

Support Vector Machines (SVMs)

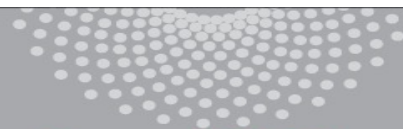
Aarti Singh

Machine Learning 10-315

Oct 20, 2021



MACHINE LEARNING DEPARTMENT



Carnegie Mellon.
School of Computer Science

Discriminative Classifiers

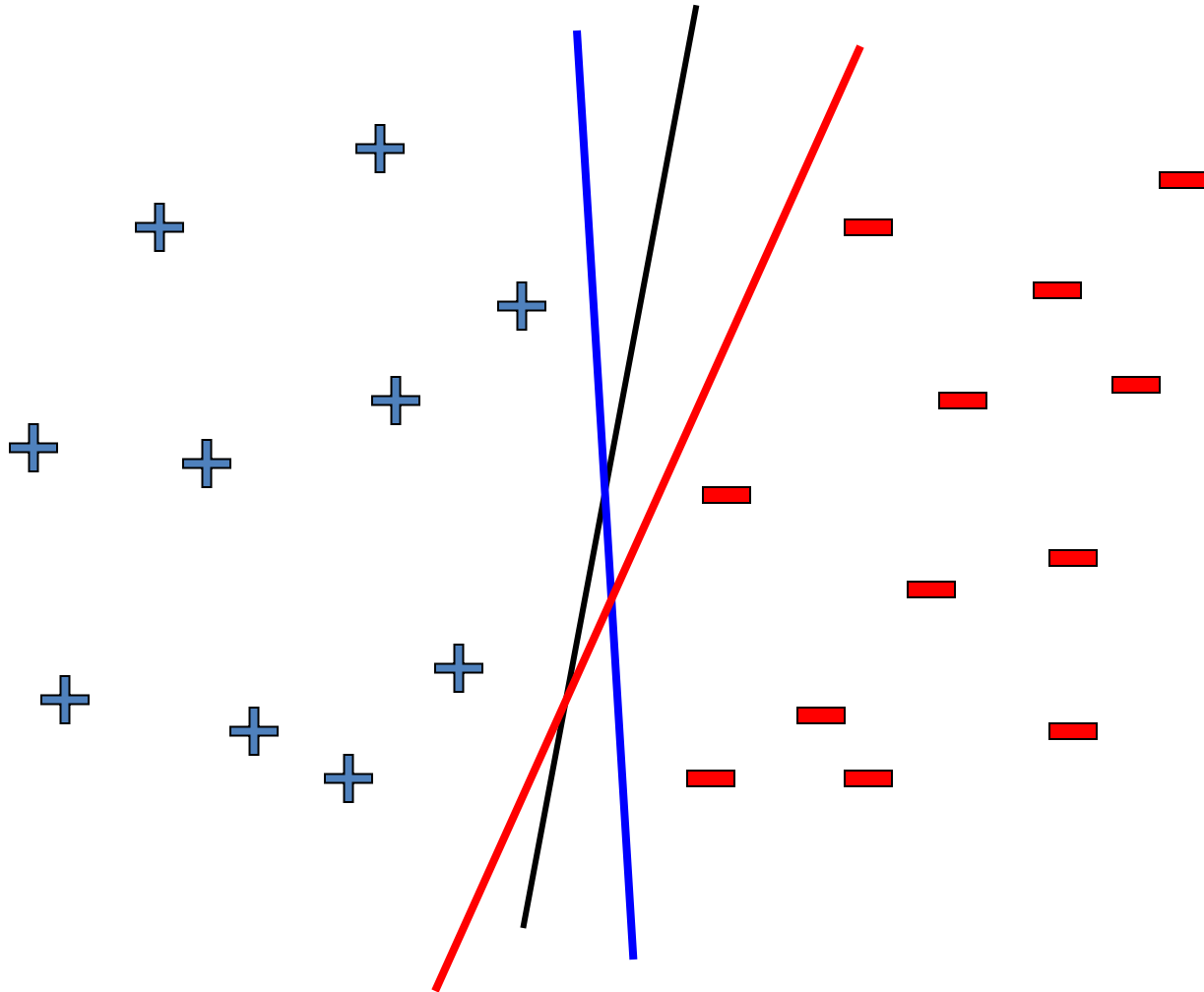
Optimal Classifier:

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

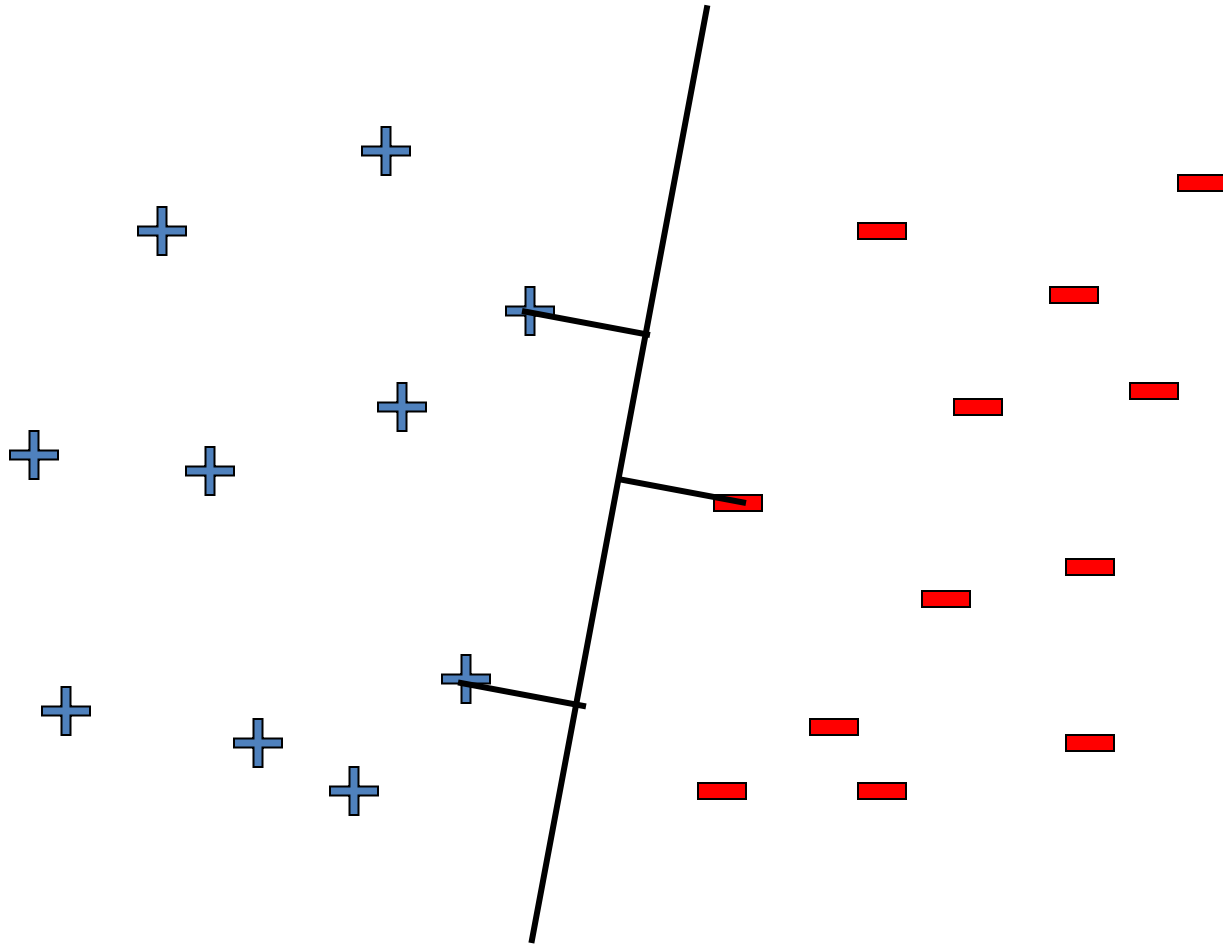
Why not learn $P(Y|X)$ directly? Or better yet, why not learn the decision boundary directly?

- 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

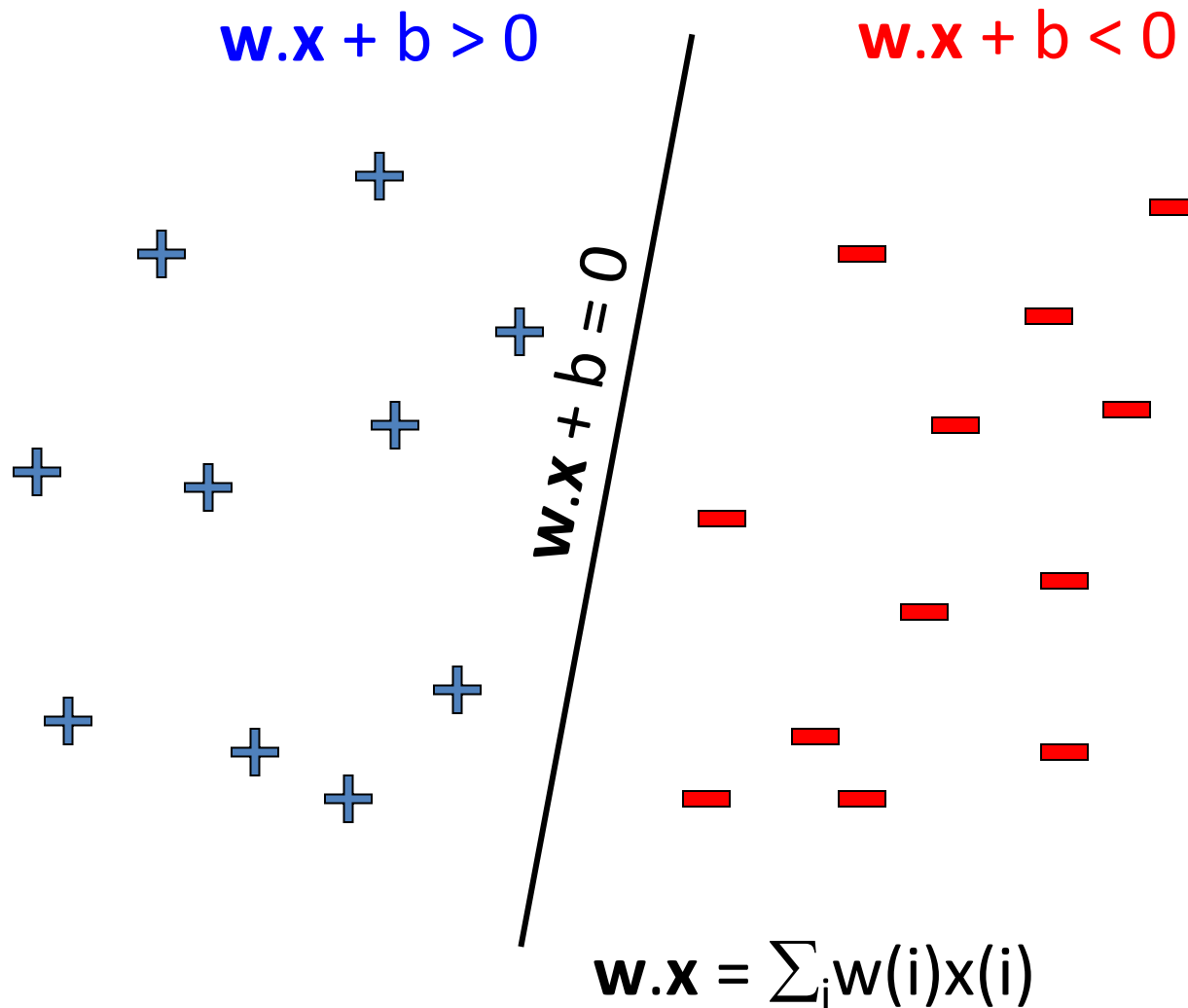
Linear classifiers – which line is better?



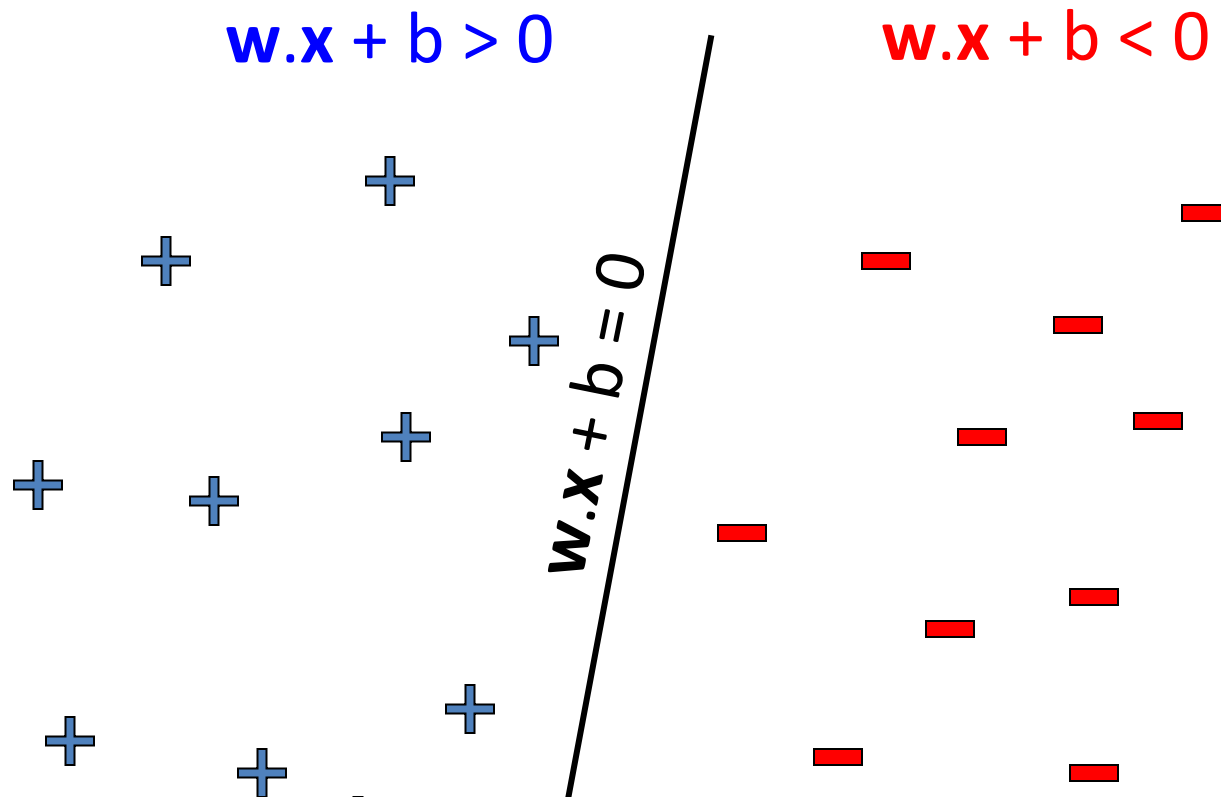
Pick the one with the largest margin!



Parameterizing the decision boundary



Parameterizing the decision boundary



$y_j \in \{-1, +1\}$ — class

$$\text{"confidence"} = (w \cdot x_j + b) y_j$$

Maximizing the margin

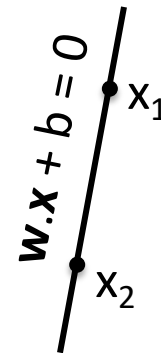
$$\mathbf{w} \cdot \mathbf{x} + b > 0$$

$$\mathbf{w} \cdot \mathbf{x} + b < 0$$

Distance of closest examples from the line/hyperplane

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

Step 1: $\mathbf{w}$ is perpendicular to lines since for any $\mathbf{x}_1, \mathbf{x}_2$ on line $\mathbf{w} \cdot (\mathbf{x}_1 - \mathbf{x}_2) = 0$



Maximizing the margin

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

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

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

Step 1: w is perpendicular to lines

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

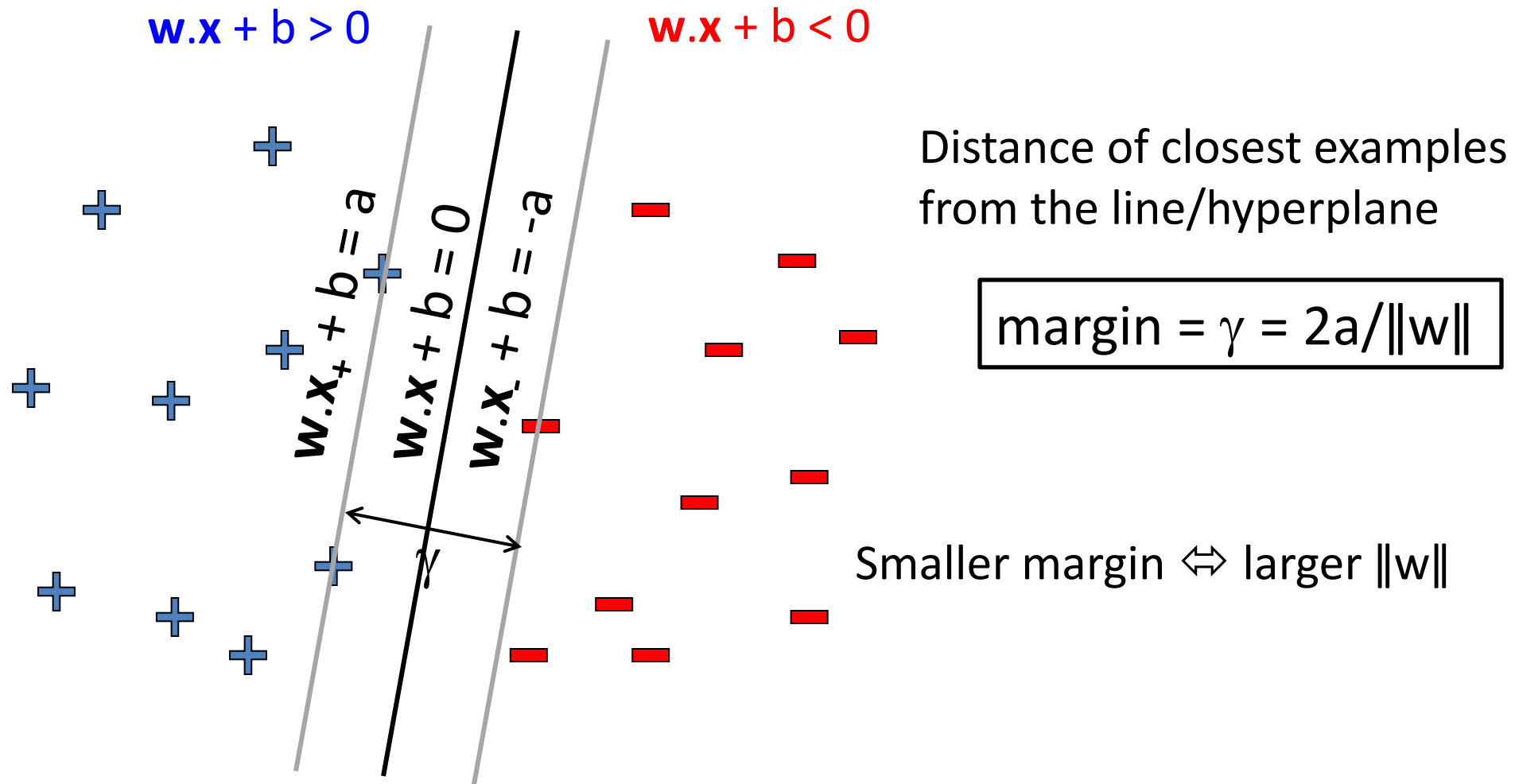
$$x_+ = x_- + \gamma w / \|w\|$$

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

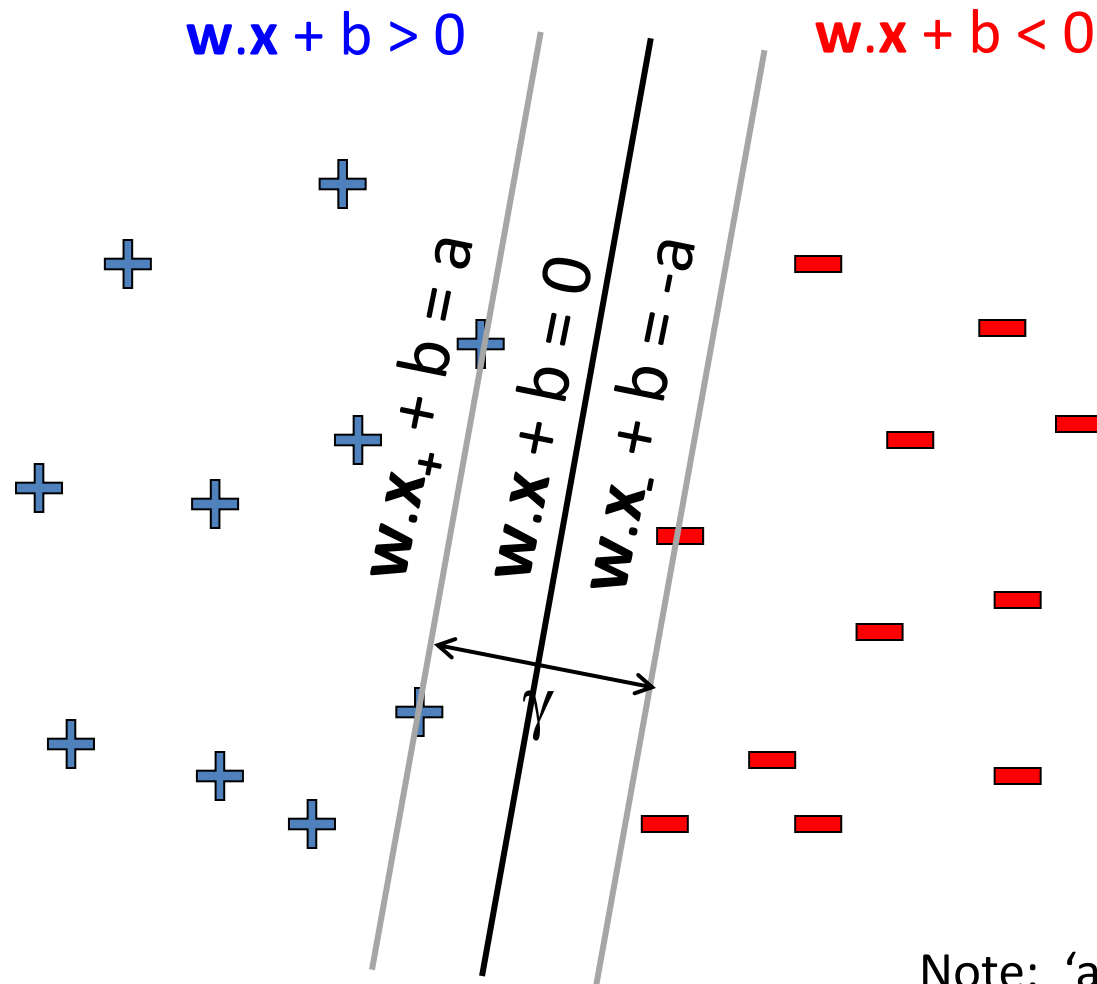
$$a - b = -a - b + \gamma \|w\|$$

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

Maximizing the margin



Maximizing the margin



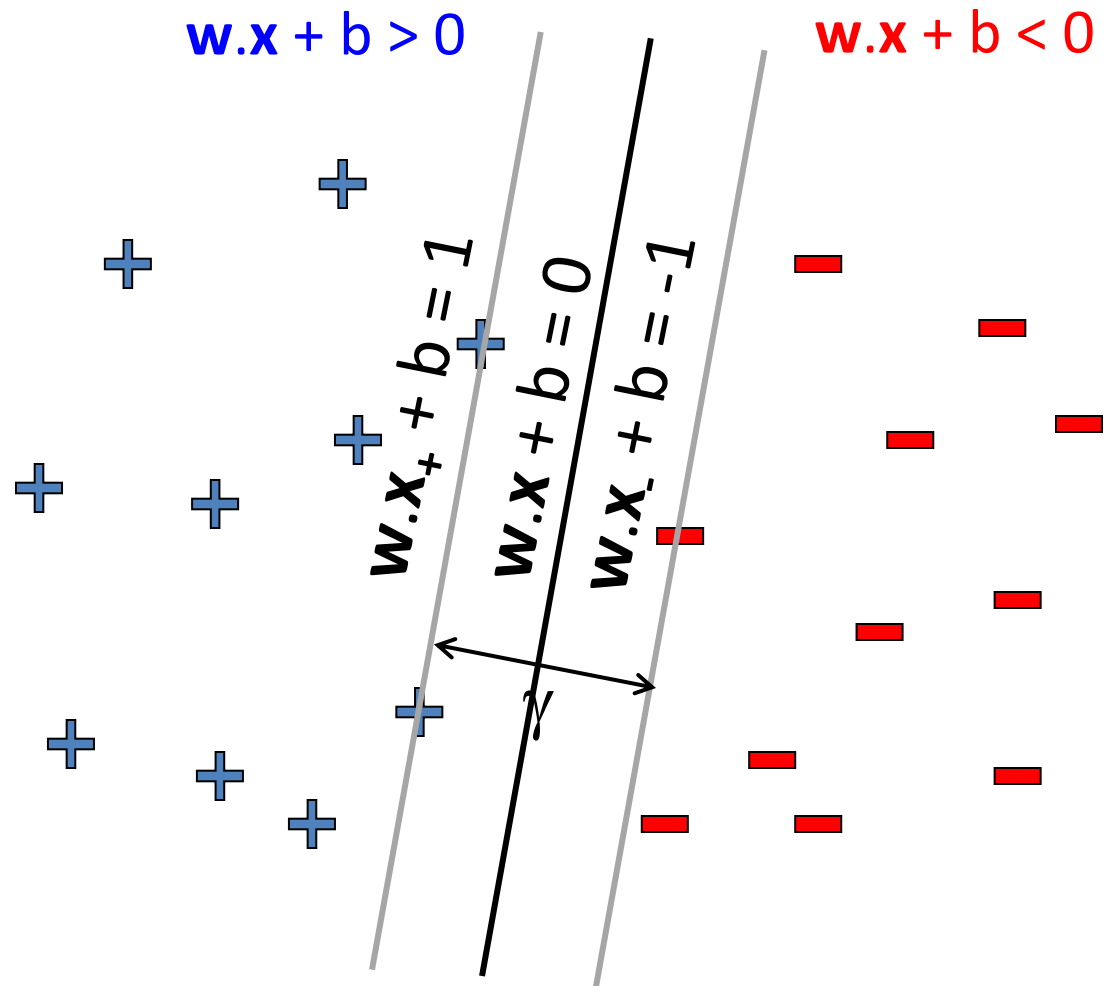
Distance of closest examples from the line/hyperplane

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

$$\begin{aligned} \max_{w, b} \quad & \gamma = 2a / \|w\| \\ \text{s.t.} \quad & (w \cdot x_j + b) y_j \geq a \quad \forall j \end{aligned}$$

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

Support Vector Machines



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

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

Solve efficiently by quadratic programming (QP)

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

Support Vectors

$$\mathbf{w} \cdot \mathbf{x} + b > 0$$

$$\mathbf{w} \cdot \mathbf{x} + b < 0$$

Linear hyperplane defined by
“support vectors”

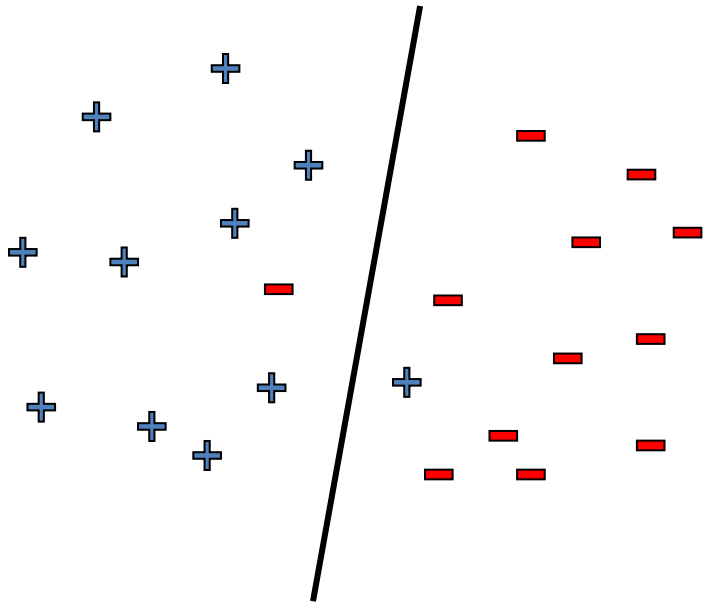
Moving other points a little
doesn't effect the decision
boundary

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

For support vectors
 $(\mathbf{w} \cdot \mathbf{x}_j + b) y_j = 1$

What if data is not linearly separable?

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

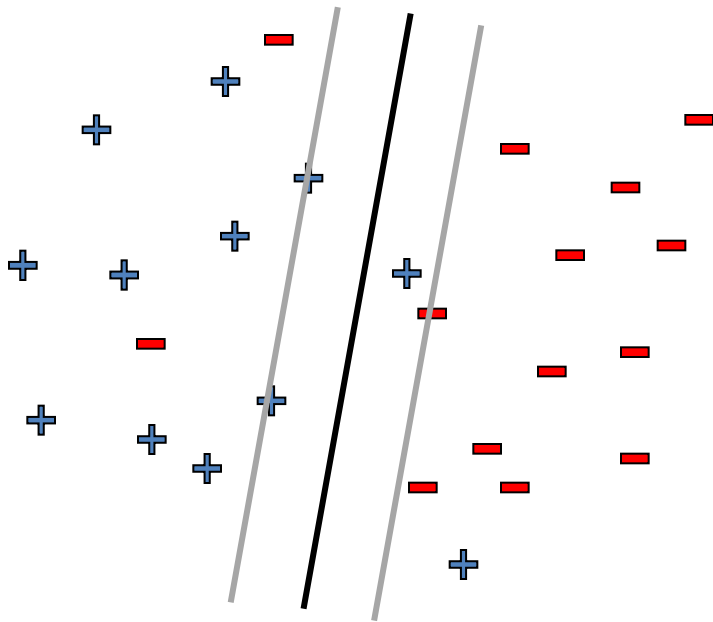


$$x_1^2, x_2^2, x_1x_2, \dots, \exp(x_1)$$

But run risk of overfitting!

What if data is still not linearly separable?

Allow “error” in classification



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

$$\begin{aligned} \min_{w,b} \quad & w \cdot w + C \# \text{mistakes} \\ \text{s.t.} \quad & (w \cdot x_j + b) y_j \geq 1 \quad \forall j \end{aligned}$$

Maximize margin and minimize
mistakes on training data

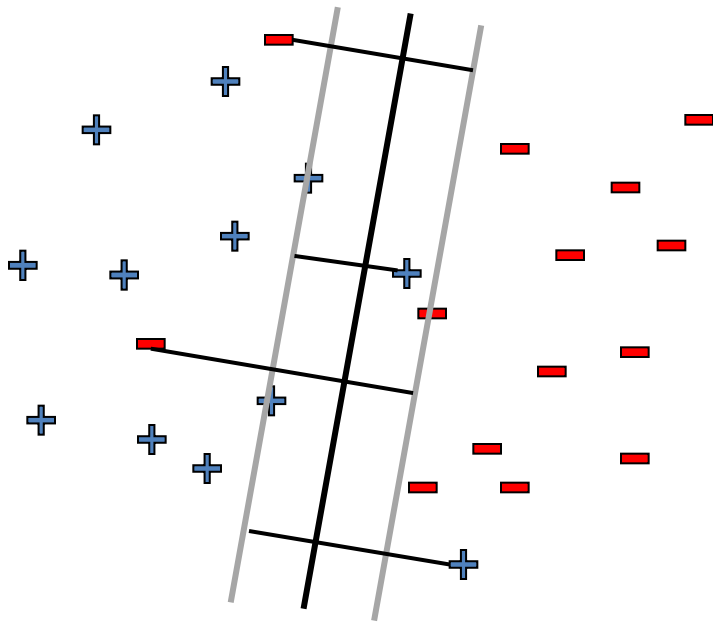
C - tradeoff parameter

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 😊

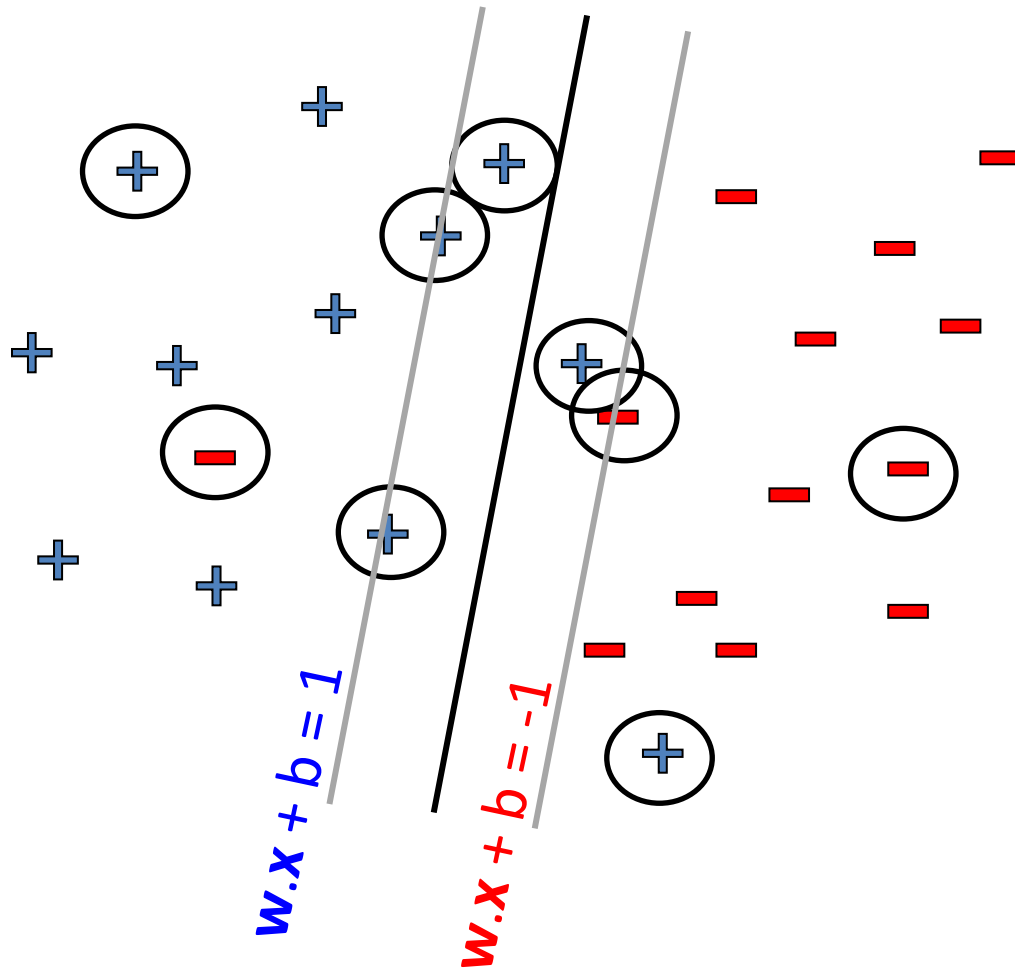
$$\begin{aligned}
 \min_{\mathbf{w}, b, \{\xi_j\}} \quad & \mathbf{w} \cdot \mathbf{w} + C \sum \xi_j \\
 \text{s.t.} \quad & (\mathbf{w} \cdot \mathbf{x}_j + b) y_j \geq 1 - \xi_j \quad \forall j \\
 & \xi_j \geq 0 \quad \forall j
 \end{aligned}$$

Slack variables

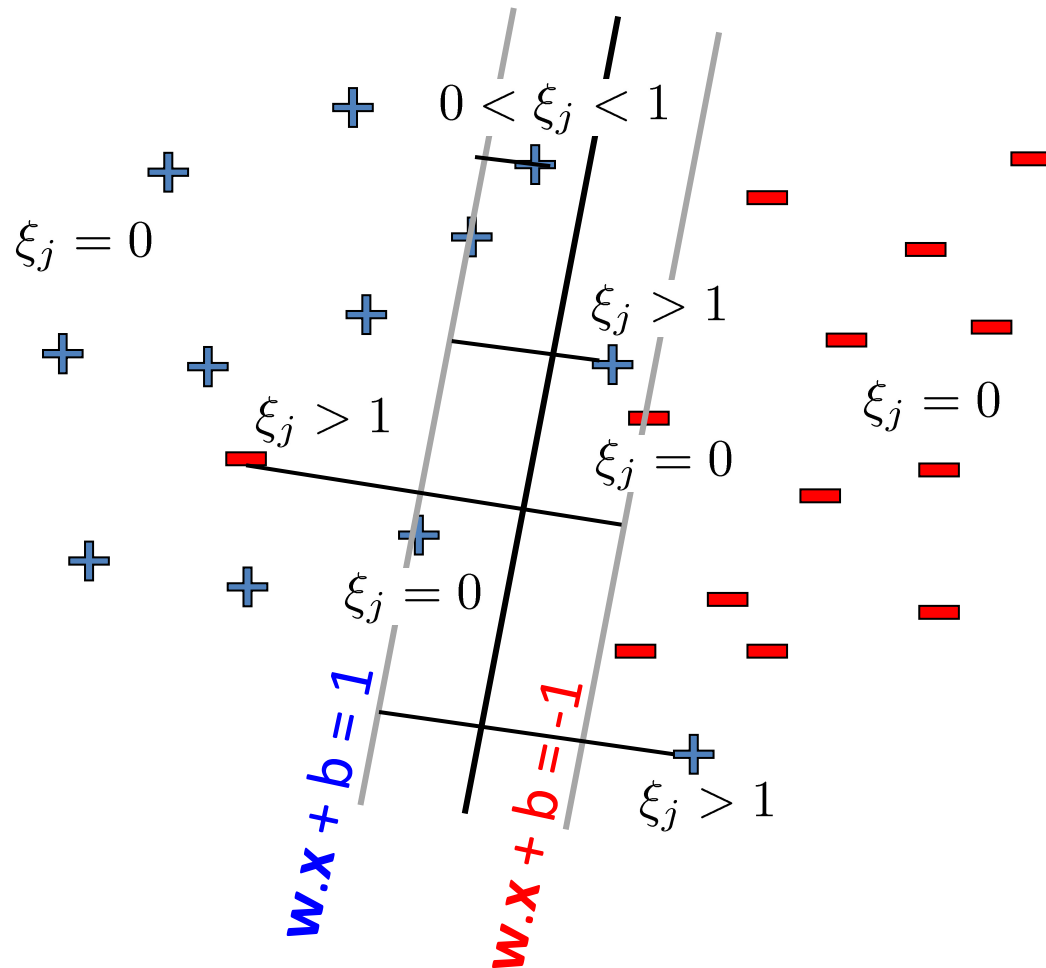
$$(\mathbf{w} \cdot \mathbf{x}_j + b) y_j \geq 1 - \xi_j \quad \forall j$$

What is the slack ξ_j for the following points?

Confidence	Slack
------------	-------

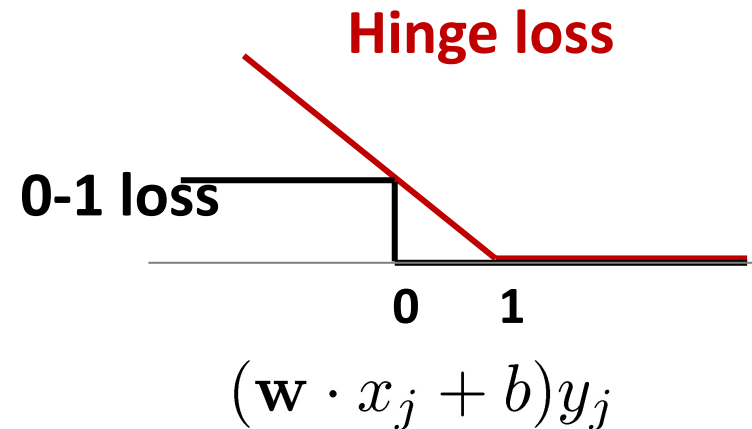


Slack variables – Hinge loss



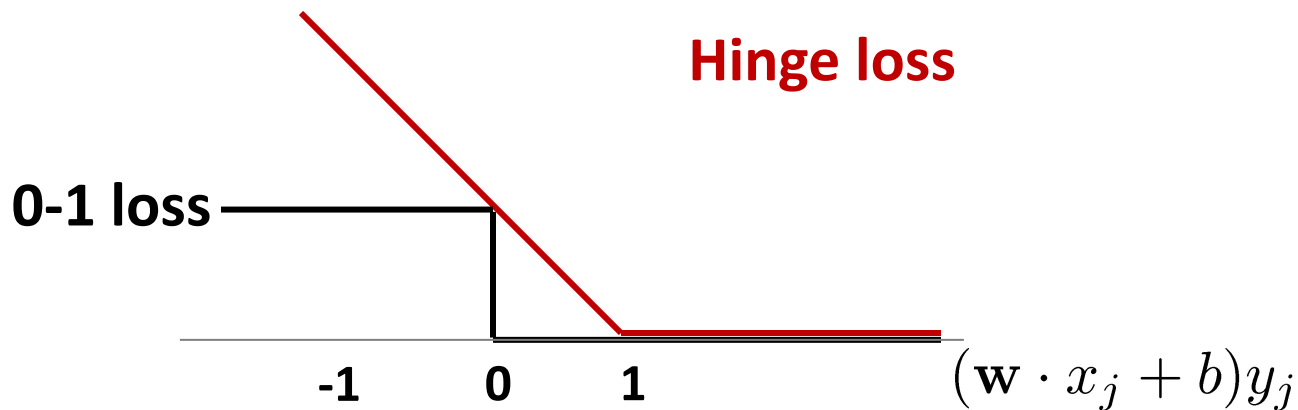
Notice that

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



Slack variables – Hinge loss

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



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

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

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



Regularized hinge loss

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

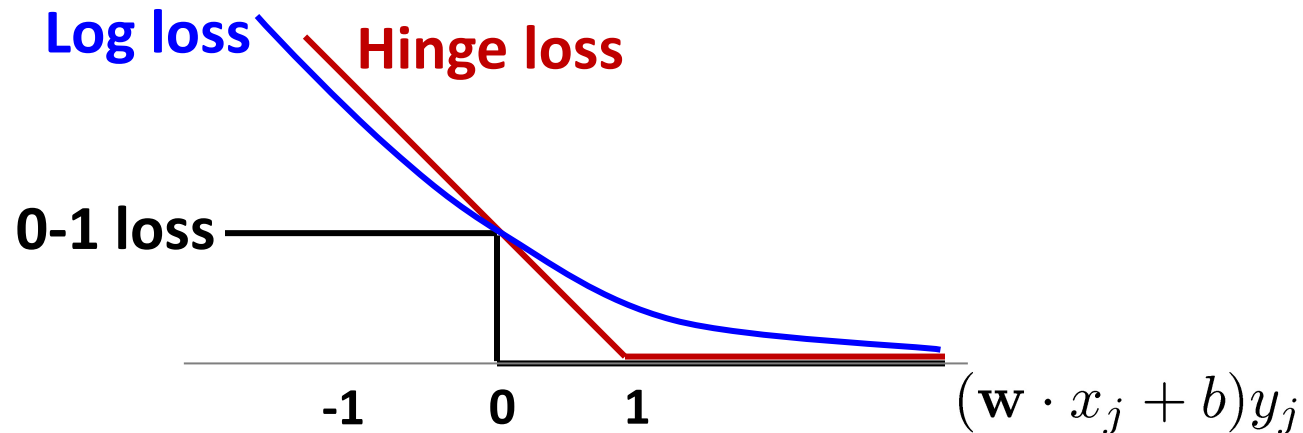
SVM vs. Logistic Regression

SVM : **Hinge loss**

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

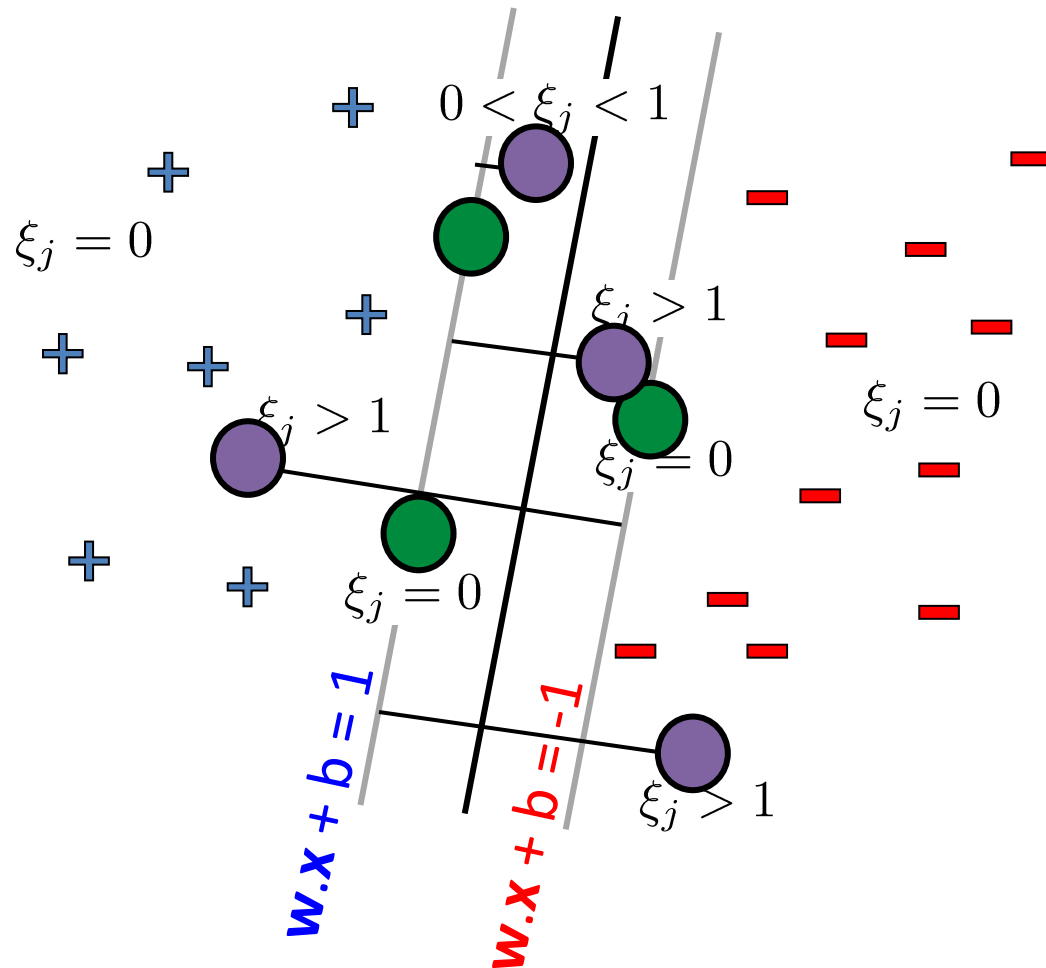
Logistic Regression : **Log loss** (-ve log conditional likelihood)

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



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

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

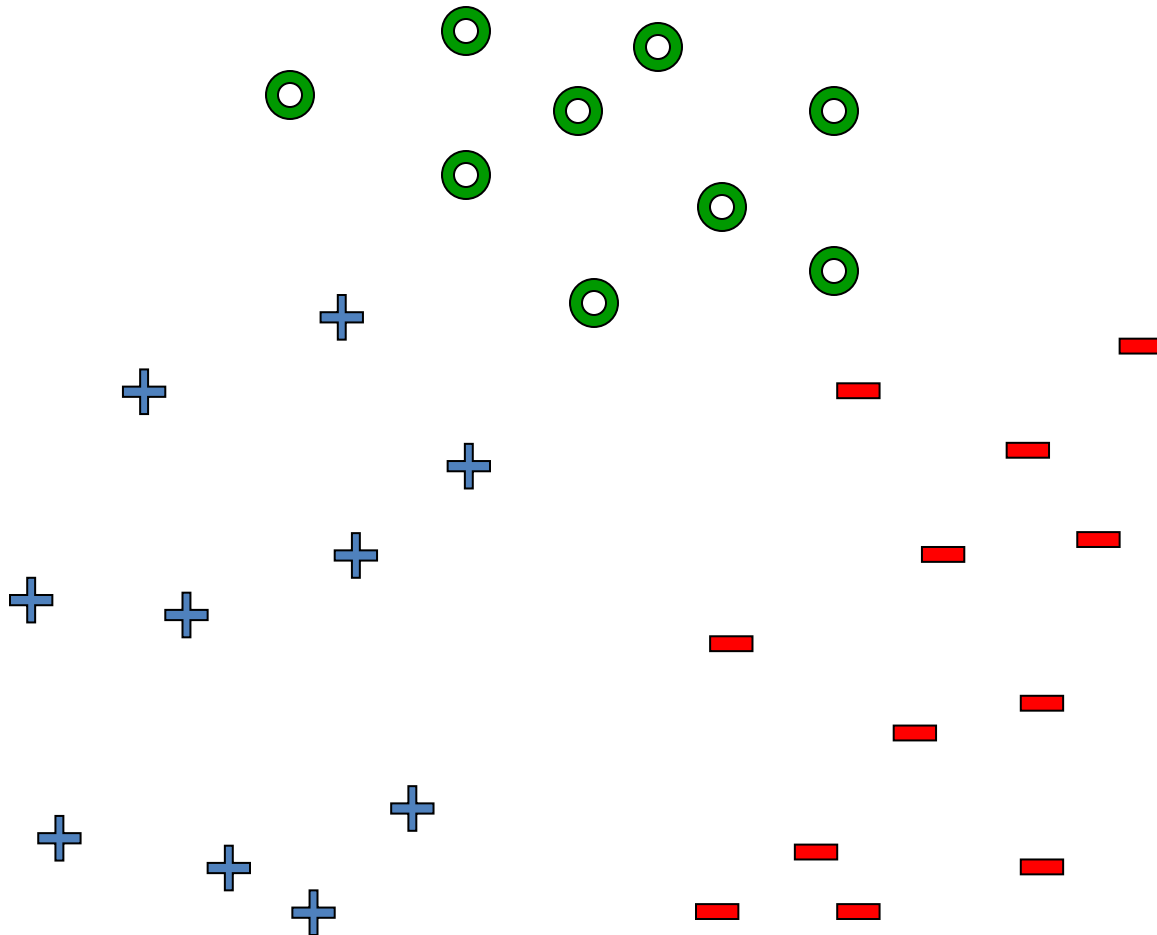
Non-margin support vectors

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

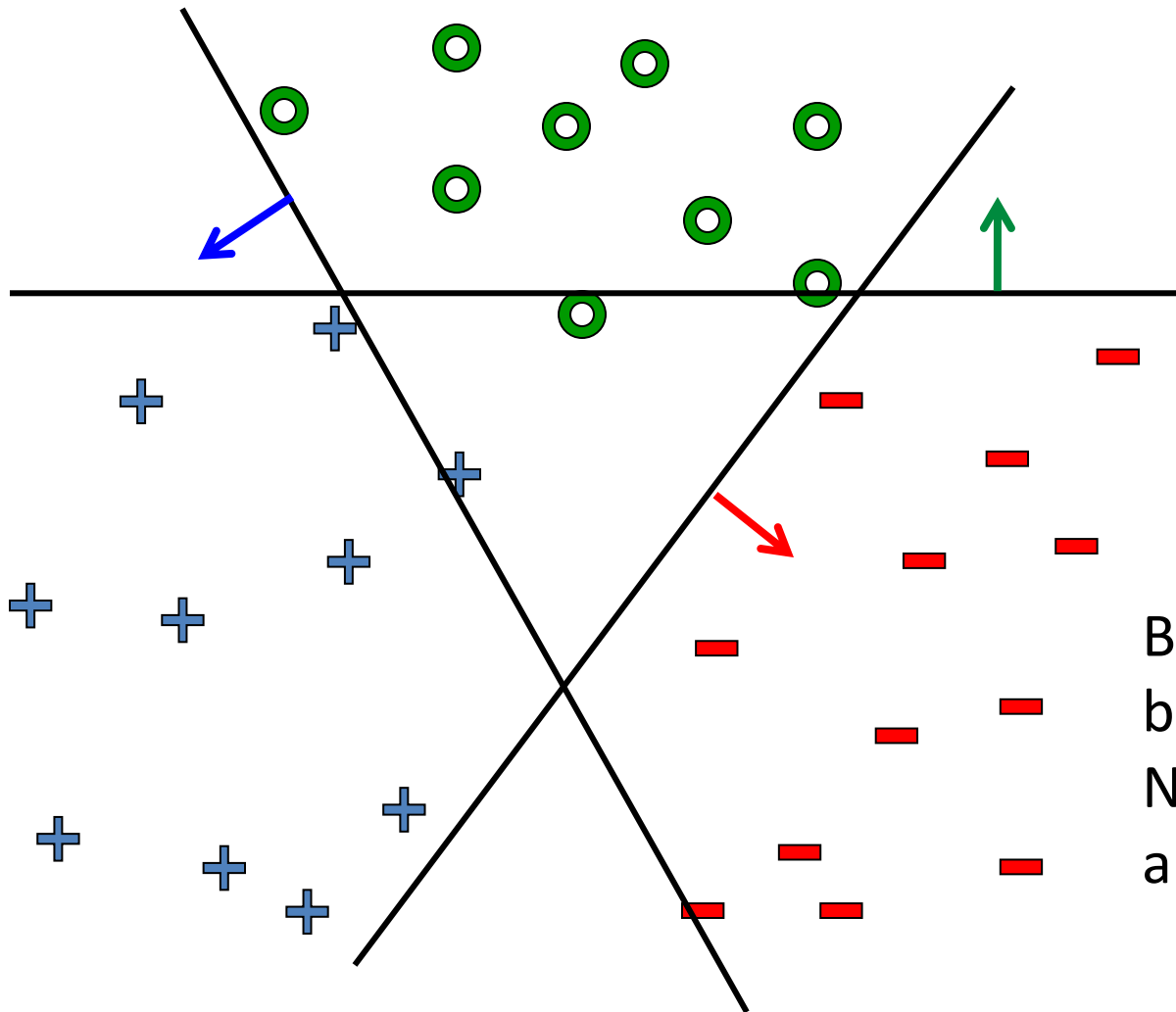
$1 > \xi_j > 0$ Correctly classified but inside margin

$\xi_j > 1$ Incorrectly classified

What about multiple classes?



One vs. rest



Learn 3 classifiers
separately:

Class k vs. rest

$$(\mathbf{w}_k, b_k)_{k=1,2,3}$$

$$y = \arg \max_k \mathbf{w}_k \cdot \mathbf{x} + b_k$$

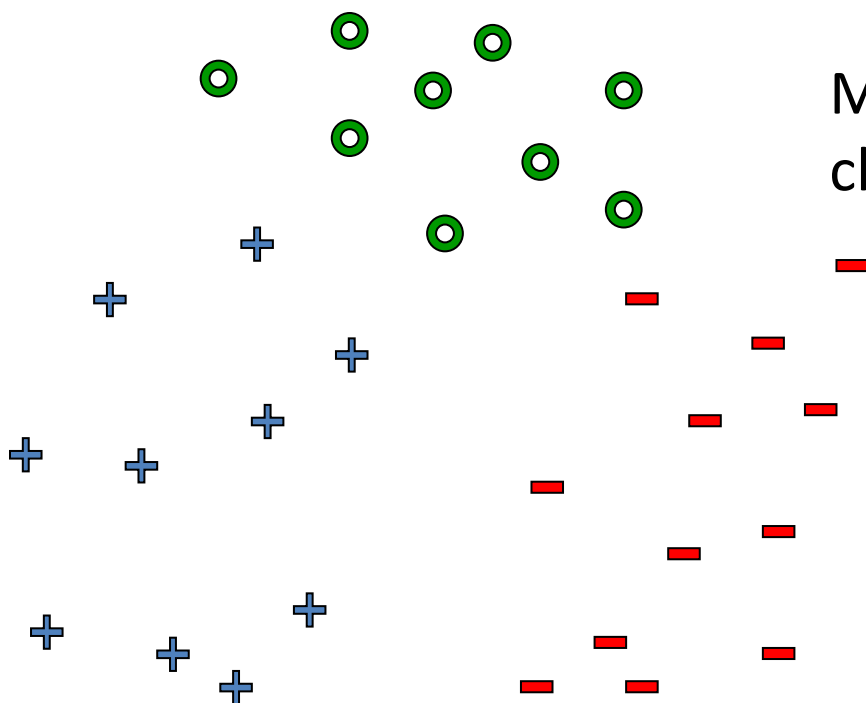
But $\mathbf{w}_k$ s may not be
based on the same scale.
Note: $(a\mathbf{w}).\mathbf{x} + (ab)$ is also
a solution

Learn 1 classifier: Multi-class SVM

Simultaneously learn 3 sets of weights

$$\min_{\{w^{(y)}\}, \{b^{(y)}\}} \sum_y w^{(y)} \cdot w^{(y)}$$

$$w^{(y_j)} \cdot x_j + b^{(y_j)} \geq w^{(y')} \cdot x_j + b^{(y')} + 1, \quad \forall y' \neq y_j, \quad \forall j$$



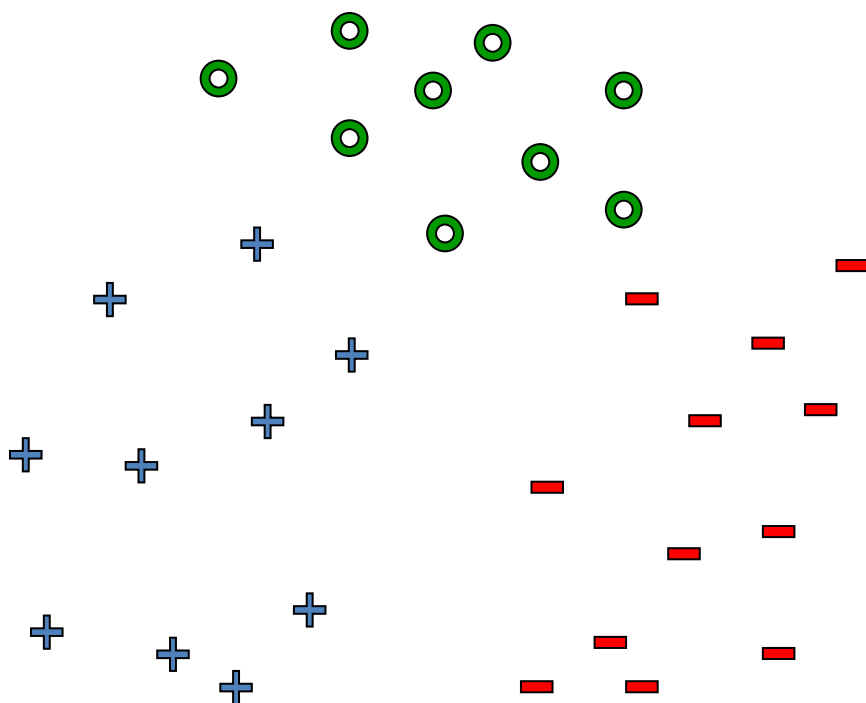
Margin - gap between correct class and nearest other class

$$y = \arg \max_k w^{(k)} \cdot x + b^{(k)}$$

Learn 1 classifier: Multi-class SVM

Simultaneously learn 3 sets of weights

$$\begin{aligned} \text{minimize} \quad & \sum_y \mathbf{w}^{(y)} \cdot \mathbf{w}^{(y)} + C \sum_j \sum_{y \neq y_j} \xi_j^{(y)} \quad \text{over } \{\mathbf{w}^{(y)}\}, \{b^{(y)}\}, \{\xi_j^{(y)}\} \\ & \mathbf{w}^{(y_j)} \cdot \mathbf{x}_j + b^{(y_j)} \geq \mathbf{w}^{(y)} \cdot \mathbf{x}_j + b^{(y)} + 1 - \xi_j^{(y)}, \quad \forall y \neq y_j, \quad \forall j \\ & \xi_j^{(y)} \geq 0, \quad \forall y \neq y_j, \quad \forall j \end{aligned}$$



$$y = \arg \max \mathbf{w}^{(k)} \cdot \mathbf{x} + b^{(k)}$$

Joint optimization: $\mathbf{w}_k$ s
have the same scale.