

COMS 4771 Lecture 13

1. Beyond binary classification
2. Beyond prediction error
3. Cross-validation

BEYOND BINARY CLASSIFICATION

REDUCTIONS

Suppose we have a good algorithm $\mathcal{A}$ for learning binary classifiers.

Question: Can we use $\mathcal{A}$ to solve other machine learning problems?

REDUCTIONS

Suppose we have a good algorithm $\mathcal{A}$ for learning binary classifiers.

Question: Can we use $\mathcal{A}$ to solve other machine learning problems?

A **reduction** for a given machine learning problem to binary classification is specified by two procedures:

- ▶ **Instance mapping procedure:** transforms samples for the given problem into samples for a binary classification problem
- ▶ **Solution mapping procedure:** transforms classifier f' for binary classification problem computed using $\mathcal{A}$ into a classifier f for the given problem.

REDUCTIONS

Suppose we have a good algorithm $\mathcal{A}$ for learning binary classifiers.

Question: Can we use $\mathcal{A}$ to solve other machine learning problems?

A **reduction** for a given machine learning problem to binary classification is specified by two procedures:

- ▶ **Instance mapping procedure:** transforms samples for the given problem into samples for a binary classification problem
- ▶ **Solution mapping procedure:** transforms classifier f' for binary classification problem computed using $\mathcal{A}$ into a classifier f for the given problem.

Desired property:

If $\mathcal{A}$ finds a good binary classifier f' , then f is a good classifier for the given problem.

EXAMPLES

1. **Problem:** finding high accuracy classifiers
 - ▶ **Reduction:** boosting
(Reduces problem to finding weak classifiers.)
2. **Problem:** importance-weighted classification
 - ▶ **Reduction:** rejection sampling
(Reduces problem to unweighted classification.)
3. **Problem:** multi-class classification
 - ▶ **Reduction:** One-Against-All
(Reduces problem to binary classification.)

IMPORTANCE-WEIGHTED CLASSIFICATION

Problem:

- ▶ **Setting:** Random triple $(X, Y, W) \sim P$ for some probability distribution P over $\mathcal{X} \times \mathcal{Y} \times \mathbb{R}_+$.

$W =$ **importance weight** for labeled example (X, Y) .

- ▶ **Goal:** Function $f: \mathcal{X} \rightarrow \mathcal{Y}$ with small **importance-weighted error**:

$$\mathbb{E}\left[W \cdot \mathbb{1}\{f(X) \neq Y\}\right].$$

IMPORTANCE-WEIGHTED CLASSIFICATION

Problem:

- ▶ **Setting:** Random triple $(X, Y, W) \sim P$ for some probability distribution P over $\mathcal{X} \times \mathcal{Y} \times \mathbb{R}_+$.

$W =$ **importance weight** for labeled example (X, Y) .

- ▶ **Goal:** Function $f: \mathcal{X} \rightarrow \mathcal{Y}$ with small **importance-weighted error**:

$$\mathbb{E}\left[W \cdot \mathbb{1}\{f(X) \neq Y\}\right].$$

IMPORTANCE-WEIGHTED CLASSIFICATION

Problem:

- ▶ **Setting:** Random triple $(X, Y, W) \sim P$ for some probability distribution P over $\mathcal{X} \times \mathcal{Y} \times \mathbb{R}_+$.

$W =$ **importance weight** for labeled example (X, Y) .

- ▶ **Goal:** Function $f: \mathcal{X} \rightarrow \mathcal{Y}$ with small **importance-weighted error**:

$$\mathbb{E}\left[W \cdot \mathbb{1}\{f(X) \neq Y\}\right].$$

Where it comes up:

- ▶ *Class-specific weights:* e.g., $W = 100 \Leftrightarrow Y = 0$ (and $W = 1$ otherwise).
- ▶ *Covariate-specific weights:* e.g., $W = 100 \Leftrightarrow X \in \mathcal{X}_0$ (and $W = 1$ o.w.).
- ▶ *Boosting, domain adaptation, causal inference, ...*

(Note: many learning algorithms natively handle importance weights.)

IMPORTANCE-WEIGHTED CLASSIFICATION

Problem:

- ▶ **Setting:** Random triple $(X, Y, W) \sim P$ for some probability distribution P over $\mathcal{X} \times \mathcal{Y} \times \mathbb{R}_+$.

$W =$ **importance weight** for labeled example (X, Y) .

- ▶ **Goal:** Function $f: \mathcal{X} \rightarrow \mathcal{Y}$ with small **importance-weighted error**:

$$\mathbb{E}\left[W \cdot \mathbb{1}\{f(X) \neq Y\}\right].$$

Where it comes up:

- ▶ *Class-specific weights:* e.g., $W = 100 \Leftrightarrow Y = 0$ (and $W = 1$ otherwise).
- ▶ *Covariate-specific weights:* e.g., $W = 100 \Leftrightarrow X \in \mathcal{X}_0$ (and $W = 1$ o.w.).
- ▶ *Boosting, domain adaptation, causal inference, ...*

(Note: many learning algorithms natively handle importance weights.)

Would like to reduce to (unweighted) classification.

THE REJECTION SAMPLING REDUCTION

Main idea: Transform samples from P so they appear to be drawn from a distribution P' , where

$$\mathbb{E}_{(X,Y,W) \sim P} \left[W \cdot \mathbb{1}\{f(X) \neq Y\} \right] \propto \mathbb{E}_{(X',Y') \sim P'} \left[\mathbb{1}\{f(X') \neq Y'\} \right].$$

THE REJECTION SAMPLING REDUCTION

Main idea: Transform samples from P so they appear to be drawn from a distribution P' , where

$$\mathbb{E}_{(X,Y,W) \sim P} \left[W \cdot \mathbb{1}\{f(X) \neq Y\} \right] \propto \mathbb{E}_{(X',Y') \sim P'} \left[\mathbb{1}\{f(X') \neq Y'\} \right].$$

Instance mapping procedure

Input Samples drawn from P

- 1: Let $w_{\max}$ be the maximum possible weight
- 2: **while** true **do**
- 3: Sample (X, Y, W) from P
- 4: Toss a coin with heads bias $\frac{W}{w_{\max}}$.
- 5: If heads, break from while loop
- 6: If tails, discard example.
- 7: **end while**
- 8: **return** (X, Y)

THE REJECTION SAMPLING REDUCTION

Main idea: Transform samples from P so they appear to be drawn from a distribution P' , where

$$\mathbb{E}_{(X,Y,W) \sim P} \left[W \cdot \mathbb{1}\{f(X) \neq Y\} \right] \propto \mathbb{E}_{(X',Y') \sim P'} \left[\mathbb{1}\{f(X') \neq Y'\} \right].$$

Instance mapping procedure

Input Samples drawn from P

- 1: Let $w_{\max}$ be the maximum possible weight
- 2: **while** true **do**
- 3: Sample (X, Y, W) from P
- 4: Toss a coin with heads bias $\frac{W}{w_{\max}}$.
- 5: If heads, break from while loop
- 6: If tails, discard example.
- 7: **end while**
- 8: **return** (X, Y)

Solution mapping procedure: identity map

MULTI-CLASS CLASSIFICATION

Problem:

- ▶ **Setting:** Random pair $(X, Y) \sim P$ for some probability distribution P over $\mathcal{X} \times \{1, 2, \dots, K\}$.
- ▶ **Goal:** Function $f: \mathcal{X} \rightarrow \mathcal{Y}$ with small prediction error $\Pr(f(X) \neq Y)$.

MULTI-CLASS CLASSIFICATION

Problem:

- ▶ **Setting:** Random pair $(X, Y) \sim P$ for some probability distribution P over $\mathcal{X} \times \{1, 2, \dots, K\}$.
- ▶ **Goal:** Function $f: \mathcal{X} \rightarrow \mathcal{Y}$ with small prediction error $\Pr(f(X) \neq Y)$.

Would like to reduce to binary classification.

ONE-AGAINST-ALL REDUCTION

Main idea: Create K binary classification problems

given $x \in \mathcal{X}$, predict whether or not $y = i$.

From each sample (x, y) drawn from P , create K samples, one for each binary classification problem:

$$(x, y) \longrightarrow \left\{ \begin{array}{ll} (x, \mathbb{1}\{y = 1\}) & \longrightarrow \text{Binary classification problem 1} \\ (x, \mathbb{1}\{y = 2\}) & \longrightarrow \text{Binary classification problem 2} \\ \vdots & \vdots \\ (x, \mathbb{1}\{y = K\}) & \longrightarrow \text{Binary classification problem } K \end{array} \right.$$

ONE-AGAINST-ALL REDUCTION

Instance mapping procedure

Input Sample $(X, Y) \in \mathcal{X} \times \{1, 2, \dots, K\}$ drawn from P

- 1: **for** each $i = 1, 2, \dots, K$ **do**
- 2: Generate sample $(X, \mathbb{1}\{Y = i\}) \in \mathcal{X} \times \{0, 1\}$ for binary classification problem i
- 3: **end for**

ONE-AGAINST-ALL REDUCTION

Instance mapping procedure

Input Sample $(X, Y) \in \mathcal{X} \times \{1, 2, \dots, K\}$ drawn from P

- 1: **for** each $i = 1, 2, \dots, K$ **do**
- 2: Generate sample $(X, \mathbb{1}\{Y = i\}) \in \mathcal{X} \times \{0, 1\}$ for binary classification problem i
- 3: **end for**

Solution mapping procedure

Input K binary predictors $f'_1, f'_2, \dots, f'_K: \mathcal{X} \rightarrow \{0, 1\}$.

return Function $f: \mathcal{X} \rightarrow \{1, 2, \dots, K\}$ where

$$f(x) = \arg \max_{i \in \{1, 2, \dots, K\}} f'_i(x) \quad (\text{breaking ties arbitrarily}).$$

- ▶ **Reductions:** reuse existing technology to solve new problems.
 - ▶ Multi-class
 - ▶ Multi-label prediction
 - ▶ Ranking
 - ▶ Sequence prediction
 - ▶ ...
- ▶ **Lots of different problems and objectives** beyond binary classification and prediction error—can be application-/domain-specific.

BEYOND PREDICTION ERROR

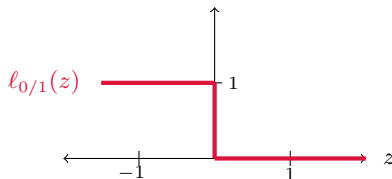
PREDICTION ERROR / ZERO-ONE LOSS

P is a distribution over $\mathcal{X} \times \{-1, +1\}$, and $(X, Y) \sim P$.

For any classifier $f: \mathcal{X} \rightarrow \{-1, +1\}$,

$$\text{err}(f) = \Pr[f(X) \neq Y] = \mathbb{E}[\ell_{0/1}(Yf(X))].$$

$$\ell_{0/1}(x) = (1 - \text{sign}(x))/2$$



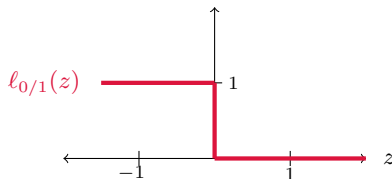
PREDICTION ERROR / ZERO-ONE LOSS

P is a distribution over $\mathcal{X} \times \{-1, +1\}$, and $(X, Y) \sim P$.

For any classifier $f: \mathcal{X} \rightarrow \{-1, +1\}$,

$$\text{err}(f) = \Pr[f(X) \neq Y] = \mathbb{E}[\ell_{0/1}(Yf(X))].$$

$$\ell_{0/1}(x) = (1 - \text{sign}(x))/2$$



Also works with **real-valued predictors** $f: \mathcal{X} \rightarrow \mathbb{R}$; where the prediction is $\text{sign}(f(x))$, for example:

- ▶ **Linear classifiers:** $x \mapsto \langle w, x \rangle - \theta$.
- ▶ **(Binary) plug-in classifiers:** $x \mapsto \widehat{\Pr}[Y = +1 | X = x] - 1/2$.

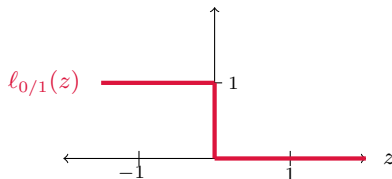
PREDICTION ERROR / ZERO-ONE LOSS

P is a distribution over $\mathcal{X} \times \{-1, +1\}$, and $(X, Y) \sim P$.

For any classifier $f: \mathcal{X} \rightarrow \{-1, +1\}$,

$$\text{err}(f) = \Pr[f(X) \neq Y] = \mathbb{E}[\ell_{0/1}(Yf(X))].$$

$$\ell_{0/1}(x) = (1 - \text{sign}(x))/2$$



Also works with **real-valued predictors** $f: \mathcal{X} \rightarrow \mathbb{R}$; where the prediction is $\text{sign}(f(x))$, for example:

- ▶ **Linear classifiers:** $x \mapsto \langle w, x \rangle - \theta$.
- ▶ **(Binary) plug-in classifiers:** $x \mapsto \widehat{\Pr}[Y = +1 \mid X = x] - 1/2$.

Often useful to adjust threshold (e.g., θ and $1/2$ above).

ADJUSTING THRESHOLDS

Often have **different costs for different kinds of mistakes**:

	$f(X) \leq t$	$f(X) > t$
$Y = -1$	0	a
$Y = +1$	b	0

ADJUSTING THRESHOLDS

Often have **different costs for different kinds of mistakes**:

	$f(X) \leq t$	$f(X) > t$
$Y = -1$	0	a
$Y = +1$	b	0

If $b \gg a$, might want to decrease t .

ADJUSTING THRESHOLDS

Often have **different costs for different kinds of mistakes**:

	$f(X) \leq t$	$f(X) > t$
$Y = -1$	0	a
$Y = +1$	b	0

If $b \gg a$, might want to decrease t .

If $a \gg b$, might want to increase t .

ADJUSTING THRESHOLDS

Also, often interested in **different performance criteria**: given test set T ,

► **Precision**:

$$\frac{|\{(\mathbf{x}, 1) \in T \text{ s.t. } f(\mathbf{x}) > t\}|}{|\{(\mathbf{x}, y) \in T \text{ s.t. } f(\mathbf{x}) > t\}|}$$

► **Recall** (a.k.a. **Sensitivity**, **True Positive Rate**):

$$\frac{|\{(\mathbf{x}, 1) \in T \text{ s.t. } f(\mathbf{x}) > t\}|}{|\{(\mathbf{x}, 1) \in T\}|}$$

► **False Positive Rate**:

$$\frac{|\{(\mathbf{x}, -1) \in T \text{ s.t. } f(\mathbf{x}) > t\}|}{|\{(\mathbf{x}, -1) \in T\}|}$$

ADJUSTING THRESHOLDS

Also, often interested in **different performance criteria**: given test set T ,

► **Precision**:

$$\frac{|\{(\mathbf{x}, 1) \in T \text{ s.t. } f(\mathbf{x}) > t\}|}{|\{(\mathbf{x}, y) \in T \text{ s.t. } f(\mathbf{x}) > t\}|}$$

► **Recall** (a.k.a. **Sensitivity**, **True Positive Rate**):

$$\frac{|\{(\mathbf{x}, 1) \in T \text{ s.t. } f(\mathbf{x}) > t\}|}{|\{(\mathbf{x}, 1) \in T\}|}$$

► **False Positive Rate**:

$$\frac{|\{(\mathbf{x}, -1) \in T \text{ s.t. } f(\mathbf{x}) > t\}|}{|\{(\mathbf{x}, -1) \in T\}|}$$

Precision generally improves if we increase t .

ADJUSTING THRESHOLDS

Also, often interested in **different performance criteria**: given test set T ,

► **Precision**:

$$\frac{|\{(\mathbf{x}, 1) \in T \text{ s.t. } f(\mathbf{x}) > t\}|}{|\{(\mathbf{x}, y) \in T \text{ s.t. } f(\mathbf{x}) > t\}|}$$

► **Recall** (a.k.a. **Sensitivity**, **True Positive Rate**):

$$\frac{|\{(\mathbf{x}, 1) \in T \text{ s.t. } f(\mathbf{x}) > t\}|}{|\{(\mathbf{x}, 1) \in T\}|}$$

► **False Positive Rate**:

$$\frac{|\{(\mathbf{x}, -1) \in T \text{ s.t. } f(\mathbf{x}) > t\}|}{|\{(\mathbf{x}, -1) \in T\}|}$$

Precision generally improves if we increase t .

Recall improves if we decrease t .

ADJUSTING THRESHOLDS

Also, often interested in **different performance criteria**: given test set T ,

► **Precision**:

$$\frac{|\{(\mathbf{x}, 1) \in T \text{ s.t. } f(\mathbf{x}) > t\}|}{|\{(\mathbf{x}, y) \in T \text{ s.t. } f(\mathbf{x}) > t\}|}$$

► **Recall** (a.k.a. **Sensitivity**, **True Positive Rate**):

$$\frac{|\{(\mathbf{x}, 1) \in T \text{ s.t. } f(\mathbf{x}) > t\}|}{|\{(\mathbf{x}, 1) \in T\}|}$$

► **False Positive Rate**:

$$\frac{|\{(\mathbf{x}, -1) \in T \text{ s.t. } f(\mathbf{x}) > t\}|}{|\{(\mathbf{x}, -1) \in T\}|}$$

Precision generally improves if we increase t .

Recall improves if we decrease t .

False positive rate improves if we increase t .

CONDITIONAL PROBABILITY ESTIMATION

Sometimes would like real-valued predictor f to be related to the **conditional probability function** η :

$$\eta(x) := \Pr[Y = +1 \mid X = x].$$

CONDITIONAL PROBABILITY ESTIMATION

Sometimes would like real-valued predictor f to be related to the **conditional probability function** η :

$$\eta(x) := \Pr[Y = +1 \mid X = x].$$

Option #1: Can use **generative models** to construct **plug-in classifier**, but **could fail if class conditional distributions are not accurate**.

CONDITIONAL PROBABILITY ESTIMATION

Sometimes would like real-valued predictor f to be related to the **conditional probability function** η :

$$\eta(x) := \Pr[Y = +1 \mid X = x].$$

- Option #1: Can use **generative models** to construct **plug-in classifier**, but **could fail if class conditional distributions are not accurate**.
- Option #2: Can use a loss function that is minimized by η (or some invertible transformation thereof).

CONDITIONAL PROBABILITY ESTIMATION

Sometimes would like real-valued predictor f to be related to the **conditional probability function** η :

$$\eta(x) := \Pr[Y = +1 \mid X = x].$$

Option #1: Can use **generative models** to construct **plug-in classifier**, but **could fail if class conditional distributions are not accurate**.

Option #2: Can use a loss function that is minimized by η (or some invertible transformation thereof).

I.e. find a function $f : \mathcal{X} \rightarrow \mathbb{R}$ that minimizes

$$\frac{1}{n} \sum_{i=1}^n \ell(y_i f(\mathbf{x}_i))$$

for some function ℓ that is different from $\ell_{0/1}$.

ELICITING CONDITIONAL PROBABILITIES

Goal: loss function that is minimized by (some invertible transformation of) the conditional probability function

$$\eta(x) = \Pr[Y = +1 \mid X = x].$$

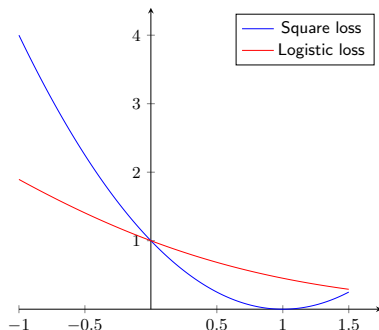
(Y takes values in $\{-1, +1\}$.)

ELICITING CONDITIONAL PROBABILITIES

Goal: loss function that is minimized by (some invertible transformation of) the conditional probability function

$$\eta(x) = \Pr[Y = +1 | X = x].$$

(Y takes values in $\{-1, +1\}$.)



- **Square loss:** $\ell_{\text{sq}}(z) = (1 - z)^2$

$\mathbb{E}[\ell_{\text{sq}}(Y f(x)) | X = x]$ is minimized by $f(x) = 2\eta(x) - 1$.

- **Logistic loss:** $\ell_{\log}(z) = \ln(1 + \exp(-z))$

$\mathbb{E}[\ell_{\log}(Y f(x)) | X = x]$ is minimized by $f(x) = \ln\left(\frac{\eta(x)}{1 - \eta(x)}\right)$.

ELICITING CONDITIONAL PROBABILITIES

Suppose we assume a linear functional form for f , i.e. $f(\mathbf{x}) = \langle \mathbf{w}, \mathbf{x} \rangle$

(assume bias incorporated into $\mathbf{w}$ by adding a constant 1 feature to examples if necessary)

Using loss functions: if loss function $\ell : \mathbb{R} \rightarrow \mathbb{R}$ is convex, then ERM (given below) is a convex problem:

$$\min_{\mathbf{w} \in \mathbb{R}^d} \frac{1}{n} \sum_{i=1}^n \ell(y_i \langle \mathbf{w}, \mathbf{x}_i \rangle)$$

ELICITING CONDITIONAL PROBABILITIES

Suppose we assume a linear functional form for f , i.e. $f(\mathbf{x}) = \langle \mathbf{w}, \mathbf{x} \rangle$

(assume bias incorporated into $\mathbf{w}$ by adding a constant 1 feature to examples if necessary)

Using loss functions: if loss function $\ell : \mathbb{R} \rightarrow \mathbb{R}$ is convex, then ERM (given below) is a convex problem:

$$\min_{\mathbf{w} \in \mathbb{R}^d} \frac{1}{n} \sum_{i=1}^n \ell(y_i \langle \mathbf{w}, \mathbf{x}_i \rangle)$$

Good news: Both ℓ_{sq} and ℓ_{log} are convex!

ELICITING CONDITIONAL PROBABILITIES

Suppose we assume a linear functional form for f , i.e. $f(\mathbf{x}) = \langle \mathbf{w}, \mathbf{x} \rangle$

(assume bias incorporated into $\mathbf{w}$ by adding a constant 1 feature to examples if necessary)

Using loss functions: if loss function $\ell : \mathbb{R} \rightarrow \mathbb{R}$ is convex, then ERM (given below) is a convex problem:

$$\min_{\mathbf{w} \in \mathbb{R}^d} \frac{1}{n} \sum_{i=1}^n \ell(y_i \langle \mathbf{w}, \mathbf{x}_i \rangle)$$

Good news: Both ℓ_{sq} and ℓ_{log} are convex!

In practice, usually incorporate (convex) regularization function R (e.g. $R(\mathbf{w}) = \lambda \|\mathbf{w}\|^2$) to avoid overfitting:

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

ELICITING CONDITIONAL PROBABILITIES

A non-example:

- ▶ **Hinge loss:** $\ell_{\text{hl}}(z) = \max\{0, 1 - z\}$

$\mathbb{E}[\ell_{\text{hl}}(Y f(x)) | X = x]$ is minimized by $f(x) = \text{sign}(2\eta(x) - 1)$.

Can't recover $\eta(x)$ from $f(x)$.

ELICITING CONDITIONAL PROBABILITIES

A non-example:

- **Hinge loss:** $\ell_{\text{hl}}(z) = \max\{0, 1 - z\}$

$\mathbb{E}[\ell_{\text{hl}}(Y f(x)) | X = x]$ is minimized by $f(x) = \text{sign}(2\eta(x) - 1)$.

Can't recover $\eta(x)$ from $f(x)$.

Caveat: Might not be possible to represent

$$\mathbf{x} \mapsto 2\eta(\mathbf{x}) - 1 \quad \text{or} \quad \mathbf{x} \mapsto \ln\left(\frac{\eta(\mathbf{x})}{1 - \eta(\mathbf{x})}\right)$$

as (say) a linear function $\mathbf{x} \mapsto \langle \mathbf{w}, \mathbf{x} \rangle$.

ELICITING CONDITIONAL PROBABILITIES

A non-example:

- **Hinge loss:** $\ell_{\text{hl}}(z) = \max\{0, 1 - z\}$

$\mathbb{E}[\ell_{\text{hl}}(Y f(x)) | X = x]$ is minimized by $f(x) = \text{sign}(2\eta(x) - 1)$.

Can't recover $\eta(x)$ from $f(x)$.

Caveat: Might not be possible to represent

$$x \mapsto 2\eta(x) - 1 \quad \text{or} \quad x \mapsto \ln\left(\frac{\eta(x)}{1 - \eta(x)}\right)$$

as (say) a linear function $x \mapsto \langle w, x \rangle$.

Common solutions: enhance the feature space via feature expansion or kernels, or use more flexible models (e.g., tree models).

CROSS VALIDATION

MODEL SELECTION

Objective

- ▶ Often necessary to consider many different models for a given problem (e.g., class conditional distributions in generative model classifiers, features in linear classifiers, kernel in kernelized classifiers).
- ▶ Sometimes “model” simply means particular setting of **hyper-parameters** (e.g., k in k -NN, λ in soft-margin SVM, number of nodes in decision tree).

Terminology

The problem of choosing a good model is called **model selection**.

EXAMPLE: SVM WITH GAUSSIAN KERNEL

Soft-margin SVM with Gaussian kernel

- ▶ Models indexed by regularization parameter λ and Gaussian kernel bandwidth $h > 0$:

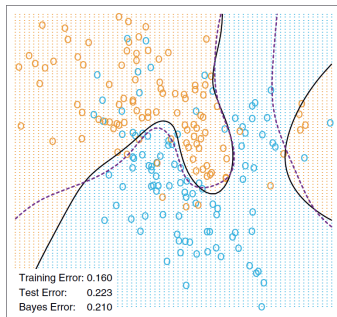
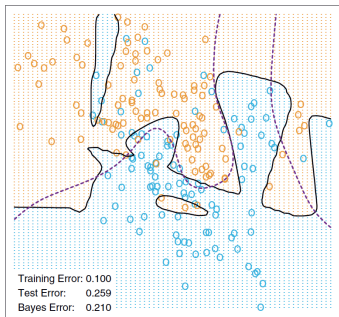
$$K(\mathbf{x}, \tilde{\mathbf{x}}) = \exp\left(-\frac{\|\mathbf{x} - \tilde{\mathbf{x}}\|_2^2}{2h}\right).$$

- ▶ Goal is to find setting of (λ, h) for which we can expect small true (generalization) error.

Naïve approach

- ▶ Also minimize over (λ, h) in SVM optimization problem.
- ▶ Leads to **overfitting**: resulting SVM classifier adapts too closely to specific properties of the training data, rather than underlying distribution.

OVERFITTING: ILLUSTRATION



- ▶ Classifier in this example has bandwidth parameter σ (similar to Gaussian kernel bandwidth).
- ▶ Small $\sigma \rightarrow$ permits curve with sharp bends
- ▶ Large $\sigma \rightarrow$ smoother curve.

MODEL SELECTION BY HOLD-OUT VALIDATION

(Henceforth, use $\mathbf{h}$ to denote particular setting of hyper-parameters / model choice.)

Hold-out validation

Model selection:

1. Randomly split data into three sets: **training**, **validation**, and **test** data.

Training	Validation	Test
----------	------------	------

2. Train classifier $\hat{f}_{\mathbf{h}}$ on **Training** data for different values of $\mathbf{h}$.
3. Compute **Validation** (“hold-out”) error for each $\hat{f}_{\mathbf{h}}$: $\text{err}(\hat{f}_{\mathbf{h}}, \text{Validation})$.
4. Selection: $\hat{\mathbf{h}}$ = value of $\mathbf{h}$ with lowest **Validation** error.
5. Train classifier $\hat{f}$ using $\hat{\mathbf{h}}$ with **Training** + **Validation** data.

Model assessment:

6. Finally: estimate the error of $\hat{f}$ using **test** data.

MAIN IDEA BEHIND HOLD-OUT VALIDATION

Training	Validation	Test
----------	------------	------

Classifier $\hat{f}_h$ trained on Training data $\rightarrow \text{err}(\hat{f}_h, \text{Validation})$.

Training + Validation	Test
-----------------------	------

Classifier $\hat{f}_h$ trained on Training + Validation data $\rightarrow \text{err}(\hat{f}_h, \text{Test})$.

MAIN IDEA BEHIND HOLD-OUT VALIDATION

Training	Validation	Test
----------	------------	------

Classifier $\hat{f}_h$ trained on Training data $\rightarrow \text{err}(\hat{f}_h, \text{Validation})$.

Training + Validation	Test
-----------------------	------

Classifier $\hat{f}_h$ trained on Training + Validation data $\rightarrow \text{err}(\hat{f}_h, \text{Test})$.

The hope is that these quantities are similar!

MAIN IDEA BEHIND HOLD-OUT VALIDATION

Training	Validation	Test
----------	------------	------

Classifier $\hat{f}_h$ trained on Training data $\rightarrow \text{err}(\hat{f}_h, \text{Validation})$.

Training + Validation	Test
-----------------------	------

Classifier $\hat{f}_h$ trained on Training + Validation data $\rightarrow \text{err}(\hat{f}_h, \text{Test})$.

The hope is that these quantities are similar!

(Making this rigorous is actually rather tricky.)

BEYOND SIMPLE HOLD-OUT VALIDATION

Standard hold-out:

Training	Validation	Test
----------	------------	------

Classifier $\hat{f}_h$ trained on Training data $\rightarrow \text{err}(\hat{f}_h, \text{Validation})$.

BEYOND SIMPLE HOLD-OUT VALIDATION

Standard hold-out:

Training	Validation	Test
----------	------------	------

Classifier $\hat{f}_h$ trained on Training data $\rightarrow \text{err}(\hat{f}_h, \text{Validation})$.

Could also:

- ▶ train $\hat{f}_h$ using Validation data, and
- ▶ evaluate $\hat{f}_h$ using Training data.

Training	Validation	Test
----------	------------	------

Classifier $\hat{f}_h$ trained on Validation data $\rightarrow \text{err}(\hat{f}_h, \text{Training})$.

BEYOND SIMPLE HOLD-OUT VALIDATION

Standard hold-out:

Training	Validation	Test
----------	------------	------

Classifier $\hat{f}_h$ trained on Training data $\rightarrow \text{err}(\hat{f}_h, \text{Validation})$.

Could also:

- ▶ train $\hat{f}_h$ using Validation data, and
- ▶ evaluate $\hat{f}_h$ using Training data.

Training	Validation	Test
----------	------------	------

Classifier $\hat{f}_h$ trained on Validation data $\rightarrow \text{err}(\hat{f}_h, \text{Training})$.

Idea: Do both, and average results as overall validation error for h .

MODEL SELECTION BY K -FOLD CROSS VALIDATION

Model selection:

1. Set aside some **test** data.
2. Of remaining data, split into K parts (“folds”, usually 5 or 10) $S_1, S_2, \dots, S_K$.
3. For each value of h :
 - ▶ For each $k \in \{1, 2, \dots, K\}$:
 - ▶ Train classifier $\hat{f}_{h,k}$ using all S_i except S_k .
 - ▶ Evaluate classifier $\hat{f}_{h,k}$ using S_k : $\text{err}(\hat{f}_{h,k}, S_k)$

Example: $K = 5$ and $k = 4$

Training	Training	Training	Validation	Training
----------	----------	----------	------------	----------

- ▶ Cross-validation error for h : $\frac{1}{K} \sum_{k=1}^K \text{err}(\hat{f}_{h,k}, S_k)$.

4. Select the value $\hat{h}$ with lowest cross-validation error.
5. Train classifier $\hat{f}$ using selected $\hat{h}$ with all $S_1, S_2, \dots, S_K$.

Model assessment:

6. Finally: estimate the error of $\hat{f}$ using **test** data.

- ▶ **Model selection:** goal is to pick best model (e.g., features, kernels, hyper-parameter settings) to achieve low true error.

- ▶ **Model selection:** goal is to pick best model (e.g., features, kernels, hyper-parameter settings) to achieve low true error.
- ▶ **Two common methods:** hold-out validation and K -fold cross validation (with $K = 5$ or $K = 10$).

- ▶ **Model selection:** goal is to pick best model (e.g., features, kernels, hyper-parameter settings) to achieve low true error.
- ▶ **Two common methods:** hold-out validation and K -fold cross validation (with $K = 5$ or $K = 10$).
- ▶ **Caution:** considering *too many* different models can lead to overfitting, even with hold-out / cross-validation.