

The background of the slide is an abstract composition. It features a series of horizontal stripes in various colors including pink, green, orange, yellow, and blue. Overlaid on these stripes are several thick, diagonal lines in colors like orange, yellow, pink, and dark brown, creating a complex, layered visual effect.

Linear Classifiers: Logistic Regression and SVM

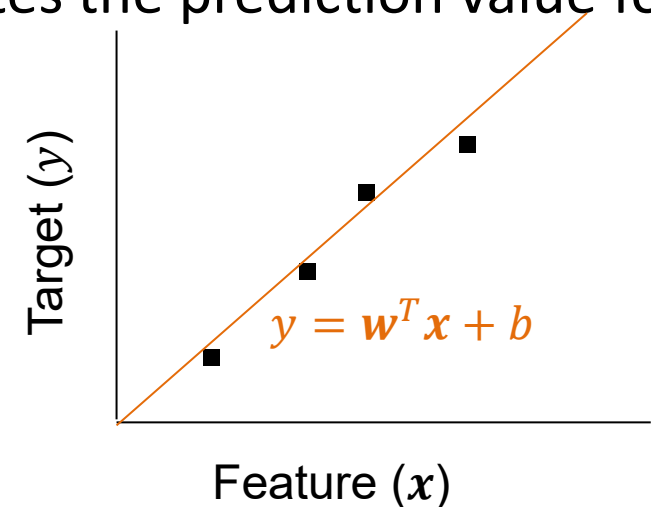
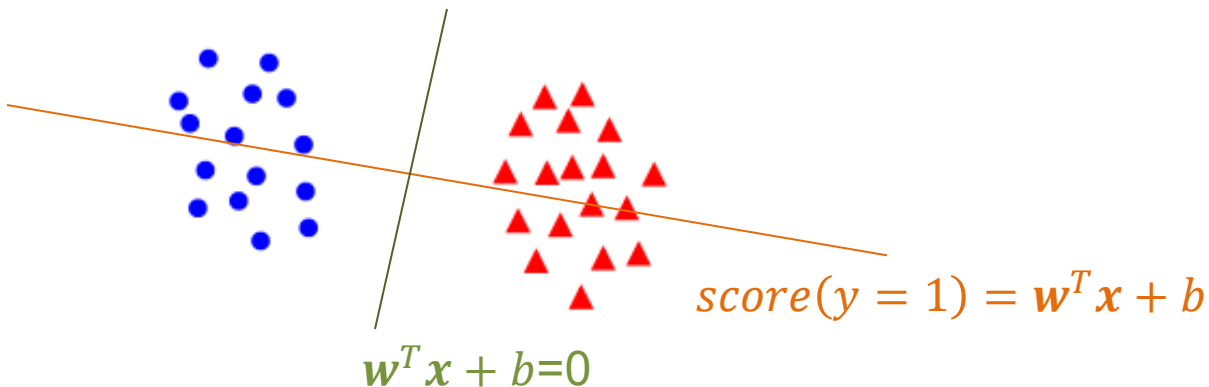
Applied Machine Learning
Derek Hoiem

Linear Models

- A model is **linear** in x if it is based on a weighted sum of the values of x (optionally, plus a constant)

$$\mathbf{w}^T \mathbf{x} + b = \left[\sum_i w_i x_i \right] + b$$

- A **linear classifier** projects the features onto a score that indicates whether the label is positive or negative (i.e., one class or the other). We often show the boundary where that score is equal to zero.
- A **linear regressor** finds a linear model that approximates the prediction value for each set of features.

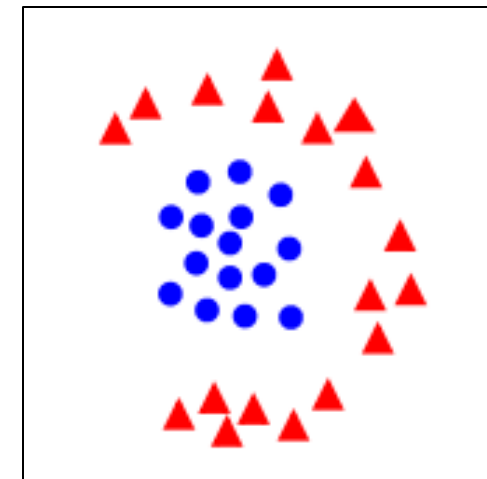
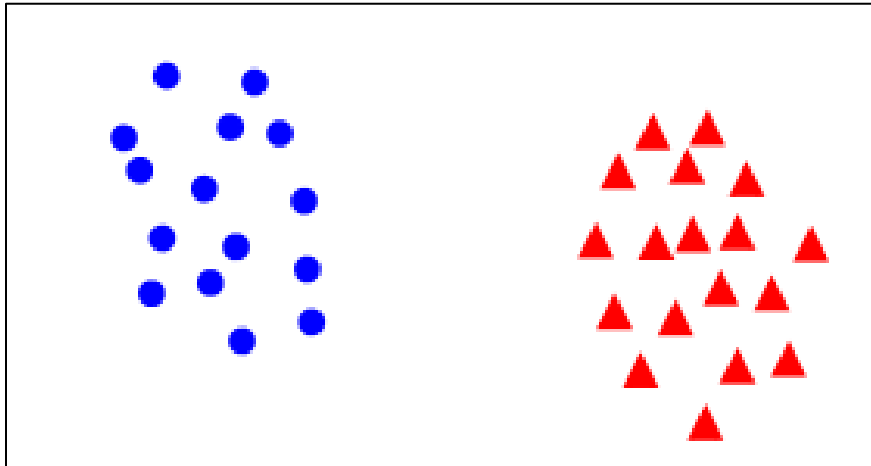


Today's Lecture

- Linear logistic regression: maximize likelihood of target labels given the features
- SVM: maximize the number of data points with confidently correct predictions

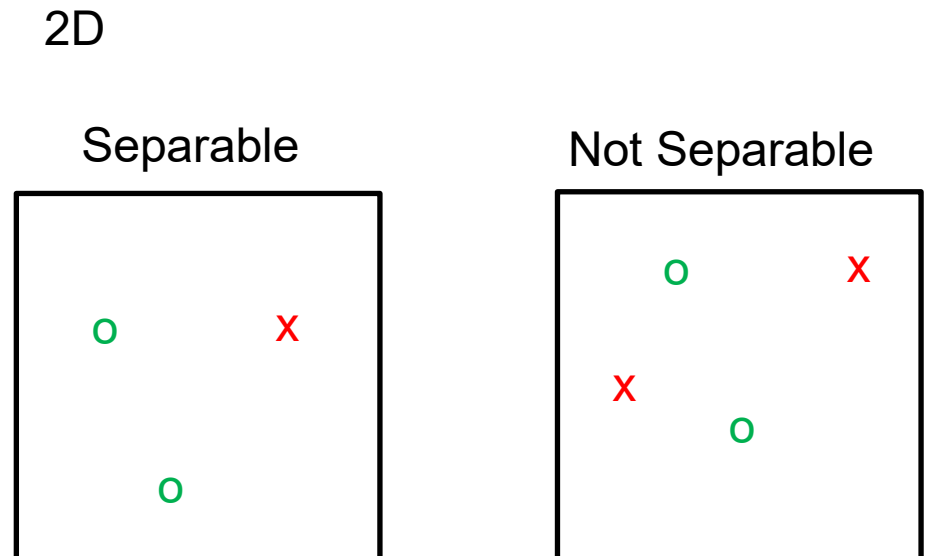
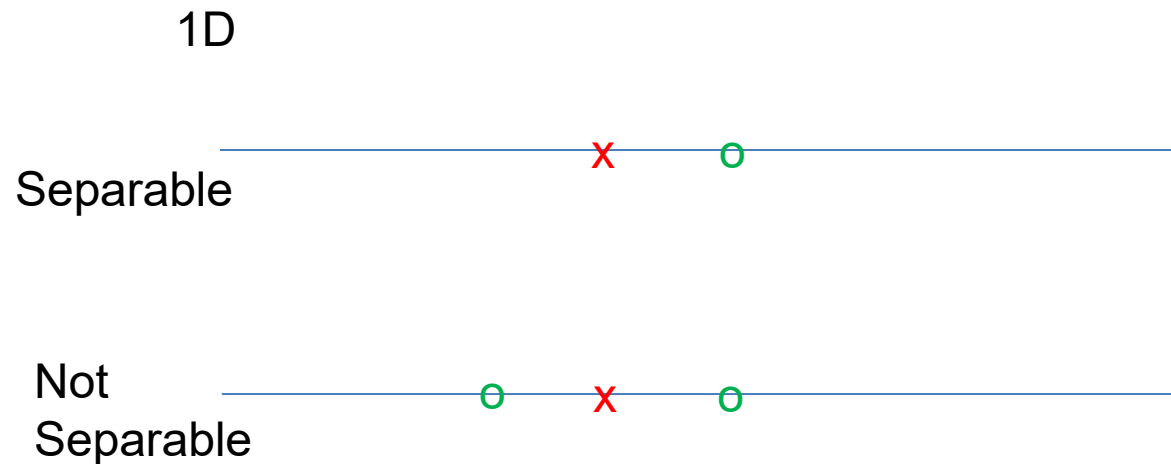
Linear Classifiers and Linear Separability

- **Linear classifier:** $y = 1$ if $\mathbf{w}^T \mathbf{x} + b > 0$
- **Linearly separable:** a line (or hyperplane) in feature space can split the two labels
- Which of these are linearly separable?



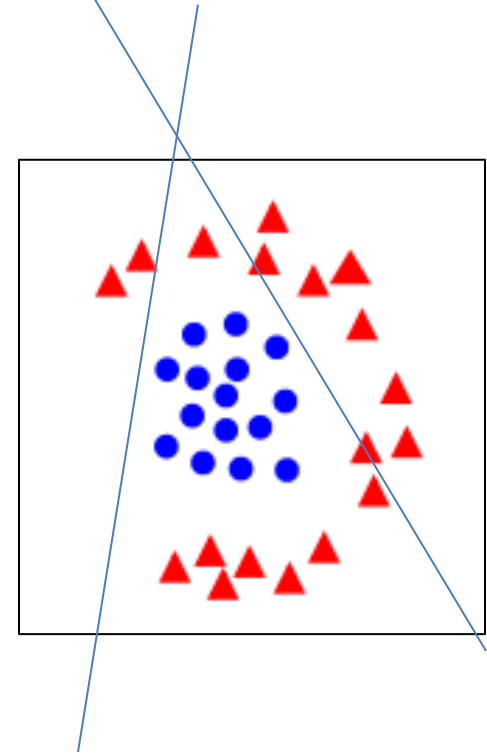
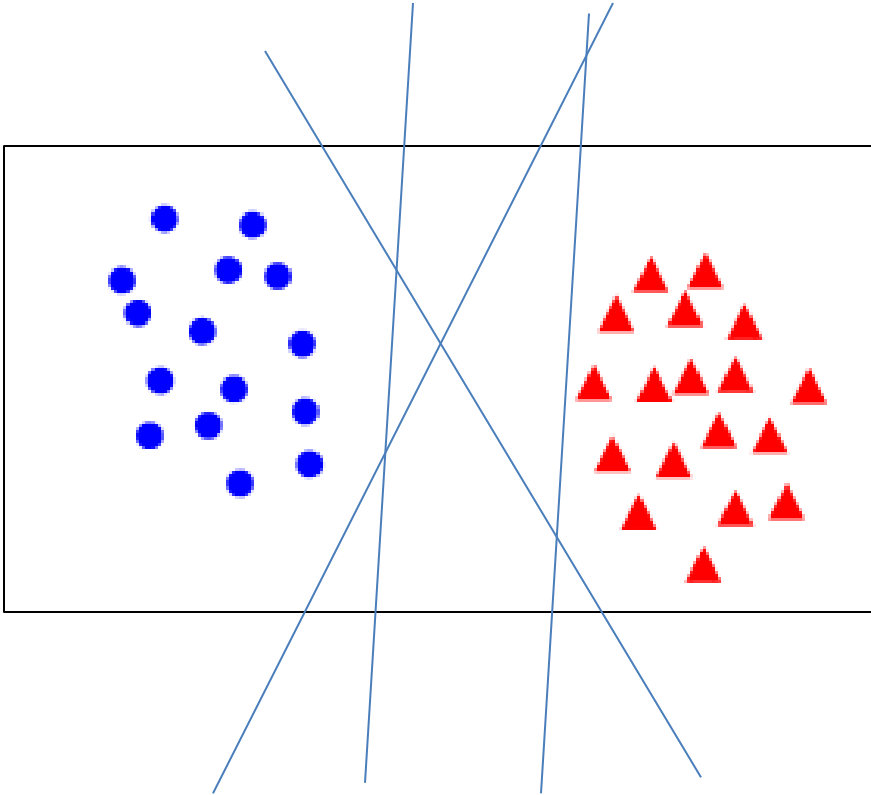
Linear Classifiers and Linear Separability

- In high dimensions, a lot more things are linearly separable
- If you have D dimensions, you can separate $D+1$ points with any arbitrary labeling



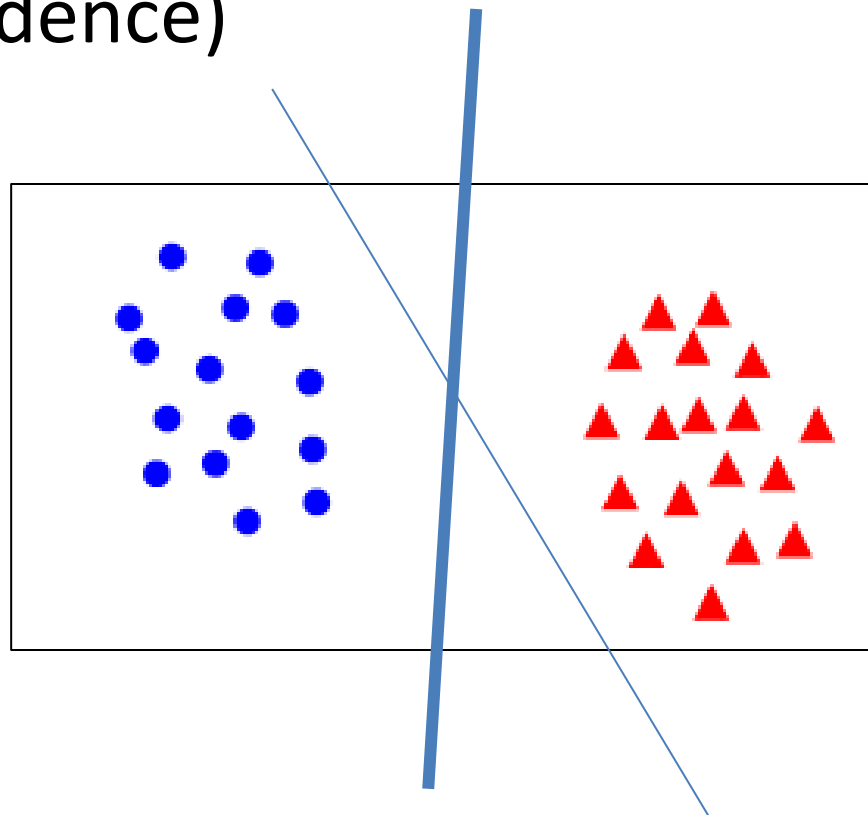
Linear Classifiers and Linear Separability

- But how do you choose which line is best?
- Different classifiers use different objectives to choose the line



Linear Classifiers and Linear Separability

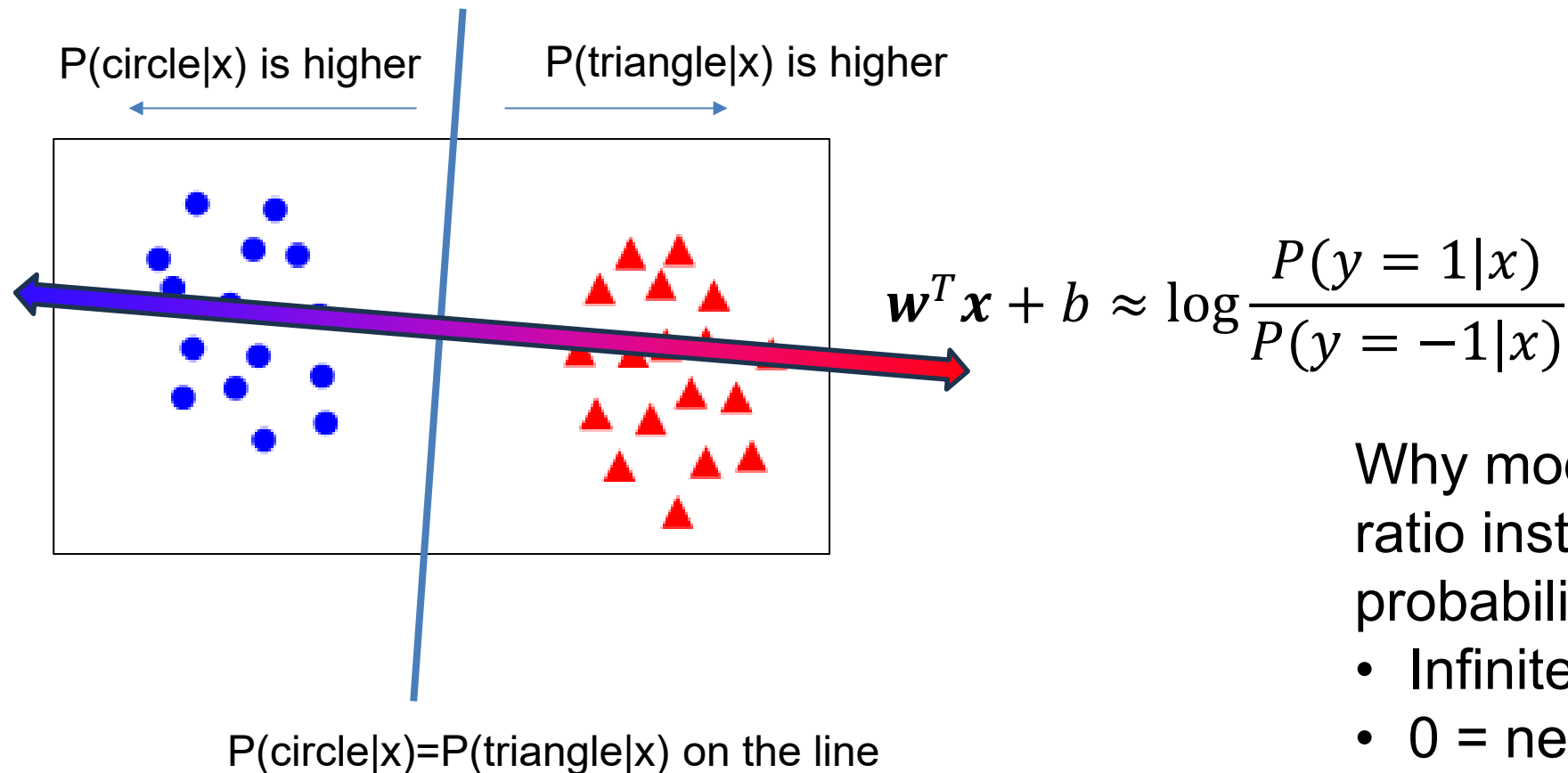
- Different classifiers use different objectives to choose the line
- Common principles are that you want training samples on the correct side of the line (low classification error) by some margin (high confidence)



Thick line is better classification function than thin line because all the examples have a good margin

Linear Logistic Regression

Regress probability by modeling the log odds ratio with a linear function of the features



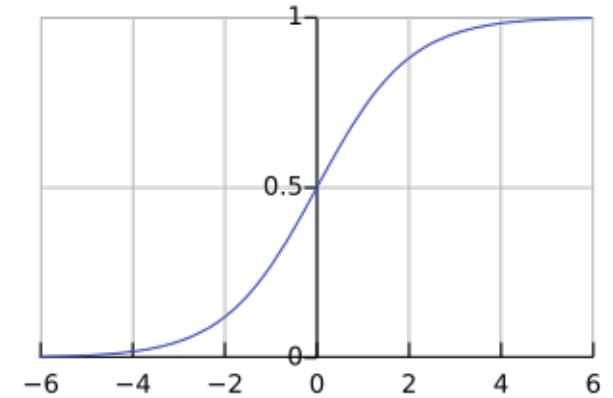
Why model the log odds ratio instead of the probability directly?

- Infinite range
- 0 = neutral prediction

(Linear) Logistic Regression Model

maximize $P(\mathbf{y}|\mathbf{x})$, $y \in \{-1, 1\}$

$$\text{where } P(y = 1|\mathbf{x}) = \frac{1}{1 + \exp(-[\mathbf{w}^T \mathbf{x} + b])}$$



← “Logistic function”

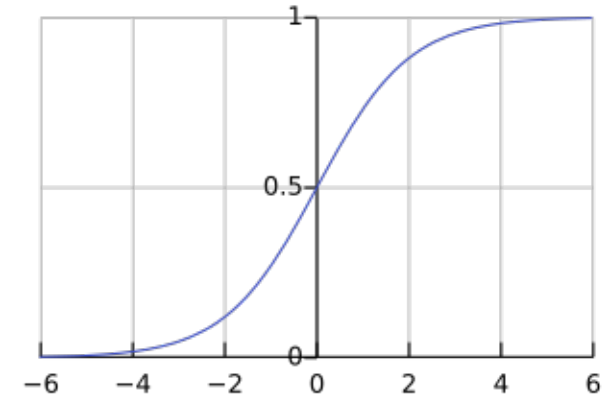
$$\mathbf{w}^T \mathbf{x} + b \approx \log \frac{P(y = 1|\mathbf{x})}{P(y = -1|\mathbf{x})} \quad \leftarrow \text{“Logit”}$$

$$\frac{1}{1 + \exp(-[\mathbf{w}^T \mathbf{x} + b])} = \frac{1}{1 + P(y = -1|\mathbf{x}) / P(y = 1|\mathbf{x})} = \frac{P(y = 1|\mathbf{x})}{P(y = -1|\mathbf{x}) + P(y = 1|\mathbf{x})} = P(y = 1|\mathbf{x})$$

Substitute. Negative flips the ratio in the log. Exp cancels the log

Law of total probability: sum of probabilities overall all values is 1.

(Linear) Logistic Regression Model



← “Logistic function”

maximize $P(\mathbf{y}|\mathbf{x})$, $y \in \{-1, 1\}$

$$\text{where } P(y = 1|\mathbf{x}) = \frac{1}{1 + \exp(-[\mathbf{w}^T \mathbf{x} + b])}$$

$$\mathbf{w}^T \mathbf{x} + b \approx \log \frac{P(y = 1|\mathbf{x})}{P(y = -1|\mathbf{x})} \quad \leftarrow \text{“Logit”}$$

* To simplify notation, I may omit the “b”, which can be avoided by adding a “1” to each feature vector

Note:

$$P(y = -1|\mathbf{x}) = 1 - P(y = 1|\mathbf{x}) = \frac{\exp(-[\mathbf{w}^T \mathbf{x} + b])}{1 + \exp(-[\mathbf{w}^T \mathbf{x} + b])} = \frac{1}{1 + \exp([\mathbf{w}^T \mathbf{x} + b])}$$
$$P(y = \hat{y}|\mathbf{x}) = \frac{1}{1 + \exp(-\hat{y}[\mathbf{w}^T \mathbf{x} + b])}$$

Deriving the loss of logistic regression

Maximize probability of correct label
given features of each data point,
assuming the data points are
conditionally independent

$$\mathbf{w}^* = \operatorname{argmax}_{\mathbf{w}} \prod_n P(y = y_n | \mathbf{x}_n; \mathbf{w})$$

Maximizing $\log(f(x))$ is the same as
maximizing $f(x)$ because $\log(x)$
monotonically increases with x

$$\mathbf{w}^* = \operatorname{argmax}_{\mathbf{w}} \sum_n \log P(y = y_n | \mathbf{x}_n; \mathbf{w})$$

Log of product is the sum of logs.

Turn it into a minimization
problem (as a convention)

$$\mathbf{w}^* = \operatorname{argmin}_{\mathbf{w}} - \sum_n \log P(y = y_n | \mathbf{x}_n; \mathbf{w})$$

Linear Logistic Regression algorithm

- Training

$$\mathbf{w}^* = \underset{\mathbf{w}}{\operatorname{argmin}} - \sum_n \log P(y = y_n | \mathbf{x}_n; \mathbf{w}) + r(\mathbf{w})$$

← regularization

log probability of all labels given features,
assuming that each example is i.i.d.

- Prediction

$$y = 1 \text{ if } \mathbf{w}^T \mathbf{x} > 0$$

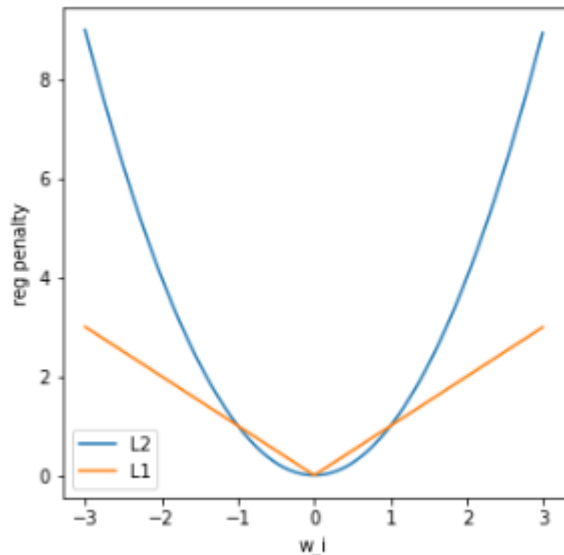
$$P(y = 1 | \mathbf{x}) = \frac{1}{1 + \exp(-\mathbf{w}^T \mathbf{x})} = \frac{\exp(\mathbf{w}^T \mathbf{x})}{\exp(\mathbf{w}^T \mathbf{x}) + 1} \quad \leftarrow \text{Binary (two-class) case}$$

$$P(y = k | \mathbf{x}) = \frac{\exp(\mathbf{w}_k^T \mathbf{x})}{\sum_j \exp(\mathbf{w}_j^T \mathbf{x})} \quad \leftarrow \text{Multiclass case (one } \mathbf{w} \text{ per class)}$$

Training Logistic Regression

$$\mathbf{w}^* = \underset{\mathbf{w}}{\operatorname{argmin}} - \sum_n \log P_{\mathbf{w}}(y = y_n | \mathbf{x}_n) + r(\mathbf{w})$$

- L2 regularization: $r(\mathbf{w}) = \lambda \|\mathbf{w}\|_2^2 = \lambda \sum_i \mathbf{w}_i^2$
- L1 regularization: $r(\mathbf{w}) = \lambda \|\mathbf{w}\|_1 = \lambda \sum_i |\mathbf{w}_i|$



L2 strongly penalizes really big weights

L1 penalizes increasing the magnitude of big and small weights the same

L1 leads to a sparse weight vector (many zeros) – why?

L1 regularization can be used to select features!

When is regularization absolutely essential?

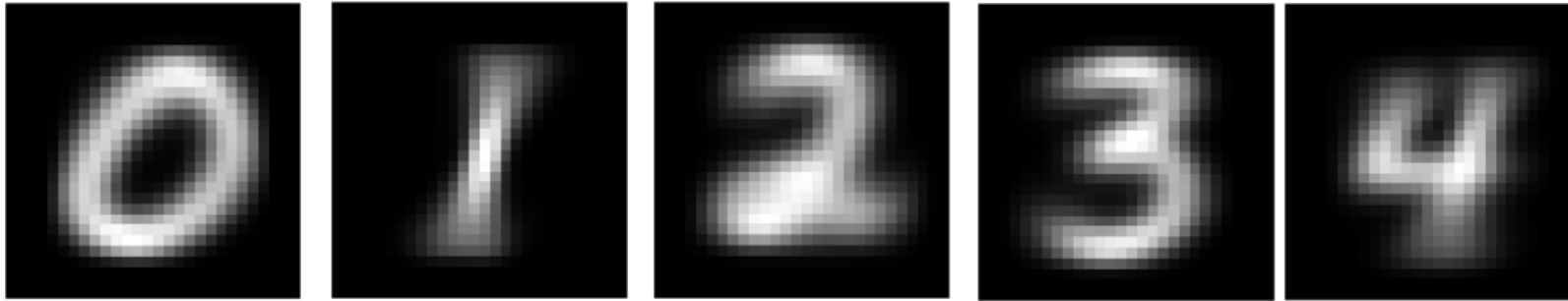
There are [many optimizers](#) for L2 and L1 logistic regression. You will want to use a library.

Inspecting weights for digits

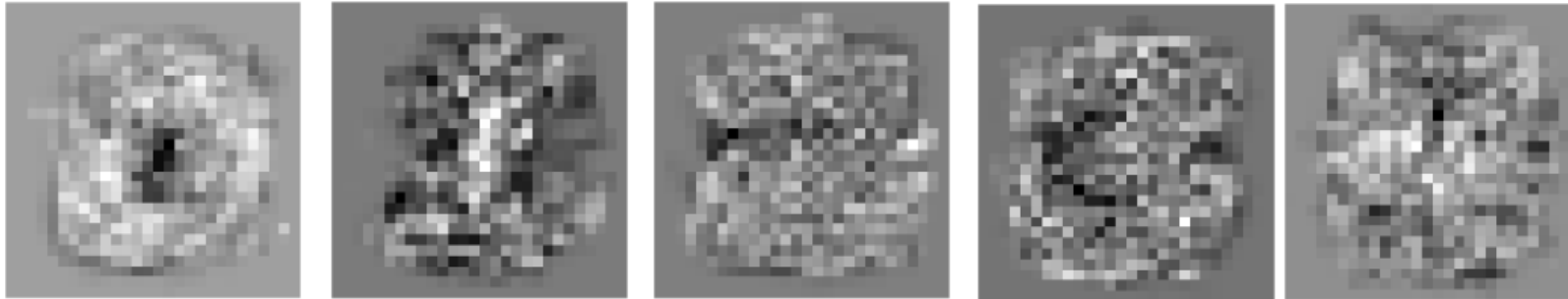
```
model_lr2 = LogisticRegression(max_iter=500, penalty='l2', C=1, verbose=2).fit(x_train[train_indices['m']], y_train[train_indices['m']])
model_lr1 = LogisticRegression(max_iter=500, penalty='l1', C=1, verbose=2, solver='saga').fit(x_train[train_indices['m']], y_train[train_indices['m']])

for k in np.arange(10):
    print(model_lr1.classes_[k])
    display_mnist(x_train[y_train==k].mean(axis=0))
    display_mnist(model_lr2.coef_[k])
    display_mnist(model_lr1.coef_[k])
```

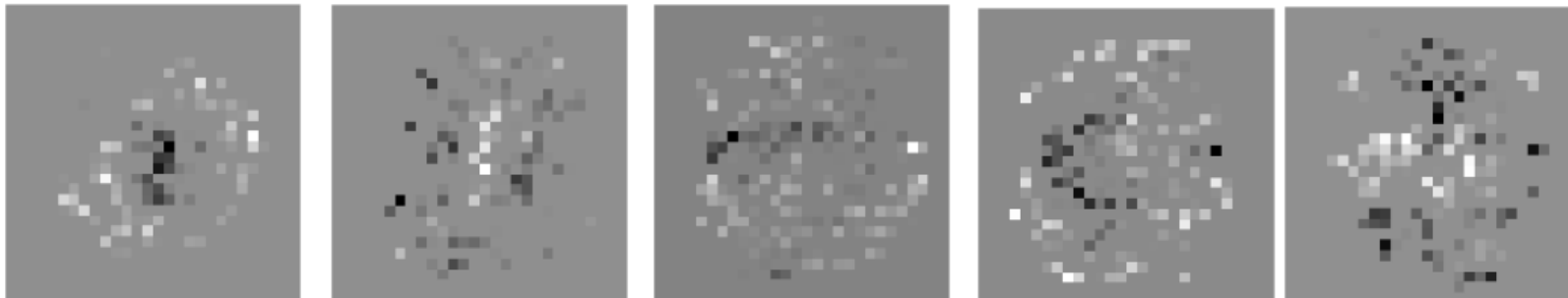
Average
Pixels



L2 weights



L1 weights



Logistic Regression Summary

- Key Assumptions
 - The log odds ratio $\log \frac{P(y = k|x)}{P(y \neq k|x)}$ can be expressed as a linear combination of features
- Model Parameters
 - One coefficient per feature per class (plus bias term)
- Designs
 - L1 or L2 or elastic (both L1 and L2) regularization weight
- When to Use / Strengths
 - Many features, some of which could be irrelevant or redundant
 - Provides a good estimate of label likelihood
- When Not to Use / Weaknesses
 - Features are low-dimensional (linear function not likely to be expressive enough)

Linear logistic regression is typically the last layer of a classification neural network

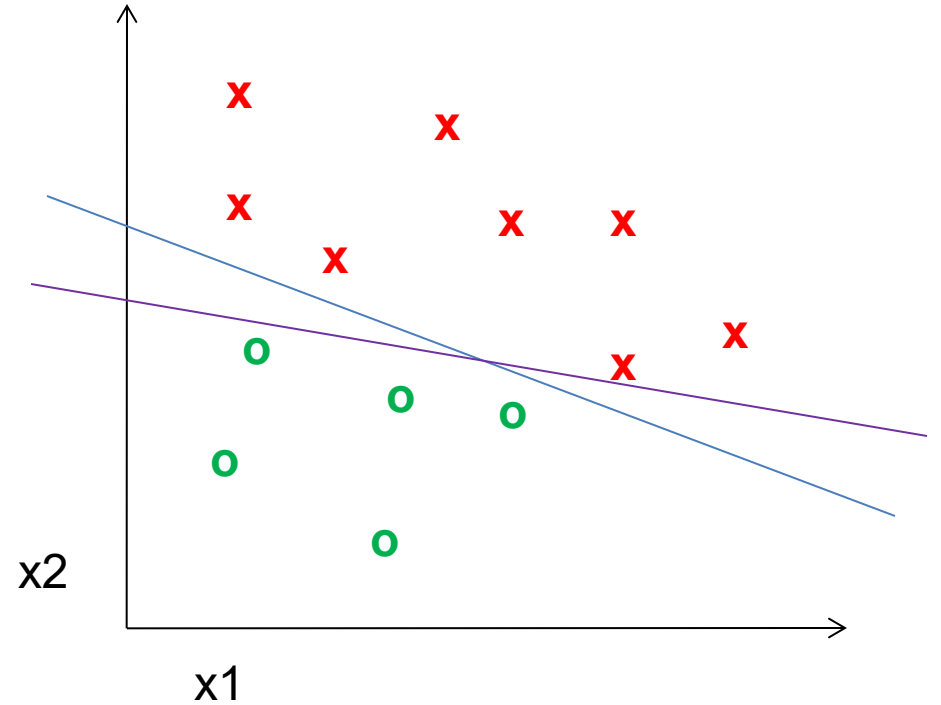
Q1-Q2

<https://tinyurl.com/AML441-L7>

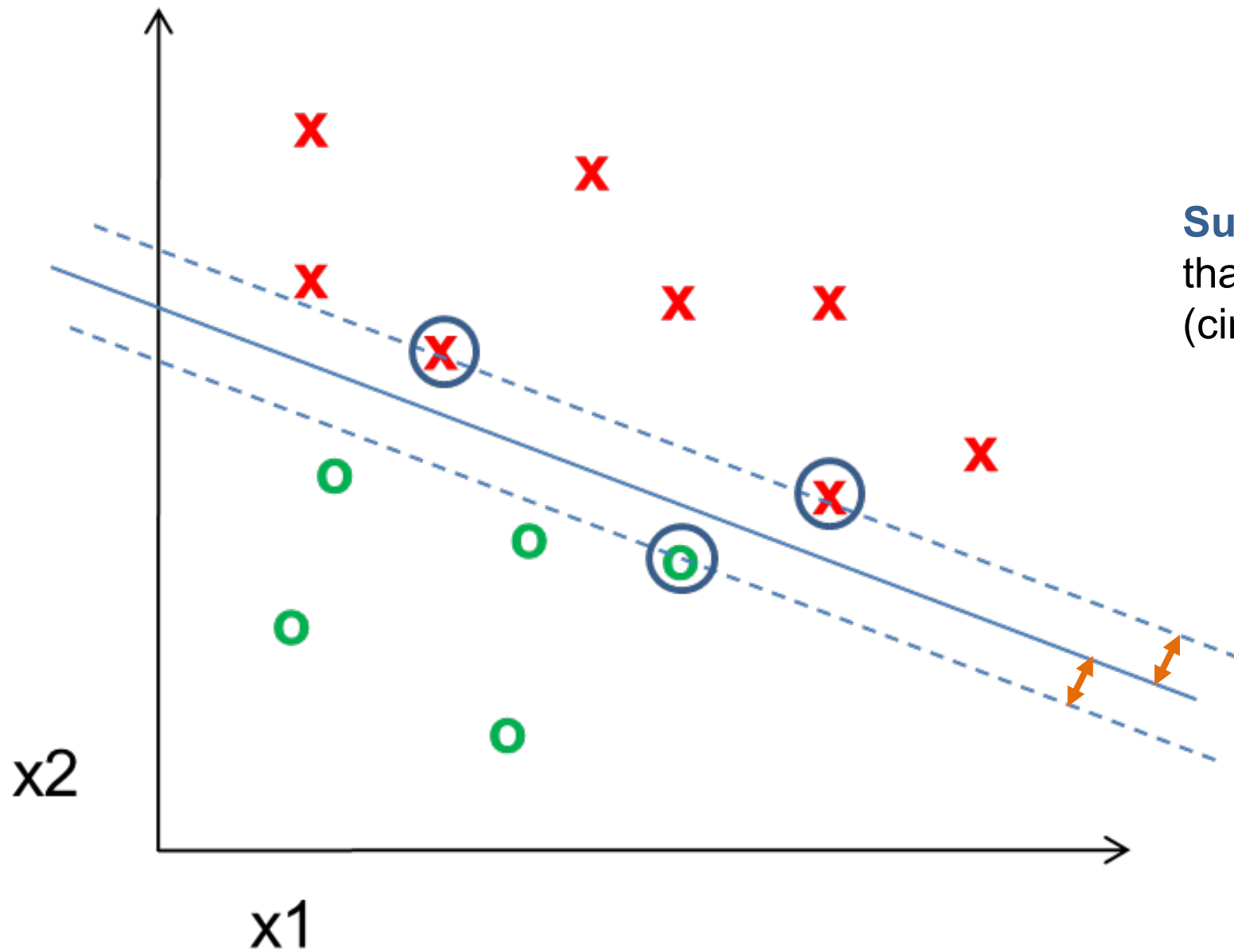


What is the best linear classifier?

- Logistic regression
 - Maximize expected likelihood of true labels given features
 - Every example contributes to loss
- SVM
 - Make all examples at least minimally confident
 - Base decision on a minimal set of examples



SVM Terminology



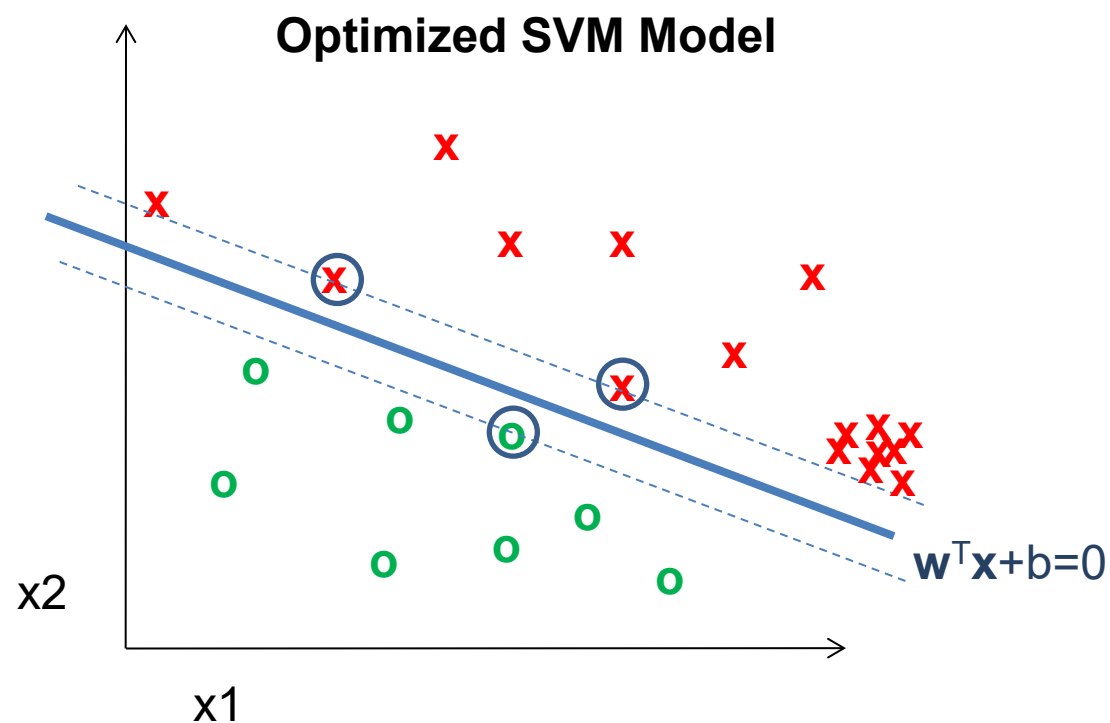
Support Vector: an example that lies on the margin (circled points)

Margin: the distance of examples (in feature space) from the decision boundary

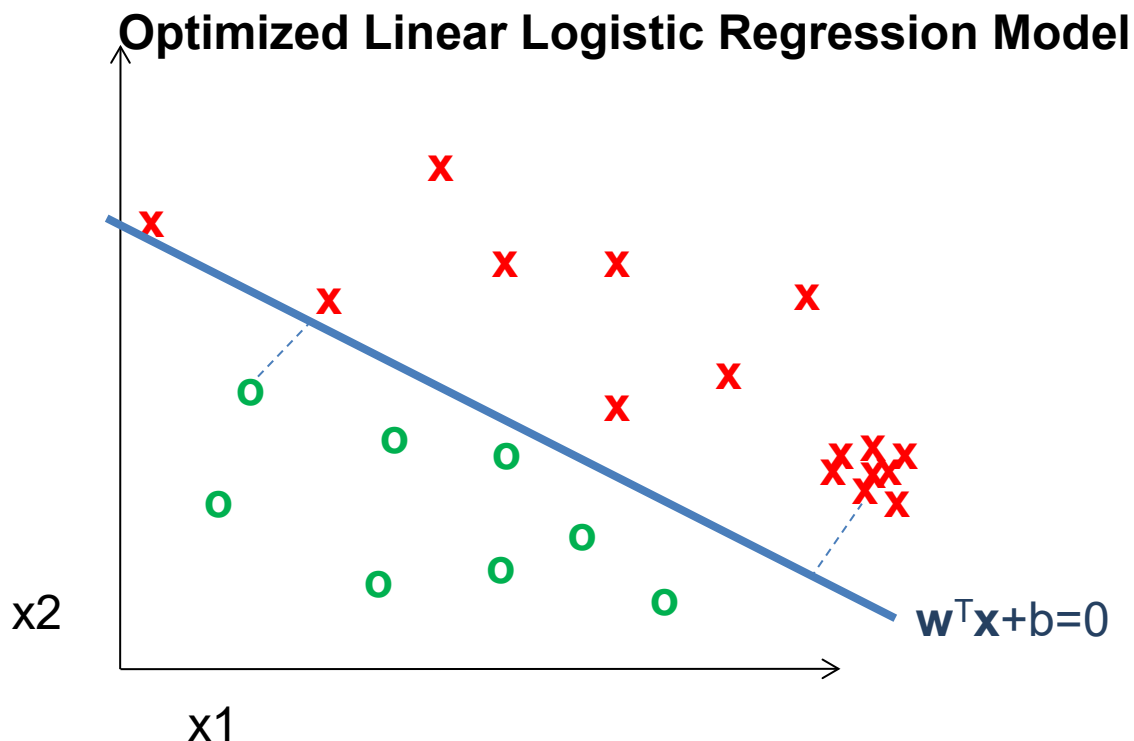
$$m(\mathbf{x}) = \frac{y(\mathbf{w}^T \mathbf{x} + b)}{\|\mathbf{w}\|}$$

$y \in \{-1, 1\}$

SVMs minimize $w^T w$ while preserving a margin of 1



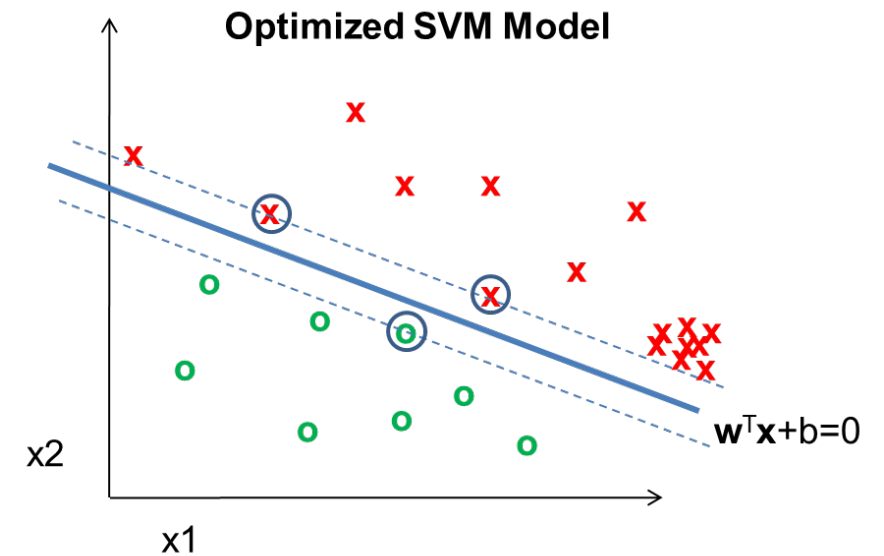
Decision boundary depends only on “support vectors” (circled)



Minimizes the sum of logistic error on all samples, so boundary should be further from dense regions

Why SVMs achieve good generalization

- Maximizing the margin – if all examples are far from the boundary, it is less likely that some test sample will end up on the wrong side of the boundary
 - If classes are linearly separable, the scores can be arbitrarily increased by scaling $\mathbf{w}$, so optimization is expressed as minimize $\mathbf{w}^T \mathbf{w}$ while preserving a margin of 1
- Dependence on few training samples – most training data points could be removed without affecting the decision boundary, which gives an upper bound on the generalization error
- E.g., expected test error is $\leq$ than the smaller of:
 - a. % of training samples that are support vectors
 - b. $D^2/m^2/N$, the diameter of the data compared to the margin divided by the number of examples(see [proof](#))



SVM in Linearly Separable Case

Prediction

$$y_n = \text{sign}(\mathbf{w}^T \mathbf{x}_n + b)$$

Optimization

$$\mathbf{w}^* = \underset{\mathbf{w}}{\text{argmin}} \|\mathbf{w}\|^2$$

subject to

$$y_n(\mathbf{w}^T \mathbf{x}_n + b) \geq 1 \text{ for all } n$$

Here, $y \in \{-1, 1\}$ which is a common convention that simplifies notation for binary classifiers

SVM in **Non**-Linearly Separable Case

Prediction

$$y_n = \text{sign}(\mathbf{w}^T x_n + b)$$

Optimization

$$w^* = \underset{\mathbf{w}}{\text{argmin}} \left[\|\mathbf{w}\|^2 + C \sum_n^N \max(0, 1 - y_n(\mathbf{w}^T x_n + b)) \right]$$

Known as “hinge loss”
Penalty is paid if margin is less than 1

Here, $y \in \{-1, 1\}$ which is a common convention that simplifies notation for binary classifiers

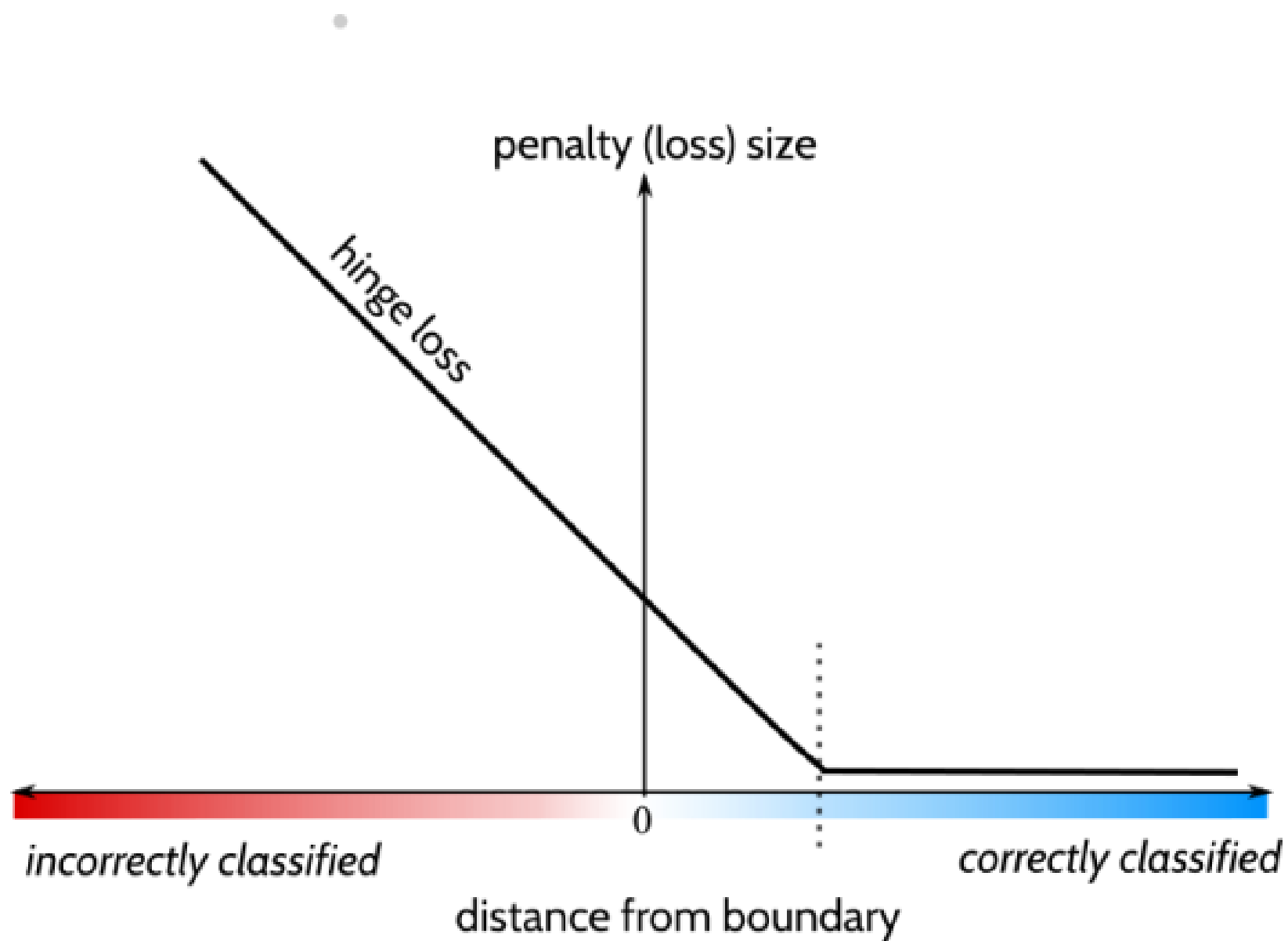
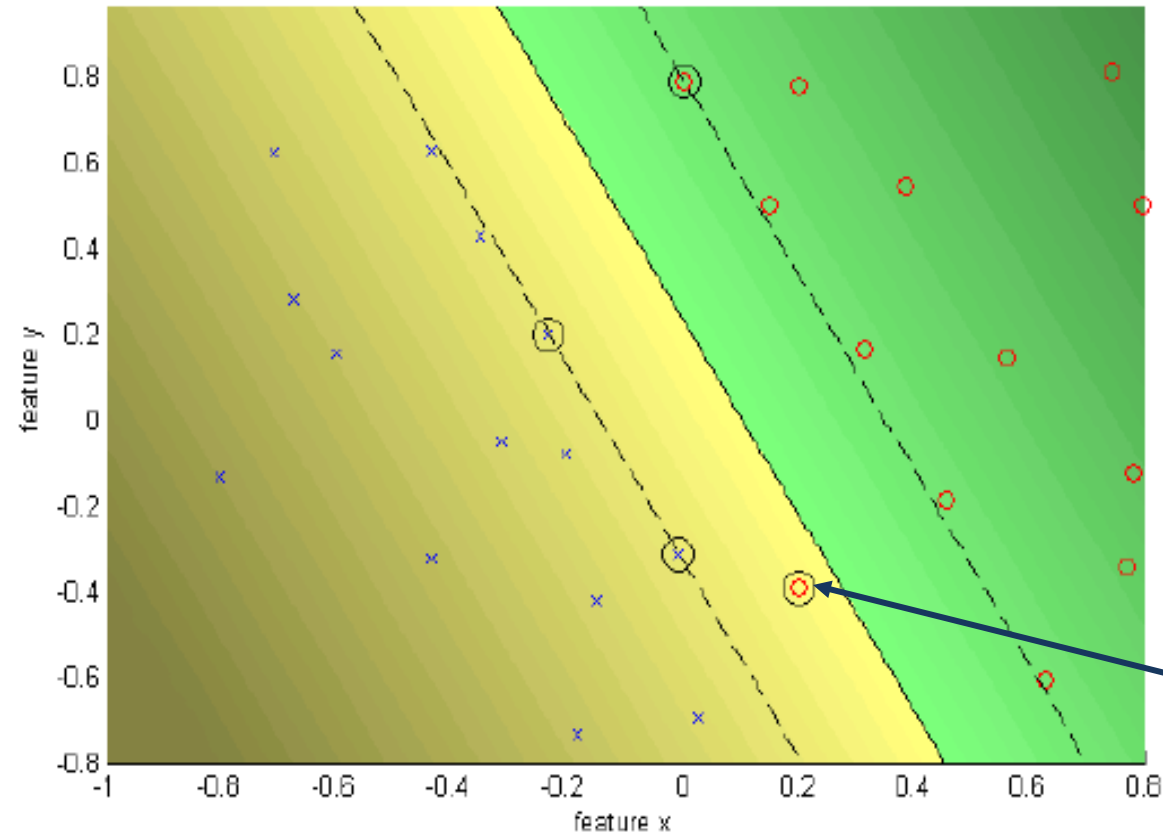


Fig [source](#)

$C = 10$ soft margin



Pays a “slack” penalty for violating the margin



Representer theorem

Optimal weights for many L2-regularized classification and regression functions can be expressed as a weighted combination of training examples

$$\mathbf{w}^* = \sum_n \alpha_n y_n \mathbf{x}_n$$
$$\alpha_n \geq 0, y_n \in \{-1, 1\}$$

So linear SVM is a kind of weighted nearest neighbor with dot product similarity

Conditions apply, e.g. function must be regularized in a Reproducing Kernel Hilbert Space ([details](#))
Does *not* apply to L1 weight regularization because that can't be expressed as a dot product of weights

Primal vs. Dual Formulations of SVM

Prediction

Training Objective

Primal

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

$$\mathbf{w}^* = \underset{\mathbf{w}}{\operatorname{argmin}} \left[\|\mathbf{w}\|^2 + C \sum_n^N \max(0, 1 - y_n(\mathbf{w}^T \mathbf{x}_n + b)) \right]$$

Dual

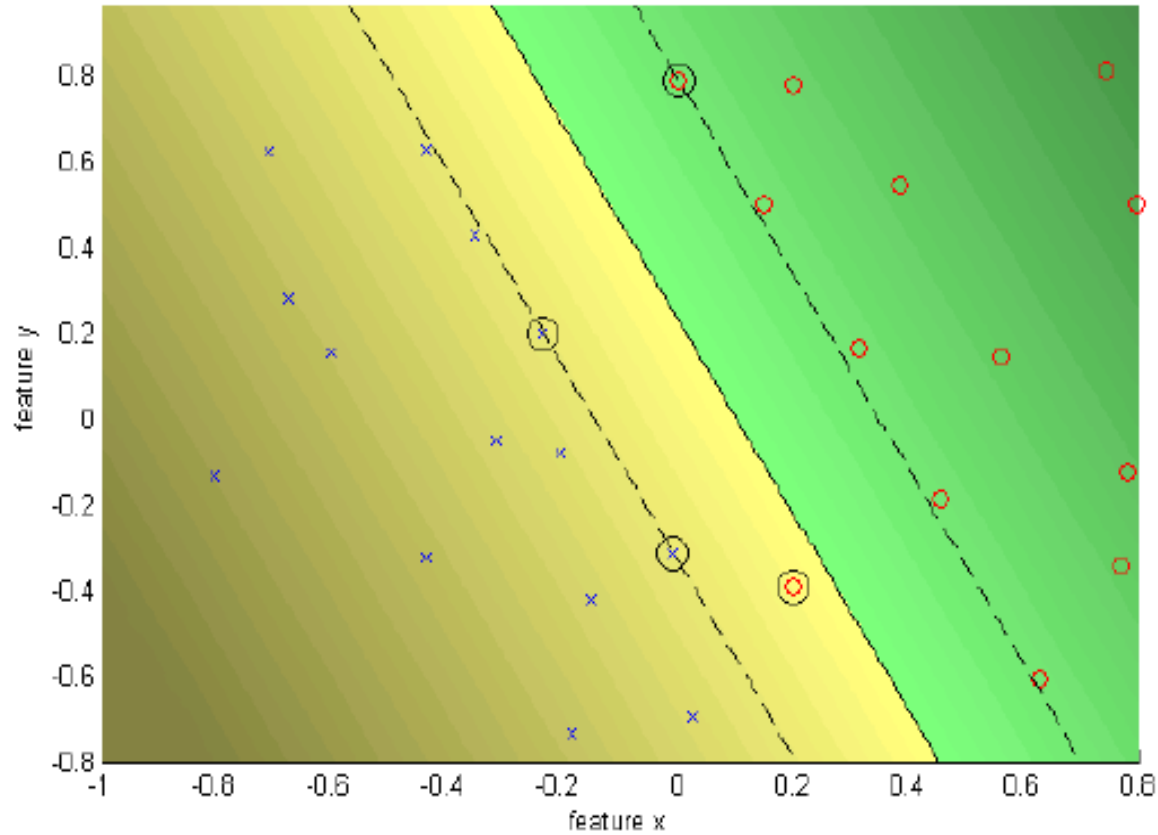
$$f(\mathbf{x}) = \sum_n \alpha_n y_n (\mathbf{x}_n^T \mathbf{x}) + b$$

$$\begin{aligned} \boldsymbol{\alpha}^* &= \underset{\boldsymbol{\alpha}}{\operatorname{argmax}} \left[\sum_i \alpha_i - \frac{1}{2} \sum_{jk} \alpha_j \alpha_k y_j y_k (\mathbf{x}_j^T \mathbf{x}_k) \right] \\ \text{s.t. } &0 \leq \alpha_i \leq C \quad \forall i \quad \text{and} \quad \sum_i \alpha_i y_i = 0 \end{aligned}$$

Primal: parameter for each feature

Dual: parameter for each training example

For SVM, α is sparse (most values are zero)



In dual, $\alpha_i > 0$ only for support vectors

$\alpha_i = 0$ for all others

Non-linear SVMs use a different similarity measure, called a “kernel function”, than dot product

- Linear: $k(\mathbf{x}_i, \mathbf{x}_j) = \mathbf{x}_i^T \mathbf{x}_j$
- Polynomial: $k(\mathbf{x}_i, \mathbf{x}_j) = (1 + \mathbf{x}_i^T \mathbf{x}_j)^d$
- Gaussian: $k(\mathbf{x}_i, \mathbf{x}_j) = \exp(-\frac{\|\mathbf{x}_i - \mathbf{x}_j\|^2}{2\sigma^2})$

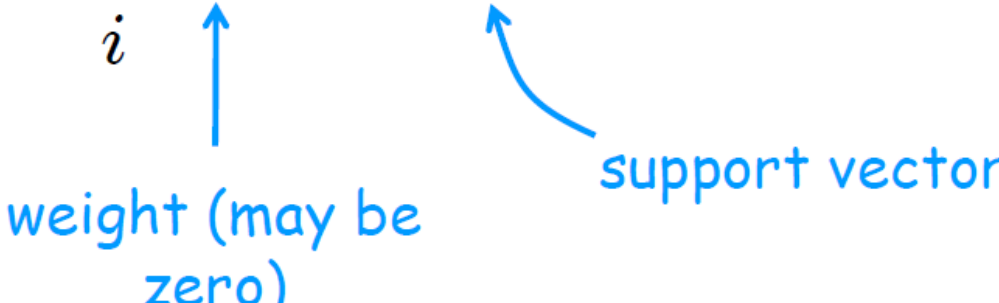
SVM classifier with Gaussian kernel

N = size of training data

$$f(\mathbf{x}) = \sum_i^N \alpha_i y_i k(\mathbf{x}_i, \mathbf{x}) + b$$

weight (may be zero)

support vector



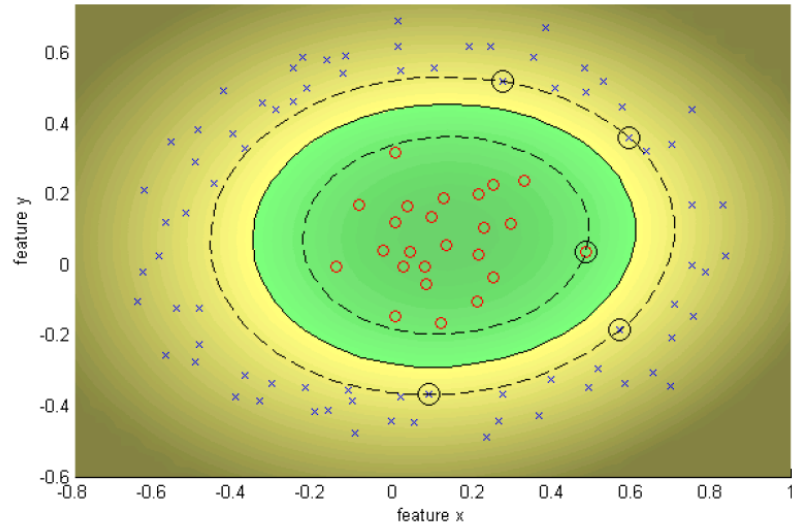
Gaussian kernel $k(\mathbf{x}, \mathbf{x}') = \exp(-||\mathbf{x} - \mathbf{x}'||^2 / 2\sigma^2)$

Radial Basis Function (RBF) SVM

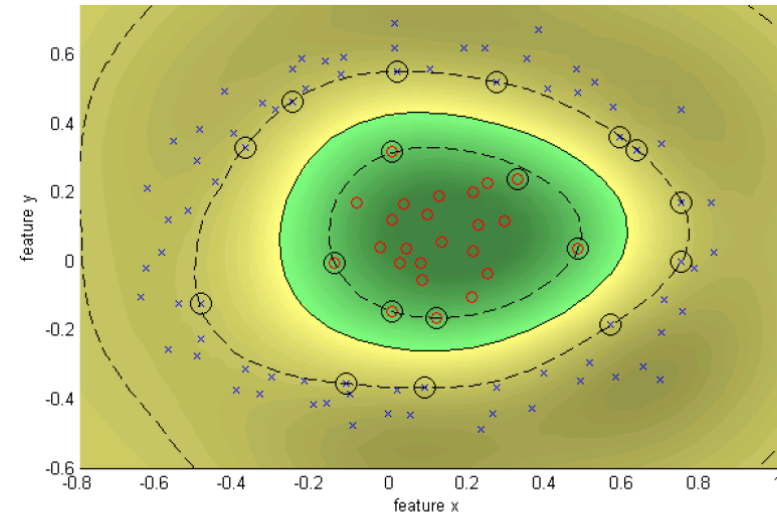
$$f(\mathbf{x}) = \sum_i^N \alpha_i y_i \exp(-||\mathbf{x} - \mathbf{x}_i||^2 / 2\sigma^2) + b$$

Decreasing sigma makes it more like nearest neighbor

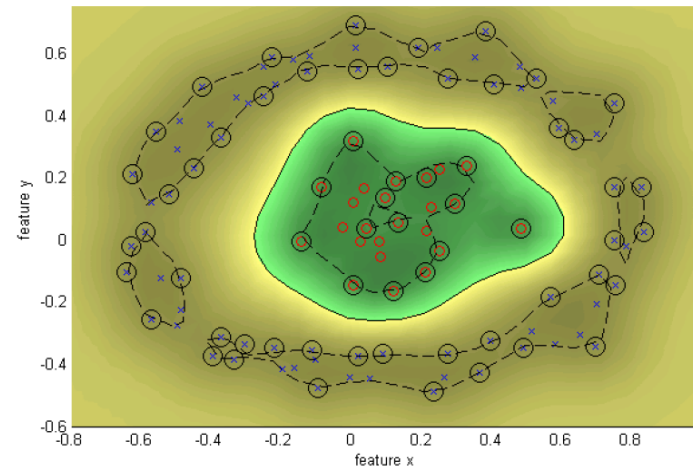
$$\sigma = 1.0 \quad C = \infty$$



$$\sigma = 0.25 \quad C = \infty$$



$$\sigma = 0.1 \quad C = \infty$$



SVM with RBF Kernel Shown

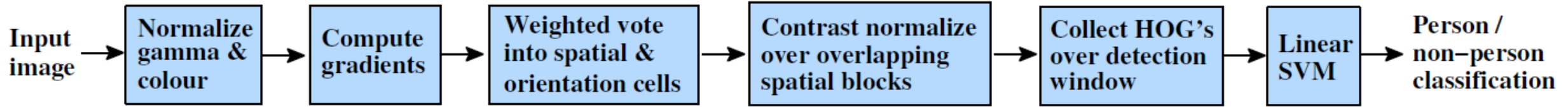
Fig credit: Zisserman [\[link\]](#)

Q3-Q4

<https://tinyurl.com/AML441-L7>



Example application of SVM: Dalal-Triggs 2005



- Detection by scanning window
 - Resize image to multiple scales and extract overlapping windows
 - Classify each window as positive or negative
- Very highly cited (40,000+) paper, mainly for HOG
- One of the best pedestrian detectors for several years



Example application of SVM: Dalal-Triggs 2005

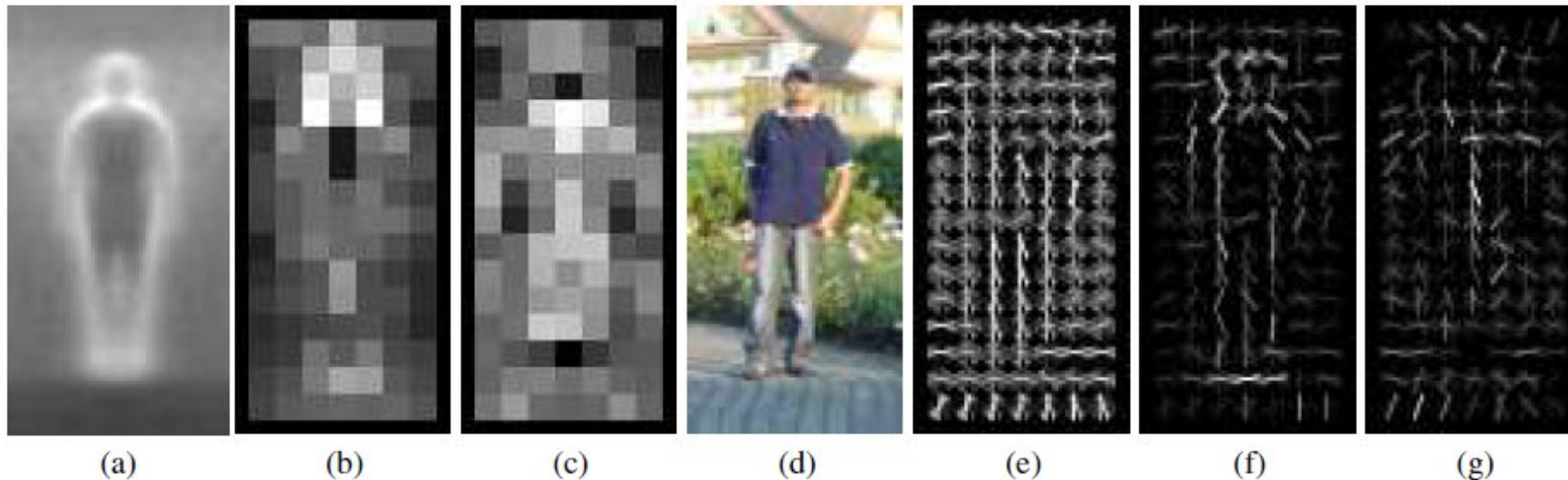
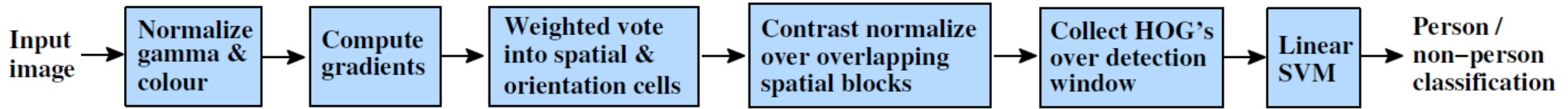
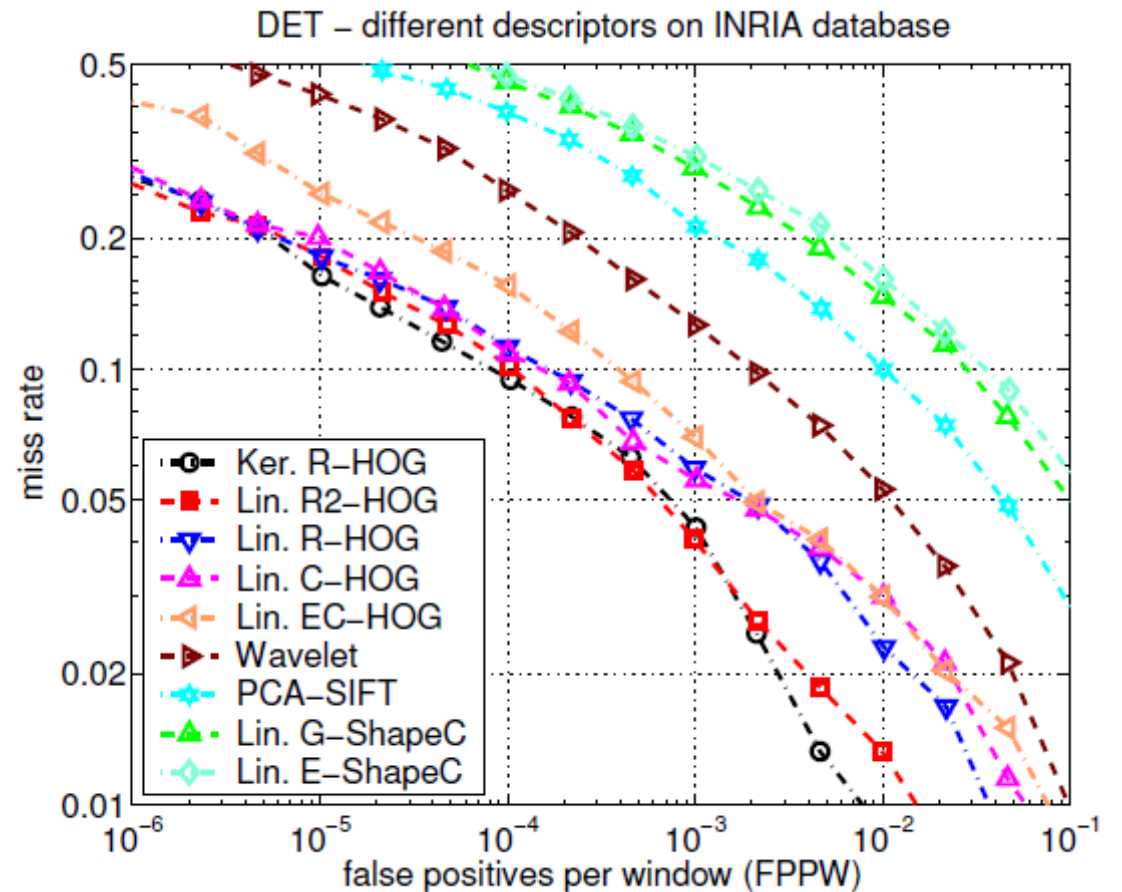
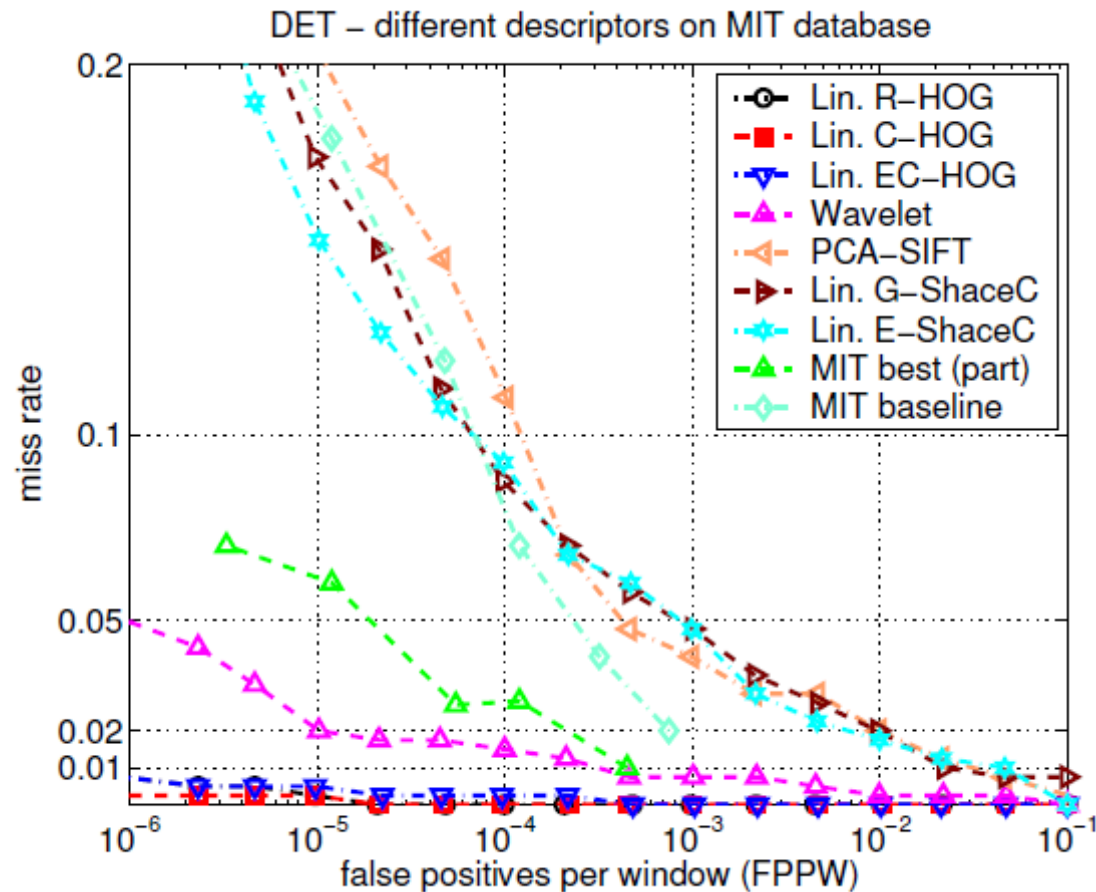
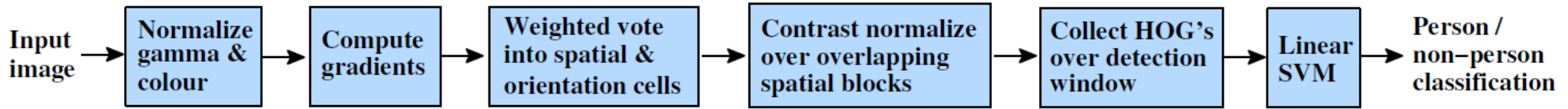


Figure 6. Our HOG detectors cue mainly on silhouette contours (especially the head, shoulders and feet). The most active blocks are centred on the image background just *outside* the contour. (a) The average gradient image over the training examples. (b) Each “pixel” shows the maximum positive SVM weight in the block centred on the pixel. (c) Likewise for the negative SVM weights. (d) A test image. (e) It's computed R-HOG descriptor. (f,g) The R-HOG descriptor weighted by respectively the positive and the negative SVM weights.

- Very highly cited (40,000+) paper, mainly for HOG
- One of the best pedestrian detectors for several years

Example application of SVM: Dalal-Triggs 2005



Using SVMs

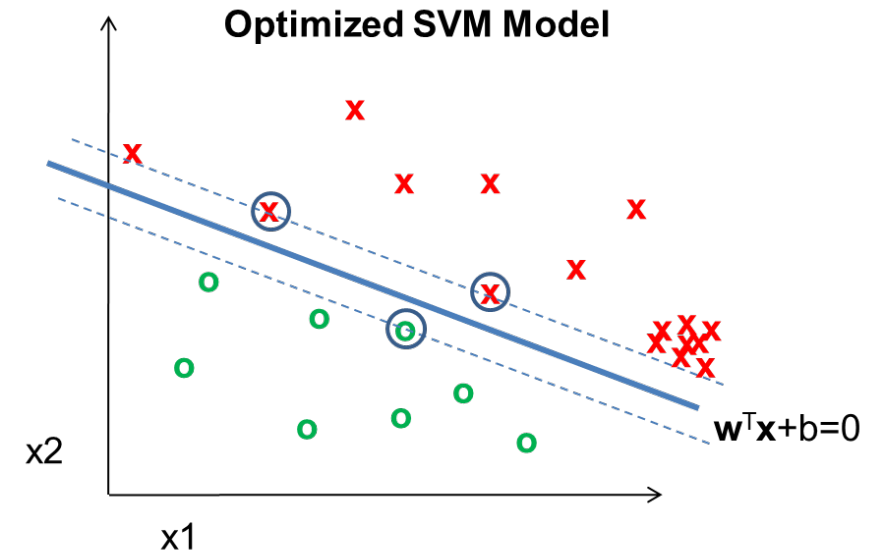
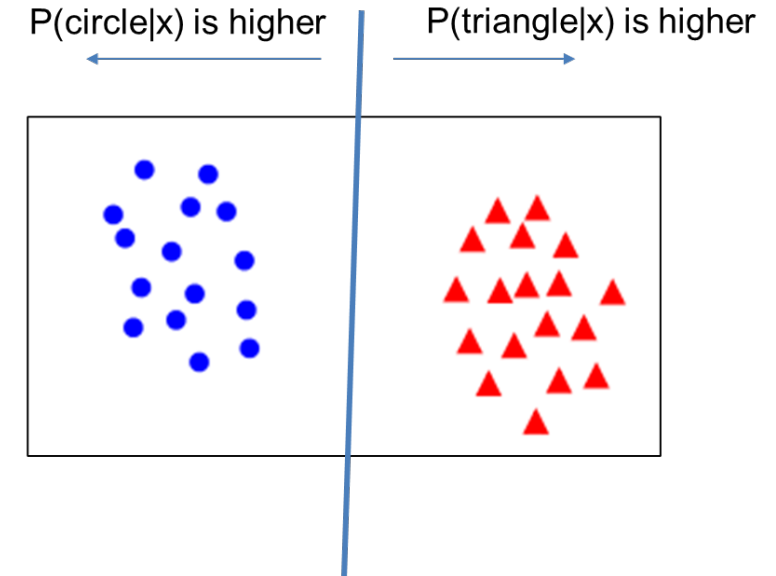
- Good broadly applicable classifier
 - Strong foundation in statistical learning theory
 - Works well with many weak features
 - Requires parameter tuning for C
 - Non-linear SVM requires defining a kernel, and slower optimization/prediction
 - RBF: related to neural networks, nearest neighbor (requires additional tuning)
 - Chi-squared, histogram intersection: good for histograms (but slower, esp. chi-squared)
 - Can learn a kernel function
- Negatives
 - Feature learning is not part of the framework (vs trees and neural nets)
 - Slow training (especially for kernels) – until Pegasos!

Recap

- Nearest neighbor is widely used
 - Super-powers: can instantly learn new classes and predict from one or many examples
- Logistic Regression is widely used
 - Super-powers: Effective prediction from high-dimensional features
- Linear Regression is widely used
 - Super-powers: Can extrapolate, explain relationships, and predict continuous values from many variables
- Almost all algorithms involve nearest neighbor, logistic regression, or linear regression
 - The main learning challenge is typically **feature learning**

Things to remember

- Linear logistic regression and linear SVM are classification techniques that aim to split features between two classes with a linear model
 - Predict categorical values with confidence
- Logistic regression maximizes confidence in the correct label, while SVM just tries to be confident enough
- Non-linear versions of SVMs can also work well and were once popular (but almost entirely replaced by deep networks)
- Nearest neighbor and linear models are the final predictors of most ML algorithms – the complexity lies in finding features that work well with NN or linear models



Next class

- Naïve Bayes Classifier