

COMS 4771 Lecture 9

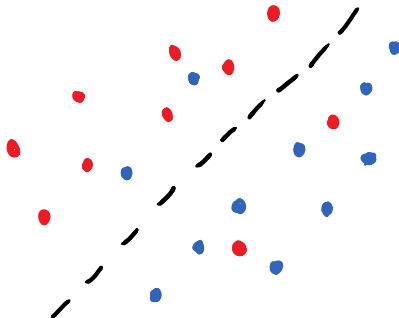
1. Soft-margin SVMs and surrogate losses
2. Convex optimization

SOFT-MARGIN SVMs AND SURROGATE LOSSES

NON-SEPARABLE CASES

Non-separable cases:

No linear classifier has zero training error on S .



But if non-separability is only due to a handful of points,
can we still find a good linear classifier?

SOFT-MARGIN SVMs (CORTES AND VAPNIK, 1995)

When $S = ((\mathbf{x}^{(1)}, y^{(1)}), \dots, (\mathbf{x}^{(n)}, y^{(n)}))$ is not linearly separable, the (primal) SVM optimization problem

$$\begin{array}{ll} \min_{\mathbf{w} \in \mathbb{R}^d, \theta \in \mathbb{R}} & \frac{1}{2} \|\mathbf{w}\|_2^2 \\ \text{s.t.} & y^{(i)} \left(\langle \mathbf{w}, \mathbf{x}^{(i)} \rangle - \theta \right) \geq 1 \quad \text{for } i = 1, 2, \dots, n \end{array}$$

has no solution.

SOFT-MARGIN SVMs (CORTES AND VAPNIK, 1995)

When $S = ((\mathbf{x}^{(1)}, y^{(1)}), \dots, (\mathbf{x}^{(n)}, y^{(n)}))$ is not linearly separable, the (primal) SVM optimization problem

$$\begin{aligned} \min_{\mathbf{w} \in \mathbb{R}^d, \theta \in \mathbb{R}} \quad & \frac{1}{2} \|\mathbf{w}\|_2^2 \\ \text{s.t.} \quad & y^{(i)} \left(\langle \mathbf{w}, \mathbf{x}^{(i)} \rangle - \theta \right) \geq 1 \quad \text{for } i = 1, 2, \dots, n \end{aligned}$$

has no solution.

Introduce **slack variables** $\xi_1, \xi_2, \dots, \xi_n \geq 0$, and a trade-off parameter $C > 0$:

$$\begin{aligned} \min_{\mathbf{w} \in \mathbb{R}^d, \theta \in \mathbb{R}, \xi \in \mathbb{R}^n} \quad & \frac{1}{2} \|\mathbf{w}\|_2^2 + C \sum_{i=1}^n \xi_i \\ \text{s.t.} \quad & y^{(i)} \left(\langle \mathbf{w}, \mathbf{x}^{(i)} \rangle - \theta \right) \geq 1 - \xi_i \quad \text{for } i = 1, 2, \dots, n \\ & \xi_i \geq 0 \quad \text{for } i = 1, 2, \dots, n \end{aligned}$$

which is **always feasible**—“soft margin” SVM.

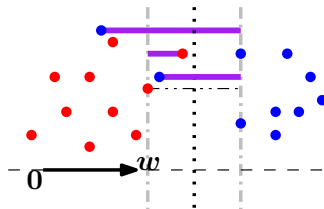
Winner of 2008 ACM Paris Kanellakis Award:

For “their revolutionary development of a highly effective algorithm known as Support Vector Machines (SVM), a set of related supervised learning methods used for data classification and regression”, which is “one of the most frequently used algorithms in machine learning, and is used in medical diagnosis, weather forecasting, and intrusion detection among many other practical applications”.

Other winners include: public key cryptography, Lempel-Ziv compression, Splay Trees, interior point method for linear programming, ... and AdaBoost (discussed later in the course).

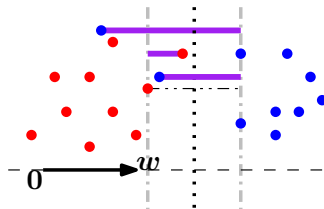
SLACK INTERPRETATION

$$\begin{aligned} \min_{\mathbf{w} \in \mathbb{R}^d, \theta \in \mathbb{R}, \boldsymbol{\xi} \in \mathbb{R}^n} \quad & \frac{1}{2} \|\mathbf{w}\|_2^2 + C \sum_{i=1}^n \xi_i \\ \text{s.t.} \quad & y^{(i)} (\langle \mathbf{w}, \mathbf{x}^{(i)} \rangle - \theta) \geq 1 - \xi_i \quad \text{for } i = 1, 2, \dots, n \\ & \xi_i \geq 0 \quad \text{for } i = 1, 2, \dots, n \end{aligned}$$



SLACK INTERPRETATION

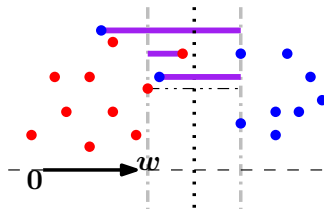
$$\begin{aligned} \min_{\mathbf{w} \in \mathbb{R}^d, \theta \in \mathbb{R}, \xi \in \mathbb{R}^n} \quad & \frac{1}{2} \|\mathbf{w}\|_2^2 + C \sum_{i=1}^n \xi_i \\ \text{s.t.} \quad & y^{(i)} (\langle \mathbf{w}, \mathbf{x}^{(i)} \rangle - \theta) \geq 1 - \xi_i \quad \text{for } i = 1, 2, \dots, n \\ & \xi_i \geq 0 \quad \text{for } i = 1, 2, \dots, n \end{aligned}$$



- For a given $(\mathbf{w}, \theta)$, $\xi_i / \|\mathbf{w}\|_2$ measures distance that $\mathbf{x}^{(i)}$ must be moved so that $y^{(i)} \langle \mathbf{w}, \mathbf{x}^{(i)} \rangle - \theta \geq 1$.

SLACK INTERPRETATION

$$\begin{aligned} \min_{\mathbf{w} \in \mathbb{R}^d, \theta \in \mathbb{R}, \boldsymbol{\xi} \in \mathbb{R}^n} \quad & \frac{1}{2} \|\mathbf{w}\|_2^2 + C \sum_{i=1}^n \xi_i \\ \text{s.t.} \quad & y^{(i)} (\langle \mathbf{w}, \mathbf{x}^{(i)} \rangle - \theta) \geq 1 - \xi_i \quad \text{for } i = 1, 2, \dots, n \\ & \xi_i \geq 0 \quad \text{for } i = 1, 2, \dots, n \end{aligned}$$



- ▶ For a given $(\mathbf{w}, \theta)$, $\xi_i / \|\mathbf{w}\|_2$ measures distance that $\mathbf{x}^{(i)}$ must be moved so that $y^{(i)} \langle \mathbf{w}, \mathbf{x}^{(i)} - \theta \rangle \geq 1$.
- ▶ C controls trade-off between slack penalties and size of margin (which is $1 / \|\mathbf{w}\|_2$).

A DIFFERENT INTERPRETATION OF SLACK

Constraints with non-negative slack variables: (using $\lambda = 1/(nC)$)

$$\begin{aligned} \min_{\mathbf{w} \in \mathbb{R}^d, \theta \in \mathbb{R}, \boldsymbol{\xi} \in \mathbb{R}^n} \quad & \frac{\lambda}{2} \|\mathbf{w}\|_2^2 + \frac{1}{n} \sum_{i=1}^n \xi_i \\ \text{s.t.} \quad & y^{(i)} \left(\langle \mathbf{w}, \mathbf{x}^{(i)} \rangle - \theta \right) \geq 1 - \xi_i \quad \text{for } i = 1, 2, \dots, n \\ & \xi_i \geq 0 \quad \text{for } i = 1, 2, \dots, n \end{aligned}$$

A DIFFERENT INTERPRETATION OF SLACK

Constraints with non-negative slack variables: (using $\lambda = 1/(nC)$)

$$\begin{aligned} \min_{\mathbf{w} \in \mathbb{R}^d, \theta \in \mathbb{R}, \boldsymbol{\xi} \in \mathbb{R}^n} \quad & \frac{\lambda}{2} \|\mathbf{w}\|_2^2 + \frac{1}{n} \sum_{i=1}^n \xi_i \\ \text{s.t.} \quad & y^{(i)} \left(\langle \mathbf{w}, \mathbf{x}^{(i)} \rangle - \theta \right) \geq 1 - \xi_i \quad \text{for } i = 1, 2, \dots, n \\ & \xi_i \geq 0 \quad \text{for } i = 1, 2, \dots, n \end{aligned}$$

Equivalent unconstrained form:

$$\min_{\mathbf{w} \in \mathbb{R}^d, \theta \in \mathbb{R}} \quad \frac{\lambda}{2} \|\mathbf{w}\|_2^2 + \frac{1}{n} \sum_{i=1}^n \left[1 - y^{(i)} \left(\langle \mathbf{w}, \mathbf{x}^{(i)} \rangle - \theta \right) \right]_+$$

Notation: $[a]_+ := \max\{0, a\}$.

A DIFFERENT INTERPRETATION OF SLACK

Constraints with non-negative slack variables: (using $\lambda = 1/(nC)$)

$$\begin{aligned} \min_{\mathbf{w} \in \mathbb{R}^d, \theta \in \mathbb{R}, \xi \in \mathbb{R}^n} \quad & \frac{\lambda}{2} \|\mathbf{w}\|_2^2 + \frac{1}{n} \sum_{i=1}^n \xi_i \\ \text{s.t.} \quad & y^{(i)} (\langle \mathbf{w}, \mathbf{x}^{(i)} \rangle - \theta) \geq 1 - \xi_i \quad \text{for } i = 1, 2, \dots, n \\ & \xi_i \geq 0 \quad \text{for } i = 1, 2, \dots, n \end{aligned}$$

Equivalent unconstrained form:

$$\min_{\mathbf{w} \in \mathbb{R}^d, \theta \in \mathbb{R}} \quad \frac{\lambda}{2} \|\mathbf{w}\|_2^2 + \frac{1}{n} \sum_{i=1}^n \left[1 - y^{(i)} (\langle \mathbf{w}, \mathbf{x}^{(i)} \rangle - \theta) \right]_+$$

Notation: $[a]_+ := \max\{0, a\}$.

The **hinge loss** of a linear classifier $f_{\mathbf{w}, \theta}$ on an example $(\mathbf{x}, y)$ is defined to be

$$\ell_{\text{hl}}(\mathbf{w}, \theta; \mathbf{x}, y) := \left[1 - y(\langle \mathbf{w}, \mathbf{x} \rangle - \theta) \right]_+.$$

A DIFFERENT INTERPRETATION OF SLACK

Constraints with non-negative slack variables: (using $\lambda = 1/(nC)$)

$$\begin{aligned} \min_{\mathbf{w} \in \mathbb{R}^d, \theta \in \mathbb{R}, \boldsymbol{\xi} \in \mathbb{R}^n} \quad & \frac{\lambda}{2} \|\mathbf{w}\|_2^2 + \frac{1}{n} \sum_{i=1}^n \xi_i \\ \text{s.t.} \quad & y^{(i)} \left(\langle \mathbf{w}, \mathbf{x}^{(i)} \rangle - \theta \right) \geq 1 - \xi_i \quad \text{for } i = 1, 2, \dots, n \\ & \xi_i \geq 0 \quad \text{for } i = 1, 2, \dots, n \end{aligned}$$

Equivalent unconstrained form:

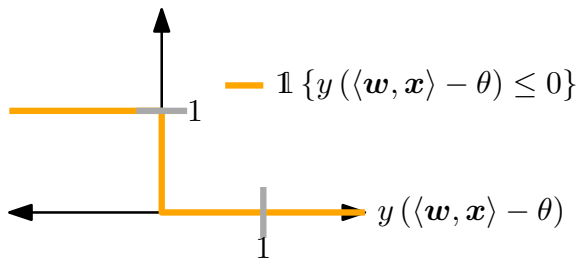
$$\min_{\mathbf{w} \in \mathbb{R}^d, \theta \in \mathbb{R}} \quad \frac{\lambda}{2} \|\mathbf{w}\|_2^2 + \frac{1}{n} \sum_{i=1}^n \ell_{\text{hl}}(\mathbf{w}, \theta; \mathbf{x}^{(i)}, y^{(i)})$$

Notation: $[a]_+ := \max\{0, a\}$.

The **hinge loss** of a linear classifier $f_{\mathbf{w}, \theta}$ on an example $(\mathbf{x}, y)$ is defined to be

$$\ell_{\text{hl}}(\mathbf{w}, \theta; \mathbf{x}, y) := \left[1 - y(\langle \mathbf{w}, \mathbf{x} \rangle - \theta) \right]_+.$$

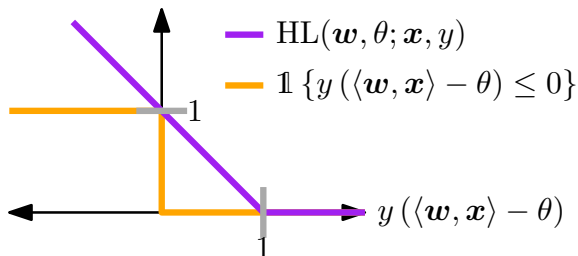
ZERO-ONE LOSS VS. HINGE LOSS



Zero-one loss: count if $f_{\mathbf{w},\theta}(\mathbf{x}) \neq y$.

$$\mathbb{1}\{y(\langle \mathbf{w}, \mathbf{x} \rangle - \theta) \leq 0\} \leq \left[1 - y(\langle \mathbf{w}, \mathbf{x} \rangle - \theta)\right]_+ = \ell_{\text{hl}}(\mathbf{w}, \theta; \mathbf{x}, y).$$

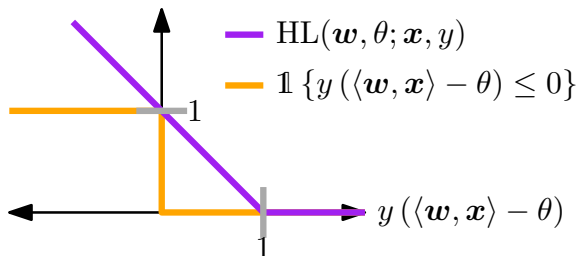
ZERO-ONE LOSS VS. HINGE LOSS



Hinge loss: an upper-bound on **zero-one loss**.

$$\mathbb{1} \{y (\langle \mathbf{w}, \mathbf{x} \rangle - \theta) \leq 0\} \leq \left[1 - y (\langle \mathbf{w}, \mathbf{x} \rangle - \theta) \right]_+ = \ell_{\text{hl}}(\mathbf{w}, \theta; \mathbf{x}, y).$$

ZERO-ONE LOSS VS. HINGE LOSS

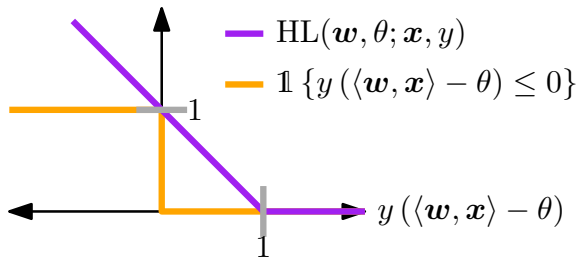


Hinge loss: an upper-bound on **zero-one loss**.

$$\mathbb{1} \{y (\langle \mathbf{w}, \mathbf{x} \rangle - \theta) \leq 0\} \leq \left[1 - y (\langle \mathbf{w}, \mathbf{x} \rangle - \theta) \right]_+ = \ell_{hl}(\mathbf{w}, \theta; \mathbf{x}, y).$$

Soft-margin SVM minimizes an upper-bound on the training error, plus a term that favors large margins.

ZERO-ONE LOSS VS. HINGE LOSS



Hinge loss: an upper-bound on **zero-one loss**.

$$\mathbb{1} \{y(\langle \mathbf{w}, \mathbf{x} \rangle - \theta) \leq 0\} \leq \left[1 - y(\langle \mathbf{w}, \mathbf{x} \rangle - \theta)\right]_+ = \ell_{hl}(\mathbf{w}, \theta; \mathbf{x}, y).$$

Soft-margin SVM minimizes an upper-bound on the training error, plus a term that favors large margins.

This is **computationally tractable** (unlike minimizing training error) because the hinge loss is a **convex function** of $(\mathbf{w}, \theta)$, and so is $\frac{\lambda}{2} \|\mathbf{w}\|_2^2$.

GENERAL FORM

Empirical risk minimization (i.e., minimize training error):

$$\min_{\mathbf{w} \in \mathbb{R}^d, \theta \in \mathbb{R}} \quad \frac{1}{n} \sum_{i=1}^n \mathbb{1} \left\{ y^{(i)} \left(\langle \mathbf{w}, \mathbf{x}^{(i)} \rangle - \theta \right) \leq 0 \right\}$$

Soft-margin SVM:

$$\min_{\mathbf{w} \in \mathbb{R}^d, \theta \in \mathbb{R}} \quad \frac{\lambda}{2} \|\mathbf{w}\|_2^2 + \frac{1}{n} \sum_{i=1}^n \ell_{\text{hl}}(\mathbf{w}, \theta; \mathbf{x}^{(i)}, y^{(i)})$$

GENERAL FORM

Empirical risk minimization (i.e., minimize training error):

$$\min_{\mathbf{w} \in \mathbb{R}^d, \theta \in \mathbb{R}} \frac{1}{n} \sum_{i=1}^n \mathbb{1} \left\{ y^{(i)} \left(\langle \mathbf{w}, \mathbf{x}^{(i)} \rangle - \theta \right) \leq 0 \right\}$$

Soft-margin SVM:

$$\min_{\mathbf{w} \in \mathbb{R}^d, \theta \in \mathbb{R}} \frac{\lambda}{2} \|\mathbf{w}\|_2^2 + \frac{1}{n} \sum_{i=1}^n \ell_{\text{hl}}(\mathbf{w}, \theta; \mathbf{x}^{(i)}, y^{(i)})$$

Generic learning objective:

$$\min_{\mathbf{w} \in \mathbb{R}^d} R(\mathbf{w}) + \frac{1}{n} \sum_{i=1}^n \ell(\mathbf{w}; \mathbf{x}^{(i)}, y^{(i)})$$

- ▶ **Regularization:** encodes “learning bias” (e.g., preference for large margins), sometimes promotes stability.
- ▶ **Data fitting/empirical loss:** how poorly does the classifier “fit” the data.

SIMILARITY TO MAXIMUM LIKELIHOOD

Generic learning objective:

$$\min_{\mathbf{w} \in \mathbb{R}^d} R(\mathbf{w}) + \frac{1}{n} \sum_{i=1}^n \ell(\mathbf{w}; \mathbf{x}^{(i)}, y^{(i)})$$

Maximum likelihood estimation for a parameteric family $\{p(\cdot; \boldsymbol{\theta}) : \boldsymbol{\theta} \in \Theta\}$ (assuming data $\mathbf{z}^{(1)}, \dots, \mathbf{z}^{(n)}$ are iid):

$$\min_{\boldsymbol{\theta} \in \mathbb{R}^d} -\frac{1}{n} \sum_{i=1}^n \log p(\mathbf{z}^{(i)}; \boldsymbol{\theta})$$

(i.e., minimize $1/n$ times negative log-likelihood of parameter $\boldsymbol{\theta}$).

SIMILARITY TO MAXIMUM LIKELIHOOD

Generic learning objective:

$$\min_{\mathbf{w} \in \mathbb{R}^d} R(\mathbf{w}) + \frac{1}{n} \sum_{i=1}^n \ell(\mathbf{w}; \mathbf{x}^{(i)}, y^{(i)})$$

Maximum likelihood estimation for a parameteric family $\{p(\cdot; \boldsymbol{\theta}) : \boldsymbol{\theta} \in \Theta\}$ (assuming data $\mathbf{z}^{(1)}, \dots, \mathbf{z}^{(n)}$ are iid):

$$\min_{\boldsymbol{\theta} \in \mathbb{R}^d} -\frac{1}{n} \sum_{i=1}^n \log p(\mathbf{z}^{(i)}; \boldsymbol{\theta})$$

(i.e., minimize $1/n$ times negative log-likelihood of parameter $\boldsymbol{\theta}$).

Sometimes generic learning objective is called “generalized (and penalized, because of $R(\mathbf{w})$) maximum likelihood estimation”.

LEARNING VIA OPTIMIZATION

- ▶ Many different choices for regularization and loss.
 - ▶ $R(\mathbf{w}) \propto \|\mathbf{w}\|_2^2$: encourage large margins
 - ▶ $R(\mathbf{w}) \propto \|\mathbf{w}\|_1$: encourage $\mathbf{w}$ to be sparse
 - ▶ $R(\mathbf{w}) \propto \sum_{i=1}^n w_i \ln w_i$: “maximum entropy” interpretation
 - ▶ $\ell(\mathbf{w}, \theta; \mathbf{x}, y) = [\langle \mathbf{w}, \mathbf{x} \rangle - \theta - y]_+^2$
 - ▶ $\ell(\mathbf{w}, \theta; \mathbf{x}, y) = \ln(1 + \exp(-y(\langle \mathbf{w}, \mathbf{x} \rangle - \theta)))$ (“logistic regression”)
 - ▶ ...
 - ▶ Also used beyond classification (e.g., regression, parameter estimation, clustering)

LEARNING VIA OPTIMIZATION

- ▶ Many different choices for regularization and loss.
 - ▶ $R(\mathbf{w}) \propto \|\mathbf{w}\|_2^2$: encourage large margins
 - ▶ $R(\mathbf{w}) \propto \|\mathbf{w}\|_1$: encourage $\mathbf{w}$ to be sparse
 - ▶ $R(\mathbf{w}) \propto \sum_{i=1}^n w_i \ln w_i$: “maximum entropy” interpretation
 - ▶ $\ell(\mathbf{w}, \theta; \mathbf{x}, y) = [\langle \mathbf{w}, \mathbf{x} \rangle - \theta - y]_+^2$
 - ▶ $\ell(\mathbf{w}, \theta; \mathbf{x}, y) = \ln(1 + \exp(-y(\langle \mathbf{w}, \mathbf{x} \rangle - \theta)))$ (“logistic regression”)
 - ▶ ...
 - ▶ Also used beyond classification (e.g., regression, parameter estimation, clustering)
- ▶ Often want ℓ be an upper-bound on zero-one loss—i.e., a **surrogate loss**.

LEARNING VIA OPTIMIZATION

- ▶ Many different choices for regularization and loss.
 - ▶ $R(\mathbf{w}) \propto \|\mathbf{w}\|_2^2$: encourage large margins
 - ▶ $R(\mathbf{w}) \propto \|\mathbf{w}\|_1$: encourage $\mathbf{w}$ to be sparse
 - ▶ $R(\mathbf{w}) \propto \sum_{i=1}^n w_i \ln w_i$: “maximum entropy” interpretation
 - ▶ $\ell(\mathbf{w}, \theta; \mathbf{x}, y) = [\langle \mathbf{w}, \mathbf{x} \rangle - \theta - y]_+^2$
 - ▶ $\ell(\mathbf{w}, \theta; \mathbf{x}, y) = \ln(1 + \exp(-y(\langle \mathbf{w}, \mathbf{x} \rangle - \theta)))$ (“logistic regression”)
 - ▶ ...
 - ▶ Also used beyond classification (e.g., regression, parameter estimation, clustering)
- ▶ Often want ℓ be an upper-bound on zero-one loss—i.e., a **surrogate loss**.
- ▶ Trade-off parameter λ : usually determine using hold out error or cross validation error.

LEARNING VIA OPTIMIZATION

- ▶ Many different choices for regularization and loss.
 - ▶ $R(\mathbf{w}) \propto \|\mathbf{w}\|_2^2$: encourage large margins
 - ▶ $R(\mathbf{w}) \propto \|\mathbf{w}\|_1$: encourage $\mathbf{w}$ to be sparse
 - ▶ $R(\mathbf{w}) \propto \sum_{i=1}^n w_i \ln w_i$: “maximum entropy” interpretation
 - ▶ $\ell(\mathbf{w}, \theta; \mathbf{x}, y) = [\langle \mathbf{w}, \mathbf{x} \rangle - \theta - y]_+^2$
 - ▶ $\ell(\mathbf{w}, \theta; \mathbf{x}, y) = \ln(1 + \exp(-y(\langle \mathbf{w}, \mathbf{x} \rangle - \theta)))$ (“logistic regression”)
 - ▶ ...
 - ▶ Also used beyond classification (e.g., regression, parameter estimation, clustering)
- ▶ Often want ℓ be an upper-bound on zero-one loss—i.e., a **surrogate loss**.
- ▶ Trade-off parameter λ : usually determine using hold out error or cross validation error.
- ▶ Computationally easier when overall objective function is **convex**: possible to efficiently find global minimizer in polynomial time.

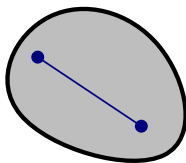
LEARNING VIA OPTIMIZATION

- ▶ Many different choices for regularization and loss.
 - ▶ $R(\mathbf{w}) \propto \|\mathbf{w}\|_2^2$: encourage large margins
 - ▶ $R(\mathbf{w}) \propto \|\mathbf{w}\|_1$: encourage $\mathbf{w}$ to be sparse
 - ▶ $R(\mathbf{w}) \propto \sum_{i=1}^n w_i \ln w_i$: “maximum entropy” interpretation
 - ▶ $\ell(\mathbf{w}, \theta; \mathbf{x}, y) = [\langle \mathbf{w}, \mathbf{x} \rangle - \theta - y]_+^2$
 - ▶ $\ell(\mathbf{w}, \theta; \mathbf{x}, y) = \ln(1 + \exp(-y(\langle \mathbf{w}, \mathbf{x} \rangle - \theta)))$ (“logistic regression”)
 - ▶ ...
 - ▶ Also used beyond classification (e.g., regression, parameter estimation, clustering)
- ▶ Often want ℓ be an upper-bound on zero-one loss—i.e., a **surrogate loss**.
- ▶ Trade-off parameter λ : usually determine using hold out error or cross validation error.
- ▶ Computationally easier when overall objective function is **convex**: possible to efficiently find global minimizer in polynomial time.
- ▶ **Next**: techniques for analyzing and solving these optimization problems.

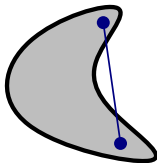
INTRODUCTION TO CONVEXITY

CONVEX SETS

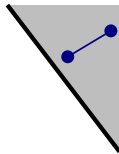
We say a set A is **convex** if, for every pair of points $\{x, x'\}$ in the set A , the line segment between the points x and x' is also contained in the set A .



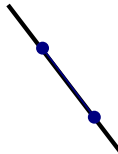
convex



not convex



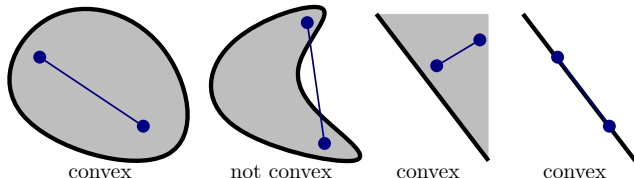
convex



convex

CONVEX SETS

We say a set A is **convex** if, for every pair of points $\{x, x'\}$ in the set A , the line segment between the points x and x' is also contained in the set A .



Examples:

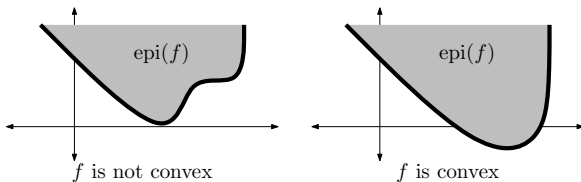
- ▶ All of $\mathbb{R}^d$.
- ▶ Empty set.
- ▶ Affine hyperplanes.
- ▶ Half-spaces: $\{a \in \mathbb{R}^d : \langle a, x \rangle - b \leq 0\}$.
- ▶ Intersections of convex sets.
- ▶ Convex hulls of points.

CONVEX FUNCTIONS

For any function $f: \mathbb{R}^d \rightarrow \mathbb{R}$, the **epigraph** of f , denoted $\text{epi}(f)$, is the set

$$\text{epi}(f) := \{(\mathbf{x}, b) \in \mathbb{R}^{d+1} : f(\mathbf{x}) \leq b\}.$$

We say a function f is **convex** if $\text{epi}(f)$ is a convex set in $\mathbb{R}^{d+1}$.

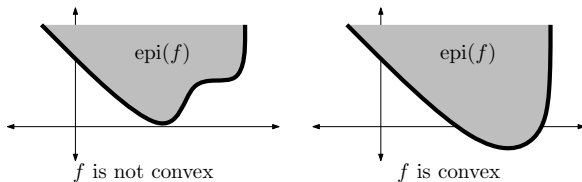


CONVEX FUNCTIONS

For any function $f: \mathbb{R}^d \rightarrow \mathbb{R}$, the **epigraph** of f , denoted $\text{epi}(f)$, is the set

$$\text{epi}(f) := \{(\mathbf{x}, b) \in \mathbb{R}^{d+1} : f(\mathbf{x}) \leq b\}.$$

We say a function f is **convex** if $\text{epi}(f)$ is a convex set in $\mathbb{R}^{d+1}$.



Examples:

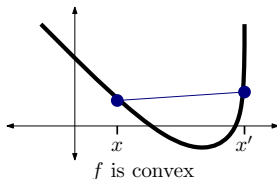
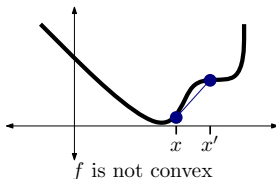
- ▶ $f(x) = c^x$ for any $c > 0$ (on $\mathbb{R}$)
- ▶ $f(x) = |x|^c$ for any $c \geq 1$ (on $\mathbb{R}$)
- ▶ $f(\mathbf{x}) = c$ for any constant $c \in \mathbb{R}$.
- ▶ $f(\mathbf{x}) = \langle \mathbf{a}, \mathbf{x} \rangle$ for any $\mathbf{a} \in \mathbb{R}^d$.
- ▶ $f(\mathbf{x}) = \|\mathbf{x}\|_p$ for any $1 \leq p \leq \infty$.
- ▶ $f(\mathbf{x}) = \langle \mathbf{x}, \mathbf{A}\mathbf{x} \rangle$ for symmetric positive semidefinite $\mathbf{A}$.

JENSEN'S INEQUALITY

Equivalent definition of convex functions

A function $f: \mathbb{R}^d \rightarrow \mathbb{R}$ is convex if and only if, for any $\mathbf{x}, \mathbf{x}' \in \mathbb{R}^d$ and $\alpha \in [0, 1]$,

$$f((1 - \alpha)\mathbf{x} + \alpha\mathbf{x}') \leq (1 - \alpha) \cdot f(\mathbf{x}) + \alpha \cdot f(\mathbf{x}').$$

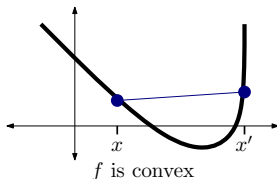
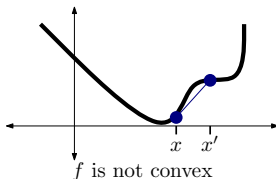


JENSEN'S INEQUALITY

Equivalent definition of convex functions

A function $f: \mathbb{R}^d \rightarrow \mathbb{R}$ is convex if and only if, for any $\mathbf{x}, \mathbf{x}' \in \mathbb{R}^d$ and $\alpha \in [0, 1]$,

$$f((1 - \alpha)\mathbf{x} + \alpha\mathbf{x}') \leq (1 - \alpha) \cdot f(\mathbf{x}) + \alpha \cdot f(\mathbf{x}').$$



Jensen's inequality

If $f: \mathbb{R}^d \rightarrow \mathbb{R}$ is convex, then for any $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n \in \mathbb{R}^d$ and $\alpha_1, \alpha_2, \dots, \alpha_n \in [0, 1]$ such that $\sum_{i=1}^n \alpha_i = 1$,

$$f\left(\sum_{i=1}^n \alpha_i \mathbf{x}_i\right) \leq \sum_{i=1}^n \alpha_i \cdot f(\mathbf{x}_i).$$

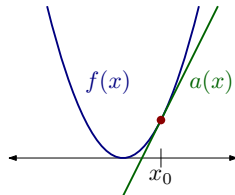
CONVEXITY OF DIFFERENTIABLE FUNCTIONS

Differentiable functions

If $f: \mathbb{R}^d \rightarrow \mathbb{R}$ is differentiable, then f is convex if and only if

$$f(\mathbf{x}) \geq f(\mathbf{x}_0) + \langle \nabla f(\mathbf{x}_0), \mathbf{x} - \mathbf{x}_0 \rangle$$

for all $\mathbf{x}, \mathbf{x}_0 \in \mathbb{R}^d$.



$$a(x) = f(x_0) + f'(x_0)(x - x_0)$$

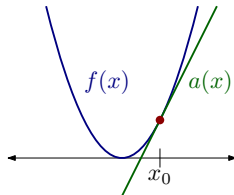
CONVEXITY OF DIFFERENTIABLE FUNCTIONS

Differentiable functions

If $f: \mathbb{R}^d \rightarrow \mathbb{R}$ is differentiable, then f is convex if and only if

$$f(\mathbf{x}) \geq f(\mathbf{x}_0) + \langle \nabla f(\mathbf{x}_0), \mathbf{x} - \mathbf{x}_0 \rangle$$

for all $\mathbf{x}, \mathbf{x}_0 \in \mathbb{R}^d$.



$$a(x) = f(x_0) + f'(x_0)(x - x_0)$$

Twice-differentiable functions

If $f: \mathbb{R}^d \rightarrow \mathbb{R}$ is twice-differentiable, then f is convex if and only if

$$\nabla^2 f(\mathbf{x}) \succeq \mathbf{0}$$

for all $\mathbf{x} \in \mathbb{R}^d$ (i.e., the Hessian, or matrix of second-derivatives, is positive semidefinite for all $\mathbf{x}$).

CONVEX OPTIMIZATION PROBLEMS

OPTIMIZATION PROBLEMS

A typical optimization problem is written as

$$\begin{array}{ll}\min_{\mathbf{x} \in \mathbb{R}^d} & f_0(\mathbf{x}) \\ \text{s.t.} & f_i(\mathbf{x}) \leq 0 \quad i = 1, 2, \dots, n\end{array}$$

where $f_0: \mathbb{R}^d \rightarrow \mathbb{R}$ is the **objective function** and $f_1, f_2, \dots, f_n: \mathbb{R}^d \rightarrow \mathbb{R}$ are the **constraint functions**.

OPTIMIZATION PROBLEMS

A typical optimization problem is written as

$$\begin{array}{ll}\min_{\mathbf{x} \in \mathbb{R}^d} & f_0(\mathbf{x}) \\ \text{s.t.} & f_i(\mathbf{x}) \leq 0 \quad i = 1, 2, \dots, n\end{array}$$

where $f_0: \mathbb{R}^d \rightarrow \mathbb{R}$ is the **objective function** and $f_1, f_2, \dots, f_n: \mathbb{R}^d \rightarrow \mathbb{R}$ are the **constraint functions**.

The inequalities $f_i(\mathbf{x}) \leq 0$ are **constraints**.

OPTIMIZATION PROBLEMS

A typical optimization problem is written as

$$\begin{array}{ll}\min_{\mathbf{x} \in \mathbb{R}^d} & f_0(\mathbf{x}) \\ \text{s.t.} & f_i(\mathbf{x}) \leq 0 \quad i = 1, 2, \dots, n\end{array}$$

where $f_0: \mathbb{R}^d \rightarrow \mathbb{R}$ is the **objective function** and $f_1, f_2, \dots, f_n: \mathbb{R}^d \rightarrow \mathbb{R}$ are the **constraint functions**.

The inequalities $f_i(\mathbf{x}) \leq 0$ are **constraints**.

The set $A := \{\mathbf{x} \in \mathbb{R}^d : f_i(\mathbf{x}) \leq 0 \text{ for all } i = 1, 2, \dots, n\}$ is the **feasible set**.

OPTIMIZATION PROBLEMS

A typical optimization problem is written as

$$\begin{array}{ll}\min_{\mathbf{x} \in \mathbb{R}^d} & f_0(\mathbf{x}) \\ \text{s.t.} & f_i(\mathbf{x}) \leq 0 \quad i = 1, 2, \dots, n\end{array}$$

where $f_0: \mathbb{R}^d \rightarrow \mathbb{R}$ is the **objective function** and $f_1, f_2, \dots, f_n: \mathbb{R}^d \rightarrow \mathbb{R}$ are the **constraint functions**.

The inequalities $f_i(\mathbf{x}) \leq 0$ are **constraints**.

The set $A := \{\mathbf{x} \in \mathbb{R}^d : f_i(\mathbf{x}) \leq 0 \text{ for all } i = 1, 2, \dots, n\}$ is the **feasible set**.

The goal: Find $\mathbf{x} \in A$ so that $f_0(\mathbf{x})$ is as small as possible.

OPTIMIZATION PROBLEMS

A typical optimization problem is written as

$$\begin{array}{ll}\min_{\mathbf{x} \in \mathbb{R}^d} & f_0(\mathbf{x}) \\ \text{s.t.} & f_i(\mathbf{x}) \leq 0 \quad i = 1, 2, \dots, n\end{array}$$

where $f_0: \mathbb{R}^d \rightarrow \mathbb{R}$ is the **objective function** and $f_1, f_2, \dots, f_n: \mathbb{R}^d \rightarrow \mathbb{R}$ are the **constraint functions**.

The inequalities $f_i(\mathbf{x}) \leq 0$ are **constraints**.

The set $A := \{\mathbf{x} \in \mathbb{R}^d : f_i(\mathbf{x}) \leq 0 \text{ for all } i = 1, 2, \dots, n\}$ is the **feasible set**.

The goal: Find $\mathbf{x} \in A$ so that $f_0(\mathbf{x})$ is as small as possible.

The **(optimal) value** of the optimization problem is the smallest such value of $f_0(\mathbf{x})$ achieved by a feasible point $\mathbf{x} \in A$.

OPTIMIZATION PROBLEMS

A typical optimization problem is written as

$$\begin{array}{ll}\min_{\mathbf{x} \in \mathbb{R}^d} & f_0(\mathbf{x}) \\ \text{s.t.} & f_i(\mathbf{x}) \leq 0 \quad i = 1, 2, \dots, n\end{array}$$

where $f_0: \mathbb{R}^d \rightarrow \mathbb{R}$ is the **objective function** and $f_1, f_2, \dots, f_n: \mathbb{R}^d \rightarrow \mathbb{R}$ are the **constraint functions**.

The inequalities $f_i(\mathbf{x}) \leq 0$ are **constraints**.

The set $A := \{\mathbf{x} \in \mathbb{R}^d : f_i(\mathbf{x}) \leq 0 \text{ for all } i = 1, 2, \dots, n\}$ is the **feasible set**.

The goal: Find $\mathbf{x} \in A$ so that $f_0(\mathbf{x})$ is as small as possible.

The **(optimal) value** of the optimization problem is the smallest such value of $f_0(\mathbf{x})$ achieved by a feasible point $\mathbf{x} \in A$.

A point $\mathbf{x} \in A$ achieving the optimal value is a **(global) minimizer** of the problem.

CONVEX OPTIMIZATION PROBLEMS

Standard form of a **convex optimization problem**:

$$\begin{array}{ll}\min_{\mathbf{x} \in \mathbb{R}^d} & f_0(\mathbf{x}) \\ \text{s.t.} & f_i(\mathbf{x}) \leq 0 \quad i = 1, 2, \dots, n\end{array}$$

where $f_0, f_1, \dots, f_n: \mathbb{R}^d \rightarrow \mathbb{R}$ are *convex functions*.

CONVEX OPTIMIZATION PROBLEMS

Standard form of a **convex optimization problem**:

$$\begin{array}{ll}\min_{\mathbf{x} \in \mathbb{R}^d} & f_0(\mathbf{x}) \\ \text{s.t.} & f_i(\mathbf{x}) \leq 0 \quad i = 1, 2, \dots, n\end{array}$$

where $f_0, f_1, \dots, f_n: \mathbb{R}^d \rightarrow \mathbb{R}$ are *convex functions*.

Fact: the feasible set $A := \{\mathbf{x} \in \mathbb{R}^d : f_i(\mathbf{x}) \leq 0 \text{ for all } i = 1, 2, \dots, n\}$ is a convex set.

CONVEX OPTIMIZATION PROBLEMS

Standard form of a **convex optimization problem**:

$$\begin{array}{ll}\min_{\mathbf{x} \in \mathbb{R}^d} & f_0(\mathbf{x}) \\ \text{s.t.} & f_i(\mathbf{x}) \leq 0 \quad i = 1, 2, \dots, n\end{array}$$

where $f_0, f_1, \dots, f_n: \mathbb{R}^d \rightarrow \mathbb{R}$ are *convex functions*.

Fact: the feasible set $A := \{\mathbf{x} \in \mathbb{R}^d : f_i(\mathbf{x}) \leq 0 \text{ for all } i = 1, 2, \dots, n\}$ is a convex set.

Fact: convex optimization problems can be solved efficiently (next lecture)

EXAMPLE: SOFT-MARGIN SVM

$$\begin{aligned} \min_{\mathbf{w} \in \mathbb{R}^d, \theta \in \mathbb{R}, \boldsymbol{\xi} \in \mathbb{R}^n} \quad & \frac{\lambda}{2} \|\mathbf{w}\|_2^2 + \frac{1}{n} \sum_{i=1}^n \xi_i \\ \text{s.t.} \quad & y^{(i)} \left(\langle \mathbf{w}, \mathbf{x}^{(i)} \rangle - \theta \right) \geq 1 - \xi_i \quad \text{for } i = 1, 2, \dots, n \\ & \xi_i \geq 0 \quad \text{for } i = 1, 2, \dots, n \end{aligned}$$

EXAMPLE: SOFT-MARGIN SVM

$$\begin{aligned} \min_{\mathbf{w} \in \mathbb{R}^d, \theta \in \mathbb{R}, \boldsymbol{\xi} \in \mathbb{R}^n} \quad & \frac{\lambda}{2} \|\mathbf{w}\|_2^2 + \frac{1}{n} \sum_{i=1}^n \xi_i \\ \text{s.t.} \quad & y^{(i)} \left(\langle \mathbf{w}, \mathbf{x}^{(i)} \rangle - \theta \right) \geq 1 - \xi_i \quad \text{for } i = 1, 2, \dots, n \\ & \xi_i \geq 0 \quad \text{for } i = 1, 2, \dots, n \end{aligned}$$

Bringing to standard form:

$$\begin{aligned} \min_{\mathbf{w} \in \mathbb{R}^d, \theta \in \mathbb{R}, \boldsymbol{\xi} \in \mathbb{R}^n} \quad & \frac{\lambda}{2} \|\mathbf{w}\|_2^2 + \frac{1}{n} \sum_{i=1}^n \xi_i \quad \text{(sum of 2 convex functions, also convex)} \\ \text{s.t.} \quad & 1 - \xi_i - y^{(i)} \left(\langle \mathbf{w}, \mathbf{x}^{(i)} \rangle - \theta \right) \leq 0 \quad \text{for } i = 1, \dots, n \quad \text{(linear)} \\ & -\xi_i \leq 0 \quad \text{for } i = 1, \dots, n \quad \text{(linear)} \end{aligned}$$