

Support Vector Machines (SVMs)

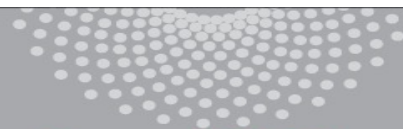
Aarti Singh

Machine Learning 10-315

Oct 20, 2021



MACHINE LEARNING DEPARTMENT



Carnegie Mellon.
School of Computer Science

Discriminative Classifiers

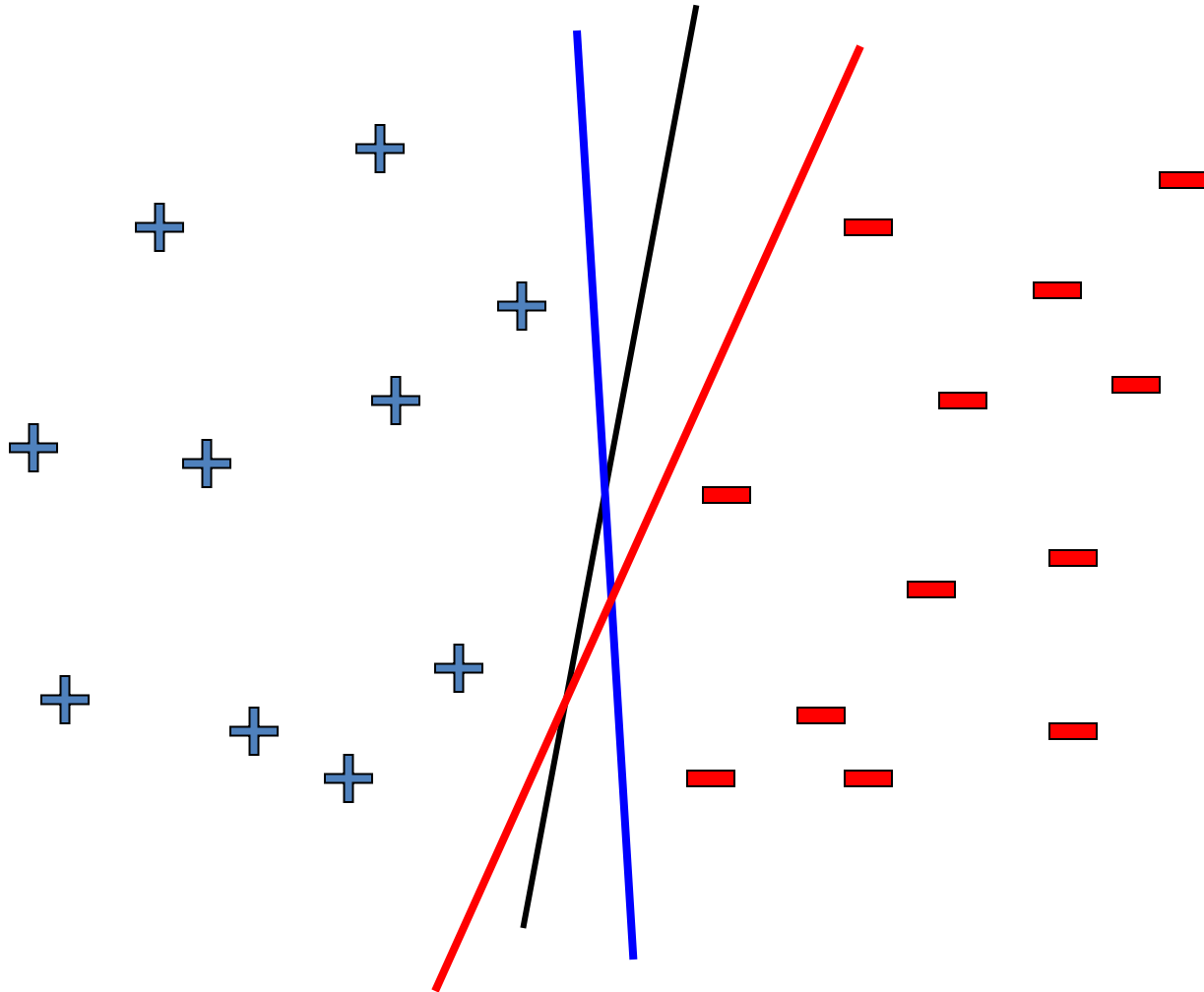
Optimal Classifier:

$$\begin{aligned} f^*(x) &= \arg \max_{Y=y} P(Y = y | X = x) \quad \checkmark \\ &= \arg \max_{Y=y} P(X = x | Y = y) P(Y = y) \quad \checkmark \quad \checkmark \end{aligned}$$

Why not learn P(Y|X) directly? Or better yet, why not learn the decision boundary directly?

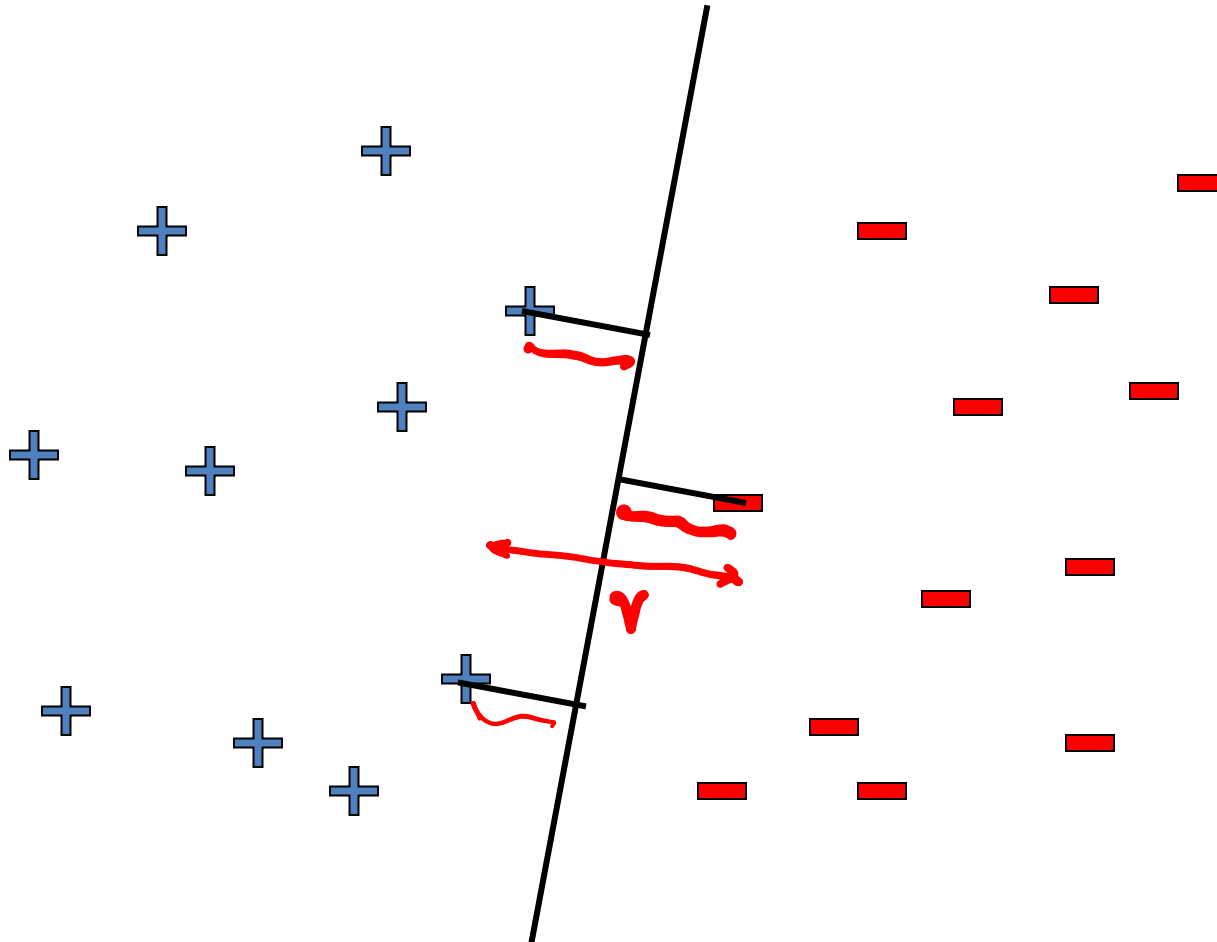
- Assume some functional form for P(Y|X) (e.g. Logistic Regression) or for the decision boundary (e.g. Neural nets, SVMs)
- Estimate parameters of functional form directly from training data

Linear classifiers – which line is better?

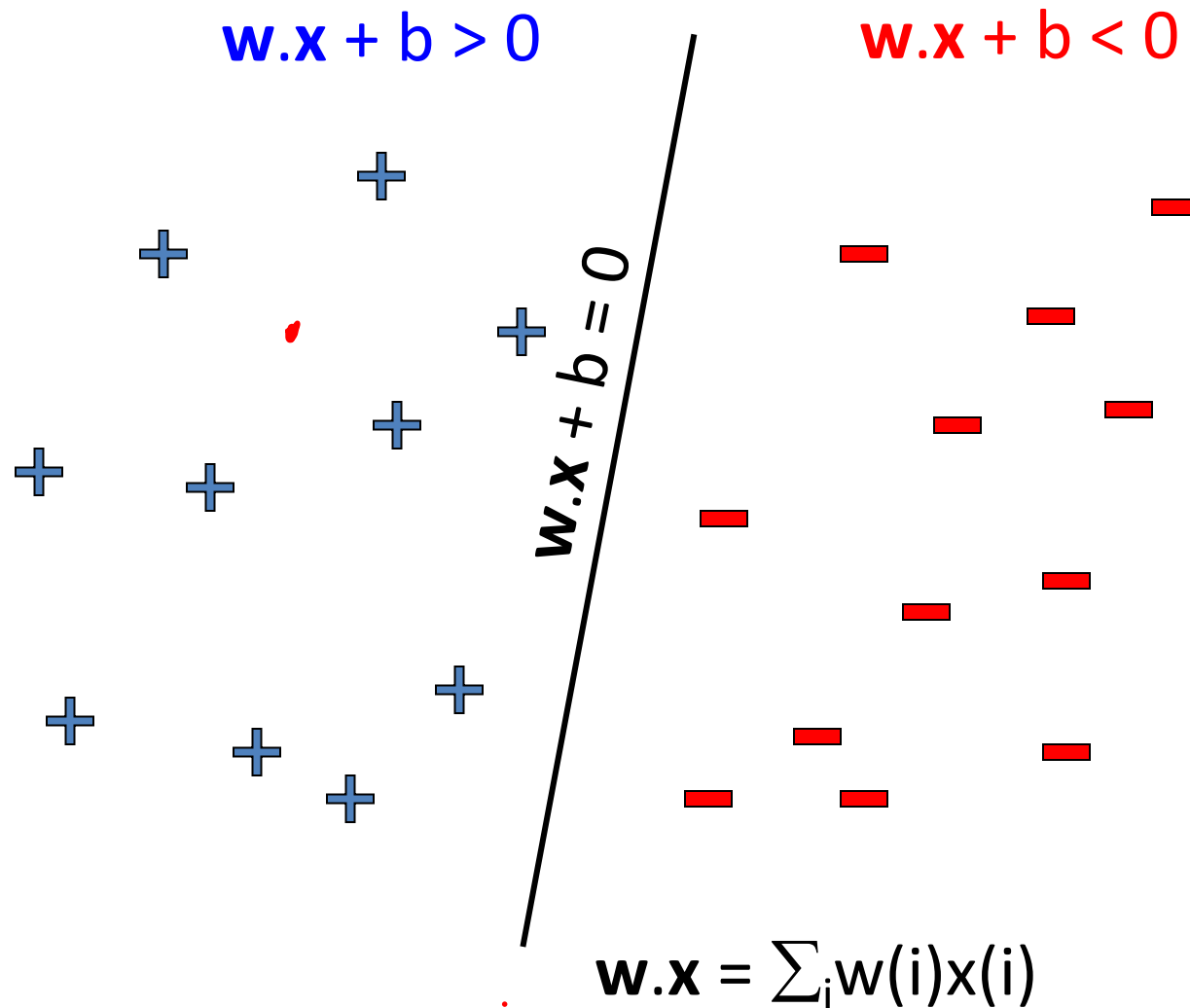


Pick the one with the largest margin!

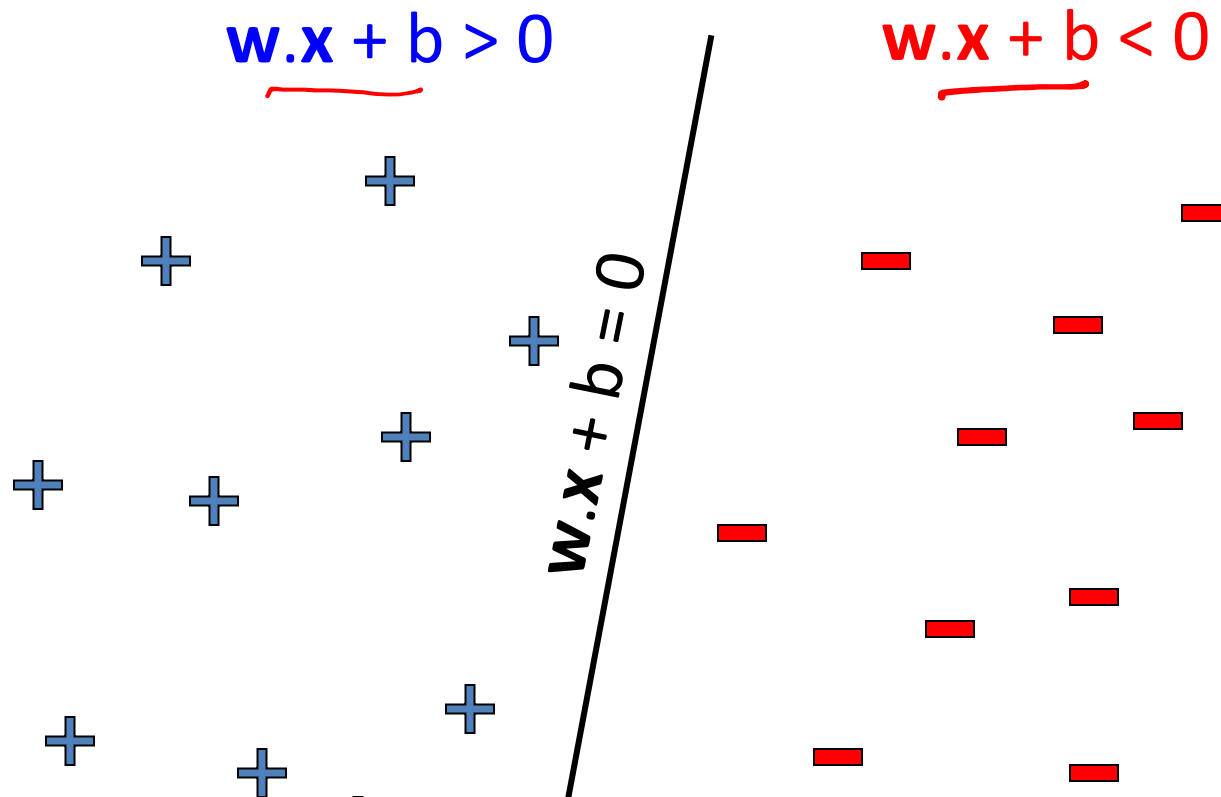
γ



Parameterizing the decision boundary



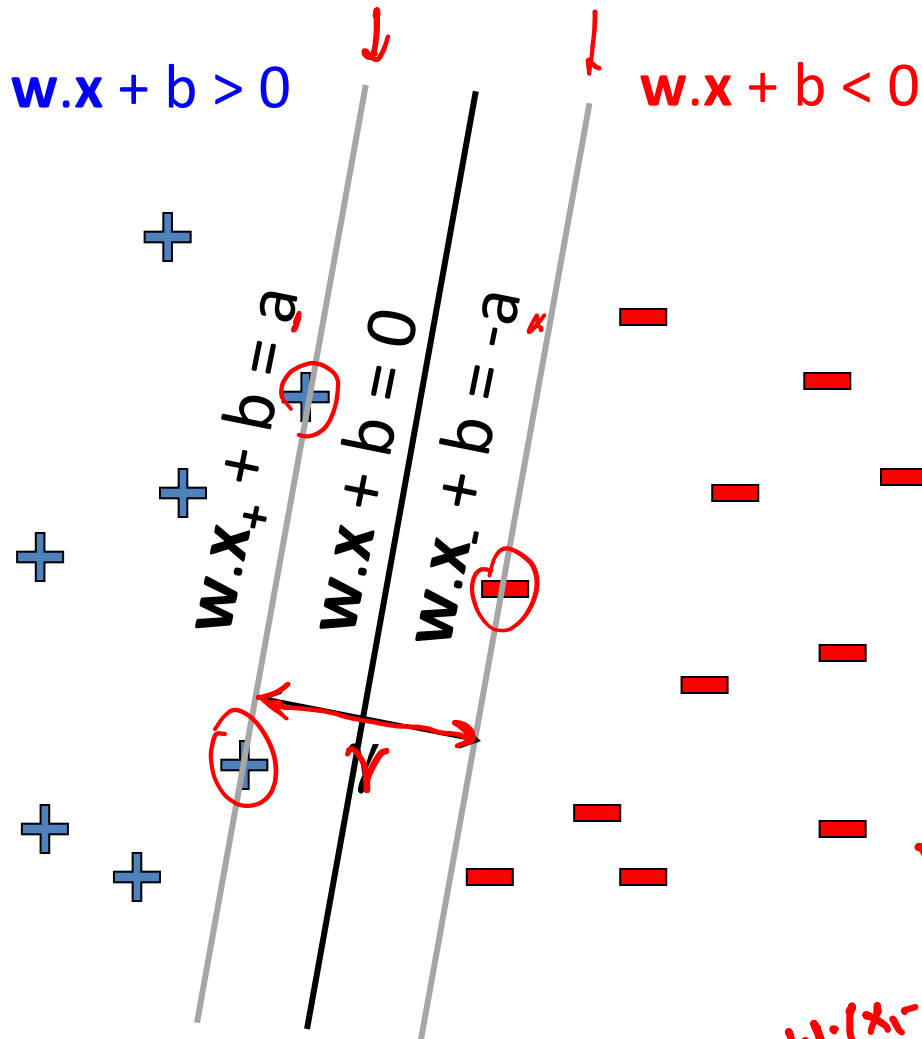
Parameterizing the decision boundary



$y_j \in \{-1, +1\}$ — class

“confidence” $= (w \cdot x_j + b) y_j$

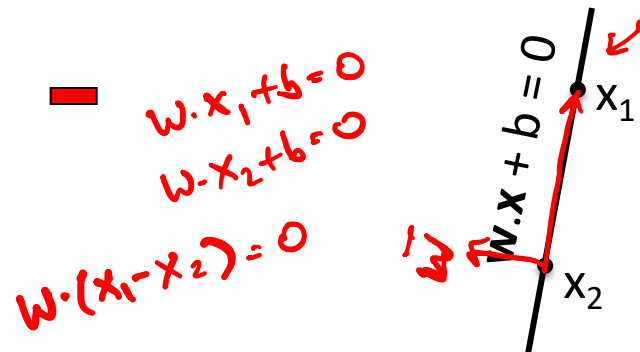
Maximizing the margin



Distance of closest examples from the line/hyperplane

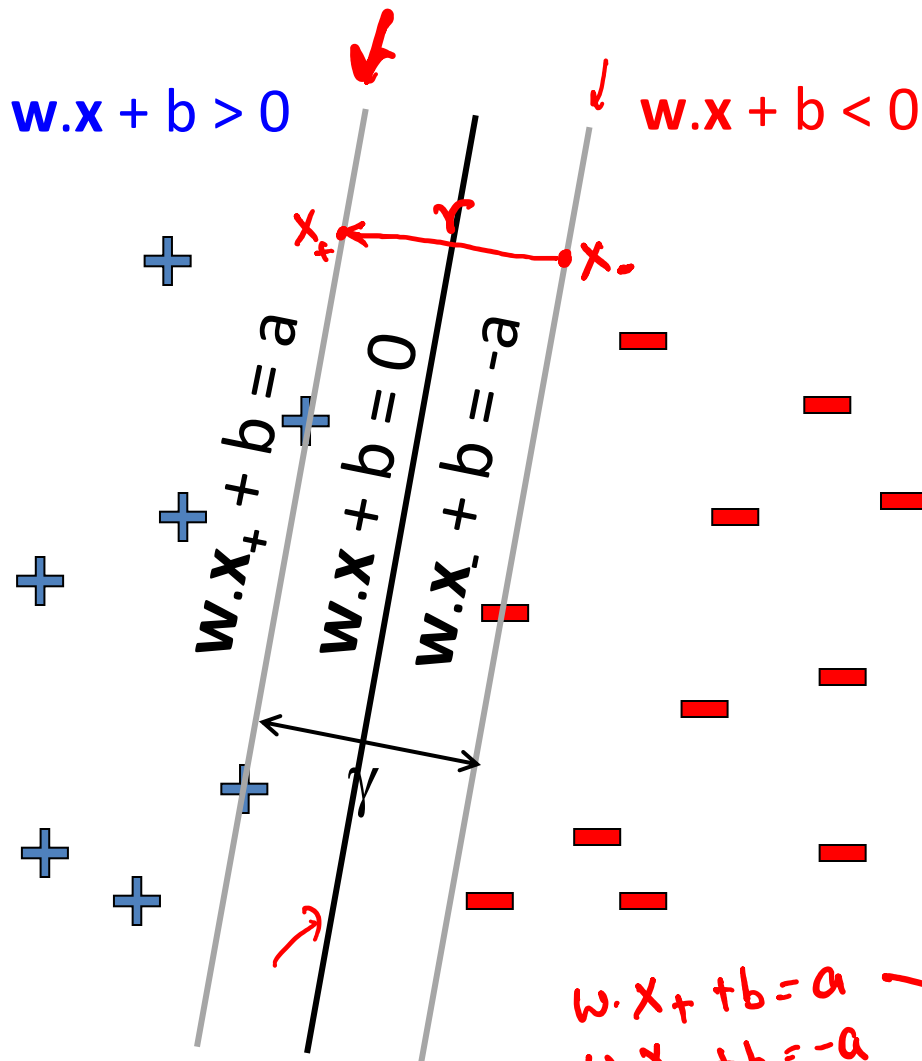
$$\text{margin} = \gamma = 2a / \|w\|$$

Step 1: w is perpendicular to lines since for any x_1, x_2 on line $w \cdot (x_1 - x_2) = 0$





Maximizing the margin



$$\text{margin} = \gamma = 2a/\|w\|$$

Step 1: w is perpendicular to lines

Step 2: Take a point x_- on $w.x + b = -a$ and move to point x_+ that is γ away on line $w.x + b = a$

$$x_+ = x_- + \gamma w / \|w\|$$

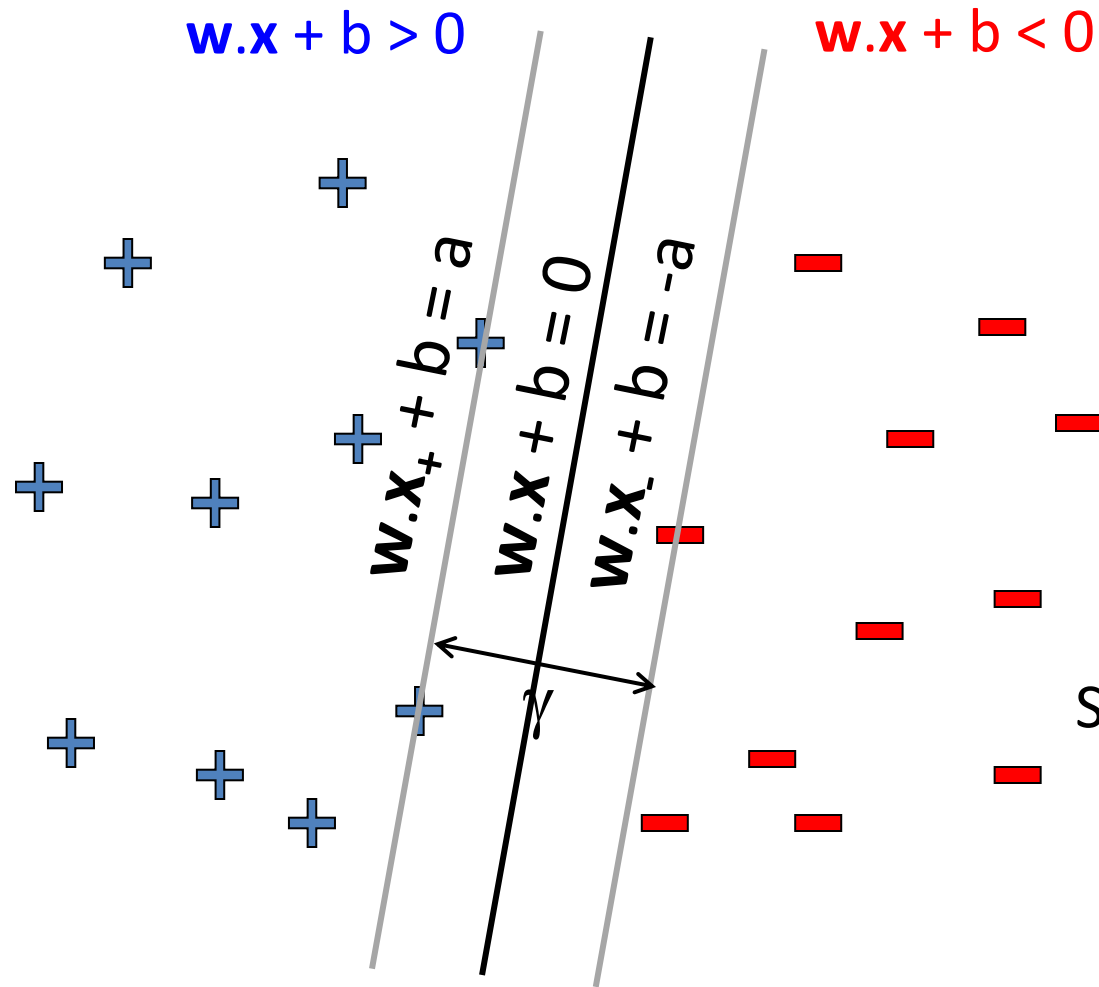
$$w \cdot x_+ = w \cdot x_- + \gamma w \cdot w / \|w\|$$

$$\begin{aligned} w \cdot x_+ + b &= a \\ w \cdot x_- + b &= -a \end{aligned}$$

$$a - b = -a - b + \gamma \|w\|$$

$$2a = \gamma \|w\|$$

Maximizing the margin



Distance of closest examples from the line/hyperplane

$$\text{margin} = \gamma = 2a / \|w\|$$

$$2 / \|w\|$$

Smaller margin $\Leftrightarrow$ larger $\|w\|$

Maximizing the margin

$$w \cdot x + b > 0$$

$$w \cdot x + b < 0$$

Distance of closest examples from the line/hyperplane

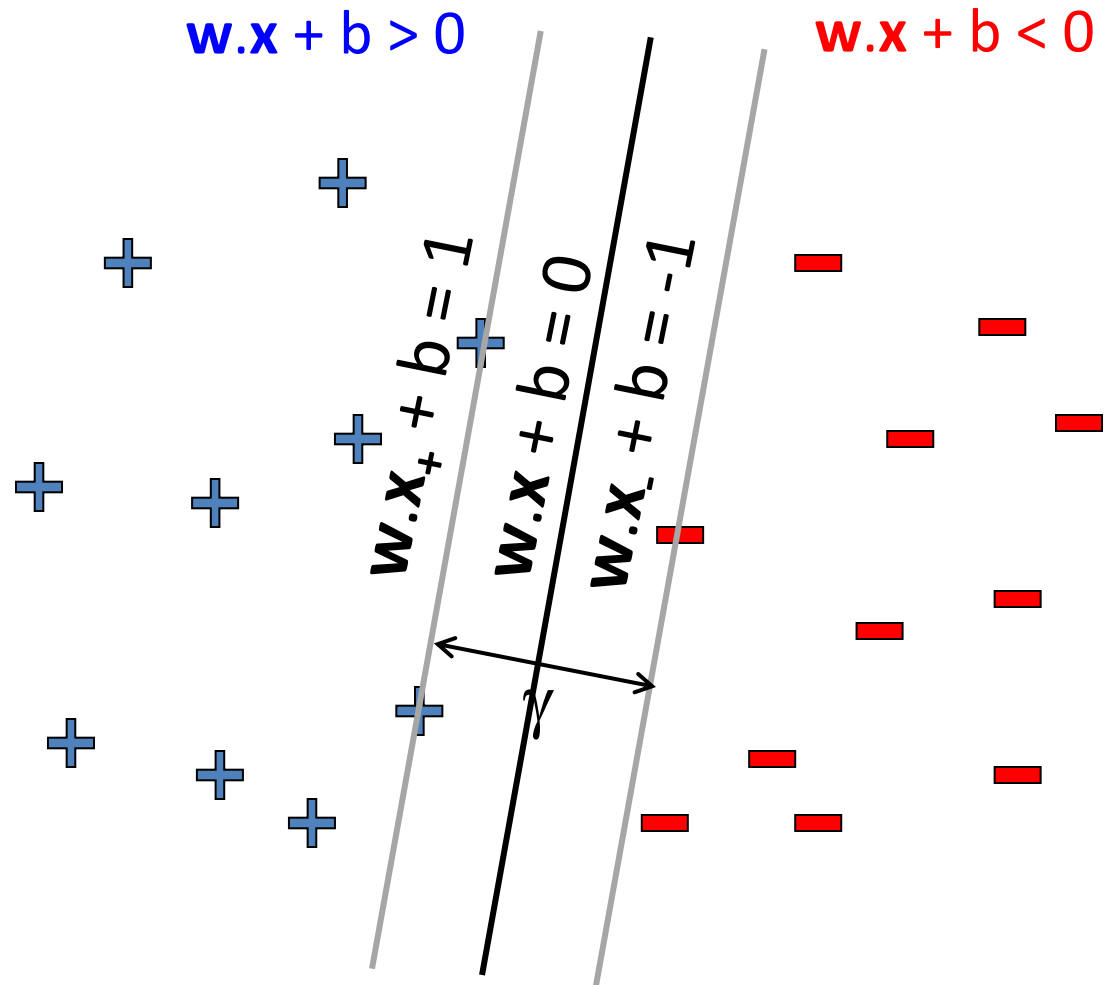
$$\text{margin} = \gamma = 2a / \|w\|$$

$$\max_{w,b} \gamma = 2a / \|w\|$$

$$\text{s.t. } (w \cdot x_j + b) y_j \geq a \quad \forall j$$

Note: 'a' is arbitrary (can normalize equations by a)

Support Vector Machines



$$\min_{\underline{w}, b} w \cdot w$$

$$\text{s.t. } (w \cdot x_j + b) y_j \geq 1 \quad \forall j$$

Solve efficiently by quadratic programming (QP)

- Quadratic objective, linear constraints
- Well-studied solution algorithms

Support Vectors

w, b $\leftarrow$
d-line

$$w \cdot x + b > 0$$

$$w \cdot x + b < 0$$

Linear hyperplane defined by
“support vectors” $\leftarrow$

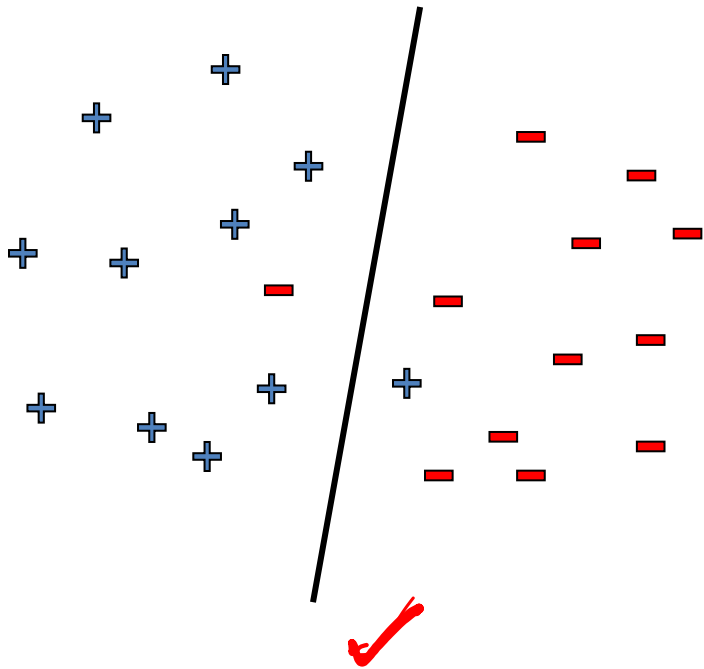
Moving other points a little
doesn't effect the decision
boundary

only need to store the
support vectors to predict
labels of new points

For support vectors
 $(w \cdot x_j + b) y_j = 1$

What if data is not linearly separable?

Use features of features
of features of features....

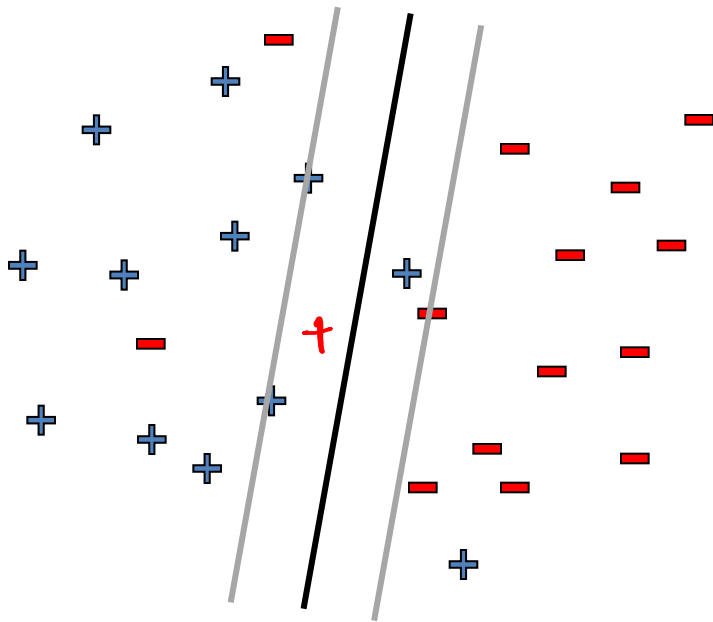


$$x_1^2, x_2^2, x_1x_2, \dots, \exp(x_1)$$

But run risk of overfitting!

What if data is still not linearly separable?

Allow “error” in classification



Smaller margin $\Leftrightarrow$ larger $\|w\|$

$$\min_{w,b} \underbrace{w \cdot w} + C \sum_j \underbrace{1_{(w \cdot x_j + b) y_j < 0}}_{\xi_j}$$

Handwritten notes: ξ_j is written above the summation term, and $\sum_j 1_{(w \cdot x_j + b) y_j < 0}$ is written to the right of the summation term.

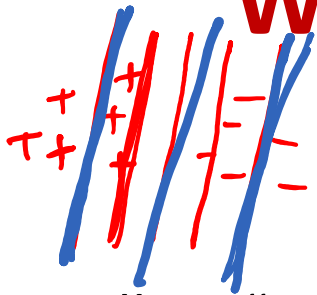
Maximize margin and minimize
mistakes on training data

C - tradeoff parameter

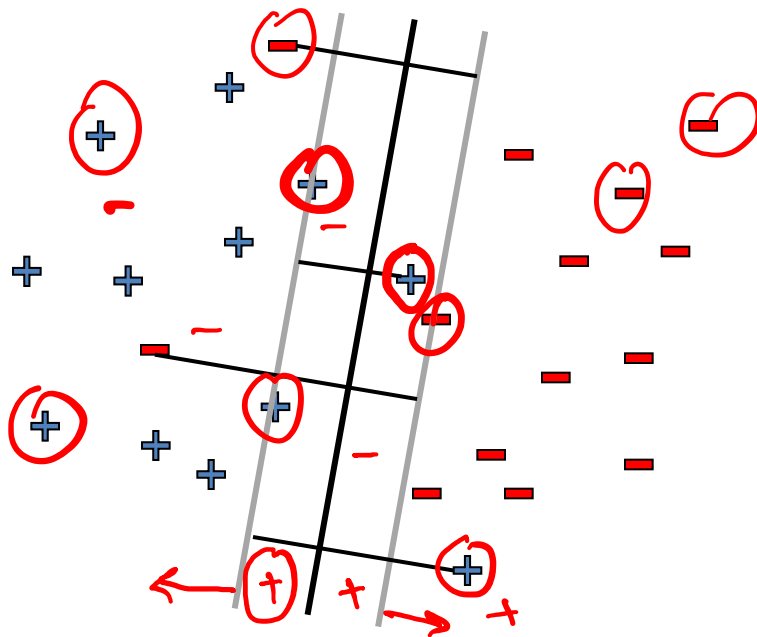
Not QP ☹

0/1 loss (doesn't distinguish between
near miss and bad mistake)

What if data is still not linearly separable?



Allow "error" in classification



Soft margin approach

$$\begin{aligned} \min_{\mathbf{w}, b, \{\xi_j\}} \quad & \mathbf{w} \cdot \mathbf{w} + C \sum_j \xi_j \\ \text{s.t.} \quad & (\mathbf{w} \cdot \mathbf{x}_j + b) y_j \geq 1 - \xi_j \quad \forall j \\ & \xi_j \geq 0 \quad \forall j \end{aligned}$$

→ ξ_j - "slack" variables
= (>1 if x_j misclassified)

pay linear penalty if mistake

C - tradeoff parameter ($C = \infty$ recovers hard margin SVM)

Still QP 😊

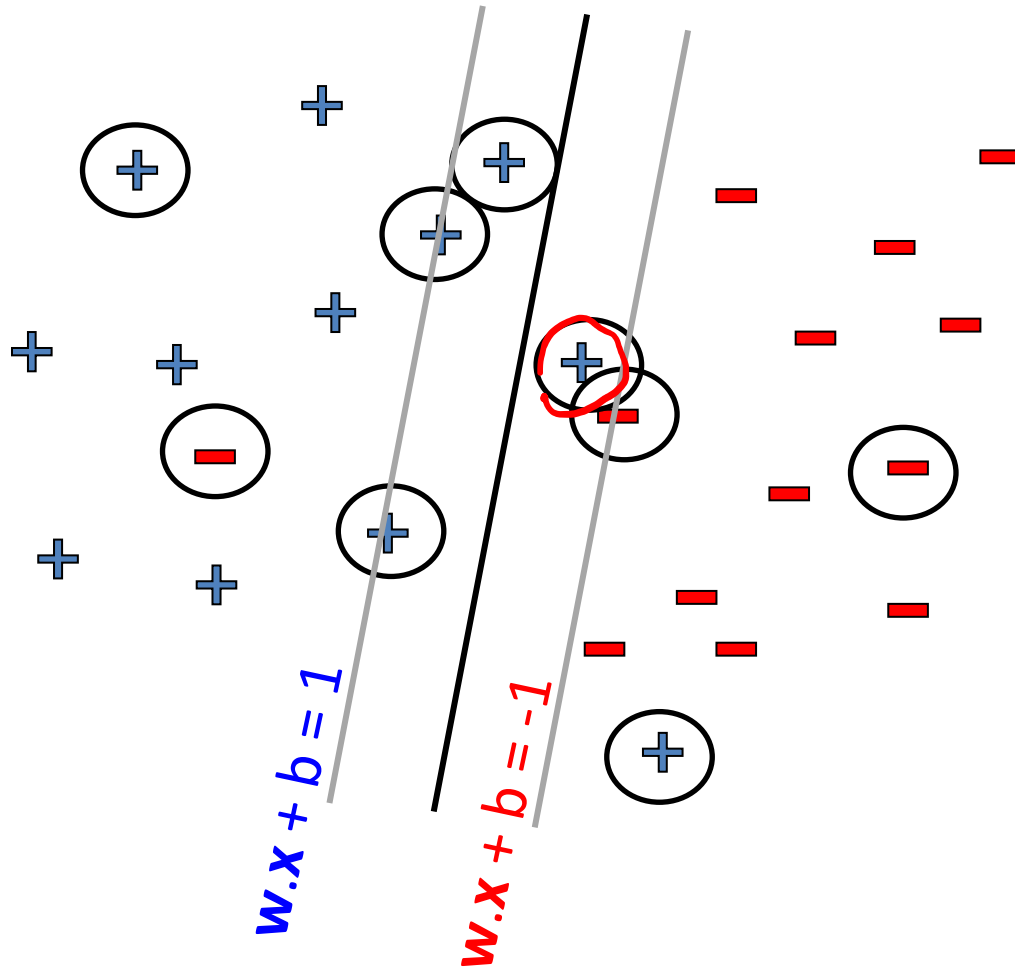
$$\begin{aligned}
 \min_{\mathbf{w}, b, \{\xi_j\}} \quad & \mathbf{w} \cdot \mathbf{w} + C \sum \xi_j \\
 \text{s.t.} \quad & (\mathbf{w} \cdot \mathbf{x}_j + b) y_j \geq 1 - \xi_j \quad \forall j \\
 & \xi_j \geq 0 \quad \forall j
 \end{aligned}$$

Slack variables

$$\begin{aligned}
 2 &\geq 1 - \xi_j \\
 1 &\geq 1 - \xi_j
 \end{aligned}$$

$$(\mathbf{w} \cdot \mathbf{x}_j + b) y_j \geq 1 - \xi_j \quad \forall j$$

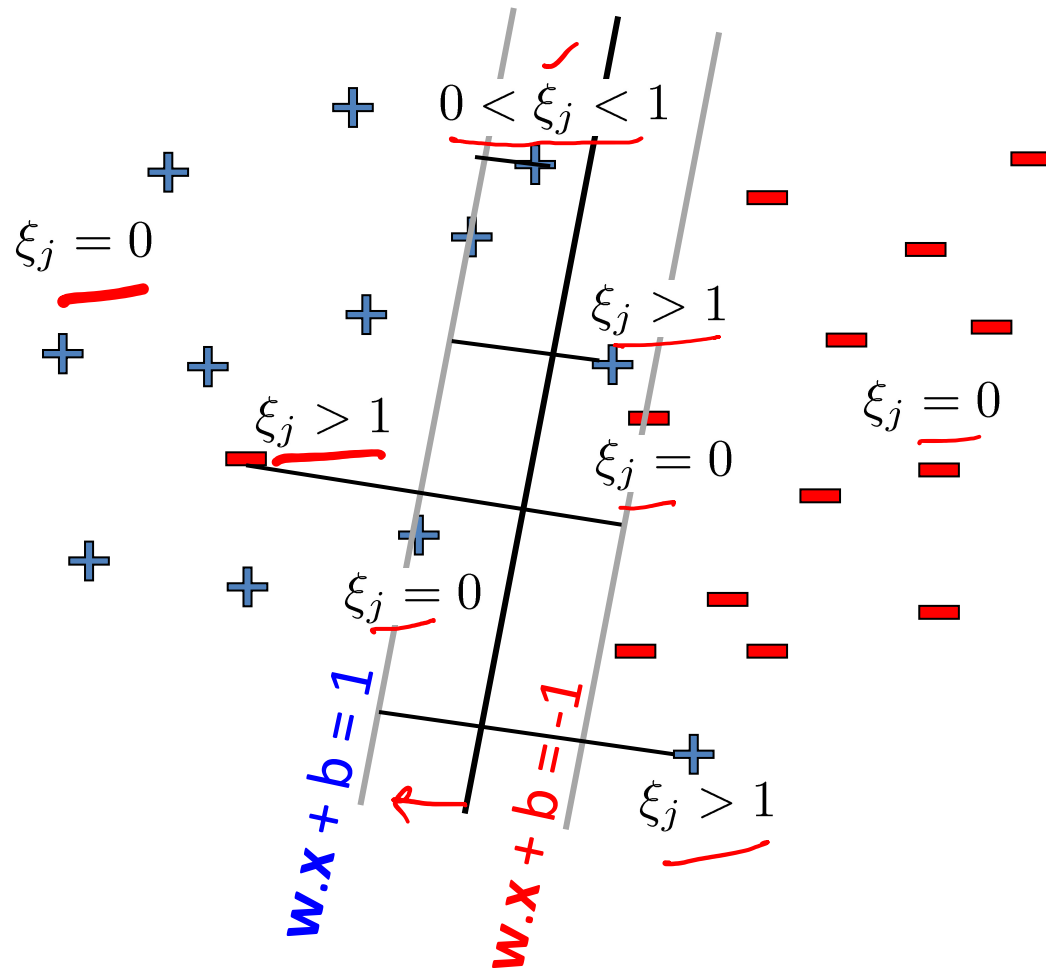
What is the slack ξ_j for the following points?



Confidence $(\mathbf{w} \cdot \mathbf{x}_j + b) y_j$	Slack ξ_j
1	0 ✓
> 1	0 ✓
0 - 1	0 - 1 ✓
< 0	> 1

$$a_+ = \max(a, 0)$$

Slack variables – Hinge loss

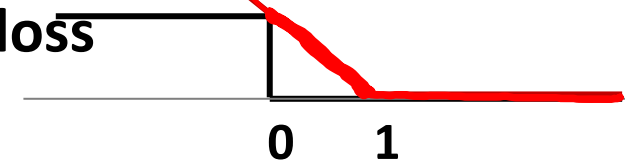


Notice that

$$\xi_j = (1 - (\mathbf{w} \cdot \mathbf{x}_j + b)y_j))_+$$

Hinge loss

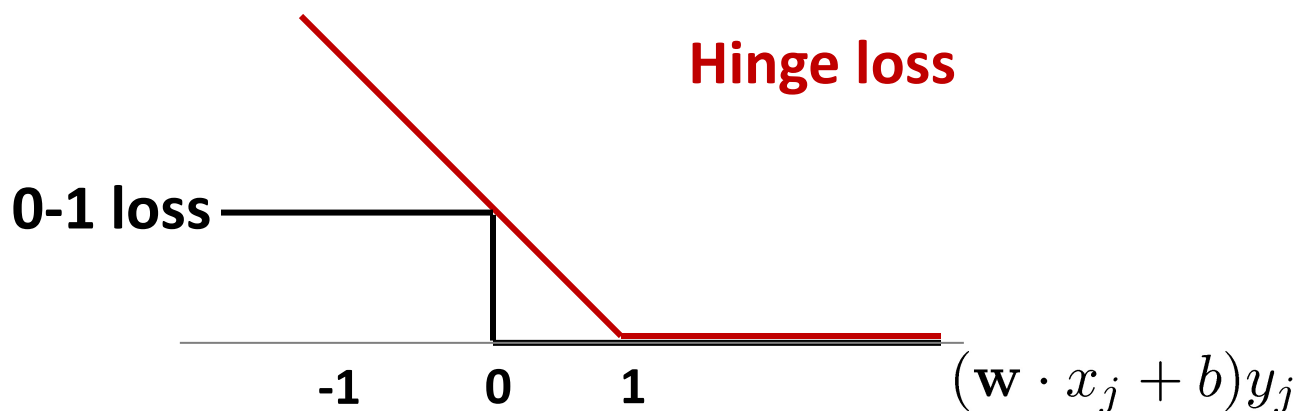
0-1 loss



$$(\mathbf{w} \cdot \mathbf{x}_j + b)y_j$$

Slack variables – Hinge loss

$$\xi_j = (1 - (\mathbf{w} \cdot \mathbf{x}_j + b)y_j))_+ \quad \leftarrow$$



$$\min_{\mathbf{w}, b, \{\xi_j\}} \mathbf{w} \cdot \mathbf{w} + C \sum_j \xi_j$$

$$\text{s.t. } (\mathbf{w} \cdot \mathbf{x}_j + b) y_j \geq 1 - \xi_j \quad \checkmark \quad \forall j$$

$$\xi_j \geq 0 \quad \checkmark \quad \forall j$$



Regularized hinge loss

$$\min_{\mathbf{w}, b} \underbrace{\mathbf{w} \cdot \mathbf{w}}_{\|\mathbf{w}\|^2} + C \sum_j \underbrace{(1 - (\mathbf{w} \cdot \mathbf{x}_j + b)y_j)_+}_{\xi_j}$$

$\sum_j \text{loss}(f_{\mathbf{w}, b}(\mathbf{x}_j), y_j)$

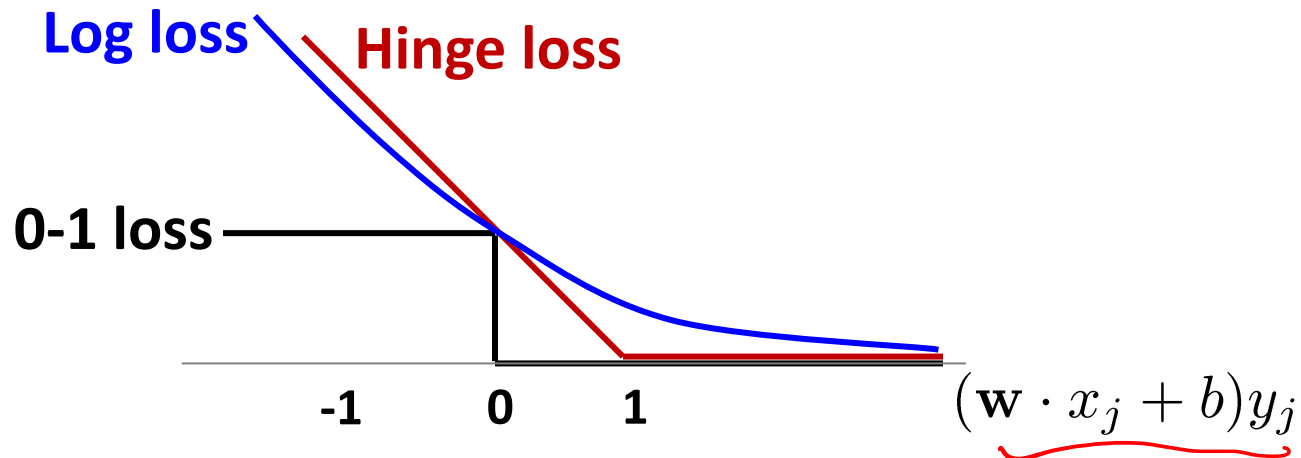
SVM vs. Logistic Regression

SVM : **Hinge loss**

$$\text{loss}(f(x_j), y_j) = (1 - (\mathbf{w} \cdot x_j + b)y_j)_+$$

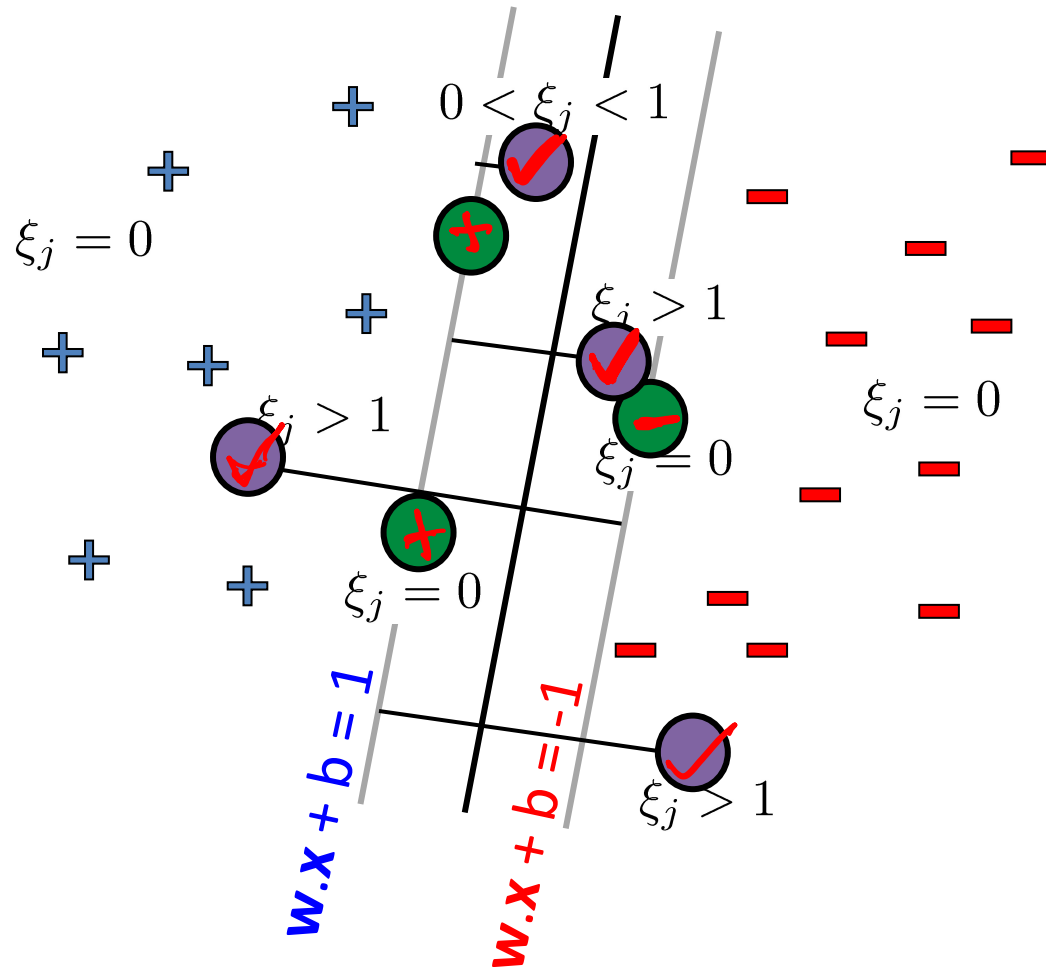
Logistic Regression : **Log loss** (-ve log conditional likelihood)

$$\text{loss}(f(x_j), y_j) = \underbrace{-\log P(y_j \mid x_j, \mathbf{w}, b)}_{\text{Log loss}} = \underbrace{\log(1 + e^{-(\mathbf{w} \cdot x_j + b)y_j})}_{\text{Log loss}}$$



$$\begin{aligned}
 \min_{\mathbf{w}, b, \{\xi_j\}} \quad & \mathbf{w} \cdot \mathbf{w} + C \sum \xi_j \\
 \text{s.t.} \quad & (\mathbf{w} \cdot \mathbf{x}_j + b) y_j \geq 1 - \xi_j \quad \forall j \\
 & \xi_j \geq 0 \quad \forall j
 \end{aligned}$$

Support Vectors



Margin support vectors

$\xi_j = 0$, $(\mathbf{w} \cdot \mathbf{x}_j + b) y_j = 1$ ✓
(don't contribute to objective but enforce constraints on solution)

Correctly classified but on margin

Non-margin support vectors

→ $\xi_j > 0$
(contribute to both objective and constraints)

$1 > \xi_j > 0$ Correctly classified but inside margin
 $\xi_j > 1$ Incorrectly classified