

Aprendizado de Máquina

Aula 9

<http://www.ic.uff.br/~bianca/aa/>

Tópicos

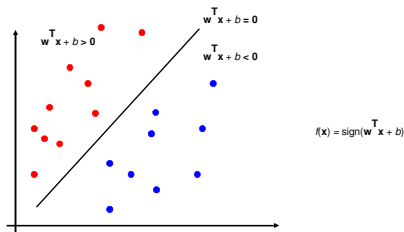
1. Introdução – Cap. 1 (16/03)
2. Classificação Indutiva – Cap. 2 (23/03)
3. Árvores de Decisão – Cap. 3 (30/03)
4. Ensembles - Artigo (13/04)
5. Avaliação Experimental – Cap. 5 (20/04)
6. Aprendizado de Regras – Cap. 10 (27/04)
7. Redes Neurais – Cap. 4 (04/05)
8. Teoria do Aprendizado – Cap. 7 (11/05)
9. **Máquinas de Vetor de Suporte – Artigo (18/05)**
10. Aprendizado Bayesiano – Cap. 6 e novo cap. online (25/05)
11. Aprendizado Baseado em Instâncias – Cap. 8 (01/05)
12. Classificação de Textos – Artigo (08/06)
13. Aprendizado Não-Supervisionado – Artigo (15/06)

Aula 9 - 25/05/2010

2

Revisão de Perceptrons: Separadores Lineares

- Classificação binária pode ser vista como a tarefa de separar as classes no espaço de características:

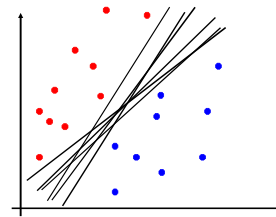


Aula 9 - 25/05/2010

3

Separadores Lineares

- Qual destes separadores é ótimo?

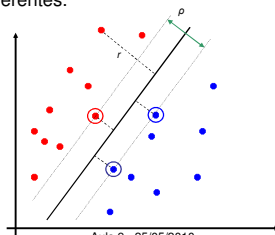


Aula 9 - 25/05/2010

4

Margem de Classificação

- Distância do exemplo x_i ao separador é $r = \frac{w^T x_i + b}{\|w\|}$
- Exemplos mais próximos ao hiperplano são **vetores de suporte**.
- Margem p** do separador é a distância entre vetores de suporte de classes diferentes.

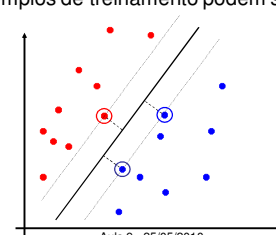


Aula 9 - 25/05/2010

5

Classificação com Máxima Margem

- Maximizar a margem é bom de acordo com a intuição e com a teoria PAC.
- Implica que só os vetores de suporte são importantes; outros exemplos de treinamento podem ser ignorados.



Aula 9 - 25/05/2010

6

SVM Linear Matematicamente

- Seja um conjunto de treinamento $\{(\mathbf{x}_i, y_i)\}_{i=1..n}$, $\mathbf{x}_i \in \mathbb{R}^d$, $y_i \in \{-1, 1\}$ separável por um hiperplano com margem ρ . Então para cada exemplo de treinamento $(\mathbf{x}_i, y_i)$:

$$\begin{aligned} \mathbf{w}^T \mathbf{x}_i + b &\leq -\rho/2 & \text{if } y_i = -1 \\ \mathbf{w}^T \mathbf{x}_i + b &\geq \rho/2 & \text{if } y_i = 1 \end{aligned} \Leftrightarrow y_i(\mathbf{w}^T \mathbf{x}_i + b) \geq \rho/2$$

- Para cada vetor de suporte $\mathbf{x}_s$ a desigualdade acima é uma igualdade. Depois de dividir $\mathbf{w}$ e b por $\rho/2$ na igualdade, temos que a distância entre cada $\mathbf{x}_s$ e o hiperplano é

$$r = \frac{y_s(\mathbf{w}^T \mathbf{x}_s + b)}{\|\mathbf{w}\|} = \frac{1}{\|\mathbf{w}\|}$$

- A margem pode então ser escrita como:

$$\rho = 2r = \frac{2}{\|\mathbf{w}\|}$$

7

SVM Linear Matematicamente

- Então podemos formular o *problema de otimização quadrática*:

$$\begin{aligned} &\text{Encontrar } \mathbf{w} \text{ e } b \text{ tais que} \\ &\rho = \frac{2}{\|\mathbf{w}\|} \text{ seja maximizado} \\ &\text{e para todo } (\mathbf{x}_i, y_i), i=1..n: \quad y_i(\mathbf{w}^T \mathbf{x}_i + b) \geq 1 \end{aligned}$$

que pode ser reformulado como:

$$\begin{aligned} &\text{Encontrar } \mathbf{w} \text{ e } b \text{ tais que} \\ &\Phi(\mathbf{w}) = \|\mathbf{w}\|^2 = \mathbf{w}^T \mathbf{w} \text{ seja minimizado} \\ &\text{e para todo } (\mathbf{x}_i, y_i), i=1..n: \quad y_i(\mathbf{w}^T \mathbf{x}_i + b) \geq 1 \end{aligned}$$

Aula 9 - 25/05/2010

8

Resolvendo o problema de otimização

$$\begin{aligned} &\text{Encontrar } \mathbf{w} \text{ e } b \text{ tais que} \\ &\Phi(\mathbf{w}) = \mathbf{w}^T \mathbf{w} \text{ seja minimizado} \\ &\text{e para todo } (\mathbf{x}_i, y_i), i=1..n: \quad y_i(\mathbf{w}^T \mathbf{x}_i + b) \geq 1 \end{aligned}$$

- Precisamos otimizar uma função quadrática sujeita a restrições lineares.
- Problemas de otimização quadrática são uma classe conhecida de problemas de programação quadrática para a qual existem muitos algoritmos (não-triviais).
- A solução envolve construir um *problema dual* onde um *multiplicador de Lagrange* α_i está associado com cada restrição de desigualdade no problema primal (original):

$$\begin{aligned} &\text{Encontrar } \alpha_1, \dots, \alpha_n \text{ tais que} \\ &\mathbf{Q}(\alpha) = \sum \alpha_i - \frac{1}{2} \sum \sum \alpha_i \alpha_j y_i y_j \mathbf{x}_i^T \mathbf{x}_j \text{ seja maximizado e} \\ &(1) \sum \alpha_i y_i = 0 \\ &(2) \alpha_i \geq 0 \text{ para todo } \alpha_i \end{aligned}$$

Aula 9 - 25/05/2010

9

A solução do problema de otimização

- Dada uma solução $\alpha_1, \dots, \alpha_n$ para o problema dual, a solução do primal é:

$$\mathbf{w} = \sum \alpha_i y_i \mathbf{x}_i \quad b = y_k - \sum \alpha_i y_i \mathbf{x}_i^T \mathbf{x}_k \quad \text{para qualquer } \alpha_k > 0$$

- Cada α_i diferente de zero indica que o $\mathbf{x}_i$ correspondente é um vetor de suporte.
- Então a função de classificação é (note que não precisamos de $\mathbf{w}$ explicitamente):

$$f(\mathbf{x}) = \sum \alpha_i y_i \mathbf{x}_i^T \mathbf{x} + b$$

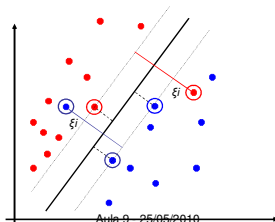
- Note que a função depende de um produto interno entre o vetor de suporte e $\mathbf{x}_i$ – voltaremos a isso depois.
- Note também que resolver o problema de otimização envolve calcular produtos internos $\mathbf{x}_i^T \mathbf{x}_j$ entre todos os exemplos de treinamento.

Aula 9 - 25/05/2010

10

Classificação “Soft margin”

- E se o conjunto de treinamento não for linearmente separável?
- Variáveis “*slack*” ξ_i podem ser adicionadas para tornar permitir o erro de exemplos difíceis ou ruidosos, resultando em uma margem chamada de “*soft*”.



Aula 9 - 25/05/2010

11

Margem “Soft” Matematicamente

- Formulação anterior:

$$\begin{aligned} &\text{Encontrar } \mathbf{w} \text{ e } b \text{ tais que} \\ &\Phi(\mathbf{w}) = \mathbf{w}^T \mathbf{w} \text{ seja minimizado} \\ &\text{e para todo } (\mathbf{x}_i, y_i), i=1..n: \quad y_i(\mathbf{w}^T \mathbf{x}_i + b) \geq 1 \end{aligned}$$

- Nova formulação:

$$\begin{aligned} &\text{Encontrar } \mathbf{w} \text{ e } b \text{ tais que} \\ &\Phi(\mathbf{w}) = \mathbf{w}^T \mathbf{w} + C \sum \xi_i \text{ seja minimizado} \\ &\text{e para todo } (\mathbf{x}_i, y_i), i=1..n: \quad y_i(\mathbf{w}^T \mathbf{x}_i + b) \geq 1 - \xi_i, \quad \xi_i \geq 0 \end{aligned}$$

- Parâmetro C pode ser visto como uma maneira de controlar o sobre-ajuste: controla a importância relativa de maximizar a margem e se ajustar aos dados de treinamento.

Aula 9 - 25/05/2010

12

Margem “Soft” Matematicamente

- Problema dual é idêntico ao separável:

Encontrar $\alpha_1, \dots, \alpha_n$ tais que
 $Q(\alpha) = \sum \alpha_i - \frac{1}{2} \sum \sum \alpha_i \alpha_j y_i y_j x_i^T x_j$ seja maximizado e
 (1) $\sum \alpha_i y_i = 0$
 (2) $0 \leq \alpha_i \leq C$ para todo α_i

- De novo, x_i com α_i diferente de zero serão vetores de suporte.

- Solução do problema dual é:

$$w = \sum \alpha_i y_i x_i \\ b = y_k(1 - \xi_k) - \sum \alpha_i y_i x_i^T x_k \text{ para qualquer } k \text{ tal que } \alpha_k > 0$$

- E função de classificação é:

$$f(x) = \sum \alpha_i y_i x_i^T x + b$$

Aula 9 - 25/05/2010

13

Justificativa Teórica para Máxima Margem

- Vapnik provou o seguinte:

A classe de separadores ótimos lineares tem dimensão VC limitada por:

$$h \leq \min \left\{ \left\lceil \frac{D^2}{\rho^2} \right\rceil, m_0 \right\} + 1$$

onde ρ é a margem, D é o diâmetro da menor esfera que cobre todos os exemplos de treinamento, e m_0 é a dimensionalidade (número de atributos).

- Intuitivamente, isto implica que não importando a dimensionalidade m_0 podemos minimizar a dimensão VC maximizando a margem ρ .
- Logo a complexidade do classificador fica pequena independentemente da dimensionalidade.

Aula 9 - 25/05/2010

14

SVM Linear: Resumo

- O classificador é um *hiperplano separador*.
- Exemplos de treinamento mais “importantes” são os vetores de suporte; eles definem o hiperplano.
- Algoritmos de otimização quadrática podem identificar que exemplos de treinamento x_i são vetores de suporte com multiplicadores de Lagrange α_i diferentes de zero.
- Tanto na formulação dual do problema quanto na solução, exemplos de treinamento aparecem apenas dentro de produtos internos:

Encontrar $\alpha_1, \dots, \alpha_n$ tais que
 $Q(\alpha) = \sum \alpha_i - \frac{1}{2} \sum \sum \alpha_i \alpha_j y_i y_j x_i^T x_j$ seja maximizado e
 (1) $\sum \alpha_i y_i = 0$
 (2) $0 \leq \alpha_i \leq C$ para todo α_i

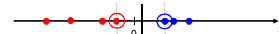
$$f(x) = \sum \alpha_i y_i x_i^T x + b$$

Aula 9 - 25/05/2010

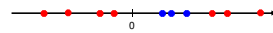
15

SVM Não-Linear

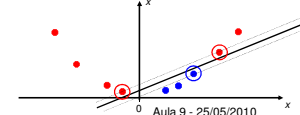
- Dados que são linearmente separáveis com algum ruído funcionam bem com margem “soft”:



- Mas o que fazer se este não é o caso?



- Que tal mapear os dados para um espaço de dimensão maior?

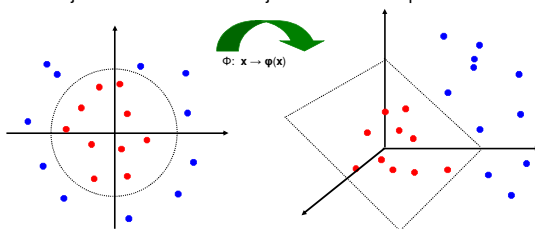


Aula 9 - 25/05/2010

16

SVM Não Linear

- Ideia geral: o espaço de atributos original pode ser mapeado em um espaço de atributos de dimensão maior onde o conjunto de treinamento seja linearmente separável:



Aula 9 - 25/05/2010

17

O “Truque” do Kernel

- O classificador linear depende do produto interno entre exemplos $K(x_i, x_j) = x_i^T x_j$.
- Se cada ponto for mapeado a um espaço de alta dimensão através de uma transformação $\Phi: x \rightarrow \phi(x)$, o produto interno fica:
 $K(x_i, x_j) = \phi(x_i)^T \phi(x_j)$
- Uma função de kernel é uma função que é equivalente a um produto interno num espaço de maior dimensionalidade.
- Exemplo:
 Vetores de 2 dimensões $x = [x_1 \ x_2]$; let $K(x_i, x_j) = (1 + x_i^T x_j)^2$.
 Precisamos mostrar que $K(x_i, x_j) = \phi(x_i)^T \phi(x_j)$:
 $K(x_i, x_j) = (1 + x_i^T x_j)^2 = 1 + x_{i1}^2 x_{j1}^2 + 2 x_{i1} x_{j1} x_{i2} x_{j2} + x_{i2}^2 x_{j2}^2 + 2 x_{i1} x_{j1} + 2 x_{i2} x_{j2}$
 $= [1 \ x_{i1}^2 \ \sqrt{2} x_{i1} x_{i2} \ x_{i2}^2 \ \sqrt{2} x_{i1} \ \sqrt{2} x_{i2}]^T [1 \ x_{j1}^2 \ \sqrt{2} x_{j1} x_{j2} \ x_{j2}^2 \ \sqrt{2} x_{j1} \ \sqrt{2} x_{j2}]$
 $= \phi(x_i)^T \phi(x_j)$, onde $\phi(x) = [1 \ x_1^2 \ \sqrt{2} x_1 x_2 \ x_2^2 \ \sqrt{2} x_1 \ \sqrt{2} x_2]$
- Não precisamos calcular $\phi(x)$ explicitamente.

Aula 9 - 25/05/2010

18

Que funções são *Kernels*?

- Para algumas funções $K(\mathbf{x}_i, \mathbf{x}_j)$ checar que $K(\mathbf{x}_i, \mathbf{x}_j) = \phi(\mathbf{x}_i)^T \phi(\mathbf{x}_j)$ pode ser difícil.
- Teorema de Mercer:
Toda função simétrica definida semi-positiva é um kernel

Aula 9 - 25/05/2010

19

Exemplos de Funções *Kernel*

- Linear: $K(\mathbf{x}_i, \mathbf{x}_j) = \mathbf{x}_i^T \mathbf{x}_j$
– $\Phi: \mathbf{x} \rightarrow \phi(\mathbf{x})$, onde $\phi(\mathbf{x}) = \mathbf{x}$
- Polinomial de expoente p : $K(\mathbf{x}_i, \mathbf{x}_j) = (1 + \mathbf{x}_i^T \mathbf{x}_j)^p$
– $\Phi: \mathbf{x} \rightarrow \phi(\mathbf{x})$, onde $\phi(\mathbf{x})$ tem $\binom{d+p}{p}$ dimensões
- Gaussiana (RBF): $K(\mathbf{x}_i, \mathbf{x}_j) = e^{-\frac{\|\mathbf{x}_i - \mathbf{x}_j\|^2}{2\sigma^2}}$
– $\Phi: \mathbf{x} \rightarrow \phi(\mathbf{x})$, onde $\phi(\mathbf{x})$ tem dimensão infinita: cada ponto é mapeado para uma função gaussiana.
- Separadores lineares correspondem a separadores não lineares no espaço de dimensão maior

Aula 9 - 25/05/2010

20

SVMs Não Lineares Matematicamente

- Formulação do problema dual:

Encontrar $\alpha_1, \dots, \alpha_n$ tais que
 $Q(\alpha) = \sum \alpha_i - \frac{1}{2} \sum \sum \alpha_i \alpha_j y_i y_j K(\mathbf{x}_i, \mathbf{x}_j)$ seja maximizado e
 (1) $\sum \alpha_i y_i = 0$
 (2) $\alpha_i \geq 0$ para todo α_i

- A solução é:

$$f(\mathbf{x}) = \sum \alpha_i y_i K(\mathbf{x}_i, \mathbf{x}) + b$$

- Técnicas de otimização para encontrar os α_i 's são as mesmas!

Aula 9 - 25/05/2010

21

SVM applications

- SVMs were originally proposed by Boser, Guyon and Vapnik in 1992 and gained increasing popularity in late 1990s.
- SVMs are currently among the best performers for a number of classification tasks ranging from text to genomic data.
- SVMs can be applied to complex data types beyond feature vectors (e.g. graphs, sequences, relational data) by designing kernel functions for such data.
- SVM techniques have been extended to a number of tasks such as regression [Vapnik *et al.* '97], principal component analysis [Schölkopf *et al.* '99], etc.
- Most popular optimization algorithms for SVMs use *decomposition* to hill-climb over a subset of α_i 's at a time, e.g. SMO [Platt '99] and [Joachims '99]
- Tuning SVMs remains a black art: selecting a specific kernel and parameters is usually done in a try-and-see manner.

Aula 9 - 25/05/2010

22