

Lecture 10: December 26, 2010 (SVM)

*Lecturer: Yishay Mansour**Scribe: Aviad Rubinstein and Uri Velelviski*

10.1 Kernels

Here we will see the advantage of the dual program, of using a number of variables proportional to the number of inputs rather than the number of features.

10.1.1 The Idea

Say we want to use SVM on data which is not linearly separable. If we could map from our initial dimension to a larger dimension, maybe in this new space we can find a linear separation for our data.

10.1.2 Example

Consider the linear separator

$$\text{sign}\left(\sum_{i=1}^n x_i \alpha_i\right).$$

Instead we can use a quadratic function

$$\text{sign}\left(\sum_{i,j} x_i x_j \alpha_{i,j}\right).$$

Our transformation:

$$(x_1 \dots x_n) \mapsto (x_1 \dots x_n, x_1^2 \dots x_n^2, x_1 x_2 \dots x_{n-1} x_n).$$

So we moved from n -dimensional space to $n + n^2$ dimensional space. We could do the same for degree 3 etc.

10.1.3 The problem and its solution

Once we move to a higher dimension, all the calculations become more complex. In fact, the complexity increases exponentially.

Definition We define a kernel as the inner products between two $\phi(\cdot)$:

$$K(x, z) = \phi(x) \cdot \phi(z)$$

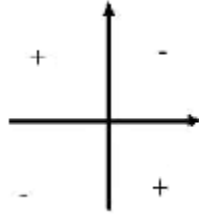


Figure 10.1: XOR

The idea is that in some cases we can calculate $K(x, z)$ without calculating $\phi(x)$, thus avoiding the added complexity.

10.1.4 Example 1 - Polynomials

Assume $n = 2$.

$$K(x, z) = (x \cdot z)^2 = (x_1 z_1 + x_2 z_2)^2 = x_1^2 z_1^2 + x_2^2 z_2^2 + 2x_1 x_2 z_1 z_2 = \underbrace{(x_1^2, x_2^2, \sqrt{2}x_1 x_2)}_{\phi(x)} \cdot \underbrace{(z_1^2, z_2^2, \sqrt{2}z_1 z_2)}_{\phi(z)}$$

We can see that it is possible to calculate $K(x, z)$ without actually calculating $\phi(x)$ or $\phi(z)$, and that makes the difference in complexity ($O(n)$ instead of $O(n^2)$).

Similarly, $K_d(x, z) = (x \cdot z + c)^d$ is mapped into $\phi(\cdot)$ with $\binom{n+d}{d} = O(n^d)$ entries, but calculating K is only $O(n)$.

10.1.5 Example 2 - XOR

XOR of two variables is obviously not linearly separable. Let's look at a general degree-2 kernel:

$$K_2(x, y) = (x \cdot y + c) = \begin{pmatrix} x_1^2 \\ x_2^2 \\ \sqrt{2}x_1 x_2 \\ \sqrt{2}c x_1 \\ \sqrt{2}c x_2 \\ c \end{pmatrix} \cdot \begin{pmatrix} y_1^2 \\ y_2^2 \\ \sqrt{2}y_1 y_2 \\ \sqrt{2}c y_1 \\ \sqrt{2}c y_2 \\ c \end{pmatrix}$$

If $c = 0$ we are left with only the top three components:

$$K_2(x, y) = \begin{pmatrix} x_1^2 \\ x_2^2 \\ \sqrt{2}x_1x_2 \end{pmatrix} \cdot \begin{pmatrix} y_1^2 \\ y_2^2 \\ \sqrt{2}y_1y_2 \end{pmatrix}$$

XOR representation will be:

$$\phi(x) \cdot \underbrace{\begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix}}_w = \sqrt{2}x_1x_2 \geq 0$$

So we have the weight vector w in the higher dimension, and we want to build it from examples in the lower dimension.

<i>label</i>	x_1	x_2	x_1^2	x_2^2	x_1x_2
+1	+1	+1	+1	+1	+1
-1	+1	-1	+1	+1	-1
-1	-1	+1	+1	+1	-1
+1	-1	-1	+1	+1	+1

$$w = +1 \cdot \phi(+1, +1) - 1 \cdot \phi(+1, -1) - 1 \cdot \phi(-1, +1) + 1 \cdot \phi(-1, -1)$$

$$\begin{aligned} h(x) &= \text{sign}(K_2((+1, +1), x) + K_2((-1, -1), x) - K_2((+1, -1), x) - K_2((-1, +1), x)) \\ &= \text{sign}\left(\phi(x) \cdot \begin{pmatrix} 0 \\ 0 \\ 4\sqrt{2} \end{pmatrix}\right) \\ &= \text{sign}(4\sqrt{2}x_1x_2) \end{aligned}$$

10.1.6 Example 3 - Gaussian Kernel

Another example is the gaussian kernel, which maps the space into an infinite new space, but is also computable in linear complexity.

$$K(x, z) = \exp\left(-\frac{\|x - z\|^2}{2\sigma^2}\right)$$

10.1.7 Example 4 - Strings Kernel

We can look at kernels also for strings, such as:

$$\begin{aligned} x, z &\in \{0, 1\}^n \\ K(x, z) &= |\{\text{common substrings}\}| \end{aligned}$$

10.1.8 To sum it up

The vector $\phi(x)$ can be thought of as the features that we learn in the new space. If the learning algorithm (SVM in this case) can be written in a way that uses only the inner products of the data, then it is possible to “kernelize” it.

10.1.9 Closure Properties of Kernel Functions

1. Closure under the following actions:

(a) Adding a constant

$$K(x, z) = K_1(x, z) + c,$$

by adding another entry to $\phi(x)$ with value $\sqrt{c}$.

(b) Multiplication by constant:

$$K(x, z) = c \cdot K_1(x, z),$$

by setting $\phi(x) = c \cdot \phi_1(x)$.

(c) Multiplication of two kernel functions:

$$K(x, z) = K_1(x, z) \cdot K_2(x, z)$$

To see this, let l_1 be the length of ϕ_1 and l_2 be the length of ϕ_2 . Define $\phi(x)$ of length $l_1 l_2$, such that the (i, j) entry in $\phi(x)$ is $\phi_{1,i}(x) \phi_{2,j}(x)$ where $\phi_{1,i}(x)$ is the i -th entry of $\phi_1(x)$.

2. Let K be a legitimate kernel function. Then it is possible to define a kernel matrix as

follows: $K_{ij} = K(x^{(i)}, x^{(j)})$. The result matrix K will be symmetric. For any z :

$$\begin{aligned}
 z^T K z &= \sum_i \sum_j z_i K_{ij} z_j \\
 &= \sum_i \sum_j z_i \phi(x^{(i)})^T \cdot \phi(x^{(j)}) z_j \\
 &= \sum_i \sum_j z_i \sum_k \phi_k(x^{(i)}) \phi_k(x^{(j)}) z_j \\
 &= \sum_k \sum_i \sum_j z_i \phi_k(x^{(i)}) \phi_k(x^{(j)}) z_j \\
 &= \sum_k \left(\sum_i z_i \phi_k(x^{(i)}) \right)^2 \geq 0
 \end{aligned}$$

This means that if K is a kernel then its matrix representation is positive-semi-definite (PSD). We have not shown the other direction, but it is also true:

Theorem 10.1 (Mercer) K is a legitimate kernel function $\leftrightarrow K$ is PSD.

(Note: For our needs it is not required to calculate $\phi(\cdot)$. It's enough that it exists.)
The following theorem characterizes the solution when we have kernel functions.

Theorem 10.2 (The Representaion Theorem) *Let K be a legitimate kernel function. H is the function space (RKHS). G is monotone increasing function. Then the solution to the maximization problem*

$$\arg \min_{h \in H} F(h) = \arg \min_{h \in H} G(\|h\|_H^2) + l(h(x_1) \dots h(x_m))$$

Is of the form:

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

Theorem 10.3 (Reminder about SVM)

The solution to the maximization problem

$$\max_{\alpha} \sum_{i=1}^m \alpha_i - \frac{1}{2} \sum_{i,j=1}^m \alpha_i \alpha_j y^{(i)} y^{(j)} K(x^{(i)}, x^{(j)}),$$

where $C \geq \alpha_i \geq 0$ and $\sum_{i=1}^m \alpha_i y^{(i)} = 0$, is of the following form:

$$h(x) = \text{sign} \left(\sum_{i=1}^m \alpha_i y^{(i)} K(x^{(i)}, x^{(j)}) + b \right).$$

10.2 Dealing with non linearly-separable data

10.2.1 The Idea

What happens if the data that we have is not linearly separable with margin 1? (Which is a common event)

Every hyperplane (w, b) has a point x_i such that $y_i(w \cdot x_i + b) < 1$. Otherwise the data would be linearly separable with margin 1. The idea is “to soften” the hard constraint (“hard margin”) and turn it into a soft constraint (“soft margin”). We would want to minimize the violations of the constraints, so we’ll change the constraint such that $y_i(w \cdot x_i + b) \geq 1 - \psi_i$ where $\psi_i \geq 0$. Now there is always a solution. So what exactly do we want to minimize?

$$\min_{w, b, \psi} \underbrace{\frac{1}{2} \|w\|^2}_{\text{margin}} + \underbrace{C}_{\text{cost}} \cdot F(\psi_1 \dots \psi_m),$$

where C is an important parameter which determines the tradeoff between accuracy and margin size. In other words - tradeoff between generalization capability (small $\|w\|$) and the classification error on the training data.

The natural choice is to calculate $F(\psi_1 \dots \psi_m)$ individually for every point and then sum everything up:

$$F(\psi_1 \dots \psi_m) = \sum F(\psi_i).$$

Ideally:

$$F_{01}(\psi) = \begin{cases} 1 & \psi > 0 \\ 0 & \psi = 0 \end{cases},$$

this will count the exact number of points that are not classified correctly by our classifier.

10.2.2 Optimization

This is a hard optimization problem (even with no margin). From a complexity point of view, even if it is guaranteed that a hyperplane exists such that it classifies 99% of the points correctly, finding a hyperplane which classifies 51% of the points correctly is an NP-complete problem (also note that F is not convex).

For this reason we write the following convex program:

$$\min_{w, b, \psi} \frac{1}{2} \|w\|^2 + C \sum \psi_i$$

Such that

$$y_i(w \cdot x_i + b) \geq 1 - \psi_i$$

$$\psi_i \geq 0.$$

For this convex program we get the following Lagrangian (α and β are the lagrange multipliers):

$$L(w, b, \psi, \alpha) = \frac{1}{2} \|w\|^2 + C \sum_{i=1}^m \psi_i - \sum_{i=1}^m \alpha_i (y_i (w \cdot x_i + b) - 1 + \psi_i) - \sum \beta_i \psi_i$$

Now we'll check the KKT conditions:

1. $\nabla L = 0$

$$(a) \nabla_w L = w - \sum \alpha_i y_i x_i = 0 \leftrightarrow w^* = \sum \alpha_i y_i x_i$$

$$(b) \nabla_b L = - \sum \alpha_i y_i = 0 \leftrightarrow \sum \alpha_i y_i = 0$$

$$(c) \nabla_{\psi_i} L = C - \alpha_i - \beta_i = 0 \leftrightarrow \alpha_i + \beta_i = c$$

2. Complementary

$$(a) \alpha_i [y_i (w \cdot x_i + b) - 1 + \psi_i] = 0 \rightarrow \alpha_i = 0 \text{ or } y_i (w \cdot x_i + b) = 1 - \psi_i$$

(when $y_i (w \cdot x_i + b) = 1 - \psi_i$ and $\alpha_i > 0$ this is a support vector, it can be classified incorrectly or be exactly on the margin)

$$(b) \beta_i \psi_i = 0 \rightarrow \beta_i = 0 \text{ or } \psi_i = 0$$

3. Non negativity

$$(a) \alpha_i \geq 0$$

$$(b) \beta_i \geq 0$$

$$(c) \psi_i \geq 0$$

Substitute these conditions back to the Lagrangian we get:

$$\begin{aligned} \frac{1}{2} \|w\|^2 &= \frac{1}{2} \sum \alpha_i \alpha_j y_i y_j (x_i \cdot x_j) \\ C \sum \psi_i - \sum \alpha_i \psi_i - \sum \beta_i \psi_i &= 0 \\ \sum \alpha_i [y_i (w \cdot x_i + b) - 1] &= \underbrace{\sum \alpha_i y_i (w^* \cdot x_i)}_{\sum_{i,j} \alpha_i \alpha_j y_i y_j (x_i \cdot x_j)} + \underbrace{\sum \alpha_i y_i b}_0 + \sum \alpha_i \end{aligned}$$

We have the following

$$L = -\frac{1}{2} \sum \alpha_i \alpha_j y_i y_j (x_i \cdot x_j) + \sum \alpha_i \quad (10.1)$$

We need to add the constraints. For the constraints $\beta_i \geq 0$ and $C \geq \beta_i + \alpha_i$ we add the constraint: $C \geq \alpha_i \geq 0$. We have in addition $\sum \alpha_i y_i = 0$. The solution:

$$h(x) = \text{sign} \left[\sum \alpha_i y_i (x_i \cdot x) + b \right],$$

where b is defined for the vector that adhere to $C > \alpha_i > 0$:

$$b = y_i - \sum_{j=1}^m \alpha_j y_j (x_j \cdot x_i).$$

If for point i we get $\psi_i > 0$, from the complementarity we have $\beta_i \psi_i = 0$ so this means $\beta_i = 0$. In the optimal solution from the KKT conditions we get $C = \alpha_i + \beta_i$ so eventually we get $\alpha_i = C$. Every point that is classified wrongly (either the margin is too small or negative), gets $\alpha_i = c$, so its effect on the decision hyperplane is limited to c . If a point is classified correctly, we get $\psi_i = 0$, then $C \geq \beta_i > 0$ and this means that $C > \alpha_i > 0$. For points that are classified correctly **and** are not support vectors we get $\alpha_i = 0$.

So again, large C means avoiding errors on the training data. Small C means more freedom in the hyperplane, and a better generalization ability.

10.3 SMO - Sequential Minimization Optimization

We shall now see a practical algorithm for solving the optimization problem (10.1). For practical Implementations, we need to add an accuracy parameter, say 0.001.

10.3.1 Single Variable Optimization

In order to optimize (10.1) we will start with a feasible but not optimal solution (e.g. $\alpha_i = 0 \forall i$) and look for local improvements. We cannot change only one variable because, given the values of $\alpha_j \forall j \neq i$ we have:

$$\begin{aligned} \alpha_i y_i &= - \sum_{j \neq i} \alpha_j y_j \\ \alpha_i &= -y_i \sum_{j \neq i} \alpha_j y_j \end{aligned}$$

(Since $y_i \in \{\pm 1\}$)

10.3.2 Two Variables Optimization

So we would like to change the values of two variables, WLOG α_1, α_2 . Thus, we need to maximize:

$$L(\alpha_1, \alpha_2) = \left(1 - y_1 \sum_{j \notin \{1,2\}} y_j \alpha_j K(x_1, x_j)\right) \alpha_1 + \left(1 - y_2 \sum_{j \notin \{1,2\}} y_j \alpha_j K(x_2, x_j)\right) \alpha_2 \\ - \left(\frac{1}{2} y_1^2 K(x_1, x_2)\right) \alpha_1^2 - (y_1 y_2 K(x_1, x_2)) \alpha_1 \alpha_2 - \left(\frac{1}{2} y_2^2 K(x_1, x_2)\right) \alpha_2^2 + \text{const.}$$

$$\text{s.t.} \quad 0 \leq \alpha_1, \alpha_2 \leq C, \\ \alpha_1 y_1 + \alpha_2 y_2 = Z = - \sum_{j \notin \{1,2\}} \alpha_j y_j$$

Substitute $\alpha_1 = y_1 (Z - \alpha_2 y_2)$ to get a quadratic:

$$L(\alpha_2) = a\alpha_2^2 + b\alpha_2 + c$$

We then need to find the value α_2^* that minimizes the L . If the value we get fits the constraints (i.e. $y_1(Z - \alpha_2^* y_2)$, $\alpha_2^* \in [0, C]$) we can use it. Otherwise, we use the best extreme values.

10.3.3 Calculating the minimum of a quadratic on a finite interval

To find the α_2^* that minimizes $L(\alpha_2)$ given the constraints, we need to calculate the constraints in terms of an upper bound UB and a lower bound LB to α_2 . Let $(\alpha_1^{old}, \alpha_2^{old})$ be the values before optimization, and let $(\alpha_1^{new}, \alpha_2^{new})$ be the values that minimize $L(\alpha_1, \alpha_2)$ without the constraints.

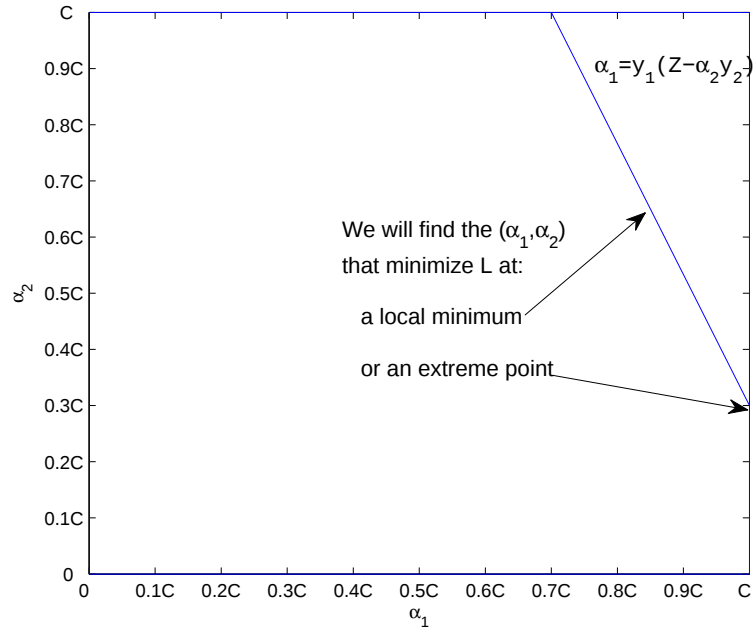
For example, if $y_1 = 1$ and $y_2 = -1$, then $Z = \alpha_1^{old} - \alpha_2^{old} = \alpha_1^{new} - \alpha_2^{new}$.

- From $\alpha_1^{new} \geq 0$ we have:

$$\alpha_2^{new} = \alpha_1^{new} - Z \geq \alpha_2^{old} - \alpha_1^{old}$$

- From $\alpha_1^{new} \leq C$ we have:

$$\alpha_2^{new} = \alpha_1^{new} - Z \leq \alpha_2^{old} - \alpha_1^{old} + C$$

Figure 10.2: Constraints on (α_1, α_2)

So we set:

$$\begin{aligned} UB &= \max \{C, \alpha_2^{old} - \alpha_1^{old} + C\} \\ LB &= \min \{0, \alpha_2^{old} - \alpha_1^{old}\} \end{aligned}$$

To get:

$$\alpha_2^* = \begin{cases} UB & \alpha_2^{new} \geq UB \\ \alpha_2^{new} & LB < \alpha_2^{new} < UB \\ LB & \alpha_2^{new} \leq LB \end{cases}$$

10.3.4 Choosing α_1, α_2

A heuristic for choosing which α_i, α_j to update. We look at the KKT conditions:

$$\begin{aligned} \alpha_i = 0 &\implies h(x_i)y_i \geq 1 \\ 0 < \alpha_i < C &\implies h(x_i)y_i = 1 \\ \alpha_i = C &\implies h(x_i)y_i \leq 1 \end{aligned}$$

Recall that for an optimal solution, all the KKT conditions are satisfied. In practice, we need to introduce an accuracy parameter, say $\epsilon = 0.001$. We will only expect the KKT conditions to be satisfied to within ϵ . To increase the speed of convergence, we will prefer to update α_i 's which are significantly far (i.e. ϵ -far) from satisfying the KKT conditions.

10.3.5 Calculating b

At each iteration, to calculate b we shall use:

$$\begin{aligned} b_1 &= y_1 - \sum \alpha_j y_j K(x_j, x_1) \\ b_2 &= y_2 - \sum \alpha_j y_j K(x_j, x_2) \\ b &= -\frac{b_1 + b_2}{2} \end{aligned}$$