

STAT 535

11/16/23

Lecture 15

SVM

optimization
C - SVM

Poll · future lectures

HW5 due

HW6 ^{TB} out

Q3 $\geq$ TxG

Lecture V: Support Vector Machines

Marina Meilă
`mmp@stat.washington.edu`

Department of Statistics
University of Washington

November, 2023

Linear SVM's

The margin and the expected classification error

Maximum Margin Linear classifiers

Linear classifiers for non-linearly separable data

✓
optimization



Non linear SVM

The “kernel trick”

Kernels

Prediction with SVM

Extensions

L_1 SVM

Multi-class and One class SVM

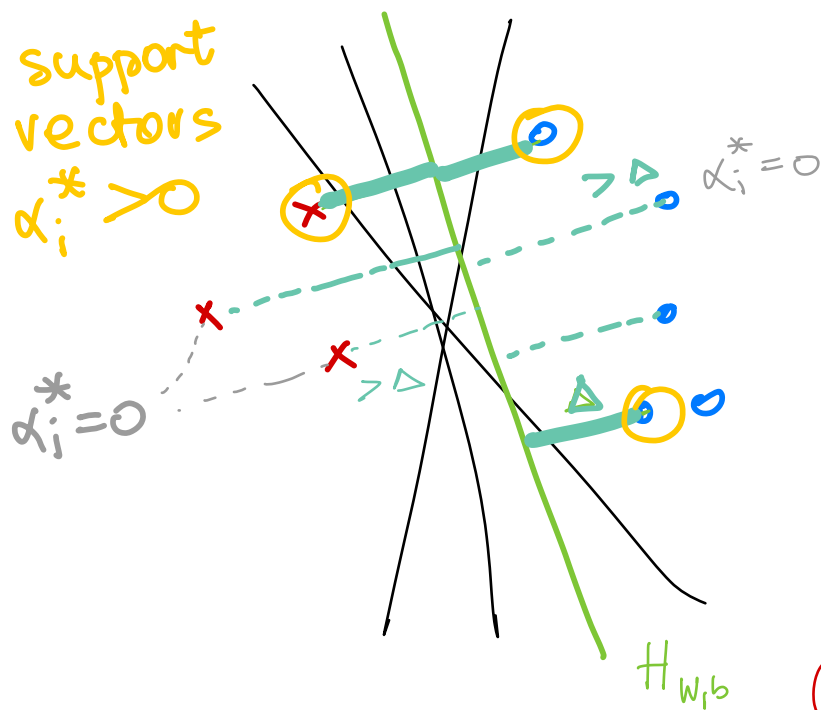
SV Regression

Reading HTF Ch.: Ch. 12.1–3, Murphy Ch.: Ch 14 (14.1,14.2–14.2.4 kernels, 14.4 and equations (14.28,14.29) kernel trick, 14.5.1.–3 Support Vector Machines), Bach Ch.: 7.1–7.4, 7.7

Additional Reading: C. Burges - “A tutorial on SVM for pattern recognition”

These notes: Appendices (convex optimization) are optional.

support
vectors
 $\alpha_i^* > 0$



Assume data
linearly separable

margin
 $\Delta = \min_{w,b} \min_i \text{dist}(x^i, H_{w,b})$

$f(x) = w^T x + b$ linear classifier

$H_{w,b} = \{x, f(x) = 0\}$

SVM pb $\max$ margin
want $\arg \max_{w,b} \Delta_{w,b}$ s.t.

(P1)

$y^i (w^T x^i + b) > 0$

classify correctly

s.t. $y^i (w^T x^i + b) \geq 1$

Remark
3. $\min_i y^i (w^T x^i + b) = 1 \Rightarrow$

Remark 1.

(P2)

Remark 2.

$d(x, H_{w,b}) = \frac{|w^T x + b|}{\|w\|}$

s.t. $y^i (w^T x^i + b) \geq 1$

(P3)

$\max_{w,b}$

$\min_i \frac{|w^T x^i + b|}{\|w\|}$

Δ

$$(P4) \quad \max_{w, b} \quad \frac{1}{\|w\|} \quad \text{s.t.} \quad y^i (w^T x^i + b) \geq 1$$

$$(P5) \quad \min_{w, b} \quad \|w\|^2 \quad \text{s.t.} \quad \text{---} \text{---} \text{---} \quad \text{CONVEX}$$

$$(P6) \quad \min \quad \|w\|^2 \quad \text{s.t.} \quad \text{---} \text{---} \text{---} \quad \text{QUADRATIC}$$

easier to analyze!

III

$$(P) \quad \min_{\underline{w}, \underline{b}} \quad \frac{1}{2} \|w\|^2 \quad \text{s.t.} \quad y^i (w^T x^i + b) \geq 1 \quad \text{for all } i=1:n$$

Primal
[opt. pb.]

Convex opt. analysis

$$f(x) = \underline{w}^T x + \underline{b}$$

$$\hat{y} = \text{sgn } f(x)$$

Optimization with Lagrange multipliers

² The **Lagrangian** of (13) is

$$L(w, b, \alpha) = \frac{1}{2} \|w\|^2 - \sum_i \alpha_i [y^i (w^T x^i + b) - 1]. \quad (16)$$

[KKT conditions]

At the optimum of (13)

$$w = \sum_i \alpha_i y^i x^i \quad \text{with } \alpha_i \geq 0 \quad (17)$$

and $b = y^i - w^T x^i$ for any i with $\alpha_i > 0$.

- ▶ **Support vector** is a data point x^i such that $\alpha_i > 0$.
- ▶ According to (17), the final decision boundary is determined by the support vectors (i.e. does not depend explicitly on any data point that is not a support vector).

²The derivations of these results are in the Appendix

(P)

Primal
[opt. pb.]

Convex opt. analysis
↓

Lagrangian fctn

$$L: L = \frac{1}{2} \| \underline{W} \|^2 +$$

$$\text{Optimum: } (\underbrace{W^*}_{\min}, \underbrace{b^*}_{\max}, \underbrace{\alpha^*}_{\max})$$

KKT conditions

$$\text{at opt} \quad \begin{cases} \frac{\partial L}{\partial W} = 0 \\ \frac{\partial L}{\partial b} = 0 \end{cases}$$

objective

$$\min_{\underline{W}, \underline{b}} \frac{1}{2} \| W \|^2 \quad \text{s.t.}$$

constraints

$$y^i (W^T x^i + b) \geq 1$$

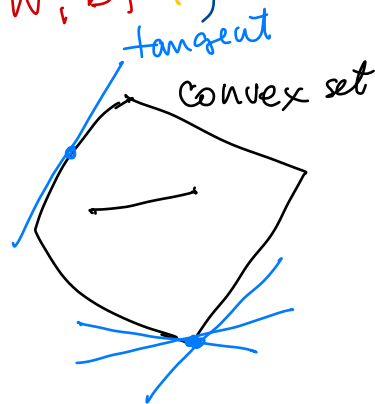
for all $i=1:w \leftarrow \alpha_i$

$$f(x) = \underline{W}^T x + \underline{b}$$

$$\hat{y} = \text{sgn } f(x)$$

$$1 - y^i (W^T x^i + b) \leq 0$$

$$\sum_{i=1}^w \underbrace{\alpha_i}_{\text{dual variable}} [1 - y^i (\underline{W}^T x^i + \underline{b})] = L(W, b, \alpha)$$



at opt $\frac{\partial L}{\partial w} = 0$

$$w + \sum_i (-y^i \alpha_i x^i)$$

$$\Rightarrow w^* = \sum_{i=1}^n \alpha_i^* y^i x^i$$

$\alpha_i^* \geq 0$ (no proof)

$$\frac{\partial L}{\partial b} = - \sum \alpha_i^* y^i = 0 \quad \text{constraint on } \alpha_i^* \text{'s}$$

$$L(\alpha) = \sum_{i=1}^n \alpha_i - \frac{1}{2} \sum_{i,j} \alpha_i \alpha_j y^i y^j (x^i)^T x^j + 0$$

$$= 1^T \alpha - \frac{1}{2} \alpha^T \bar{G} \alpha = g(\alpha)$$

dual pb

(D)

$$\max_{\alpha \in \mathbb{R}^n} g(\alpha) \quad \text{s.t.} \quad \alpha_i \geq 0$$

$$\sum_{i=1}^n \alpha_i y^i = 0$$

solution α^*

$$L = \frac{1}{2} \|w\|^2 + \sum_{i=1}^n \alpha_i [1 - y^i (w^T x^i + b)]$$

linear

$w^* = \text{linear combination of } x^i$!!

$$\|w\|^2 = w^T w$$

$$-w^T \sum \alpha_i y^i x^i = -w^T w$$

$$\alpha = \begin{bmatrix} \alpha_1 \\ \vdots \\ \alpha_n \end{bmatrix}$$

concave, quadratic

$$G = [(x^i)^T x^j]_{i,j=1:n}$$

Gram matrix

$$\bar{G} = [y^i y^j (x^i)^T x^j]_{i,j}$$

$$1 = \begin{bmatrix} 1 \\ \vdots \\ 1 \end{bmatrix} \in \mathbb{R}^n$$

$$- [y^i \vdots] G [y^j \vdots] \leq 0$$

(P) $\min_{w, b} \frac{1}{2} \|w\|^2$ s.t. $y^i (w^T x^i + b) \geq 1$

objective constraints

for all $i = 1:n$ $\leftarrow \alpha_i$

Primal
[opt. pb.]

(D) $\max_{\alpha \in \mathbb{R}^n} g(\alpha)$ s.t. $\alpha_i \geq 0$

$\sum_{i=1}^n \alpha_i y^i = 0$ ||

linear

solution α^*

$w^* = \sum_{i=1}^n \alpha_i^* y^i x^i$

$\alpha_i^* > 0$

b^*

for S.V: $y^i (w^T x^i + b) = 1$

$\rightarrow$ solve for $\underline{b^*}$

$w^T x^i + b = y^i$

STATISTICS, finally

$\sim$ depends on subset of (x^i, y^i)

directly $\rightarrow$ complexity

depends on data

params = # S.V.

$\Rightarrow$ low variance

Theory

$\alpha_i^* > 0$ tight

$\alpha_i^* = 0$ slack constraints

support vectors

(y^i, x^i)

= x^i 's on the margin

SVM Training

1. $\{(x^i, y^i)\} \rightarrow$ compute $G, \bar{G}$
2. Solve (D) dual pb $\Rightarrow \alpha_i^*, i=1:n$
3. $W = \sum_{i=1}^n \alpha_i^* y^i x^i$
 $\alpha_i^* > 0$
4. for i support vector: solve for b
 $\alpha_i^* > 0$

Prediction

new x : $\hat{y} = \text{sgn } f(x)$

$$f(x) = W^T x + b = \sum_{i, \alpha_i^* > 0} \alpha_i^* y^i \underline{\underline{x^i}}^T x + b$$

- So far: linear f , data linearly separable ✓
- Next: —||—, data NOT —||— —||—
- —||— : NonLinear f

Computation efficient
(D) $\alpha \in \mathbb{R}^n$ for $n \ll d$
 n variables
($n+1$) constraints

(P) $d+1$ variables
 $W \in \mathbb{R}^d$ n constraints
 $b \in \mathbb{R}$

efficient for d small
 n large

Dual SVM optimization problem

- ▶ Any convex optimization problem has a **dual** problem. In SVM, it is both illuminating and practical to solve the dual problem.
- ▶ The dual to problem (13) is

$$\max_{\alpha_{1:n}} \sum_i \alpha_i - \frac{1}{2} \sum_i \alpha_i \alpha_j y^i y^j x^i x^j \text{ s.t } \alpha_i \geq 0 \text{ for all } i \text{ and } \sum_i \alpha_i y^i = 0. \quad (18)$$

- ▶ This is a **quadratic** problem with n variables on a convex domain.
- ▶ Dual problem in matrix form

- ▶ Denote $\alpha = [\alpha_i]_{i=1:n}$, $y = [y^i]_{i=1:n}$, $G_{ij} = x^i x^j$, $\bar{G}_{ij} = y^i y^j x^i x^j$,
 $G = [G_{ij}] \in \mathbb{R}^{n \times n}$, $\bar{G} = [\bar{G}_{ij}] \in \mathbb{R}^{n \times n}$.

$$\max_{\alpha \in \mathbb{R}^n} 1^T \alpha - \frac{1}{2} \alpha^T \bar{G} \alpha \text{ s.t } \alpha \succeq 0 \text{ and } y^T \alpha = 0. \quad (19)$$

- ▶ $g(\alpha) = 1^T \alpha - \frac{1}{2} \alpha^T \bar{G} \alpha$ is the **dual objective function**
- ▶ G is called the **Gram matrix** of the data. Note that $\bar{G} = \text{diag}\{y^{1:n}\}^T G \text{diag}\{y^{1:n}\}$.
- ▶ At the dual optimum
 - ▶ $\alpha_i > 0$ for constraints that are satisfied with equality, i.e. **tight**
 - ▶ $\alpha_i = 0$ for the **slack** constraints

Non-linearly separable problems and their duals

The C-SVM

$C \uparrow$ less slack
 $\Rightarrow \hat{\text{L}} \Rightarrow \text{bias} \searrow$
 $\text{var} \nearrow$

minimize w, b, ξ

s.t.

$$\frac{1}{2} \|w\|^2 + C \sum_i \xi_i \quad (20)$$

$C > 0$ regularization
as little slack as possible

$$y^i (w^T x^i + b) \geq 1 - \xi_i$$

slack variables

$$\xi_i \geq 0$$

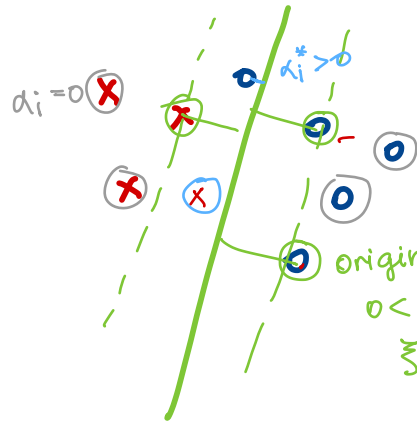
In the above, ξ_i are the slack variables. Dual³:

$$\begin{aligned} \text{maximize}_{\alpha} \quad & \sum_i \alpha_i - \frac{1}{2} \sum_i \alpha_i \alpha_j y^i y^j x^i x^j \\ \text{s.t.} \quad & C \geq \alpha_i \geq 0 \text{ for all } i \\ & \sum_i \alpha_i y^i = 0 \end{aligned} \quad (21)$$

$\Rightarrow$ two types of SV

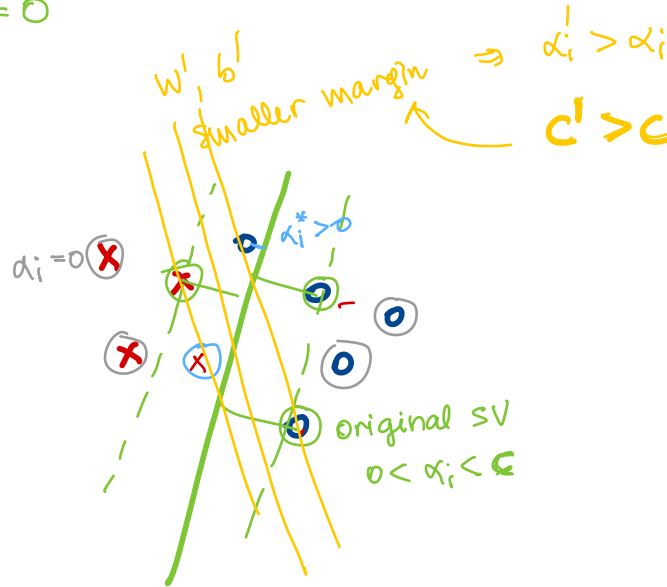
- ▶ $\alpha_i < C$ data point x^i is "on the margin" $\Leftrightarrow y^i (w^T x^i + b) = 1$ (original SV)
- ▶ $\alpha_i = C$ data point x^i cannot be classified with margin 1 (**margin error**)
 $\Leftrightarrow y^i (w^T x^i + b) < 1$

³Lagrangian $L(w, b, \xi, \alpha, \mu) = \frac{1}{2} \|w\|^2 + C \sum_i \xi_i - \sum_i \alpha_i [y^i (w^T x^i + b) - 1 + \xi_i] - \sum_i \mu_i \xi_i$ with $\alpha_i \geq 0, \xi_i \geq 0, \mu_i \geq 0$



margin errors $\xi_i > 0 \iff \alpha_i^* = C$

In general:
 $\alpha_i^* \uparrow \iff$ point i harder
 to classify



$C' > C \Rightarrow$ less robust
 = higher var
 $\Rightarrow$ less bias

Lectures 1.5-2

Boosting SVM

—||— [kernel machines] $\rightarrow$ R Fourier F. $\rightarrow$ Double $\rightarrow$ Descent

< 1 $\hookrightarrow$ Neuro-tangent* + ...

1.5 — Clust K-means + EM

1.5 — Spectral clustering

< 1 Model selection AIC, BIC, struct Risk Min $\leftarrow$ VC dim

Kernel tricks

RKHS $\mathcal{H}$

4.5 lectures

The ν -SVM

$$\text{minimize}_{w,b,\xi,\rho} \quad \frac{1}{2} \|w\|^2 - \nu\rho + \frac{1}{n} \sum_i \xi_i \quad (22)$$

$$\text{s.t.} \quad y^i(w^T x^i + b) \geq \rho - \xi_i \quad (23)$$

$$\xi_i \geq 0 \quad (24)$$

$$\rho \geq 0 \quad (25)$$

where $\nu \in [0, 1]$ is a parameter.

Dual⁴:

$$\text{maximize}_{\alpha} \quad -\frac{1}{2} \sum_i \alpha_i \alpha_j y^i y^j x^i{}^T x^j \quad (26)$$

$$\text{s.t.} \quad \frac{1}{n} \geq \alpha_i \geq 0 \text{ for all } i \quad (27)$$

$$\sum_i \alpha_i y^i = 0 \quad (28)$$

$$\sum_i \alpha_i \geq \nu \quad (29)$$

Properties If $\rho > 0$ then:

- ▶ ν is an upper bound on $\# \text{margin errors} / n$ (if $\sum_i \alpha_i = \nu$)
- ▶ ν is a lower bound on $\#(\text{original support vectors} + \text{margin errors}) / n$
- ▶ ν -SVM leads to the same w, b as C-SVM with $C = 1/\nu$

⁴Lagrangian $L(w, b, \xi, \rho, \alpha, \mu, \delta) = \frac{1}{2} \|w\|^2 - \nu\rho + \frac{1}{n} \sum_i \xi_i - \sum_i \alpha_i [y^i(w^T x^i + b) - \rho + \xi_i] - \sum_i \mu_i \xi_i - \delta\rho$
with $\alpha_i \geq 0, \delta \geq 0, \mu_i \geq 0$