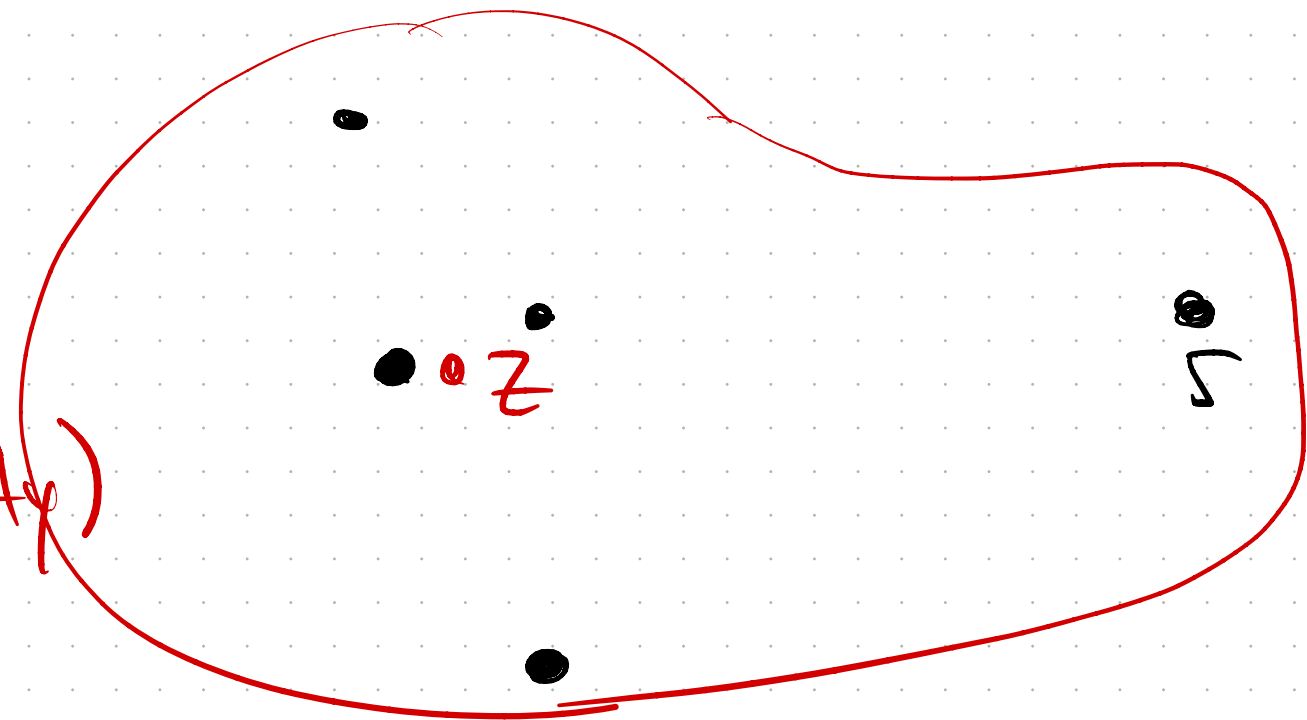
KNN 3 variants

- $K = \text{fixed}$; ex $\Rightarrow K = 5$. N_Z always has 5 training points.
 $N_Z = \{ \text{closest 5 to } Z \}$

— rank all training points
by $K z_i$ similarity

— select top $K \rightarrow N_Z$

disadvantage: some Z don't have
5 close neighbors (low density)

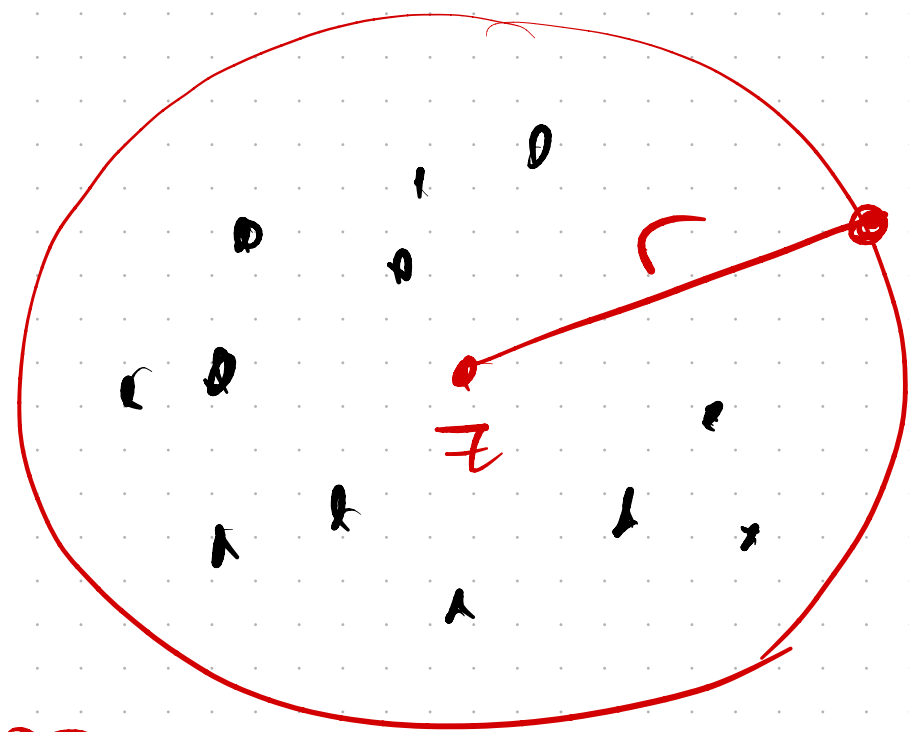


- range = fixed & variable

$$N_Z = \{ x_i \text{ training} \mid K_{iz} \geq r \}$$
$$\text{dist}(i, z) \leq r$$

disadvantage: some Z ^{very} few neighbors
in range

some Z : many neighbors in range



• use all training points weighted by similarity $k_{iz} = \text{sim}(x_i, z)$

regression predict $(z) = \frac{\sum_{i=1}^N k_{iz} \cdot y_i}{\sum_{i=1}^N k_{iz}} \rightarrow \text{labels / targets}$

weights

classification

class l
label $\uparrow$

score $(z, y_l) = \sum_{i=1}^N k_{iz} \cdot \boxed{1[y_i = y_l]}$ filter

N_z = all training set, but weighted.

$k_{iz} = \text{high} \Rightarrow x_i \sim z \Rightarrow y_i$ counts a lot

$k_{iz} = \text{low} \Rightarrow x_i$ not similar $\Rightarrow y_i$ doesn't count much to z

Kernel PCA

- rewrite PCA primal $\Rightarrow$ dual form.

dual variables.

PCA word = $X \cdot \begin{bmatrix} e_1 & e_2 & \dots & e_p \\ | & | & & | \\ 1 & 2 & & p \end{bmatrix}$

eigen vectors (Σ covar)
sorted by desc eigen val

$\stackrel{?}{=} \boxed{X X^T} \cdot \boxed{B}$

K

- choose K
- can use any valid kernel K (not just linear $X X^T$)
- ex $K = \text{gaussian} = e^{-\|x_i - x_j\|^2 / 2\sigma^2}$

- compute B for that kernel.

compute

$$B = \text{eigenvec}(K)$$

Intuition we want eigen vectors
(primal var) $= X^T \beta$ dual.

detail (skip): show that $e = X^T \beta$ possible for every e

calculate/find β ?: e eigen vector of $\Sigma \Rightarrow \Sigma e = \lambda e$
eigen val

derivation: $\Sigma e = \lambda e$

$$\Sigma = \frac{1}{N} X^T X \quad \text{covar estimation}$$

$$e = X^T \beta \quad \rightarrow \text{want}$$

$$\left(\frac{1}{N} X^T X \right) X^T \beta = \lambda X^T \beta$$

$$X^T X X^T \beta = N \lambda X^T \beta$$

get rid of X left: $X X^T X X^T \beta = N \lambda X X^T \beta$

get rid of X ? $K \cdot K \cdot \beta = N \lambda K \cdot \beta$
 $K \cdot \beta = N \lambda \cdot \beta$ scalar

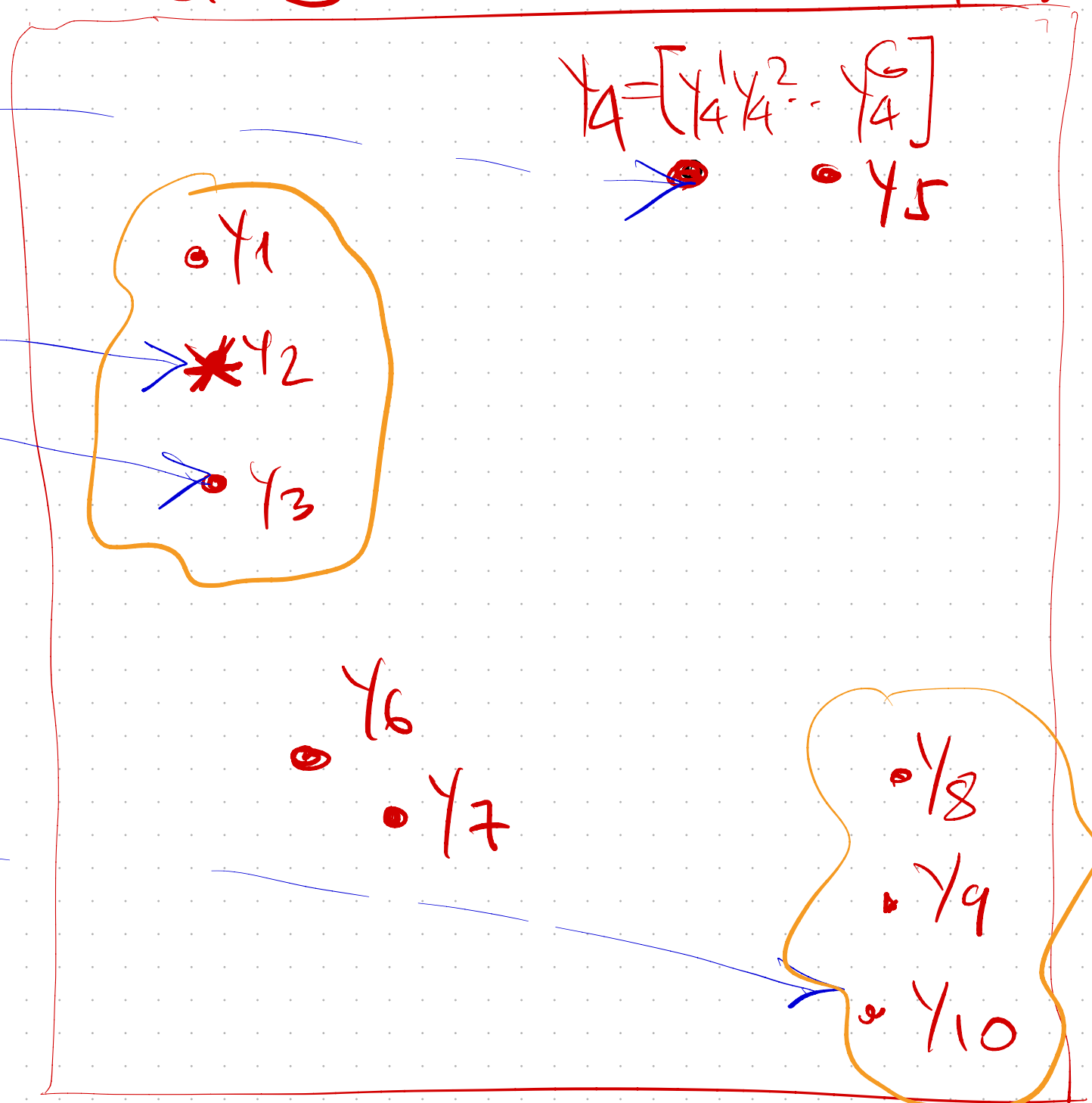
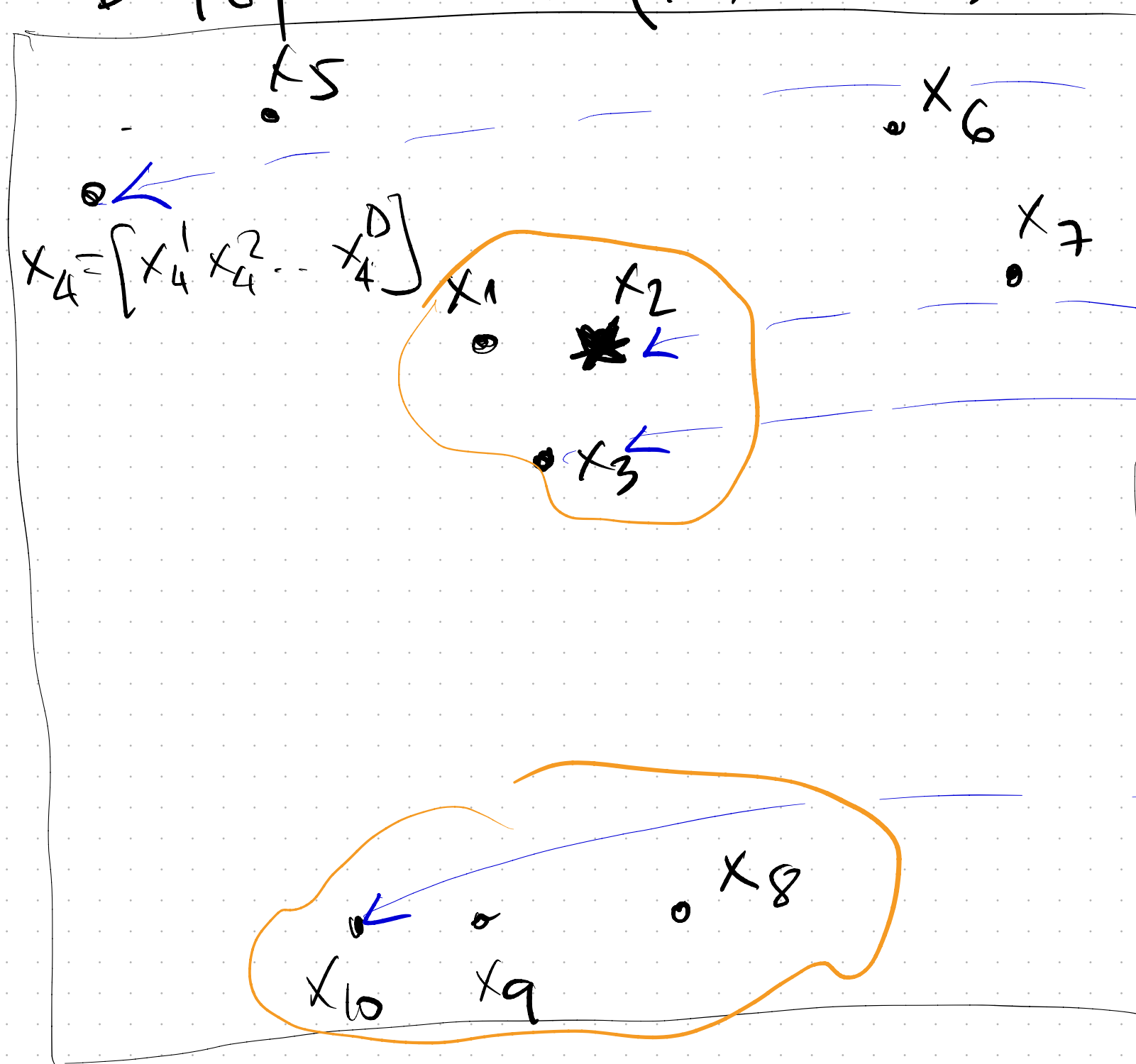
$$\beta = \text{eigenvector}(K) \\ N \lambda = \text{corresp eigenval}$$

t-SNE representation of data $D\text{-dim} \Rightarrow G\text{-dim} (G \ll D)$

$D\text{-dim}$ $X\text{-space (original)}$
 $D=784$
 $X = (x^1 x^2 \dots x^D)$

$$X_i \Rightarrow Y_i$$

$G\text{-dim}$ $Y\text{-space (new)}$
 $G=3$
 $Y = (y^1 y^2 \dots y^G)$



- do not worry about exact distances / positioning (in algebra)

$\text{dist}(x_1, x_5)$?? $\text{dist}(y_1, y_5)$

$\text{angle}(x_1, x_9, x_7)$?? $\text{angle}(y_1, y_9, y_7)$

= care about: preserve "close" vs "far" neighbors.

one time
Fixed x_2 much closer to x_3 than to x_{10} $\iff$

y_2 much closer to y_3 than to y_{10}

Q = distribution of distances (x)
 $\neq$ not gaussian. use t -distrib.

$$P_{ij} = \frac{\exp(-\|x_i - x_j\|^2 / 2\sigma^2)}{\text{normalized}}$$

$N \times N$

$$Q_{ij} = \frac{1}{[1 + \|y_i - y_j\|^2]} \quad t\text{-distrib}$$

$N \times N$
normalized

not fixed, updated with Y

- Given $X_{N \times D}$ compute $\underline{P}_{N \times N}$ (search for "width" σ^2 by perplexity)

- want $Y_{N \times G}$ such that $Q(Y) \underset{\text{dist}}{\approx} \underset{\text{dist}}{P}(X)$

MIN OBJ = $KL(P \parallel Q) = \sum_{i,j} P_{ij} \log \frac{P_{ij}}{Q_{ij}}$

P_{ij} fixed Q_{ij} variable

$P = \text{fixed}$
 $Q = \text{variable}(Y)$

- init $Y = \text{random (better cluster)}$

- iterate (Grad descent)

- compute $Q = t\text{-dist}(Y \text{ positions})$

- update Y based Grad Decent Rule

$$\frac{\partial \text{OBJ}}{\partial Y_i} = 4 \sum_j (P_{ij} - Q_{ij}) (y_i - y_j)$$

Q_{ij} numerator

$$\frac{1}{1 + \|y_i - y_j\|^2}$$

easy to compute (fast using Laplacian - kind of comp on graph degrees)

$$y_i^{\text{new}} = y_i + \eta \frac{\partial \text{OBJ}}{\partial y_i} + \beta (y_i - y_i^{\text{prev}})$$