

Kernel functions

Linear predictors may potentially suffer from a large approximation error because they are always described by a number of coefficients which can not be larger than the number of features. A popular technique to reduce this bias is feature expansion, which adds new parameters by constructing new features through nonlinear combinations of the base features. Formally, this can be viewed in terms of a function $\phi : \mathbb{R}^d \rightarrow \mathcal{H}$ mapping data points $\mathbf{x} \in \mathbb{R}^d$ to a higher-dimensional space $\mathcal{H}$. By training a linear predictor on a feature-expanded training set, one actually learns a more complex nonlinear predictor in the original space.

For example, consider the quadratic feature-expansion map $\phi : \mathbb{R}^2 \rightarrow \mathbb{R}^6$ defined by $\phi(x_1, x_2) = (1, x_1^2, x_2^2, x_1, x_2, x_1x_2)$. Recall that a homogeneous hyperplane in $\mathbb{R}^6$ with coefficients given by $\mathbf{w} = (w_1, \dots, w_6)$ is the set of points $\{\mathbf{z} \in \mathbb{R}^6 : \mathbf{w}^\top \mathbf{z} = 0\}$. Now note that $\mathbf{w}^\top \phi(\mathbf{x}) = w_1 + w_2x_1^2 + w_3x_2^2 + w_4x_1 + w_5x_2 + w_6x_1x_2$. Sets of the form $\{\mathbf{x} \in \mathbb{R}^2 : w_1 + w_2x_1^2 + w_3x_2^2 + w_4x_1 + w_5x_2 + w_6x_1x_2 = 0\}$ describe second-degree curves on the plane $\mathbb{R}^2$. These surfaces include ellipses, parabolas, and hyperbolas.

In general, we may consider polynomial feature expansion maps $\phi : \mathbb{R}^d \rightarrow \mathcal{H}$, where $\mathcal{H} \equiv \mathbb{R}^N$, that use monomial features of the form $\phi(\mathbf{x})_{\mathbf{k}} = \prod_{s=1}^d x_s^{k_s}$ where

$$\mathbf{k} \in \mathcal{M}(d, n) = \left\{ (k_1, \dots, k_d) \in \mathbb{N}^d : k_1 + \dots + k_d \leq n \right\}$$

is the set of monomial feature indices (the previous example is a special case for $d = 2$ and $n = 2$). Now consider the classifier $h : \mathbb{R}^d \rightarrow \{-1, 1\}$ defined by

$$h(\mathbf{x}) = \text{sgn}(\mathbf{w}^\top \phi(\mathbf{x})) \quad \text{where} \quad \mathbf{w}^\top \phi(\mathbf{x}) = \sum_{\mathbf{k} \in \mathcal{M}(d, n)} w_{\mathbf{k}} \phi(\mathbf{x})_{\mathbf{k}}$$

This nonlinear classifier in $\mathbb{R}^d$ corresponds to a linear classifier in the space $\mathbb{R}^N$ where

$$N = |\mathcal{M}(d, n)| = \sum_{k=0}^n \binom{d+k-1}{k} = \binom{d+n}{n}$$

This implies that N is exponential in the degree n , and computing ϕ becomes infeasible even for moderately large n .

This computational barrier can be fully sidestepped using the so-called **kernel trick**, which can be applied to many algorithms for learning linear predictors. For example, recall the Perceptron update rule $\mathbf{w}_{t+1} = \mathbf{w}_t + y_t \mathbf{x}_t$ where $\mathbf{w}_1 = \mathbf{0}$. Then, linear classifiers learned through the Perceptron are of the form

$$h(\mathbf{x}) = \text{sgn} \left(\sum_{s \in S} y_s \mathbf{x}_s^\top \mathbf{x} \right)$$

where S is the set of indices s of training examples $(\mathbf{x}_s, y_s)$ on which the Perceptron made an update. If we run the Perceptron in the space $\mathbb{R}^N$, the linear classifier h becomes

$$h_\phi(\mathbf{x}) = \text{sgn} \left(\sum_{s \in S} y_s \phi(\mathbf{x}_s)^\top \phi(\mathbf{x}) \right) .$$

Hence, in order to compute $h_\phi(\mathbf{x})$ quickly we need a way of computing quickly each term $\phi(\mathbf{x}_s)^\top \phi(\mathbf{x})$. The kernel trick helps us find an efficiently computable kernel function $K : \mathbb{R}^d \times \mathbb{R}^d \rightarrow \mathbb{R}$ such that

$$K(\mathbf{x}, \mathbf{x}') = \phi(\mathbf{x})^\top \phi(\mathbf{x}') \quad \text{for all } \mathbf{x}, \mathbf{x}' \in \mathbb{R}^d . \quad (1)$$

For example, the quadratic kernel, corresponding to the quadratic feature-expansion map

$$\phi(x_1, x_2) = (1, x_1^2, x_2^2, \sqrt{2} x_1, \sqrt{2} x_2, \sqrt{2} x_1 x_2)$$

is $K(\mathbf{x}, \mathbf{x}') = (1 + \mathbf{x}^\top \mathbf{x}')^2$, as one can easily verify (the presence of the $\sqrt{2}$ coefficients, which is needed for the math, does not change the class of linear predictors that are representable using the mapping).

Given a kernel K , we can then write the linear classifier generated by the Perceptron as

$$h_K(\mathbf{x}) = \text{sgn} \left(\sum_{s \in S} y_s K(\mathbf{x}_s, \mathbf{x}) \right) . \quad (2)$$

Below here, we give the pseudo-code of the Kernel Perceptron algorithm (run on a stream of data points).

Algorithm: Kernel Perceptron

Let S be the empty list.

For all $t = 1, 2, \dots$

1. Get next example $(\mathbf{x}_t, y_t)$

2. Compute $\hat{y}_t = \text{sgn} \left(\sum_{s \in S} y_s K(\mathbf{x}_s, \mathbf{x}_t) \right)$

3. If $\hat{y}_t \neq y_t$ add t to the list S

The polynomial kernel $K_n(\mathbf{x}, \mathbf{x}') = (1 + \mathbf{x}^\top \mathbf{x}')^n$ for all $n \in \mathbb{N}$ generalizes the quadratic kernel defined earlier. Using Newton's binomial theorem, we can explicitly compute the map ϕ such that $K_n(\mathbf{x}, \mathbf{x}') = \phi(\mathbf{x})^\top \phi(\mathbf{x}')$,

$$(1 + \mathbf{x}^\top \mathbf{x}')^n = \sum_{i=0}^n \binom{n}{i} (\mathbf{x}^\top \mathbf{x}')^i . \quad (3)$$

Now observe that, using the multinomial theorem,

$$(\mathbf{x}^\top \mathbf{x}')^i = \left(\sum_{j=1}^d x_j x'_j \right)^i = \sum_{\substack{\mathbf{k} \in \mathbb{N}^d \\ k_1 + \dots + k_d = i}} c_{\mathbf{k}} \left(\prod_{s=1}^d x_s^{k_s} (x_s^{k_s})' \right)$$

where $c_{\mathbf{k}}$ are multinomial coefficients. Therefore,

$$\phi(\mathbf{x}) = \left(\sqrt{c'_{\mathbf{k}}} \prod_{s=1}^d x_s^{k_s} \right)_{\mathbf{k} \in \mathcal{M}(d,n)} \quad (4)$$

where $c'_{\mathbf{k}}$ is a numerical coefficient resulting from the combination of the binomial theorem with the multinomial theorem. In other words, the feature map $\phi : \mathbb{R}^d \rightarrow \mathbb{R}^N$ associated with the polynomial kernel K_n sends each $\mathbf{x} \in \mathbb{R}^d$ to a vector whose components are indexed by all monomials of degree at most n and weighted by roots of binomial and multinomial coefficients.

Another type of kernel is the Gaussian kernel,

$$K_{\gamma}(\mathbf{x}, \mathbf{x}') = \exp \left(-\frac{1}{2\gamma} \|\mathbf{x} - \mathbf{x}'\|^2 \right) \quad \gamma > 0 .$$

In order to derive the map ϕ_{γ} associated with K_{γ} we proceed as follows,

$$\begin{aligned} \exp \left(-\frac{1}{2\gamma} \|\mathbf{x} - \mathbf{x}'\|^2 \right) &= \exp \left(-\frac{1}{2\gamma} (\|\mathbf{x}\|^2 + \|\mathbf{x}'\|^2) \right) \exp \left(\frac{1}{\gamma} (\mathbf{x}^{\top} \mathbf{x}') \right) \\ &= \exp \left(-\frac{\|\mathbf{x}\|^2}{2\gamma} \right) \exp \left(-\frac{\|\mathbf{x}'\|^2}{2\gamma} \right) \sum_{n=0}^{\infty} \frac{1}{n!} \frac{(\mathbf{x}^{\top} \mathbf{x}')^n}{\gamma^n} \end{aligned} \quad (5)$$

where we used the Taylor series expansion $e^x = 1 + x + \frac{x^2}{2!} + \dots$. A closer look at (5) reveals that the Gaussian kernel is a linear combination of infinitely many polynomial kernels (3) of increasing degree, each weighted by the reciprocal of the factorial of its degree. The parameter γ is a scaling factor for the products $\mathbf{x}^{\top} \mathbf{x}'$. Finally, note that $K_{\gamma}(\mathbf{x}, \mathbf{x}) = 1$ for all $\mathbf{x} \in \mathbb{R}^d$ because $\|\mathbf{x} - \mathbf{x}\|^2 = 0$.

Note that predictors of the form (2) that utilize Gaussian kernels predict any point $\mathbf{x}$ using a linear combination (with coefficients y_t) of Gaussians $e^{-z^2/(2\gamma)}$ centered on $\mathbf{x}_s$ for $s \in S$ and evaluated at $\mathbf{x}$.

The Gaussian kernel is *universal* in the following sense: for each $\gamma > 0$, for each continuous function $f : \mathbb{R}^d \rightarrow \mathbb{R}$, and for each $\varepsilon > 0$, there exists a function g of the form

$$g = \sum_{i=1}^N \alpha_i K_{\gamma}(\mathbf{x}_i, \cdot)$$

for some $\mathbf{x}_1, \dots, \mathbf{x}_N \in \mathcal{X}$, $\alpha_1, \dots, \alpha_N \in \mathbb{R}$, and $N \in \mathbb{N}$ such that¹

$$\sup_{\mathbf{x} \in \mathcal{X}} |f(\mathbf{x}) - g(\mathbf{x})| \leq \varepsilon$$

An important consequence of this fact is that learning algorithms that use Gaussian kernels can be consistent (that is, the expected risk of their predictors converges to the Bayes risk as the sample size grows to infinity).

Given a data space $\mathcal{X}$ (not necessarily $\mathbb{R}^d$) and a symmetric function $K : \mathcal{X} \times \mathcal{X} \rightarrow \mathbb{R}$ how can we check whether K is a kernel? In other words, we want to know whether there exists a feature map

¹Technically, we need $\mathcal{X}$ to be a compact subset of $\mathbb{R}^d$.

$\phi_K : \mathcal{X} \rightarrow \mathcal{H}_K$, where $\mathcal{H}_K$ is a linear space with inner product $\langle \cdot, \cdot \rangle_K$, such that $\langle \phi_K(\mathbf{x}), \phi_K(\mathbf{x}') \rangle_K = K(\mathbf{x}, \mathbf{x}')$ for all $\mathbf{x}, \mathbf{x}' \in \mathcal{X}$. The space $\mathcal{H}_K$ must be linear because our algorithms build linear predictors using additive updates. Luckily, there is a very simple characterization: K is a kernel if and only if for all $m \in \mathbb{N}$ and for all $\mathbf{x}_1, \dots, \mathbf{x}_m \in \mathcal{X}$, the $m \times m$ matrix $\mathbf{K}$ such that $\mathbf{K}_{i,j} = K(\mathbf{x}_i, \mathbf{x}_j)$ is positive semidefinite. That is, $\mathbf{z}^\top \mathbf{K} \mathbf{z} \geq 0$ for all $\mathbf{z} \in \mathbb{R}^m$.

The above result tells us the conditions under which ϕ_K exists for a given K . On the other hand, it does not tell us how ϕ_K looks like. In fact, there is no unique representation of the pair $(\phi_K, \mathcal{H}_K)$ for a given kernel K . However, it can be shown that if K is a kernel, then we can always represent $\phi_K : \mathcal{X} \rightarrow \mathcal{H}_K$ as $\phi_K(\mathbf{x}) = K(\mathbf{x}, \cdot)$. Note that, in this representation, $\phi_K(\mathbf{x})$ is a function from $\mathcal{X}$ to $\mathbb{R}$. Now, since $\mathcal{H}_K$ is a linear space, it contains all linear combinations of $K(\mathbf{x}, \cdot)$ for arbitrary choices of the coefficients and of the points $\mathbf{x} \in \mathcal{X}$. One can indeed show that any element of $\mathcal{H}_K$ can be written in this way,

$$\mathcal{H}_K \equiv \left\{ \sum_{i=1}^N \alpha_i K(\mathbf{x}_i, \cdot) : \mathbf{x}_1, \dots, \mathbf{x}_N \in \mathcal{X}, \alpha_1, \dots, \alpha_N \in \mathbb{R}, N \in \mathbb{N} \right\}$$

An element $f \in \mathcal{H}_K$ is then any function $f : \mathcal{X} \rightarrow \mathbb{R}$ such that

$$f(\mathbf{x}) = \sum_{i=1}^N \alpha_i K(\mathbf{x}_i, \mathbf{x})$$

for some $\mathbf{x}_1, \dots, \mathbf{x}_N \in \mathcal{X}$, $\alpha_1, \dots, \alpha_N \in \mathbb{R}$, and $N \in \mathbb{N}$.

A learning algorithm producing linear predictors $g \in \mathcal{H}_K$ becomes nonparametric when K is the Gaussian kernel and N (the number of terms in the sum defining any $g \in \mathcal{H}_K$) is free to grow unbounded as the sample size increases. If, instead, K is a kernel such that ϕ_K maps $\mathcal{X}$ to a finite dimensional space, then any $g \in \mathcal{H}_K$ can be always represented with a fixed number of parameters, and therefore any algorithm choosing predictors from $\mathcal{H}_K$ is parametric.

For linear predictors $\mathbf{w} \in \mathbb{R}^d$, we compute $\mathbf{w}^\top \mathbf{x}$ using the standard notion of Euclidean inner product. For linear predictors $f \in \mathcal{H}_K$, instead, we compute $\langle f, \phi_K(\mathbf{x}) \rangle_K$ using the inner product for $\mathcal{H}_K$. Now, since ϕ_K is the feature map for K , we must have $\langle \phi_K(\mathbf{x}), \phi_K(\mathbf{x}') \rangle_K = K(\mathbf{x}, \mathbf{x}')$. Thus, recalling that the inner product is a bilinear operator, we can write

$$\langle f, \phi_K(\mathbf{x}) \rangle_K = \sum_{i=1}^N \alpha_i \langle K(\mathbf{x}_i, \cdot), \phi_K(\mathbf{x}) \rangle_K = \sum_{i=1}^N \alpha_i \langle \phi_K(\mathbf{x}_i), \phi_K(\mathbf{x}) \rangle_K = \sum_{i=1}^N \alpha_i K(\mathbf{x}_i, \mathbf{x}) = f(\mathbf{x})$$

The equality $\langle f, \phi_K(\mathbf{x}) \rangle_K = f(\mathbf{x})$ is known as *reproducing property*. For this reason, $\mathcal{H}_K$ (or, more precisely, an appropriate completion of $\mathcal{H}_K$) is also known as *reproducing kernel Hilbert space* (RKHS).

Next, we see how the inner product $\langle f, g \rangle_K$ between two arbitrary $f, g \in \mathcal{H}_K$ is computed, where

$$f = \sum_{i=1}^N \alpha_i K(\mathbf{x}_i, \cdot) \quad \text{and} \quad g = \sum_{j=1}^M \beta_j K(\mathbf{x}'_j, \cdot)$$

Using once more the bilinearity of inner products,

$$\begin{aligned}
\langle f, g \rangle_K &= \left\langle \sum_{i=1}^N \alpha_i K(\mathbf{x}_i, \cdot), \sum_{j=1}^M \beta_j K(\mathbf{x}'_j, \cdot) \right\rangle_K \\
&= \sum_{i=1}^N \alpha_i \left\langle K(\mathbf{x}_i, \cdot), \sum_{j=1}^M \beta_j K(\mathbf{x}'_j, \cdot) \right\rangle_K \\
&= \sum_{i=1}^N \sum_{j=1}^M \alpha_i \beta_j \langle K(\mathbf{x}_i, \cdot), K(\mathbf{x}'_j, \cdot) \rangle_K \\
&= \sum_{i=1}^N \sum_{j=1}^M \alpha_i \beta_j K(\mathbf{x}_i, \mathbf{x}'_j)
\end{aligned}$$

We can lift to any RKHS the bounds we derived for learning algorithms that can be run with kernels. For instance recall the bound on the number of mistakes provided by the Perceptron convergence theorem,

$$\|\mathbf{u}\|^2 \left(\max_t \|\mathbf{x}_t\|^2 \right)$$

which holds for any $\mathbf{u} \in \mathbb{R}^d$ such that $y_t \mathbf{u}^\top \mathbf{x}_t \geq 1$ for $t = 1, \dots, m$.

In a generic RKHS $\mathcal{H}_K$, the linear separator $\mathbf{u}$ is some $g \in \mathcal{H}_K$ such that $y_t g(\mathbf{x}_t) \geq 1$ for $t = 1, \dots, m$. The squared norm $\|\mathbf{x}_t\|^2 = \mathbf{x}_t^\top \mathbf{x}_t$ becomes $\|\phi_K(\mathbf{x})\|_K^2 = \langle K(\mathbf{x}, \cdot), K(\mathbf{x}, \cdot) \rangle_K = K(\mathbf{x}, \mathbf{x})$. Finally $\|\mathbf{u}\|^2$ is replaced by

$$\|f\|_K^2 = \left\| \sum_{i=1}^N \alpha_i K(\mathbf{x}_i, \cdot) \right\|_K^2 = \left\langle \sum_{i=1}^N \alpha_i K(\mathbf{x}_i, \cdot), \sum_{j=1}^N \alpha_j K(\mathbf{x}_j, \cdot) \right\rangle_K = \sum_{i,j=1}^N \alpha_i \alpha_j K(\mathbf{x}_i, \mathbf{x}_j) .$$

A more complex linear predictor that can be kernelized is ridge regression, $\mathbf{w} = (\alpha I_d + X^\top X)^{-1} X^\top \mathbf{y}$, where X is the $m \times d$ matrix whose rows are the training points $\mathbf{x}_1, \dots, \mathbf{x}_m \in \mathbb{R}^d$ and $\mathbf{y} = (y_1, \dots, y_m)$ is the vector of training labels $y_t \in \mathbb{R}$. Using the identity

$$(\alpha I_d + X^\top X)^{-1} X^\top = X^\top (\alpha I_m + X X^\top)^{-1} \quad (6)$$

we can represent the ridge regression predictor in a generic RKHS by $\mathbf{y}^\top (\alpha I + \mathbf{K})^{-1} \mathbf{k}(\cdot)$, where $\mathbf{K}$ is the $m \times m$ matrix with entries $\mathbf{K}_{i,j} = K(\mathbf{x}_i, \mathbf{x}_j)$ and $\mathbf{k}(\cdot)$ is the vector $(K(\mathbf{x}_1, \cdot), \dots, K(\mathbf{x}_m, \cdot))$ of functions $K(\mathbf{x}_t, \cdot) = \langle \phi_K(\mathbf{x}_t), \phi_K(\cdot) \rangle_K$. Indeed, using (6),

$$\mathbf{w}^\top \mathbf{x} = \mathbf{y}^\top X (\alpha I + X^\top X)^{-1} \mathbf{x} = \mathbf{y}^\top (\alpha I + X X^\top)^{-1} X \mathbf{x}$$

Now, the elements of $X X^\top$ are $\mathbf{x}_i^\top \mathbf{x}_j$, that in kernel space become $K(\mathbf{x}_i, \mathbf{x}_j)$. Similarly, the components of $X \mathbf{x}$ are $\mathbf{x}_i^\top \mathbf{x}$, that correspond to $K(\mathbf{x}_i, \mathbf{x})$ in kernel space. Hence, the ridge regression prediction $\mathbf{w}^\top \mathbf{x} = \mathbf{y}^\top X (\alpha I + X^\top X)^{-1} \mathbf{x}$ in kernel space becomes $\langle g, \phi_K(\mathbf{x}) \rangle_K = \mathbf{y}^\top (\alpha I + \mathbf{K})^{-1} \mathbf{k}(\mathbf{x})$.

Linear predictors in RKHS can incur overfitting. For example, by increasing the degree n of a polynomial kernel K_n we reduce the training error because higher-degree curves can be used to

separate the training points. If the degree is too high, the predictor will overfit. A similar reasoning applies to Gaussian kernels K_γ . The γ parameter corresponds to the width of the Gaussians centered on training points $\mathbf{x}_s$. If γ is small relative to the typical squared distances $\|\mathbf{x}_s - \mathbf{x}\|^2$ between training and test points, then the classification of a test point $\mathbf{x}$ is essentially determined by the training point closest to it. This implies a training error equal or close to zero, because —similarly to 1-NN classifiers— the training points are never misclassified. Once again, the resulting predictor is likely to overfit. On the other hand, for values of γ that are large with respect to the squared distances $\|\mathbf{x}_s - \mathbf{x}\|^2$, the Gaussians centered on training points are very wide, and the resulting predictors are similar to k -NN classifiers when k is chosen close to the training set size, which is likely to cause underfitting.

We established that any symmetric function $K : \mathcal{X} \times \mathcal{X} \rightarrow \mathbb{R}$ is a kernel if and only if the kernel matrix $\mathbf{K}$ is positive semidefinite. This is an important result, because it holds irrespective to the choice of the data space $\mathcal{X}$. We can therefore define kernels on any set $\mathcal{X}$: be it a set of matrices, sequences, trees, graphs, and so on. Kernels can be viewed as a way of encapsulating the data space, offering a uniform interface to a learning algorithm that can be efficiently run in the corresponding RKHS. In order to guide the intuition when designing a kernel function on a given data space $\mathcal{X}$, one should recall that $K(\mathbf{x}, \mathbf{x}')$ implements an inner product between $K(\mathbf{x}, \cdot)$ and $K(\mathbf{x}', \cdot)$ in some RKHS. We can then first define a notion of similarity on $\mathcal{X}$, and then adjust it so that arbitrary kernel matrices are always positive semidefinite.