

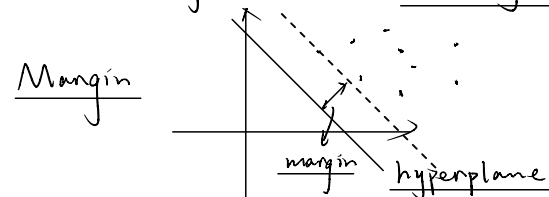
SVM: Supported Vector Machine

training set: $D = \{x_i, y_i\}_{i=1}^N$ where $y_i \in \{-1, +1\}$

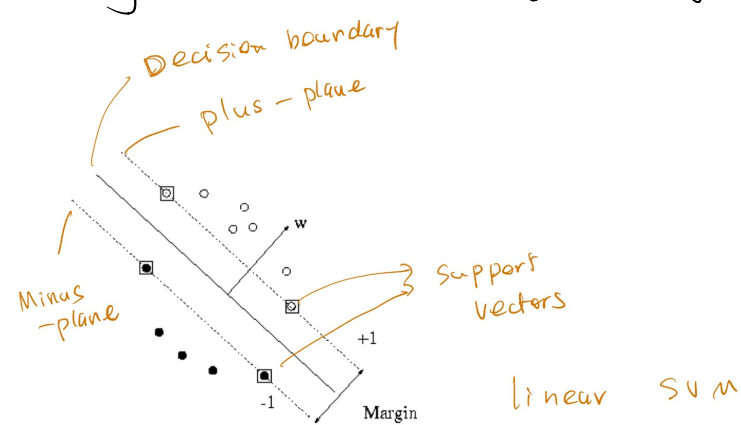
$$\Rightarrow \hat{y} = \text{sign}(w^T x + b) \quad (\text{sign}(\cdot) = \begin{cases} +1 \\ -1 \end{cases})$$

$$x = (x_1 \dots x_D)^T \quad w = (w_1 \dots w_D)^T$$

Hard-margin SVM $\Rightarrow$ linearly separable data



Theory: hyperplane with largest margin generalise best

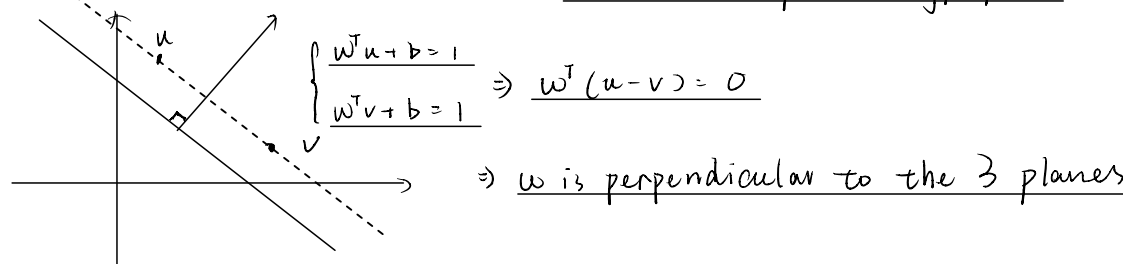


since: $\hat{y}_i = \text{sign}(w^T x_i + b)$, Decision Boundary: $w^T x + b = 0$

let: $\begin{cases} \text{plus-plane: } w^T x + b = c \\ \text{minus-plane: } w^T x + b = -c \end{cases}$

$$\Rightarrow \text{rewrite: } \begin{cases} \frac{1}{c} w^T x + \frac{b}{c} = 1 \\ \frac{1}{c} w^T x + \frac{b}{c} = -1 \end{cases} \Rightarrow \begin{cases} \text{plus-plane: } w^T x + b = 1 \\ \text{minus-plane: } w^T x + b = -1 \end{cases}$$

"Canonical Separable Hyperplane"



- Let x^* be a point (not necessarily a training example) on the minus plane.
- Let $x^\dagger$ be a point on the plus plane that is closest to x^* .
- The $\vec{w} = (x^\dagger - x^*) / \|x^\dagger - x^*\|$ length of the vector from x^* to $x^\dagger$.

$$\begin{cases} w^T x^* + b = -1 \\ w^T x^\dagger + b = 1 \end{cases}$$

$$\vec{x}^* = \vec{x}^\dagger - \lambda \vec{w}, \text{ where } \lambda \text{ is a constant}$$

$$w^T (\vec{x}^\dagger - \lambda \vec{w}) = -1 \Rightarrow w^T \vec{x}^\dagger - \lambda w^T \vec{w} = -1 \Rightarrow \lambda = \frac{\frac{1}{w^T \vec{w}}}{\frac{1}{w^T \vec{w}}} = \frac{1}{\|w\|^2}$$

$$\Rightarrow |M| \text{ (length of margin)} = \| \lambda \vec{w} \| = \frac{1}{\|w\|}$$

Optimal Margin Classifier

$$\Rightarrow \text{Max: } |M| = \frac{1}{\|w\|} \text{ with constraints: } w^T x_i + b \begin{cases} \geq 1 & \text{if } y_i = 1 \\ \leq -1 & \text{if } y_i = -1 \end{cases}$$

$$\min_{w, b} \frac{\|w\|^2}{2} \text{ subject to } y_i (w^T x_i + b) \geq 1$$

"The Primal Optimization Problem"

Lagrangian multiplier

$$\mathcal{L}(w, b, \alpha) = \frac{1}{2} \|w\|^2 - \sum_{i=1}^N \alpha_i [y_i (w^T x_i + b) - 1]$$

where: $w = (w_1 \dots w_D)$ bc $b = b_1 \dots b_n$ $\alpha = (\alpha_1 \dots \alpha_n)$

$$\begin{cases} \text{primal problem} \\ \min_{w, b} \mathcal{L}(w, b) = \min_{w, b} \max_{\alpha} \mathcal{L}(w, b, \alpha) \end{cases}$$

dual problem:

$$\max_{\alpha} \mathcal{L}_d(\alpha) = \max_{\alpha} \min_{w, b} \mathcal{L}(w, b, \alpha)$$

$$\max_{\alpha} \min_{w, b} \mathcal{L}(w, b, \alpha) \leq \max_{\alpha} \min_{w, b} \max_{\alpha} \mathcal{L}(w, b, \alpha) = \min_{w, b} \max_{\alpha} \mathcal{L}(w, b, \alpha)$$

Dual Problem $\Rightarrow$ "A Lower Bound"

$$d^* \leq p^*$$

When training example is linear separable

$$\exists w^*, b^*, \alpha^*, \quad d^* = \mathcal{L}(w^*, b^*, \alpha^*) = p^*$$

Strengthen KKT Conditions:

$$\begin{cases} \text{Optimality: } \begin{cases} \frac{\partial \mathcal{L}}{\partial w} = 0 \\ \frac{\partial \mathcal{L}}{\partial b} = 0 \end{cases} \\ \text{Dual Complementary Condition: } \begin{cases} \alpha_i [y_i (w^T x_i + b) - 1] = 0 \\ y_i (w^T x_i + b) \geq 1 \\ \alpha_i \geq 0 \end{cases} \end{cases}$$

Solving Dual Problem

$$\max_{\alpha} \mathcal{L}_d(\alpha) = \max_{\alpha: \alpha_i \geq 0} \min_{w, b} \mathcal{L}(w, b, \alpha) = \frac{1}{2} \|w\|^2 - \sum_{i=1}^N \alpha_i [y_i (w^T x_i + b) - 1]$$

$$\begin{cases} \frac{\partial \mathcal{L}}{\partial w_k} = w_k - \sum_{i=1}^N \alpha_i y_i x_{ik} = 0 \\ \frac{\partial \mathcal{L}}{\partial b} = - \sum_{i=1}^N \alpha_i y_i = 0 \end{cases} \Rightarrow \begin{cases} w_k = \sum_{i=1}^N \alpha_i y_i x_{ik} \\ \sum_{i=1}^N \alpha_i y_i = 0 \end{cases}$$

$$\therefore \max_{\alpha} \mathcal{L}(\alpha) = \sum_{i=1}^N \alpha_i - \frac{1}{2} \sum_{i,j=1}^N \alpha_i \alpha_j y_i y_j x_i^T x_j$$

How to find b ?

N_s : # of support vector.

$$\text{KKT: } y_i (w^T x_i + b) = 1$$

$$y_i y_i (w^T x_i + b) = y_i$$

$$b = y_i - w^T x_i = \frac{y_i}{\sum_{i=1}^{N_s} (y_i - w^T x_i)}$$

Prediction

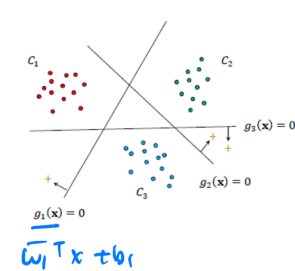
$$\hat{y} = \text{sign}(w^T x + b) = \text{sign}(\sum_{i=1}^{N_s} \alpha_i y_i x_i^T x + b)$$

Multiple Classes

- One way to handle multiple (K) classes is to define K two-class problems, each separating one class from all other classes combined.
- Let $g(x)$ denote $w_j^T x + b_j$ (Discriminant function)
- Classification rule:

$$\hat{y} = \arg \max_j g_j(x)$$

Note: Not using $\text{sign}(g_j(x))$



Soft-Margin SVM: not linearly separable

$$\text{Relaxation: } y_i (w^T x_i + b) \geq 1 - \xi_i \quad (\xi_i \geq 0)$$

Slack Variable

$$\text{Errors} \leq \sum_{i=1}^N \xi_i$$

The larger the constant C , the more we want to minimize error, the more complex the decision boundary.

$$\Rightarrow \min \frac{1}{2} \|w\|^2 + C \sum_{i=1}^N \xi_i \text{ subject to } y_i (w^T x_i + b) - 1 \geq -\xi_i$$

$$\mathcal{L}(w, b, \xi, \alpha, \gamma) = \frac{1}{2} \|w\|^2 + C \sum_{i=1}^N \xi_i - \sum_{i=1}^N \alpha_i (y_i (w^T x_i + b) - 1 + \xi_i) - \sum_{i=1}^N \gamma_i \xi_i$$

Constraint: $\xi_i \geq 0$

- Original problem equivalent to (Primal Problem):

$$\min_{w, b, \xi} \max_{\alpha, \gamma} \max_{\alpha_i \geq 0, \gamma_i \geq 0} \mathcal{L}(w, b, \xi, \alpha, \gamma)$$

- Dual Problem:

$$\max_{\alpha, \gamma} \max_{\alpha_i \geq 0, \gamma_i \geq 0} \min_{w, b, \xi} \mathcal{L}(w, b, \xi, \alpha, \gamma)$$

Solving the dual: $\min_{w, b, \xi} \mathcal{L}(w, b, \xi, \alpha, \gamma)$

$$\frac{\partial \mathcal{L}}{\partial w} = 0 \Rightarrow w = \sum_{i=1}^N \alpha_i y_i x_i$$

$$\frac{\partial \mathcal{L}}{\partial b} = 0 \Rightarrow \sum_{i=1}^N \alpha_i y_i = 0$$

$$\frac{\partial \mathcal{L}}{\partial \xi_i} = 0 \Rightarrow C - \alpha_i - \gamma_i = 0$$

Dual Problem:

$$\begin{aligned} & \text{maximize} \quad \sum_{i=1}^N \alpha_i - \frac{1}{2} \sum_{i,j=1}^N \alpha_i \alpha_j y_i y_j x_i^T x_j \\ & \text{subject to} \quad C \geq \alpha_i \geq 0, \quad i = 1, \dots, N \\ & \quad \quad \quad \sum_{i=1}^N \alpha_i y_i = 0 \end{aligned}$$

Impact of C :

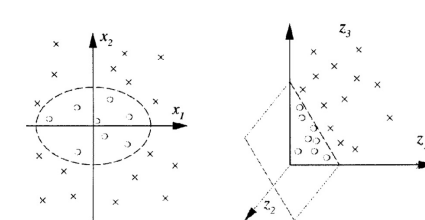
- Small C : wider-margin, more misclassified points
- Large C : minimize misclassification errors, narrower margin

$$\text{minimize } \frac{1}{2} \|w\|^2 + C \sum_i \xi_i$$

Relationship with Hinge Loss

$$\begin{aligned} & \min \frac{1}{2} \|w\|^2 + C \sum_i \xi_i \quad \text{s.t. } y_i (w^T x_i + b) \geq 1 - \xi_i, \quad \xi_i \geq 0 \\ & \left\{ \begin{aligned} \xi_i & \geq 1 - y_i (w^T x_i + b) \\ \xi_i & \geq 0 \end{aligned} \right. \Rightarrow \xi_i = \max \{0, 1 - y_i (w^T x_i + b)\} \text{ where } z_i = w^T x_i + b \\ & \text{Recall: Hinge loss: } \mathcal{L}_{\text{hinge}}(y, b) = \max \{0, 1 - yb\} \\ & \Rightarrow \min \frac{1}{2} \|w\|^2 + C \sum_i \max \{0, 1 - y_i (w^T x_i + b)\} \end{aligned}$$

Kernel Extension



- Mapping $\varphi: \mathbb{R}^m \rightarrow \mathcal{H}, x \mapsto \varphi(x)$
- $\mathbb{R}^m$: input space (attributes)
- $\mathcal{H}$: feature space (features)
- Transform the data with the mapping

$$(x_1, y_1), \dots, (x_N, y_N) \Rightarrow (\varphi(x_1), y_1), \dots, (\varphi(x_N), y_N)$$

not linearly separable $\xrightarrow[\text{dimension}]{\text{kernel}}$ linearly separable

feature vector

Dual Problem on Features:

$$\begin{aligned} & \text{maximize} \quad \sum_{i=1}^N \alpha_i - \frac{1}{2} \sum_{i,j=1}^N \alpha_i \alpha_j y_i y_j \kappa(x_i, x_j) \\ & \text{subject to} \quad C \geq \alpha_i \geq 0, \quad i = 1, \dots, N \\ & \quad \quad \quad \sum_{i=1}^N \alpha_i y_i = 0 \end{aligned}$$

Kernel Trick:

$$\text{let } K(x_i, x_j) = \varphi(x_i)^T \varphi(x_j)$$

$$\begin{aligned} & \text{Problem: } \max \sum_{i=1}^N \alpha_i - \frac{1}{2} \sum_{i,j=1}^N \alpha_i \alpha_j y_i y_j K(x_i, x_j) \\ & \text{s.t. } 0 \leq \alpha_i \leq C, \quad i = 1, \dots, N \\ & \quad \quad \quad \sum_{i=1}^N \alpha_i y_i = 0 \end{aligned}$$

The trick:

- No need to explicitly calculate φ and the inner product $\varphi(x_i)^T \varphi(x_j)$
- Just need to specify the kernel function $K(x_i, x_j)$
- K is cheaper to calculate
- Allow one to use very high dimensional feature space.

Common Kernels

- $K(x, z) = (x^T z)^d$, polynomial kernel

- Corresponding feature transformation when $d = 2$:

$$\varphi(x_1, x_2, \dots, x_n) = (x_1 x_1, x_1 x_2, x_1 x_3, \dots, x_n x_n)$$

- Inhomogeneous polynomial: $K(x, z) = (x^T z + 1)^d$

- The corresponding feature transformation?

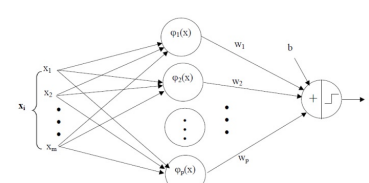
- Gaussian: $K(x, z) = \exp(-\|x - z\|^2 / 2\sigma^2)$

- Radial basis function (RBF) network

- Corresponds to an infinite-dimensional feature space

- Intuitively, $K(x, z)$ represents our desired notion of similarity between data x and z and this is from our prior knowledge

Classification with SVM



- Choose nonlinear transformation $\varphi = (\phi_1, \phi_2, \dots, \phi_p)$ (implicitly via $K(x, z)$)

- Solve the dual optimization problem on features, get α

- Calculate w and b from α .

- Classify future examples as follows:

$$\text{sign} \left(\sum_{i=1}^{N_s} \alpha_i y_i K(x_i, x) + b \right)$$

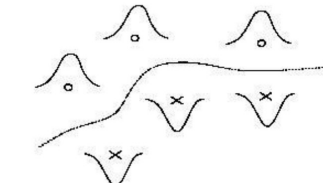
Gaussian Kernel

Recall that the decision rule uses

$$\text{sign}(w^T \varphi(x) + b) = \text{sign} \left(\sum_{i=1}^{N_s} \alpha_i y_i K(x_i, x) + b \right)$$

- Gaussian kernel: $K(x_1, x_2) = \exp(-\|x_1 - x_2\|^2 / 2\sigma^2)$

- Amounts to putting bumps of various sizes on the training set



- Small σ : Narrow Gaussian bell curve. Each support vector affects only nearby points. Decision boundary is complex and wiggly. Result: Low bias, high variance (overfitting risk)
- Large σ : Wide Gaussian bell curve. Each support vector affects a large neighborhood. Decision boundary is smooth and simple (near-linear). Result: High bias, low variance (underfitting risk)
- Small C : Soft margin (wide margin allowed). Many points can be misclassified or inside margin. Low penalty for training errors. Result: Simpler model, underfitting possible
- Large C : Hard margin (narrow margin enforced). Few misclassifications tolerated. High penalty for training errors. Result: Complex model, overfitting possible

When to use SVM instead of deep learning?

- Small datasets - SVM performs well with limited samples; deep learning needs large data.
- When training time must be short - SVM trains faster than deep networks for small/medium GEs
- When interpretability matters - SVMs are simpler and easier to analyze than deep nets.
- When the decision boundary is complex but data is not huge - kernel tricks give SVM strong nonlinear power.
- Risk of overfitting is high - SVM's margin maximization naturally controls overfitting better on small data.