

Classifiers that attempt to separate instances of two classes by hyperplanes in the original linear or a special high-dimensional space while maximizing the margin of separation between the classes.

$$y = \{+1, -1\}$$

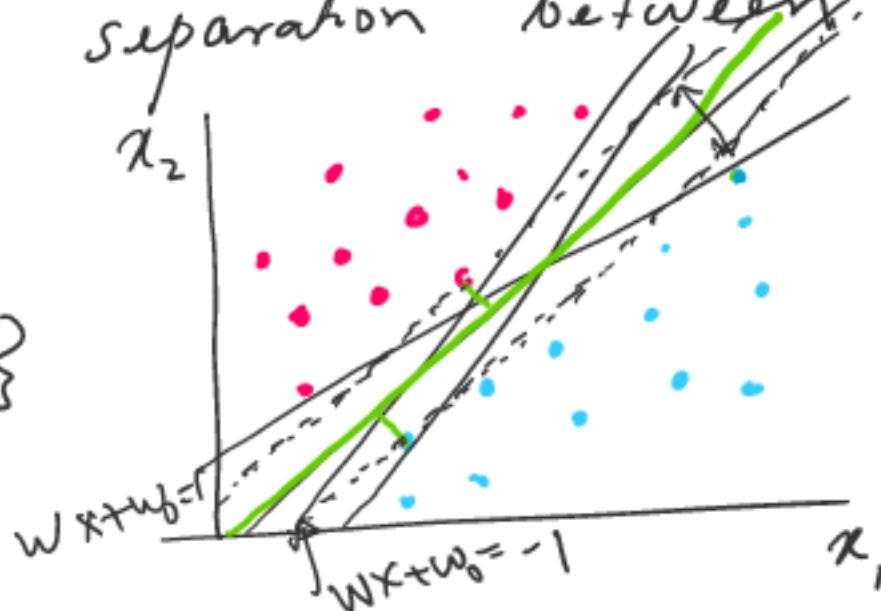
$$D = \{(x^i, y^i) : i = 1 \dots N\}$$

Hyperplanes in linear space:

$$w x + w_0 \equiv w_1 x_1 + w_2 x_2 + \dots + w_d x_d + w_0 = 0 \quad d=2, k=2$$

1) Separable case

no overlapping points



1) given

W s.t. $\forall$ positive points

$$\left. \begin{aligned} W^T x^i + w_0 &\geq 1 \quad \forall i : y_i = +1 \\ W^T x^i + w_0 &\leq -1 \quad \forall i : y_i = -1 \end{aligned} \right\} i=1 \dots N$$

What is the geometric distance between

$$W^T x^i + w_0 - 1 = 0 \quad \& \quad W^T x^i + w_0 + 1 = 0 \quad \text{is} \quad \frac{2}{\|W\|}$$

$$\max_{W, w_0} \frac{2}{\|W\|} \quad \text{s.t.}$$

$$y_i (W^T x^i + w_0) \geq 1 \quad \forall i$$

$$\begin{aligned} &= \min_{W, w_0} \frac{1}{2} \|W\|^2 \\ &\quad \text{s.t.} \\ &\quad y_i (W^T x^i + w_0) \geq 1 \end{aligned}$$

Case II Non-separable data

$$\min_{W, \xi_i, w_0} \frac{1}{2} \|W\|^2 + C \sum_{i=1}^N \xi_i$$

s.t.

$\dots \dots \dots \rightarrow C \leftarrow$ slack variable



$$\begin{cases} y_i (w^T x_i + w_0) \geq 1 - \xi_i \\ \xi_i \geq 0 \end{cases}$$

C is a hyperparameter



$\xi_i \equiv$ Hinge loss with

$$f(x) = w^T x + w_0$$

$$\text{Hinge loss}(y f(x)) = \max(1 - y f(x), 0)$$

Proof:

$$\rightarrow \xi_i \geq 0 \quad \& \quad \xi_i \geq 1 - y_i (w^T x_i + w_0)$$

$$\xi_i \geq \max(0, 1 - y_i f(x_i))$$

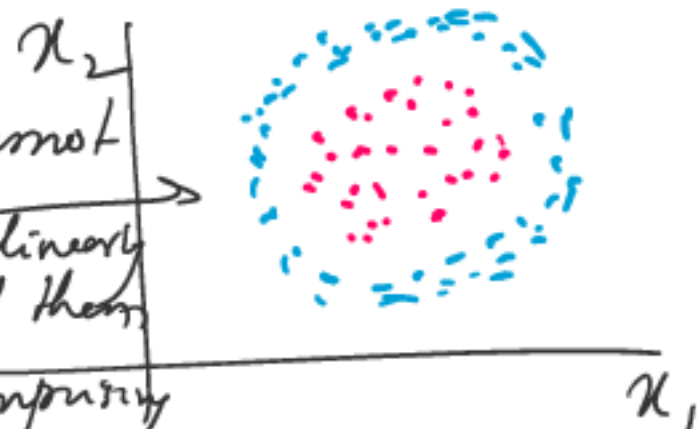
$$\xi_i \geq \text{Hinge loss}(y_i f(x_i))$$

Since the objective minimizes w.r.t ξ_i , ξ_i appears as a standalone positive linear term in objective.

as a standard
the optimal $\xi_i = \text{Hinge loss}(y_i f(x_i))$
Linear SVM for the non-separable case
 $\equiv c \cdot \text{Hinge loss} + \text{square regularizer}.$

Non-linear SVMs

these points cannot
be separated linearly
but if we embed them
in a 5-D space computing
 $\vec{x} \rightarrow (x_1, x_2, x_1^2, x_2^2, x_1 x_2)$ then they can
be easily separated linearly in 5-D space.
Instead explicitly representing points in a
transformed space, in SVMs you work with
the original space.



transformed space, kernels to stay in the original space.

Kernel: $k(x, x') = \phi(x) \cdot \phi(x')$

$\phi(x) \equiv$ embedding vector of x

- Example: kernel function: $k(x, x') = (x \cdot x' + 1)^2$
To show: when $\phi(x) = (x_1, x_2, x_1^2, x_2^2, x_1 x_2)$, then $(x \cdot x' + 1)^2 = \phi(x) \cdot \phi(x')$

$$\begin{aligned} X &\equiv (x_1, x_2) & \phi(x) &\equiv (x_1, x_2, x_1^2, x_2^2, x_1 x_2) \\ X' &\equiv (x'_1, x'_2) & \phi(x') &\equiv (x'_1, x'_2, x'^2_1, x'^2_2, x'_1 x'_2) \end{aligned}$$

$$\begin{aligned} &(x'_1, x'_2, x'^2_1, x'^2_2, x'_1 x'_2) \\ &= x_1 x'_1 + x_2 x'_2 + x_1^2 x'^2_1 + x_2^2 x'^2_2 + x_1 x_2 x'_1 x'_2 \end{aligned}$$

$$\begin{aligned} k(x, x') &\equiv (x \cdot x' + 1)^2 \\ &\equiv ((x_1, x_2) \cdot (x'_1, x'_2) + 1)^2 \\ &\equiv (x \cdot x' + x_2 x'_2 + 1)^2 \end{aligned}$$

the same
set of
terms.

$$= (x_1 x_1' + x_2 x_2' + 1) \quad //$$

$$= x_1^2 x_1'^2 + x_2^2 x_2'^2 + 1 + 2 x_1 x_1' x_2 x_2' + 2 x_1 x_1' + 2 x_2 x_2'$$

Types of kernels

$$K(x, x') = x \cdot x' \quad \leftarrow \text{linear kernel}$$

$$= (x \cdot x' + 1)^d \quad \leftarrow \text{polynomial of degree } d \text{ kernel}$$

$$= e^{-\sigma(x-x')^2} \quad \leftarrow \text{Gaussian kernel or the RBF kernel.}$$

the SVM classifier is defined in terms of the kernel function as

$$\frac{N}{2} \dots$$

$$w, \quad f(x) \equiv \left[\sum_{i=1}^N \alpha_i k(x^i, x) y_i + w_0 \right]$$

Say $k(x^i, x) = (x^i \cdot x)$

$$f(x) \equiv \sum_{i=1}^N \alpha_i (x^i \cdot x) y_i + w_0$$

$$\equiv \left(\sum_{i=1}^N \alpha_i x^i y_i \right) \cdot x + w_0$$

$$w^T x + w_0$$

similarity among points

Linear SVMs Constrained OP $\Rightarrow$ Hinge-loss formulation

$$\min_{W, w_0, \xi_i} \frac{1}{2} \|W\|^2 + C \sum_{i=1}^N \xi_i$$

$$\text{s.t. } \forall i: (W^T x^i + w_0) \geq 1 - \xi_i \quad i=1, \dots, N$$

$$\xi_i \geq 0$$

$$\min_{W, w_0} \frac{1}{2} \|W\|^2 + C$$

$$\sum_{i=1}^N \max(0, 1 - y_i (W^T x^i + w_0))$$

$$F(W, w_0)$$

Two methods

1) Hinge-loss formulation

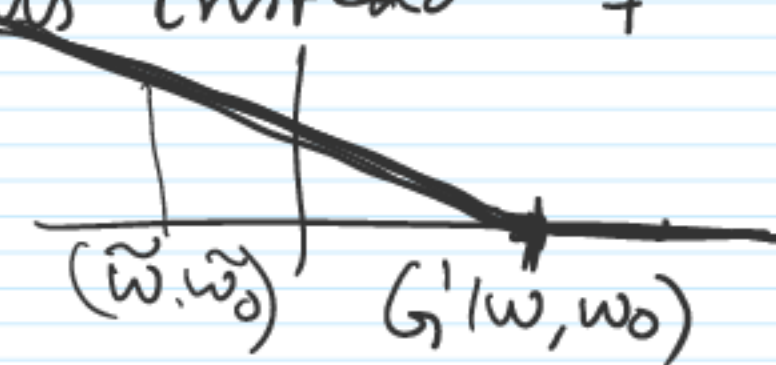
Apply stochastic gradient descent type algorithms

But Hinge loss is not differentiable in W .

So, we will work with sub-gradients instead of gradient.

$$\min_{W, w_0} G(W, w_0) \equiv \max(0, G'(W, w_0))$$

$\nearrow$ non-differentiable $\nearrow$ differentiable



sub-gradient at $(\tilde{W}, \tilde{w}_0) = \lambda$ if $G'(\tilde{W}, \tilde{w}_0) \leq 0$

subgradient at $(\tilde{w}, \tilde{w}_0) \equiv 0$ if $G(w, w_0) \leq 0$

$$\equiv \nabla_{w, w_0} G'(\tilde{w}, \tilde{w}_0) \text{ otherwise}$$

Run gradient descent algorithm
with sub-gradients instead of gradient.

II) Primal-dual method

Starting with the constrained QP formulation:
Primal (P) $\min_w F(w)$ &

s.t.

$$g_k(w) \leq 0; k=1 \dots K$$

$g_k(w)$ is convex
 $F(w)$ is convex

Lagrangian dual of (P)

$$\theta(\alpha_1, \alpha_2, \dots, \alpha_K) = \min_w F(w) + \sum_{k=1}^K \alpha_k g_k(w)$$

s.t. $\underline{\alpha_k} \geq 0 \quad \forall k=1 \dots K$

Then (Weak duality)

Thm (Weak duality)

Let $\hat{w}$ be a feasible solution of the primal P
 Then, $\forall \alpha_k \geq 0$, the dual $\theta(\alpha_1, \dots, \alpha_k) \equiv \theta(\alpha) \leq F(\hat{w})$

$$\theta(\alpha) = \min_w F(w) + \sum_k \alpha_k g_k(w)$$

$$\leq F(\hat{w}) + \sum_k \alpha_k g_k(\hat{w})$$

$\hat{w}$ is feasible $g_k(\hat{w}) \leq 0$

$\alpha_k \geq 0$

$$\leq F(\hat{w})$$

$\equiv$ Primal objective value at $\hat{w}$.

dual feasible

duality gap

primal feasible values

$\theta(\alpha^1) \theta(\alpha^2)$

$\theta(\alpha^*)$

$F(w^*)$

$F(\hat{w})$

$F(w^4)$

$F(w^1)$

optimal w for primal

$\forall \alpha$

$\theta(\alpha)$

$\leq$

$F(\hat{w})$

$\forall$

feasible $\hat{w}$

s.t

$g_k(\hat{w}) \leq 0$

primal.

$$\begin{array}{l} \max_{\alpha_1 \dots \alpha_k} \theta(\alpha) \\ \alpha_i \geq 0 \end{array} \quad \begin{array}{l} \text{Lagrangian} \\ \text{Dual} \end{array} \quad \begin{array}{l} \equiv \\ \Rightarrow \end{array} \quad \begin{array}{l} \boxed{\begin{array}{l} \max_{\alpha_1 \dots \alpha_k} \min_W F(W) + \sum_{k=1}^K \alpha_k g_k(W) \\ \alpha_k \geq 0 \quad \forall k=1 \dots K \end{array}} \end{array}$$

$| g_k(\hat{w}) \leq 0 \text{ primal}$

Thm (Strong duality)

If $F(W)$, $g_k(W)$ are convex functions, $\exists W$ s.t. $g_k(W) < 0 \quad \forall k$, then there exists (α^*, W^*) s.t.

$F(W^*)$ is primal optimal and W^* is primal feasible and $\theta(\alpha^*)$ is dual feasible and optimum and $\theta(\alpha^*) = F(W^*)$

$\Rightarrow$ duality gap is 0.

Primal

$$\min_{W \text{ s.t.}} F(W)$$

Dual

$$\max_{\alpha_1 \dots \alpha_k} \min_W F(W) + \sum_{k=1}^K \alpha_k g_k(W) \\ \alpha_k \geq 0 \quad \forall k=1 \dots K$$

$$\begin{aligned} & \text{s.t.} \\ & g_k(w) \leq 0 \quad \forall k=1 \dots K \end{aligned}$$

$$\alpha_k \geq 0 \quad \forall k=1 \dots K.$$

when $f(w)$ & $g_k(w)$ are convex.

Applying to SVMs

$$\begin{aligned} & \min_{w, w_0, \xi_i} \frac{1}{2} \|w\|^2 + C \sum_{i=1}^N \xi_i \\ & w \in \mathbb{R}^d \quad \left\{ \begin{aligned} & 1 - \xi_i - y_i (w^T x^i + w_0) \leq 0 \\ & -\xi_i \leq 0 \end{aligned} \right. \quad \forall i=1 \dots N \end{aligned}$$

$$\begin{aligned} & \max_{\alpha_i, \mu_i} \left[\min_{w, w_0, \xi_i} \left\{ \frac{1}{2} \|w\|^2 + C \sum_{i=1}^N \xi_i \right. \right. \\ & \quad \left. \left. + \sum_{i=1}^N \alpha_i (1 - \xi_i - y_i (w^T x^i + w_0)) \right. \right. \\ & \quad \left. \left. + \sum_{i=1}^N \mu_i (-\xi_i) \right\} \right] \end{aligned}$$

$$\text{s.t. } \alpha_i \geq 0; \mu_i \geq 0 \quad \forall i=1 \dots N$$

Solving

$$\begin{aligned} & \max_{\alpha_i, \mu_i} \left[\min_{w, w_0, \xi_i} L(w, w_0, \xi_i, \alpha_i, \mu_i) \right] \\ & \text{convex in } w, w_0, \xi_i \end{aligned}$$

Step 1

$$\nabla_w L \equiv w^* + \sum_{i=1}^N \alpha_i^* (-y_i x^i) = 0$$

$$\Rightarrow w^* = \sum_{i=1}^N \alpha_i^* y_i x^i$$

$$\nabla_{w_0} L \equiv -\sum_{i=1}^N \alpha_i^* y_i = 0$$

$$\Rightarrow \sum_{i=1}^N \alpha_i^* y_i = 0$$

$$\nabla_{\alpha_i} L = C - \alpha_i^* - u_i = 0$$

$$\Rightarrow \boxed{\alpha_i^* + u_i^* = C}$$

$$\sum_{i=1}^N \alpha_i^* = \sum_{\substack{i=1: \\ y_i = -1}}^N \alpha_i^* + \sum_{\substack{i=1: \\ y_i = +1}}^N \alpha_i^*$$

Step-2

Substitute $w = \sum_{i=1}^N \alpha_i y_i x^i$ in $L(w, w_0, \alpha_i, u_i)$ and simplify and that will give us finally

$$\min_{w, w_0, \alpha_i} L(w, w_0, \alpha_i, u_i) \Leftrightarrow -\frac{1}{2} \sum_{i=1}^N \sum_{j=1}^N \alpha_i \alpha_j y_i y_j x^i \cdot x^j + \sum_{i=1}^N \alpha_i$$

$$\text{s.t. } \alpha_i \geq 0, u_i \geq 0$$

$$\left. \begin{array}{l} \text{s.t. } \sum_{i=1}^N \alpha_i y_i = 0 ; \alpha_i + u_i = C \quad \forall i \\ \alpha_i \geq 0 ; u_i \geq 0 \end{array} \right\}$$

Step-3

Eliminate u_i

$$\max_{\alpha_1, \dots, \alpha_N} -\frac{1}{2} \sum_{i=1}^N \sum_{j=1}^N \alpha_i \alpha_j y_i y_j x^i \cdot x^j + \sum_{i=1}^N \alpha_i$$

$$\text{s.t. } \sum_{i=1}^N \alpha_i y_i = 0, \alpha_i \geq 0 ; \alpha_i \leq C$$

$$s.t \sum_{i=1}^N \alpha_i y_i = 0, \quad \alpha_i \geq 0; \quad \alpha_i \leq C$$

Last class Primal ✓

$$\min_{W, w_0, \xi_i} \frac{1}{2} \|W\|^2 + C \sum_{i=1}^N \xi_i$$

$$s.t \quad y_i (W X^i + w_0) \geq 1 - \xi_i \\ \forall i = 1 \dots N$$

Dual ✓

$$\max_{\alpha_1, \dots, \alpha_N} - \frac{1}{2} \sum_{i=1}^N \sum_{j=1}^N \alpha_i \alpha_j y_i y_j \boxed{X^i \cdot X^j} + \sum_{i=1}^N \alpha_i$$

s.t

$$0 \leq \alpha_i \leq C \leftarrow \text{box constraint} \\ \sum_{i=1}^N \alpha_i y_i = 0$$

Extending the dual to arbitrary kernels using "kernel trick"

1) In the dual, X^i & X^j only appear as $X^i \cdot X^j$

2) For a non-linear embedding $\phi(X)$, we can use $\phi_K(X^i) \cdot \phi_K(X^j) = \underline{\underline{K(X^i, X^j)}}$

Now the kernelized dual objective becomes:

view the problem as

$$\max_{\alpha_1, \dots, \alpha_N} -\frac{1}{2} \sum_i \sum_j \alpha_i \alpha_j y_i y_j k(x^i, x^j) + \sum_i \alpha_i \quad \theta(\alpha)$$

s.t. $0 \leq \alpha_i \leq C$; $\sum_i \alpha_i y_i = 0$

$$W = \sum_{i=1}^N \alpha_i y_i \underline{x^i}$$

Your classifier: $W^T x + w_0$

$$\left(\sum_{i=1}^N \alpha_i y_i x^i \right) \cdot (x) + w_0$$

$$= \sum_{i=1}^N \alpha_i y_i k(x^i, x) + w_0$$

$\theta(\alpha) = \theta(\alpha_1, \dots, \alpha_N)$ is concave in $\alpha_1, \dots, \alpha_N$

1) $\max_{\alpha_1, \dots, \alpha_N} \theta(\alpha_1, \dots, \alpha_N)$ can be solved using gradient ascent subject to the simple constraints

Cannot use direct gradient ascent since the objective is quadratic in the number of training examples.

- 2) Single co-ordinate ascent will not work because of the constraint $\sum \alpha_i y_i = 0$
- 2) SMO Algorithm - change two α -s in each step.

Initially: $t=0$

$\rightarrow \alpha_i^t = 0 \quad \forall i = 1 \dots N$

$\rightarrow$ while (change) {
 (α_i^t, α_j^t) ← pick two α -s to change

choose a change amount: λ

$\rightarrow \alpha_i^{t+1} \leftarrow \alpha_i^t + \lambda$

$\alpha_j^{t+1} \leftarrow \alpha_j^t \mp \lambda$

depend on whether $y_i = y_j$ or not.

Pick λ so as to maximize

$$\lambda^* = \arg \max_{\lambda} \Phi(\alpha_1^t, \dots, \alpha_i^t + \lambda, \dots, \alpha_j^t \mp \lambda, \dots, \alpha_N^t)$$

λ^* so as to be closest

can be solved by closed form making gradient to 0

based on heuristics

$\lambda^t \leftarrow$ so as to be closest to λ^* while satisfying the constraints

close
equating
of θ w.r.t λ to
be 0.

$$\begin{cases} 0 \leq \alpha_i^t + \lambda^t \leq C \\ 0 \leq \alpha_j^t + \lambda^t \leq C \end{cases} \quad t = t+1$$

Example }
 $t=0 \alpha_1 \dots \alpha_N \equiv 0$; α_1, α_2 & α_i, α_j ; $y_1 \neq y_2$

$$\theta(\lambda) = -\frac{2}{2}(\lambda)(\lambda)k(x_1, x_2)y_1y_2 + 2\lambda$$

$$-\frac{1}{2}\lambda^2 k(x_1, x_1)y_1^2 - \frac{1}{2}\lambda^2 k(x_2, x_2)y_2^2$$

$$\frac{\partial \theta}{\partial \lambda} \equiv -2\lambda k(x_1, x_2)(-1) - \lambda k(x_1, x_1) - \lambda k(x_2, x_2) + 2$$

$$\Rightarrow \lambda^* = \frac{-2}{2k(x_1, x_2) - k(x_1, x_1) - k(x_2, x_2)}$$

Say $\rightarrow k(x_1, x_2) = -1$;

$k(x_1, x_1) = k(x_2, x_2) = 1$

$C = 1/3$

$$\Rightarrow \lambda = \frac{-2}{-2-1-1} = \frac{1}{2}$$

$\lambda^t \equiv 1/2$ to ensure them $\alpha_i + \lambda^t \leq C$

$$\lambda \equiv 1/3 \text{ to ensure that } \alpha_1, \alpha_2, \dots$$

$$\Rightarrow \alpha_1 = \frac{1}{3}, \alpha_2 = \frac{1}{3}$$

Thm (KKT conditions)

when $F(w)$, $g_k(w)$ are convex
then a (α^*, w^*) are optimal
solutions for the dual &
primal respectively if and only if

$$1) \nabla_w F(w^*) + \sum_{k=1}^K \alpha_k^* \nabla_w g_k(w^*) = 0$$

$$2) \alpha_k g_k(w) = 0 \quad \forall k = 1 \dots K \quad \text{---} \quad \sum_k \alpha_k^* g_k(w^*) = 0$$

$$3) \alpha_k^* \geq 0$$

$$4) g_k(w^*) \leq 0$$

Primal
 $\min_w F(w)$
s.t. $g_k(w) \leq 0$
 $\forall k = 1 \dots K$

Dual
 $\max_{\alpha_1, \dots, \alpha_K} \left(\min_w F(w) + \sum_{k=1}^K \alpha_k g_k(w) \right)$
 $\alpha_k \geq 0 \quad \forall k = 1 \dots K$

Applying KKT conditions to SVMs

$$w^* = \sum_{i=1}^N \alpha_i^* y_i x_i \quad \left. \vphantom{\sum_{i=1}^N} \right\} \text{1st KKT condition}$$

$$\sum_{i=1}^N \alpha_i y_i = 0$$

$$\alpha_i + \mu_i = C$$

KK condition

$$\mu_i \geq 0$$

$$\alpha_i \geq 0$$

$$\xi_i \geq 0$$

primal
dual
feasibility

$$\alpha_i (1 - y_i (w^T x_i + w_0) - \xi_i) = 0 \quad \forall i$$

2nd
KK condition.

$$\mu_i \xi_i = 0$$

Case 1: consider $i : y_i (w^T x_i + w_0) > 1$, $\xi_i = 0$

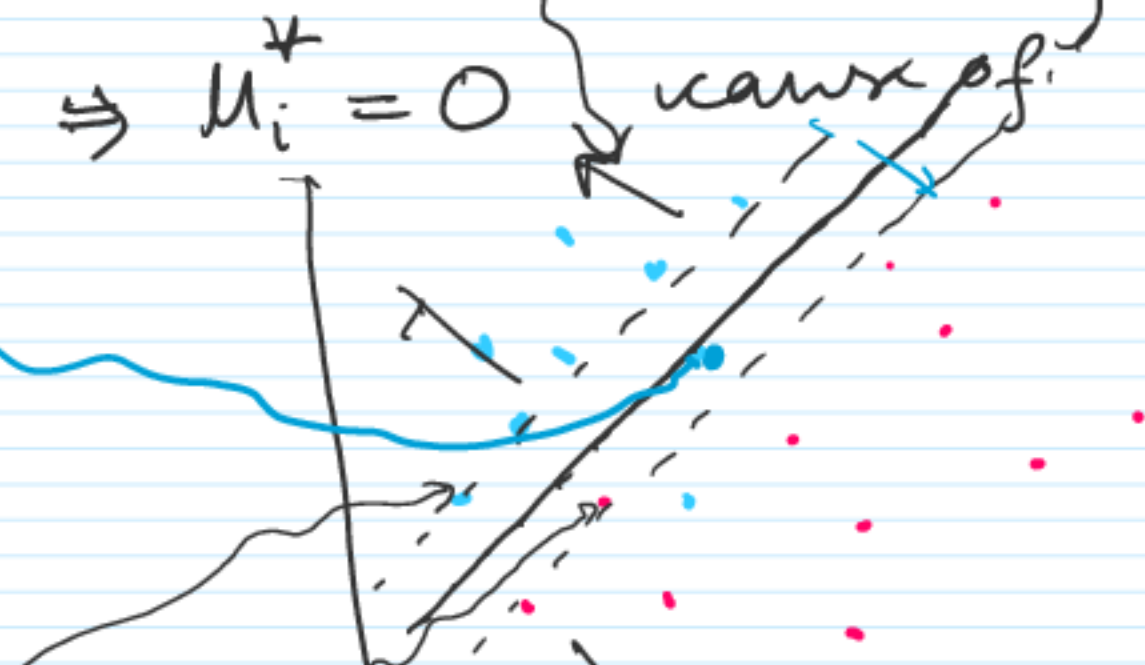
$$\Rightarrow \alpha_i = 0 \text{ from}$$

Case 2: consider $i : \xi_i > 0 \Rightarrow \mu_i = 0$ cause of

$$\Rightarrow \alpha_i = C$$

Case 3: i s.t.

$$0 < \alpha_i < C$$



$$0 < \underline{\alpha_i} < C$$

$$\Rightarrow \mu_i > 0; \quad \eta_i = 0$$

$$\Rightarrow y_i (w^* x^i + w_0) = 1$$

$$w^* = \sum_{i=1}^N \alpha_i y_i x^i$$

$$= \sum \alpha_i y_i x^i$$

$i \in \text{support vectors}$

$\downarrow$ those examples for which $\alpha_i > 0$

$$f(x) = w^* x + w_0$$

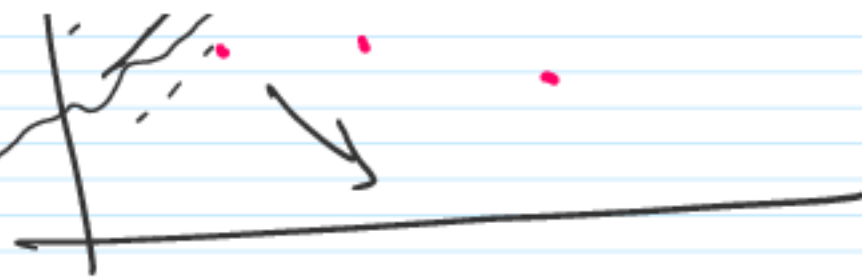
$$= \sum \alpha_i y_i k(x^i, x) + w_0$$

$i \in \text{support vectors}$

How to compute w_0

Find an example i for which α_i is $< C$ & > 0 .

$y_i (w^* x^i + w_0) = 1$ $\Rightarrow$ solve this to find w_0



$y_i \in \{1, \dots, K\}$

find w_0

Extending SVMs to
multi-class classification.

Suppose we have K classes.

Like in logistic regression classifiers we allocate
parameters w^y, w_0^y for each class y .

$$w^y \in \mathbb{R}^d, w_0^y \in \mathbb{R}.$$

Score for an example being in class y as

$$S(x, y) = w^y x + w_0^y$$

$$\min_{\{w^y, w_0^y\}, \xi_i} - \frac{1}{2} \sum_{y=1}^K \|w^y\|^2 + C \sum_{i=1}^N \xi_i$$

$$\begin{aligned} (x^i, y_i) \quad \text{s.t.} \quad & w^{y_i} x^i + w_0^{y_i} \geq w^y x^i + w_0^y + 1 - \xi_i \quad \forall i=1 \dots N \\ & \forall y \neq y_i \\ & \xi_i \geq 0 \end{aligned}$$

Hinge loss

$$\max(0, 1 - (w^y x^i + w_0^y))$$

Hinge loss

$$\max_{y \neq y_i} (0, 1 - [\underline{s(x^i, y_i)} - \underline{s(x^i, y)}])$$