

ML & Optimization Lecture 16

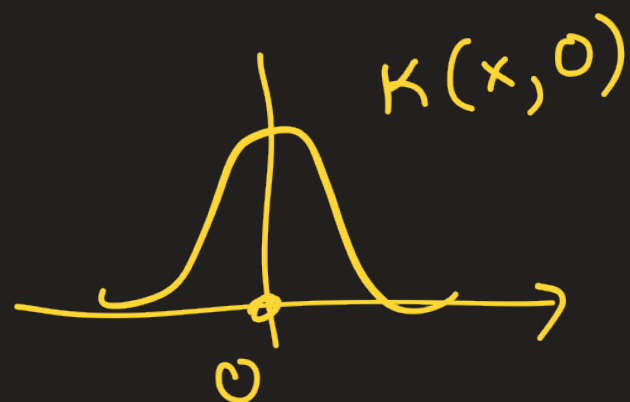
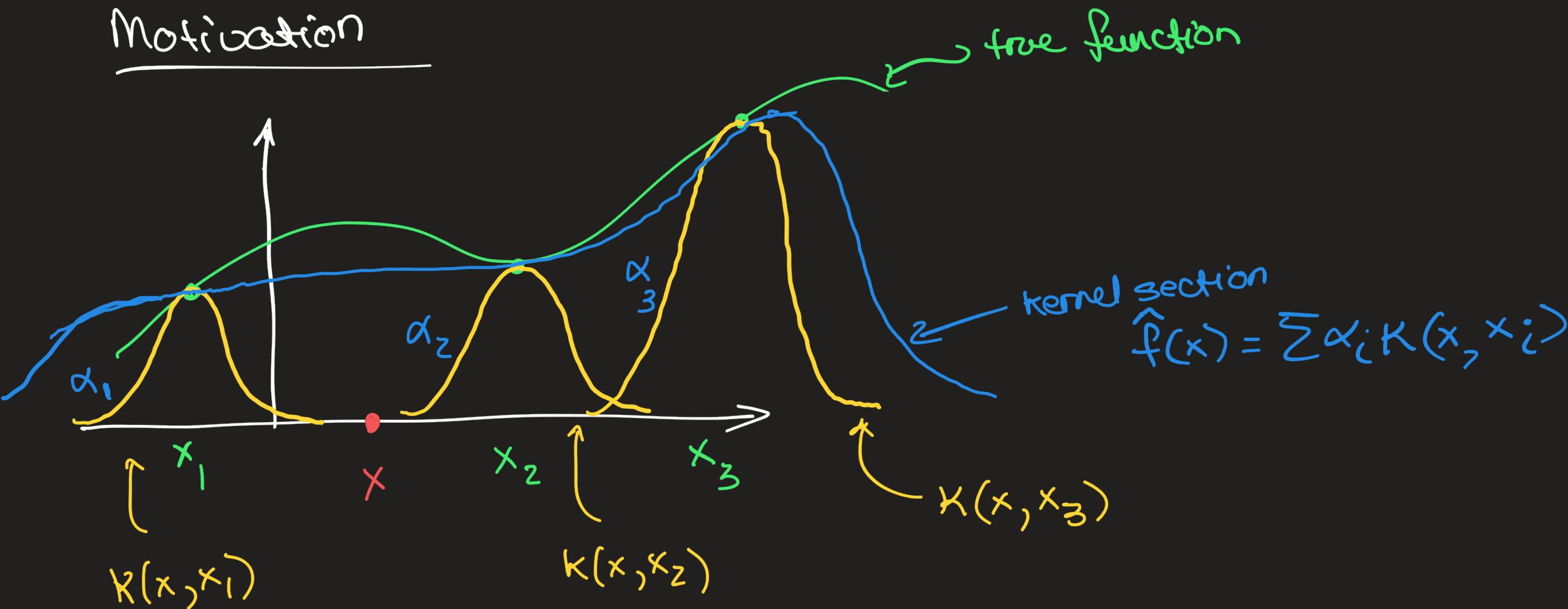
- Kernel Learning
- Approximate Kernel Learning

Brief History of ML

- linear models, learn w such that $w^T x$ is useful for a task:
 - SVMs: $\text{sgn}(w^T x + b)$
 - Logistic Regression: $\text{softmax}(w x + b)$
 - linear regression: $y = A x$
- feature-engineering: domain-dependent design $\phi: \mathbb{R}^d \rightarrow \mathbb{R}^D$ so that $w^T \phi(x)$ is useful
- neural networks: learn ϕ as a composition of multiple simpler functions, e.g. $\phi = \phi_3 \circ \phi_2 \circ \phi_1$, where $\phi_i: x \mapsto \sigma(w_i x + b_i)$, so that $w^T \phi(x)$
- kernel methods Fix a feature map ϕ . learn weights w s.t. $f(x) = \sum_{i=1}^n w_i \langle \phi(x), \phi(x_i) \rangle$ is useful.

- Neural Networks redux (Deep Learning)
 - Neural networks again, but with:
 - more data
 - better architectures
 - more computational capabilities
 - more sophisticated optim algs.

Motivation



In general we define a feature map $\phi: \mathbb{R}^d \rightarrow \mathbb{R}^D$
then the kernel function

$$k(x, y) = \langle \phi(x), \phi(y) \rangle$$

Notice that k is a PSD, symmetric function. That is,
for any $\{x_1, \dots, x_n\}$, then the corresponding kernel matrix

$$K = \left[k(x_i, x_j) \right]_{i,j=1}^n = \left[\begin{array}{cccc} \langle \phi(x_1), \phi(x_1) \rangle & \dots & \langle \phi(x_1), \phi(x_n) \rangle \\ \vdots & \ddots & \vdots \\ \langle \phi(x_n), \phi(x_1) \rangle & \dots & \langle \phi(x_n), \phi(x_n) \rangle \end{array} \right]$$
$$= \begin{bmatrix} \phi(x_1)^T \\ \vdots \\ \phi(x_n)^T \end{bmatrix} \begin{bmatrix} \phi(x_1) & \dots & \phi(x_n) \end{bmatrix} = \underline{\Phi} \underline{\Phi}^T$$

where $\underline{\Phi} \in \mathbb{R}^{n \times D}$

More generally (because we don't want to design ϕ) we pick a kernel function directly. This corresponds to some choice of ϕ .

Any function $K(x, y)$ that is PSD is a valid kernel.

Common choices:

- polynomial kernel of order r

$$K_{\text{poly}}(x, y) = (c \langle x, y \rangle + d)^r$$

corresponds to

$$\phi(x) = \begin{bmatrix} \text{all the monomials in the } d \text{ variables} \\ x_1, \dots, x_d \text{ and a constant term} \\ \text{of order up to and including } r \end{bmatrix}$$

$$D = \mathcal{O}(d^r)$$

- Gaussian kernel (RBF kernel, squared-exponential kernel)

$$K_{\text{Gauss}}(x, y) = \exp\left(-\frac{\|x - y\|_2^2}{2\sigma^2}\right)$$

$$\phi(x) = [\text{all monomials in } x, \text{ appropriately scaled}]$$

Most used in ML because of its universality

- Laplacian kernel

$$K_{\text{Laplacian}}(x, y) = \exp\left(-\frac{\|x - y\|_1}{\sigma}\right)$$

- many other standard kernels for dealing w/ data in non-Euclidean domains: graphs, strings, etc.

Motivation for working directly w/ kernels instead of ϕ

Consider an NLP problem where we use BOW encodings of input text spans.

"the dog are my homework" $\mapsto$ $[0 \ 1 \ 0 \dots 0 \ 1 \ 0 \ 1 \ 0 \ 0 \ 0 \ 0 \ 1 \ 0 \dots]$

dog cat are homework my

$\downarrow$ $\downarrow$ $\downarrow$ $\downarrow$ $\downarrow$

$\in \mathbb{R}^D$ where D is size of vocab

We want to consider interactions of pairs of words. Using ERM we would fit this model

$$v_* = \operatorname{argmin}_v \frac{1}{n} \sum_{i=1}^n (y_i - v^T \phi(x_i))^2 + \lambda \|v\|_2^2$$

where

$$\phi\left(\begin{bmatrix} x_1 \\ \vdots \\ x_d \end{bmatrix}\right) = \begin{bmatrix} x_1 \\ \vdots \\ x_d \\ x_1^2 \\ x_1 x_2 \\ \vdots \\ x_d^2 \end{bmatrix} \in \mathbb{R}^{1 + 2d + \binom{d}{2}}$$

In practice $d \approx 70k$ (say on Wikipedia) and that means $D = \mathcal{O}(100M)$. Let's see what happens if we use the standard RR formulation

$$v_* = \underset{v}{\operatorname{argmin}} \frac{1}{n} \|y - \bar{\Phi} v\|_2^2 + \lambda \|v\|_2^2$$

$$= \underbrace{(\bar{\Phi}^T \bar{\Phi} + n\lambda I)^{-1}}_{D \times D} \bar{\Phi}^T y$$

$D \times D$ cost $\mathcal{O}(D^3)$ to invert

so we cannot find v_* .

We only care about predicting. so we will use the kernel trick : we only need to compute

$$K(x_i, x_j) = \langle \phi(x_i), \phi(x_j) \rangle$$

Notice

$$\begin{aligned} f(x_{\text{new}}) &= v_*^T \phi(x_{\text{new}}) = \phi(x_{\text{new}})^T v_* \\ &= \phi(v_{\text{new}})^T (\bar{\Phi}^T \bar{\Phi} + n\lambda I)^{-1} \bar{\Phi}^T y \end{aligned}$$

observe (from the SVD) that for any matrix M and scalar $c > 0$

$$(M^T M + cI)^{-1} M^T = M^T (M M^T + cI)^{-1}$$

Take $M = \bar{\Phi}$ and $c = n\lambda$, then

$$\begin{aligned} f(x_{\text{new}}) &= \phi(x_{\text{new}})^T (M^T M + cI)^{-1} M^T y \\ &= \phi(x_{\text{new}})^T M^T (M M^T + cI)^{-1} y \\ &= \phi(x_{\text{new}})^T \bar{\Phi}^T (\bar{\Phi} \bar{\Phi}^T + cI)^{-1} y \end{aligned}$$

Notice that $K = \Phi \Phi^T = [K_{\text{poly}}(x_i, x_j)]_{i,j=1}^n$

and

$$\begin{aligned}\phi(x_{\text{new}})^T \Phi^T &= \phi(x_{\text{new}})^T \begin{bmatrix} \phi(x_1) \\ \vdots \\ \phi(x_n) \end{bmatrix} \\ &= [K(x_{\text{new}}, x_1), \dots, K(x_{\text{new}}, x_n)]\end{aligned}$$

so our prediction is

$$\begin{aligned}f(x_{\text{new}}) &= [K(x_{\text{new}}, x_1), \dots, K(x_{\text{new}}, x_n)] \underbrace{(\underbrace{K + n\lambda I}^{:= \alpha})^{-1} y}_{:= \alpha} \\ &= \sum_{i=1}^n \alpha_i K(x_{\text{new}}, x_i)\end{aligned}$$

Kernel Learning Paradigm

- Choose our kernel K

- This corresponds to a function space (RKHS)

$$\mathcal{H} = \overline{\text{span} \{ K(\cdot, x) : x \in \mathbb{R}^d \}}$$

$$= \left\{ \sum_{i=1}^n \alpha_i K(\cdot, x_i) : x_i \in \mathbb{R}^d, \alpha_i \in \mathbb{R} \right\}$$

This space can be very large.

- Do ERM on this function space

$$f^* = \underset{f \in \mathcal{H}}{\operatorname{argmin}} \quad \frac{1}{n} \sum_{i=1}^n \ell(f(x_i), y_i) + \lambda \|f\|_{\mathcal{H}}^2$$

- this is CVX.

- practically we have a Representer Theorem

The ERM

$$f^* = \arg \min_{f \in \mathcal{H}} \frac{1}{n} \sum_{i=1}^n \ell(f(x_i), y_i) + \frac{\lambda}{2} \|f\|_{\mathcal{H}}^2$$

has a solution of the form

$$f^*(x) = \sum_{i=1}^n K(x, x_i) \alpha_i$$

where K is the kernel associated with $\mathcal{H}$.

Drawbacks:

- dealing w/ $K \in \mathbb{R}^{n \times n}$ can be expensive when $n \gg 1$
- applying f requires you to store the training data set.

Mitigating approaches:

- Nyström approximation
- Random feature maps (Recht & Rahimi)

- because of the representer theorem we know
in particular

$$f^*(x_i) = \sum_{j=1}^n K(x_i, x_j) \alpha_j = [K \alpha^*]_i$$

so we can fit f^* by finding α^*

$$\alpha^* = \underset{\alpha \in \mathbb{R}^n}{\operatorname{argmin}} \quad \frac{1}{n} \sum_{i=1}^n \ell([K \alpha]_i, y_i) + \lambda \alpha^T K \alpha$$