



Support Vector Machines

Instructor: Jessica Wu -- Harvey Mudd College

The instructor gratefully acknowledges Andrew Ng (Stanford), Eric Eaton (UPenn), David Sontag (NYU), Piyush Rai (Utah) and the many others who made their course materials freely available online.

Robot Image Credit: Viktoriya Sukhanova © 123RF.com

SVM Basics

Learning Goals

- Describe the goal of SVMs
 - Start from perceptron and build graphical intuition
- Setup formal SVM optimization problem

Recall the Perceptron

Assume there is a linear classifier

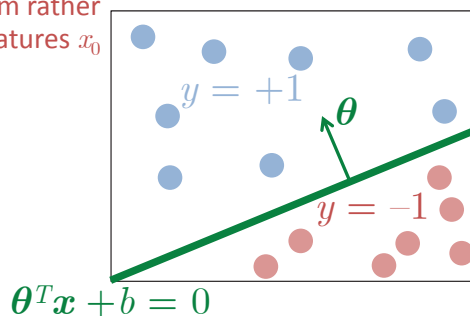
Start with simple learner and analyze what it does

Let $y \in \{-1, +1\}$ use explicit bias term rather than folding into features x_0

$$h_{\theta}(\mathbf{x}) = g(\theta^T \mathbf{x} + b)$$

where

$$g(z) = \text{sgn}(z) \begin{cases} +1, & z \geq 0 \\ -1, & z < 0 \end{cases}$$

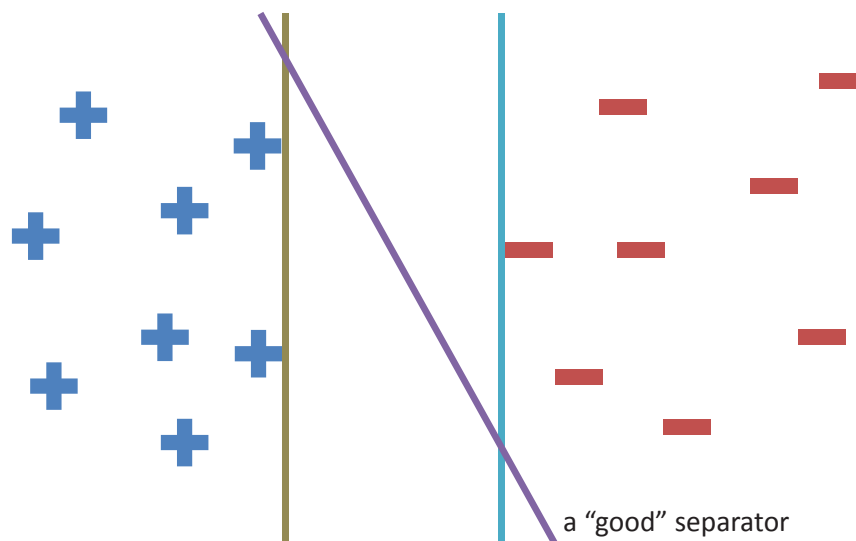


So if $\theta^T \mathbf{x} + b \geq 0$, then $y = +1$

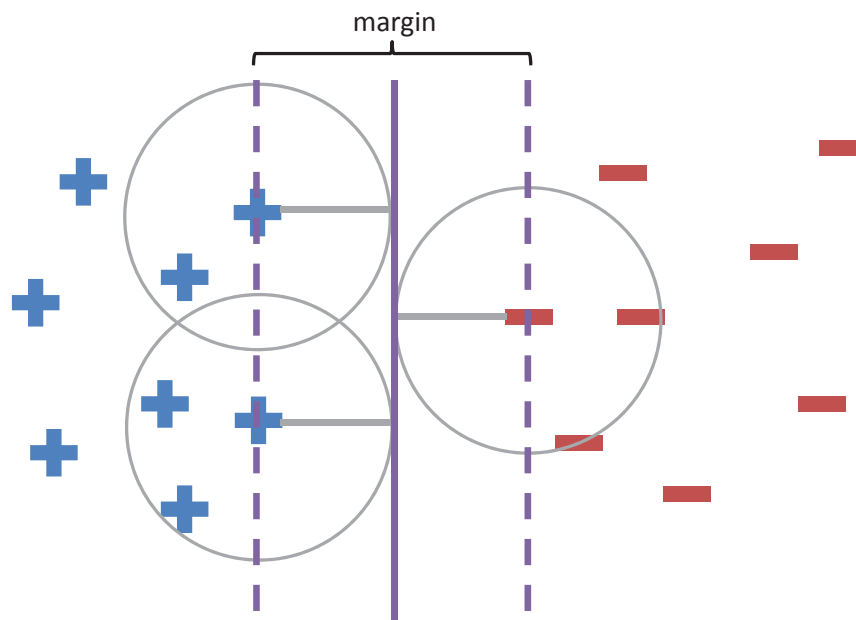
if $\theta^T \mathbf{x} + b < 0$, then $y = -1$

Optimal Linear Separator

Which linear separator is optimal?



Maximizing the Margin



Based on slide by Eric Eaton [originally by Tim Oates]

Summary

- Perceptron finds one of many possible hyperplanes separating data (if one exists)
- Of possible choices, find one with **maximum margin**
⇒ Support Vector Machines (SVMs)
- Why?
 - Good according to intuition, theory, and practice
- Some history
 - SVM became famous when, in handwriting recognition task using images as input, it gave accuracy comparable to neural network with hand-designed features

Review: Margins

For training set $\{(\mathbf{x}^{(i)}, y^{(i)})\}_{i=1}^n$ and parameter vector $\boldsymbol{\theta}$:

- **functional margin** of i^{th} example

$$\hat{\gamma}^{(i)} = y^{(i)} (\boldsymbol{\theta}^T \mathbf{x} + b)$$

- positive if $\boldsymbol{\theta}$ classifies $\mathbf{x}^{(i)}$ correctly
- absolute value = “confidence” in predicted label (or “mis-confidence”)

- **geometric margin** of i^{th} example

$$\gamma^{(i)} = \frac{\hat{\gamma}^{(i)}}{\|\boldsymbol{\theta}\|} = \frac{y^{(i)} (\boldsymbol{\theta}^T \mathbf{x} + b)}{\|\boldsymbol{\theta}\|}$$

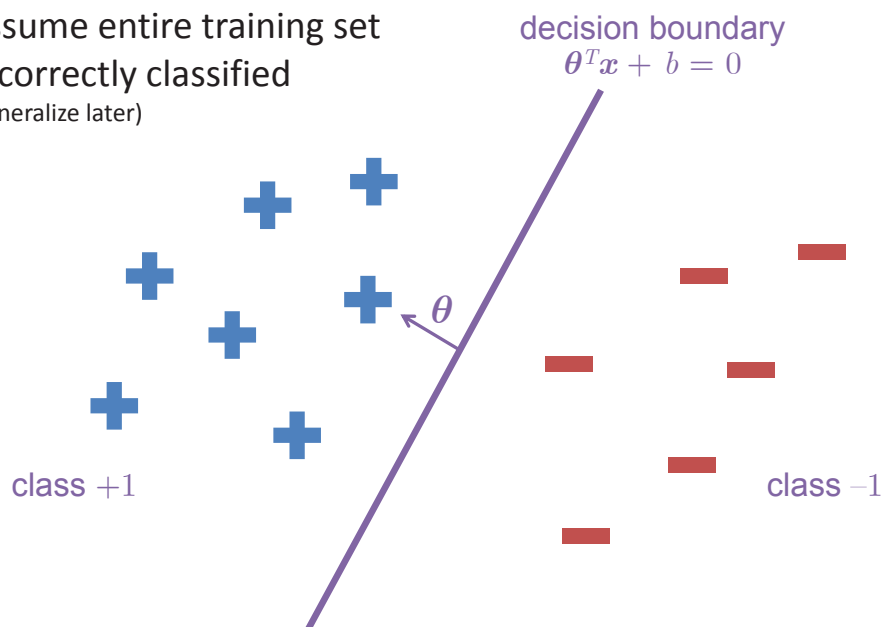
- signed distance of $\mathbf{x}^{(i)}$ to hyperplane (positive if classified correctly)

- **margin** over training set

- minimum functional margin $\hat{\gamma} = \min_i \hat{\gamma}^{(i)}$
- minimum geometric margin $\gamma = \min_i \gamma^{(i)}$

Maximum-Margin Classifier

Assume entire training set
is correctly classified
(generalize later)



enforce $\hat{\gamma} = 1$ (functional margin)

then $\gamma = \frac{\hat{\gamma}}{\|\boldsymbol{\theta}\|} = \frac{1}{\|\boldsymbol{\theta}\|}$ (geometric margin)



SVM Exercise

Consider the following training data:

x_1	x_2	label
1	1	+
2	2	+
2	0	+
0	0	-
1	0	-
0	1	-

- Plot these points. Are the classes $\{+, -\}$ linearly separable?
- Plot the maximum-margin hyperplane by inspection. What is the associated weight vector θ of this hyperplane? Identify the support vectors.
- If you remove one of the support vectors, does the size of the optimal margin decrease, stay the same, or increase?
- *Extra:* Is your answer above true for any data set? Provide a counterexample or give a short proof.

Example adapted from Russell Norvig



(This slide intentionally left blank.)

Generalization

- Large margin $\Rightarrow$ good generalization

- Can we justify this more formally?

reminder: margin $\gamma = 1/\|\theta\|$

large margin

$\Rightarrow$

- Want more formality?

– learning theory ☺



Optimization Problem

$$\min_{\theta, b} \frac{1}{2} \|\theta\|^2$$

$$\begin{aligned} \text{maximum margin} &= \max 1/\|\theta\| \\ &\Rightarrow \min \|\theta\| \text{ (non-convex)} \\ &\Rightarrow \min \frac{1}{2} \|\theta\|^2 \end{aligned}$$

$$\text{s.t. } y^{(i)} (\theta^T x^{(i)} + b) \geq 1 \text{ for } i = 1, \dots, n$$

Note

- convex quadratic objective with linear constraints (n of them)

$\Rightarrow$ Quadratic program (QP)

Polynomial-time algorithms exist for solving QPs

SVMs

Learning Goals

- Solving the SVM optimization problem
 - using Lagrange multipliers (leads to dual problem)
- Allowing misclassified examples
 - using slack variables (leads to soft-margin SVM)
- Describe the SVM loss function
- Allowing non-linear decision boundaries
 - using kernels

Primal and Dual SVMs

Primal

weights
associated
with features

$$\min_{\boldsymbol{\theta}, b} \frac{1}{2} \|\boldsymbol{\theta}\|^2$$
$$\text{s.t. } y^{(i)} \left(\boldsymbol{\theta}^T \mathbf{x}^{(i)} + b \right) \geq 1 \text{ for } i = 1, \dots, n$$

⇓ **MATH! (Lagrange multipliers)**

Dual

weights
associated
with examples

$$\max_{\boldsymbol{\alpha}} \sum_{i=1}^n \alpha_i - \frac{1}{2} \sum_{i=1}^n \sum_{j=1}^n \alpha_i \alpha_j y^{(i)} y^{(j)} \underbrace{\langle \mathbf{x}^{(i)}, \mathbf{x}^{(j)} \rangle}_{\text{dot product } \langle \mathbf{x}, \mathbf{y} \rangle = \mathbf{x}^T \mathbf{y}}$$
$$\text{s.t. } \alpha_i \geq 0 \quad \text{for } i = 1, \dots, n$$
$$\sum_{i=1}^n \alpha_i y^{(i)} = 0$$

Math Details

For math to work,

Karush-Kuhn-Tucker (KKT) conditions must hold:

for optimal θ , b , α ,

$$\alpha_i \left[1 - y^{(i)} (\theta^T \mathbf{x}^{(i)} + b) \right] = 0 \quad \text{complementary slackness}$$

$\Rightarrow$ if $\alpha_i > 0$, then $y^{(i)} (\theta^T \mathbf{x}^{(i)} + b) = 1$

i.e. if weight for example $i > 0$,

then $\mathbf{x}^{(i)}$ lies on a margin boundary ($\mathbf{x}^{(i)}$ is a SV)

$\Rightarrow$ if $y^{(i)} (\theta^T \mathbf{x}^{(i)} + b) > 1$, then $\alpha_i = 0$

i.e. if $\mathbf{x}^{(i)}$ does not lie a margin boundary ($\mathbf{x}^{(i)}$ is not a SV)

then weight for example $i = 0$

Solving SVM Problem

Solve dual in lieu of primal

- Solve for α directly
- How?
 - coordinate descent, sequential minimal optimization (SMO)

Compute primal solution from dual

- to obtain θ ,

$$\theta = \sum_{i=1}^n \alpha_i y^{(i)} \mathbf{x}^{(i)}$$

- can sum over SVs since only these examples have $\alpha_i > 0$
- remember, we have not added $x_0 = 1$

- to obtain b , find a SV $\mathbf{x}^{(i)}$, solve

$$y^{(i)} (\theta^T \mathbf{x}^{(i)} + b) = 1$$

$$b = y^{(i)} - \theta^T \mathbf{x}^{(i)}$$

Sparsity and Generalization

$\alpha_i > 0$ only for SVs

SVs $\ll$ # training examples

$\Rightarrow$ sparseness leads to better generalization

Why?

What is maximum LOOCV error?



(This slide intentionally left blank.)

(This slide intentionally left blank.)

SVMs

Learning Goals

- ✓ Solving the SVM optimization problem
 - using Lagrange multipliers (leads to dual problem)
- Allowing misclassified examples
 - using slack variables (leads to soft-margin SVM)
- Describe the SVM loss function
- Allowing non-linear decision boundaries
 - using kernels

Ideas?

What if Data is Not Linearly Separable?

Idea 1: Add more features

$$\mathbf{x} \rightarrow \phi(\mathbf{x})$$

Problems

- What ϕ ?
- Increasing dimensionality could lead to overfitting

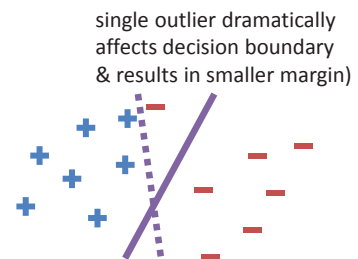
What if Data is Not Linearly Separable?

Idea 2: Allow some training data to be misclassified

(be robust to outliers)

$$\begin{aligned} \min_{\boldsymbol{\theta}, b} \quad & \frac{1}{2} \|\boldsymbol{\theta}\|^2 + C \sum_{i=1}^n \mathbb{I} \left[\left[y^{(i)} \left(\boldsymbol{\theta}^T \mathbf{x}^{(i)} + b \right) \leq 0 \right] \right] \\ \text{s.t.} \quad & y^{(i)} \left(\boldsymbol{\theta}^T \mathbf{x}^{(i)} + b \right) \geq 1 \quad \text{for some subset } i \end{aligned}$$

number of mistakes
trade-off mistakes



Problems

- 0/1 loss not QP and NP-hard
- Does not distinguish near misses from really bad mistakes

Soft-Margin SVM

Idea 2 revised: Allow **slack**

$$\min_{\theta, b, \xi} \frac{1}{2} \|\theta\|^2 + C \sum_{i=1}^n \xi_i$$

Greek xi (zai, sai)

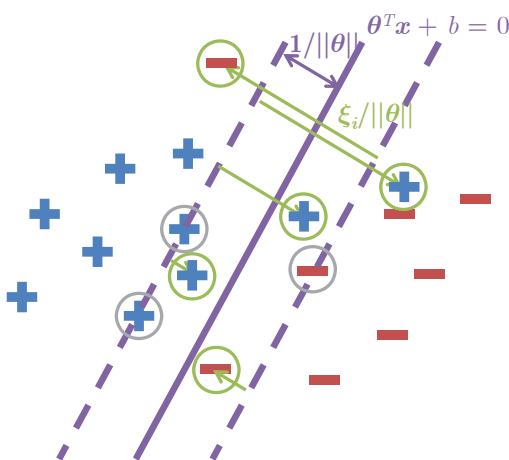
new variable ξ

relax constraints using
slack variables ξ_i

$$\text{s.t. } y^{(i)} (\theta^T x^{(i)} + b) \geq 1 - \xi_i \quad \text{for } i = 1, \dots, n$$

$$\xi_i \geq 0 \quad \text{for } i = 1, \dots, n$$

Soft-Margin SVM



- examples can have (functional) margin < 1
- if $\xi_i > 0$, pay penalty $C\xi_i$
- C is **slack penalty** [$C \geq 0$]
trade-off between...
 - minimizing $\|\theta\|^2$ (maximizing margin)
 - ensuring most examples have functional margin of at least 1
- small C [$\min \frac{1}{2} \|\theta\|^2$]

- margin boundaries still at $\theta^T x + b = \pm 1$
- training example ...
 - within margin region when $0 < \xi_i < 1$
 - misclassified when $\xi_i > 1$
- large C [$\min \sum_i \xi_i$]



Dual Formulation

$$\begin{aligned} \max_{\alpha} \quad & \sum_{i=1}^n \alpha_i - \frac{1}{2} \sum_{i=1}^n \sum_{j=1}^n \alpha_i \alpha_j y^{(i)} y^{(j)} \langle \mathbf{x}^{(i)}, \mathbf{x}^{(j)} \rangle \\ \text{s.t.} \quad & 0 \leq \alpha_i \leq C \quad \text{for } i = 1, \dots, n \\ & \sum_{i=1}^n \alpha_i y^{(i)} = 0 \end{aligned}$$

new constraint: add upper bound $\alpha_i \leq C$

intuition

- w/o slack
 α_i can increase without bound to prevent misclassification
- w/ slack
upper bound of C limits α_i to allow misclassifications

Math Details

Different KKT conditions

$$\begin{aligned} \alpha_i = 0 & \Rightarrow y^{(i)}(\boldsymbol{\theta}^T \mathbf{x}^{(i)} + b) > 1 \\ \alpha_i = C & \Rightarrow y^{(i)}(\boldsymbol{\theta}^T \mathbf{x}^{(i)} + b) < 1 \\ 0 < \alpha_i < C & \Rightarrow y^{(i)}(\boldsymbol{\theta}^T \mathbf{x}^{(i)} + b) = 1 \end{aligned}$$

Types of support vectors

- **hard-margin SVM has only one type of SV**
 - points on margin boundaries
 - **soft-margin SVM has three types of SVs**
 - points on margin boundaries
 $\xi_i = 0, \quad 0 < \alpha_i < C$
 - points within margin region but still on correct side
 $0 < \xi_i < 1$
 - points on wrong side of hyperplane (misclassified)
 $\xi_i \geq 1$
- } margin violators
 $\alpha_i = C$

SVMs

Learning Goals

- ✓ Solving the SVM optimization problem
 - using Lagrange multipliers (leads to dual problem)
- ✓ Allowing misclassified examples
 - using slack variables (leads to soft-margin SVM)
- Describe the SVM loss function
- Allowing non-linear decision boundaries
 - using kernels

Deriving the Loss Function

Recall soft-margin SVM optimization problem

$$\min_{\boldsymbol{\theta}, b, \boldsymbol{\xi}} \frac{1}{2} \|\boldsymbol{\theta}\|^2 + C \sum_{i=1}^n \xi_i$$

$$\text{s.t. } y^{(i)} (\boldsymbol{\theta}^T \mathbf{x}^{(i)} + b) \geq 1 - \xi_i \quad \text{for } i = 1, \dots, n$$
$$\xi_i \geq 0 \quad \text{for } i = 1, \dots, n$$

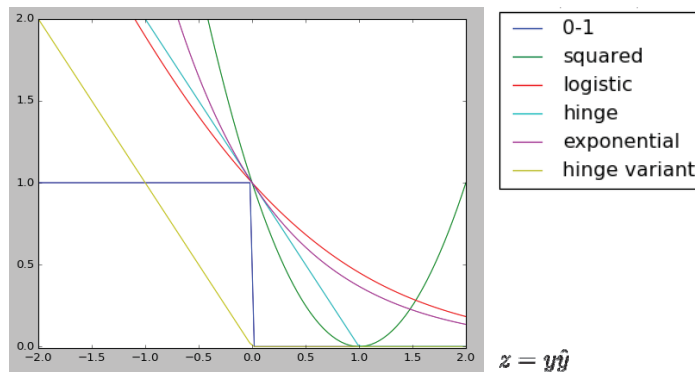
- if $y^{(i)} (\boldsymbol{\theta}^T \mathbf{x}^{(i)} + b) \geq 1$,
then $\xi_i = 0$ (no penalty)
 - if $y^{(i)} (\boldsymbol{\theta}^T \mathbf{x}^{(i)} + b) < 1$,
then $\xi_i = 1 - y^{(i)} (\boldsymbol{\theta}^T \mathbf{x}^{(i)} + b)$ (linear penalty)
- $$\Rightarrow \xi_i = \max(0, 1 - y^{(i)} (\boldsymbol{\theta}^T \mathbf{x}^{(i)} + b))$$

hinge loss (encoded in slack variables)

Loss Functions

$\hat{y}$ = raw prediction

- 0/1 $L(y, \hat{y}) = H(-y\hat{y})$ Heaviside step function
- squared $L(y, \hat{y}) = (1 - y\hat{y})^2$
- logistic $L(y, \hat{y}) = \frac{1}{\log 2} \log(1 + e^{-y\hat{y}})$
- hinge $L(y, \hat{y}) = \max(0, 1 - y\hat{y})$
- exponential $L(y, \hat{y}) = e^{-y\hat{y}}$
- (variant of hinge loss) $L(y, \hat{y}) = \max(0, -y\hat{y})$



Soft-Margin SVM as Regularization

Let

$$L(y, \hat{y}) = \max(0, 1 - y\hat{y})$$

note:

$\hat{y} = \boldsymbol{\theta}^T \mathbf{x} + b$ is "raw" output,
not predicted class

So

$$\min_{\boldsymbol{\theta}, b} \underbrace{\frac{1}{2} \|\boldsymbol{\theta}\|^2}_{\text{regularization}} + C \sum_{i=1}^n \underbrace{L(y^{(i)}, \hat{y}^{(i)})}_{\text{empirical risk using hinge loss}}$$

Standard form of regularization

$$\min_{\boldsymbol{\theta}, b} \frac{\lambda}{2} \|\boldsymbol{\theta}\|^2 + \frac{1}{n} \sum_{i=1}^n L(y^{(i)}, \hat{y}^{(i)})$$

$$\text{So } C \propto \frac{1}{\lambda}$$

low regularization
small λ , large C

high regularization
large λ , small C